



Stratégie de vérification et validation de modèles dédiée au calcul de quantités d'intérêt en ingénierie mécanique

Ludovic Chamoin

► To cite this version:

Ludovic Chamoin. Stratégie de vérification et validation de modèles dédiée au calcul de quantités d'intérêt en ingénierie mécanique. Sciences de l'ingénieur [physics]. Ecole Normale Supérieure de Cachan, 2013. tel-01579865

HAL Id: tel-01579865

<https://hal.science/tel-01579865>

Submitted on 31 Aug 2017

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.

Public Domain

**HABILITATION A DIRIGER DES RECHERCHES
DE L'ÉCOLE NORMALE SUPERIEURE DE CACHAN**

Présentée par

Ludovic CHAMOIN

Discipline :
Sciences pour l'ingénieur

Titre :
**Stratégie de vérification et validation de modèles dédiée
au calcul de quantités d'intérêt en ingénierie mécanique**

Soutenue à Cachan le 11 octobre 2013 devant le jury composé de :

N. MOËS	École Centrale de Nantes	Président
D. AUBRY	École Centrale de Paris	Rapporteur
M. BONNET	ENSTA ParisTech	Rapporteur
K. RUNESSON	Université Chalmers (Suède)	Rapporteur
P. CRESTA	EADS Innovation Works	Examineur
P. LADEVÈZE	ENS de Cachan	Examineur
F. LEGOLL	École des Ponts ParisTech	Examineur
M. VOHRALIK	INRIA Paris-Rocquencourt	Examineur

Remerciements

Je tiens à remercier chaleureusement D. Aubry, M. Bonnet et K. Runesson d'avoir accepté de rapporter sur ce travail. J'exprime aussi toute ma reconnaissance envers les autres membres du jury qui ont porté leurs différents regards sur la thématique de recherche abordée : P. Cresta, F. Legoll, M. Vohralik et N. Moës qui m'a aussi fait l'honneur de présider ce jury.

J'ai évidemment une profonde gratitude envers P. Ladevèze qui a accompagné une grande partie de mes activités de recherche depuis la thèse, toujours dans un fort esprit de convivialité, de confiance et d'épanouissement scientifique ; il est un artisan essentiel de ma situation actuelle. Je remercie également toutes les autres personnes qui ont participé aux contributions de cette HDR, et avec qui j'ai collaboré avec bonheur : la formidable équipe de l'institut ICES d'Austin (J.T. Oden et S. Prudhomme notamment), E. Florentin, C. Rey, F. Louf, L. Desvilletes, H. Ben Dhia, R. Cottureau, et bien sûr les brillants étudiants de thèse (J. Waeytens, F. Pled, C. Zaccardi, J. Marchais) et de Master 2 (V. Visseq, R. Bouclier, K. Kergrene, P.E. Allier, ...) que j'ai eu le plaisir d'encadrer.

J'ai enfin une affectueuse pensée pour l'ensemble des membres du LMT-Cachan, pour mes parents, pour ma sœur, ainsi que pour Agnès et Clément qui m'apportent chaque jour une joie immense.

Table des matières

Remerciements	3
Table des matières	i
Introduction	1
1 Concept d'adjoint pour le traitement de quantités d'intérêt	5
1 Espace dual et opérateur adjoint	6
1.1 Espace dual	6
1.2 Opérateur adjoint	7
2 Application à l'estimation d'erreur sur des quantités d'intérêt	7
2.1 Fonction de Green	8
2.2 Estimation d'erreur <i>a posteriori</i>	8
2.3 Calcul adaptatif	9
3 Application au contrôle optimal	10
3.1 Principe de l'état adjoint	10
3.2 Lien avec l'étude de quantités d'intérêt	11
2 Contrôle de l'erreur locale de discrétisation	13
1 Contributions à l'obtention de bornes garanties	14
1.1 Méthode générale utilisant l'ERC	14
1.1.1 Cas de l'élasticité linéaire	14
1.1.2 Cas des problèmes dissipatifs	17
1.2 Extension aux problèmes stochastiques	19
1.2.1 Calcul des bornes	20
1.2.2 Estimation des différentes sources d'erreur	23
1.3 Extension aux problèmes résolus par la XFEM	24
1.3.1 Présentation du problème	24
1.3.2 Construction de solutions admissibles	26
1.4 Extension aux quantités d'intérêt non-linéaires	32
2 Contributions à l'obtention de bornes précises	37
2.1 Nouvelle technique de majoration en espace	37
2.2 Prise en compte des effets d'histoire	40
2.3 Résolution fine du problème adjoint	43
3 Contributions à l'obtention de bornes pratiques	46

3.1	Procédure simplifiée pour le calcul des champs admissibles . . .	46
3.1.1	Approche SPET ou <i>flux-free</i>	47
3.1.2	Nouvelle approche EESPT	49
3.2	Technique non-intrusive pour la résolution de l'adjoint	54
3.2.1	Présentation générale	54
3.2.2	Application aux quantités ponctuelles	57
4	Conclusions partielles et perspectives	60
3	Contrôle de l'erreur locale de réduction	61
1	Maîtrise de la procédure de couplage entre modèles	62
1.1	Couplage particulaire/continu par la méthode Arlequin	63
1.1.1	Présentation des modèles concurrents	63
1.1.2	Processus de couplage	65
1.2	Choix de l'opérateur de couplage	67
1.3	Couplage dans le cadre stochastique	70
1.4	Prise en compte des incompatibilités entre modèles	80
1.4.1	Cas de couplages en statique	80
1.4.2	Cas de couplages en dynamique	86
2	Stratégie adaptative de couplage	92
2.1	Méthode générale	92
2.1.1	Estimation d'erreur locale	93
2.1.2	Adaptation du modèle couplé	95
2.2	Application au couplage particulaire/continu	97
2.3	Application au couplage transport/diffusion	101
2.4	Application au couplage stochastique/déterministe	108
2.4.1	Estimation de l'erreur totale	109
2.4.2	Séparation des sources d'erreur	112
2.5	Procédure optimale d'adaptation	115
3	Contrôle des modèles PGD	121
3.1	Présentation de la PGD	123
3.2	Estimation d'erreur locale par l'ERC	125
3.3	Séparation des sources et adaptation	127
4	Conclusions partielles et perspectives	131
4	Contrôle de l'erreur locale de modèle	133
1	Recalage de modèle par l'ERC modifiée	134
2	Vers la validation en temps réel	137
2.1	Construction d'un modèle réduit de processus d'usinage	138
2.2	Contrôle en temps réel par recalage du modèle réduit	141
2.2.1	Etape de localisation	142
2.2.2	Etape de correction	142
2.3	Application	143
3	Recalage dédié à des quantités d'intérêt	146
3.1	Définition des nouvelles fonctions coût	147
3.1.1	Cas d'une quantité d'intérêt non mesurée	147
3.1.2	Cas d'une quantité d'intérêt mesurée	149

3.2	Sensibilité vis-à-vis des mesures	152
3.3	Application	154
4	Conclusions partielles et perspectives	156
Conclusion générale		157
Bibliographie		159

Introduction

“Essentially, all models are wrong but some are still useful” (George E.P. Box)

Les différents domaines des sciences et de l'ingénierie ont vécu une révolution importante avec l'avènement de la simulation numérique, fruit du développement constant des outils informatiques et des méthodes numériques. Ce constat est particulièrement vrai dans le domaine de la Mécanique des Structures que nous considérons dans ce travail ; à partir des années 80, les progrès à la fois dans les capacités de calcul et dans la richesse des modèles sont devenus tels qu'il était possible de traiter par le calcul des problèmes réalistes. La simulation numérique est alors apparue progressivement comme un outil d'analyse et de conception indispensable, permettant de prédire le comportement des structures mécaniques dans leur environnement.

Nécessité du contrôle des modèles Encore aujourd'hui, l'impact de la modélisation et de la simulation ne cesse de grandir. Les puissances informatiques extraordinaires actuellement disponibles et la pression de l'industrie conduisent à aborder numériquement des problèmes de plus en plus complexes, dans lesquels l'utilisation de modèles pluridisciplinaires devient souvent indispensable. Un enjeu industriel majeur est le *Virtual Testing*, i.e. le remplacement des essais coûteux par des essais virtuels simulés. De plus, la simulation n'est plus simplement envisagée comme un moyen de compréhension ou de prédiction, mais aussi comme un moyen de décision. Cette utilisation de l'ordinateur comme laboratoire virtuel appelle donc naturellement, en plus de méthodes numériques et d'architectures de machines adaptées, à un besoin croissant de connaître la précision des résultats obtenus, le caractère qualitatif de ces résultats ne suffisant plus. Ainsi, un objectif important des sciences numériques modernes est la prédiction quantifiable, i.e. la prédiction systématique du comportement des systèmes physiques avec des métriques quantifiables sur la confiance et les incertitudes.

Pour rendre l'outil numérique crédible et prédictif, la sélection et le contrôle des divers modèles utilisés comme base de la simulation est primordiale : d'une part, les modèles mathématiques utilisés doivent être en concordance avec les phénomènes

physiques observés ; d'autre part, les modèles numériques traités par l'ordinateur doivent fournir une approximation relativement fine de la réponse prédite par ces modèles mathématiques. La thématique scientifique du contrôle des modèles, connue sous le nom de *Vérification et Validation des modèles (V&V)*, a été initiée à la fin des années 70 et fait encore l'objet de nombreux développements dans la recherche et l'industrie [Roache 1998, Babuška et Strouboulis 2001, Oberkampf *et al.* 2003, Stein 2003, Ladevèze et Pelle 2004, Babuška et Oden 2005, SBE 2006]. La vérification est une procédure exclusivement mathématique qui permet d'étudier la qualité de la résolution numérique ; la validation est quant à elle tournée vers la physique et donne le degré de pertinence des modèles mathématiques comme abstraction du réel.

Par conséquent, la maîtrise des outils de simulation pour la Mécanique des Structures nécessite l'interaction entre trois grandes familles de modèles :

- les modèles physiques qui sont des reflets (partiels) du réel décrits à partir de mesures expérimentales ;
- les modèles mathématiques (ou conceptuels) qui sont des représentations théoriques de ce réel généralement établies à l'aide d'équations aux dérivées partielles (EDP). A noter que la panoplie des modèles mathématiques mis en œuvre peut être considérablement large pour certaines applications, notamment avec l'émergence des méthodes de réduction de modèles qui visent à remplacer des modèles mathématiques complexes par d'autres plus grossiers mais aussi plus simples à manipuler ; cette démarche fait alors apparaître une cascade de modèles mathématiques hiérarchisés selon leur complexité ;
- les modèles numériques, fréquemment liés à la méthode des éléments finis, qui sont directement manipulés par l'outil informatique.

Dans la transition du réel vers le virtuel, la V&V nécessite le contrôle de différentes sources d'erreur afin de garantir la pertinence des résultats de simulation. On peut classer ces sources en plusieurs catégories (Figure 1), chacune étant caractérisée par la référence choisie pour définir l'erreur : (i) les **erreurs de mesure**, directement liées à la précision et la fiabilité des mesures expérimentales faites pour définir le modèle physique ; (ii) les **erreurs de modèle**, traduisant l'écart de représentation entre le modèle physique et le modèle mathématique ; (iii) les **erreurs de réduction de modèle**, générées au passage entre un modèle mathématique initial et un autre modèle mathématique plus grossier ; (iv) les **erreurs de discrétisation**, liées à l'utilisation de méthodes numériques pour résoudre de façon approchée les modèles mathématiques.

Contrôle de modèles ciblé sur des quantités d'intérêt Dans la grande majorité des applications des sciences et de l'ingénierie qui utilisent la simulation numérique, l'analyse de la solution sur l'ensemble du domaine d'étude n'est nécessaire que d'un point de vue qualitatif. Le côté quantitatif est quant à lui dirigé vers quelques caractéristiques localisées de cette solution, appelées quantités d'intérêt, et servant directement au dimensionnement ou à la prise de décision. En Mécanique des Structures, ces quantités peuvent être la résultante ou le moment d'effort exercé sur une partie (bord) d'une pièce, mais elles sont plus souvent liées à des données locales comme la température ou la contrainte de Mises maximale, les facteurs d'intensité de contrainte, etc. . .

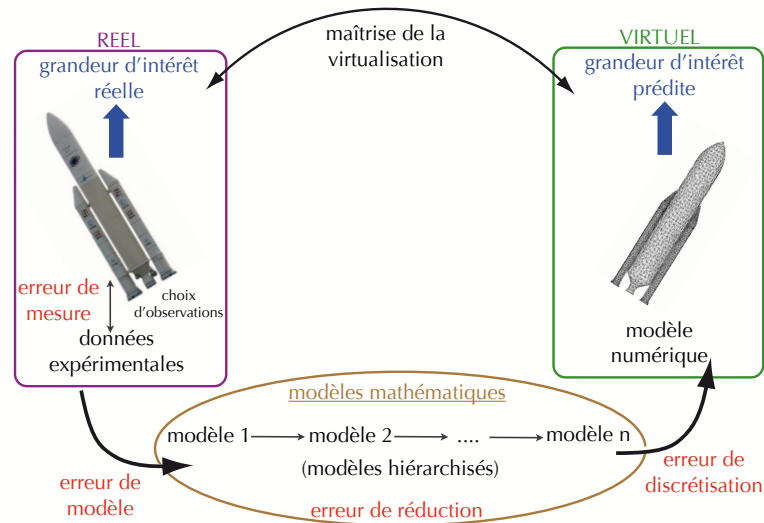


Figure 1: Définition des différentes sources d'erreur dans le passage du réel vers le virtuel

Il est alors légitime de chercher à définir des modèles de simulation optimisés, c'est-à-dire d'une complexité minimale tout en restant suffisamment précis, pour le calcul de telles quantités. La cohérence de cette approche, qualifiée de *goal-oriented* parmi la communauté scientifique, réside dans le fait que tout modèle n'est par essence qu'une représentation incomplète, avec une validité limitée. Tout comme la légendaire carte à l'échelle 1/1 de Borges, seul un modèle identique au système pourrait être considéré comme une représentation exacte de ce dernier ; la pertinence d'un modèle n'est donc en pratique nécessaire que pour l'utilisation envisagée. Un exemple illustratif est encore celui de la cartographie, les modèles (cartes) ne reflétant qu'une partie des informations du réel suivant l'application souhaitée et à des échelles différentes selon les besoins : cartes routières, topographiques, géologiques, climatiques, etc. . . Dans le domaine de la V&V, la vision *goal-oriented* se traduit par le contrôle des paramètres de simulation qui sont les plus influents pour la quantité d'intérêt étudiée, menant à une maîtrise des modèles partielle et spécifique à l'objectif visé. La démarche de V&V est alors simplifiée par rapport à la démarche classique, souvent fastidieuse et complexe, qui vise à contrôler l'ensemble des modèles. La question initiale à se poser est donc :

Quel est l'objectif de la simulation ?

Le challenge consiste alors à déterminer les modèles optimaux pour atteindre cet objectif, en fixant le niveau d'erreur acceptable sur les divers modèles impliqués. Ce challenge est encore loin d'être atteint, en particulier au niveau industriel, au regard des résultats d'un récent *benchmark* proposé sur le problème de Girkmann [Pitkaranta *et al.* 2009].

Le cadre mathématique adapté au traitement de quantités d'intérêt est celui des

espaces duaux, lié au fait qu'une quantité d'intérêt peut être définie à partir de la solution globale en utilisant des fonctions d'extraction [Greenberg 1948, Washizu 1953]. D'autre part, la notion d'adjoint est essentielle pour étudier la sensibilité des quantités d'intérêt vis-à-vis des erreurs commises depuis l'expérience jusqu'à la simulation proprement dite [Lions 1971, Giles et Süli 2002, Estep 2004]; elle est liée à l'introduction d'un problème annexe à résoudre dont la solution détermine la sensibilité des quantités d'intérêt vis-à-vis des divers paramètres des modèles. Ces différents aspects sont brièvement rappelés dans le Chapitre 1 de ce document.

Travaux abordés dans le manuscrit Nous présentons dans les Chapitres 2, 3 et 4 quelques contributions ayant pour objectif de construire des modèles de simulation optimisés en vue du calcul d'une quantité d'intérêt, ou d'un ensemble de quantités d'intérêt. Une majorité des outils présentés ont été développés au LMT-Cachan, au sein de l'UTR *Vérification et Validation des modèles* que j'anime, et à partir des concepts de base considérés au LMT liés à l'erreur en relation de comportement. Une autre partie des outils, en particulier ceux liés à la vérification des modèles réduits, a été développée dans le cadre de collaborations avec l'Université du Texas à Austin (Institut ICES).

Le Chapitre 2 présente les travaux liés à la vérification "classique" des modèles, c'est-à-dire au contrôle de l'erreur de discrétisation, pour les problèmes résolus par la méthode des éléments finis. Nous montrons comment il est possible d'obtenir, pour une large gamme de problèmes de mécanique, des bornes d'erreur locales à la fois garanties, précises et relativement simples à calculer et à implémenter pour l'ingénieur. Ce sujet a fait l'objet de plusieurs travaux novateurs développés notamment dans les thèses de J. Panetier, J. Waeytens et F. Pled [Prudhomme *et al.* 2004, Chamoin et Ladevèze 2007 - 2008 - 2009, Cottureau *et al.* 2010, Ladevèze *et al.* 2010, Ladevèze et Chamoin 2010, Panetier *et al.* 2010, Pled *et al.* 2011, Chamoin *et al.* 2012, Pled *et al.* 2012, Waeytens *et al.* 2012, Ladevèze *et al.* 2013].

Le Chapitre 3 traite du contrôle des modèles réduits, très fréquemment utilisés aujourd'hui pour aborder les modèles complexes. Dans ce contexte, nous analysons les deux cas du couplage de modèles et de l'application de la PGD (*Proper Generalized Decomposition*), en mettant en place des stratégies d'estimation d'erreur et d'adaptation de modèle pour le contrôle de quantités d'intérêt. Une partie des outils développés a fait l'objet de mon post-doctorat aux USA et des relations ultérieures avec l'institut ICES (Austin, USA), tandis qu'une autre partie a fait l'objet des thèses de C. Zaccardi et J. Marchais [Chamoin *et al.* 2008, Prudhomme *et al.* 2009, Chamoin *et al.* 2010, Oden *et al.* 2010, Ben Dhia *et al.* 2011, Ladevèze et Chamoin 2011, Prudhomme *et al.* 2012, Chamoin et Ladevèze 2012, Ladevèze et Chamoin 2012, Chamoin et Desvillettes 2013, Marchais *et al.* 2013, Zaccardi *et al.* 2013]. Enfin, le Chapitre 4 aborde des travaux plus récents et donc moins matures sur la validation en temps réel et vis-à-vis d'une quantité d'intérêt [Bouclier *et al.* 2013, Chamoin *et al.* 2013].

Concept d'adjoint pour le traitement de quantités d'intérêt

Sommaire

1	Espace dual et opérateur adjoint	6
1.1	Espace dual	6
1.2	Opérateur adjoint	7
2	Application à l'estimation d'erreur sur des quantités d'intérêt	7
2.1	Fonction de Green	8
2.2	Estimation d'erreur <i>a posteriori</i>	8
2.3	Calcul adaptatif	9
3	Application au contrôle optimal	10
3.1	Principe de l'état adjoint	10
3.2	Lien avec l'étude de quantités d'intérêt	11

La notion d'adjoint, conjointement à celle de dualité, a une longue histoire d'application dans l'étude de sensibilité et l'optimisation [Lions 1971, Giles et Süli 2002, Estep 2004]. Plus récemment, elle a été utilisée pour l'estimation *a posteriori* de l'erreur sur des quantités d'intérêt issues de la résolution numérique d'équations différentielles [Becker et Rannacher 2001]. Dans cette partie, on donne une vision générale de cette notion et de ses applications.

On considère un modèle décrit par une équation aux dérivées partielles scalaire, de solution u , mise sous la forme générique :

$$Lu = f \quad (1.1)$$

et associée à des conditions limites fixées. $L \in \mathcal{L}(\mathcal{U}, \mathcal{F})$ est supposé être un opérateur différentiel linéaire, et f un chargement suffisamment régulier. On suppose de plus que $\mathcal{U}$ et $\mathcal{F}$ sont des espaces vectoriels normés sur $\mathbb{R}$, et on note $\|\cdot\|$ la norme sur $\mathcal{U}$. Enfin, on suppose que le problème précédemment décrit est bien posé au sens de Hadamard, i.e. la solution u existe, est unique et est stable par rapport à f .

On s'intéresse dans la suite à une quantité scalaire produite par ce modèle, i.e. une fonctionnelle de sortie appelée quantité d'intérêt et notée $Q(u)$. Elle est définie comme une forme, supposée continue et linéaire, de $\mathcal{U}$ dans $\mathbb{R}$.

1 Espace dual et opérateur adjoint

1.1 Espace dual

Il est commode d'écrire la quantité d'intérêt $Q(u)$ à partir du crochet de dualité :

$$Q(u) = \langle u, Q \rangle \quad (1.2)$$

où Q appartient à l'espace des formes linéaires et continues sur $\mathcal{U}$, i.e. l'espace dual de $\mathcal{U}$. Cet espace vectoriel, noté $\mathcal{U}^*$, est équipé de la norme (dite duale ou triple) définie par :

$$\|Q\|_* = \sup_{w \in \mathcal{U}, w \neq 0} \frac{|Q(w)|}{\|w\|} \quad (1.3)$$

garantissant l'inégalité de Cauchy généralisée suivante :

$$|\langle u, Q \rangle| \leq \|u\| \cdot \|Q\|_* \quad (1.4)$$

Le terme $Q \in \mathcal{U}^*$ permet d'extraire la quantité d'intérêt de la solution globale u ; il est donc souvent relié à une fonction d'extraction appelée aussi extracteur. En guise d'illustration, considérons le cas où $\mathcal{U}$ est un espace de fonctions définies sur un domaine spatial Ω . Le crochet de dualité est alors défini par :

$$\langle p, q \rangle = \int_{\Omega} p(\mathbf{x})q(\mathbf{x})d\Omega \quad \forall p \in \mathcal{U}, \forall q \in \mathcal{U}^* \quad (1.5)$$

et Q peut être caractérisé par une fonction $q(\mathbf{x}) \in \mathcal{U}^*$. Un choix simple pour q peut être $q(\mathbf{x}) = \chi_{\omega}(\mathbf{x})/|\omega|$ afin d'étudier la valeur moyenne de u sur un sous-domaine

$\omega \subset \Omega$; χ_ω (resp. $|\omega|$) désigne ici la fonction caractéristique (resp. la mesure) de ω .

A noter que l'espace dual $\mathcal{U}^*$ est généralement plus régulier que l'espace vectoriel $\mathcal{U}$ d'origine (ou primal). Par exemple, $\mathcal{U}^*$ est un espace de Banach même si $\mathcal{U}$ n'en est pas un. Par ailleurs, si $\mathcal{U}$ est un espace de Hilbert, le théorème de représentation de Riesz indique que le crochet de dualité correspond au produit scalaire et par conséquent $\mathcal{U}^* = \mathcal{U}$.

1.2 Opérateur adjoint

On note $\mathcal{F}^*$ l'espace dual de $\mathcal{F}$. L'opérateur adjoint de L , noté L^* , est défini comme l'opérateur satisfaisant l'identité bilinéaire suivante :

$$\langle Lu, v \rangle = \langle u, L^*v \rangle \quad \forall u \in \mathcal{U}, \forall v \in \mathcal{F}^* \quad (1.6)$$

$L^* \in \mathcal{L}(\mathcal{F}^*, \mathcal{U}^*)$ est un opérateur linéaire.

L'adjoint L^* est en pratique obtenu par intégration par parties, avec un processus en deux étapes :

- on calcule l'adjoint formel en supposant que u et v ont un support compact et en ignorant les termes de bord. Par exemple :

$$\begin{aligned} Lu &= -\nabla \cdot [a(\mathbf{x}) \nabla u] + b(\mathbf{x}) \nabla u + c(\mathbf{x})u \\ \implies L^*v &= -\nabla \cdot [a(\mathbf{x}) \nabla v] - b(\mathbf{x}) \nabla v + c(\mathbf{x})v \end{aligned} \quad (1.7)$$

- on fixe les conditions limites de l'adjoint comme celles minimales permettant de vérifier l'identité bilinéaire (1.6). A noter que seule la forme des conditions limites imposées est importante, et leur valeur est généralement prise nulle. Pour les problèmes d'évolution (avec conditions initiales nulles) sur un intervalle temporel $[0, T]$, $\langle Lu, v \rangle$ contient des termes de la forme $\int_0^T \frac{\partial^k u}{\partial t^k} v(t) dt$ et l'intégration par parties nécessite d'associer des conditions finales (à $t = T$) à l'opérateur L^* ; le problème qui en découle est alors rétrograde en temps.

Remarque : Lorsque l'opérateur L est non-linéaire, il n'y a pas d'adjoint "naturel". On définit généralement $G(e) = L(u + e) - L(u)$, et un opérateur adjoint $G^(e)$ associé vérifiant :*

$$\langle G(e), v \rangle = \langle e, G^*(e)v \rangle \quad \forall e \in \mathcal{U}, \forall v \in \mathcal{F}^* \quad (1.8)$$

2 Application à l'estimation d'erreur sur des quantités d'intérêt

Les opérateurs adjoints et arguments de dualité sont utilisés dans l'analyse numérique des EDP, notamment pour l'estimation d'erreur sur des quantités d'intérêt. Dans le cadre de l'estimation d'erreur *a priori*, leur utilisation remonte aux travaux

menés dans [Aubin 1967, Nitsche 1968]. Pour l'estimation d'erreur *a posteriori*, leur utilisation est plus récente ; elle fut abordée dans [Gartland 1984] puis développée pleinement dans de nombreux travaux à la fin des années 90 [Paraschivoiu *et al.* 1997, Rannacher et Suttmeier 1997, Cirak et Ramm 1998, Ladevèze *et al.* 1999, Oden et Prudhomme 1999]. L'ingrédient de base est la résolution d'un problème auxiliaire qui utilise l'adjoint L^* de l'opérateur linéaire L , et dont les données proviennent de la quantité d'intérêt considérée.

2.1 Fonction de Green

Il y a une connexion intime entre le calcul de quantité d'intérêt et les notions de dualité et d'adjoint vues précédemment. Supposons qu'on souhaite calculer la quantité $Q(u) = \langle u, Q \rangle$ avec u solution de (1.1) ou de la forme faible associée :

$$A(u, v) = \langle Lu - f, v \rangle = 0 \quad \forall v \in \mathcal{F}^* \quad (1.9)$$

On introduit dans ce cas un problème, dit dual ou adjoint, qui se base sur L^* et consiste à trouver $\tilde{u} \in \mathcal{F}^*$ tel que :

$$L^* \tilde{u} = Q \quad \text{ou} \quad \langle v, L^* \tilde{u} - Q \rangle = 0 \quad \forall v \in \mathcal{U} \quad (1.10)$$

avec des conditions limites homogènes.

L'analyse variationnelle conduit alors à la formule de représentation suivante :

$$Q(u) = \langle u, Q \rangle = \langle u, L^* \tilde{u} \rangle = \langle Lu, \tilde{u} \rangle = \langle f, \tilde{u} \rangle \quad (1.11)$$

La solution adjointe $\tilde{u}$, appelée *fonction de Green généralisée*, permet donc de calculer la quantité d'intérêt $Q(u)$ pour divers chargements f du problème initial, et sans connaître explicitement u . Elle peut être vue comme une mesure de la sensibilité de $Q(u)$ vis-à-vis de f .

Remarque : dans le cas où le chargement Q du problème adjoint est construit à partir de la fonction de Dirac δ , la solution adjointe $\tilde{u}$ devient la véritable fonction de Green.

2.2 Estimation d'erreur *a posteriori*

Dans le cas où on a une solution approchée u_h de u , on s'intéresse à l'erreur $Q(u) - Q(u_h)$ occasionnée sur la quantité d'intérêt. En utilisant la solution adjointe $\tilde{u}$ définie précédemment, on obtient directement :

$$Q(u) - Q(u_h) = \langle f - Lu_h, \tilde{u} \rangle = \langle R(u_h), \tilde{u} \rangle \quad (1.12)$$

et il en découle la majoration :

$$|Q(u) - Q(u_h)| \leq \|R(u_h)\|_{\mathcal{F}} \|\tilde{u}\|_{\mathcal{F}^*} \quad (1.13)$$

où $R(u_h) = f - Lu_h$ est le résidu (calculable). On voit donc que $\tilde{u}$, en jouant un rôle de pondération, permet de lier l'erreur globale sur u à celle locale sur $Q(u)$;

le scalaire $\|\tilde{u}\|_{\mathcal{F}^*}$ est un facteur de stabilité pour la quantité d'intérêt [Eriksson et al. 1995]. Un exemple illustratif est donné sur la Figure 1.1 pour un cas élasto-dynamique de poutre en traction sur le domaine spatio-temporel $[0, L] \times [0, T]$. La quantité d'intérêt est la vitesse moyenne en $x = L/2$ sur la plage temporelle $[t_0, t_1]$, et le second membre de (1.12) s'écrit alors $\int_0^T \int_0^L (\dot{u} - \dot{u}_h) \tilde{N} dx dt$ où $\tilde{N}$ est l'effort normal associé au problème adjoint.

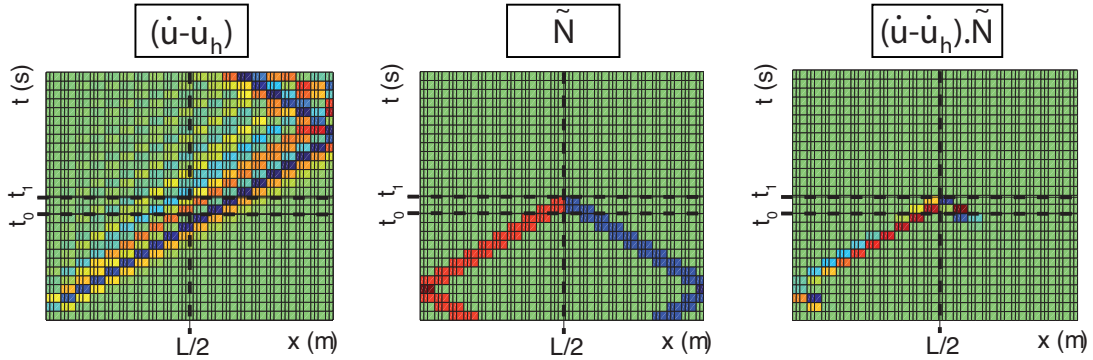


Figure 1.1: Visualisation de la pondération du résidu par la solution adjointe pour une poutre en dynamique transitoire

Une première possibilité pour estimer l'erreur sur la quantité d'intérêt est de calculer une approximation numérique (fine) de $\tilde{u}$ ou $\|\tilde{u}\|_{\mathcal{F}^*}$. Une autre possibilité, lorsque u_h est obtenue par la méthode des éléments finis, consiste à utiliser la propriété d'orthogonalité de Galerkin :

$$\langle R(u_h), v_h \rangle = 0 \quad \forall v_h \in \mathcal{F}_h^* \quad (1.14)$$

$\mathcal{F}_h^*$ désignant l'espace éléments finis des fonctions test. On obtient alors, en notant $\tilde{u}_h \in \mathcal{F}_h^*$ l'approximation éléments finis de $\tilde{u}$:

$$\begin{aligned} Q(u) - Q(u_h) &= \langle R(u_h), \tilde{u} - \tilde{u}_h \rangle = \langle f, \tilde{u} - \tilde{u}_h \rangle \\ \implies |Q(u) - Q(u_h)| &\leq \|R(u_h)\|_{\mathcal{F}} \cdot \|\tilde{u} - \tilde{u}_h\|_{\mathcal{F}^*} \end{aligned} \quad (1.15)$$

et une majoration de $|Q(u) - Q(u_h)|$ est possible à partir d'estimateurs d'erreur globaux majorant $\|R(u_h)\|_{\mathcal{F}}$ et $\|\tilde{u} - \tilde{u}_h\|_{\mathcal{F}^*}$.

Remarque : dans le cas où l'opérateur L est non-linéaire, on utilise généralement l'opérateur adjoint issu d'une linéarisation (cf. (1.8)), menant à la résolution d'un problème adjoint linéaire qui ne permet cependant pas d'obtenir une majoration garantie de l'erreur sur la quantité d'intérêt.

2.3 Calcul adaptatif

La procédure d'estimation d'erreur précédente peut aussi être utilisée pour optimiser la discrétisation dans une stratégie adaptative. L'idée générale est de séparer

l'estimateur en contributions par zones, par exemple sur les N éléments d'un maillage éléments finis :

$$|Q(u) - Q(u_h)| \leq \sum_{E=1}^N |\langle R(u_h), \tilde{u} \rangle| \quad (1.16)$$

et d'utiliser un algorithme glouton en raffinant (ou déraffinant) les zones en fonction de leur contribution à l'erreur. A noter que cette pratique permet le contrôle adaptatif de la borne d'erreur, mais pas de l'erreur elle-même à cause des possibles compensations d'erreur entre les zones. Une alternative consiste à classer ces zones selon qu'elles apportent une contribution d'erreur positive ou négative, et à mener une approche probabiliste de raffinement de zone pour équilibrer les contributions et annuler ainsi l'erreur.

3 Application au contrôle optimal

En sciences de l'ingénieur, l'utilisation des équations adjointes est très forte dans le cadre de la théorie du contrôle optimal et de l'analyse de sensibilité [Lions 1971, Stengel 1994]. Nous en donnons ici les grandes idées, et établissons un lien avec le contrôle des quantités d'intérêt évoqué précédemment.

On considère ici que le modèle décrit dans (1.1), à la fois au niveau de l'opérateur L et du chargement f , dépend d'un ensemble de paramètres noté $\mathbf{p}$ et appartenant à un espace $\mathcal{P}$. Ce modèle se réécrit donc sous la forme :

$$L(\mathbf{p})u = f(\mathbf{p}) \quad \text{ou} \quad A(\mathbf{p}, u, v) = \langle L(\mathbf{p})u - f(\mathbf{p}), v \rangle = 0 \quad \forall v \in \mathcal{F}^* \quad (1.17)$$

La solution $u(\mathbf{p})$ du problème direct (1.17), supposé bien posé, dépend implicitement de la valeur de $\mathbf{p}$.

3.1 Principe de l'état adjoint

On suppose ici que les paramètres dans $\mathbf{p}$ sont variables et on en cherche leur valeur optimale pour atteindre un objectif fixé. En Mécanique des Structures, l'application classique consiste à identifier ou recalibrer ces paramètres par rapport à des mesures expérimentales, afin par exemple de reconstruire numériquement le champ solution u [Nassiopoulos et Bourquin 2010]. Ceci définit un problème inverse [Bonnet et Constantinescu 2005], généralement mal posé, et une démarche pour l'aborder consiste à chercher $\mathbf{p}_{sol}$ tel que :

$$\mathbf{p}_{sol} = \underset{\mathbf{p} \in \mathcal{P}}{\operatorname{argmin}} \mathcal{J}(\mathbf{p}, u(\mathbf{p})) \quad (1.18)$$

où $\mathcal{J}(\mathbf{p}, u(\mathbf{p}))$ est une fonction coût faisant intervenir les données mesurées. La résolution de (1.18) se fait généralement par une méthode de descente nécessitant le gradient (ou la sensibilité) de la fonction coût $\mathcal{J}(\mathbf{p}, u(\mathbf{p}))$ par rapport à l'ensemble $\mathbf{p}$ des paramètres. Ce gradient s'écrit, en supposant $\mathcal{J}$ différentiable :

$$\mathrm{d}_{\mathbf{p}} \mathcal{J}(\mathbf{p}, u(\mathbf{p}); \delta \mathbf{p}) = \mathcal{J}'_{\mathbf{p}}(\mathbf{p}, u(\mathbf{p}); \delta \mathbf{p}) + \mathcal{J}'_u(\mathbf{p}, u(\mathbf{p}); du) \quad (1.19)$$

où $d_{\mathbf{p}}\mathcal{J}$ (resp. $\mathcal{J}'_{\mathbf{p}}$) désigne la dérivée totale (resp. partielle), au sens de Gâteaux, par rapport à $\mathbf{p}$ et dans la direction $\delta\mathbf{p} \in \mathcal{P}_0$, et $du = d_{\mathbf{p}}u(\mathbf{p}, \delta\mathbf{p})$. Ce dernier terme, donnant la sensibilité de u par rapport à $\mathbf{p}$ est le plus complexe et coûteux à calculer en pratique.

Les deux approches naturelles de calcul de $d_{\mathbf{p}}\mathcal{J}$, basées sur la méthode des différences finies ou une différentiation directe de (1.17), sont coûteuses car elles nécessitent de multiples résolutions. Une approche alternative, appelée méthode de l'état adjoint, reformule (1.18) comme un problème de minimisation sous contrainte. Elle introduit alors le lagrangien $\mathcal{L}(\mathbf{p}, u, \lambda) = \mathcal{J}(\mathbf{p}, u) - A(\mathbf{p}, u, \lambda)$ dont la recherche du point de stationnarité $\mathbf{z} \equiv (\mathbf{p}, u, \lambda)$ aboutit au système :

$$\begin{aligned}\mathcal{L}'_{\mathbf{p}}(\mathbf{z}; \delta\mathbf{p}) &= \mathcal{J}'_{\mathbf{p}}(\mathbf{p}, u; \delta\mathbf{p}) - A'_{\mathbf{p}}(\mathbf{p}, u, \lambda; \delta\mathbf{p}) = 0 \quad \forall \delta\mathbf{p} \in \mathcal{P}_0 \\ \mathcal{L}'_u(\mathbf{z}; \delta u) &= \mathcal{J}'_u(\mathbf{p}, u; \delta u) - A(\mathbf{p}, \delta u, \lambda) = 0 \quad \forall \delta u \in \mathcal{U}_0 \\ \mathcal{L}'_{\lambda}(\mathbf{z}; \delta\lambda) &= -A(\mathbf{p}, u, \delta\lambda) = 0 \quad \forall \delta\lambda \in \mathcal{F}^*\end{aligned}\tag{1.20}$$

La dernière équation est l'équation d'état (1.17), tandis que la seconde équation est appelée équation adjointe et peut se réécrire sous forme :

$$\langle L(\mathbf{p})\delta u, \lambda \rangle = \langle \delta u, L^*(\mathbf{p})\lambda \rangle = \mathcal{J}'_u(\mathbf{p}, u; \delta u) \quad \forall \delta u \in \mathcal{U}_0\tag{1.21}$$

Pour un couple $(u, \mathbf{p})$ vérifiant l'équation d'état, on a $\mathcal{L}(\mathbf{p}, u, \lambda) = \mathcal{J}(\mathbf{p}, u)$. Si de plus, λ est solution de l'équation adjointe, alors :

$$d_{\mathbf{p}}\mathcal{J}(\mathbf{p}, u(\mathbf{p}); \delta\mathbf{p}) = d_{\mathbf{p}}\mathcal{L}(\mathbf{p}, u, \lambda; \delta\mathbf{p}) = \mathcal{L}'_{\mathbf{p}}(\mathbf{p}, u, \lambda; \delta\mathbf{p})\tag{1.22}$$

et $\mathcal{L}'_{\mathbf{p}}(\mathbf{p}, u, \lambda; \delta\mathbf{p})$ se calcule assez simplement.

Remarque : La valeur de $d_{\mathbf{p}}\mathcal{J}$ s'annule lorsque la première équation du système (1.20), appelée équation de contrôle, est vérifiée.

Par conséquent, le calcul de $d_{\mathbf{p}}\mathcal{J}(\mathbf{p}, u(\mathbf{p}); \delta\mathbf{p})$ peut se faire à partir de la solution $u(\mathbf{p})$ et de celle d'un problème annexe (problème adjoint), valable pour tout chargement $f(\mathbf{p})$. De plus, le coût de calcul du gradient est indépendant du nombre de paramètres dans $\mathbf{p}$ avec cette approche. C'est pourquoi elle est très utilisée en pratique.

3.2 Lien avec l'étude de quantités d'intérêt

L'étude de la sensibilité d'une quantité d'intérêt Q par rapport à $\mathbf{p}$ peut également être menée avec la vision précédente issue du contrôle optimal et l'introduction d'un lagrangien [Becker et Rannacher 2001]. Il suffit pour cela de définir la minimisation sous contrainte triviale :

$$Q(u(\mathbf{p})) = \min_{u^* \in \mathcal{U}_{sol}} Q(u^*) \quad \text{avec} \quad \mathcal{U}_{sol} = \{w \in \mathcal{U}, A(\mathbf{p}, w, v) = 0 \quad \forall v \in \mathcal{F}^*\}\tag{1.23}$$

L'introduction du lagrangien associé $\mathcal{L}(\mathbf{p}, u, \lambda) = Q(u) - A(\mathbf{p}, u, \lambda)$ et la recherche de son point de stationnarité $\mathbf{z} \equiv (\mathbf{p}, u, \lambda)$ aboutit à un système semblable à (1.20), en remplaçant $\mathcal{J}$ par Q . Le nouveau problème adjoint s'écrit alors (cf. (1.21)) :

$$\langle L(\mathbf{p})\delta u, \lambda \rangle = \langle \delta u, L^*(\mathbf{p})\lambda \rangle = \langle \delta u, Q \rangle \quad \forall \delta u \in \mathcal{U}_0 \quad \implies \quad L^*(\mathbf{p})\lambda = Q\tag{1.24}$$

et est similaire à celui trouvé dans (1.10), pour $\mathbf{p}$ fixé. Sa solution $\lambda(\mathbf{p})$ permet donc de calculer la sensibilité de la quantité d'intérêt vis-à-vis des paramètres, sous la forme :

$$\mathrm{d}_{\mathbf{p}}Q(u(\mathbf{p}); \delta\mathbf{p}) = \mathcal{L}'_{\mathbf{p}}(\mathbf{p}, u, \lambda; \delta\mathbf{p}) = -\mathcal{A}'_{\mathbf{p}}(\mathbf{p}, u, \lambda; \delta\mathbf{p}) \quad (1.25)$$

pour $(u, \mathbf{p})$ vérifiant (1.17). Elle permet en outre, à $\mathbf{p}$ fixé et comme indiqué dans la Section 2.2, d'estimer l'erreur locale de discrétisation.

Contrôle de l'erreur locale de discrétisation

Sommaire

1	Contributions à l'obtention de bornes garanties	14
1.1	Méthode générale utilisant l'ERC	14
1.2	Extension aux problèmes stochastiques	19
1.3	Extension aux problèmes résolus par la XFEM	24
1.4	Extension aux quantités d'intérêt non-linéaires	32
2	Contributions à l'obtention de bornes précises	37
2.1	Nouvelle technique de majoration en espace	37
2.2	Prise en compte des effets d'histoire	40
2.3	Résolution fine du problème adjoint	43
3	Contributions à l'obtention de bornes pratiques	46
3.1	Procédure simplifiée pour le calcul des champs admissibles . .	46
3.2	Technique non-intrusive pour la résolution de l'adjoint	54
4	Conclusions partielles et perspectives	60

Les techniques développées dans cette partie cherchent à répondre aux attentes des industriels pour la conception robuste ; elles visent à construire, pour une large gamme de problèmes de mécanique, un outil de post-traitement fournissant des bornes d'erreur locale de discrétisation à la fois **garanties**, **précises**, et **pratiques** à utiliser en termes d'implémentation et de coût. L'objectif à terme est d'avoir une pleine confiance envers les résultats de simulation et de réaliser le meilleur compromis entre coût et précision.

La bibliographie sur l'estimation d'erreur de discrétisation locale est assez large (voir [Paraschivoiu *et al.* 1997, Rannacher et Suttmeier 1997, Cirak et Ramm 1998, Aubry *et al.* 1999, Ladevèze *et al.* 1999, Oden et Prudhomme 1999] pour les premiers travaux) et est encore un axe de recherche très actif. Néanmoins, une grande majorité des méthodes proposées ne fournit pas de bornes garanties et/ou précises. Un point clé pour obtenir des bornes garanties est, outre l'introduction du problème adjoint, le calcul de champs de contraintes équilibrés qui est un point central du concept d'erreur en relation de comportement (ERC) utilisé dans tous les travaux présentés dans ce chapitre. Pour avoir également des bornes précises, ce qui est souvent vu comme une caractéristique incompatible avec l'aspect garanti (exemple des bornes asymptotiques), l'approche suivie consiste essentiellement à profiter du cadre de l'ERC et résoudre judicieusement le problème adjoint. Enfin, le caractère pratique des outils d'estimation d'erreur proposés ici est obtenu par des techniques numériques adaptées aux outils de calcul actuels.

Dans toute la suite, seule l'erreur de discrétisation de la solution est considérée ; les autres sources d'erreur (code, quadrature, approximation de la géométrie ou du chargement, arrondi numérique,...), généralement négligeables ou contrôlables à partir de solutions manufacturées, ne sont pas prises en compte.

1 Contributions à l'obtention de bornes garanties

1.1 Méthode générale utilisant l'ERC

1.1.1 Cas de l'élasticité linéaire

On considère une structure définie par un domaine borné $\Omega \subset \mathbb{R}^d$ de frontière $\partial\Omega$ (Figure 2.1). Elle est soumise à un chargement statique défini par un déplacement imposé $\mathbf{u}_d$ sur $\partial_1\Omega \subset \partial\Omega$ ($\partial_1\Omega \neq \emptyset$), une densité volumique d'effort $\mathbf{f}_d$ dans Ω et une densité surfacique d'effort $\mathbf{F}_d$ sur $\partial_2\Omega$, avec $\overline{\partial_1\Omega \cup \partial_2\Omega} = \partial\Omega$ et $\partial_1\Omega \cap \partial_2\Omega = \emptyset$. Le matériau composant Ω est supposé élastique linéaire, et on note $\mathbf{K}$ le tenseur de Hooke. Le problème associé consiste à trouver le couple déplacement-contrainte $(\mathbf{u}, \boldsymbol{\sigma})$ vérifiant les contraintes cinématiques, les équations d'équilibre et la relation de comportement qui s'écrivent respectivement :

$$\boldsymbol{\varepsilon} = \frac{1}{2}(\nabla \mathbf{u} + \nabla^T \mathbf{u}) \quad \text{dans } \Omega \quad ; \quad \mathbf{u} = \mathbf{u}_d \quad \text{sur } \partial_1\Omega \quad (2.1)$$

$$\operatorname{div} \boldsymbol{\sigma} + \mathbf{f}_d = \mathbf{0} \quad \text{dans } \Omega \quad ; \quad \boldsymbol{\sigma} \mathbf{n} = \mathbf{F}_d \quad \text{sur } \partial_2\Omega \quad (2.2)$$

$$\boldsymbol{\sigma} = \mathbf{K} \boldsymbol{\varepsilon} \quad \text{dans } \Omega \quad (2.3)$$

ε désignant le tenseur des déformations linéarisées et $\mathbf{n}$ la normale sortante au domaine.

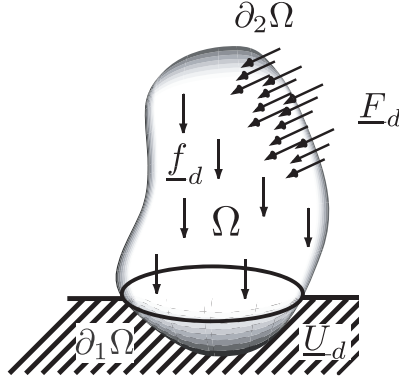


Figure 2.1: Structure et son environnement

On introduit les espaces $\mathcal{U}$ et $\mathcal{S}$ des champs de déplacement cinématiquement admissibles (champs de vecteurs $\mathbf{v} \in [H^1(\Omega)]^d$ vérifiant (2.1)) et des champs de contrainte statiquement admissibles (champs de tenseurs du second ordre symétriques $\pi \in [L^2(\Omega)]^{d(d+1)/2}$ vérifiant (2.2)), respectivement. Le problème peut alors être réécrit sous la forme faible suivante :

$$\text{Trouver } \mathbf{u} \in \mathcal{U} \text{ tel que } A(\mathbf{u}, \mathbf{v}) = a(\mathbf{u}, \mathbf{v}) - l(\mathbf{v}) = 0 \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (2.4)$$

où $\mathcal{U}_0$ est l'espace vectoriel associé à $\mathcal{U}$, et les formes bilinéaires a et linéaires l sont définies par :

$$a(\mathbf{u}, \mathbf{v}) = \int_{\Omega} \mathbf{K} \varepsilon(\mathbf{u}) : \varepsilon(\mathbf{v}) d\Omega \quad ; \quad l(\mathbf{v}) = \int_{\Omega} \mathbf{f}_d \cdot \mathbf{v} d\Omega + \int_{\partial_2 \Omega} \mathbf{F}_d \cdot \mathbf{v} dS \quad (2.5)$$

On utilise aussi les notations :

$$\begin{aligned} \langle \mathbf{u}_1, \mathbf{u}_2 \rangle_K &= \int_{\Omega} \mathbf{K} \varepsilon(\mathbf{u}_1) : \varepsilon(\mathbf{u}_2) d\Omega \quad ; \quad \|\mathbf{u}\|_K^2 = \langle \mathbf{u}, \mathbf{u} \rangle_K \\ (\mathbf{u}_1, \mathbf{u}_2) &= \int_{\Omega} \mathbf{u}_1 \cdot \mathbf{u}_2 d\Omega \quad ; \quad \|\mathbf{u}\|^2 = (\mathbf{u}, \mathbf{u}) \\ \langle \sigma_1, \sigma_2 \rangle_{K^{-1}} &= \int_{\Omega} \mathbf{K}^{-1} \sigma_1 : \sigma_2 d\Omega \quad ; \quad \|\sigma\|_{K^{-1}}^2 = \langle \sigma, \sigma \rangle_{K^{-1}} \quad ; \quad \langle \sigma, \varepsilon \rangle = \int_{\Omega} \sigma : \varepsilon d\Omega \end{aligned} \quad (2.6)$$

En pratique, une approximation $(\mathbf{u}_h, \sigma_h) \in \mathcal{U}_h \times \mathcal{S}_h$ (avec $\sigma_h = \mathbf{K} \varepsilon(\mathbf{u}_h)$) de la solution exacte $(\mathbf{u}, \sigma)$ est calculée par la Méthode des Eléments Finis associée à un maillage spatial $\mathcal{M}_h$. L'erreur de discrétisation qui en résulte est notée $\mathbf{e}_h = \mathbf{u} - \mathbf{u}_h$.

On suppose qu'on sait construire, à partir de $(\mathbf{u}_h, \sigma_h)$ une solution $(\hat{\mathbf{u}}_h, \hat{\sigma}_h) \in \mathcal{U} \times \mathcal{S}$ admissible i.e. vérifiant (2.1) et (2.2). Dans la suite, on fera le choix simple $\hat{\mathbf{u}}_h = \mathbf{u}_h$, et le calcul de $\hat{\sigma}_h$ peut être mené par post-traitement de σ_h avec diverses techniques (voir Section 3.1). On définit alors l'erreur en relation de comportement e_{ERC} :

$$e_{ERC}^2(\mathbf{u}_h, \hat{\sigma}_h) = \frac{1}{2} \|\hat{\sigma}_h - \mathbf{K} \varepsilon(\mathbf{u}_h)\|_{K^{-1}}^2 \quad (2.7)$$

qui est associée aux propriétés suivantes [Prager et Synge 1947, Ladevèze et Pelle 2004] :

$$\begin{aligned} \|\mathbf{u} - \mathbf{u}_h\|_K^2 + \|\sigma - \hat{\sigma}_h\|_{K^{-1}}^2 &= 2e_{ERC}^2(\mathbf{u}_h, \hat{\sigma}_h) \\ \|\sigma - \hat{\sigma}_h^m\|_{K^{-1}}^2 &= \frac{1}{2}e_{ERC}^2(\mathbf{u}_h, \hat{\sigma}_h) \end{aligned} \quad (2.8)$$

avec $\hat{\sigma}_h^m = \frac{1}{2}(\hat{\sigma}_h + \sigma_h)$. Une illustration graphique des relations (2.8) est donnée sur la Fig. 2.2. La première relation mène à la majoration garantie $\|\mathbf{e}_h\|_K \leq \sqrt{2}e_{ERC}(\mathbf{u}_h, \hat{\sigma}_h)$ de l'erreur de discrétisation globale (norme énergétique) ; la qualité de cette majoration dépend de celle de $\hat{\sigma}_h$.

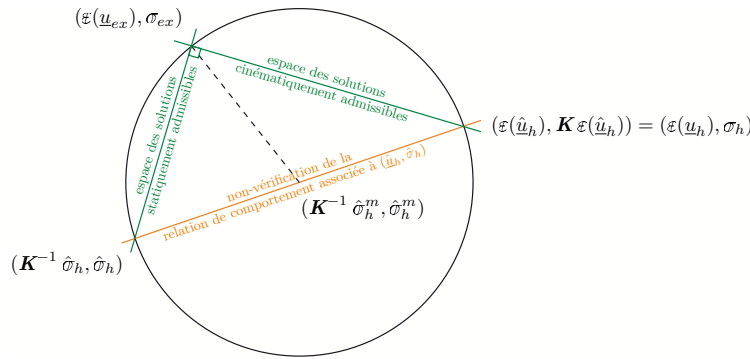


Figure 2.2: Représentation des propriétés de Prager-Synge

On considère une mesure locale $Q(\mathbf{u}) - Q(\mathbf{u}_h)$ de l'erreur de discrétisation, définie sur une quantité d'intérêt Q . Généralement, la quantité d'intérêt est écrite sous la forme globale :

$$Q(\mathbf{u}) = \int_{\Omega} (\varepsilon(\mathbf{u}) : \sigma_{\Sigma} + \mathbf{u} \cdot \mathbf{f}_{\Sigma}) d\Omega + \int_{\partial_2 \Omega} \mathbf{u} \cdot \mathbf{F}_{\Sigma} dS \quad (2.9)$$

où σ_{Σ} , $\mathbf{f}_{\Sigma}$ et $\mathbf{F}_{\Sigma}$ sont des fonctions d'extraction, données explicitement ou implicitement, qui s'interprètent respectivement comme une précontrainte, un pré-chargement volumique et un pré-chargement surfacique. On introduit le problème adjoint (cf. 1.10), qui consiste à trouver le déplacement $\tilde{\mathbf{u}} \in \mathcal{U}_0$ vérifiant :

$$a(\mathbf{v}, \tilde{\mathbf{u}}) = Q(\mathbf{v}) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (2.10)$$

L'opérateur $\mathbf{K}$ impliqué ici étant auto-adjoint, le problème adjoint a une structure similaire au problème direct (2.4) ; la contrainte liée à $\tilde{\mathbf{u}}$ est notée $\tilde{\sigma} = \mathbf{K}\varepsilon(\tilde{\mathbf{u}})$ et l'équation d'équilibre associée s'écrit :

$$\int_{\Omega} \varepsilon(\mathbf{v}) : (\tilde{\sigma} - \sigma_{\Sigma}) d\Omega = \int_{\Omega} \mathbf{v} \cdot \mathbf{f}_{\Sigma} d\Omega + \int_{\partial_2 \Omega} \mathbf{v} \cdot \mathbf{F}_{\Sigma} dS \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (2.11)$$

Une approximation $(\tilde{\mathbf{u}}_h, \tilde{\sigma}_h)$ de la solution adjointe peut être calculée par la MEF (avec le même maillage $\mathcal{M}_h$), et une solution admissible $(\hat{\tilde{\mathbf{u}}}_h, \hat{\tilde{\sigma}}_h)$, avec le choix $\hat{\tilde{\mathbf{u}}}_h = \tilde{\mathbf{u}}_h$, est obtenue par post-traitement. On obtient alors, à partir de (1.15) et (2.8), la majoration garantie :

$$\begin{aligned} |Q(\mathbf{u}) - Q(\mathbf{u}_h)| &= |Q(\mathbf{u} - \mathbf{u}_h)| = a(\mathbf{u} - \mathbf{u}_h, \tilde{\mathbf{u}} - \tilde{\mathbf{u}}_h) \\ &\leq \|\mathbf{u} - \mathbf{u}_h\|_K \cdot \|\tilde{\mathbf{u}} - \tilde{\mathbf{u}}_h\|_K \\ &\leq 2e_{ERC}(\mathbf{u}_h, \hat{\sigma}_h) \cdot e_{ERC}(\tilde{\mathbf{u}}_h, \hat{\tilde{\sigma}}_h) \end{aligned} \quad (2.12)$$

Un résultat plus précis peut être obtenu en introduisant $\hat{\sigma}_h^m$ et les propriétés des champs admissibles. On peut en effet écrire :

$$\begin{aligned} a(\mathbf{u} - \mathbf{u}_h, \tilde{\mathbf{u}} - \tilde{\mathbf{u}}_h) &= \langle \sigma - \sigma_h, \tilde{\sigma} - \tilde{\sigma}_h \rangle_{K^{-1}} \\ &= \langle \sigma - \sigma_h, \hat{\tilde{\sigma}}_h - \tilde{\sigma}_h \rangle_{K^{-1}} \\ &= \langle \sigma - \hat{\sigma}_h^m, \hat{\tilde{\sigma}}_h - \tilde{\sigma}_h \rangle_{K^{-1}} + \langle \hat{\sigma}_h^m - \sigma_h, \hat{\tilde{\sigma}}_h - \tilde{\sigma}_h \rangle_{K^{-1}} \end{aligned} \quad (2.13)$$

et une nouvelle majoration se met sous la forme :

$$|Q(\mathbf{u}) - Q(\mathbf{u}_h) - Q_{corr}| \leq e_{ERC}(\mathbf{u}_h, \hat{\sigma}_h) \cdot e_{ERC}(\tilde{\mathbf{u}}_h, \hat{\tilde{\sigma}}_h) \quad (2.14)$$

où $Q_{corr} = \langle \hat{\sigma}_h^m - \sigma_h, \hat{\tilde{\sigma}}_h - \tilde{\sigma}_h \rangle_{K^{-1}}$, entièrement calculable, peut être interprété comme un terme de correction de $Q(\mathbf{u}_h)$. Avec cette nouvelle majoration, non classique, la qualité de l'estimation de l'erreur sur Q est similaire à celle utilisant l'égalité du parallélogramme [Oden et Prudhomme 1999].

1.1.2 Cas des problèmes dissipatifs

Le concept d'erreur en relation de comportement introduit précédemment pour les problèmes d'élasticité peut être étendu à des problèmes plus complexes, notamment en introduisant le concept d'erreur en dissipation basé sur des pseudo-potentiels [Ladevèze et Moës 1998]. Associé à la résolution d'un problème adjoint, et en exploitant les propriétés de convexité des lois de comportement, ce concept a permis l'obtention des premières bornes d'erreur sur Q garanties pour les problèmes dissipatifs [Chamoin et Ladevèze 2007, Waeytens *et al.* 2012], avec prise en compte des erreurs de discrétisation spatiale et temporelle.

On se place ici dans le cas d'un problème de dynamique sur l'intervalle temporel $[0, T]$, avec un matériau standard viscoélastique sans variable interne supplémentaire (voir [Ladevèze 2006 - 2008] pour le cas général des problèmes avec relation de comportement monotone). En séparant la déformation sous la forme $\varepsilon(\mathbf{u}) = \varepsilon^e + \varepsilon^p$, avec ε^e (resp. ε^p) la partie élastique (resp. anélastique), le comportement $\sigma = \mathcal{A}(\dot{\varepsilon}(\tau), \tau \leq t)$ du matériau est représenté par deux relations :

- l'équation d'état $\sigma = \frac{\partial \phi}{\partial \varepsilon^e} = \mathbf{K} \varepsilon^e$ avec ϕ le potentiel thermodynamique ;
- la loi d'évolution $\dot{\varepsilon}^p = \frac{\partial \psi}{\partial \sigma} = \mathbf{B} \sigma$, liée à la dissipation intrinsèque $d = \sigma : \dot{\varepsilon}^p$, avec ψ un pseudo-potentiel de dissipation et $\mathbf{B}$ un opérateur symétrique défini positif.

Le problème direct, constitué en outre des équations d'équilibre :

$$\int_{\Omega} \left(\sigma : \varepsilon(\dot{\mathbf{v}}) + \rho \frac{\partial^2 \mathbf{u}}{\partial t^2} \cdot \dot{\mathbf{v}} \right) d\Omega = \int_{\Omega} \mathbf{f}_d \cdot \dot{\mathbf{v}} d\Omega + \int_{\partial_2 \Omega} \mathbf{F}_d \cdot \dot{\mathbf{v}} dS \quad \forall \mathbf{v} \in \mathcal{U}_0, \forall t \in [0, T] \quad (2.15)$$

et des conditions initiales nulles sur $\mathbf{u}$, $\dot{\mathbf{u}}$, σ et ε^p , peut s'écrire sous une forme faible intégrée en espace et temps du type :

$$\text{Trouver } \mathbf{u} \in \mathcal{U} \text{ tel que } a(\mathbf{u}, \dot{\mathbf{v}}) = l(\dot{\mathbf{v}}) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (2.16)$$

Une solution approchée de ce problème, calculée par la MEF en espace (maillage $\mathcal{M}_h$) et un schéma en temps (maillage $\mathcal{M}_{\Delta t}$), est notée $(\mathbf{u}_{h,\Delta}, \sigma_{h,\Delta})$. On introduit de plus les notations :

$$\begin{aligned} ((\mathbf{u}_1, \mathbf{u}_2)) &= \int_0^T (\mathbf{u}_1, \mathbf{u}_2) dt \quad ; \quad |||\mathbf{u}|||^2 = ((\mathbf{u}, \mathbf{u})) \\ \langle\langle \sigma_1, \sigma_2 \rangle\rangle_X &= \int_0^T \langle \sigma_1, \sigma_2 \rangle_X dt \quad ; \quad |||\sigma|||_X^2 = \langle\langle \sigma, \sigma \rangle\rangle_X \quad ; \quad \langle\langle \sigma, \varepsilon \rangle\rangle = \int_0^T \langle \sigma, \varepsilon \rangle dt \end{aligned} \quad (2.17)$$

Pour appliquer le concept d'ERC dans ce cadre, la notion de solution admissible est redéfinie. Une telle solution doit vérifier, en plus des équations de liaison et d'équilibre, l'équation d'état et les conditions initiales ; la seule équation du problème qui n'est pas vérifiée est l'équation d'évolution. On introduit alors, pour une solution admissible quelconque $(\hat{\mathbf{u}}, \hat{\sigma})$, l'erreur en dissipation comme une mesure de la non-vérification de la loi d'évolution :

$$e_{ERC}^2(\dot{\hat{\varepsilon}}^p, \hat{\sigma}) = \int_0^T \eta_t(\dot{\hat{\varepsilon}}^p, \hat{\sigma}) dt \quad \text{avec} \quad \eta_t(\dot{\hat{\varepsilon}}^p, \hat{\sigma}) \equiv \frac{1}{2} ||\dot{\hat{\varepsilon}}^p - \mathbf{B}\hat{\sigma}||_{B^{-1}}^2 \quad (2.18)$$

De façon plus générale, $\eta_t(\dot{\hat{\varepsilon}}^p, \hat{\sigma})$ est construit à partir de ψ et de son potentiel dual (au sens de la transformée de Legendre-Fenchel). On a alors les relations [Ladevèze et Pelle 2004] :

$$\begin{aligned} ||\dot{\hat{\varepsilon}}^p - \dot{\hat{\varepsilon}}^p||_{B^{-1}}^2 + ||\sigma - \hat{\sigma}||_B^2 + \frac{d}{dt} E_l(\sigma - \hat{\sigma}) &= 2\eta_t(\dot{\hat{\varepsilon}}^p, \hat{\sigma}) \\ ||\sigma - \hat{\sigma}^m||_B^2 + \frac{1}{2} \frac{d}{dt} E_l(\sigma - \hat{\sigma}) &= \frac{1}{2} \eta_t(\dot{\hat{\varepsilon}}^p, \hat{\sigma}) \end{aligned} \quad (2.19)$$

avec $\hat{\sigma}^m = \frac{1}{2}(\hat{\sigma} + \mathbf{B}^{-1}\dot{\hat{\varepsilon}}^p)$ et $E_l = (\rho(\dot{\mathbf{u}} - \dot{\hat{\mathbf{u}}}), \dot{\mathbf{u}} - \dot{\hat{\mathbf{u}}}) + ||\sigma - \hat{\sigma}||_{K^{-1}}^2$ l'énergie libre. Une intégration en temps donne la relation fondamentale :

$$|||\sigma - \hat{\sigma}|||_B^2 + \frac{1}{2} E_l(\sigma - \hat{\sigma})|_T = \frac{1}{2} e_{ERC}^2(\dot{\hat{\varepsilon}}^p, \hat{\sigma}) \quad (2.20)$$

qui connecte la mesure de l'erreur en dissipation (calculable) à une erreur en solution constituée d'une partie liée à l'élasticité à l'instant final T , et d'une partie impliquant toute l'histoire du chargement.

La construction d'une solution admissible $(\hat{\mathbf{u}}_{h,\Delta}, \hat{\sigma}_{h,\Delta})$ à partir de $(\mathbf{u}_{h,\Delta}, \sigma_{h,\Delta})$ est menée de la façon suivante : (i) on définit $\ddot{\hat{\mathbf{u}}}_{h,\Delta}$ comme l'interpolation linéaire de l'accélération approchée ; (ii) on construit $\hat{\sigma}_{h,\Delta}$ en équilibre avec $\ddot{\hat{\mathbf{u}}}_{h,\Delta}$ et le chargement extérieur $(\mathbf{f}_d, \mathbf{F}_d)$ par les procédures classiques ; (iii) on en déduit $\hat{\varepsilon}_{h,\Delta}^e = \mathbf{K}^{-1}\hat{\sigma}_{h,\Delta}$ et $\hat{\varepsilon}_{h,\Delta}^p = \hat{\varepsilon}_{h,\Delta} - \hat{\varepsilon}_{h,\Delta}^e$.

On définit la quantité d'intérêt sous la forme :

$$Q(\dot{\mathbf{u}}) = \int_0^T \left(\int_{\Omega} (\varepsilon(\dot{\mathbf{u}}) : \sigma_{\Sigma} + \dot{\mathbf{u}} \cdot \rho \mathbf{\Gamma}_{\Sigma}) d\Omega + \int_{\partial_2 \Omega} \dot{\mathbf{u}} \cdot \mathbf{F}_{\Sigma} dS \right) dt \quad (2.21)$$

où Γ_Σ correspond à une pré-accélération. On introduit alors le problème adjoint associé (cf. 1.10) qui consiste à trouver le déplacement $\tilde{\mathbf{u}} \in \mathcal{U}_0$ vérifiant :

$$a(\mathbf{v}, \dot{\tilde{\mathbf{u}}}) = a^*(\dot{\mathbf{v}}, \tilde{\mathbf{u}}) = Q(\dot{\mathbf{v}}) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (2.22)$$

Pour $\sigma|_{t=0} = \tilde{\sigma}|_{t=T} = 0$, on a :

$$\langle \langle \dot{\tilde{\mathbf{e}}}, \tilde{\sigma} \rangle \rangle = \langle \langle \mathbf{K}^{-1} \dot{\tilde{\sigma}} + \mathbf{B} \sigma, \tilde{\sigma} \rangle \rangle = -\langle \langle \sigma, \mathbf{K}^{-1} \dot{\tilde{\sigma}} - \mathbf{B} \tilde{\sigma} \rangle \rangle \quad (2.23)$$

ce qui montre que le problème adjoint est inverse en temps, avec des conditions finales nulles, et que la loi d'évolution s'écrit $\dot{\tilde{\mathbf{e}}}^p = -\mathbf{B} \tilde{\sigma}$. Cependant, ce problème a une structure identique au problème primal si on le décrit avec la variable temporelle $\tau = T - t$; seul le chargement $(\Gamma_\Sigma, \sigma_\Sigma, \mathbf{F}_\Sigma)$ est nouveau. Le problème adjoint est résolu de façon approchée sur le maillage $\mathcal{M}_h \times \mathcal{M}_{\Delta t}$ et une solution admissible $(\hat{\mathbf{u}}_{h,\Delta}, \hat{\sigma}_{h,\Delta})$ est obtenue par post-traitement. On obtient alors :

$$\begin{aligned} Q(\dot{\mathbf{u}}) - Q(\dot{\hat{\mathbf{u}}}_{h,\Delta}) &= \langle \langle \dot{\tilde{\mathbf{e}}} - \dot{\hat{\mathbf{e}}}_{h,\Delta}, \tilde{\sigma}_{h,\Delta} \rangle \rangle + ((\dot{\mathbf{u}} - \dot{\hat{\mathbf{u}}}_{h,\Delta}, \rho \ddot{\hat{\mathbf{u}}}_{h,\Delta})) \\ &= \langle \langle \dot{\tilde{\mathbf{e}}} - \mathbf{K}^{-1} \dot{\hat{\sigma}}_{h,\Delta} - \mathbf{B} \hat{\sigma}_{h,\Delta}, \tilde{\sigma}_{h,\Delta} \rangle \rangle + ((\dot{\mathbf{u}} - \dot{\hat{\mathbf{u}}}_{h,\Delta}, \rho \ddot{\hat{\mathbf{u}}}_{h,\Delta})) + C_{h,\Delta} \\ &= \langle \langle \sigma - \hat{\sigma}_{h,\Delta}, -\mathbf{K}^{-1} \dot{\hat{\sigma}}_{h,\Delta} + \mathbf{B} \hat{\sigma}_{h,\Delta} \rangle \rangle + ((\dot{\mathbf{u}} - \dot{\hat{\mathbf{u}}}_{h,\Delta}, \rho \ddot{\hat{\mathbf{u}}}_{h,\Delta})) + C_{h,\Delta} \\ &= \langle \langle \sigma - \hat{\sigma}_{h,\Delta}, -\dot{\hat{\mathbf{e}}}_{h,\Delta} + \dot{\hat{\mathbf{e}}}_{h,\Delta}^p + \mathbf{B} \hat{\sigma}_{h,\Delta} \rangle \rangle + ((\dot{\mathbf{u}} - \dot{\hat{\mathbf{u}}}_{h,\Delta}, \rho \ddot{\hat{\mathbf{u}}}_{h,\Delta})) + C_{h,\Delta} \end{aligned} \quad (2.24)$$

avec $C_h = \langle \langle \mathbf{K}^{-1} \dot{\hat{\sigma}}_{h,\Delta} + \mathbf{B} \hat{\sigma}_{h,\Delta} - \dot{\hat{\mathbf{e}}}_{h,\Delta}, \tilde{\sigma}_{h,\Delta} \rangle \rangle$. En notant que :

$$((\dot{\mathbf{u}} - \dot{\hat{\mathbf{u}}}_{h,\Delta}, \rho \ddot{\hat{\mathbf{u}}}_{h,\Delta})) = -((\ddot{\mathbf{u}} - \ddot{\hat{\mathbf{u}}}_{h,\Delta}, \rho \dot{\hat{\mathbf{u}}}_{h,\Delta})) = \langle \langle \sigma - \hat{\sigma}_{h,\Delta}, \dot{\hat{\mathbf{e}}}_{h,\Delta} \rangle \rangle \quad (2.25)$$

l'erreur de discrétisation sur Q s'écrit :

$$\begin{aligned} Q(\dot{\mathbf{u}}) - Q(\dot{\hat{\mathbf{u}}}_{h,\Delta}) &= \langle \langle \sigma - \hat{\sigma}_{h,\Delta}, \dot{\hat{\mathbf{e}}}_{h,\Delta}^p + \mathbf{B} \hat{\sigma}_{h,\Delta} \rangle \rangle + C_{h,\Delta} \\ &= \langle \langle \sigma - \hat{\sigma}_{h,\Delta}^m, \dot{\hat{\mathbf{e}}}_{h,\Delta}^p + \mathbf{B} \hat{\sigma}_{h,\Delta} \rangle \rangle + C'_{h,\Delta} \end{aligned} \quad (2.26)$$

et on obtient la majoration garantie :

$$\begin{aligned} |Q(\dot{\mathbf{u}}) - Q(\dot{\hat{\mathbf{u}}}_{h,\Delta}) - Q_{corr}| &\leq |||\sigma - \hat{\sigma}_{h,\Delta}^m|||_{B \cdot} |||\mathbf{B}^{-1} \dot{\hat{\mathbf{e}}}_{h,\Delta}^p + \hat{\sigma}_{h,\Delta}|||_B \\ &\leq e_{ERC}(\dot{\hat{\mathbf{e}}}_{h,\Delta}^p, \hat{\sigma}_{h,\Delta}) \cdot e_{ERC}(\dot{\hat{\mathbf{e}}}_{h,\Delta}^p, \hat{\sigma}_{h,\Delta}) \end{aligned} \quad (2.27)$$

avec $Q_{corr} = C'_{h,\Delta}$ un terme correctif calculable. Ce résultat, qui n'utilise pas de propriété d'orthogonalité de Galerkin (voir Section 2.3), permet d'encadrer la valeur exacte $Q(\dot{\mathbf{u}})$ de la quantité d'intérêt (ou l'erreur $Q(\dot{\mathbf{u}}) - Q(\dot{\hat{\mathbf{u}}}_{h,\Delta})$).

1.2 Extension aux problèmes stochastiques

A présent, de nombreux modèles impliquant des paramètres stochastiques sont utilisés pour prendre en compte l'aléa ou les méconnaissances dans les simulations [Schueller 1997, Matthies et Bucher 1999, Sudret *et al.* 2004]. La vérification de tels modèles s'est essentiellement concentrée sur l'erreur globale [Ghanem et Pelissetti 2002, Ladevèze 2003]. Pour l'estimation d'erreur locale, les méthodes proposées

jusqu'à présent [Oden *et al.* 2005, Ladevèze et Florentin 2006] s'appliquaient à des quantités d'intérêt spécifiques et ne fournissaient pas des bornes d'erreur strictes, mais seulement des indicateurs d'erreur obtenus avec des arguments heuristiques. Dans [Chamoïn *et al.* 2012], la démarche présentée dans la Section 1.1 a été étendue aux modèles stochastiques. Deux aspects sont alors revisités : (i) la construction des champs admissibles dans un cadre stochastique, (ii) la séparation des deux sources d'erreur (discrétisations dans l'espace physique et dans l'espace stochastique) pour mener la procédure d'adaptation efficacement.

1.2.1 Calcul des bornes

On considère ici un matériau élastique linéaire dont les propriétés fluctuent aléatoirement de sorte que le tenseur de Hooke est modélisé par un champ aléatoire $\mathbf{K}(\mathbf{x}, \theta) \in [L^2(\Theta, C^0(\Omega))]^{d^4}$; $(\Theta, \mathcal{F}, P)$ est l'espace probabiliste de Kolmogorov, avec Θ l'ensemble des possibles, $\mathcal{F}$ une σ -algèbre des événements (sous-ensembles de Θ), et $P : \mathcal{F} \rightarrow [0, 1]$ une mesure de probabilité. On suppose que le champ $\mathbf{K}(\mathbf{x}, \theta)$ est borné et uniformément coercif, i.e $\exists(K_{min}, K_{max}) \in]0, +\infty[^2$ tel que :

$$0 < K_{min} \leq |\mathbf{K}(\mathbf{x}, \theta)| \leq K_{max} \quad \forall \mathbf{x} \in \Omega, \text{ presque sûrement} \quad (2.28)$$

En pratique, en suivant la décomposition de Karhunen-Loeve, la description stochastique de $\mathbf{K}$ est limitée à un nombre fini M de variables stochastiques indépendantes $\xi_i(\theta) : \Theta \rightarrow \mathbb{R}$ telles que :

$$\mathbf{K}(\mathbf{x}, \theta) \approx \overline{\mathbf{K}}(\mathbf{x}) + \sum_{i=1}^M \sqrt{\lambda_i} \xi_i(\theta) \mathbf{Z}_i(\mathbf{x}) \quad (2.29)$$

où $\overline{\mathbf{K}}(\mathbf{x}) = E_{\Theta}[\mathbf{K}(\mathbf{x})] = \int_{\Theta} \mathbf{K}(\mathbf{x}) dP$ est la valeur moyenne ou espérance de $\mathbf{K}(\mathbf{x})$, tandis que les couples $\{\mathbf{Z}_i, \lambda_i\}$ sont les vecteurs/valeurs propres de l'opérateur de covariance associé à $\mathbf{K}$. On définit le produit scalaire L^2 sur Θ :

$$\langle \alpha_1(\theta), \alpha_2(\theta) \rangle \equiv \int_{\Theta} \alpha_1(\theta) \alpha_2(\theta) dP \quad (2.30)$$

ainsi que les notations suivantes sur $\Omega \times \Theta$:

$$\begin{aligned} \langle \mathbf{u}_1, \mathbf{u}_2 \rangle_{K, \Theta} &= E_{\Theta}[\langle \mathbf{u}_1, \mathbf{u}_2 \rangle_K] & ; & \quad \|\mathbf{u}\|_{K, \Theta}^2 = \langle \mathbf{u}, \mathbf{u} \rangle_{K, \Theta} \\ \langle \sigma_1, \sigma_2 \rangle_{K^{-1}, \Theta} &= E_{\Theta}[\langle \sigma_1, \sigma_2 \rangle_{K^{-1}}] & ; & \quad \|\sigma\|_{K^{-1}, \Theta}^2 = \langle \sigma, \sigma \rangle_{K^{-1}, \Theta} \end{aligned} \quad (2.31)$$

Les nouvelles équations d'admissibilité cinématique et statique s'écrivent alors respectivement :

$$\mathbf{u} \in \mathcal{U} \quad ; \quad \mathbf{u}|_{\partial_1 \Omega} = \mathbf{u}_d \quad \text{presque sûrement} \quad (2.32)$$

$$\sigma \in \mathcal{S} \quad ; \quad E_{\Theta} \left[\int_{\Omega} \sigma : \varepsilon(\mathbf{v}) d\Omega - \int_{\Omega} \mathbf{f}_d \cdot \mathbf{v} d\Omega - \int_{\partial_2 \Omega} \mathbf{F}_d \cdot \mathbf{v} dS \right] = 0 \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (2.33)$$

où $\mathcal{U} = [L^2(\Theta, H^1(\Omega))]^d$ et $\mathcal{S} = \left\{ \pi ; \pi = \pi^T, \pi \in [L^2(\Theta, L^2(\Omega))]^{d^2} \right\}$. Une solution approchée $(\mathbf{u}_{h,s}, \sigma_{h,s})$, avec $\sigma_{h,s} = \mathbf{K} \varepsilon(\mathbf{u}_{h,s})$, est calculée par la MEF ; sans perte

de généralité, on considère sur le domaine stochastique une technique non-intrusive basée sur l'interpolation (paramètre de discrétisation s) d'un ensemble de réalisations calculées. La solution approchée $\mathbf{u}_{h,s}$ est donc écrite sous la forme :

$$\mathbf{u}_{h,s}(\mathbf{x}, \theta) = \sum_k \mathbf{u}_{h,s}^k(\mathbf{x}) \cdot \Psi_k(\theta) \quad ; \quad \sigma_{h,s}(\mathbf{x}, \theta) = \sum_k \sigma_{h,s}^k(\mathbf{x}) \cdot \Psi_k(\theta) \quad (2.34)$$

où Ψ_k est la base polynomiale associée à l'ensemble $\{\xi_i(\theta)\}_{i=1}^M$ des variables aléatoires, définie par $\Psi_k = \prod_{i=1}^M H_{k_i}(\xi_i)$ avec $H_{k_i}(\xi_i)$ des polynômes orthogonaux au sens du produit scalaire défini dans (2.30).

La définition de l'ERC dans le cadre stochastique :

$$e_{ERC}(\hat{\mathbf{u}}, \hat{\sigma}) = \|\hat{\sigma} - \mathbf{K}\varepsilon(\hat{\mathbf{u}})\|_{\mathbf{K}^{-1}, \Theta} \quad (2.35)$$

associée à une solution admissible au sens de (2.32) et (2.33) permet d'étendre naturellement les propriétés vues dans la Section 1.1.1, notamment :

$$\|\sigma - \hat{\sigma}^m\|_{\mathbf{K}^{-1}, \Theta}^2 = \frac{1}{2} e_{ERC}^2(\hat{\mathbf{u}}, \hat{\sigma}) \quad (2.36)$$

Une solution admissible $(\hat{\mathbf{u}}_{h,s}, \hat{\sigma}_{h,s})$ est calculée par post-traitement de la solution approchée $(\mathbf{u}_{h,s}, \sigma_{h,s})$ à disposition. D'une part, on prend par simplicité $\hat{\mathbf{u}}_{h,s} = \mathbf{u}_{h,s}$. D'autre part, à partir des composantes $\sigma_{h,s}^k(\mathbf{x})$, on peut construire les champs admissibles associés $\hat{\sigma}_{h,s}^k(\mathbf{x})$ avec les techniques classiquement utilisées dans le cas déterministe [Ladevèze *et al.* 2010, Pled *et al.* 2011]. Néanmoins, le champ $\sum_k \hat{\sigma}_{h,s}^k(\mathbf{x}) \cdot \Psi_k(\theta)$ n'est généralement pas admissible au sens de (2.33), car le chargement $(\mathbf{f}_d, \mathbf{F}_d)$ est généralement constant ou multi-linéaire par rapport à l'ensemble $\{\xi_i(\theta)\}_{i=1}^M$ des variables aléatoires tandis que $\Psi_k(\theta)$ dépend non-linéairement de ces variables.

On propose donc d'introduire une base linéaire dédiée et de projeter $\sigma_{h,s}(\mathbf{x}, \theta)$ sur cette base avant d'appliquer les techniques de construction de champ admissible. On utilise la base EF classique des fonctions de forme N_k linéaires par morceaux ; chaque direction de la grille régulière associée correspond à une variable aléatoire ξ_i . On obtient alors, après calcul des coefficients $\sigma_{h,s}^{N_k}$ par projection, une nouvelle contrainte approchée et la solution admissible associée :

$$\sigma_{h,s}^P(\mathbf{x}, \theta) = \sum_k \sigma_{h,s}^{N_k}(\mathbf{x}) \cdot N_k(\theta) \quad \implies \quad \hat{\sigma}_{h,s}(\mathbf{x}, \theta) = \sum_k \hat{\sigma}_{h,s}^{N_k}(\mathbf{x}) \cdot N_k(\theta) \quad (2.37)$$

Le champ $\hat{\sigma}_{h,s}$ est admissible par combinaison linéaire de champs déterministes admissibles $\hat{\sigma}_{h,s}^{N_k}(\mathbf{x})$.

Pour l'estimation de l'erreur sur une quantité d'intérêt, définie de façon globale sous la forme :

$$Q(\mathbf{u}) = \mathbb{E}_\Theta \left[\int_\Omega \left\{ \varepsilon(\mathbf{u}) : \tilde{\sigma}_\Sigma + \mathbf{u} \cdot \tilde{\mathbf{f}}_\Sigma \right\} d\Omega \right] \quad (2.38)$$

on calcule une solution approchée du problème adjoint (stochastique) associé puis une solution $(\tilde{\mathbf{u}}_{h,s}, \hat{\tilde{\sigma}}_{h,s})$ admissible. On obtient alors la majoration d'erreur suivante,

dont la démonstration est identique à celle de la Section 1.1.1 :

$$\begin{aligned} |Q(\mathbf{u}) - Q(\mathbf{u}_{h,s}) - Q_{corr}| &= |\langle \langle \sigma - \hat{\sigma}_{h,s}, \hat{\hat{\sigma}}_{h,s} - \tilde{\sigma}_{h,s} \rangle \rangle_{K^{-1}, \Theta}| \\ &\leq e_{ERC}(\mathbf{u}_{h,s}, \hat{\sigma}_{h,s}) \cdot e_{ERC}(\tilde{\mathbf{u}}_{h,s}, \hat{\hat{\sigma}}_{h,s}) \end{aligned} \quad (2.39)$$

Le calcul de la borne précédente est facile à implémenter, avec des calculs analytiques possibles dans la dimension stochastique. Une autre majoration consiste à appliquer l'inégalité de Cauchy-Schwarz sur le domaine spatial seulement, avant d'intégrer sur la dimension stochastique ; cela fournit une borne plus précise mais plus complexe à calculer.

Nous illustrons les performances de la méthode sur le cas-test décrit sur la Figure 2.3. Le module de Young est supposé mal connu dans un sous-domaine Ω_1 de la structure et est décrit par une variable aléatoire $E_1(\theta)$ de la forme :

$$E_1(\theta) = \overline{E_1} \cdot [1 + \delta_1 g(\xi(\theta))] \quad \text{avec } g(x) = \frac{2 \arcsin(\text{Erf}(\frac{x}{\sqrt{2}}))}{\sqrt{\pi^2 - 8}} \quad (2.40)$$

où $\overline{E_1}$ (resp. δ_1) est la moyenne (resp. variance) de $E_1(\theta)$ et $\xi(\theta)$ est une variable aléatoire centrée gaussienne. La fonction non-linéaire g permet d'avoir une densité de probabilité $E_1(\theta)$ à support borné. Dans la zone Ω_2 complémentaire de Ω_1 , le module de Young E_2 est déterministe. La structure est encastree à sa base, et le chargement consiste en une densité de force de traction $\mathbf{F}_d$ sur le bord supérieur droit et un déplacement $\mathbf{u}_d$ sur le bord du haut. On prend $\overline{E_1} = 1$, $\delta_1 = 0, 1$, $E_2 = 2$, $\mathbf{F}_d \cdot \mathbf{x} = -1, 5$, $\mathbf{u}_d \cdot \mathbf{y} = -2$, $a = 25$, $b = 20$, $c = 10$ et $d = 10$.

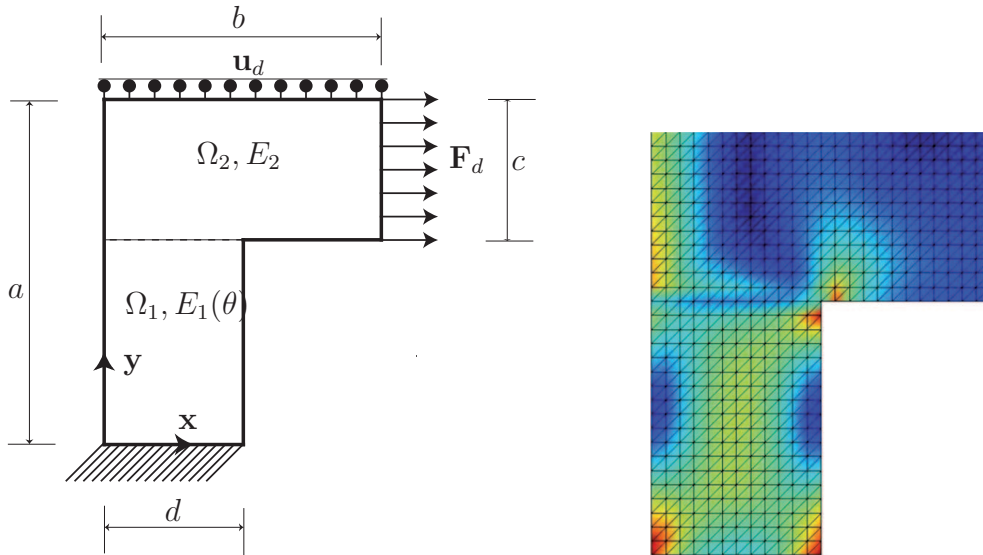


Figure 2.3: Définition du problème (gauche), et contrainte de Mises EF pour une réalisation de E_1 (droite)

La quantité d'intérêt étudiée est la moyenne du déplacement horizontal sur la zone d'application de $\mathbf{F}_d$:

$$Q(\mathbf{u}) = \mathbb{E}_\Theta \left[\int_\Omega \mathbf{u} \cdot \mathbf{f}_\Sigma \, d\Omega \right] \quad ; \quad \mathbf{f}_\Sigma = \frac{1}{c} \delta_{x=b} \mathbf{x} \quad (2.41)$$

$\delta_{x=b}$ étant la fonction de Dirac qui localise Q dans l'espace physique. Une valeur quasi-exacte $Q(\mathbf{u}) = -168.819$ de cette quantité est calculée avec un calcul très précis (maillage EF très fin, méthode de Monte-Carlo à convergence). Une valeur approchée $Q(\mathbf{u}_{h,s}) = -185.591$ est quant à elle calculée avec un calcul plus grossier (Figure 2.3) : en espace, on prend un maillage uniforme d'éléments finis $Q4$ avec $h = 12$ éléments sur le côté $y = 0$ de la structure ; dans la dimension stochastique, on prend $s = 11$ points de calcul. En écrivant l'encadrement donné par (2.39) sous la forme :

$$Q^- \leq Q(\mathbf{u}) \leq Q^+ \quad ; \quad \eta_Q^\pm = Q^\pm / Q(\mathbf{u}) \quad (2.42)$$

on obtient les bornes adimensionnées $\eta_Q^- = 0,9674$ et $\eta_Q^+ = 1,2313$.

L'évolution de ces bornes en fonction de la finesse du maillage spatial (variation de h) et du nombre de réalisations (variation de s) est présentée dans le Tableau 2.1.

s	h	η_Q^-	η_Q^+
11	6	0,542	1,282
11	12	0,834	1,098
11	24	0,938	1,026
11	48	0,974	1,004
11	72	0,988	1,005

s	h	η_{loc}^-	η_{loc}^+
3	12	0,817	1,099
5	12	0,831	1,098
11	12	0,834	1,098
21	12	0,832	1,095
41	12	0,832	1,095
81	12	0,832	1,095

Tableau 2.1: Evolution de (η^-, η^+) en fonction de h (gauche) et s (droite)

1.2.2 Estimation des différentes sources d'erreur

L'erreur de discrétisation $Q(\mathbf{u}) - Q(\mathbf{u}_{h,s})$ provient de deux sources : (i) la discrétisation en espace avec un maillage éléments finis ; (ii) la discrétisation du domaine stochastique. On estime ici les contributions de ces deux sources, afin d'obtenir une information pertinente pour mener les procédures adaptatives.

L'erreur locale peut être réécrite sous la forme :

$$Q(\mathbf{u}) - Q(\mathbf{u}_{h,s}) = [Q(\mathbf{u}) - Q(\mathbf{u}_h)] + [Q(\mathbf{u}_h) - Q(\mathbf{u}_{h,s})] = \Delta Q_{spa} + \Delta Q_{sto} \quad (2.43)$$

où ΔQ_{spa} (resp. ΔQ_{sto}) est la contribution à l'erreur locale due à la discrétisation dans l'espace physique (resp. dans la dimension stochastique).

D'un côté, la contribution ΔQ_{sto} de l'erreur peut être estimée avec la méthode précédente en considérant un modèle de référence déjà discrétisé en espace (maillage $\mathcal{M}_h$). Pour ce nouveau modèle, $\mathbf{u}_h$ est la solution exacte et $\mathbf{u}_{h,s}$ une solution approchée obtenue après discrétisation dans la dimension stochastique. Dans ce contexte, une solution admissible notée $(\hat{\mathbf{u}}_s, \hat{\sigma}_s)$ doit être définie au sens du nouveau modèle de référence ; en particulier, $\hat{\sigma}_s$ doit simplement vérifier l'équilibre au sens EF sur Θ . En pratique, $(\hat{\mathbf{u}}_s, \hat{\sigma}_s)$ peut être automatiquement obtenu comme un simple post-traitement de la solution approchée $(\mathbf{u}_{h,s}, \sigma_{h,s})$ à disposition : on prend $\hat{\mathbf{u}}_s = \mathbf{u}_{h,s}$ et on construit $\hat{\sigma}_s$ comme :

$$\hat{\sigma}_s(\mathbf{x}, \theta) = \sum_k \sigma_{h,s}^{N_k}(\mathbf{x}) \cdot N_k(\theta) \quad (2.44)$$

Une construction similaire est menée pour obtenir une solution admissible $(\hat{\mathbf{u}}_s, \hat{\sigma}_s)$ du problème adjoint. Dès lors, le terme ΔQ_{sto} peut être estimé à partir de la borne $e_{ERC}(\hat{\mathbf{u}}_s, \hat{\sigma}_s) \cdot e_{ERC}(\hat{\mathbf{u}}_s, \hat{\sigma}_s)$.

De la même façon, la contribution $\Delta Q_{spa} \approx Q(\mathbf{u}_s) - Q(\mathbf{u}_{h,s})$ peut être estimée en considérant un modèle de référence déjà discrétisé dans la dimension stochastique. Pour ce nouveau modèle, $\mathbf{u}_s$ est la solution exacte et $\mathbf{u}_{h,s}$ une solution approchée obtenue après discrétisation en espace. Une solution admissible $(\hat{\mathbf{u}}_h, \hat{\sigma}_h)$ pour ce nouveau modèle est obtenue avec un simple post-traitement de la solution approchée $(\mathbf{u}_{h,s}, \sigma_{h,s})$ à disposition : on prend $\hat{\mathbf{u}}_h = \mathbf{u}_{h,s}$ et on construit $\hat{\sigma}_h$, vérifiant l'équilibre uniquement pour les réalisations calculées, sous la forme :

$$\hat{\sigma}_h(\mathbf{x}, \theta) = \sum_k \hat{\sigma}_{h,m}^k(\mathbf{x}) \cdot \Psi_k(\theta) \quad (2.45)$$

La construction des champs admissibles $(\hat{\mathbf{u}}_h, \hat{\sigma}_h)$ pour le problème adjoint est similaire. La contribution d'erreur ΔQ_{spa} peut être estimée à partir de la borne $e_{ERC}(\hat{\mathbf{u}}_h, \hat{\sigma}_h) \cdot e_{ERC}(\hat{\mathbf{u}}_h, \hat{\sigma}_h)$.

Cette méthode d'estimation des deux sources d'erreur est appliquée au cas-test précédent. En notant ΔQ_{spa}^{est} (resp. ΔQ_{sto}^{est}) l'estimateur de ΔQ_{spa} (resp. ΔQ_{sto}), on donne dans le Tableau 2.2 les résultats obtenus pour différentes valeurs de h et s .

s	h	ΔQ_{spa}^{est}	ΔQ_{sto}^{est}
11	6	62,048	$1, 15 \cdot 10^{-4}$
11	12	22,285	$1, 21 \cdot 10^{-4}$
11	24	7,419	$1, 23 \cdot 10^{-4}$
11	48	2,525	$1, 22 \cdot 10^{-4}$

s	h	ΔQ_{spa}^{est}	ΔQ_{sto}^{est}
3	12	22,181	0,174
5	12	22,311	0,004
11	12	22,285	$1, 12 \cdot 10^{-4}$
21	12	22,287	$8, 35 \cdot 10^{-6}$
41	12	22,287	$4, 28 \cdot 10^{-7}$
81	12	22,285	$3, 01 \cdot 10^{-8}$

Tableau 2.2: Estimation des sources d'erreur en fonction de h (gauche) et de s (droite)

Les résultats montrent que l'erreur de discrétisation stochastique est, pour cet exemple, très faible même pour un nombre de réalisations limité. Les estimateurs ΔQ_{spa}^{est} et ΔQ_{sto}^{est} peuvent être utiles pour l'adaptation du modèle numérique, car ils donnent une information sur la dimension (stochastique ou spatiale) dans laquelle la discrétisation doit être raffinée en priorité pour diminuer l'erreur sur Q .

1.3 Extension aux problèmes résolus par la XFEM

1.3.1 Présentation du problème

La méthode XFEM a été introduite dans [Moës *et al.* 1999] comme une généralisation de la FEM permettant de capter précisément et facilement, dans un milieu fissuré (Figure 2.4), les caractéristiques locales de la solution.

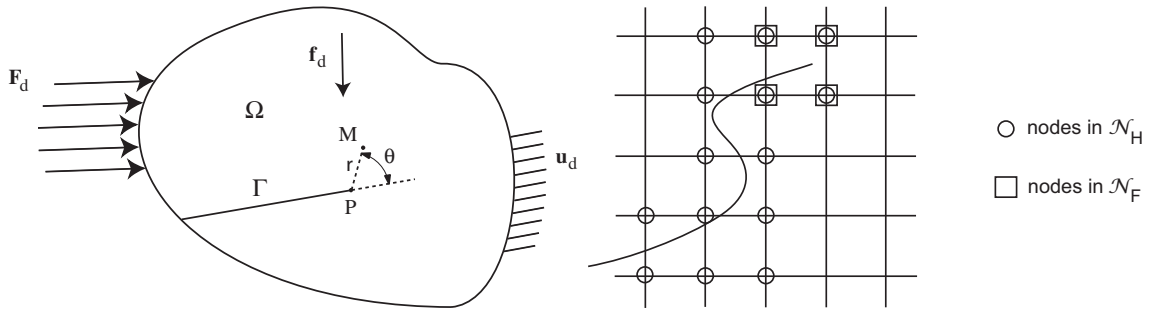


Figure 2.4: Structure fissurée (gauche) et ensembles $\mathcal{N}_H$ et $\mathcal{N}_F$ des nœuds enrichis dans un maillage régulier (droite)

Afin d'améliorer le taux de convergence et de pouvoir utiliser un maillage non conforme avec la fissure Γ (supposée libre d'effort), la méthode XFEM consiste à enrichir l'approximation EF classique à l'aide de la partition d'unité définie par les fonctions de forme $N_i(\mathbf{x})$ du maillage $\mathcal{M}_h$, et en impliquant deux types de fonction d'enrichissement :

- la fonction de Heaviside $H(\mathbf{x})$ qui permet de représenter le saut de déplacement à travers Γ . Cette fonction vaut ± 1 sur les deux bords de la fissure et est en pratique obtenue avec des fonctions level-set. L'enrichissement avec $H(\mathbf{x})$ est appliqué à l'ensemble $\mathcal{N}_H$ des nœuds associés aux éléments coupés par la fissure (Figure 2.4) ;
- l'ensemble $\{F_j(\mathbf{x}), 1 \leq j \leq 4\}$ des fonctions de base générant la solution asymptotique singulière au voisinage de la pointe de fissure P :

$$\left\{ \sqrt{r} \cos\left(\frac{\theta}{2}\right), \sqrt{r} \sin\left(\frac{\theta}{2}\right), \sqrt{r} \cos\left(\frac{\theta}{2}\right) \sin\left(\frac{\theta}{2}\right), \sqrt{r} \sin\left(\frac{\theta}{2}\right) \sin\left(\frac{\theta}{2}\right) \right\} \quad (2.46)$$

L'enrichissement avec les fonctions $F_j(\mathbf{x})$ est appliqué à l'ensemble $\mathcal{N}_F$ des nœuds associés aux éléments entourant la pointe de fissure (Figure 2.4).

Dès lors, en notant $\mathcal{N}$ l'ensemble des nœuds de $\mathcal{M}_h$, la solution approchée XFEM en déplacement est cherchée sous la forme :

$$\mathbf{u}_h(\mathbf{x}) = \sum_{i \in \mathcal{N}} N_i(\mathbf{x}) \mathbf{u}_i + \sum_{i \in \mathcal{N}_F} N_i(\mathbf{x}) (F_j(\mathbf{x}) \mathbf{a}_i^j) + \sum_{i \in \mathcal{N}_H} N_i(\mathbf{x}) H(\mathbf{x}) \mathbf{b}_i \quad (2.47)$$

où $\{\mathbf{u}_i\}$ est l'ensemble des degrés de liberté EF standards tandis que $\{\mathbf{a}_i^j, \mathbf{b}_i\}$ est l'ensemble des degrés de liberté additionnels.

Dans la plupart des cas, comparé à la MEF classique, l'enrichissement introduit avec la XFEM améliore considérablement la solution approchée et les valeurs des quantités d'intérêt qui en sont issues. Néanmoins, l'estimation de l'erreur de discrétisation pour ces quantités reste primordiale pour le calcul robuste. En mécanique de la rupture, la vérification des calculs vis-à-vis de quantités d'intérêt a été abordée dans [Stone et Babuska 1998, Gallimard et Panetier 2006] dans le cadre de la MEF classique ; dans [Bordas et Duflot 2007], une méthode de vérification a été proposée pour la XFEM mais ne fournit pas de bornes d'erreur garanties. Dans [Panetier *et al.* 2010], une procédure garantie d'estimation d'erreur locale

dans le cadre XFEM a été introduite et analysée pour les problèmes 2D. Basée sur l'ERC, le point technique de cette procédure concerne la construction de champs admissibles dans le contexte XFEM, et on montre que cette construction peut être menée en généralisant la procédure classique détaillée dans la Section 3.1.

On choisit comme quantités d'intérêt les facteurs d'intensité des contraintes (FIC) K_I et K_{II} , et utilisons les fonctions d'extraction asymptotiques proposées dans [Babuška et Miller 1984] et définies sur une couronne arbitraire ω entourant la pointe de fissure :

$$K_\alpha = \int_\omega \varepsilon(\mathbf{u}) : (\mathbf{K} \varepsilon(\phi \mathbf{v}_\alpha) - \phi \sigma_\alpha) d\omega - \int_\omega \mathbf{u} \cdot (\sigma_\alpha \nabla \phi) d\omega \quad \alpha = I, II \quad (2.48)$$

Dans (2.48), $\mathbf{v}_\alpha$ et σ_α sont les solutions singulières du problème de référence tandis que ϕ est une fonction continue, définie dans la couronne ω , qui s'annule le long du bord interne $\mathcal{C}_1$ et vaut 1 le long du bord externe $\mathcal{C}_2$. Dans la suite, on considère une couronne circulaire et on utilise une fonction linéaire ϕ définie en coordonnées polaires comme :

$$\begin{cases} \phi(r, \theta) &= \frac{R_2 - r}{R_2 - R_1} & \text{pour } R_1 \leq r \leq R_2 \\ &= 0 & \text{sinon} \end{cases} \quad (2.49)$$

R_1 et R_2 représentent les rayons des cercles $\mathcal{C}_1$ et $\mathcal{C}_2$ respectivement (Figure 2.5).

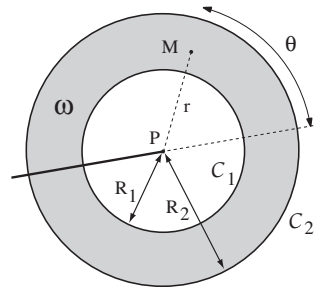
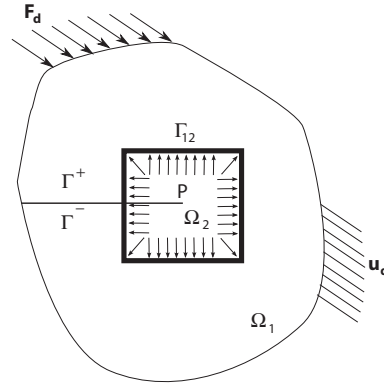


Figure 2.5: Définition de la couronne ω entourant la pointe de fissure

L'estimation d'erreur (2.14) reste valable pourvu que des solutions admissibles $(\hat{\mathbf{u}}_h, \hat{\sigma}_h)$ et $(\hat{\hat{\mathbf{u}}}_h, \hat{\hat{\sigma}}_h)$ soient construites à partir des solutions approchées XFEM.

1.3.2 Construction de solutions admissibles

Le champ cinématique $\hat{\mathbf{u}}_h$ (resp. $\hat{\hat{\mathbf{u}}}_h$) est pris égal au déplacement approché $\mathbf{u}_h$ (resp. $\tilde{\mathbf{u}}_h$) car il est possible dans l'approche XFEM de vérifier exactement les conditions de Dirichlet imposées sur $\partial_1 \Omega$ [Moës et al. 2006]. La construction de $\hat{\sigma}_h$ (ou $\hat{\hat{\sigma}}_h$) est quant à elle menée en divisant la structure Ω en deux zones complémentaires Ω_1 et Ω_2 (Figure 2.6) afin de prendre en compte séparément les deux types d'enrichissement introduits dans la XFEM. Une approche similaire a été examinée dans [Panetier et al. 2009] avec un maillage conforme à la fissure.

Figure 2.6: Définition des deux zones Ω_1 et Ω_2

La zone Ω_2 , qui entoure la pointe de fissure, contient l'ensemble des nœuds enrichis par les fonctions $F_j(\mathbf{x})$. Dans cette zone, $\hat{\sigma}_h$ est construit à partir des fonctions d'Airy, i.e. en exprimant les composantes de $\hat{\sigma}_h$ dans la base polaire sous la forme :

$$\begin{cases} \hat{\sigma}_{h,rr} &= \frac{1}{r^2}\phi_{,\theta\theta} + \frac{1}{r}\phi_{,r} \\ \hat{\sigma}_{h,\theta\theta} &= \phi_{,rr} \\ \hat{\sigma}_{h,r\theta} &= -(\frac{1}{r}\phi_{,\theta})_{,r} \end{cases} ; \quad \phi(r, \theta) = \sum_{i=1}^n r^{\beta_i+2} \gamma_i(\theta) \quad (2.50)$$

avec $\beta_i = \frac{1}{2}(i - 2)$. Les fonctions γ_i ($i = 1, \dots, n$), qui correspondent aux solutions asymptotiques, s'écrivent :

$$\gamma_i(\theta) = A_i \sin(\beta_i \theta) + B_i \cos(\beta_i \theta) + C_i \sin((\beta_i + 2)\theta) + D_i \cos((\beta_i + 2)\theta) \quad (2.51)$$

et les constantes A_i , B_i , C_i et D_i vérifient :

$$\begin{cases} B_i + D_i = 0 & \text{et} & A_i i + C_i(i + 2) = 0 & \text{pour } \beta_i = 0, 1, 2, \dots \\ A_i + C_i = 0 & \text{et} & B_i(i + 1/2) + D_i(i + 5/2) = 0 & \text{pour } \beta_i = -1/2, 1/2, \dots \end{cases} \quad (2.52)$$

Les valeurs optimales des constantes A_i , B_i , C_i et D_i sont alors déterminées en résolvant un problème de minimisation sur Ω_2 :

$$\{A_i, B_i, C_i, D_i\} = \underset{\{A'_i, B'_i, C'_i, D'_i\}}{\operatorname{argmin}} \quad \|\hat{\sigma}_h(\{A'_i, B'_i, C'_i, D'_i\}) - \sigma_h\|_{\mathbf{K}^{-1}|\Omega_2} \quad (2.53)$$

En pratique, on conserve seulement le premier ordre $\beta = -1/2$ ($i = 1$) dont seules 4 constantes doivent être déterminées. A la fin de cette étape, pour assurer la continuité du vecteur contrainte à l'interface Γ_{12} entre Ω_1 et Ω_2 , on résout sur Ω_1 un nouveau problème avec des conditions limites de Neuman $\bar{\mathbf{F}}_d = \hat{\sigma}_h|_{\Omega_2} \mathbf{n}$ sur Γ_{12} .

Dans la zone Ω_1 , la construction de $\hat{\sigma}_h$ est faite en étendant au cadre XFEM la procédure classique [Ladevèze et Pelle 2004] basée sur une condition de prolongement et utilisant les propriétés de σ_h (voir Section 3.1). On prend donc en compte par ce biais les discontinuités introduites avec l'enrichissement par la fonction $H(\mathbf{x})$, sans aucune restriction sur le maillage $\mathcal{M}_h$.

La procédure classiquement utilisée dans la MEF, qui consiste à construire des densités d'effort équilibrées $\hat{\mathbf{F}}_h$ sur le bord ∂E de chaque élément E de $\mathcal{M}_h$ avant de résoudre des problèmes locaux par élément, est directement étendue au cadre XFEM en observant que la XFEM consiste simplement à ajouter de nouvelles fonctions à la base $\{N_i(\mathbf{x})\}$. La condition de prolongement pour les éléments enrichis devient alors :

$$\int_E (\hat{\sigma}_h - \sigma_h) \nabla N_i dE = 0 \quad ; \quad \int_E (\hat{\sigma}_h - \sigma_h) \nabla (N_i H) dE = 0 \quad (2.54)$$

et mène aux deux relations :

$$\begin{aligned} \int_{\partial E} \eta_E \hat{\mathbf{F}}_h N_i dS &= \int_E (\sigma_h \nabla N_i - \mathbf{f}_d N_i) dE = \mathbf{Q}_E(i) \\ \int_{\partial E} \eta_E \hat{\mathbf{F}}_h N_i H dS &= \int_E (\sigma_h \nabla (N_i H) - \mathbf{f}_d N_i H) dE = \mathbf{Q}_E^H(i) \end{aligned} \quad (2.55)$$

Ecrire (2.55) pour tous les éléments connectés au nœud i définit un problème local. Cela est illustré sur le cas de la Figure 2.7 où deux des quatre éléments connectés au nœud i sont coupés par la fissure ; on a alors :

$$\begin{cases} \mathbf{b}_{14} - \mathbf{b}_{21} = \mathbf{Q}_{E_1}(i) \\ \mathbf{b}_{21} - \mathbf{b}_{32} = \mathbf{Q}_{E_2}(i) \\ \mathbf{b}_{32} - \mathbf{b}_{43} = \mathbf{Q}_{E_3}(i) \\ \mathbf{b}_{43} - \mathbf{b}_{14} = \mathbf{Q}_{E_4}(i) \end{cases} \quad \text{et} \quad \begin{cases} \mathbf{b}_{14} - \mathbf{b}_{21}^H = \mathbf{Q}_{E_1}^H(i) \\ \mathbf{b}_{21}^H - \mathbf{b}_{32} = \mathbf{Q}_{E_2}^H(i) \end{cases} \quad (2.56)$$

avec les inconnues $\mathbf{b}_{kl} = \int_{\Gamma_{kl}} \eta_{E_k} \hat{\mathbf{F}}_h N_i dS$ et $\mathbf{b}_{kl}^H = \int_{\Gamma_{kl}} \eta_{E_k} \hat{\mathbf{F}}_h N_i H dS$.

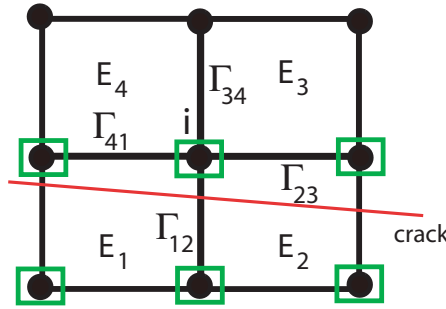


Figure 2.7: Patch d'éléments autour du nœud i , avec les éléments E_1 et E_2 coupés par la fissure. Les nœuds enrichis sont encadrés

Les propriétés de σ_h imposent $\sum_E \mathbf{Q}_E(i) = \sum_E \mathbf{Q}_E^H(i) = \mathbf{0}$ et assurent que les problèmes locaux du type (2.56) sont bien posés.

Excepté pour les nœuds situés sur des bords soumis à des conditions limites de Neuman, la solution des problèmes locaux issus de (2.55) n'est pas unique ; il est alors usuel de minimiser l'écart entre les projections $\mathbf{b}_{kl}$ et $\mathbf{b}_{kl}^H$ recherchées et celles issues de la solution XFEM σ_h . Cependant, une procédure explicite pour le calcul des $\mathbf{b}_{kl}^H$ peut être mise en place. On définit tout d'abord :

$$\overline{\mathbf{Q}}_E(i) = \int_{\partial E} \eta_E \hat{\mathbf{F}}_h (N_i H - N_i) dS = \int_E (\sigma_h \nabla (N_i H - N_i) - \mathbf{f}_d (N_i H - N_i)) dE \quad (2.57)$$

En revenant par exemple au cas de la Figure 2.7 et en observant que la fonction $N_i H - N_i$ s'annule partout sauf sur le bord où $H = -1$ (e.g. sous la fissure), on obtient :

$$\overline{\mathbf{Q}}_{E_1}(i) = \int_{\Gamma_{12}} \eta_{E_1} \hat{\mathbf{F}}_h(N_i H - N_i) dS \quad ; \quad \overline{\mathbf{Q}}_{E_2}(i) = \int_{\Gamma_{12}} \eta_{E_2} \hat{\mathbf{F}}_h(N_i H - N_i) dS \quad (2.58)$$

et l'unicité de la solution est garantie par la relation $\overline{\mathbf{Q}}_{E_1}(i) + \overline{\mathbf{Q}}_{E_2}(i) = \mathbf{0}$. On peut alors déterminer $\mathbf{b}_{21}(i)$ et $\mathbf{b}_{21}(j)$ en utilisant la technique classique, puis calculer :

$$\begin{aligned} \overline{\mathbf{Q}}_{E_1}(i) &= -\overline{\mathbf{Q}}_{E_2}(i) = \int_{E_1} (\sigma_h \nabla(N_i H - N_i) - \mathbf{f}_d(N_i H - N_i)) dE \\ \overline{\mathbf{Q}}_{E_1}(j) &= -\overline{\mathbf{Q}}_{E_2}(j) = \int_{E_1} (\sigma_h \nabla(N_j H + N_j) - \mathbf{f}_d(N_j H + N_j)) dE \end{aligned} \quad (2.59)$$

qui fournit les projections :

$$\mathbf{b}_{21}^H(i) = \overline{\mathbf{Q}}_{E_1}(i) + \mathbf{b}_{21}(i) \quad ; \quad \mathbf{b}_{21}^H(j) = \overline{\mathbf{Q}}_{E_1}(i) - \mathbf{b}_{21}(j) \quad (2.60)$$

Regardons pour finir deux configurations particulières (Figure 2.8). La première est celle où la fissure coupe deux côtés adjacents. Dans ce cas :

$$\overline{\mathbf{Q}}_{E_1}(i) + \overline{\mathbf{Q}}_{E_3}(i) + \overline{\mathbf{Q}}_{E_4}(i) = \mathbf{0} \quad (2.61)$$

En observant qu'aux endroits où $\Psi_i = N_i H \pm N_i$ a été introduit, on a :

$$\overline{\mathbf{Q}}_{E_4}(i) = \int_{\Gamma_{34}} \sigma_{h|E_4} \mathbf{n} \Psi_i dS + \int_{\Gamma_{41}} \sigma_{h|E_4} \mathbf{n} \Psi_i dS \quad (2.62)$$

on lève l'indétermination en prenant :

$$\overline{\mathbf{Q}}_{E_3}(i) = - \int_{\Gamma_{34}} \sigma_{h|E_4} \mathbf{n} \Psi_i dS \quad ; \quad \overline{\mathbf{Q}}_{E_1}(i) = - \int_{\Gamma_{41}} \sigma_{h|E_4} \mathbf{n} \Psi_i dS \quad (2.63)$$

L'autre configuration est celle d'une fissure passant par un nœud du maillage (Fi-

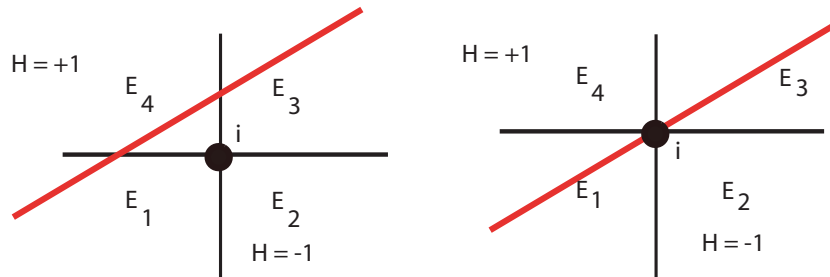


Figure 2.8: Cas d'une fissure coupant deux côtés adjacents (gauche) ou passant par un nœud (droite)

gure 2.8). Dans ce cas, on n'a pas besoin de prendre en compte de discontinuité pour calculer les densités. On écrit simplement :

$$\begin{aligned} \mathbf{b}_{14}(i) &= -\frac{1}{2} (\mathbf{Q}_{E_4}(i) + \mathbf{Q}_{E_4}^H(i)) \quad ; \quad \mathbf{b}_{21}(i) = \frac{1}{2} (\mathbf{Q}_{E_1}(i) + \mathbf{Q}_{E_1}^H(i)) \\ \mathbf{b}_{32}(i) &= \frac{1}{2} (\mathbf{Q}_{E_2}^H(i) - \mathbf{Q}_{E_2}(i)) \quad ; \quad \mathbf{b}_{43}(i) = -\frac{1}{2} (\mathbf{Q}_{E_3}^H(i) - \mathbf{Q}_{E_3}(i)) \end{aligned} \quad (2.64)$$

et on peut voir que le calcul est entièrement explicite.

Après avoir calculé les projections $\mathbf{b}_{ij}$ et $\mathbf{b}_{ij}^H$, on construit les densités le long des bords des éléments. Dans le cas de bords coupés par la fissure (par exemple Γ_{12} dans la Figure 2.7), on cherche les densités sous la forme :

$$\hat{\mathbf{F}}_h = \hat{\mathbf{F}}_i N_i + \hat{\mathbf{F}}_j N_j + \hat{\mathbf{F}}_i^H N_i H + \hat{\mathbf{F}}_j^H N_j H \quad (2.65)$$

i.e. comme la somme d'une partie linéaire continue et d'une partie discontinue. Les quantités $\hat{\mathbf{F}}_i$, $\hat{\mathbf{F}}_j$, $\hat{\mathbf{F}}_i^H$ et $\hat{\mathbf{F}}_j^H$ peuvent être facilement déterminées en inversant un petit système linéaire impliquant les projections calculées précédemment. Pour les autres bords d'élément, on cherche les densités sous la forme linéaire continue classique.

Un champ de contrainte admissible $\hat{\sigma}_{h|E}$ est alors construit sur chaque élément E à partir des densités $\hat{\mathbf{F}}_h$. Les problèmes locaux, qui s'écrivent :

$$\int_E \hat{\sigma}_{h|E} : \varepsilon(\mathbf{v}) dE = \int_E \mathbf{f}_d \cdot \mathbf{v} dE + \int_{\partial E} \eta_E \hat{\mathbf{F}}_h \cdot \mathbf{v} dS \quad \forall \mathbf{v} \in \mathcal{U}_\Gamma^E, \forall E \quad (2.66)$$

sont bien posés grâce aux conditions de prolongement (2.55). On impose en outre la condition de bord libre $\hat{\sigma}_{h|E} \mathbf{n} = \mathbf{0}$ sur les lèvres de la fissure. En pratique, les problèmes locaux (2.66) sont résolus numériquement avec des fonctions de forme de degré $p + 3$ et en introduisant un enrichissement avec les fonctions HN_i pour les éléments coupés par la fissure. Cela mène à la solution d'un petit problème EF sur un quadrangle enrichi à 16 nœuds (Figure 2.9).

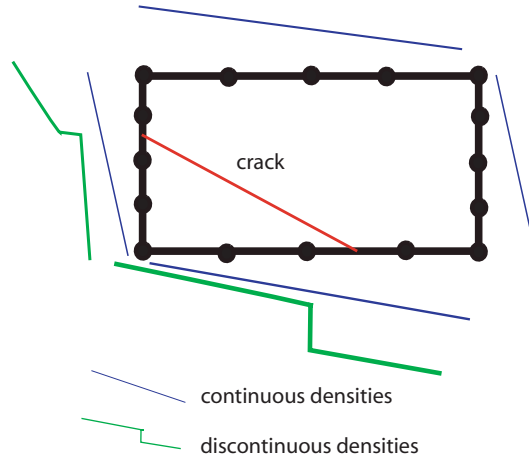


Figure 2.9: Problème local à résoudre sur chaque élément E

Pour illustration, on considère une structure (Figure 2.10) soumise à une densité d'effort de traction uniforme σ (mode I) ou de cisaillement uniforme τ (mode mixte). Les paramètres sont $\sigma = 1$, $\tau = 1$, $E = 210$, $\nu = 0,3$, $L = 16$ et $w = 7$. La structure est discrétisée avec un maillage régulier comprenant 1071 éléments Q4. Dans le cas du mode I, la quantité d'intérêt considérée est K_I , tandis que les quantités d'intérêt

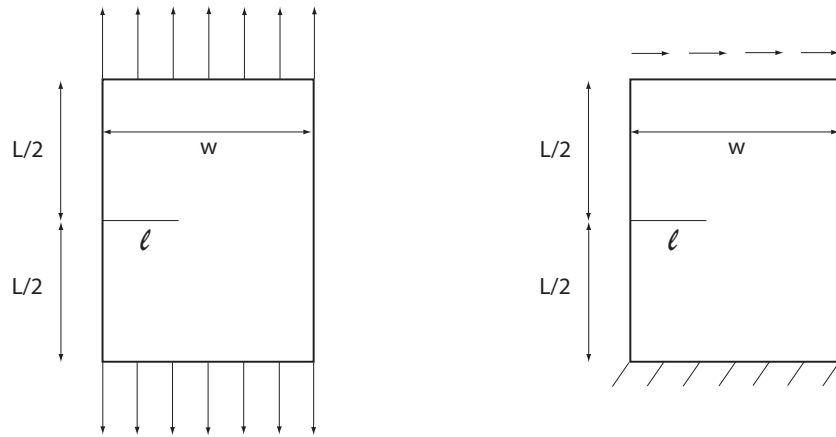
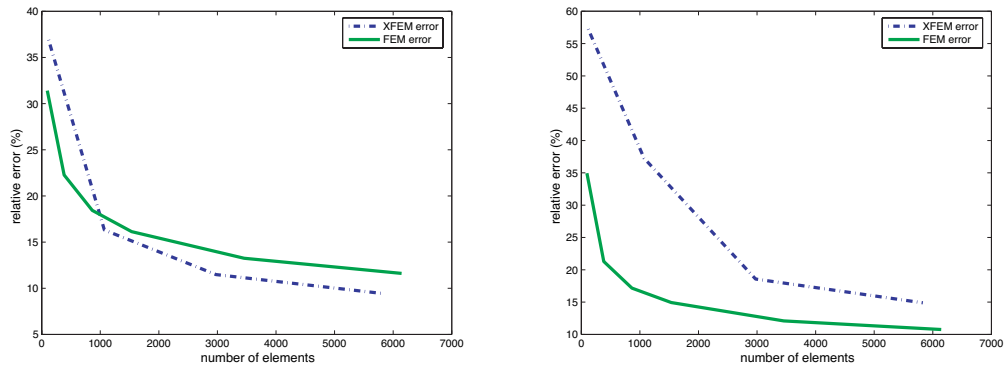


Figure 2.10: Structure avec chargement en mode I (gauche) ou mode mixte (droite)

K_I et K_{II} sont considérées dans le cas du mode mixte. La Figure 2.11 montre l'erreur globale relative des problèmes adjoints associés en fonction du nombre d'éléments, comparée à une résolution par la MEF standard.

Figure 2.11: Convergence de l'erreur du problème adjoint pour K_I en mode I (gauche) et pour K_{II} en mode mixte (droite)

Les Tableaux 2.3, 2.4, et 2.5 montrent quant à eux les valeurs des bornes obtenues pour les quantités d'intérêt considérées. Les quantités $K_{I,corr}$ et $K_{II,corr}$ sont les termes correctifs, tandis que $K_I^\pm$ et $K_{II}^\pm$ sont les bornes supérieure et inférieure de la valeur exacte des quantités d'intérêt. Les valeurs de référence, calculées avec un maillage très fin, sont $K_{I,ref} = 9,33$ en mode I et $[K_{I,ref}, K_{II,ref}] = [33, 93; 4, 53]$ en mode mixte.

Les quantités extraites pour K_I sont de très bonne qualité, et le terme de correction couplé avec des bornes très précises permet d'obtenir une estimation correcte de la valeur de référence. Dans le cas du mode II, les quantités extraites pour K_{II} sont de moins bonne qualité, comme cela est souvent le cas dans les applications pratiques. Le terme de correction ne réussit pas ici à améliorer les résultats. Les bornes obtenues sont acceptables, mais nécessiteraient d'être encore améliorées (voir Section 2.3).

nombre d'éléments	$K_{I,h}$	$K_{I,h} + K_{I,corr}$	K_I^-	K_I^+
1071	9,2137	9,2416	9,0223	9,4609
2975	9,3162	9,3167	9,2123	9,4141
5831	9,3231	9,3472	9,3977	9,2847

Tableau 2.3: Bornes de K_I pour le problème XFEM en mode I

nombre d'éléments	$K_{I,h}$	$K_{I,h} + K_{I,corr}$	K_I^-	K_I^+
1071	33,3374	33,6305	32,5714	34,6896
2975	33,3390	33,9100	33,3398	34,4803
5831	33,8579	33,8742	33,5845	34,5845

Tableau 2.4: Bornes de K_I pour le problème XFEM en mode mixte

1.4 Extension aux quantités d'intérêt non-linéaires

La méthode d'estimation d'erreur locale proposée dans la Section 1.1 fournit des bornes garanties uniquement pour des quantités d'intérêt $Q(\mathbf{u})$ linéaires par rapport à $\mathbf{u}$. En effet, $Q(\mathbf{u}) - Q(\mathbf{u}_h) \neq Q(\mathbf{e}_h)$ dans le cas contraire. Une approche classique pour aborder les quantités d'intérêt non-linéaires consiste alors à linéariser, au voisinage de $\mathbf{u}_h$, ces quantités [Oden et Prudhomme 2002], sous la forme :

$$Q(\mathbf{u}) - Q(\mathbf{u}_h) = \mathcal{Q}'(\mathbf{u}_h; \mathbf{e}_h) + O(\|\mathbf{e}_h\|_{\text{loc}}^2) \quad (2.67)$$

où $\mathcal{Q}'$ est la dérivée au sens de Gâteaux de Q , et $\|\bullet\|_{\text{loc}}$ est la norme L^2 définie dans une région locale $\Omega_{\text{loc}} \subset \Omega$ où la quantité Q est définie. Les termes en $\|\mathbf{e}_h\|_{\text{loc}}$ d'ordre supérieur à 1 sont alors négligés, et on considère les extracteurs tangents $(\sigma_{\Sigma|\mathbf{u}_h}^T, \mathbf{f}_{\Sigma|\mathbf{u}_h}^T, \mathbf{F}_{\Sigma|\mathbf{u}_h}^T)$, tels que :

$$Q(\mathbf{u}) - Q(\mathbf{u}_h) \approx \mathcal{Q}'(\mathbf{u}_h; \mathbf{e}_h) = \int_{\Omega} (\varepsilon(\mathbf{e}_h) : \sigma_{\Sigma|\mathbf{u}_h}^T + \mathbf{e}_h \cdot \mathbf{f}_{\Sigma|\mathbf{u}_h}^T) d\Omega + \int_{\partial_2\Omega} \mathbf{e}_h \cdot \mathbf{F}_{\Sigma|\mathbf{u}_h}^T dS \quad (2.68)$$

comme chargement du problème adjoint. Cette approche donne des bornes relativement précises lorsque l'erreur est petite, mais l'aspect garanti de ces bornes n'est plus valide.

Dans [Ladevèze et Chamoin 2010], une technique alternative a été introduite pour conserver une majoration garantie de l'erreur sur des quantités d'intérêt $Q(\mathbf{u})$ non-linéaires, sous réserve qu'elles soient ponctuelles en espace. Cette technique est basée sur une décomposition de $Q(\mathbf{u})$ avec des propriétés de projection pour prendre en compte les termes d'ordre élevé dans les bornes. Elle mène à l'introduction d'un ensemble d'extracteurs, donc à la résolution d'un ensemble de problèmes adjoints.

On illustre la nouvelle technique en prenant l'exemple de la contrainte de Mises :

$$Q = \sigma_{eq} = \sqrt{k_0 \mathbb{S} : \mathbb{S}} \equiv \sqrt{k_0} |\mathbb{S}| \quad (2.69)$$

nombre d'éléments	$K_{II,h}$	$K_{II,h} + K_{II,corr}$	K_{II}^-	K_{II}^+
1071	4,3963	5,3204	3,8389	6,8019
2975	4,3681	5,2492	4,0783	6,4200
5831	4,3964	4,8800	4,2227	5,5373

Tableau 2.5: Bornes de K_{II} pour le problème XFEM en mode mixte

où $\mathbb{S} \equiv \sigma - \frac{1}{3}(\sigma : \mathbb{I})\mathbb{I}$ est la partie déviatorique de σ et k_0 un paramètre scalaire de normalisation. La linéarisation de cette quantité au premier ordre :

$$Q(\mathbf{u}) = \sqrt{k_0}|\mathbb{S}(\mathbf{u}_h) + \mathbb{S}(\mathbf{e}_h)| = \sigma_{eq}(\mathbf{u}_h) + k_0 \frac{\mathbb{S}(\mathbf{u}_h) : \mathbb{S}(\mathbf{e}_h)}{\sigma_{eq}(\mathbf{u}_h)} \quad (2.70)$$

mène à des bornes non garanties, et précises uniquement lorsque $\|\mathbf{e}_h^t\|_{loc}$ est petit. Néanmoins, on peut écrire :

$$\begin{aligned} \frac{1}{k_0} (Q^2(\mathbf{u}) - Q^2(\mathbf{u}_h)) &= |\mathbb{S} - \mathbb{S}_h + \mathbb{S}_h|^2 - |\mathbb{S}_h|^2 \\ &= 2\mathbb{S}_h : (\mathbb{S} - \mathbb{S}_h) + |\mathbb{S} - \mathbb{S}_h|^2 \end{aligned} \quad (2.71)$$

le terme $|\mathbb{S} - \mathbb{S}_h|^2$ étant négligé dans la procédure classique de linéarisation. Ici, nous proposons de décomposer le tenseur $\mathbb{S} - \mathbb{S}_h$ sur une base orthonormée $\{\mathbb{R}_{h,0}, \mathbb{R}_{h,1}, \dots, \mathbb{R}_{h,J}\}$ de l'espace des matrices symétriques déviatoriques, avec $J = 1$ (resp. $J = 2$) en dimension 2 (resp. dimension 3). Pour des raisons de simplicité, on prend $\mathbb{R}_{h,0} = \mathbb{S}_h/|\mathbb{S}_h|$ et on construit les autres éléments $(\mathbb{R}_{h,j}, j = 1, \dots, J)$ de la base à partir des composantes de $\mathbb{S}_h$. On obtient alors :

$$\mathbb{S} - \mathbb{S}_h = \sum_{j=0}^J \alpha_{h,j} \mathbb{R}_{h,j} \quad (2.72)$$

avec $\alpha_{h,0} = \frac{\mathbb{S}_h}{|\mathbb{S}_h|} : (\mathbb{S} - \mathbb{S}_h)$ et $\alpha_{h,j} = \mathbb{R}_{h,j} : (\mathbb{S} - \mathbb{S}_h)$ pour $1 \leq j \leq J$. De ce fait, le terme du second ordre dans (2.71) s'écrit $\sum_{j=0}^J \alpha_{h,j}^2$ et on obtient au final :

$$\begin{aligned} \frac{1}{k_0} (Q^2(\mathbf{u}) - Q^2(\mathbf{u}_h)) &= 2\mathbb{S}_h : (\mathbb{S} - \mathbb{S}_h) + \sum_{j=0}^J \alpha_{h,j}^2 \\ &= 2\mathbb{S}_h : (\sigma - \sigma_h) + \sum_{j=0}^J [\mathbb{R}_{h,j} : (\sigma - \sigma_h)]^2 \\ &= 2\mathbf{K}\mathbb{S}_h : \varepsilon(\mathbf{u} - \mathbf{u}_h) + \sum_{j=0}^J [\mathbf{K}\mathbb{R}_{h,j} : \varepsilon(\mathbf{u} - \mathbf{u}_h)]^2 \end{aligned} \quad (2.73)$$

On considère alors un ensemble de problèmes adjoints définis par des extracteurs donnés explicitement : l'extracteur $\sigma_{\Sigma,0} = \frac{1}{|\mathbb{S}_h|} \mathbf{K}\mathbb{S}_h \delta(\mathbf{x})$ est utilisé pour calculer des

bornes garanties sur les deux premiers termes de (2.73) tandis que les extracteurs $\sigma_{\Sigma,j} = \mathbf{K} \mathbb{R}_{h,j} \delta(\mathbf{x})$, $1 \leq j \leq J$ sont utilisés pour calculer des bornes garanties sur les autres termes. $\delta(\mathbf{x})$ est la distribution de Dirac au point $\mathbf{x}$. La résolution approchée de ces $J + 1$ problèmes adjoints à solutions singulières, ainsi que la construction de solutions admissibles associées, est faite en utilisant les techniques d'enrichissement local présentées dans la Section 3.2. On obtient alors les encadrements suivants :

$$\begin{aligned} \chi_{h,0}^{inf} &\leq \frac{\mathbb{S}_h}{|\mathbb{S}_h|} : (\sigma - \sigma_h) = \sigma_{\Sigma,0} : \varepsilon(\mathbf{u} - \mathbf{u}_h) \leq \chi_{h,0}^{sup} \\ \chi_{h,j}^{inf} &\leq \mathbb{R}_{h,j} : (\sigma - \sigma_h) = \sigma_{\Sigma,j} : \varepsilon(\mathbf{u} - \mathbf{u}_h) \leq \chi_{h,j}^{sup} \quad \forall j = 1, \dots, J \end{aligned} \quad (2.74)$$

et des premières bornes sur $Q^2(\mathbf{u}) - Q^2(\mathbf{u}_h)$ (donc sur $Q(\mathbf{u})$) peuvent être déduites :

$$2|\mathbb{S}_h| \chi_{h,0}^{inf} + \sum_{j=0}^J \theta_{h,j}^{inf} \leq \frac{1}{k_0} (Q^2(\mathbf{u}) - Q^2(\mathbf{u}_h)) \leq 2|\mathbb{S}_h| \chi_{h,0}^{sup} + \sum_{j=0}^J \theta_{h,j}^{sup} \quad (2.75)$$

dans laquelle les paramètres $\theta_{h,j}^{inf}$ et $\theta_{h,j}^{sup}$ sont définis par :

$$\begin{aligned} \theta_{h,j}^{inf} &= \begin{cases} 0 & \text{si } \text{sign}(\chi_{h,j}^{inf}) \cdot \text{sign}(\chi_{h,j}^{sup}) = -1 \\ \min(|\chi_{h,j}^{inf}|^2, |\chi_{h,j}^{sup}|^2) & \text{si } \text{sign}(\chi_{h,j}^{inf}) \cdot \text{sign}(\chi_{h,j}^{sup}) = 1 \end{cases} \\ \theta_{h,j}^{sup} &= \max(|\chi_{h,j}^{inf}|^2, |\chi_{h,j}^{sup}|^2) \end{aligned} \quad (2.76)$$

Cependant, des bornes plus précises peuvent être obtenues en écrivant $\mathbb{S} - \mathbb{S}_h$ comme :

$$\mathbb{S} - \mathbb{S}_h = [\mathbb{V}_h : (\mathbb{S} - \mathbb{S}_h)] \mathbb{V}_h \quad (2.77)$$

où $\mathbb{V}_h$ est une matrice de norme unité définie comme la combinaison linéaire des matrices $\mathbb{R}_{h,j}$ ($j = 0, \dots, J$) :

$$\mathbb{V}_h = \sum_{j=0}^J a_j \mathbb{R}_{h,j} \quad ; \quad \sum_{j=0}^J a_j^2 = 1 \quad (2.78)$$

les valeurs scalaires a_j étant inconnues. De ce fait, on peut écrire :

$$\begin{aligned} \frac{1}{k_0} (Q^2(\mathbf{u}) - Q^2(\mathbf{u}_h)) &= 2\mathbb{S}_h : (\sigma - \sigma_h) + [\mathbb{V}_h : (\sigma - \sigma_h)]^2 \\ &= 2\mathbb{S}_h : (\sigma - \sigma_h) + \left[\sum_{j=0}^J a_j \mathbb{R}_{h,j} : (\sigma - \sigma_h) \right]^2 \end{aligned} \quad (2.79)$$

A partir des bornes définies dans (2.74), on obtient :

$$\sum_{j=0}^J a_j \chi_{h,j}^{inf} \leq \sum_{j=0}^J a_j \mathbf{R}_{h,j} : (\sigma - \sigma_h) \leq \sum_{j=0}^J a_j \chi_{h,j}^{sup} \quad (2.80)$$

qui permet de définir les bornes d'erreur locale :

$$2|\mathbb{S}_h| \chi_{h,0}^{inf} + \bar{\theta}_h^{inf} \leq \frac{1}{k_0} (Q^2(\mathbf{u}) - Q^2(\mathbf{u}_h)) \leq 2|\mathbb{S}_h| \chi_{h,0}^{sup} + \bar{\theta}_h^{sup} \quad (2.81)$$

avec

$$\bar{\theta}_h^{inf} = 0 \quad \text{et} \quad \bar{\theta}_h^{sup} = \max \left[\sup_{\substack{a_0, a_1, \dots, a_J \\ \sum_{j=0}^J a_j^2 = 1}} \left| \sum_{j=0}^J a_j \chi_{h,j}^{inf} \right|^2, \sup_{\substack{a_0, a_1, \dots, a_J \\ \sum_{j=0}^J a_j^2 = 1}} \left| \sum_{j=0}^J a_j \chi_{h,j}^{sup} \right|^2 \right] \quad (2.82)$$

Au final, les bornes optimales sur $\frac{1}{k_0}(Q^2(\mathbf{u}) - Q^2(\mathbf{u}_h))$ sont ξ_{opt}^{inf} et ξ_{opt}^{sup} définies par :

$$\begin{cases} \xi_{opt}^{inf} &= 2|\mathbb{S}_h|\chi_{h,0}^{inf} + \sum_{j=0}^J \theta_{h,j}^{inf} \\ \xi_{opt}^{sup} &= 2|\mathbb{S}_h|\chi_{h,0}^{sup} + \bar{\theta}_h^{sup} \end{cases} \quad (2.83)$$

En guise de résultat numérique, nous considérons une éprouvette 2D en contrainte plane, avec $E = 1$ et $\nu = 0,3$ (Figure 2.12). Elle est soumise à une densité surfacique d'effort normal $N = \pm 1$ et à un déplacement tangentiel $U = \pm 1$ sur les faces de gauche et droite; les autres bords sont libres. On résout le problème avec la MEF (maillage constitué d'éléments T3).

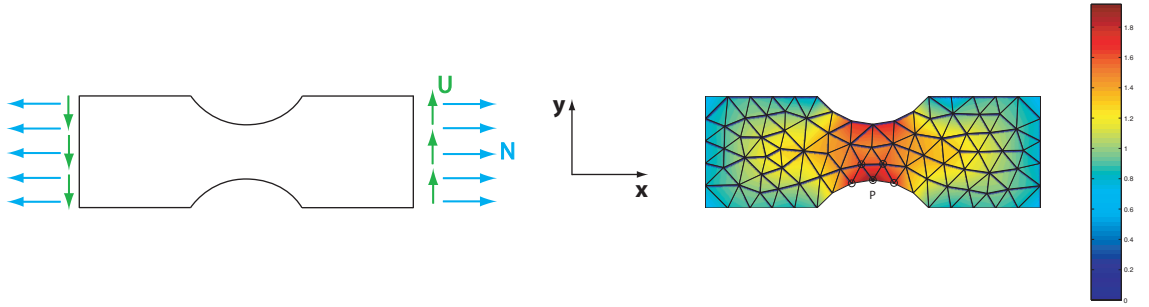


Figure 2.12: Structure soumise au chargement de traction/cisaillement (gauche) et position du point P d'intérêt (droite)

La quantité d'intérêt considérée est la contrainte de Mises au point P où cette contrainte est maximale (Figure 2.12). On introduit alors comme indiqué précédemment les deux extracteurs :

$$\sigma_{\Sigma,0} = \frac{1}{|\mathbb{S}_h(P)|} \mathbf{K} \mathbb{S}_h(P) \delta(\mathbf{x} - \mathbf{x}_P) \quad ; \quad \sigma_{\Sigma,1} = \mathbf{K} \mathbb{R}_{h,1}(P) \delta(\mathbf{x} - \mathbf{x}_P) \quad (2.84)$$

En exprimant $\mathbb{S}_h(P)$ sous la forme :

$$\mathbb{S}_h(P) = \begin{bmatrix} S_h^{xx} & S_h^{xy} \\ S_h^{xy} & -S_h^{xx} \end{bmatrix} \quad (2.85)$$

on obtient

$$\mathbb{R}_{h,1}(P) = \frac{1}{\sqrt{2[(S_h^{xx})^2 + (S_h^{xy})^2]}} \begin{bmatrix} -S_h^{xy} & S_h^{xx} \\ S_h^{xx} & S_h^{xy} \end{bmatrix} \quad (2.86)$$

La solution locale du problème adjoint associé à l'extracteur $\sigma_{\Sigma,0}$ est montrée sur la Figure 2.13.

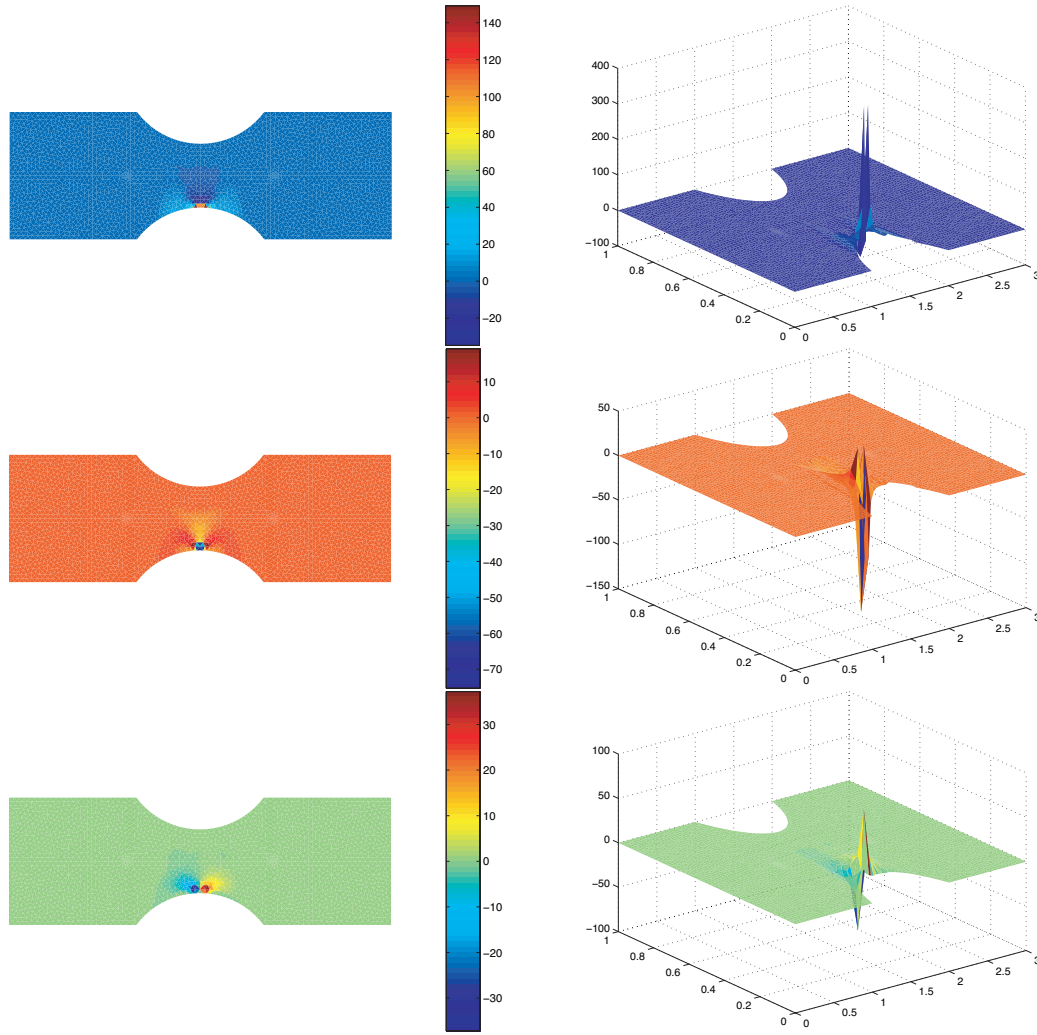


Figure 2.13: Composantes xx (haut), yy (centre) et xy (bas) de la solution en contrainte pour le problème adjoint associé à $\sigma_{\Sigma,0}$

Après post-traitement des solutions approchées des problèmes primal et adjoint, on obtient :

$$\begin{aligned} \chi_{h,0}^{inf} &= 0,0604 & ; & \quad \chi_{h,0}^{sup} = 0,0971 & \quad (\text{bornes liées au problème adjoint avec } \sigma_{\Sigma,0}) \\ \chi_{h,1}^{inf} &= 0,0726 & ; & \quad \chi_{h,1}^{sup} = 0,0811 & \quad (\text{bornes liées au problème adjoint avec } \sigma_{\Sigma,1}) \end{aligned} \quad (2.87)$$

et les bornes optimales définies dans (2.83) sont :

$$\begin{aligned} \xi_{opt}^{inf} &= 2|\mathbb{S}_h|\chi_{h,0}^{inf} + \sum_{j=0}^1 \theta_{h,j}^{inf} = 0,1384 + 0,0199 \\ \xi_{opt}^{sup} &= 2|\mathbb{S}_h|\chi_{h,0}^{sup} + \bar{\theta}_h^{sup} = 0,1927 + 0,0103 \end{aligned} \quad (2.88)$$

Finalement, avec $Q(\mathbf{u}_h) = 1,9601$, on obtient les bornes sur la valeur inconnue $Q(\mathbf{u})$:

$$2,4487 \leq Q(\mathbf{u}) \leq 2,4923 \quad (2.89)$$

qui peut être comparée à la valeur quasi-exacte $Q(\mathbf{u}) \approx 2,4662$ obtenue avec un maillage très fin. Les bornes sur $Q(\mathbf{u})$ sont donc dans une marge de 2%.

2 Contributions à l'obtention de bornes précises

2.1 Nouvelle technique de majoration en espace

La majoration d'erreur proposée dans (2.14) (ou (2.27)), basée sur l'utilisation de l'inégalité de Cauchy-Schwarz de façon globale (i.e. sur le domaine Ω entier), a l'inconvénient de fournir des bornes très pessimistes et donc peu pertinentes lorsque les sources d'erreur des problème primal et adjoint se situent dans des zones spatiales disjointes, ou lorsque le problème primal comporte plusieurs zones d'erreur importante. La majoration correspond dans ce cas à appliquer l'inégalité de Cauchy-Schwarz pour deux vecteurs quasi-orthogonaux, comme illustré sur la Figure 2.14 avec une structure fissurée. Pour cette structure, l'erreur associée au problème primal est concentrée en pointe de fissure. De ce fait, lorsque la quantité d'intérêt se situe au niveau de la fissure (FIC par exemple), appliquer l'inégalité de Cauchy-Schwarz de façon globale fournit une majoration correcte; lorsque la quantité d'intérêt se situe loin de la fissure, cette inégalité donne une forte surmajoration et occasionne des bornes d'erreur de faible qualité.

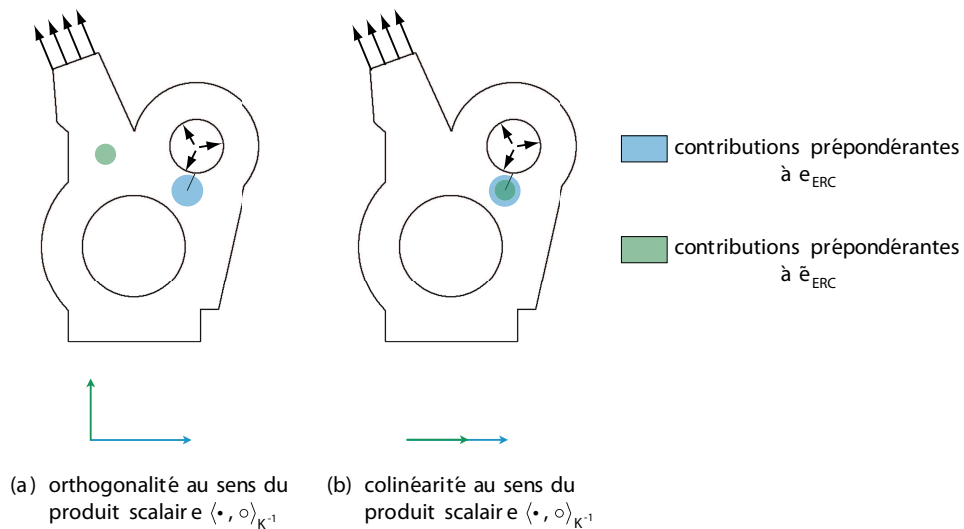


Figure 2.14: Visualisation des contributions d'erreur disjointes

Dans [Ladevèze *et al.* 2013], une nouvelle approche de majoration a été étudiée pour éviter la surmajoration évoquée précédemment. L'idée est de décomposer le domaine Ω en deux zones disjointes : (i) une zone ω_λ , paramétrée par le scalaire λ , qui entoure la zone de définition de la quantité d'intérêt ; (ii) la zone complémentaire Ω/ω_λ . Dès lors, en reprenant la définition de l'erreur de discrétisation et la relation

(2.13), on peut écrire :

$$\begin{aligned}
 Q(\mathbf{u}) - Q(\mathbf{u}_h) - Q_{corr} &= \langle \sigma - \hat{\sigma}_h^m, \hat{\tilde{\sigma}}_h - \tilde{\sigma}_h \rangle_{K^{-1}} \\
 &= \langle \sigma - \hat{\sigma}_h^m, \tilde{\tilde{\sigma}}_h - \tilde{\sigma}_h \rangle_{K^{-1}|\omega_\lambda} + \langle \sigma - \hat{\sigma}_h^m, \tilde{\tilde{\sigma}}_h - \tilde{\sigma}_h \rangle_{K^{-1}|\Omega/\omega_\lambda} \\
 &= q_{\omega_\lambda} + q_{\Omega/\omega_\lambda}
 \end{aligned} \tag{2.90}$$

L'encadrement du terme $q_{\Omega/\omega_\lambda}$ est mené de façon simple et précise à l'aide de l'inégalité de Cauchy-Schwarz appliquée sur Ω/ω_λ , car $\tilde{e}_{ERC|\Omega/\omega_\lambda}$ reste en pratique très petit. Celui du terme q_{ω_λ} est plus fin et technique ; il repose sur l'inégalité suivante liée au principe de Saint-Venant :

$$\|\sigma - \hat{\sigma}_h\|_{K^{-1}|\omega_\lambda} \leq \left(\frac{\lambda}{\bar{\lambda}}\right)^{1/k} \|\sigma - \hat{\sigma}_h\|_{K^{-1}|\omega_{\bar{\lambda}}} + \gamma_{\lambda, \bar{\lambda}} \tag{2.91}$$

Dans cette inégalité, $\omega_{\bar{\lambda}}$ est un domaine homothétique de ω_λ , paramétré par le scalaire $\bar{\lambda} \geq \lambda$ (Figure. 2.15), k est la constante de Steklov dépendant de la géométrie de ω_λ et obtenue analytiquement ou numériquement en résolvant un problème local annexe (problème aux valeurs propres), tandis que $\gamma_{\lambda, \bar{\lambda}}$ est un terme connu. En pratique, $\bar{\lambda}$ est choisi le plus grand possible tout en assurant $\omega_{\bar{\lambda}} \subset \Omega$, et λ le plus petit possible avec ω_λ englobant la zone d'intérêt.

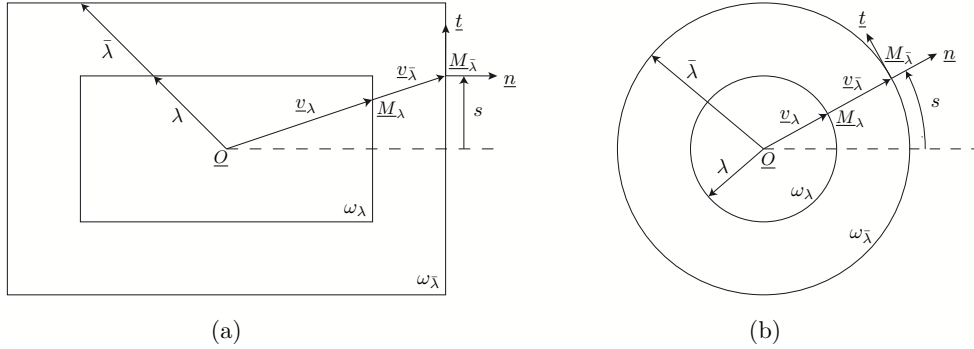


Figure 2.15: Exemples de domaines homothétiques

La décroissance exponentielle en fonction de $\lambda/\bar{\lambda}$ dans (2.91) est le point clé pour éviter la surmajoration. A partir de (2.91) et des propriétés de l'ERC, on en déduit :

$$\|\sigma - \hat{\sigma}_h^m\|_{K^{-1}|\omega_\lambda} \leq \left(\frac{\lambda}{\bar{\lambda}}\right)^{1/k} \frac{1}{2} (e_{ERC} + e_{ERC|\omega_{\bar{\lambda}}})^2 + \gamma_{\lambda, \bar{\lambda}} \tag{2.92}$$

puis, en utilisant (2.90) :

$$|Q(\mathbf{u}) - Q(\mathbf{u}_h) - \overline{Q}_{corr}| \leq \tilde{e}_{ERC|\omega_\lambda} \cdot \delta_{\lambda, \bar{\lambda}} + e_{ERC} \cdot \tilde{e}_{ERC|\Omega/\omega_\lambda} \tag{2.93}$$

avec

$$\begin{aligned}
 \overline{Q}_{corr} &= Q_{corr} + \frac{1}{2} \langle \hat{\sigma}_h - \sigma_h, \hat{\tilde{\sigma}}_h - \tilde{\sigma}_h \rangle_{K^{-1}} \\
 \delta_{\lambda, \bar{\lambda}} &= \left[\left(\frac{\lambda}{\bar{\lambda}}\right)^{1/k} (e_{ERC} + e_{ERC|\omega_{\bar{\lambda}}})^2 + \gamma_{\lambda, \bar{\lambda}} \right]^{1/2}
 \end{aligned} \tag{2.94}$$

Les bornes de $Q(\mathbf{u})$ qui peuvent en être déduites sont notées χ_{inf} et χ_{sup} .

Une seconde majoration, plus performante, a aussi été introduite dans [Ladevèze et al. 2013]. Elle repose sur la propriété suivante :

$$\|\mathbf{u} - \mathbf{u}_\lambda^h\|_{K|\omega_\lambda}^2 \leq \left(\frac{\lambda}{\bar{\lambda}}\right)^\kappa \|\mathbf{u} - \mathbf{u}_{\bar{\lambda}}^h\|_{K|\omega_{\bar{\lambda}}}^2 \quad (2.95)$$

où $\mathbf{u}_\lambda^h$ est solution du problème d'équilibre sur ω_λ avec des conditions de Dirichlet $\mathbf{u}_\lambda^h = \mathbf{u}_h$ sur $\partial\omega_\lambda$, et κ est une autre constante déduite de la résolution d'un problème local. On peut en déduire la nouvelle majoration :

$$|Q(\mathbf{u}) - Q(\mathbf{u}_h) - Q_{corr}| \leq e_{ERC} \cdot (\tilde{\theta}_\lambda^2 + \tilde{e}_{ERC|\Omega/\omega_\lambda}^2)^{1/2} + e_{ERC,|\omega_{\bar{\lambda}}|} \cdot (\tilde{\theta}_{\bar{\lambda}} + \tilde{e}_{ERC|\omega_{\bar{\lambda}}}^2) \quad (2.96)$$

avec $\tilde{\theta}_\lambda = \|\hat{\sigma}_h - \tilde{\sigma}_h\|_{K^{-1}|\partial\omega_\lambda} \frac{2\kappa^{1/2}}{h_2 + n + 1}$ et $n = 1$ (resp. $n = 2$) en 2D (resp. 3D). Les bornes de $Q(\mathbf{u})$ qui peuvent en être déduites sont notées ζ_{inf} et ζ_{sup} .

Pour la structure fissurée présentée sur la Figure 2.14, on s'intéresse à deux quantités d'intérêt :

- la moyenne, sur un élément E situé loin de la fissure, de la composante σ_{xx} du tenseur des contraintes : $Q_1 = \frac{1}{|E|} \int_E \sigma_{xx} d\Omega$;
- le premier facteur d'intensité des contraintes au niveau de la fissure : $Q_2 = K_I$.

Le maillage utilisé est constitué de 7751 éléments T3. On montre sur la Figure 2.16 les domaines homothétiques ω_λ et $\omega_{\bar{\lambda}}$ considérés pour l'estimation d'erreur sur Q_1 et Q_2 .

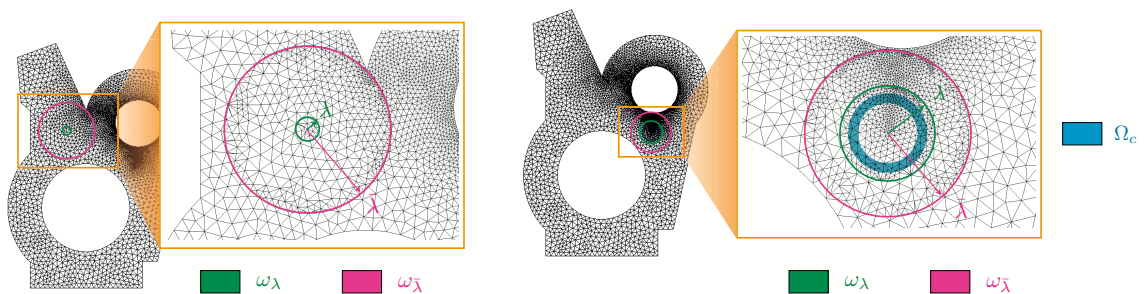
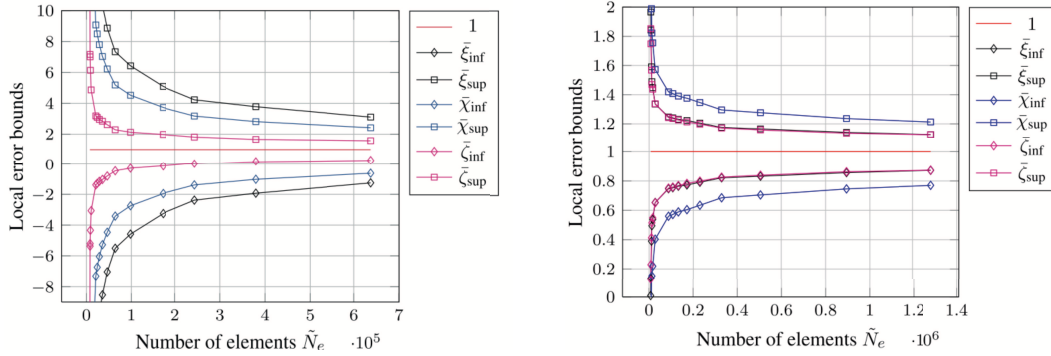


Figure 2.16: Domaines homothétiques utilisés pour Q_1 (gauche) et Q_2 (droite)

Pour la quantité Q_1 (resp. Q_2), la forme de ces domaines est circulaire pleine (resp. circulaire fendue) avec $k = 0,769$ et $\kappa = 2,0$ (resp. $k = 0,798$ et $\kappa = 2,0$) en contraintes planes. On représente sur la Figure 2.17 l'évolution des bornes adimensionnées classiques $\bar{\xi} = \xi/Q(\mathbf{u})$, et améliorées $\bar{\chi} = \chi/Q(\mathbf{u})$ et $\bar{\zeta} = \zeta/Q(\mathbf{u})$, pour diverses discrétisations du problème adjoint.


 Figure 2.17: Améliorations obtenues pour Q_1 (gauche) et Q_2 (droite)

On observe une nette amélioration des bornes pour Q_1 et une légère amélioration pour Q_2 , l'erreur sur cette dernière quantité étant déjà relativement bien estimée avec l'application globale de l'inégalité de Cauchy-Schwarz.

2.2 Prise en compte des effets d'histoire

Pour les problèmes d'évolution décrits par variables internes, certaines quantités physiques deviennent asymptotiquement indépendantes de l'histoire du chargement et par conséquent ne dépendent que du chargement à l'instant t (mémoire courte). Dans le cas de la viscoélasticité exposé dans la Section 1.1.2, ces quantités sont la contrainte σ ainsi que le taux de déformation anélastique $\dot{\varepsilon}^p$. L'estimation d'erreur classique (2.27) pour de telles quantités, basée sur l'erreur en dissipation, devient alors rapidement imprécise car elle cumule inutilement les erreurs au cours du temps. La prise en compte des effets d'histoire dans l'estimation d'erreur locale a donc été étudiée dans [Chamoin et Ladevèze 2007], en introduisant un terme temporel de pondération qui évite le cumul néfaste des erreurs en temps et permet donc de conserver des bornes d'erreur locale pertinentes.

On repart du résultat de représentation de l'erreur sur une quantité $Q(\mathbf{u})$ obtenue dans (2.26) :

$$Q(\dot{\mathbf{u}}) - Q(\dot{\mathbf{u}}_{h,\Delta}) - Q_{corr} = \langle \langle \sigma - \hat{\sigma}_{h,\Delta}^m, \dot{\varepsilon}_{h,\Delta}^p + \mathbf{B}\hat{\sigma}_{h,\Delta} \rangle \rangle \quad (2.97)$$

et du lien, à tout instant $t \in [0, T]$, entre l'erreur vraie et l'erreur en dissipation obtenu dans (2.19) :

$$\|\sigma - \hat{\sigma}^m\|_B^2 + \frac{1}{2} \frac{d}{dt} E_l(\sigma - \hat{\sigma}) = \frac{1}{2} \eta_t(\dot{\varepsilon}^p, \hat{\sigma}) \quad (2.98)$$

On introduit alors une fonction du temps arbitraire $a(t)$ vérifiant :

$$a(t) \text{ continue sur } [0, T] \quad ; \quad a(t) \geq 0 \quad ; \quad a(T) = 0 \quad (2.99)$$

Après multiplication de (2.98) par la fonction $a(t)$ et intégration sur le domaine temporel, on aboutit à la relation :

$$\int_0^T \left[a(t) \|\sigma - \hat{\sigma}^m\|_B^2 - \frac{1}{2} \dot{a}(t) E_l(\sigma - \hat{\sigma}) \right] dt = \frac{1}{2} e_{ERC,a}^2(\dot{\varepsilon}^p, \hat{\sigma}) \quad (2.100)$$

où $e_{ERC,a}^2 = \int_0^T a(t) \eta_t dt$ peut être vue comme une erreur en dissipation pondérée. Le choix de la fonction $a(t)$ permet de prendre en compte des contributions plus ou moins locales de l'erreur totale e_{ERC} (voir ci-après).

Le terme de gauche dans (2.100) peut se réécrire sous la forme d'une fonction quadratique :

$$\int_0^T \int_{\Omega} g(\mathbb{X}, \mathbb{Z}_{h,\Delta}) d\Omega dt \quad ; \quad g(\mathbb{X}, \mathbb{Z}_{h,\Delta}) = g_2(\mathbb{X}) + g_1(\mathbb{X}, \mathbb{Z}_{h,\Delta}) + g_0(\mathbb{Z}_{h,\Delta}) \quad (2.101)$$

avec $\mathbb{X} = \sigma - \hat{\sigma}_{h,\Delta}^m$ et $\mathbb{Z}_{h,\Delta} = \hat{\sigma}_{h,\Delta} - \hat{\sigma}_{h,\Delta}^m$. En introduisant la fonction quadratique duale $f(\mathbb{Y}, \mathbb{Z}_{h,\Delta})$, définie au sens de Legendre-Fenchel par :

$$f(\mathbb{Y}, \mathbb{Z}_{h,\Delta}) = f_2(\mathbb{Y}) + f_1(\mathbb{Y}, \mathbb{Z}_{h,\Delta}) + f_0(\mathbb{Z}_{h,\Delta}) \equiv \sup_{\mathbb{X}} [\mathbb{X} : \mathbb{Y} - g(\mathbb{X}, \mathbb{Z}_{h,\Delta})] \quad \forall \mathbb{Y} \quad (2.102)$$

on obtient alors, par définition et pour tout $\mu \in \mathbf{R}^*$:

$$\mathbb{X} : \mathbb{Y} = \frac{1}{\mu} \mathbb{X} : (\mu \mathbb{Y}) \leq \frac{1}{\mu} (g(\mathbb{X}, \mathbb{Z}_{h,\Delta}) + f(\mu \mathbb{Y}, \mathbb{Z}_{h,\Delta})) \quad \forall \mathbb{X}, \forall \mathbb{Y} \quad (2.103)$$

La minimisation par rapport à μ donne, en notant $F = \int_0^T \int_{\Omega} f d\Omega dt$ et $G = \int_0^T \int_{\Omega} g d\Omega dt$:

$$\langle \mathbb{X}, \mathbb{Y} \rangle \leq 2 [G(\mathbb{X}, \mathbb{Z}_{h,\Delta}) + F_0(\mathbb{Z}_{h,\Delta})]^{1/2} [F_2(\mathbb{Y})]^{1/2} + F_1(\mathbb{Y}, \mathbb{Z}_{h,\Delta}) \quad \forall \mathbb{X}, \forall \mathbb{Y} \quad (2.104)$$

A partir de (2.97) et en appliquant (2.104) avec $\mathbb{Y} = \pm (\dot{\hat{\sigma}}_{h,\Delta}^p + \mathbf{B} \hat{\sigma}_{h,\Delta})$, on obtient la nouvelle majoration garantie :

$$|Q(\dot{\mathbf{u}}) - Q(\dot{\mathbf{u}}_{h,\Delta}) - \overline{Q}_{corr}| \leq 2 \left[\frac{1}{2} e_{ERC,a}^2 (\dot{\hat{\sigma}}_{h,\Delta}^p, \hat{\sigma}_{h,\Delta}) + F_0(\mathbb{Z}_{h,\Delta}) \right]^{1/2} \left[F_2(\dot{\hat{\sigma}}_{h,\Delta}^p + \mathbf{B} \hat{\sigma}_{h,\Delta}) \right]^{1/2} \quad (2.105)$$

avec $\overline{Q}_{corr} = Q_{corr} + F_1(\dot{\hat{\sigma}}_{h,\Delta}^p + \mathbf{B} \hat{\sigma}_{h,\Delta}, \mathbb{Z}_{h,\Delta})$. La quantité $Q(\dot{\mathbf{u}}_{h,\Delta}) + \overline{Q}_{corr}$ peut être vue comme une nouvelle approximation de $Q(\dot{\mathbf{u}})$.

Le choix de la fonction $a(t)$ vérifiant (2.99) est arbitraire et peut être optimisé en fonction de la quantité d'intérêt étudiée. En pratique, on considère deux types de fonctions (Figure 2.18) :

- pour les quantités d'intérêt qui sont asymptotiquement indépendantes de l'histoire, on prend la fonction :

$$a_1(t) = \begin{cases} \exp[k(t - T_0)] & \text{sur } [0, T_0] \\ (T - t)/(T - T_0) & \text{sur } [T_0, T] \end{cases} \quad (2.106)$$

où k et T_0 sont des paramètres à fixer. Le paramètre k doit être suffisamment grand tout en assurant la convexité de g i.e. l'opérateur $a(t)\mathbf{B} - \frac{1}{2}\dot{a}(t)\mathbf{K}^{-1}$ doit rester défini positif; un choix naturel est $k = 1/\tau$ où τ , tel que $\mathbf{B}^{-1} = \tau\mathbf{K}$, est le temps caractéristique des effets visqueux;

- pour les quantités d'intérêt dépendant de l'histoire, on prend la fonction :

$$a_2(t) = 1 - t/T \quad (2.107)$$

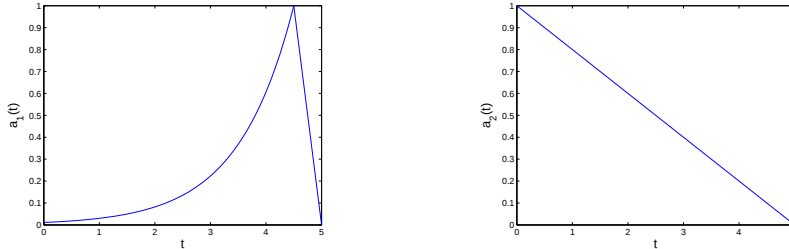


Figure 2.18: Evolution des fonctions $a_1(t)$ et $a_2(t)$ sur le domaine temporel $[0, T]$

Pour illustration, on applique la nouvelle méthode de majoration à la structure représentée sur la Figure 2.19. Il s'agit d'une plaque trouée, en contraintes planes, fixée à sa base et soumise sur sa partie supérieure à un déplacement normal cyclique. On prend $T = 20$ et $\tau = 0, 1$. Le problème est résolu avec la MEF (maillage régulier composé d'éléments Q4) et un schéma temporel d'Euler implicite.

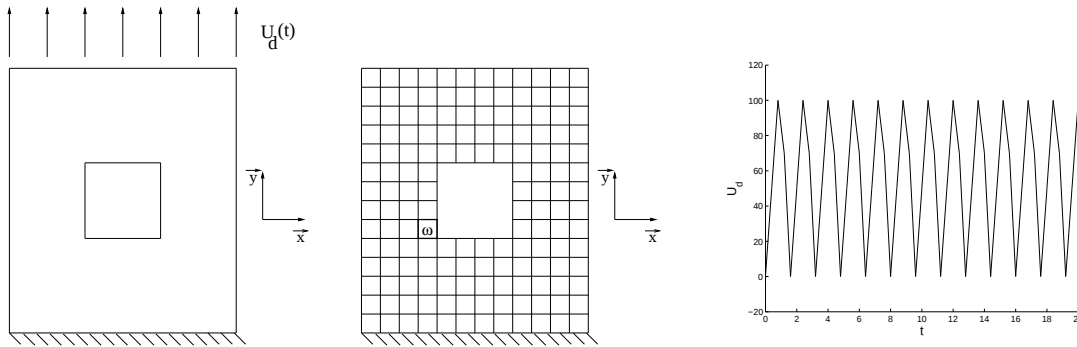


Figure 2.19: Données du problème : géométrie (gauche), maillage (centre), évolution temporelle du chargement (droite)

On prend comme quantité d'intérêt la moyenne sur une zone $\omega \subset \Omega$ située dans la région la plus sollicitée de la structure (voir Figure 2.19), à l'instant $t_1 \in [0, T]$, de la composante $\dot{\varepsilon}_{yy}^p$ du taux de déformation anélastique : $Q = \frac{1}{|\omega|} \int_{\omega} \dot{\varepsilon}_{yy}^p|_{t_1} d\Omega$. Cette quantité est asymptotiquement indépendante de l'histoire, et tend vers un cycle limite. On représente ce cycle sur la Figure 2.20, ainsi que les contributions à chaque pas de temps de l'erreur en dissipation pour les problème primal et adjoint (pour $t_1 = T$). On observe que les contributions sont distribuées pour le problème primal dû au chargement cyclique, tandis qu'elles sont très localisées pour le problème adjoint dû à la faible dépendance de Q à l'histoire.

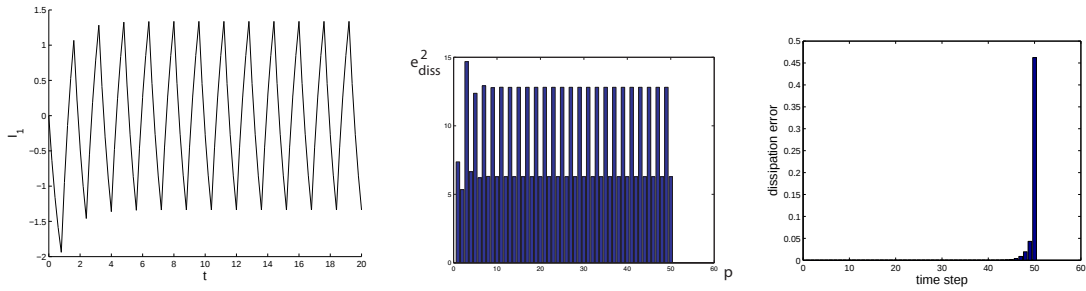


Figure 2.20: Evolution temporelle de Q vers un cycle limite (gauche), et contributions à chaque pas de temps de l'erreur en dissipation pour les problèmes primal (centre) et adjoint (droite)

Les résultats de majoration (bornes Q^- et Q^+ de $Q(\dot{\mathbf{u}})$) obtenus avec l'encadrement classique (2.27) et le nouvel encadrement (2.105) sont montrés sur la Figure 2.21 ; on représente les bornes normées $Q^\pm/Q(\dot{\mathbf{u}})$ en fonction de l'instant t_1 d'estimation de la quantité d'intérêt.

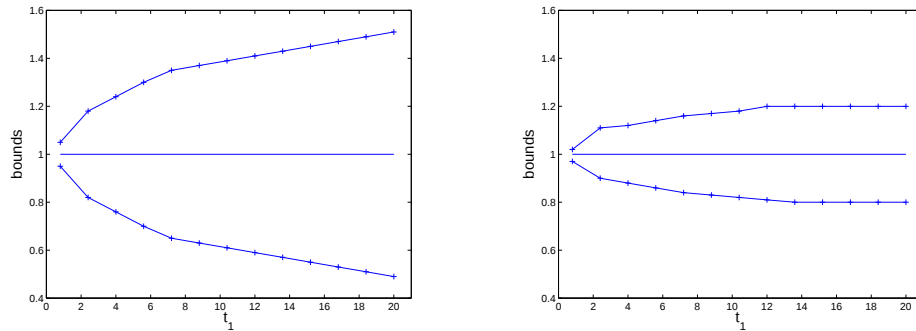


Figure 2.21: Bornes normées sur Q données par l'encadrement classique (gauche) et le nouvel encadrement (droite)

On observe que le cumul d'erreur et la dégradation des bornes avec le temps, observés avec l'encadrement classique, sont évités avec le nouvel encadrement qui fournit des bornes pertinentes pour tout instant t_1 .

2.3 Résolution fine du problème adjoint

Des travaux initiés dans [Chamoin et Ladevèze 2007] et appliqués à divers cas par la suite se sont concentrés sur une nouvelle vision de l'encadrement d'erreur de discrétisation locale. Celle-ci exploite une majoration n'utilisant pas la propriété d'orthogonalité de Galerkin et découplant donc complètement les discrétisations des problèmes primal et adjoint. L'intérêt est alors de pouvoir améliorer localement la qualité de la discrétisation du problème adjoint, dont la solution est très souvent localisée en espace et en temps, sans changer celle du problème initial ; la qualité

des bornes sur Q est alors améliorée à un coût raisonnable.

On présente la nouvelle vision dans le cas de l'élasticité linéaire, bien qu'elle soit aussi valable pour les problèmes dépendant du temps. Soit $(\tilde{\mathbf{u}}_{hh}, \tilde{\sigma}_{hh})$ une solution approchée du problème adjoint calculée avec un maillage $\mathcal{M}_{hh}$ plus fin que $\mathcal{M}_h$. On peut alors écrire :

$$Q(\mathbf{u}) - Q(\mathbf{u}_h) = a(\mathbf{u} - \mathbf{u}_h, \tilde{\mathbf{u}}) = a(\mathbf{u} - \mathbf{u}_h, \tilde{\mathbf{u}} - \tilde{\mathbf{u}}_{hh}) + a(\mathbf{u} - \mathbf{u}_h, \tilde{\mathbf{u}}_{hh}) \quad (2.108)$$

De façon similaire à (2.13), le premier terme peut s'écrire :

$$a(\mathbf{u} - \mathbf{u}_h, \tilde{\mathbf{u}} - \tilde{\mathbf{u}}_{hh}) = \langle \sigma - \hat{\sigma}_h^m, \hat{\sigma}_{hh} - \tilde{\sigma}_{hh} \rangle_{K^{-1}} + \langle \hat{\sigma}_h^m - \sigma_h, \hat{\sigma}_{hh} - \tilde{\sigma}_{hh} \rangle_{K^{-1}} \quad (2.109)$$

Le second terme est quant à lui calculable et se met sous la forme :

$$a(\mathbf{u} - \mathbf{u}_h, \tilde{\mathbf{u}}_{hh}) = l(\tilde{\mathbf{u}}_{hh}) - a(\mathbf{u}_h, \tilde{\mathbf{u}}_{hh}) = \langle \hat{\sigma}_h - \sigma_h, \tilde{\sigma}_{hh} \rangle_{K^{-1}} \quad (2.110)$$

Par conséquent, on obtient :

$$|Q(\mathbf{u}) - Q(\mathbf{u}_h) - Q_{hh}| = |\langle \sigma - \hat{\sigma}_h^m, \hat{\sigma}_{hh} - \tilde{\sigma}_{hh} \rangle_{K^{-1}}| \leq e_{ERC}(\mathbf{u}_h, \hat{\sigma}_h) \cdot e_{ERC}(\tilde{\mathbf{u}}_{hh}, \hat{\sigma}_{hh}) \quad (2.111)$$

où $Q_{hh} = \frac{1}{2} \langle \hat{\sigma}_h - \sigma_h, \hat{\sigma}_{hh} + \tilde{\sigma}_{hh} \rangle_{K^{-1}}$ est un terme correctif calculable, fonction des solutions admissibles des deux problèmes. A nouveau, $Q(\mathbf{u}_h) + Q_{hh}$ représente une nouvelle approximation de $Q(\mathbf{u})$.

Une méthode efficace pour obtenir des bornes précises de $Q(\mathbf{u})$ à partir de (2.111) consiste à raffiner localement le maillage du problème adjoint seul ; le terme $e_{ERC}(\tilde{\mathbf{u}}_{hh}, \hat{\sigma}_{hh})$ tend alors vers 0 et $Q(\mathbf{u}_h) + Q_{hh}$ converge vers $Q(\mathbf{u})$. Nous illustrons ce propos en reprenant l'exemple de la plaque trouée de la Section 2.2. On raffine le maillage du problème adjoint dans la zone localisée autour de ω en espace et $t = t_1$ en temps, en jouant sur trois paramètres (Figure 2.22) : (i) le rapport r_h entre la taille d'un élément du maillage initial et celle d'un élément du maillage raffiné ; (ii) le nombre N_L de couches d'éléments raffinés autour de ω ; (iii) le rapport r_t entre le pas de temps initial et celui raffiné. Les bornes adimensionnées $\bar{\xi}_{inf}$ et $\bar{\xi}_{sup}$ obtenues pour différents degrés de raffinement sont données dans le Tableau 2.6.

$r_h/N_L/r_t$	$\bar{\xi}_{inf}$	$\bar{\xi}_{sup}$	$r_h/N_L/r_t$	$\bar{\xi}_{inf}$	$\bar{\xi}_{sup}$	$r_h/N_L/r_t$	$\bar{\xi}_{inf}$	$\bar{\xi}_{sup}$
1/2/1	0,60	1,38	3/0/1	0,65	1,31	1/2/1	0,60	1,38
2/2/1	0,74	1,22	3/1/1	0,77	1,20	1/2/2	0,74	1,22
3/2/1	0,83	1,15	3/2/1	0,83	1,15	1/2/5	0,81	1,17
4/2/1	0,85	1,14	3/3/1	0,84	1,15	1/2/10	0,86	1,12

Tableau 2.6: Bornes sur $Q(\mathbf{u})$ en fonction du raffinement local du maillage du problème adjoint

On observe que le raffinement local de la discrétisation du problème adjoint permet d'améliorer considérablement la qualité des bornes sur $Q(\mathbf{u})$.

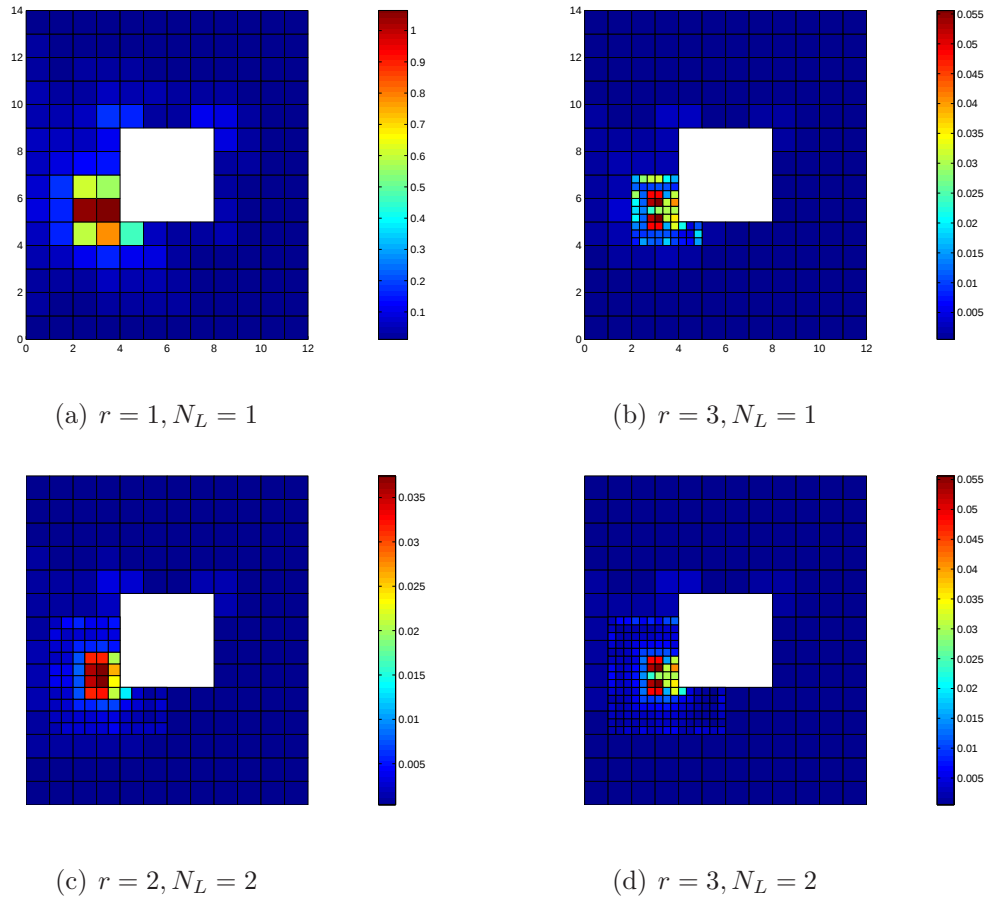


Figure 2.22: Distribution spatiale de $\tilde{e}_{ERC}$ pour divers raffinements du maillage du problème adjoint

On reprend aussi l'exemple de la plaque fissurée résolue par la méthode XFEM dans la Section 1.3. On se place dans le cas du mode mixte et les quantités d'intérêt choisies sont K_I et K_{II} . La Figure 2.23 montre l'évolution des bornes adimensionnées obtenues en fonction du nombre d'éléments utilisés pour résoudre le problème adjoint ; la discrétisation du problème primal reste inchangée. Deux méthodes d'approximation sont utilisées : la MEF standard et la XFEM avec enrichissement des nœuds proches de la pointe de fissure. On peut voir une nette amélioration des bornes obtenues pour les deux modes de sollicitation. L'enrichissement du problème adjoint permet d'avoir des bornes plus précises en gardant la même solution pour le problème de référence. Cela peut être intéressant si on souhaite réaliser une analyse sur une structure complexe : on peut commencer par une analyse EF assez grossière, et se focaliser sur la résolution précise du problème adjoint en utilisant un raffinement quand cela est nécessaire.

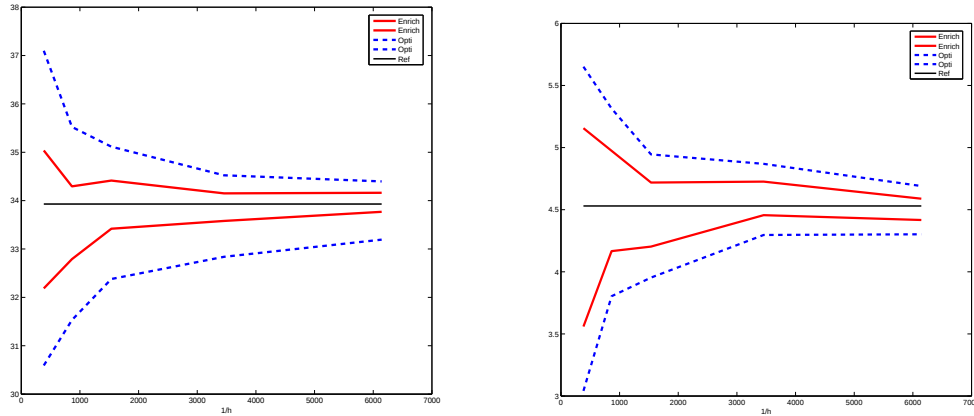


Figure 2.23: Comparaison des bornes de K_I (gauche) et K_{II} (droite) en mode mixte en fonction de la résolution du problème adjoint (classique ou enrichie)

3 Contributions à l'obtention de bornes pratiques

3.1 Procédure simplifiée pour le calcul des champs admissibles

La construction de solution admissible (champ de contrainte en particulier) est, comme illustré dans les sections précédentes, un point clé des méthodes d'estimation d'erreur basées sur l'ERC. Elle permet en particulier d'obtenir des bornes d'erreur garanties pour une large gamme de problèmes de mécanique, et sa qualité influe directement sur celle des bornes. Dans cet objectif, l'approche duale (MEF en contrainte) est optimale mais inadaptée pour la vérification car elle demande un second calcul complet de la solution, avec des outils peu présents dans les codes industriels. Une approche générale de construction, basée sur le post-traitement de la contrainte EF σ_h , a été introduite dans [Ladevèze et Leguillon 1983, Ladevèze et Pelle 2004]. Cette approche, récemment renommée EET (Element Equilibration Technique), se décompose en deux étapes :

1. des densités d'effort $\hat{\mathbf{F}}_h$, en équilibre avec le chargement extérieur, sont construites sur les bords des éléments du maillage ;
2. sur chaque élément E , un champ de contrainte $\hat{\sigma}_{h|E}$ vérifiant l'équilibre :

$$\mathbf{div} \hat{\sigma}_h + \mathbf{f}_d = \mathbf{0} \text{ dans } E \quad ; \quad \hat{\sigma}_h \mathbf{n} = \eta_E \hat{\mathbf{F}}_h \text{ sur } \partial E \quad (2.112)$$

est calculé, avec $\mathbf{n}$ la normale sortante à E et $\eta_E = \pm 1$ un scalaire assurant la continuité du vecteur contrainte. Le problème local associé est en pratique résolu par une technique quasi-explicite avec une base polynomiale, ou par une approche duale avec enrichissement de type p (degré $p + k$ pour les fonctions de forme).

La première étape s'appuie sur la condition d'équivalence énergétique (dite condition de prolongement) suivante :

$$\int_E (\hat{\sigma}_h - \sigma_h) \nabla \phi_i dE = 0 \quad \forall i \quad (2.113)$$

où ϕ_i est la fonction de forme EF du nœud i . Cette condition, qui assure l'équilibre de $\mathbf{F}_h$ sur E , aboutit à la résolution d'un système de la forme :

$$\sum_{r=1}^{R_n} \mathbf{b}_n^r(i) = \mathbf{Q}_{E_n}(i) \quad \forall n = 1, \dots, N \quad (2.114)$$

sur l'ensemble des N éléments connectés à chaque nœud i (Figure 2.24). R_n est le nombre de bords de l'élément E_n connectés au nœud i , $\mathbf{Q}_{E_n}(i) = \int_{E_n} (\sigma_h \nabla \phi_i - \mathbf{f}_d \phi_i) dE$, et les inconnues $\mathbf{b}_n^r(i)$ sont les projections des densités définies par $\hat{\mathbf{b}}_n^r(i) = \int_{\Gamma_{E_n}^r} \eta_{E_n} \hat{\mathbf{F}}_h \phi_i dS$. L'existence de solution pour chaque système est assurée par la propriété d'équilibre au sens des EF de σ_h , et l'unicité est obtenue éventuellement par minimisation d'une fonction coût.

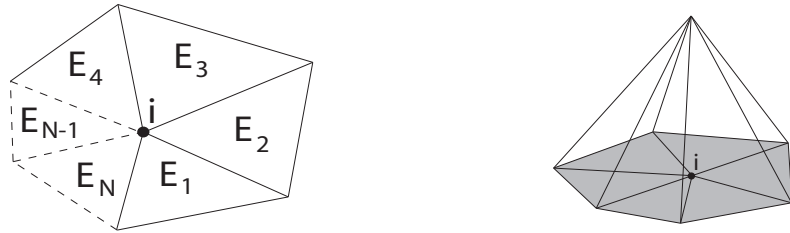


Figure 2.24: Patch Ω_i de N éléments autour du nœud sommet i (gauche) et fonction de forme linéaire φ_i associée (droite)

L'utilisation pratique de cette méthode est cependant limitée par sa complexité d'implémentation ; en effet, une technicité majeure réside dans le calcul des densités $\hat{\mathbf{F}}_h$ et la résolution spécifique des systèmes locaux suivant le type du nœud i (nœud sommet, de bord, ...). Les approches détaillées dans la suite visent à contourner cette difficulté.

3.1.1 Approche SPET ou *flux-free*

L'approche SPET (Star Patch Equilibration Technique) a la particularité d'être une technique de calcul de contrainte admissible qui n'a pas recours au calcul de densité sur le bord des éléments (d'où son nom de *flux-free*) ; elle a été initiée dans [Morin et al. 2001] pour les problèmes scalaires et nécessite de résoudre des problèmes locaux par patch d'éléments. Dans [Prudhomme et al. 2004], cette technique a été analysée et utilisée comme estimateur d'erreur basé sur les résidus (par sous-domaine) pour les problèmes elliptiques scalaires (équation de Poisson), avec calcul de bornes inférieure et supérieure de l'erreur. L'approche SPET a été étendue aux

problèmes d'élasticité dans [Pares *et al.* 2006] en introduisant dans l'équation des résidus le projecteur Π sur l'espace EF :

$$a(\mathbf{e}_h, \mathbf{v}) = \int_{\Omega} \mathbf{f}_d \cdot \mathbf{v} d\Omega + \int_{\partial_2 \Omega} \mathbf{F}_d \cdot \mathbf{v} dS - \int_{\Omega} \sigma_h : \varepsilon(\mathbf{v}) d\Omega = \mathcal{R}_h(\mathbf{v}) = \mathcal{R}_h(\mathbf{v} - \Pi \mathbf{v}) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (2.115)$$

La méthode SPET consiste à utiliser la partition d'unité $\sum_i \varphi_i(\mathbf{x}) = 1$ formée par les fonctions de forme φ_i linéaires :

$$a(\mathbf{e}_h, \mathbf{v}) = \sum_i \mathcal{R}_h(\varphi_i(\mathbf{v} - \Pi \mathbf{v})) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (2.116)$$

On cherche alors, sur chaque patch Ω_i constitué des éléments connectés au nœud sommet i (Figure 2.24), la solution $\mathbf{e}_i \in \mathcal{U}_{0|\Omega_i}$ vérifiant :

$$a_{|\Omega_i}(\mathbf{e}_i, \mathbf{v}) = \mathcal{R}_h(\varphi_i(\mathbf{v} - \Pi \mathbf{v})) \quad \forall \mathbf{v} \in \mathcal{U}_{0|\Omega_i} \quad (2.117)$$

Le problème local (2.117) est bien posé, sans avoir recours à une technique d'équilibrage, car l'opérateur Π assure l'annulation du second membre pour les fonctions tests $\mathbf{v}$ liées aux mouvements de corps rigide. La résolution de (2.117) est en général menée par la MEF avec une discrétisation fine (enrichissement de type h ou p), en bloquant éventuellement des mouvements pour assurer l'unicité. On peut alors en déduire le champ de contrainte statiquement admissible suivant :

$$\hat{\sigma}_h = \sigma_h + \mathbf{K} \varepsilon \left(\sum_i \mathbf{e}_i \right) \quad (2.118)$$

Dans [Cottreau *et al.* 2010], l'aspect strictement admissible de $\hat{\sigma}_h$ est conservé en utilisant une approche duale pour la résolution des problèmes par patch. Cette approche duale, valide pour un degré EF $p \geq 2$, consiste à trouver sur chaque patch Ω_i un champ π^i vérifiant :

$$\begin{aligned} \operatorname{div} \pi^i + \varphi_i(\mathbf{f}_d + \operatorname{div} \sigma_h) &= 0 \quad \text{dans } \Omega_i \\ [|\pi^i \mathbf{n}|] &= -\varphi_i[|\sigma_h \mathbf{n}|] \quad \text{sur } \Gamma_i - \partial\Omega_i \\ \pi^i \mathbf{n} &= \varphi_i(\mathbf{F}_d - \sigma_h \mathbf{n}) \quad \text{sur } \partial_2 \Omega \cap \partial\Omega_i \\ \pi^i \mathbf{n} &= \mathbf{0} \quad \text{sur } \partial\Omega_i - \partial_2 \Omega \end{aligned} \quad (2.119)$$

où Γ_i représente l'union des frontières des éléments composant Ω_i . Le champ π^i est cherché sur une base polynomiale, et on obtient $\hat{\sigma}_h = \sigma_h + \sum_i \pi_i$.

La méthode a été appliquée en viscoélasticité sur une structure rectangulaire avec rétrécissement local, en contraintes planes, avec un chargement normal de type rampe (Figure 2.25). Les paramètres matériau sont $E = 1$ Pa, $\nu = 0,3$ et on prend $\tau = 0,5$ s pour temps caractéristique de la viscosité. Le calcul de l'erreur en dissipation est réalisé avec trois discrétisations spatiales différentes i.e. 117, 419 et 1581 nœuds (ou 186, 744 et 2976 éléments). On donne sur la Figure 2.26 les visualisations spatiale et temporelle de l'erreur en dissipation cumulée $\int_0^T \int_{\Omega} e_{ERC}^2 d\Omega dt$.

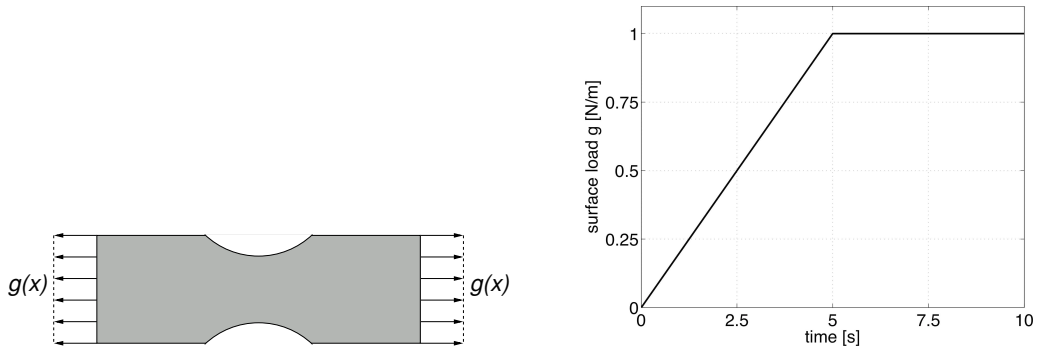


Figure 2.25: Structure étudiée (gauche) et évolution en temps du chargement (droite)

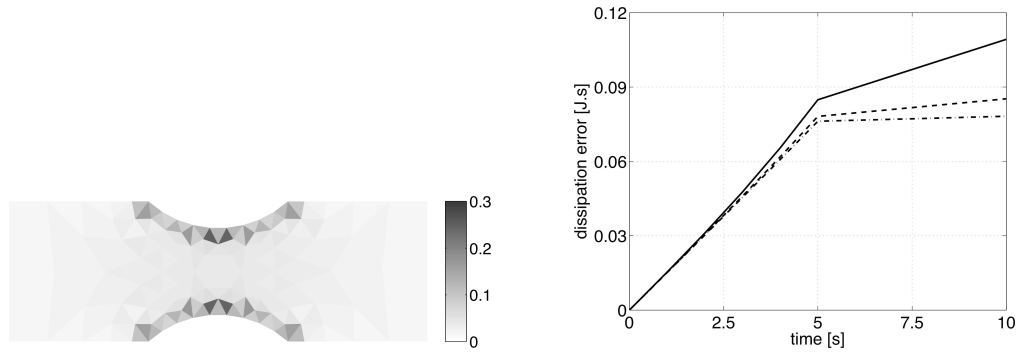


Figure 2.26: Contributions spatiales de l'erreur en dissipation à $t = 10$ s pour le maillage avec 117 nœuds (gauche) et évolution temporelle de l'erreur en dissipation cumulée pour les trois discrétisations spatiales (droite)

3.1.2 Nouvelle approche EESPT

Malgré ses atouts indéniables, la méthode SPET souffre du fait qu'elle nécessite de résoudre finement des problèmes à l'échelle de patches d'éléments, ce qui peut être très coûteux en 3D. Dans [Ladevèze *et al.* 2010], une nouvelle procédure hybride de construction de champs de contrainte admissibles a été introduite. Cette nouvelle procédure appelée EESPT (Element Equilibration+Star Patch Technique), qui est un intermédiaire entre les méthodes EET et SPET, réalise un bon compromis entre qualité des champs calculés, coût de calcul et facilité d'implémentation en ingénierie. De façon similaire à la méthode EET, l'approche EESPT conserve deux étapes et repose sur la construction de densités $\hat{\mathbf{F}}_h$ en équilibre sur le bord des éléments. Le principal changement concerne la façon de construire ces densités, avec une flexibilité accrue apportée par l'utilisation de la PUM ; cela mène à des problèmes par patch résolus de façon automatique et non-intrusive, avec les outils classiques des codes industriels. Le calcul de $\hat{\sigma}_h$, sur chaque élément et à partir des densités $\hat{\mathbf{F}}_h$, reste quant à lui inchangé et peut être parallélisé.

La condition de prolongement (2.113), qui assure l'équilibre des densités $\hat{\mathbf{F}}_h$, est réécrite sous la forme :

$$\int_{\Omega} (\hat{\sigma}_h - \sigma_h) : \varepsilon(\mathbf{v}_h) d\Omega = 0 \quad \forall \mathbf{v}_h \in \mathcal{U}_h^c \quad (2.120)$$

où $\mathcal{U}_h^c$ est l'espace de fonctions polynomiales linéaires par morceaux, non nécessairement continues à l'interface entre les éléments (on considère ici $p = 1$ mais l'extension est directe pour un degré plus élevé). Par définition des $\hat{\mathbf{F}}_h$, on obtient alors :

$$\sum_E \left[\int_{\partial E} \eta_E \hat{\mathbf{F}}_h \cdot \mathbf{v}_h dS - \int_E (\sigma_h : \varepsilon(\mathbf{v}_h) - \mathbf{f}_d \cdot \mathbf{v}_h) d\Omega \right] = 0 \quad \forall \mathbf{v}_h \in \mathcal{U}_h^c \quad (2.121)$$

L'idée majeure de la méthode EESPT est d'introduire à ce stade la partition d'unité formée par les fonctions de forme linéaires φ_i . Les densités sont alors mises sous la forme $\hat{\mathbf{F}}_h = \sum_i \hat{\mathbf{F}}_h^{(i)}$ où $\hat{\mathbf{F}}_h^{(i)}$, associé au nœud sommet i , vérifie le problème local suivant sur le patch Ω_i :

$$\sum_{E \subset \Omega_i} \int_{\partial E} \eta_E \hat{\mathbf{F}}_h^{(i)} \cdot \mathbf{v}_h dS = \int_{\Omega_i} (\sigma_h : \varepsilon(\varphi_i \mathbf{v}_h) - \mathbf{f}_d \cdot \varphi_i \mathbf{v}_h) d\Omega = F_{\Omega_i}(\varphi_i \mathbf{v}_h) \quad \forall \mathbf{v}_h \in \mathcal{U}_h^c \quad (2.122)$$

Les quantités $\hat{\mathbf{F}}_h^{(i)}$ sont cherchées sous forme polynomiale, de même degré que $\mathbf{u}_h$. Il est montré dans [Ladevèze *et al.* 2010] que les problèmes locaux (2.122) sont bien posés ; le cas $p = 1$ nécessite cependant un changement du second membre en $F_{\Omega_i}(\varphi_i \mathbf{v}_h(\mathbf{x}_i))$. Le système issu de la discrétisation de (2.122) peut s'écrire :

$$\hat{\mathbf{f}}_h^{(i)T} \mathbb{A}^{(i)} \mathbf{X}^{(i)} = \mathbf{R}^{(i)T} \mathbf{X}^{(i)} \quad \forall \mathbf{X}^{(i)} \quad (2.123)$$

où $\hat{\mathbf{f}}_h^{(i)}$ (resp. $\mathbf{X}^{(i)}$) recense les inconnues nodales de $\hat{\mathbf{F}}_h^{(i)}$ (resp. $\mathbf{v}_h$). Un intérêt majeur de la méthode EESPT est que le noyau de la matrice $\mathbb{A}^{(i)}$ est connu explicitement ; il correspond à l'ensemble des fonctions continues sur le patch Ω_i . Ce noyau est en pratique annulé en fixant de façon automatique certains degrés de liberté du problème local (Figure 2.27).

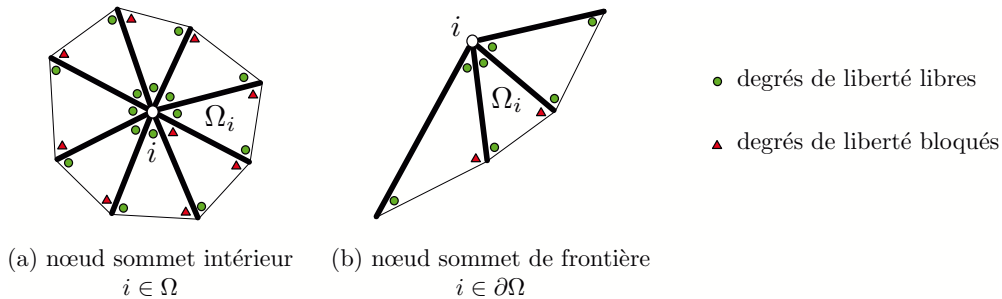


Figure 2.27: Visualisation du blocage de nœuds

L'unicité de $\hat{\mathbf{F}}_h^{(i)}$ (ou $\hat{\mathbf{f}}_h^{(i)}$) est quant à elle obtenue par minimisation d'une fonction coût de la forme :

$$\frac{1}{2} (\hat{\mathbf{f}}_h^{(i)} - \mathbf{f}_h^{(i)})^T \mathbb{P} (\hat{\mathbf{f}}_h^{(i)} - \mathbf{f}_h^{(i)}) \quad (2.124)$$

avec $\mathbb{P}$ une matrice diagonale; le problème final est alors un problème de minimisation de (2.124) sous la contrainte (2.123), qui est assez peu coûteux. Des détails sur l'implémentation pratique de la méthode, sur le choix de la fonction coût et sur le coût de calcul global en fonction du nombre de nœuds sommets du maillage peuvent être trouvés dans [Pled *et al.* 2011].

Une comparaison entre les méthodes EET, SPET et EESPT est également faite dans [Pled *et al.* 2011] sur plusieurs applications industrielles. Une de ces applications considère comme structure le moyeu de rotor principal de l'hélicoptère NH90 (Figure 2.28). Celui-ci est encastré sur un bord (plan coloré sur la Figure 2.28) et est soumis à une densité d'effort normal unitaire sur le bord opposé. Le matériau est élastique linéaire isotrope homogène avec $E = 1$ et $\nu = 0,3$.

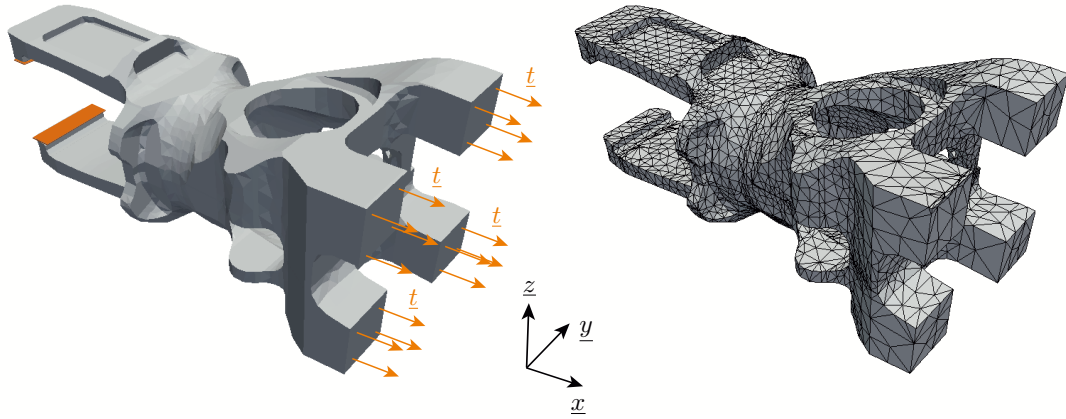


Figure 2.28: Problème de référence (gauche) et maillage EF associé (droite)

On s'intéresse spécifiquement à la zone centrale du moyeu, pour laquelle la solution EF et la carte d'erreur vraie sont données sur la Figure 2.29.

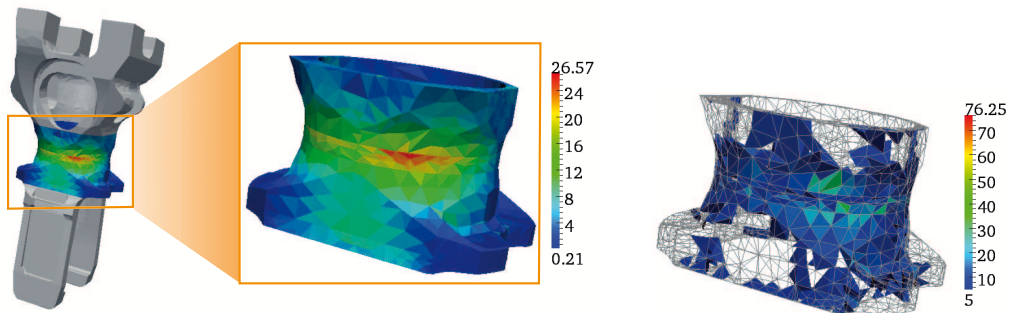


Figure 2.29: Amplitude de la contrainte de Von Mises FE (gauche) et distribution spatiale des contributions locales pertinentes à l'erreur vraie (droite)

Les méthodes EET, SPET et EESPT sont alors mises en œuvre pour déterminer une solution admissible $\hat{\sigma}_h$ à partir de σ_h , et ainsi définir un estimateur d'erreur global basé sur l'ERC. Les résultats sont donnés sur la Figure 2.30 et dans le Tableau 2.7. On observe que la méthode SPET est plus précise que les méthodes EET

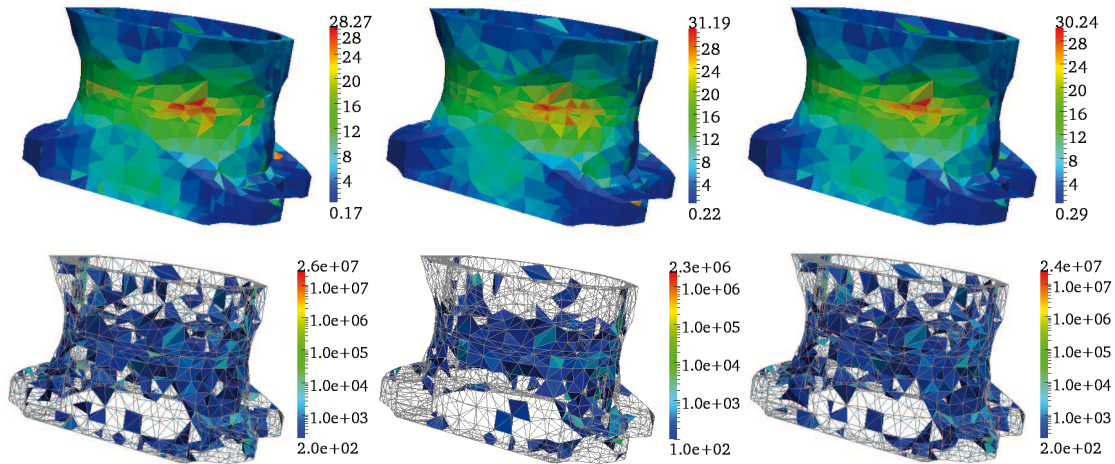


Figure 2.30: Amplitude de la contrainte de Von Mises admissible (haut) et distribution spatiale des contributions locales pertinentes aux estimateurs d'erreur (bas) calculées avec les méthodes EET (gauche), SPET (centre) et EESPT (droite)

Tableau 2.7: Comparaison des estimateurs d'erreur donnés par la EET, la SPET, et la EESPT

Méthode	Estimateur θ	Indice d'efficacité η	Temps CPU normalisé
EET	58 061	15,071	1,00
SPET	18 948	4,918	4,27
EESPT	42 078	10,922	1,00

et EESPT, mais qu'elle est aussi associée à un coût de calcul plus grand. La méthode EESPT, qui donne des résultats comparables à la méthode EET, semble quant à elle être un bon compromis entre qualité, coût de calcul et facilité d'implémentation.

Dans [Pled *et al.* 2012], une version améliorée de la méthode EESPT a été étudiée. Elle reprend les idées développées dans [Florentin *et al.* 2002] en considérant une condition de prolongement (dite affaiblie) ne portant que sur les fonctions de haut degré (nœuds non sommets). Cela résulte en une minimisation locale de l'énergie complémentaire et aboutit à des densités d'effort optimisées dans des régions sélectionnées, particulièrement celles avec des éléments distordus ou à fort gradient. La version améliorée de la méthode EESPT occasionnant un coût de calcul élevé, des critères géométriques ou d'erreur sont introduits pour sélectionner les zones dans lesquelles cette version doit être appliquée pour obtenir un bon compromis entre qualité et coût.

Un des exemples traités est celui d'une plaque trouée chargée avec une force de traction unitaire (Figure 2.31).

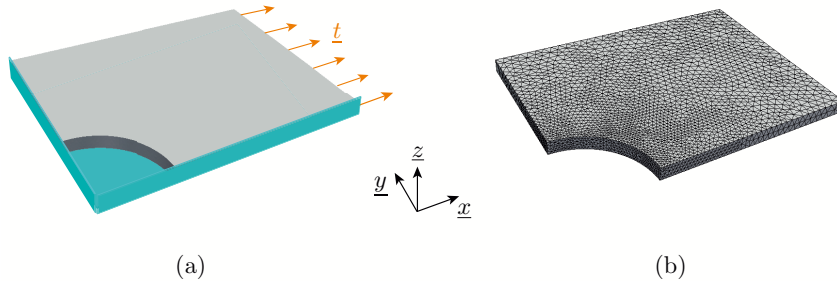


Figure 2.31: Modèle de plaque avec trou (gauche) et maillage EF associé (droite)

Les méthodes EET, EESPT standard et EESPT améliorée sont mises en œuvre pour calculer, à partir de σ_h , un champ de contrainte SA $\hat{\sigma}_h$ (Figure 2.32) et en déduire un estimateur d'erreur associé.

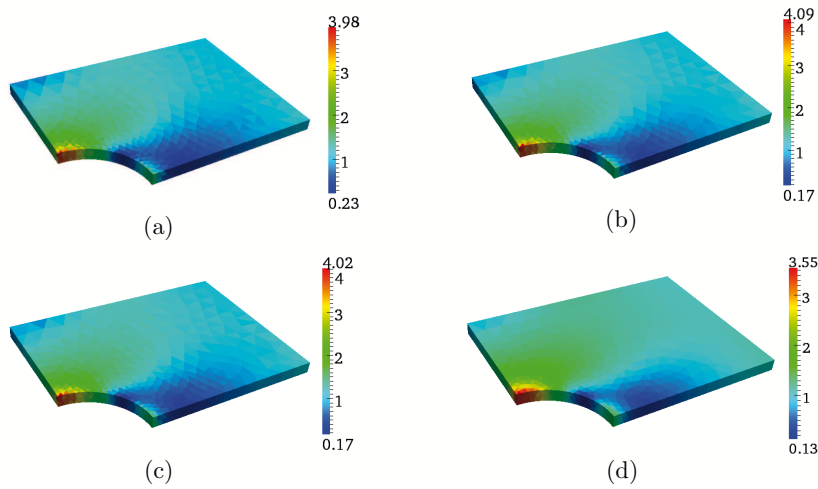


Figure 2.32: Norme du champ de contrainte EF (a), et des champs admissibles calculés avec la méthode EET (b), la méthode SPET standard (c) et sa version améliorée (d)

Deux critères sont alors introduits pour détecter les zones dans lesquelles il faut utiliser la version améliorée de la méthode EESPT. Le premier critère est basé sur la forme de l'élément i.e. son degré de distorsion ; ce critère est donc lié à la qualité locale du maillage. Le second critère considère les contributions locales d'erreur. Les qualités des estimations d'erreur obtenues et les coûts de calcul engendrés avec ces deux critères sont représentés sur les Figures 2.33 et 2.34. On observe d'une part que la version améliorée de la méthode EESPT augmente considérablement la qualité des bornes d'erreur calculées, mais à un surcoût non négligeable. D'autre part, on remarque qu'un compromis peut être trouvé en fixant convenablement les paramètres liés aux deux critères. Le meilleur compromis semble en outre être réalisé avec le second critère.

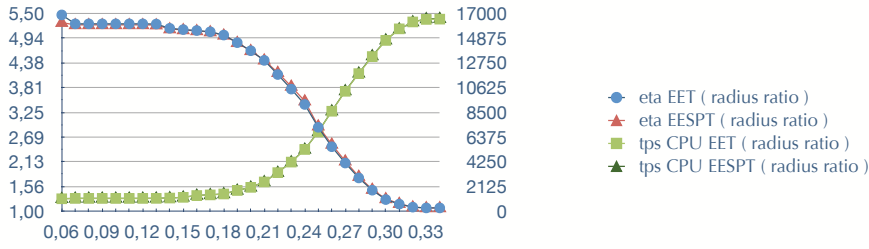


Figure 2.33: Indice d'efficacité et temps CPU normé en fonction du rapport de rayon et du nombre d'éléments

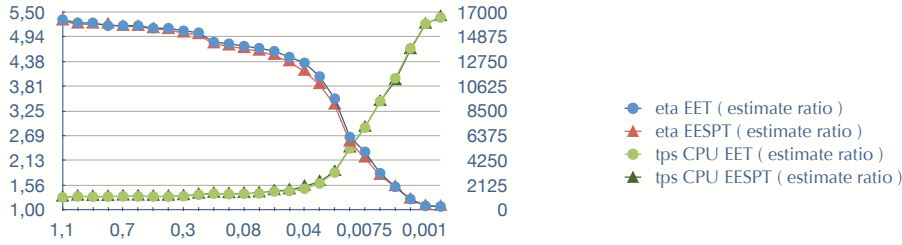


Figure 2.34: Indice d'efficacité et temps CPU normé en fonction de l'estimateur d'erreur local relatif

3.2 Technique non-intrusive pour la résolution de l'adjoint

3.2.1 Présentation générale

La méthode de résolution fine du problème adjoint par raffinement local, présentée dans la Section 2.3, a l'inconvénient de nécessiter un remaillage de la structure. Une alternative à ce remaillage, qualifiée de non-intrusive car ne modifiant pas le maillage initial, a été proposée dans [Chamoin et Ladevèze 2009, Ladevèze et Chamoin 2010]. Elle consiste en un enrichissement local, à l'aide de la PUM, de la solution adjointe avec des fonctions connues dites fonctions *handbook* qui visent à reproduire la partie à haut gradient de $(\tilde{\mathbf{u}}, \tilde{\sigma})$; l'approche est donc similaire à celle proposée dans la XFEM ou GFEM [Moës et al. 1999, Strouboulis et al. 2000] excepté qu'aucun degré de liberté supplémentaire n'est introduit ici, la singularité venant du chargement. On présente la méthode dans le cadre de l'élasticité linéaire, mais l'extension aux problèmes dépendant du temps est aisée [Chamoin et Ladevèze 2009, Waeytens et al. 2012].

Les fonctions d'enrichissement utilisées, notées par la suite $(\tilde{\mathbf{u}}^{hand}, \tilde{\sigma}^{hand})$, correspondent en pratique aux fonctions de Green généralisées et représentent une solution (quasi-)exacte dans un milieu (semi-)infini soumis au chargement de l'adjoint. Elles sont déterminées soit analytiquement (quand cela est possible), soit pré-calculées numériquement avec un maillage suffisamment fin et un domaine suffisamment large.

Un exemple de telle fonction, correspondant à un chargement de précontrainte localisée $\tilde{\sigma}_\Sigma$, est donné sur la Figure 2.35.

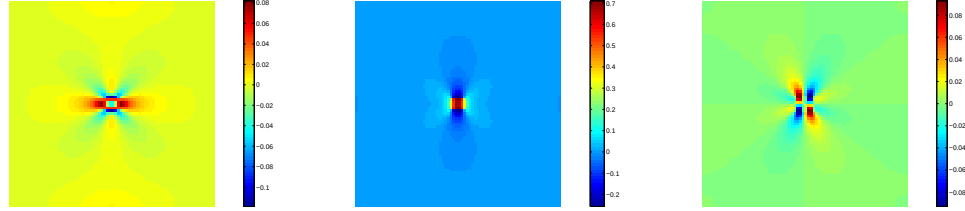


Figure 2.35: Champ de contrainte quasi-exact dans un domaine (quasi-) infini soumis à une précontrainte locale dans une région carrée : $\tilde{\sigma}_{xx}^{hand}$ (gauche), $\tilde{\sigma}_{yy}^{hand}$ (centre), $\tilde{\sigma}_{xy}^{hand}$ (droite)

Les fonctions *handbook* sont insérées dans la solution adjointe avec la partition d'unité définie par les fonctions de forme EF linéaires φ_i associées aux nœuds sommets i du maillage $\mathcal{M}_h$ ($\varphi_i(\mathbf{x}_j) = \delta_{ij}$). L'enrichissement est introduit localement, et par conséquent, seulement un ensemble de n^{PUM} nœuds sommets sont utilisés. La région d'enrichissement $\Omega^{PUM} \subset \Omega$ est définie par $\{\mathbf{x} \in \Omega, \sum_{i=1}^{n^{PUM}} \varphi_i(\mathbf{x}) \neq 0\}$, et peut être divisée en deux régions disjointes Ω_1^{PUM} et Ω_2^{PUM} telles que (Figure 2.36) :

$$\sum_{i=1}^{n^{PUM}} \varphi_i(\mathbf{x}) = \begin{cases} 1 & \text{dans } \Omega_1^{PUM} \\ a \in]0, 1[& \text{dans } \Omega_2^{PUM} \\ 0 & \text{dans } \Omega / (\overline{\Omega_1^{PUM} \cup \Omega_2^{PUM}}) \end{cases} \quad (2.125)$$

En pratique, Ω_1^{PUM} est telle qu'elle contient la zone d'intérêt Ω_Σ dans laquelle la quantité $Q(\mathbf{u})$ est définie, i.e. la région qui supporte les fonctions d'extraction.

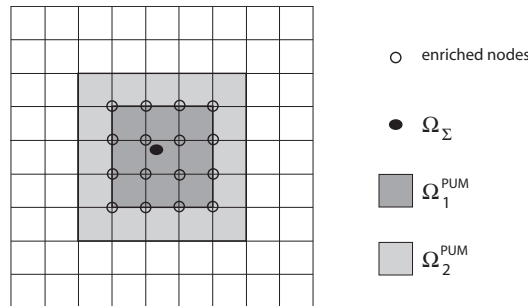


Figure 2.36: Définition des régions Ω_1^{PUM} , Ω_2^{PUM} et Ω_Σ dans la région enrichie Ω^{PUM}

De ce fait, la solution en déplacement du problème adjoint est cherchée sous la forme :

$$\tilde{\mathbf{u}}(\mathbf{x}) = \sum_{i=1}^{n^{PUM}} \tilde{\mathbf{u}}^{hand}(\mathbf{x}) \varphi_i(\mathbf{x}) + \tilde{\mathbf{u}}^{res}(\mathbf{x}) \quad (2.126)$$

où $\tilde{\mathbf{u}}^{res}$ est une solution résiduelle, généralement très régulière, à calculer. La partie d'enrichissement $\sum_{i=1}^{n^{PUM}} \tilde{\mathbf{u}}^{hand} \varphi_i$ permet de reproduire les forts gradients locaux de $\tilde{\mathbf{u}}$

tandis que la partie résiduelle $\tilde{\mathbf{u}}^{res}$ permet de corriger la partie d'enrichissement afin de vérifier les conditions limites du problème adjoint sur $\partial\Omega$. La nouvelle expression de $\tilde{\sigma}$ est déduite de (2.126) :

$$\tilde{\sigma}(\mathbf{x}) = \tilde{\sigma}_{PUM}^{hand}(\mathbf{x}) + \tilde{\sigma}^{res}(\mathbf{x}) \quad (2.127)$$

avec $\tilde{\sigma}_{PUM}^{hand} = \mathbf{K}\varepsilon(\sum_{i=1}^{n^{PUM}} \tilde{\mathbf{u}}^{hand}_i \varphi_i)$ et $\tilde{\sigma}^{res} = \mathbf{K}\varepsilon(\tilde{\mathbf{u}}^{res})$; évidemment, $\tilde{\sigma}_{PUM}^{hand} = \tilde{\sigma}^{hand}$ dans Ω_1^{PUM} .

Une fois que l'ensemble des n^{PUM} nœuds sommets d'enrichissement est défini, le nouveau problème adjoint consiste à trouver $\tilde{\mathbf{u}}^{res} \in \mathcal{U}_0$ tel que :

$$a(\mathbf{v}, \tilde{\mathbf{u}}^{res}) = Q(\mathbf{v}) - a(\mathbf{v}, \sum_{i=1}^{n^{PUM}} \tilde{\mathbf{u}}^{hand}_i \varphi_i) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (2.128)$$

La contrainte résiduelle $\tilde{\sigma}^{res} = \mathbf{K}\varepsilon(\tilde{\mathbf{u}}^{res})$ vérifie alors l'équation d'équilibre :

$$\begin{aligned} \int_{\Omega} \tilde{\sigma}^{res} : \varepsilon(\mathbf{v}) d\Omega &= \int_{\Omega} \left(\tilde{\sigma}_{\Sigma} : \varepsilon(\mathbf{v}) + \tilde{\mathbf{f}}_{\Sigma} \cdot \mathbf{v} \right) d\Omega - \int_{\Omega} \tilde{\sigma}_{PUM}^{hand} : \varepsilon(\mathbf{v}) d\Omega \\ &= \int_{\Omega_1^{PUM}} \left(\tilde{\sigma}_{\Sigma} : \varepsilon(\mathbf{v}) + \tilde{\mathbf{f}}_{\Sigma} \cdot \mathbf{v} - \tilde{\sigma}^{hand} : \varepsilon(\mathbf{v}) \right) d\Omega - \int_{\Omega_2^{PUM}} \tilde{\sigma}_{PUM}^{hand} : \varepsilon(\mathbf{v}) d\Omega \\ &= - \int_{\partial\Omega_1^{PUM}} \tilde{\sigma}^{hand} \mathbf{n}_{12} \cdot \mathbf{v} dS - \int_{\Omega_2^{PUM}} \tilde{\sigma}_{PUM}^{hand} : \varepsilon(\mathbf{v}) d\Omega \quad \forall \mathbf{v} \in \mathcal{U}_0 \end{aligned} \quad (2.129)$$

où $\mathbf{n}_{12}$ est la normale sortante de Ω_1^{PUM} vers Ω_2^{PUM} (voir Figure 2.36). Le chargement consiste en un chargement surfacique $-\tilde{\sigma}^{hand} \mathbf{n}_{12}$ sur $\partial\Omega_1^{PUM}$ et une précontrainte $-\tilde{\sigma}_{PUM}^{hand}$ dans Ω_2^{PUM} .

Une approximation fine $(\tilde{\mathbf{u}}_h^{res}, \tilde{\sigma}_h^{res})$ de la solution résiduelle peut être obtenue avec le maillage initial $\mathcal{M}_h$; la technique d'enrichissement est donc non-intrusive dans le sens où les opérateurs (matrice de raideur, table de connectivité) définis pour le problème primal peuvent être réutilisés sans changement pour résoudre le problème adjoint ; seul le chargement doit être modifié. Le calcul d'une contrainte admissible $\hat{\tilde{\sigma}}_h$ est également mené d'une façon non-intrusive : on définit tout d'abord un champ de contrainte $\hat{\tilde{\sigma}}_h^{res}$ vérifiant (2.129), avec la même méthode que celle utilisée pour calculer $\hat{\sigma}_h$; on définit alors :

$$\hat{\tilde{\sigma}}_h(\mathbf{x}) = \tilde{\sigma}_{PUM}^{hand}(\mathbf{x}) + \hat{\tilde{\sigma}}_h^{res}(\mathbf{x}) \quad (2.130)$$

Au final, on aboutit à la majoration suivante (cf. (2.14)) :

$$|Q(\mathbf{u}) - Q(\mathbf{u}_h) - Q_{corr}| \leq e_{ERC}(\mathbf{u}_h, \hat{\sigma}_h) \cdot e_{ERC}(\tilde{\mathbf{u}}_h^{res}, \hat{\tilde{\sigma}}_h^{res}) \quad (2.131)$$

où le second membre est indépendant de l'enrichissement $(\tilde{\mathbf{u}}^{hand}, \tilde{\sigma}^{hand})$. En pratique, l'erreur sur la solution résiduelle $e_{ERC}(\tilde{\mathbf{u}}_h^{res}, \hat{\tilde{\sigma}}_h^{res})$ est petite, et (2.131) fournit des bornes très précises sur l'erreur locale sans avoir besoin de changer le maillage. Cette méthode a été appliquée à divers cas, notamment pour l'estimation des facteurs d'intensité de contrainte en pointe de fissure (Figure 2.37) où la solution est singulière autour de la couronne ω (voir Section 1.3).

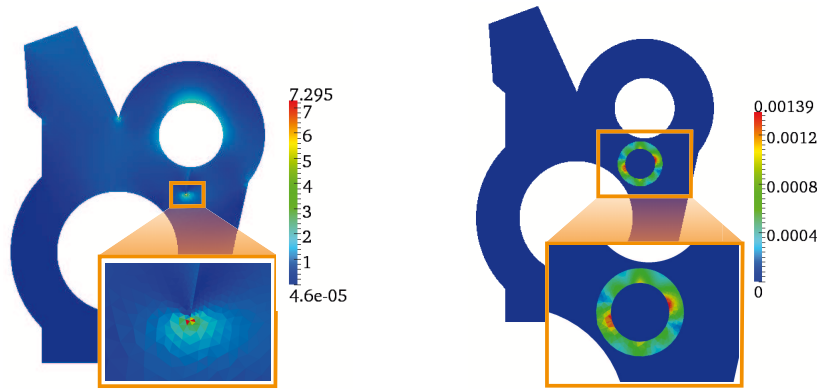


Figure 2.37: Distribution de l'erreur de discrétisation dans une structure fissurée (gauche) et solution adjointe (droite)

3.2.2 Application aux quantités ponctuelles

Lorsqu'on considère une quantité d'intérêt ponctuelle en espace (et/ou en temps), le chargement du problème adjoint associé est défini à partir de fonctions de Dirac. La solution adjointe correspondante est très singulière et souvent à énergie infinie [Cao et Kelly 2003]; il est de ce fait obsolète d'en calculer une solution approchée par la MEF. Une alternative classique consiste à régulariser la quantité d'intérêt, en utilisant par exemple des fonctions de pondération [Prudhomme et Oden 1999], de façon à préserver la régularité de la solution adjointe; la quantité ponctuelle est alors remplacée par une moyenne locale pondérée.

La technique d'enrichissement non-intrusive présentée précédemment permet, par extension directe, de considérer l'évaluation de l'erreur de discrétisation sur des quantités véritablement ponctuelles, sans aucune régularisation. Il suffit pour cela de prendre comme fonctions d'enrichissement les fonctions de Green. Sous certaines hypothèses, ces fonctions peuvent être déterminées analytiquement dans un domaine (semi-)infini [Mindlin 1936, Love 1944, Sneddon et Hill 1964, Graff 1975, Vijayakumar et Cormack 1987]; un exemple de telle fonction est donné dans la Figure 2.38. Dans des situations plus complexes (matériau anisotrope par exemple), les fonctions de Green peuvent parfois être obtenues implicitement, comme dans le cas des matériaux isotropes transverses [Dahan et Predeleanu 1980].

La majoration (2.131) est encore valable pour des quantités $Q(\mathbf{u})$ ponctuelles et les bornes peuvent être calculées malgré le fait que $(\tilde{\mathbf{u}}^{hand}, \tilde{\sigma}^{hand})$ puisse être à énergie infinie; en effet, seule la partie résiduelle $(\tilde{\mathbf{u}}_h^{res}, \hat{\sigma}_h^{res})$ de la solution adjointe admissible est utilisée pour obtenir les bornes d'erreur. Néanmoins, une certaine technicité est nécessaire pour déterminer Q_{corr} et le terme de chargement du problème adjoint qui impliquent la solution *handbook* singulière $(\tilde{\mathbf{u}}^{hand}, \tilde{\sigma}^{hand})$; cela est fait en pratique par sur-intégration.

Une application présentée dans [Ladevèze et Chamoin 2010] a concerné l'estimation de bornes d'erreur sur l'énergie de déformation en un point P ($Q(\mathbf{u}) = \frac{1}{2}\sigma :$

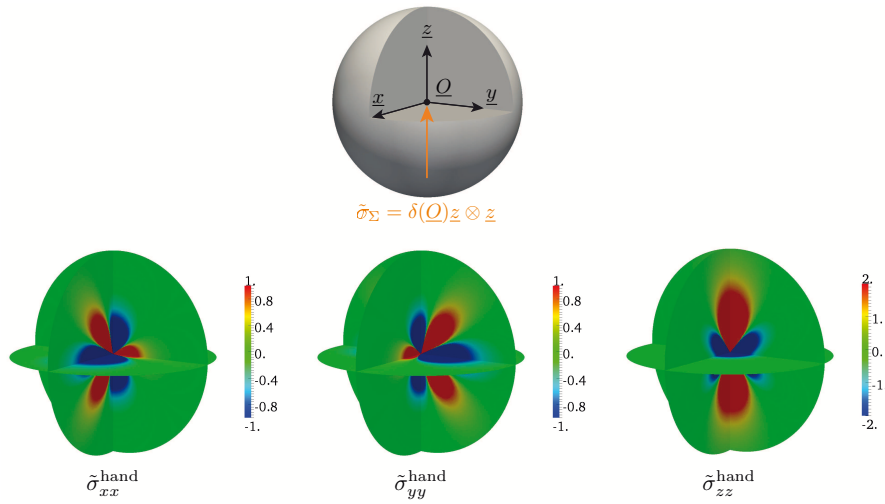


Figure 2.38: Champ de contrainte associé à une précontrainte ponctuelle dans un domaine infini en 3D

$\mathbf{K}^{-1}\sigma|_P$) dans une zone de singularité (Figure 2.39). En pratique, on considère P à une distance r du coin et on examine les bornes sur $Q(r)$ lorsque r tend vers 0.

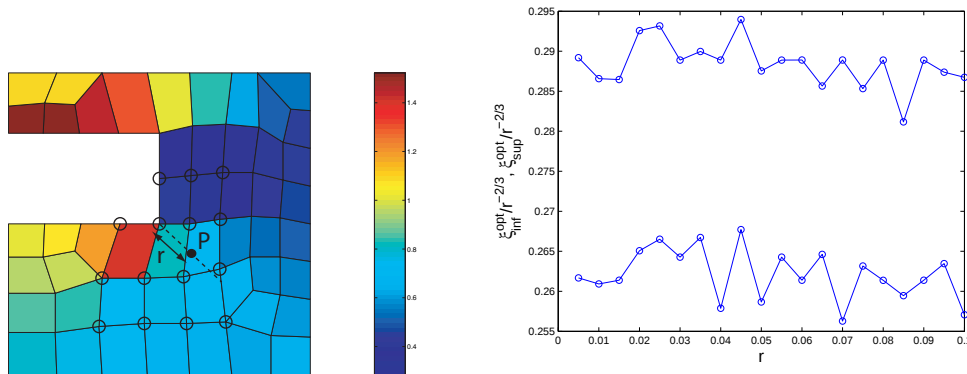


Figure 2.39: Définition de la distance radiale r avec l'ensemble des nœuds enrichis (cercles) pour la résolution du problème adjoint (gauche) et évolution des bornes inférieure et supérieure de $Q(r)/r^{-2/3}$ en fonction de r (droite)

L'enrichissement du problème adjoint utilise une fonction *handbook* analytique prenant en compte les conditions locales de bord libre, et est appliqué à 15 nœuds au voisinage du coin (voir Figure 2.39). La largeur de la zone d'enrichissement permet de considérer que la partie résiduelle ($\tilde{\mathbf{u}}^{res}, \tilde{\sigma}^{res}$) de la solution adjointe ne varie pas significativement avec la position du point P , et par conséquent qu'un seul calcul de la partie résiduelle est nécessaire. Seule la partie enrichie varie de ce fait avec P , et les bornes sont alors calculées comme :

$$Q_h(P) + Q_{corr}^{inf}(P) - \nu^{inf} \leq Q(P) \leq Q_h(P) + Q_{corr}^{sup}(P) + \nu^{sup} \quad (2.132)$$

où

- les termes correctifs Q_{corr}^{inf} et Q_{corr}^{sup} , qui dépendent de la solution enrichie du problème adjoint, sont calculés pour chaque position de P ;
- les termes ν^{inf} et ν^{sup} , qui dépendent de l'ERC donc de la solution résiduelle seule, sont indépendants de la position de P .

Avec cette hypothèse, qui est valide quand r est petit, le calcul des bornes en fonction de r est peu coûteux. Théoriquement, l'angle de la singularité implique que l'énergie de déformation $Q(r)$ varie en $r^{-2/3}$. De ce fait, on trace sur la Figure 2.39 l'évolution des bornes sur $Q(r)/r^{-2/3}$ en fonction de r . Ces bornes garanties peuvent fournir une information utile sur des quantités liées à la singularité, telles que les facteurs d'intensité de contrainte.

Dans [Waeytens *et al.* 2012], l'encadrement garanti de quantités ponctuelles a été mené pour des problèmes de visco-élastodynamique. Le traitement de la singularité de la solution adjointe en espace et en temps a été traité par enrichissement avec la fonction de Green associée à Q . Le calcul de cette fonction, en équilibre dynamique dans un milieu (semi-)infini, a été fait à partir du principe de correspondance [Alfrey 1944] et de la transformée de Laplace. Il a été montré que les bornes d'erreur locale se détériorent quand le modèle tend vers un modèle élastodynamique car dans ce cas la zone d'influence de la fonction de Green occupe tout le domaine Ω et il faut représenter correctement cette fonction après réflexion contre le bord $\partial\Omega$; lorsque les effets visqueux sont fortement présents, la zone d'influence de la fonction de Green est au contraire localisée car l'amplitude tend rapidement vers zéro (voir Figure 2.40 où on observe les fronts d'onde P et S sans et avec amortissement de 20%, avec un chargement ponctuel de type Dirac en espace et temps). Les bornes

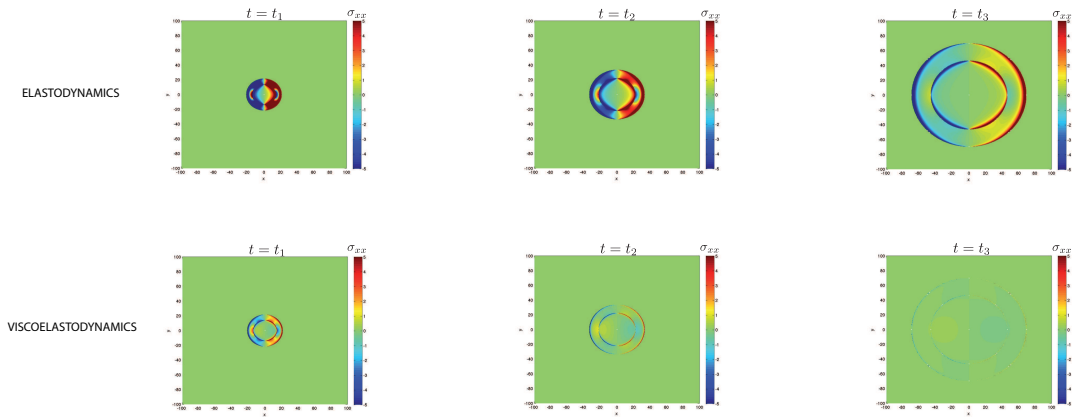


Figure 2.40: Evolution de la composante xx de la contrainte *handbook* pour un modèle sans (haut) et avec (bas) amortissement

d'erreur locale se dégradent également lorsque la quantité d'intérêt devient de plus en plus localisée en temps (singularité en temps de la solution adjointe de plus en plus sévère) ; cela illustre que considérer une quantité ponctuelle en temps pour les problèmes de dynamique n'a pas vraiment de sens physique.

4 Conclusions partielles et perspectives

On a montré dans ce chapitre des travaux réalisés pour tendre vers des méthodes d'estimation d'erreur locale de discrétisation à la fois robustes et utilisables. Plusieurs axes de développement futur peuvent être identifiés.

Tout d'abord, toutes les études se sont placées dans le cadre de modèles matériau linéaires, et il convient d'étendre ces études à des comportements plus complexes. Un cadre assez général basé sur l'ERC et valide pour les comportements matériaux non-linéaires monotones (exemple de la (visco-)plasticité) a été donné dans [Ladevèze 2008] et analysé dans [Ladevèze *et al.* 2012], même si des travaux supplémentaires restent nécessaires pour rendre les bornes utilisables. Néanmoins, encore aucun de travail de recherche ne fournit à l'heure actuelle des bornes d'erreur dans le cas de matériaux adoucissants (endommagement), d'instabilités (flambage par exemple) ou de grandes transformations ; dans tous ces cas, seuls des outils heuristiques existent du fait de la perte de convexité. D'autres terrains inexploités, plus abordables *a priori*, concernent l'analyse isogéométrique récemment développée et en plein essor, les méthodes de volumes finis, ou les problèmes de contact. Enfin, des actions sont en cours pour l'analyse modale (accord-cadre EDF) ou la théorie des plaques/coques (collaboration avec P.T. Mendonça).

Un deuxième axe de développement porte sur la construction de champs de contrainte admissibles qui est une étape essentielle pour avoir des bornes d'erreur garanties. Des optimisations des méthodes de construction actuelles sont sans doute encore possibles, et il serait bénéfique d'analyser en détail les outils proposés par la communauté de mathématiques appliquées notamment ceux utilisant le cadre des méthodes de Galerkin discontinues (éléments de Raviart-Thomas) qui ne nécessitent que des polynômes de bas degré [Cochez-Dhondt et Nicaise 2008, Vohralik 2011]. Enfin, le transfert des outils dans les codes industriels est une finalité majeure. Des premières collaborations ont été instaurées pour une implémentation dans les codes SAMCEF et ASTER, mais il reste beaucoup de travail pour que les outils de vérification soient utilisés largement ; cela passe par un dialogue accru entre le monde de la recherche et celui de l'industrie sur l'utilité du contrôle des calculs, ainsi que par des démonstrations des outils de recherche sur des applications industrielles complexes.

Sommaire

1	Maîtrise de la procédure de couplage entre modèles	62
1.1	Couplage particulaire/continu par la méthode Arlequin	63
1.2	Choix de l'opérateur de couplage	67
1.3	Couplage dans le cadre stochastique	70
1.4	Prise en compte des incompatibilités entre modèles	80
2	Stratégie adaptative de couplage	92
2.1	Méthode générale	92
2.2	Application au couplage particulaire/continu	97
2.3	Application au couplage transport/diffusion	101
2.4	Application au couplage stochastique/déterministe	108
2.5	Procédure optimale d'adaptation	115
3	Contrôle des modèles PGD	121
3.1	Présentation de la PGD	123
3.2	Estimation d'erreur locale par l'ERC	125
3.3	Séparation des sources et adaptation	127
4	Conclusions partielles et perspectives	131

Nous exposons dans ce chapitre des travaux réalisés pour contrôler la réduction de modèle i.e. le passage d'un modèle mathématique fin, pris comme référence mais difficile voire impossible à traiter, à un modèle mathématique plus abordable et considéré pour la simulation numérique. On distingue ici deux classes de réduction de modèle : (i) celle par changement d'échelle i.e. par remplacement, dans une grande majorité de la structure, du modèle initial fin par un modèle plus grossier issu de techniques d'homogénéisation ; (ii) celle par décomposition spectrale de la solution. Pour le premier cas, on s'intéresse à la maîtrise du couplage multiéchelle de modèles, en mettant en place des procédures propres de couplage puis une méthode d'estimation d'erreur et d'adaptation vis-à-vis de quantités d'intérêt. Pour le second cas, on aborde le contrôle d'une technique de représentation de solution par séparation de variables qui a connu récemment un essor important sous le nom de PGD (*Proper Generalized Decomposition*).

1 Maîtrise de la procédure de couplage entre modèles

Il y a des familles de problèmes pour lesquelles le modèle de simulation initial reste intraitable avec les capacités numériques actuelles. Un modèle plus grossier, souvent associé à une procédure d'homogénéisation ou d'analyse limite, est alors introduit et occasionne une approche multiéchelle du problème. Dans cette partie, on se focalise essentiellement sur les modèles particuliers qui sont de plus en plus utilisés pour décrire le comportement des matériaux à petite échelle (nanotechnologies, fissuration,...) ; un exemple est donné sur la Figure 3.1 dans le cadre de la modélisation de structures polymères.

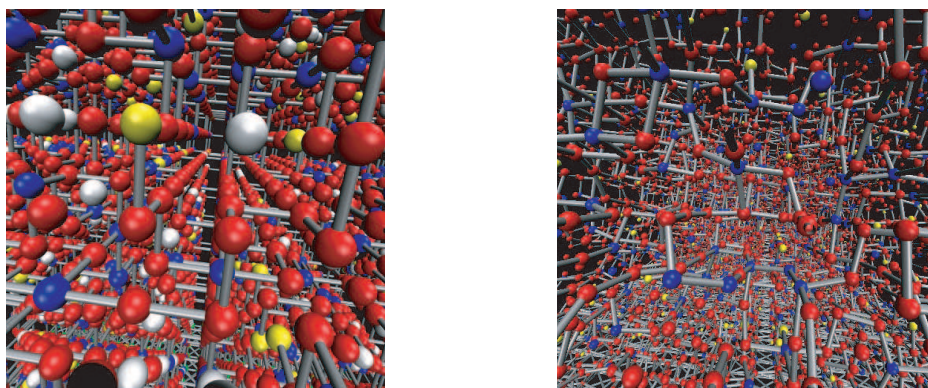


Figure 3.1: Modélisation d'une structure polymère 3D

Pour ce genre de modèle, qui implique généralement un nombre prohibitif d'inconnues, l'approche multiéchelle est vitale ; elle est réalisée par un couplage avec un modèle continu classique qui est placé loin des zones d'intérêt locales du problème [Curtin et Miller 2003, Liu *et al.* 2004, Fish 2006]. De ce fait, la taille du problème est réduite en conservant le modèle particulier uniquement dans des zones critiques où les phénomènes à petite échelle sont importants et où l'utilisation de

modèles particuliers est pertinente; cela permet de mener des simulations à des échelles spatiales et temporelles plus grandes. La problématique traitée dans cette section concerne la procédure de couplage, dont la qualité peut affecter grandement la qualité de la simulation.

1.1 Couplage particulaire/continu par la méthode Arlequin

Le couplage de modèles particulaire et continu concurrents a été intensivement et diversement traité; citons pour exemples les méthodes de couplage *Quasi-Continuum* [Tadmor *et al.* 1996], *handshake* [Broughton *et al.* 1999], *bridging-scale* [Wagner et Liu 2003], ou *bridging-domain* [Xiao et Belytschko 2004]. On considère ici un processus de couplage basé sur la méthode Arlequin, initialement développée pour le couplage entre deux modèles continus [Ben Dhia 1998, Ben Dhia et Rateau 2001 - 2005] puis étendue au couplage particulaire/continu [Bauman *et al.* 2008, Prudhomme *et al.* 2008, Bauman *et al.* 2009, Prudhomme *et al.* 2009, Chamoïn *et al.* 2010], qui mélange les modèles concurrents dans une région de couplage volumique. L'approche est particulièrement attirante et adaptée pour le couplage particulaire/continu car elle ne nécessite pas de vérifier l'hypothèse de Cauchy-Born au niveau de la transition entre les modèles.

1.1.1 Présentation des modèles concurrents

On considère un ensemble $\mathcal{N}$ fait de n particules distribuées dans un réseau régulier de $\mathbb{R}^d$ ($d = 1, 2$, ou 3) comme illustré sur la Figure 3.2. La position de la particule i dans la configuration de référence est donnée par les coordonnées $\mathbf{x}_i$ et $\mathbf{z}_i$ représente le déplacement de cette particule après déformation. On note $\mathcal{N}_\partial$ le sous-ensemble des particules de $\mathcal{N}$ qui sont sur les bords du réseau (ou proche, dans le cas d'interactions à longue distance), et on suppose que les particules dans $\mathcal{N}_\partial$ satisfont des conditions de Dirichlet homogènes. Chaque particule dans $\mathcal{N} \setminus \mathcal{N}_\partial$ est quant à elle soumise à une force ponctuelle externe $\mathbf{f}_i$ (nulle pour la plupart des particules).

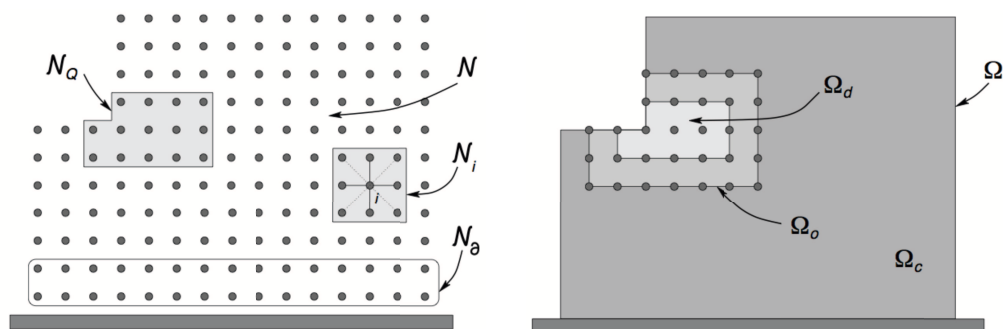


Figure 3.2: Modèle particulaire avec région d'intérêt (gauche) et couplage entre le modèle particulaire et le modèle continu (droite)

Les interactions entre les particules i et j sont décrites par des potentiels de

paires V_{ij} , tels que le potentiel harmonique :

$$V_{ij}(\mathbf{z}_i, \mathbf{z}_j) = \frac{1}{2} k_{ij} (|\mathbf{x}_j + \mathbf{z}_j - \mathbf{x}_i - \mathbf{z}_i| - |\mathbf{x}_j - \mathbf{x}_i|)^2 \quad (3.1)$$

où $|\bullet|$ est la norme euclidienne et k_{ij} un coefficient de raideur. En pratique, on limite les interactions de chaque particule i avec le sous-ensemble $\mathcal{N}_i \subset \mathcal{N}$ de particules voisines en négligeant les interactions à très grande distance.

En définissant l'énergie de déformation $E(\mathbf{z})$ du système particulaire, avec $\mathbf{z} = [\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_n]$ le vecteur déplacement global, ainsi que l'espace vectoriel $\mathcal{Y}$:

$$\mathcal{Y} = \{\mathbf{y} \in \mathbb{R}^{d \times n}; \mathbf{y}_i = \mathbf{0} \text{ pour chaque particule } i \text{ dans } \mathcal{N}_\partial\} \quad (3.2)$$

le problème consiste à trouver la configuration déformée du système à l'équilibre i.e. $\mathbf{y} \in \mathcal{Y}$ tel que :

$$\mathbf{z} = \operatorname{argmin}_{\mathbf{y} \in \mathcal{Y}} \left[E(\mathbf{y}) - \sum_{i=1}^n \mathbf{f}_i \cdot \mathbf{y}_i \right] \quad ; \quad E(\mathbf{y}) = \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n V_{ij}(\mathbf{y}_i, \mathbf{y}_j) \quad (3.3)$$

Le problème peut être écrit sous la forme faible suivante :

$$\text{Trouver } \mathbf{z} \in \mathcal{Y} \text{ tel que } a(\mathbf{z}; \mathbf{y}) = l(\mathbf{y}) \quad \forall \mathbf{y} \in \mathcal{Y} \quad (3.4)$$

où $a(\mathbf{z}; \mathbf{y}) = \sum_{i=1}^n \frac{\partial E}{\partial \mathbf{z}_i}(\mathbf{z}) \cdot \mathbf{y}_i$ et $l(\mathbf{y}) = \sum_{i=1}^n \mathbf{f}_i \cdot \mathbf{y}_i$ sont des formes semi-linéaire et linéaire, respectivement.

Résoudre (3.4) est intraitable dans la plupart des applications pratiques. Néanmoins, on s'intéresse généralement à des caractéristiques spécifiques de la solution globale i.e. des quantités d'intérêt $Q(\mathbf{z})$ (moyenne locale de déplacement, elongation d'une liaison, ...) définies dans une région critique de la structure incluant des défauts (fissure, trou, dislocation) ou de fortes fluctuations de potentiels; cette région est par la suite représentée par un sous-ensemble $\mathcal{N}_Q$ de particules (Figure 3.2). Par conséquent, on conserve le modèle particulaire dans une région locale $\Omega_d \subset \Omega$, composée d'un sous-ensemble $\mathcal{N}_d \subset \mathcal{N}$ de $m \ll n$ particules avec $\mathcal{N}_Q \subset \mathcal{N}_d$. L'énergie de déformation associée à $\mathcal{N}_d$ est alors donnée par :

$$E_d(\mathbf{z}) = \frac{1}{2} \sum_{i=1}^m \sum_{j=1}^m V_{ij}(\mathbf{z}_i, \mathbf{z}_j) \quad ; \quad \mathbf{z} \in \mathcal{W} = \mathbb{R}^{d \times m} \quad (3.5)$$

Dans le reste du domaine (sous-région $\Omega_c \subset \Omega$ avec $\Omega = \Omega_c \cup \Omega_d$), on remplace le modèle particulaire par un modèle plus grossier continu dont le rôle est de propager l'information à grande échelle seulement. Ce dernier modèle, souhaité compatible avec le modèle particulaire au sens de l'homogénéisation, est en pratique obtenu par calibration à partir de tests sur un volume élémentaire représentatif (VER). Le choix de comportement pour le modèle continu est large et reste encore un sujet de recherche ouvert (voir [Blanc *et al.* 2007]). Par exemple, un modèle hyperélastique

isotrope de type Mooney-Rivlin est considéré dans [Bauman *et al.* 2009] pour décrire le comportement à grande échelle d'un polymère ; l'énergie de déformation associée est de la forme :

$$E_c = C_1(I_1 - 3) + C_2(I_2 - 3) + C_3(J - 1)^2 - (2C_1 + 4C_2) \ln J \quad (3.6)$$

avec I_i les invariants principaux du tenseur de déformation de Green $\mathbb{C}$ et $J = \sqrt{I_3}$. On emploie dans la suite un modèle élastique linéaire, avec pour énergie de déformation :

$$E_c(\mathbf{u}) = \frac{1}{2} \int_{\Omega_c} \varepsilon(\mathbf{u}) : \mathbf{K} \varepsilon(\mathbf{u}) d\Omega \quad (3.7)$$

en supposant que les déformations dans Ω_c sont petites. On suppose par simplicité que $\partial\Omega \subset \partial\Omega_c$; l'espace des déplacements admissibles est donc défini comme :

$$\mathcal{U} = \{\mathbf{v} \in (H^1(\Omega_c))^d; \mathbf{v} = \mathbf{0} \text{ sur } \partial\Omega_c \cap \partial\Omega\} \quad (3.8)$$

1.1.2 Processus de couplage

Le couplage entre les modèles particulaire et continu est réalisé avec le cadre Arlequin [Ben Dhia 1998, Ben Dhia et Rateau 2005, Bauman *et al.* 2008]. Ce cadre nécessite une région de recouvrement $\Omega_o = \Omega_c \cap \Omega_d \neq \emptyset$ (Figure 3.2) et est basé sur les deux principes suivants :

1. Une énergie de déformation globale E_g du système couplé est définie en pondérant l'énergie de chaque modèle à l'aide d'une partition d'unité (en moyenne). On introduit alors la fonction poids $\alpha_c(\mathbf{x})$ associée au modèle continu et la fonction poids α_d^{ij} associée à chaque liaison entre les particules i et j :

$$\alpha_c(\mathbf{x}) = \begin{cases} 1 & \forall \mathbf{x} \in \Omega_c \setminus \Omega_o \\ 0 & \forall \mathbf{x} \in \Omega_d \setminus \Omega_o \end{cases} ; \quad \alpha_d^{ij} = 1 - \frac{1}{|S_{ij}|} \int_{S_{ij}} \alpha_c(\mathbf{x}) d\Omega \quad (3.9)$$

S_{ij} étant la région physique d'influence de la liaison (Figure 3.3). D'une part,

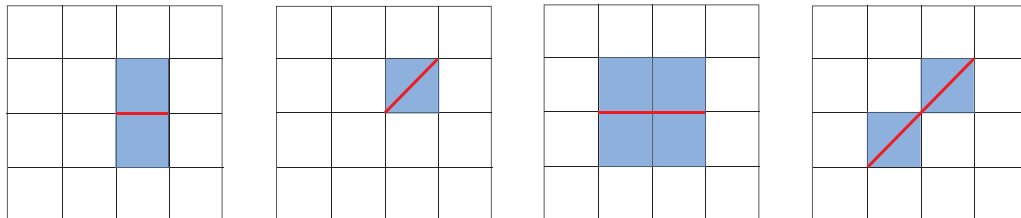


Figure 3.3: Exemples de régions S_{ij} : le segment épais représente la liaison tandis que le domaine grisé représente la région S_{ij}

la forme de la fonction α_c dans Ω_o est libre ; on la choisit généralement sous forme polynomiale de bas degré (voir Figure 3.4).

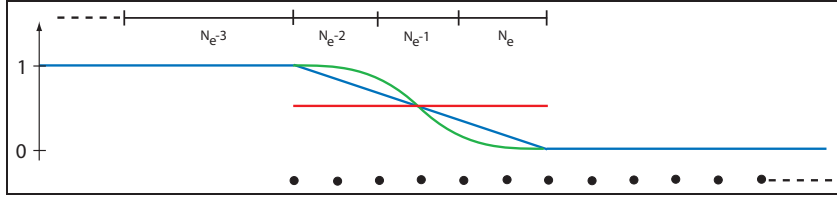


Figure 3.4: Couplage 1D avec fonction poids $\alpha_c(x)$ constante (rouge), linéaire (bleu), ou cubique (vert)

D'autre part, la définition de α_d^{ij} par une intégrale est une façon naturelle de pondérer l'énergie des particules, en définissant implicitement une densité d'énergie à l'échelle particulaire [Prudhomme *et al.* 2009].

L'énergie de déformation E_g est ensuite exprimée sous la forme :

$$E_g(\mathbf{u}, \mathbf{z}) = \frac{1}{2} \int_{\Omega_c} \alpha_c(\mathbf{x}) \boldsymbol{\varepsilon}(\mathbf{u}) : \mathbf{K} \boldsymbol{\varepsilon}(\mathbf{u}) d\Omega + \sum_{i=1}^m \sum_{j=1}^m \alpha_d^{ij} V_{ij}(\mathbf{z}_i, \mathbf{z}_j) \quad (3.10)$$

2. Les solutions $\mathbf{u}$ et $\mathbf{z}$ des modèles continu et particulaire, respectivement, sont mises en correspondance sur la région de recouvrement Ω_o . On introduit pour cela un opérateur d'interpolation linéaire Π_o qui transforme le vecteur déplacement discret $\mathbf{z}$ en un champ de déplacement continu et linéaire par morceaux $\Pi_o \mathbf{z}$ sur Ω_o , ainsi qu'un opérateur de restriction R_o qui restreint le déplacement continu $\mathbf{u}$ à Ω_o . On impose alors la contrainte :

$$\|R_o \mathbf{u} - \Pi_o \mathbf{z}\|_o = 0 \quad (3.11)$$

où la norme $\|\bullet\|_o$ est souvent celle associée à la forme bilinéaire :

$$b(\boldsymbol{\mu}, \mathbf{v}) = \int_{\Omega_o} (\boldsymbol{\mu} \cdot \mathbf{v} + \beta \nabla \boldsymbol{\mu} : \nabla \mathbf{v}) d\Omega \quad \forall \boldsymbol{\mu}, \mathbf{v} \in \mathcal{X} \quad (3.12)$$

définie sur $\mathcal{X} = (H^1(\Omega_o))^d$, β étant un facteur d'échelle.

En supposant que le chargement externe $\mathbf{f}_i$ est appliqué aux particules dans la région critique $\Omega_d \setminus \Omega_o$ seulement, le problème issu du couplage Arlequin consiste à trouver $(\mathbf{u}, \mathbf{z}) \in \mathcal{U} \times \mathcal{W}$ minimisant l'énergie potentielle globale sous la contrainte (3.11) :

$$(\mathbf{u}, \mathbf{z}) = \underset{\substack{(\mathbf{v}, \mathbf{w}) \in \mathcal{U} \times \mathcal{W} \\ \|R_o \mathbf{v} - \Pi_o \mathbf{w}\|_o = 0}}{\operatorname{argmin}} \left[E_g(\mathbf{v}, \mathbf{w}) - \sum_{i=1}^m \mathbf{f}_i \cdot \mathbf{w}_i \right] \quad (3.13)$$

L'introduction du lagrangien associé :

$$\mathcal{L}(\mathbf{v}, \mathbf{w}, \boldsymbol{\mu}) = \left[E_g(\mathbf{v}, \mathbf{w}) - \sum_{i=1}^m \mathbf{f}_i \cdot \mathbf{w}_i \right] - b(\boldsymbol{\mu}, R_o \mathbf{v} - \Pi_o \mathbf{w}) \quad (3.14)$$

permet de reformuler (3.13) sous forme faible ; celle-ci consiste à trouver $(\mathbf{u}, \mathbf{z}, \boldsymbol{\lambda}) \in \mathcal{U} \times \mathcal{W} \times \mathcal{X}$ tel que :

$$a_0((\mathbf{u}, \mathbf{z}, \boldsymbol{\lambda}); (\mathbf{v}, \mathbf{w}, \boldsymbol{\mu})) = l(\mathbf{w}) \quad \forall (\mathbf{v}, \mathbf{w}, \boldsymbol{\mu}) \in \mathcal{U} \times \mathcal{W} \times \mathcal{X} \quad (3.15)$$

avec

$$\begin{aligned} a_0((\mathbf{u}, \mathbf{z}, \boldsymbol{\lambda}); (\mathbf{v}, \mathbf{w}, \boldsymbol{\mu})) = & \int_{\Omega_c} \alpha_c(\mathbf{x}) \varepsilon(\mathbf{u}) : \mathbf{K} \varepsilon(\mathbf{v}) d\Omega + \sum_{k=1}^m \frac{\partial}{\partial \mathbf{z}_k} \left(\sum_{i=1}^m \sum_{j=1}^m \alpha_d^{ij} V_{ij}(\mathbf{z}_i, \mathbf{z}_j) \right) \cdot \mathbf{w}_k \\ & - b(\boldsymbol{\lambda}, R_o \mathbf{v} - \Pi_o \mathbf{w}) - b(\boldsymbol{\mu}, R_o \mathbf{u} - \Pi_o \mathbf{z}) \end{aligned} \quad (3.16)$$

Le problème couplé (3.15) est ensuite résolu par la MEF en introduisant les espaces $\mathcal{U}_h \subset \mathcal{U}$ et $\mathcal{X}_h \subset \mathcal{X}$.

1.2 Choix de l'opérateur de couplage

Un paramètre important du couplage par la méthode Arlequin est l'opérateur utilisé pour définir la contrainte de couplage (3.11) ainsi que le choix de la discrétisation de l'espace $\mathcal{X}$ des multiplicateurs de Lagrange. D'une part, il est prouvé dans [Bauman *et al.* 2008] que le problème (3.15) et sa version discrétisée ne sont pas bien posés pour une norme de type L^2 . D'autre part, les analyses numériques menées dans [Guidault et Belytschko 2007, Bauman *et al.* 2008] montrent que la discrétisation associée à $\mathcal{X}_h$ doit être au moins aussi grossière que celle associée à $\mathcal{U}_h$ et doit être liée à la taille du VER utilisé pour la calibration du modèle continu, afin d'éviter les phénomènes de verrouillage dans la zone de couplage Ω_o . Si la discrétisation de $\mathcal{X}_h$ est trop fine, la solution dans le milieu continu ne reproduit pas le comportement à grande échelle et pollue la solution du modèle couplé (voir illustration sur la Figure 3.5) ; il en résulte que la solution de (3.15) dépend fortement du choix de $\mathcal{X}_h$.

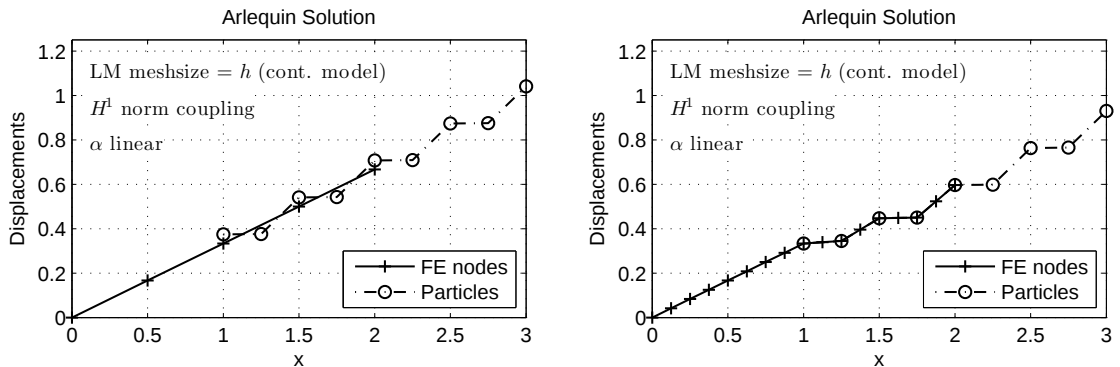


Figure 3.5: Solutions du problème couplé Arlequin 1D avec différentes discrétisations pour définir $\mathcal{X}_h$ (tailles d'éléments identiques pour $\mathcal{X}_h$ et $\mathcal{U}_h$) : non verrouillage lorsque la taille des éléments est prise égale à la taille du VER (gauche) et verrouillage lorsque la taille des éléments est prise égale à la distance entre les particules (droite)

Pour palier à ce dernier problème, un nouveau terme de couplage a été introduit dans [Prudhomme *et al.* 2012] et ses propriétés mathématiques ont été étudiées. Ce nouveau terme est construit à l'aide d'un opérateur intégral (ou de moyenne) défini sur un VER, dont la taille détermine dans un certain sens l'échelle à laquelle le modèle continu et le modèle particulaire peuvent échanger l'information. La nouvelle formulation du problème couplé qui en découle est consistante et virtuellement indépendante du maillage ; elle est présentée ci-après dans le cas 1D avec des potentiels harmoniques (raideurs k_i , distance inter-particulaire ℓ_i).

A la vue du processus d'homogénéisation, un choix naturel est de faire correspondre : (i) les moyennes, sur un VER, des déformations $(R_o u)'$ et $(\Pi_o \mathbf{z})'$ issues des deux modèles concurrents ; (ii) les moyennes des déplacements $R_o u$ et $\Pi_o \mathbf{z}$ sur Ω_o afin de contraindre les mouvements de solide rigide. La définition de ces moyennes est directe sauf sur les frontières de la zone de couplage Ω_o ; on propose donc de définir les opérateurs de moyenne suivants, avec ξ la taille du VER et $\Omega_o = [x_a, x_b]$ (voir Figure 3.6), pour $v \in H^1(\Omega_o)$:

$$v^*(x) = \begin{cases} \frac{1}{\xi} \int_{x_a}^{x_a+\xi} v' dy = \frac{v(x_a + \xi) - v(x_a)}{\xi} & \forall x \in [x_a, x_a + \xi/2] \\ \frac{1}{\xi} \int_{x-\xi/2}^{x+\xi/2} v' dy = \frac{v(x + \xi/2) - v(x - \xi/2)}{\xi} & \forall x \in [x_a + \xi/2, x_b - \xi/2] \\ \frac{1}{\xi} \int_{x_b-\xi}^{x_b} v' dy = \frac{v(x_b) - v(x_b - \xi)}{\xi} & \forall x \in [x_b - \xi/2, x_b] \end{cases} \quad (3.17)$$

On introduit aussi la moyenne sur Ω_o d'une fonction $v \in H^1(\Omega_o)$: $\bar{v} = \frac{1}{|\Omega_o|} \int_{\Omega_o} v dx$.

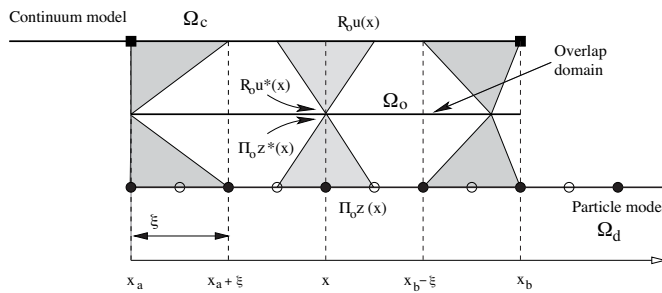


Figure 3.6: Domaine de définition des opérateurs de moyenne

On définit alors la mesure suivante de l'écart sur Ω_o entre les solutions des deux modèles :

$$\mathcal{G}(R_o u - \Pi_o \mathbf{z}) = \beta_0 |\overline{(R_o u - \Pi_o \mathbf{z})}|^2 + \beta_1 \int_{\Omega_o} |(R_o u - \Pi_o \mathbf{z})^*|^2 dx \quad (3.18)$$

où (β_0, β_1) sont des paramètres de dimension non-négatifs. $\mathcal{G}$ définit une semi-norme sur $H^1(\Omega_o)$ car elle n'est pas forcément définie, i.e. le sous-espace :

$$M_0 = \{\mu_0 \in H^1(\Omega_o) : \bar{\mu}_0 = 0 \text{ et } \mu_0^*(x) = 0, \forall x \in \Omega_o\} \quad (3.19)$$

n'est pas vide. En se restreignant au cas où Ω_o couvre exactement un VER, on montre que les fonctions $\mu_0 \in M_0$ peuvent être représentées par des combinaisons linéaires de la forme :

$$\mu_0(x) = \sum_{k=1}^{\infty} b_k \left[\sin \left(2k\pi \frac{x - x_a}{\xi} \right) \right] + c_k \left[\sin \left((2k-1)\pi \frac{x - x_a}{\xi} \right) - \frac{2/\pi}{(2k-1)} \right] \quad (3.20)$$

et que toute fonction dans $H^1(\Omega_o)$ peut se décomposer sous la forme :

$$\mu(x) = a_0 + a_1 x + \mu_0(x) \quad ; \quad \{a_0, a_1\} \in \mathbb{R}^2 \quad (3.21)$$

L'espace vectoriel des multiplicateurs de Lagrange est donc choisi comme $M = H^1(\Omega_o)/M_0$ avec pour espace vectoriel et norme associés :

$$(\lambda, \mu)_M = \beta_0 \bar{\lambda} \bar{\mu} + \beta_1 \int_{\Omega_o} \lambda^* \mu^* dx \quad ; \quad \|\mu\|_M = \sqrt{(\mu, \mu)_M} \quad (3.22)$$

Les normes sur les espaces vectoriels $\mathcal{U}$ et $\mathcal{W}$ associés aux modèles continu et discret, respectivement, sont quant à elles choisies comme :

$$\|v\|_{\mathcal{U}} = \sqrt{\int_{\Omega_c} E |v'|^2 dx} \quad ; \quad \|\mathbf{w}\|_{\mathcal{W}} = \sqrt{|\mathbf{w}|_{\mathcal{W}}^2 + \delta |\bar{\mathbf{w}}|^2} \quad (3.23)$$

avec E le module de Young du modèle continu, $|\mathbf{w}|_{\mathcal{W}}^2 = \sum_{i=1}^m k_i (w_i - w_{i-1})^2$, $\bar{\mathbf{w}} = \overline{\Pi_o \mathbf{w}}$ et δ une constante consistante dimensionnellement.

On définit enfin la forme bilinéaire $b(\bullet, \bullet)$ sur $M \times X$ telle que :

$$b(\mu, V) = (\mu, R_o v - \Pi_o \mathbf{w})_M \quad (3.24)$$

où, par simplicité de notation, on a introduit l'espace produit $X = \mathcal{U} \times \mathcal{W}$ avec la norme associée $\|V\|_X = \sqrt{\|v\|_{\mathcal{U}}^2 + \|\mathbf{w}\|_{\mathcal{W}}^2}$. L'espace noyau $X_0 \subset X$ est tel que :

$$X_0 = \{V \in X; b(\mu, V) = 0 \quad \forall \mu \in M\} \quad (3.25)$$

Le problème couplé Arlequin 1D étudié ici consiste donc à chercher $U = (u, \mathbf{z}) \in X$ minimisant l'énergie potentielle totale sous la contrainte $\|R_o u - \Pi_o \mathbf{z}\|_M = 0$. Ce problème peut à nouveau être écrit comme un problème de point-selle :

$$\text{Trouver } U \in X, \lambda \in M \text{ tel que } L(U, \lambda) = \inf_{V \in X} \sup_{\mu \in M} L(V, \mu) \quad (3.26)$$

où le lagrangien s'écrit sous la forme $L(V, \mu) = \frac{1}{2} a(V, V) - l(V) - b(\mu, V)$ avec :

$$a(U, V) = \int_{\Omega_c} \alpha_c E u' v' dx + \sum_{i=1}^m \alpha_d k_i (z_i - z_{i-1}) (w_i - w_{i-1}) \quad ; \quad l(V) = \sum_{i=1}^m f_i w_i \quad (3.27)$$

Le problème couplé peut donc être mis sous la forme :

Trouver $U \in X, \lambda \in M$ tel que : $\begin{aligned} a(U, V) - b(\lambda, V) &= l(V) & \forall V \in X \\ b(\mu, U) &= 0 & \forall \mu \in M \end{aligned}$	(3.28)
--	--------

Il est montré dans [Prudhomme *et al.* 2012] que (3.28) est bien posé pour $\beta_0 > 0$ et $\beta_1 > 0$ et sous certaines conditions pour α_c . Les continuités des formes a et l ayant été analysées dans [Bauman *et al.* 2008], la démonstration porte sur la continuité du terme de couplage b , sa vérification de la condition inf-sup de Babuška-Brezzi [Babuška 1973], ainsi que la coercivité de a . On rappelle ici les résultats majeurs.

- Continuité de b : pour tout $\mu \in M$, $V = (v, \mathbf{w}) \in X$, on a :

$$|b(\mu, V)| \leq M_b \|\mu\|_M \|V\|_X \quad ; \quad M_b = \sqrt{\beta_0 \left(\frac{|\Omega_c|^2}{2E|\Omega_o|} + \frac{1}{\delta} \right) + \beta_1 \left(\frac{1}{E} + \frac{1}{\min_i k_i \ell_i} \right)} \quad (3.29)$$

- Condition inf-sup pour b : pour $\beta_1 > 0$, on a :

$$\inf_{\mu \in M} \sup_{V \in X} \frac{|b(\mu, V)|}{\|\mu\|_M \|V\|_X} \geq \gamma_b \quad ; \quad \gamma_b = \min \left(\sqrt{\frac{\beta_0}{2\delta}}, \sqrt{\frac{2\beta_1}{2E + \delta|\Omega_o|}} \right) \quad (3.30)$$

- Coercivité de a : pour $\alpha_c = \alpha_d = 1/2$ sur Ω_o , on a :

$$\begin{cases} \inf_{U \in X_0} \sup_{V \in X_0} \frac{|a(U, V)|}{\|U\|_X \|V\|_X} > \gamma_a \\ \sup_{U \in X_0} a(U, V) > 0 \quad \forall V \in X_0, V \neq 0 \end{cases} \quad ; \quad \gamma_a = \frac{1}{2} \min_i \left(\frac{E}{k_i \ell_i} \right) \min_i \left(\frac{k_i \ell_i}{E} \right) \min \left(\frac{1}{2}, \frac{E}{\delta |\Omega_c|^2} \right) \quad (3.31)$$

On présente maintenant quelques résultats numériques obtenus dans [Prudhomme *et al.* 2012] pour un couplage particulaire/continu 1D avec une distribution hétérogène de potentiels harmoniques et une distance inter-particulaire $\ell = 0,5$ constante. Le domaine d'étude est $\Omega = (0; 3)$; le modèle continu est utilisé dans le sous-domaine $\Omega_c = (0; 2)$ tandis que le modèle particulaire est utilisé dans le sous-domaine $\Omega_d = (1; 3)$. La force f appliquée en $x = 3$ est choisie de telle façon que le déplacement en ce même point, lorsqu'on utilise le modèle continu dans tout Ω , est égal à 1.

La distribution des potentiels harmoniques est choisie périodique avec deux paramètres de raideur : $k_1 = 100$ et $k_2 = 1$; la structure particulaire est alors construite avec $k_{2j-1} = k_1$ et $k_{2j} = k_2$, et le module de Young homogénéisé correspondant est donné explicitement par $E = 2\ell \frac{k_1 k_2}{k_1 + k_2}$. La taille des éléments utilisés pour discrétiser M ainsi que celle du domaine de couplage Ω_o sont choisies égales à celle du VER, i.e. 2ℓ . On étudie l'effet de la taille du maillage associée à $\mathcal{U}_h$ sur la solution Arlequin. Les résultats obtenus pour quatre différentes tailles d'élément ($h = 2\ell$, $h = \ell/2$, $h = \ell/4$ et $h = \ell/32$) sont montrées sur la Figure 3.7. Comme attendu, la solution couplée se montre indépendante de la taille du maillage ; le déplacement de la particule en $x = 3$ est égal à 1,0817 pour les quatre cas.

1.3 Couplage dans le cadre stochastique

La plupart des méthodes abordant le couplage particulaire/continu se placent dans un cadre déterministe. Cependant, les modèles particuliers introduisent

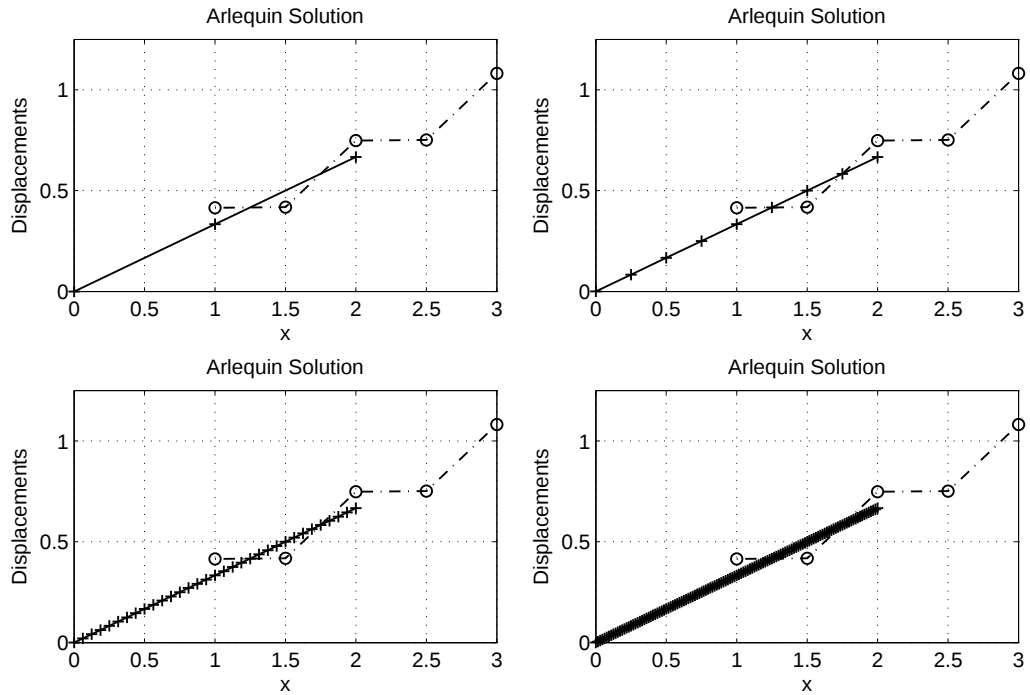


Figure 3.7: Solution Arlequin pour différentes discrétisations du modèle continu

de manière intrinsèque des paramètres incertains (présence de défauts, structure du réseau,...) dont l'influence doit être étudiée. Dans [Chamoïn *et al.* 2008], le couplage multiéchelle basé sur la méthode Arlequin et détaillé dans la Section 1.1 a été étendu aux analyses probabilistes, permettant de mener des simulations de Monte-Carlo raisonnablement coûteuses pour des systèmes particuliers soumis à des phénomènes aléatoires. En particulier, la construction du modèle continu stochastique équivalent et son couplage avec le modèle particulière ont été abordés à l'aide des outils (décomposition de Karhunen-Loeve, projection sur le Chaos Polynomial) classiquement utilisés dans la méthode des éléments finis stochastiques [Ghanem et Spanos 1991].

La méthode est décrite sur une chaîne 1D de particules connectées avec des potentiels d'interaction harmoniques (Figure 3.8). La chaîne est encadrée à ses deux extrémités $x = x_0$ et $x = x_n$, et on suppose que la région d'intérêt est au centre de la chaîne. La région continue Ω_c (divisée en Ω_c^1 et Ω_c^2), la région particulière Ω_d et la région de couplage Ω_o (divisée en Ω_o^1 et Ω_o^2) sont conformes avec le réseau de particules et sont définies par :

$$\Omega_c = (x_0, x_b) \cup (x_c, x_n) \quad ; \quad \Omega_d = (x_a, x_d) \quad ; \quad \Omega_o = \Omega_d \cap \Omega_c \quad (3.32)$$

En supposant que le chargement est uniquement appliqué dans Ω_d , le couplage Arlequin (3.15) (ou (3.28)) dans le cadre déterministe mène après discrétisation à la

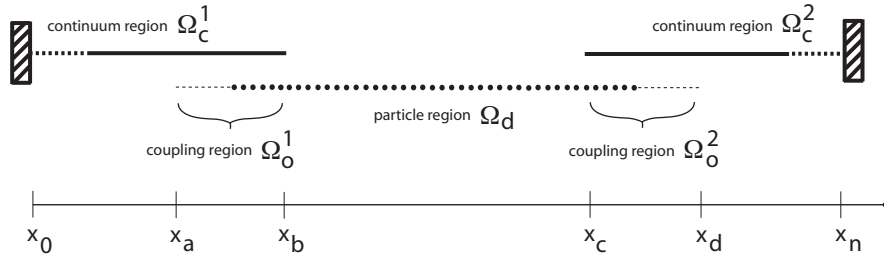


Figure 3.8: Les différentes régions du problème de couplage 1D

résolution d'un système de la forme :

$$\begin{pmatrix} \mathbb{K}_1 & 0 & 0 & \mathbb{A}_1^T & 0 \\ 0 & \mathbb{K}_Z & 0 & -\mathbb{A}_Z^T & \mathbb{D}_Z^T \\ 0 & 0 & \mathbb{K}_2 & 0 & -\mathbb{D}_2^T \\ \mathbb{A}_1 & -\mathbb{A}_Z & 0 & 0 & 0 \\ 0 & \mathbb{D}_Z & -\mathbb{D}_2 & 0 & 0 \end{pmatrix} \begin{pmatrix} \mathbf{U}_1 \\ \mathbf{Z} \\ \mathbf{U}_2 \\ \boldsymbol{\Lambda}_1 \\ \boldsymbol{\Lambda}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{0} \\ \mathbf{F} \\ \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \end{pmatrix} \quad (3.33)$$

où $(\mathbf{U}_1, \mathbf{U}_2)$ (resp. $\mathbf{Z}$ et $(\boldsymbol{\Lambda}_1, \boldsymbol{\Lambda}_2)$) regroupe les inconnues liées à la solution dans le milieu continu (resp. à la solution dans le milieu particulaire et aux multiplicateurs de Lagrange).

Dans le cas de la simulation de structures polymères par exemple, les paramètres de raideur k_i des potentiels harmoniques sont stochastiques et décrits par des variables aléatoires. Cet aspect stochastique peut venir soit de l'échantillon (dans ce cas les k_i appartiennent à un ensemble fini de valeurs connues) soit des paramètres eux mêmes (dans ce cas les k_i peuvent fluctuer autour d'une valeur moyenne). On introduit alors l'espace stochastique $(\Theta, \mathcal{F}, P)$, au sens de Kolmogorov, équipé du produit scalaire $\langle \bullet, \bullet \rangle$.

La construction du modèle continu stochastique est faite par homogénéisation à partir de n_{rea} réalisations sur un VER. Evidemment, les caractéristiques stochastiques obtenues pour le modèle continu dépendent de la taille du VER choisi ; lorsque la taille du VER augmente, la moyenne des paramètres matériau (module de Young E ici) se stabilise et la variance diminue. A partir des valeurs n_{rea} valeurs E_n obtenues, une densité de probabilité continue est ensuite construite, sous la forme $E(\theta) = \overline{E} + \sigma_E \xi(\theta)$ avec $\overline{E}$ (resp. σ_E) la moyenne (resp. l'écart-type) de $E(\theta)$ et $\xi(\theta)$ une variable aléatoire centrée réduite. Cette construction est menée en pratique à l'aide de la théorie des noyaux de lissage qui fournit une densité de probabilité bornée $\hat{p}(E)$ de la forme :

$$\hat{p}(E) = \frac{1}{h n_{rea}} \sum_{n=1}^{n_{rea}} \kappa \left(\frac{E - E_n}{h} \right) \quad (3.34)$$

avec κ la fonction noyau utilisée et h un paramètre (largeur de fenêtre) à fixer. On choisit ici le noyau d'Epanechnikov [Epanechnikov 1969] construit à partir d'un polynôme de degré 2 à moyenne unitaire et variance nulle (Figure 3.9).

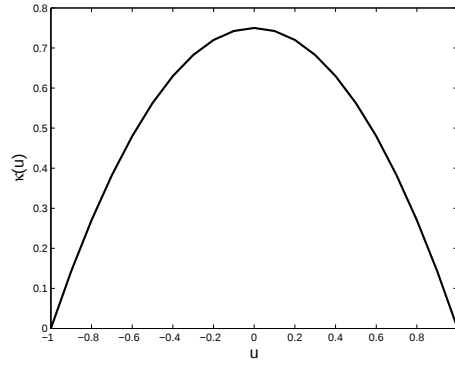
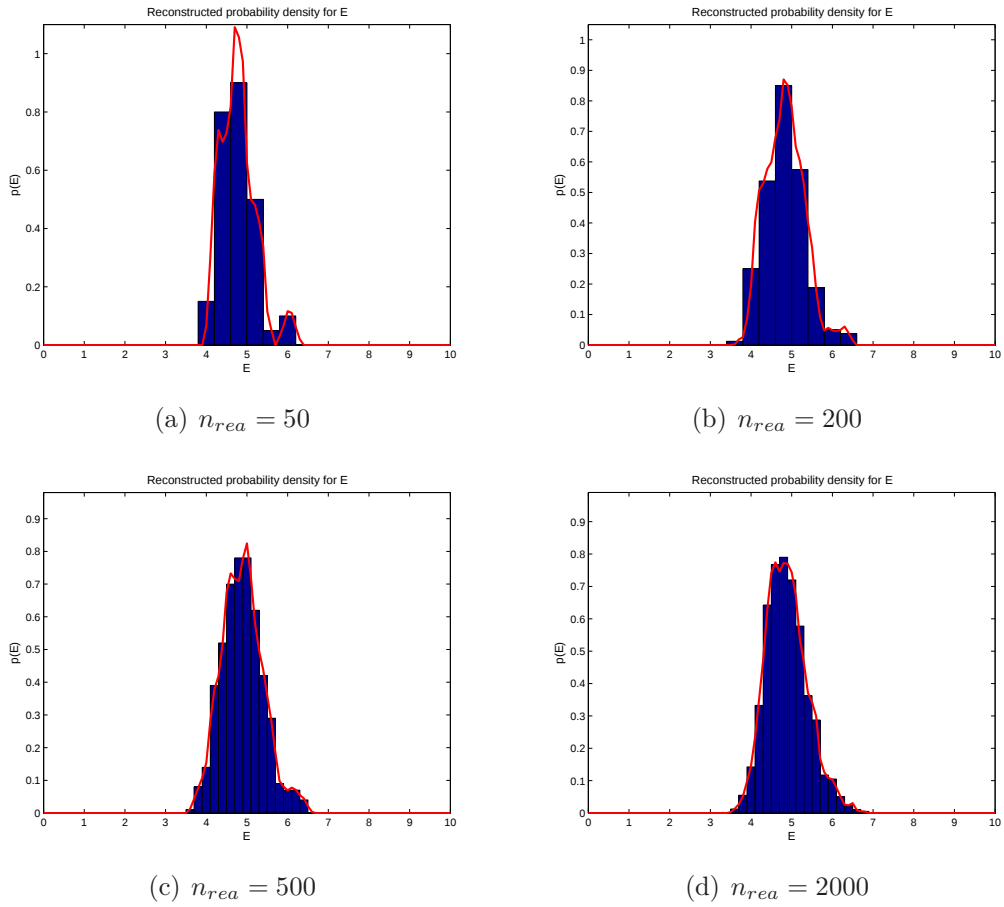


Figure 3.9: Représentation du noyau d'Epanechnikov

Une illustration de la construction de la densité de probabilité $\hat{p}(E)$ avec le noyau d'Epanechnikov est donnée sur la Figure 3.10, à partir de divers nombres n_{rea} de réalisations pour E .

Figure 3.10: Densités de probabilité $\hat{p}(E)$ obtenues par le noyau de lissage d'Epanechnikov pour différentes valeurs de n_{rea}

On définit à présent la méthode de couplage stochastique. Comme on doit comparer deux modèles compatibles dans un sens probabiliste, la définition de $E(\theta)$ (moyenne $\overline{E}$, écart-type σ_E) dans la région de couplage doit être liée à la taille des éléments de discrétisation des multiplicateurs de Lagrange i.e. à la taille du VER sur lequel les deux modèles concurrents sont mis en relation. Par conséquent, on considère la définition stochastique suivante pour la région continue Ω_c^i ($i = 1, 2$) (Figure 3.11) :

- dans la région exclusivement continue (x_0, x_a) (ou (x_d, x_n)), $E(\theta)$ est exprimé au moyen d'une variable aléatoire unique $\xi_1(\theta)$:

$$E_1(\theta) = \overline{E}_1 + \sigma_1 \xi_1(\theta) \quad (3.35)$$

où $\overline{E}_1$ et σ_1 sont calculés en fonction de la taille L_1 de cette région ;

- dans la région de couplage Ω_o^1 (ou Ω_o^2), $E(\theta)$ est donné dans chaque élément i ($i \in \{2, \dots, N_e + 1\}$) au moyen d'une variable aléatoire correspondante $\xi_i(\theta)$:

$$E_i(\theta) = \overline{E}_i + \sigma_i \xi_i(\theta) \quad (3.36)$$

où $\overline{E}_i$ et σ_i sont calculés en fonction de la taille L_i de l'élément i .

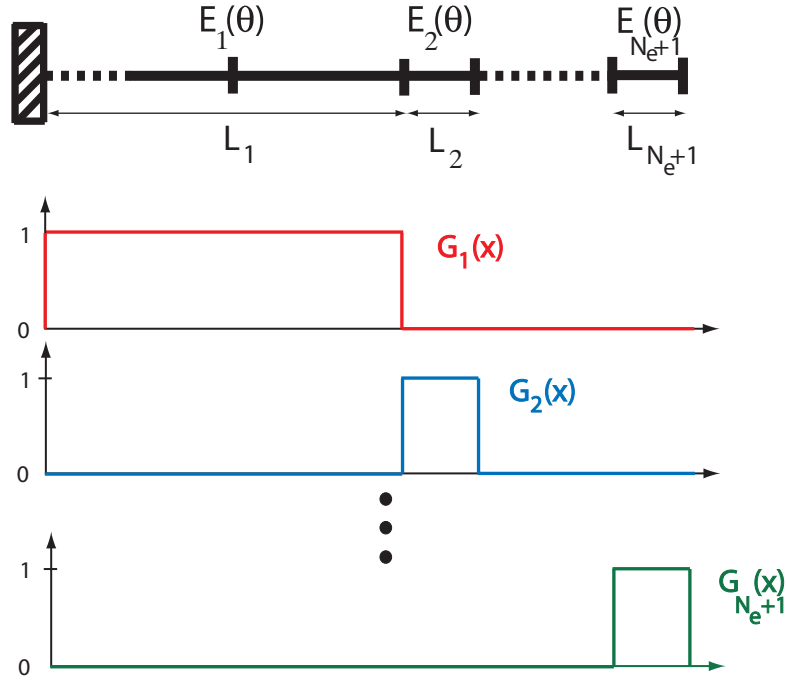


Figure 3.11: Définition stochastique de $E(x, \theta)$ pour le modèle continu

De ce fait, dans la région continue, $E(x, \theta)$ est exprimé par une expansion de type Karhunen-Loeve :

$$E(x, \theta) = \overline{E}(x) + \sum_{i=1}^{N_e+1} \sigma_i \xi_i(\theta) G_i(x) \quad (3.37)$$

où, pour i donné ($i = 1, \dots, N_e + 1$), $G_i(x)$ est la fonction indicatrice de la zone correspondante (Figure 3.11).

Afin de diminuer le nombre de variables aléatoires du problème, et donc le coût de calcul, on réduit à un le nombre de VER (ou éléments) dans chaque région de couplage Ω_o^i , $i = 1, 2$. De ce fait, quatre variables aléatoires définissent l'aspect stochastique de la région continue (voir Figure 3.12) et le champ stochastique de module de Young est donné par :

$$\begin{aligned} E(\theta) &= \bar{E}_1 + \sigma_1 \xi_1(\theta) \quad \text{dans } \Omega_c^1/\Omega_o^1 & ; & \quad E(\theta) = \bar{E}_2 + \sigma_2 \xi_2(\theta) \quad \text{dans } \Omega_o^1 \\ E(\theta) &= \bar{E}_3 + \sigma_3 \xi_3(\theta) \quad \text{dans } \Omega_c^2/\Omega_o^2 & ; & \quad E(\theta) = \bar{E}_4 + \sigma_4 \xi_4(\theta) \quad \text{dans } \Omega_o^2 \end{aligned} \quad (3.38)$$



Figure 3.12: Visualisation des quatre variables aléatoires associées au modèle continu

Le problème Arlequin déterministe (3.28) est alors reformulé sous une forme stochastique qui consiste à trouver $U(\theta) \in X_P$, $\lambda(\theta) \in M_P$ tels que :

$$\int_{\Theta} \left[a(U, V) - l(V) - b(\lambda, V) - b(\mu, U) \right] dP = 0 \quad \forall V \in X_P, \forall \mu \in M_P \quad (3.39)$$

où X_P (resp. M_P) est l'extension au cadre stochastique de l'espace X (resp. M).

Pour un échantillon donné des paramètres de raideur k_i dans la zone particulière Ω_d , la solution du système couplé peut s'exprimer à partir des n_v variables aléatoires $\xi_1, \xi_2, \dots, \xi_{n_v}$ de la région continue ($n_v = 4$ dans l'exemple traité) en utilisant une approche spectrale. Pour cela, on projette chaque variable $q(x, \theta)$ du problème (ici q est u , $\mathbf{z}$, ou λ) sur la base du Chaos Polynomial :

$$q(x, \theta) = \sum_{k=0}^{\infty} q^{(k)}(x) \Gamma^{(k)}(\{\xi_i\}) \quad (3.40)$$

où $\{\Gamma^{(k)}\}$ est une base orthogonale ($\langle \Gamma^{(i)}, \Gamma^{(j)} \rangle = 0 \quad \forall i \neq j$) et normée discrétisant la dimension stochastique Θ . Dans la décomposition (3.40), tronquée en pratique à l'ordre N , $q^{(0)}$ correspond à la moyenne de q et $q^{(1)}, q^{(2)}, \dots, q^{(N)}$ peuvent être vus comme des modes additionnels. Du fait que les variables aléatoires $\xi_i(\theta)$ ont leur propre région spatiale associée, ces variables sont découplées en espace et les termes croisés $\xi_i^n \xi_j^m$ ($i \neq j$ et $n, m \geq 1$) sont nuls dans l'expansion du Chaos Polynomial. On peut donc réécrire (3.40) sous la forme :

$$\begin{aligned} q &= \hat{q}^{(0)} \\ &+ \hat{q}^{(1)} \xi_1 + \hat{q}^{(2)} \xi_2 + \dots + \hat{q}^{(n_v)} \xi_{n_v} \\ &+ \hat{q}^{(n_v+1)} \left(\xi_1^2 - \frac{m_{31}}{m_{21}} \xi_1 - m_{21} \right) + \hat{q}^{(n_v+2)} \left(\xi_2^2 - \frac{m_{32}}{m_{22}} \xi_2 - m_{22} \right) + \dots \\ &+ \dots \end{aligned} \quad (3.41)$$

où $m_{k_i} = m_k(\xi_i)$ est le moment d'ordre k de la variable ξ_i . De ce fait, en notant d_i le degré maximal du polynôme de la variable aléatoire ξ_i , on a $N = \sum_{i=1}^{n_v} d_i$.

Le cadre classique de la méthode des éléments finis stochastiques [Ghanem et Spanos 1991], basée sur une approche de type Galerkin dans les espaces physique et stochastique, est ensuite utilisé pour résoudre (3.39). Dû à l'orthogonalité des polynômes $\Gamma^{(k)}$ qui constituent la base du Chaos Polynomial, $u^{(0)}$ et $\mathbf{z}^{(0)}$ sont couplés ensemble par $\lambda^{(0)}$, $u^{(1)}$ et $\mathbf{z}^{(1)}$ sont couplés ensemble par $\lambda^{(1)}$, etc. On obtient ainsi un système similaire à celui donné dans (3.33), de la forme :

$$\begin{pmatrix} \mathbb{K}_1^\diamond & 0 & 0 & \mathbb{A}_1^{\diamond T} & 0 \\ 0 & \mathbb{K}_Z^\diamond & 0 & -\mathbb{A}_Z^{\diamond T} & \mathbb{D}_Z^{\diamond T} \\ 0 & 0 & \mathbb{K}_2^\diamond & 0 & -\mathbb{D}_2^{\diamond T} \\ \mathbb{A}_1^\diamond & -\mathbb{A}_Z^\diamond & 0 & 0 & 0 \\ 0 & \mathbb{D}_Z^\diamond & -\mathbb{D}_2^\diamond & 0 & 0 \end{pmatrix} \begin{pmatrix} \mathbf{U}_1^\diamond \\ \mathbf{Z}^\diamond \\ \mathbf{U}_2^\diamond \\ \mathbf{\Lambda}_1^\diamond \\ \mathbf{\Lambda}_2^\diamond \end{pmatrix} = \begin{pmatrix} \mathbf{0} \\ \mathbf{F}^\diamond \\ \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \end{pmatrix} \quad (3.42)$$

mais à présent $(\mathbf{U}_1^\diamond, \mathbf{Z}^\diamond, \mathbf{U}_2^\diamond, \mathbf{\Lambda}_1^\diamond, \mathbf{\Lambda}_2^\diamond, \mathbf{F}^\diamond)$ et $(\mathbb{K}_1^\diamond, \mathbb{K}_Z^\diamond, \mathbb{K}_2^\diamond, \mathbb{A}_1^\diamond, \mathbb{A}_Z^\diamond, \mathbb{D}_Z^\diamond, \mathbb{D}_2^\diamond)$ sont des vecteurs blocs (taille N) et matrices blocs (taille $N \times N$) s'écrivant respectivement :

$$\mathbf{B}^\diamond = [\mathbf{B}^{(0)}, \mathbf{B}^{(1)}, \dots, \mathbf{B}^{(N)}]^T \quad ; \quad \mathbf{F}^\diamond = [\mathbf{F}, \mathbf{0}, \dots, \mathbf{0}]^T \quad ; \quad \mathbb{B}_{ij}^\diamond = \langle \Gamma^{(i)}, \mathbb{B} \Gamma^{(j)} \rangle \quad (3.43)$$

On obtient ainsi une solution du problème couplé qui dépend explicitement de l'ensemble $\{\xi_i\}$ décrivant les aspects aléatoires de la région continue Ω_c , ainsi que des paramètres k_i dans la région particulière Ω_d . Il est ensuite possible de dériver la densité de probabilité approchée d'une quantité d'intérêt donnée Q définie dans la région où le modèle particulière est conservé ; l'idée est d'utiliser une simulation de Monte-Carlo dans cette région uniquement, tandis que la représentation stochastique utilisée au niveau du modèle continu reste inchangée et fournit des conditions limites au modèle particulière. Pour chaque tirage i dans la simulation de Monte-Carlo, la méthode fournit une distribution de Q sous la forme :

$$Q_i(\theta) = \sum_{k=0}^N Q_i^{(k)} \Gamma^{(k)}(\xi_1, \xi_2, \dots, \xi_{n_v}) \quad (3.44)$$

qui prend en compte l'aléatoire induit par la région continue. A partir des n_t tirages et des n_t distributions $Q_i(\theta)$ associées, on peut obtenir une distribution finale $Q(\theta)$ de la quantité d'intérêt qui prend en compte à la fois l'aléa de la région particulière et celui de la région continue. Des valeurs approchées de la moyenne et de la variance de $Q(\theta)$ sont données respectivement par :

$$\bar{Q} \approx \frac{1}{n_t} \sum_{i=1}^{n_t} Q_i^{(0)} \quad ; \quad \sigma_Q^2 \approx \frac{1}{n_t} \sum_{i=1}^{n_t} \left[(Q_i^{(0)} - \bar{Q})^2 + \sum_{k=1}^N Q_i^{(k)^2} m_2(\Gamma^{(k)}) \right] \quad (3.45)$$

Ces valeurs approchées convergent vers les valeurs exactes quand le nombre de tirages augmente.

Pour exemple, on considère la structure 1D de la Figure 3.13, composée de 531 particules et encastree à ses deux extrémités. Dans la configuration initiale, les particules sont séparées par des liaisons de longueur $\ell = 0,1$. On suppose une distribution uniforme de quatre types de liaisons, notées a, b, c, d , dans Ω avec des constantes de raideur associées $\{k_a, k_b, k_c, k_d\} = \{25, 50, 75, 100\}$. La structure est affaiblie par un défaut local, modélisé par une modification de raideur de la forme $k_i^* = \frac{1}{1+g(x_{i/2})}k_i$ avec $g(x) = 10e^{-7x^2}$, centré autour de la particule $P = 316$ (Figure 3.13). Le chargement consiste en une force ponctuelle F appliquée en P .

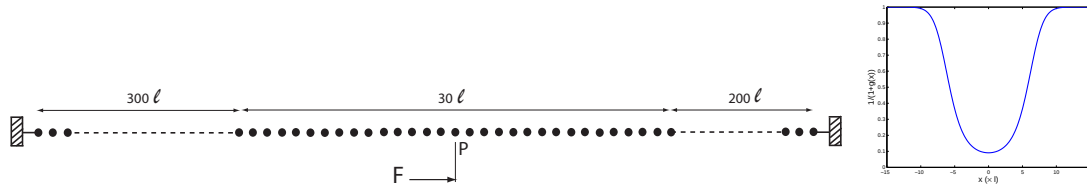


Figure 3.13: Structure 1D aléatoire (gauche) et fonction poids utilisée pour modéliser le défaut (droite)

Pour la modélisation multiéchelle, on conserve le modèle particulaire dans une zone contenant 31 particules autour de la particule P et on utilise un modèle continu dans le reste de la structure. Les régions de couplage Ω_o^1 et Ω_o^2 s'étendent sur une longueur de 20ℓ chacune, longueur qui est aussi prise pour le VER; par conséquent, le modèle continu stochastique est décrit par quatre variables aléatoires ξ_i ($i = 1, \dots, 4$).

La quantité d'intérêt considérée est le déplacement $z_P(\theta)$ de la particule P , et un résultat de référence peut être obtenu à partir d'une simulation de Monte-Carlo appliquée au modèle particulaire initial. La Figure 3.14 montre la moyenne du déplacement dans la structure couplée pour une réalisation du processus de Monte-Carlo. La ligne continue correspond au déplacement moyen dans la région continue ($u^{(0)}$) et les cercles correspondent au déplacement moyen des particules ($\mathbf{z}^{(0)}$). La projection sur le Chaos Polynomial est faite avec $N = 12$ i.e. le degré maximal des polynômes de chaque variable ξ_i est 3. À partir de l'expression stochastique du déplacement donnée dans (3.41), les harmoniques peuvent être séparés en quatre groupes, dépendant de la variable aléatoire à laquelle ils sont reliés (polynômes en ξ_1, ξ_2, ξ_3 , ou ξ_4). On trace sur la Figure 3.15, pour chaque ξ_i , les trois premiers harmoniques correspondant aux polynômes de degré 1, 2, et 3. On observe que les variations du module de Young autour de la moyenne dans la région continue ont une influence non-négligeable sur la valeur de la quantité d'intérêt.

On analyse à présent l'effet de l'ordre choisi pour le Chaos Polynomial. On note d_i le degré maximal des polynômes en ξ_i , et on prend $d_1 = d_3 = d_{13}$ et $d_2 = d_4 = d_{24}$ pour que les phénomènes aléatoires soient représentés avec la même précision dans les deux régions continues et dans les deux zones de couplage. Le Tableau 3.1 montre les valeurs approchées obtenues pour la moyenne $\bar{Q}$ et l'écart-type σ_Q de Q ,

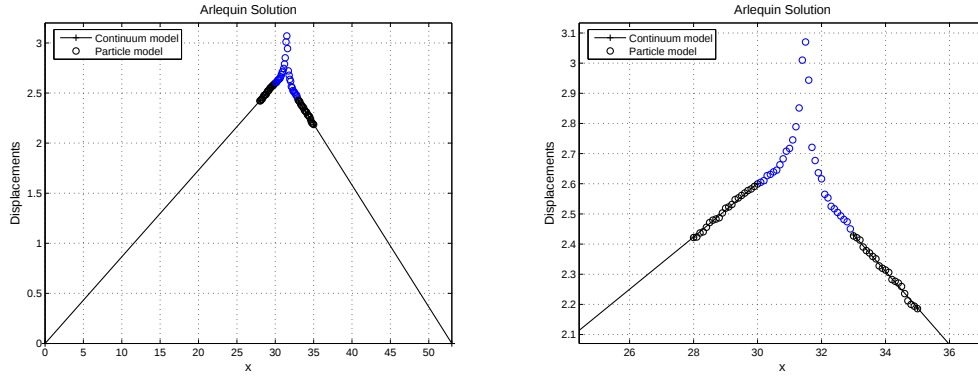


Figure 3.14: Valeur moyenne du déplacement obtenu avec la méthode de couplage : solution totale (gauche) et zoom local autour de la particule P (droite)

en fonction de d_{13} et d_{24} dans le cas de 100 tirages. On observe qu'il est inutile de

$d_{13} \backslash d_{24}$	0	1	2	3
0	3,0106/0,0880	3,0118/0,0886	3,0118/0,0886	3,0118/0,0886
1	3,0129/0,1097	3,0141/0,1102	3,0141/0,1102	3,0141/0,1102
2	3,0129/0,1097	3,0141/0,1102	3,0141/0,1102	3,0141/0,1102
3	3,0129/0,1097	3,0141/0,1102	3,0141/0,1102	3,0141/0,1102

Tableau 3.1: Evolution de $\overline{Q}$ (terme de gauche) et σ_Q (terme de droite) en fonction de d_{13} et d_{24}

prendre d_{13} et d_{24} supérieurs à 1 car la plupart de l'information stochastique de la région continue est concentrée dans les premiers harmoniques pour chaque variable aléatoire ξ_i ($i = 1, \dots, 4$).

Enfin, on montre sur la Figure 3.16 les résultats de convergence obtenus pour $\overline{Q}$ et σ_Q , en utilisant la solution de référence (simulation de Monte-Carlo sur le modèle particulaire complet) d'une part, et la méthode de couplage d'autre part.

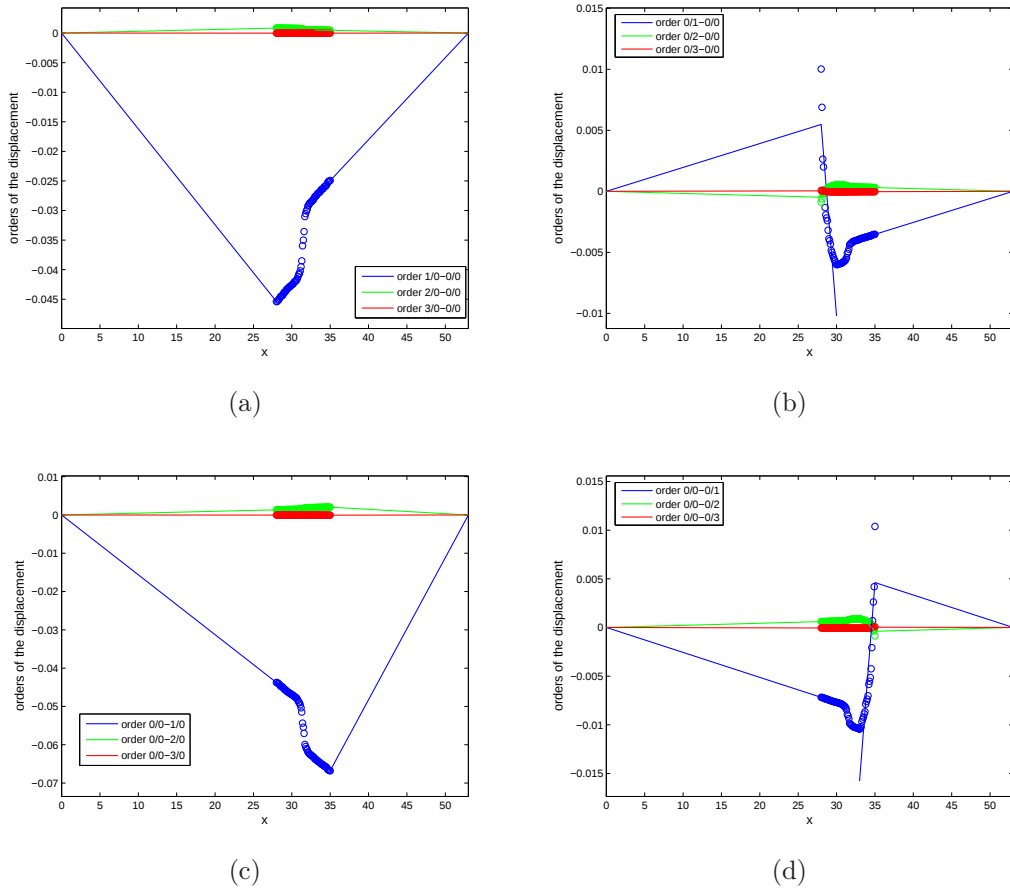


Figure 3.15: Influence de ξ_1 (a), ξ_2 (b), ξ_3 (c) et ξ_4 (d) sur les harmoniques du déplacement

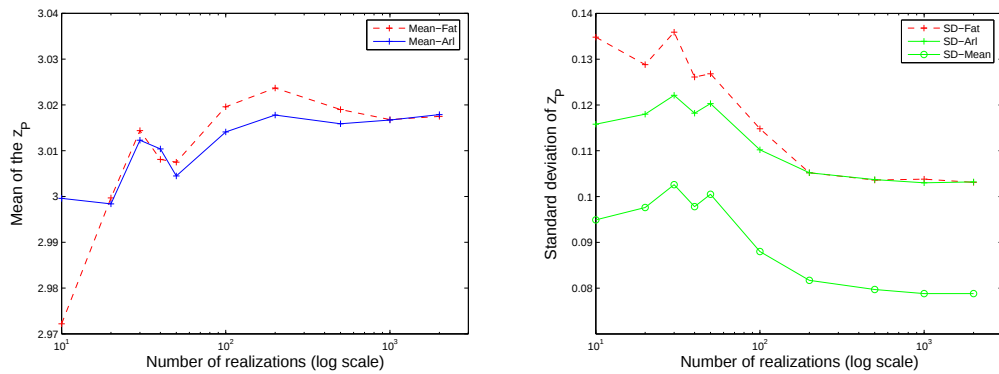


Figure 3.16: Convergence de $\overline{Q}$ et σ_Q

On observe que les deux méthodes donnent des résultats similaires asymptotiquement. Sur la seconde courbe est également représentée la convergence de σ_Q calculée seulement à partir des variations de $Q_i^{(0)}$ i.e. sans considérer les effets des harmo-

niques $Q_i^{(1)}$, $Q_i^{(2)}$, $Q_i^{(3)}$, etc. On voit clairement que l'écart-type réel est alors sous-estimé ce qui permet de conclure que les harmoniques ne doivent pas être négligés pour obtenir une estimation fiable sur la distribution de Q . Néanmoins, la moyenne $\overline{Q}$ peut être calculée avec une bonne approximation même avec $d_{13} = d_{24} = 0$, i.e. en couplant seulement les moyennes des différents champs du problème.

1.4 Prise en compte des incompatibilités entre modèles

1.4.1 Cas de couplages en statique

Les méthodes de couplage particulaire/continu peuvent exhiber des effets non-physiques au voisinage de la région de couplage qui dégradent la précision de la solution numérique. Ces effets, majoritairement dus aux forces fantômes qui apparaissent dans les couplages basés sur la minimisation d'une énergie globale, furent initialement analysés et traités dans le cadre de la méthode Quasi-Continuum [Shenoy *et al.* 1999, Knap et Ortiz 2001, Miller et Tadmor 2002, Curtin et Miller 2003]. Dans [Chamoïn *et al.* 2010], un travail similaire a été réalisé dans le cadre du couplage avec la méthode Arlequin; les effets non-physiques (autres que le possible verrouillage décrit dans la Section 1.2) et leur cause ont été identifiés, puis une technique corrective a été proposée. Cette dernière technique repose sur un post-traitement de la solution numérique avec introduction de forces fictives (ou mortes) qui peuvent être systématiquement évaluées et insérées de façon consistante dans la formulation Arlequin.

On analyse ici les effets non-physiques sur un modèle 1D avec potentiels harmoniques tel que celui utilisé dans la Section 1.2 (voir aussi Figure 3.4); le problème couplé Arlequin consiste donc à résoudre (3.28). Le domaine d'étude est $\Omega = (0; 8)$ et on prend $\Omega_c = (0; 4, 8)$, $\Omega_d = (3, 2; 8)$ et $\Omega_o = \Omega_a \cap \Omega_c = (3, 2; 4, 8)$. On note $\ell = 0, 2$ la distance inter-particulaire et $h = s\ell$ la taille des éléments finis utilisés pour discrétiser le modèle continu dans Ω_c . La chaîne de particules est fixée en $x = 0$ et soumise à une force de traction $F = E^{eq}/L$ en $x = 8$, avec E^{eq} le module de Young homogénéisé relatif au milieu continu. La fonction poids α_c dans la région de couplage Ω_o sera alternativement choisie comme constante, linéaire, ou cubique (voir Figure 3.4).

On considère tout d'abord le cas idéal où seules les interactions avec les plus proches voisins sont impliquées, avec des coefficients de raideur homogènes $k_{ij} = k = 100$, et $s = 1$; les résultats sont représentés sur la Figure 3.17. Comme attendu, le déplacement en L est égal à 1 pour les trois types de fonction poids et le champ de déformation est constant. On répète à présent l'expérience avec $s = 2$ et 8 dans le cas d'une fonction poids α^c cubique. Les résultats de la Figure 3.18 montrent que la solution se détériore au fur et à mesure que s augmente. Cela est dû au fait que le problème couplé devient localement mal posé, l'énergie des particules étant pondérée par un coefficient qui tend vers zéro proche du bord gauche de Ω_o . Cette perte de coercivité, analysée dans [Ben Dhia 2008] pour le couplage Arlequin entre deux modèles continus, fait alors apparaître des modes libres.

Un autre effet, occasionné par les forces fantômes, apparaît lorsque le modèle particulaire implique des interactions à longue distance; il est le résultat de la non-localité du domaine particulaire qui crée, au niveau de la région de couplage Ω_o ,

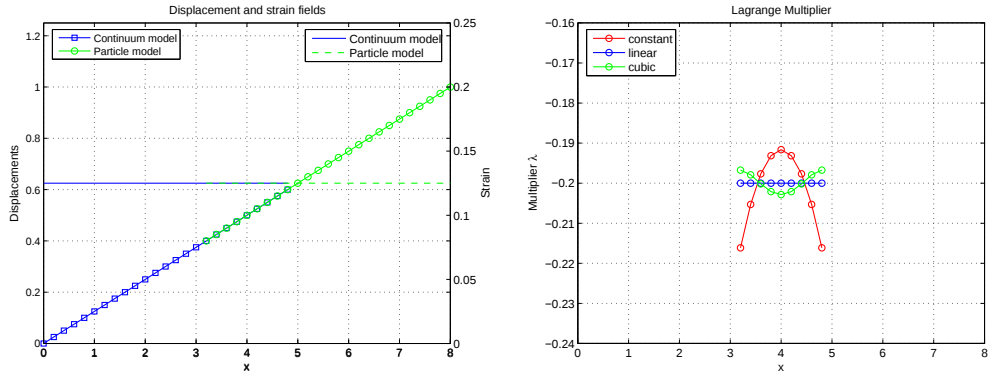


Figure 3.17: Champs de déplacement et de déformation (gauche) et multiplicateur de Lagrange en fonction de la fonction poids α_c (droite)

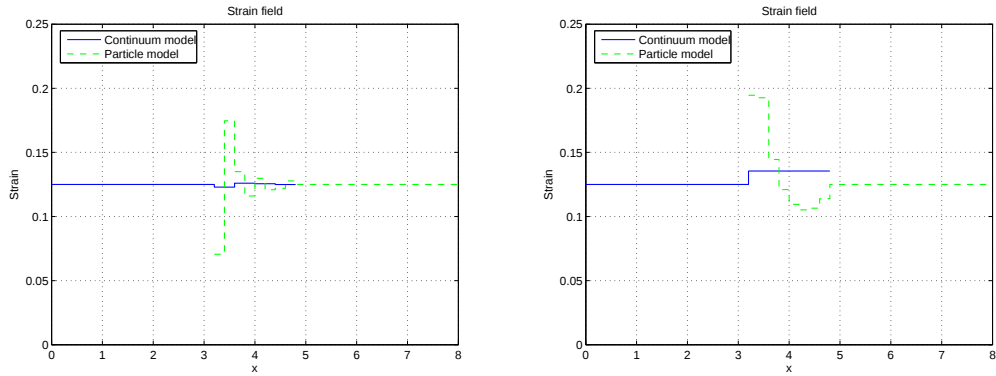


Figure 3.18: Champ de déformation obtenu avec une fonction poids cubique, dans les cas $s = 2$ (gauche) et $s = 8$ (droite)

une dissymétrie dans la définition de l'énergie totale du système couplé (l'énergie du modèle continu est locale). Par exemple, la Figure 3.19 montre les résultats obtenus, toujours avec $s = 1$, mais en considérant cette fois des interactions avec les premiers et seconds voisins (constantes de raideur associées égales à 100 et 50 respectivement). Au delà des effets observés vers le bord gauche de Ω_o et le bord droit de Ω , qui sont des perturbations de surface dues à la troncature du modèle particulaire, on observe sur le bord droit de Ω_o des effets attribués aux forces fantômes issues d'une distribution d'énergie non symétrique. On remarque également que choisir une fonction poids α_c linéaire ou cubique permet d'atténuer ces effets mais reste insuffisant.

Afin d'éviter les effets de surface, on introduit des particules supplémentaires à énergie nulle au voisinage de l'interface entre les modèles (voir Figure 3.20) à la manière des *pad atoms* introduits dans la méthode Quasi-Continuum. La position de ces particules est déterminée par duplication du champ de déformation dans le VER le plus proche. D'autre part, on définit des forces mortes pour chaque nœud du maillage EF dans Ω_c et chaque particule dans Ω_d . Ces forces sont construites comme la différence entre les forces agissant sur les nœuds/particules dans la configuration

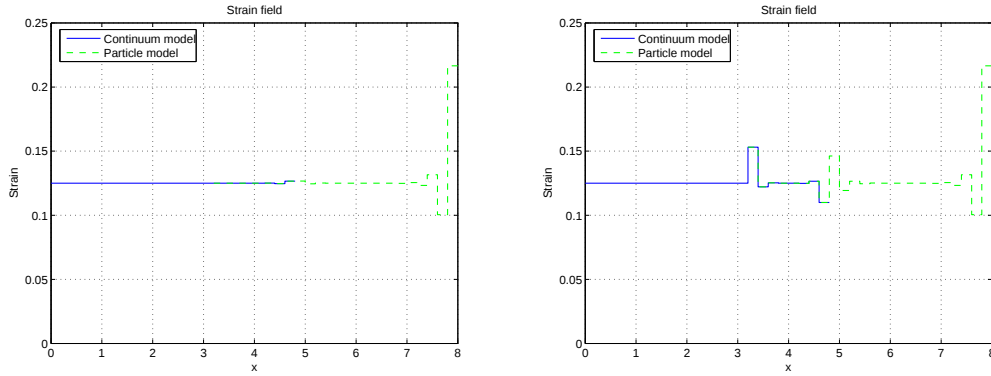


Figure 3.19: Champ de déformation obtenu avec $h = 1$, pour des fonctions poids linéaire (gauche) et constante (droite)

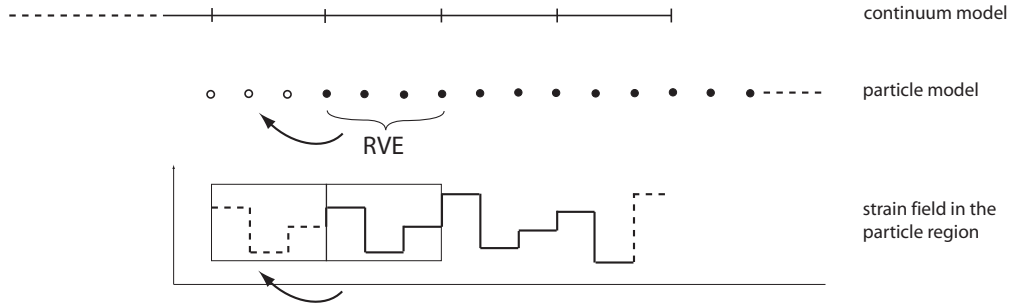


Figure 3.20: Duplication du champ de déformation du dernier VER pour obtenir la position des particules fantômes

actuelle du problème couplé (ces forces sont nulles à l'équilibre) et celles qui agiraient si les nœuds/particules étaient dans un modèle totalement continu/particulaire. Les forces mortes, notées respectivement g_k^c et g_i^d , s'écrivent pour le nœud k et la particule i :

$$g_k^c = - \sum_{j \in \{k-1, k+1\}} \frac{E^{eq}}{h} (u_h^j - u_h^k) \quad ; \quad g_i^d = - \sum_{j \in \mathcal{N}_i} k_{ij} (z_j - z_i) - f_i \quad (3.46)$$

En respectant la philosophie de l'approche Arlequin, on introduit les forces mortes (3.46) d'une manière consistante avec la PUM, i.e. les fonctions poids α^c et α^d . Le lagrangien associé au problème Arlequin est alors modifié et devient, pour $V = (v, \mathbf{w}) \in \mathcal{U} \times \mathcal{W}$:

$$\mathcal{L}(V, \mu) = \frac{1}{2} a(V, V) - l(V) + b(\mu, V) - \sum_k \alpha^c(\mathbf{x}_k) \mathbf{g}_k^c v_h^k - \sum_i \alpha^d(\mathbf{x}_i) \mathbf{g}_i^d w_i \quad (3.47)$$

La procédure de correction est appliquée de manière itérative pour prendre en compte les effets de relaxation ; le lagrangien à l'itération $n + 1$ implique les forces mortes calculées avec les déplacements des nœuds et des particules à l'itération n . Cette procédure permet de tendre vers la solution physique du problème couplé dès

lors que les modèles concurrents sont compatibles et que l'espace de discrétisation des multiplicateurs de Lagrange est correctement choisi.

Les performances de la méthode de correction par forces mortes sont illustrées ci-après, avec une fonction poids α_c constante dans Ω_o . La Figure 3.21 montre pour les itérations 1, 2 et 4 les forces mortes pondérées et le champ de déformation corrigé obtenu, sur le cas-test étudié précédemment pour mettre en évidence les modes libres (voir Figure 3.18). La Figure 3.22 montre quant à elle les champs de défor-

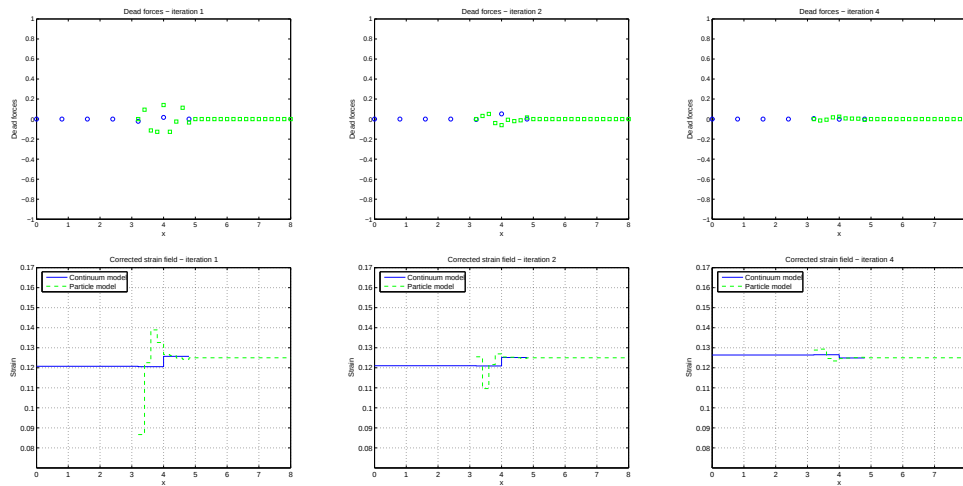


Figure 3.21: Forces mortes pondérées (haut) et champ de déformation (bas) obtenus aux itérations 1, 2 et 4 (de gauche à droite) de la correction des modes libres

mation corrigés obtenus après les trois premières itérations sur le cas-test étudié précédemment pour mettre en évidence les modes libres (voir Figure 3.19).

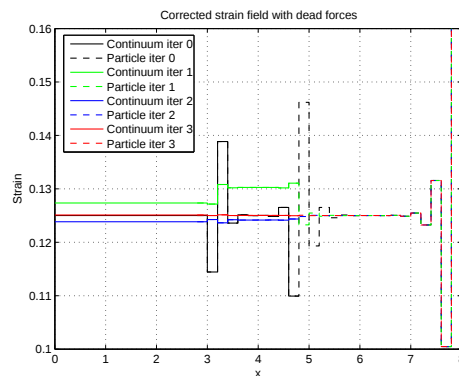


Figure 3.22: Champ de déformation obtenu aux itérations 0, 1, 2 et 3 de la correction des forces fantômes

On applique pour finir la méthode de correction dans le cas de coefficients de raideur k_i hétérogènes. La longueur du VER étant 2ℓ , on prend $s = 2$ pour avoir

un couplage consistant quand on discrétise le champ de multiplicateurs de Lagrange sur le maillage continu. On considère des interactions avec les plus proches voisins (raideurs égales à 100 et 70 distribuées périodiquement) et avec les seconds voisins (raideur égale à 50). La fonction poids α_c est choisie constante dans Ω_o . Les résultats obtenus et présentés sur la Figure 3.23 montrent que la méthode de correction est également efficace dans ce cas de structure hétérogène et permet de tendre vers la solution physique du problème couplé. Sur cette même Figure 3.23, on montre les résultats obtenus dans le cas où le couplage est surcontraint de telle sorte que la méthode de correction n'arrive pas à converger vers la solution physique. En prenant $s = 1$, la discrétisation du champ de multiplicateurs de Lagrange force le verrouillage dans Ω_o entre les déplacements des modèles continu et particulaire. En d'autres termes, une telle contrainte n'est pas compatible avec la nature hétérogène du modèle moléculaire.

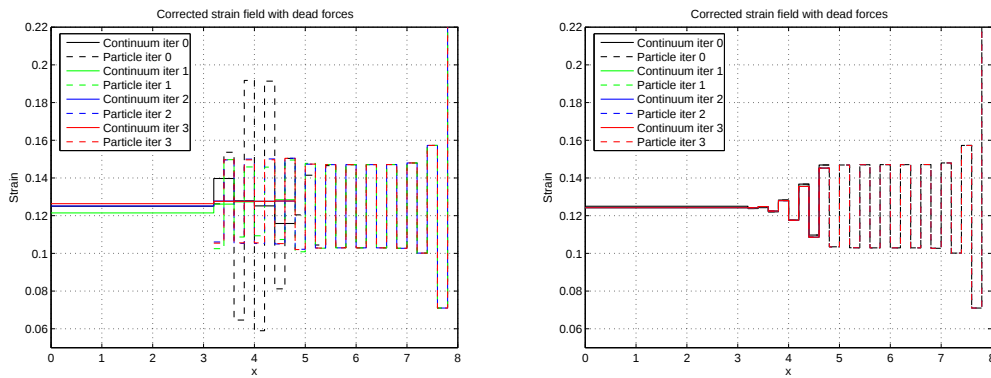


Figure 3.23: Champs de déformation corrigés aux itérations 0, 1, 2 et 3 pour $s = 2$ (gauche) et $s = 1$ (droite)

Dans [Chamoin *et al.* 2010], la technique de correction a aussi été appliquée à un problème 2D. On considère pour cela la structure de la Figure 3.24 faite d'un réseau carré régulier Ω de 51×51 particules avec une longueur caractéristique $\ell = 0,1$. La structure est encastree sur son bord et une force ponctuelle $\mathbf{F}$ est appliquée à la particule centrale. Cela définit une région localisée Ω_d où les déformations sont importantes et dans laquelle le modèle particulaire est conservé ; dans le reste du domaine, on remplace le réseau particulaire par un modèle continu équivalent (Figure 3.24). On considère des interactions avec les premiers, seconds et troisièmes voisins, avec des coefficients de raideur homogènes valant 10, 5 et 10, respectivement. Le modèle continu issu de l'homogénéisation est un modèle d'élasticité linéaire avec symétrie cubique, avec $E = 56,18$, $\nu = 0,215$, et $G = 24,17$. Dans la région de couplage Ω_o , le champ de multiplicateurs de Lagrange est discrétisé sur le maillage EF introduit dans Ω_c , fait d'éléments Q4 de taille $2\ell \times 2\ell$.

La correction avec les forces mortes est effectuée en introduisant une couche de particules supplémentaires vers l'interface. Les forces mortes calculées à l'itération 5 et interpolées sur Ω_c et Ω_d sont montrées sur la Figure 3.25. On observe que ces

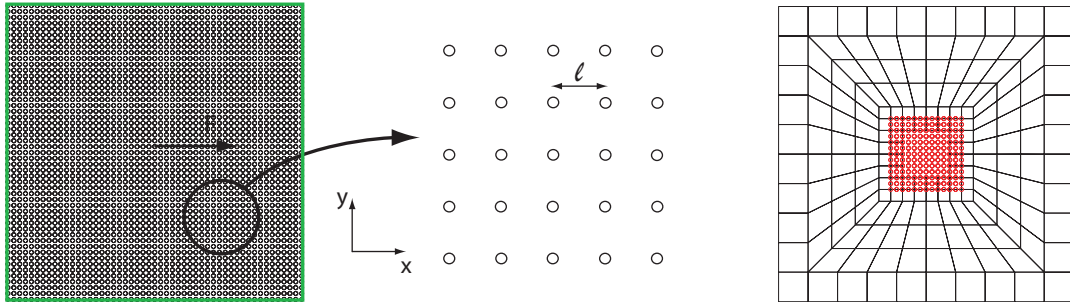


Figure 3.24: Réseau 2D définissant le modèle initial (gauche) et le modèle de couplage (droite)

forces mortes, qui diminuent de deux ordres entre les itérations 1 et 5, sont non nulles seulement dans un voisinage local de la région de couplage Ω_o .

Pour conclure cette partie, on compare les performances de la technique de correction par rapport à une procédure classique d'adaptation de modèle avec élargissement de la zone Ω_d jusqu'à atteindre un critère d'erreur donné. On reprend une structure 1D similaire à celle de la Figure 3.13, faite de 171 particules distantes de $\ell = 0,2$. Les interactions se font avec les premiers et seconds voisins, avec des coefficients de raideur uniformes de valeur 100 et 50, respectivement. Le défaut local modélisé par un affaiblissement de la raideur affecte environ 20 particules autour de la particule centrale P . On s'intéresse à deux quantités d'intérêt :

1. le déplacement de la particule P : $Q_1 = z_P$;
2. l'étirement de la liaison entre la particule P et la particule $P - 1$: $Q_2 = (z_P - z_{P-1})/\ell$

Dans le problème couplé, on discrétise le milieu continu et les multiplicateurs de Lagrange avec un maillage EF de taille caractéristique $h = 4\ell$ ($s = 4$) et on prend une largeur égale à 8ℓ pour les deux régions de couplage. La fonction poids $\alpha_c(x)$ est choisie linéaire dans ces régions de couplage.

On compare les valeurs approchées de Q_1 et Q_2 données avec ou sans introduction de la technique de correction dans le couplage Arlequin, en fonction de la taille de la région particulière Ω_d ; la Figure 3.26 montre la solution Arlequin en utilisant une largeur égale à 2ℓ ou 42ℓ pour Ω_d . Les résultats de comparaison sont donnés dans le Tableau 3.2 ; ils montrent l'erreur relative obtenue avec la méthode Arlequin sans aucune correction, puis après 5 et 10 itérations avec la technique de correction.

On observe que la technique de correction permet d'améliorer la précision sur les valeurs prédites des quantités Q_1 et Q_2 . Cependant, en fonction du coût de calcul et pour un critère de précision donné, il peut être plus efficace d'élargir la région particulière Ω_d plutôt que d'utiliser la technique de correction ; il est donc pertinent de faire un compromis entre coût de calcul et précision, ainsi que de développer des techniques d'adaptation de modèle associées à des estimations d'erreur (voir Section 2).

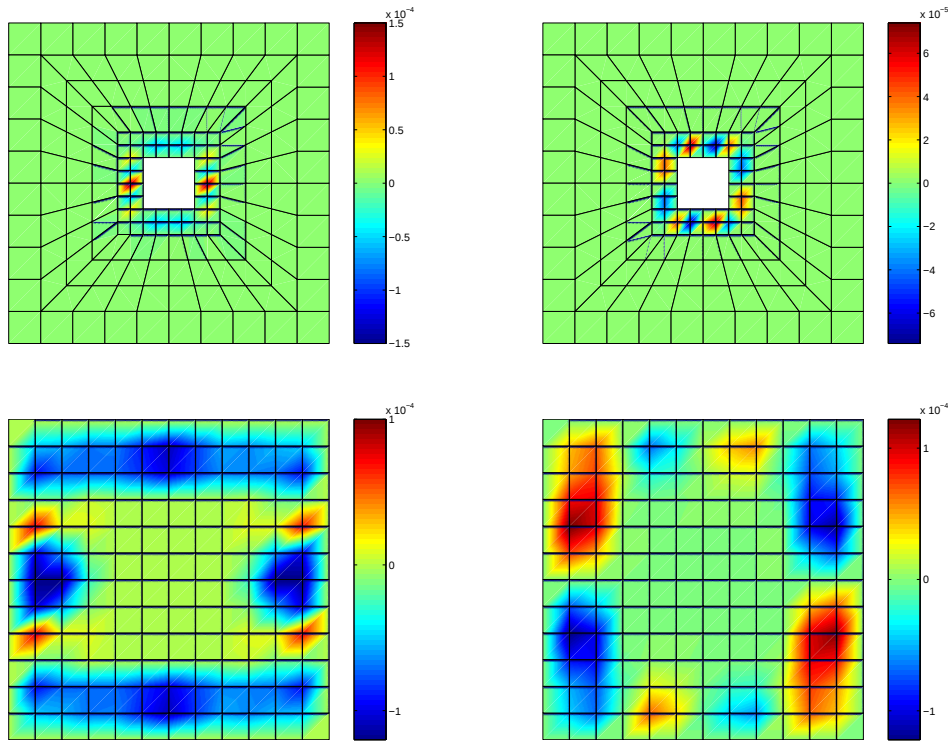


Figure 3.25: Carte d'interpolation des forces mortes pondérées appliquées aux nœuds dans Ω_c (haut) et aux particules dans Ω_d (bas) à l'itération 5. La composante x (resp. y) de ces forces mortes est représentée à gauche (resp. droite)

1.4.2 Cas de couplages en dynamique

La maîtrise du couplage particulaire/continu en dynamique nécessite des outils supplémentaires par rapport à ceux présentés précédemment. En effet, la nature propagative de la solution et l'existence de nouvelles sources d'effets non-physiques (liées à d'autres causes d'incompatibilités entre les modèles) rendent insuffisante la correction par forces mortes. Le phénomène non-physique majeur engendré par le couplage en dynamique est la réflexion d'ondes sur l'interface de couplage, ce phénomène ayant initialement été mis en exergue dans [Bazant 1978] pour le couplage de deux milieux continus avec différentes discrétisations.

Comme indiqué dans [E *et al.* 2006], les natures des incompatibilités rencontrées dans le couplage particulaire/continu en dynamique peuvent être séparées en deux types : (i) celles associées au passage entre un modèle particulaire non-local (interactions à longue distance) vers un modèle particulaire local (interactions avec les plus proches voisins seulement) ; (ii) celles associées au couplage entre un modèle particulaire local et un modèle continu issu d'un processus d'homogénéisation. Les travaux menés dans la thèse de J. Marchais, dont le cadre est la simulation numérique de la fissuration dans les bétons fibrés avec un modèle particulaire non-local (Figure 3.27), ont suivi cette philosophie et ont visé à traiter les deux sources d'incompatibilité séparément.

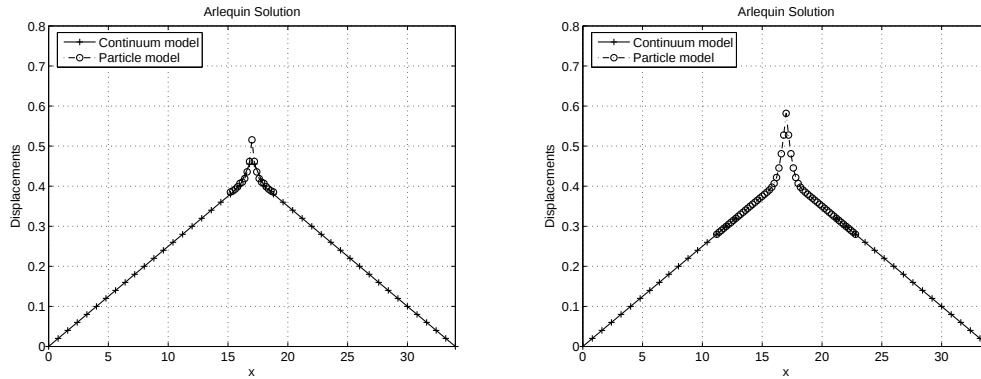


Figure 3.26: Déplacements approchés obtenus dans le cadre Arlequin avec différentes largeurs pour Ω_d : 2ℓ (gauche), 42ℓ (droite)

taille de $\Omega_d (\times \ell)$	erreur sur Q_1 (%)				erreur sur Q_2 (%)		
	iter. 0	iter. 5	iter. 10		iter. 0	iter. 5	iter. 10
2	8,72	1,14	0,26		11,66	2,12	0,97
10	0,13	$6,0 \cdot 10^{-4}$	$2,2 \cdot 10^{-5}$		$7,8 \cdot 10^{-3}$	$6,1 \cdot 10^{-5}$	$9,4 \cdot 10^{-6}$
18	$3,5 \cdot 10^{-5}$	$9,1 \cdot 10^{-7}$	$1,3 \cdot 10^{-7}$		$1,5 \cdot 10^{-5}$	$4,5 \cdot 10^{-8}$	$1,2 \cdot 10^{-8}$
26	$3,9 \cdot 10^{-7}$	$7,4 \cdot 10^{-9}$	$1,1 \cdot 10^{-9}$		$7,5 \cdot 10^{-8}$	$8,3 \cdot 10^{-9}$	$2,3 \cdot 10^{-10}$
34	$9,1 \cdot 10^{-9}$	$8,3 \cdot 10^{-10}$	$4,5 \cdot 10^{-10}$		$6,8 \cdot 10^{-9}$	$4,4 \cdot 10^{-10}$	$8,2 \cdot 10^{-11}$
42	$1,6 \cdot 10^{-10}$	$8,8 \cdot 10^{-11}$	$3,1 \cdot 10^{-11}$		$2,9 \cdot 10^{-10}$	$7,4 \cdot 10^{-11}$	$9,2 \cdot 10^{-12}$

Tableau 3.2: Erreur sur les valeurs de Q_1 et Q_2 en fonction de la largeur de Ω_d et du nombre d'itérations dans la technique de correction

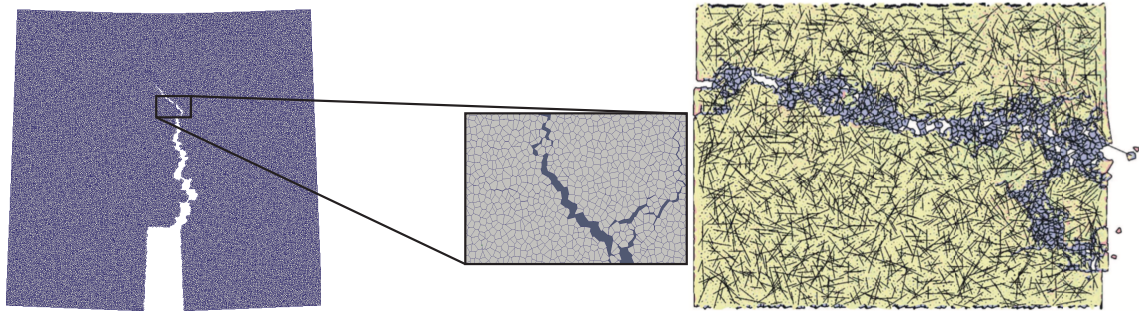


Figure 3.27: Simulation numérique d'un béton fibré avec un modèle particulaire (source : A. Delaplace & F.K. Wittel)

Le second type d'incompatibilité, lié au processus d'homogénéisation et connecté à une transition d'échelle, fait l'objet de travaux en cours (filtrage partiel d'ondes notamment) et n'est pas abordé ici. On se focalise dans cette section sur le premier type d'incompatibilité qui provient des natures géométriquement différentes entre l'interface du modèle local (interface surfacique) et celle du modèle non-local (in-

terface volumique), et est donc fortement lié aux conditions limites. Dans le cadre de la méthode Quasi-Continuum (*QC*), des techniques de correction basées sur une redéfinition de l'énergie ont été récemment proposées sur le sujet. Dans [Shimokawa *et al.* 2004] par exemple, un nouveau genre de particules (appelées *quasi-nonlocales*) a été introduit au niveau de l'interface, l'énergie associée étant un mélange entre les formulations énergétiques des particules locales et non-locales; cette technique corrige les effets non-physiques pour des distances d'interaction relativement courtes et en statique. Des extensions ont aussi été apportées dans [E *et al.* 2006, Shapeev 2011, Ortner et Zhang 2011] mais elles restent limitées à des applications 2D avec des formes particulières d'interface, et n'abordent pas non plus le cas de la dynamique.

Dans [Marchais *et al.* 2013], une technique consistante et générique, basée sur l'énergie, a été mise en place pour faire la transition entre un modèle particulaire avec des interactions non-locales et un autre avec des interactions avec les premiers voisins seulement. Pour cela, on se focalise sur la construction d'une zone de transition, libre de forces fantômes et avec conservation d'énergie, qui définit une interface volumique entre les deux modèles et permet de traiter les incompatibilités. Sur cette zone de transition, une approximation géométriquement consistante (au sens de [E *et al.* 2006]) de l'énergie du système particulaire est définie et construite de façon automatique. Deux schémas possibles de reconstruction sont présentés et analysés.

En repartant de la définition de consistance géométrique donnée dans [E *et al.* 2006], on définit une reconstruction d'énergie explicite et automatique valable pour toute dimension spatiale. Cette reconstruction agit dans la zone de transition entre les modèles particuliers local et non-local et introduit une interface volumique définie rigoureusement (ensemble de particules encadrées sur la Figure 3.28) sans modifier les modèles concurrents initiaux.

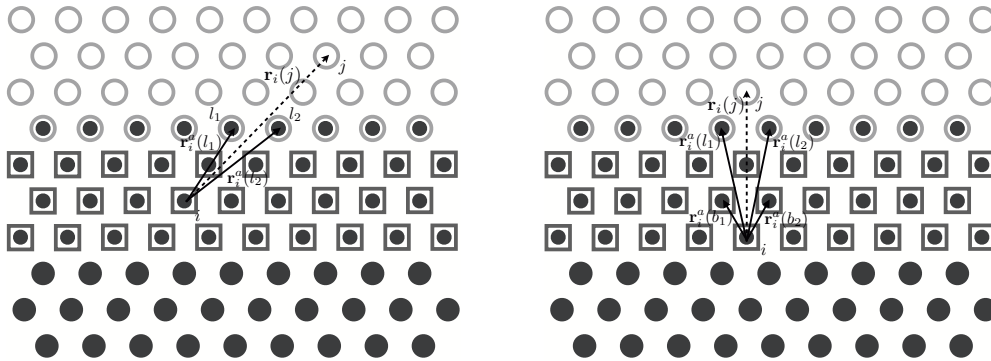


Figure 3.28: Illustration des schémas de reconstruction *NL2L* et *NL2L-loc*

Dans la zone de transition volumique, l'énergie de chaque particule i est reconstruite par redéfinition des distances d'interaction inter-particulaires avec les particules j , notées $\mathbf{r}_i(j)$, à partir des distances d'interaction $\mathbf{r}_i^a(k) = \mathbf{x}_k + \mathbf{z}_k - \mathbf{x}_i - \mathbf{z}_i$ avec un ensemble donné de particules k situées dans un voisinage de i . La nouvelle formulation suivie dans [Marchais *et al.* 2013] offre une flexibilité qui permet de définir deux schémas de reconstruction possibles. Le premier schéma, noté *NL2L*, conserve l'énergie globalement et est équivalent à celui présenté dans [Shapeev

2011]; pour ce schéma, les particules k servant à la reconstruction sont celles situées sur la couche frontière avec le modèle local (particules entourées sur la Figure 3.28). Le second schéma, noté *NL2L-loc*, conserve l'énergie localement; pour ce schéma, les particules k servant à la reconstruction sont situées dans toute la zone de transition.

On analyse tout d'abord les performances de ces deux schémas sur un exemple 2D en statique (Figure 3.29(a)). On considère un réseau hexagonal régulier de particules, avec un pas $r_0 = 1$. L'échantillon contient 121×121 particules, et les interactions entre les particules i et j sont décrites par un potentiel de Lennard-Jones de la forme :

$$4\epsilon_{ij} \left[\left(\frac{\sigma_{ij}}{r_{ij}} \right)^{12} - \left(\frac{\sigma_{ij}}{r_{ij}} \right)^6 \right] \quad (3.48)$$

avec $\epsilon_{ij} = 1/72$, $\sigma_{ij} = (\frac{1}{2})^{1/6} |\mathbf{x}_j - \mathbf{x}_i|$, et $r_{ij} = |\mathbf{x}_j + \mathbf{z}_j - \mathbf{x}_i - \mathbf{z}_i|$. On utilise une des-

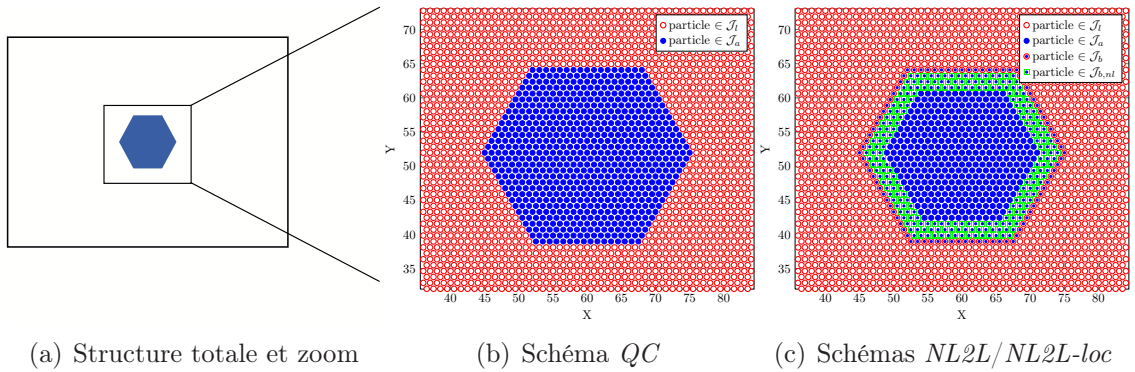


Figure 3.29: Représentation du problème couplé 2D

cription non-locale de l'énergie (interactions à longue distance) pour un ensemble de particules situées au centre de l'échantillon (Figure 3.29(a)); les interactions incluent alors jusqu'aux huitièmes voisins (rayon d'interaction $r_c = 4r_0$). On donne la distribution des particules locales et non-locales pour le schéma classique *QC* (Figure 3.29(b)) et les nouveaux schémas *NL2L* et *NL2L-loc* (Figure 3.29(c)); pour ce dernier cas, la zone de transition est indiquée en vert.

On suppose que la structure présente un défaut local, modélisé en enlevant un ensemble de 7 particules au centre de la structure. En appliquant un chargement de traction, on trace sur la Figure 3.30 l'erreur relative obtenue dans le domaine non-local avec les schémas *QC*, *NL2L* et *NL2L-loc*, respectivement. La taille du domaine non-local choisie est telle que 47 particules non-locales se trouvent le long de la direction horizontale. On observe que l'erreur décroît de deux ordres de grandeur quand on utilise les schémas *NL2L* ou *NL2L-loc* vis-à-vis d'un schéma *QC* classique. On remarque d'autre part des concentrations d'erreur (négligeables) dans les coins de l'interface avec le schéma *NL2L-loc* car ce schéma n'est pas totalement libre de forces fantômes près des coins. On représente finalement sur la Figure 3.31 l'évolution de l'erreur relative en fonction de la taille du domaine non-local : sur la Figure 3.31(a), l'erreur est mesurée sur un domaine fixe comprenant 18 particules au voisinage du

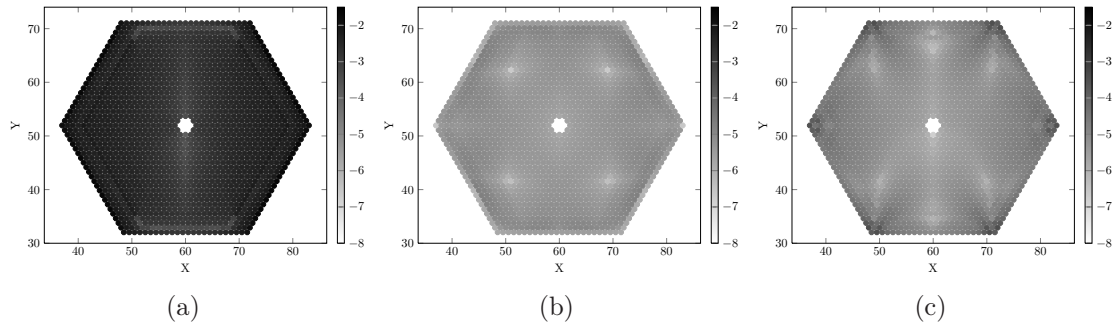
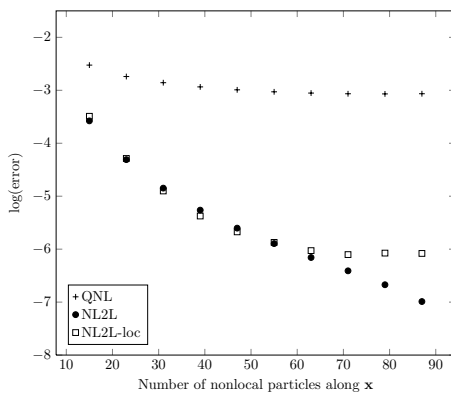


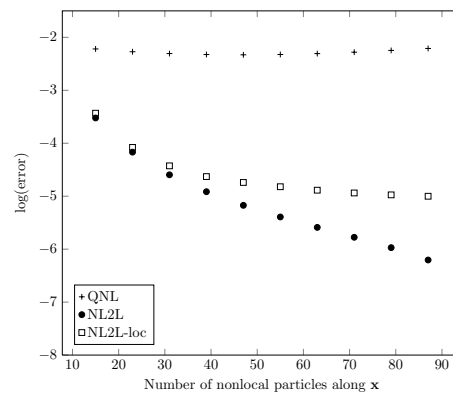
Figure 3.30: Erreur (log) sur les particules non-locales avec les schémas QC (a), $NL2L$ (b) et $NL2L-loc$ (c)

centre de l'échantillon ; sur la Figure 3.31(b), l'erreur est mesurée sur l'ensemble total des particules non-locales. Plusieurs observations peuvent être faites :

- comme prévu, pour un domaine non-local trop petit (partie gauche des courbes), tous les schémas de couplage mènent à une erreur non négligeable, bien qu'elle soit plus faible pour les schémas $NL2L$ et $NL2L-loc$. Ceci est dû au fait que dans ce cas, la déformation n'est pas suffisamment régulière dans la zone de transition ;
- quand on utilise le schéma QC , l'erreur ne décroît pas significativement avec l'augmentation de la taille du domaine non-local, dû au fort impact des forces fantômes ;
- l'erreur obtenue avec le schéma $NL2L-loc$ tend vers une valeur asymptotique quand on augmente la taille du domaine non-local, dû aux effets de pollution apportés par les forces fantômes près des coins de l'interface. Néanmoins, cette valeur asymptotique reste très faible ;
- l'erreur obtenue avec le schéma $NL2L$ diminue de façon monotone avec la taille du domaine non-local.



(a) $\log(\text{erreur})$ sur 18 particules



(b) $\log(\text{erreur})$ sur les particules non-locales

Figure 3.31: Evolution de l'erreur (log) en fonction de la taille du domaine non-local

On utilise à présent les schémas proposés sur un problème de dynamique pour illustrer le contrôle des réflexions d'onde induites par les interfaces non-local/local.

On considère une chaîne 1D, de pas $r_0 = 1$, faite de 300 particules. Les interactions sont à nouveau décrites par un potentiel de Lennard-Jones et incluent les particules jusqu'aux huitièmes voisins ($r_c = 8r_0$). La masse des particules est $m = 10$ et on utilise un schéma explicite en temps (schéma de Verlet) avec un pas de temps $\Delta t = 0,1$. Dans le problème couplé, la région non-locale est centrée et constituée de 60 particules.

On impose un déplacement initial au centre de la chaîne générant une onde se propageant de part et d'autre du domaine; l'énergie de la solution de référence est représentée en fonction du temps sur la Figure 3.32(a), avec séparation sur les deux ensembles de particules qui respectivement passent locales ou restent non-locales quand on utilise les schémas *QC*, *NL2L* ou *NL2L-loc*. On compare alors cette solution de référence, en terme d'énergie, avec celles obtenues avec le schéma *QC* et les schémas *NL2L* et *NL2L-loc* proposés (Figure 3.32). Initialement, toute l'énergie est concentrée dans la région non-locale jusqu'à ce que l'onde s'étende sur la région locale. Néanmoins, comme on peut le voir sur la Figure 3.32(b), une part importante de l'énergie (20%) reste piégée dans la région non-locale à cause des réflexions aux interfaces quand on utilise le schéma *QC*. Quand on utilise les schémas *NL2L* ou *NL2L-loc*, toute l'énergie est transmise aux régions locales (Figures 3.32(c) et 3.32(d)), sans réflexion à l'interface.

La même simulation est menée pour diverses valeurs de rayon de coupure r_c et diverses longueurs d'onde $\lambda = n \times r_c$. Les résultats présentés sur la Figure 3.33 montrent que les schémas *NL2L* et *NL2L-loc* occasionnent une erreur en énergie bien inférieure à celle du schéma *QC* pour une large gamme de fréquences (hautes fréquences en particulier). Ils montrent aussi, comme prévu, qu'aucun schéma n'est capable de propager une onde de longueur plus petite que le rayon de coupure ($n < 1$).

Pour finir, on regarde les performances des schémas de couplage quand une onde transite du modèle local vers le modèle non-local, comme cela peut être le cas après réflexion sur les frontières de la structure globale. On garde pour cela la même chaîne 1D mais on inverse les modèles locaux et non-locaux. L'énergie de la solution de référence, toujours séparée en deux parties, est représentée en fonction du temps sur la Figure 3.34. On compare cette solution de référence pleinement non-locale, en terme d'énergie, avec celles obtenues par les schémas *QC*, *NL2L*, et *NL2L-loc* (Figure 3.34). On observe que toute l'énergie va dans le domaine non-local quand on utilise les schémas *QC* ou *NL2L-loc*, mais qu'une part importante de l'énergie (25%) reste piégée dans la région locale quand on utilise le schéma *NL2L* (voir Figure 3.34(c)). Quand on réalise la même simulation avec diverses valeurs de r_c et diverses longueurs d'onde (Figure 3.35), les mêmes observations peuvent être faites. D'un côté, cela montre que le schéma *NL2L* n'est pas efficace quand une onde va d'un modèle particulière local vers un modèle particulière non-local, même pour de grandes longueurs d'onde comparées au rayon de coupure r_c . D'autre part, en générant une distribution d'énergie régulière et progressive entre les plus proches particules et les particules d'interface dans la zone de transition, le schéma *NL2L-loc* permet d'obtenir une propagation d'onde propre du modèle local vers le modèle non-local, pourvu que la longueur d'onde soit suffisamment grande par rapport à r_c .

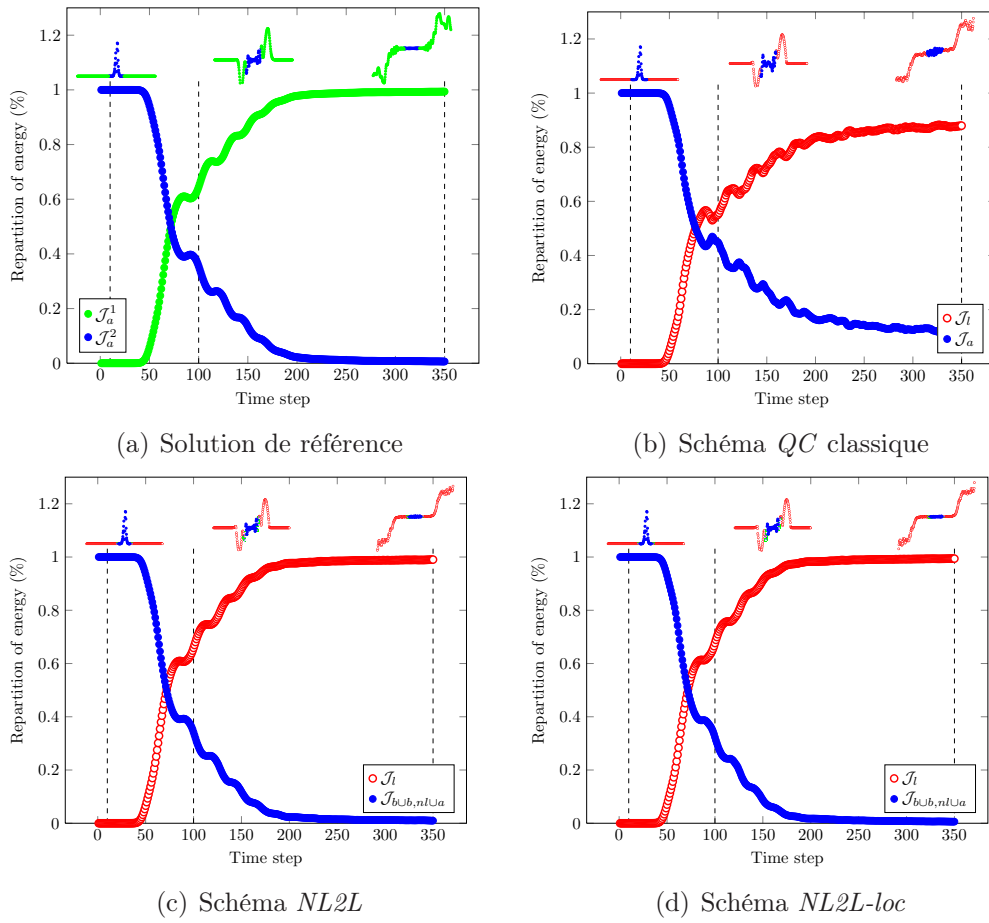


Figure 3.32: Répartition de l'énergie en fonction du temps pour une propagation d'onde du modèle non-local vers le modèle local

2 Stratégie adaptative de couplage

Une difficulté de la modélisation multiéchelle avec des modèles concurrents est d'identifier les régions les plus appropriées pour placer chacun des deux modèles, i.e. de déterminer la position optimale de la région de couplage qui n'est pas connue *a priori* et dépend généralement de l'objectif de la simulation. On présente dans cette partie des outils permettant d'une part de contrôler l'erreur créée par l'utilisation d'un modèle approché multiéchelle sur la prédiction d'une quantité d'intérêt, et d'autre part d'adapter le modèle multiéchelle si besoin pour atteindre une tolérance d'erreur souhaitée. Evidemment, la discrétisation des modèles influe sur la qualité du modèle multiéchelle ; sauf cas contraire, l'erreur de discrétisation sera ici supposée préalablement maîtrisée et négligeable par rapport aux autres sources d'erreur.

2.1 Méthode générale

L'estimation d'erreur sur une quantité d'intérêt dans le cadre du remplacement d'un modèle fin initial, pris comme référence, par un modèle réduit plus grossier a été abordée assez récemment. Après des premiers travaux visant à estimer l'erreur

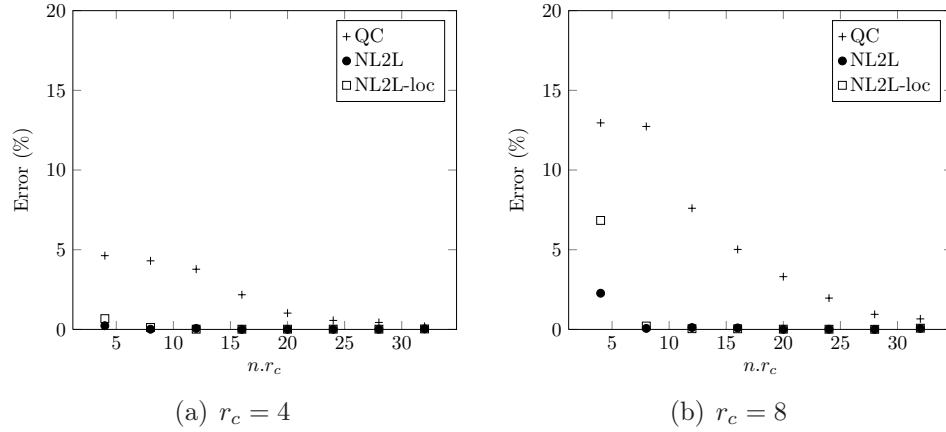


Figure 3.33: Déviation sur la répartition de l'énergie en fonction de la longueur d'onde pour une propagation d'onde du modèle non-local vers le modèle local

globale [Oden et Zohdi 1997], une extension aux quantités locales a ensuite été proposée et appliquée à divers problèmes tels que ceux impliquant des matériaux hétérogènes [Oden *et al.* 1999, Oden et Vemaganti 2000, Vemaganti et Oden 2001] ou la propagation d'ondes dans les matériaux élastiques [Romkes et Oden 2004]. Les bases de la méthode employée, dont une revue générale peut être trouvée dans [Oden et Prudhomme 2002, Oden *et al.* 2005 - 2006], sont brièvement présentées ci-après. Le processus d'adaptation de modèle qui est associé, utilisant un algorithme glouton, est présenté dans un second temps.

2.1.1 Estimation d'erreur locale

On suppose l'existence d'un modèle mathématique, pris pour référence, suffisamment fin pour capturer tous les phénomènes physiques d'intérêt. Ce problème est écrit sous la forme faible générique suivante :

$$\text{Trouver } \mathbf{u} \in \mathcal{U} \text{ tel que } a(\mathbf{u}; \mathbf{v}) = l(\mathbf{v}) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (3.49)$$

avec a une forme linéaire par rapport à $\mathbf{v}$ mais potentiellement non-linéaire par rapport à $\mathbf{u}$ (la notation utilisée implique une linéarité par rapport aux variables situées à droite du signe “;”). Le problème (3.49) étant supposé inabordable, il est remplacé par un modèle réduit obtenu par un couplage multiéchelle entre le modèle fin initial et un modèle plus grossier. Ce nouveau modèle est écrit sous la forme :

$$\text{Trouver } \mathbf{u}_0 \in \mathcal{U} \text{ tel que } a_0(\mathbf{u}_0; \mathbf{v}) = l(\mathbf{v}) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (3.50)$$

où a_0 est l'opérateur obtenu par le processus multiéchelle.

On identifie une quantité d'intérêt représentable par une fonctionnelle $Q(\mathbf{u})$ de la solution fine, et supposée linéaire par rapport à $\mathbf{u}$. On introduit alors le problème adjoint (voir Chapitre 1) :

$$\text{Trouver } \tilde{\mathbf{u}} \in \mathcal{U}_0 \text{ tel que } a'(\mathbf{u}; \mathbf{v}, \tilde{\mathbf{u}}) = a'^*(\tilde{\mathbf{u}}, \mathbf{u}; \mathbf{v}) = Q(\mathbf{v}) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (3.51)$$

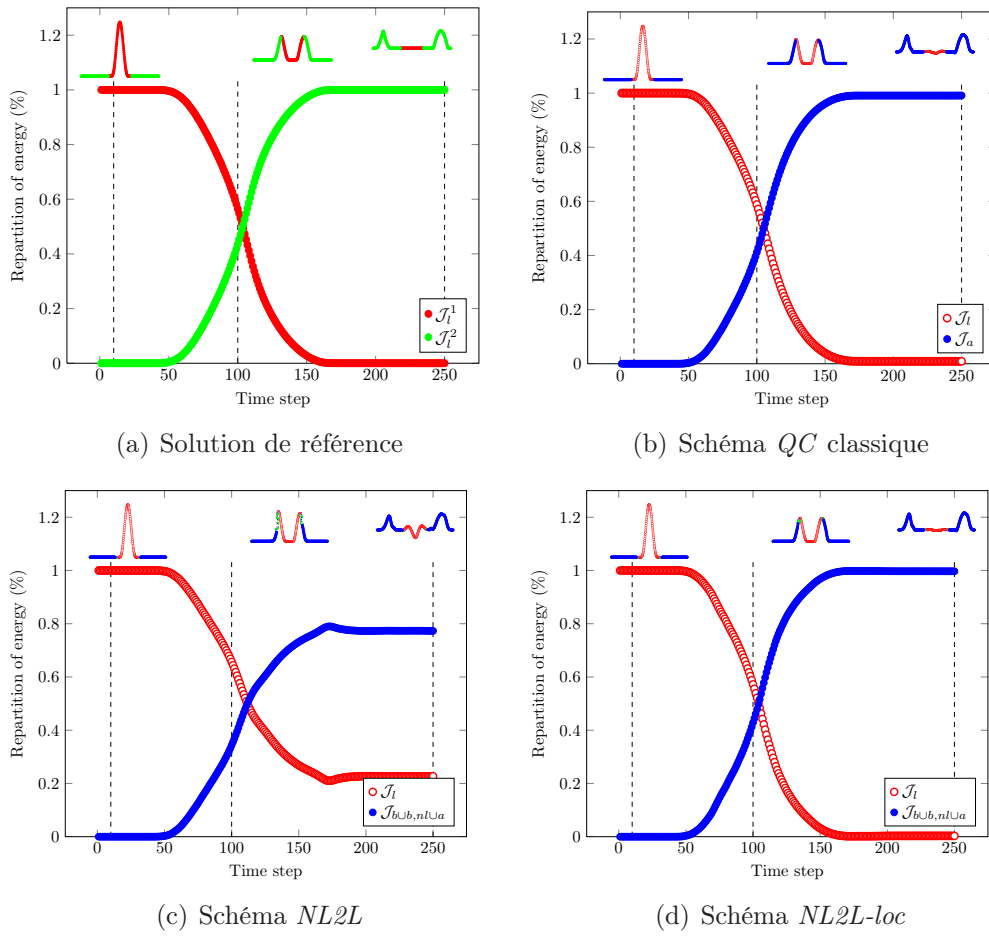


Figure 3.34: Répartition de l'énergie en fonction du temps pour une propagation d'onde du modèle local vers le modèle non-local

avec

$$a'(\mathbf{u}; \mathbf{v}, \tilde{\mathbf{u}}) = \lim_{\theta \rightarrow 0} \frac{1}{\theta} [a(\mathbf{u} + \theta \mathbf{v}; \tilde{\mathbf{u}}) - a(\mathbf{u}; \tilde{\mathbf{u}})] \quad (3.52)$$

Le problème adjoint (3.51), bien que linéaire, reste en pratique inabordable pour deux raisons : (i) la solution $\mathbf{u}$ n'est pas connue ; (ii) la dimension du problème est identique à celle du problème de référence (3.49). On calcule donc une solution approchée de $\tilde{\mathbf{u}}$ par un processus multiéchelle menant au problème :

$$\text{Trouver } \tilde{\mathbf{u}}_0 \in \mathcal{U}_0 \text{ tel que } a'_0(\mathbf{u}_0; \mathbf{v}, \tilde{\mathbf{u}}_0) = Q(\mathbf{v}) \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (3.53)$$

Dès lors, il est montré dans [Oden et Prudhomme 2002] que l'erreur sur la quantité d'intérêt Q peut s'exprimer sous la forme suivante :

$$Q(\mathbf{u}) - Q(\mathbf{u}_0) = \mathcal{R}(\mathbf{u}_0; \tilde{\mathbf{u}}) + \Delta = \mathcal{R}(\mathbf{u}_0; \tilde{\mathbf{u}}_0) + \mathcal{R}(\mathbf{u}_0; \tilde{\mathbf{u}} - \tilde{\mathbf{u}}_0) + \Delta \quad (3.54)$$

où $\mathcal{R}(\bullet; \bullet)$ est la fonctionnelle des résidus : $\mathcal{R}(\mathbf{u}_0; \tilde{\mathbf{u}}) = F(\tilde{\mathbf{u}}) - B(\mathbf{u}_0; \tilde{\mathbf{u}})$, et Δ regroupe tous les termes d'ordre élevé des erreurs $\mathbf{e}_0 = \mathbf{u} - \mathbf{u}_0$ et $\tilde{\mathbf{e}}_0 = \tilde{\mathbf{u}} - \tilde{\mathbf{u}}_0$. Si a

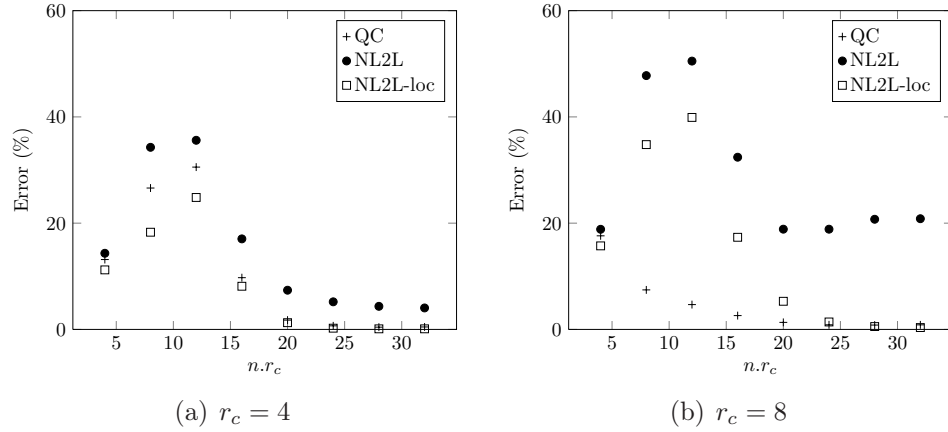


Figure 3.35: Déviation sur la répartition de l'énergie en fonction de la longueur d'onde pour une propagation d'onde du modèle local vers le modèle non-local

est dérivable trois fois, Δ peut être écrit explicitement :

$$\begin{aligned} \Delta = & \frac{1}{2} \int_0^1 a''(\mathbf{u}_0 + s\mathbf{e}_0; \mathbf{e}_0, \mathbf{e}_0, \tilde{\mathbf{u}}_0 + s\tilde{\mathbf{e}}_0) ds \\ & - \frac{1}{2} \int_0^1 (3a''(\mathbf{u}_0 + s\mathbf{e}_0; \mathbf{e}_0, \mathbf{e}_0, \tilde{\mathbf{e}}_0) + a'''(\mathbf{u}_0 + s\mathbf{e}_0; \mathbf{e}_0, \mathbf{e}_0, \mathbf{e}_0, \tilde{\mathbf{u}}_0 + s\tilde{\mathbf{e}}_0))(s-1) s ds \end{aligned} \quad (3.55)$$

Généralement, ce terme est négligé en considérant $\|\mathbf{e}_0\|$ et $\|\tilde{\mathbf{e}}_0\|$ petits.

Dès lors, deux stratégies d'estimation d'erreur apparaissent selon le problème traité :

- s'il est possible d'avoir un encadrement de $\mathcal{R}(\mathbf{u}_0; \tilde{\mathbf{u}} - \tilde{\mathbf{u}}_0)$, avec la méthode de l'ERC par exemple, on obtient des bornes garanties (au terme Δ près) sur l'erreur $Q(\mathbf{u}) - Q(\mathbf{u}_0)$;
- sinon, une estimation simple non garantie est $Q(\mathbf{u}) - Q(\mathbf{u}_0) \approx \mathcal{R}(\mathbf{u}_0; \tilde{\mathbf{u}}_0)$. Cette estimation est pertinente si $\tilde{\mathbf{u}}_0$ est calculé avec un modèle multiéchelle plus fin que celui utilisé pour $\mathbf{u}_0$.

2.1.2 Adaptation du modèle couplé

A partir de l'estimateur d'erreur locale précédemment mis en place, l'objectif est à présent de dériver un algorithme adaptatif qui génère automatiquement une séquence de modèles multiéchelle k ($k = 1, 2, \dots$), de solutions $(\mathbf{u}_0^{(k)}, \tilde{\mathbf{u}}_0^{(k)})$ associées, réduisant l'erreur sur la quantité d'intérêt Q . L'algorithme est appliqué jusqu'à atteindre, pour un modèle K , une valeur d'erreur tolérée γ_{tol} :

$$|Q(\mathbf{u}) - Q(\mathbf{u}_0^{(K)})| \leq \gamma_{\text{tol}} Q(\mathbf{u}_0^{(K)}) \quad (3.56)$$

A chaque itération, la quantité $\mathcal{R}(\mathbf{u}_0^{(k)}; \tilde{\mathbf{u}}_0^{(k)})$ est évaluée et comparée à la valeur γ_{tol} . Si le critère d'erreur souhaitée n'est pas validé, le modèle est enrichi localement en

élargissant la zone d'application du modèle fin aux zones dans lesquelles les sources d'erreur sont prépondérantes. Cette stratégie est rendue possible en observant que l'estimateur d'erreur $\eta_{\text{est}} = \mathcal{R}(\mathbf{u}_0^{(k)}; \tilde{\mathbf{u}}_0^{(k)})$ est défini de manière globale sur l'ensemble du domaine Ω d'étude, et peut être décomposé en contributions locales $\eta_{c,i}$ définies sur des sous-domaines i de Ω :

$$\eta_{\text{est}} = \sum_{i=1}^N \eta_{c,i} \quad (3.57)$$

Un choix naturel et possible de sous-domaines est l'ensemble des éléments du maillage EF utilisé pour la discrétisation du modèle grossier continu, bien que d'autres choix puissent être envisagés.

On suit alors un processus du maximum (ou glouton) dans lequel, après avoir défini un paramètre $\gamma_a \in]0, 1[$, les sous-domaines i dont la contribution d'erreur $\eta_{c,i}$ vérifie :

$$\frac{|\eta_{c,i}|}{\max_j |\eta_{c,j}|} > \gamma_a \quad (3.58)$$

deviennent associés au modèle fin. L'algorithme global d'adaptation est résumé dans la Figure 3.36. Cet algorithme détermine automatiquement une configuration de cou-

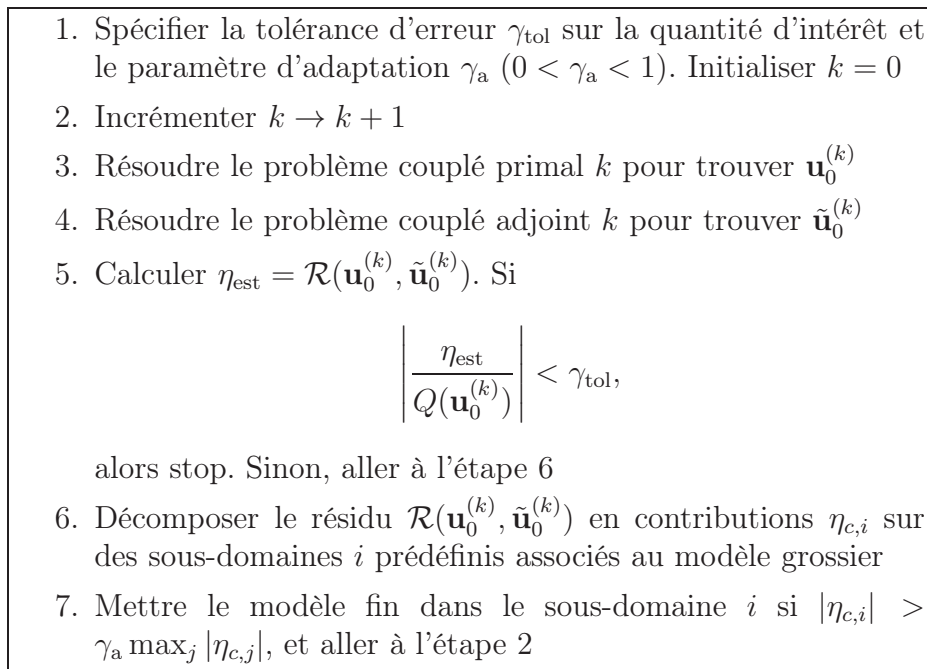


Figure 3.36: Algorithme adaptatif pour le contrôle de l'erreur locale de réduction de modèle

plage permettant de prédire des quantités d'intérêt avec une tolérance fixée. Il est nécessairement convergent, car dans le pire cas, le modèle fin est utilisé dans tout le domaine. Cependant, rien ne garantit une décroissance monotone de l'erreur, ni la convergence vers le modèle couplé optimal. De plus, des compensations d'erreur entre les sous-domaines peuvent exister ; elles rendent la stratégie d'adaptation généralement pessimiste vis-à-vis de la qualité du modèle couplé.

2.2 Application au couplage particulaire/continu

Les premiers travaux utilisant l'estimation d'erreur et les procédures adaptatives pour les couplages particulaire/continu sont apparus dans [Prudhomme *et al.* 2006, Arndt et Luskin 2007], dans le cadre de la méthode Quasi-Continuum. Dans [Prudhomme *et al.* 2008, Bauman *et al.* 2009, Prudhomme *et al.* 2009], les outils de contrôle ont été appliqués au couplage par la méthode Arlequin décrit dans la Section 1.1 ; on en donne ici les principales caractéristiques.

On reprend le problème de référence (3.4), qui consiste à trouver $\mathbf{z} \in \mathcal{Y}$ tel que :

$$a(\mathbf{z}; \mathbf{y}) = l(\mathbf{y}) \quad \forall \mathbf{y} \in \mathcal{Y} \quad (3.59)$$

ainsi que le problème couplé (3.15) qui consiste à trouver $(\mathbf{u}, \mathbf{z}, \boldsymbol{\lambda}) \in \mathcal{U} \times \mathcal{W} \times \mathcal{X}$ tel que :

$$a_0((\mathbf{u}, \mathbf{z}, \boldsymbol{\lambda}); (\mathbf{v}, \mathbf{w}, \boldsymbol{\mu})) = l(\mathbf{w}) \quad \forall (\mathbf{v}, \mathbf{w}, \boldsymbol{\mu}) \in \mathcal{U} \times \mathcal{W} \times \mathcal{X} \quad (3.60)$$

A partir de $(\mathbf{u}, \mathbf{z}) \in \mathcal{U} \times \mathcal{W}$, on redéfinit une solution approchée $\mathbf{z}_0 \in \mathcal{Y}$ avec :

$$\mathbf{z}_{0,i} = \begin{cases} \mathbf{u}(\mathbf{x}_i) & \text{si } \mathbf{x}_i \in \Omega_c \setminus \Omega_o \\ \mathbf{z}_i & \text{sinon} \end{cases} \quad (3.61)$$

bien que d'autres constructions soient possibles. La reconstruction (3.61) est particulièrement pertinente pour des matériaux avec potentiels homogènes et en petites déformations ; l'hypothèse de Cauchy-Born est valide dans ce cas.

En considérant une quantité d'intérêt $Q(\mathbf{z})$, le problème adjoint consiste à trouver $\tilde{\mathbf{z}} \in \mathcal{Y}$ tel que :

$$a'(\mathbf{z}; \mathbf{y}, \tilde{\mathbf{z}}) = Q(\mathbf{y}) \quad \forall \mathbf{y} \in \mathcal{Y} \quad (3.62)$$

avec $a'(\mathbf{z}; \mathbf{y}, \tilde{\mathbf{z}}) = \sum_{i=1}^n \sum_{j=1}^n \mathbf{y}_i \cdot \frac{\partial^2 E(\mathbf{z})}{\partial \mathbf{z}_i \partial \mathbf{z}_j} \cdot \tilde{\mathbf{z}}_j$ (opérateur hessien). Une approximation $\tilde{\mathbf{z}}_0$

de $\tilde{\mathbf{z}}$ est obtenue par un couplage multiéchelle Arlequin, en considérant un domaine particulaire Ω_d plus large que celui utilisé pour calculer $\mathbf{z}_0$. On définit alors l'estimateur d'erreur sur Q :

$$\eta_{\text{est}} = \mathcal{R}(\mathbf{z}_0; \tilde{\mathbf{z}}_0) \quad (3.63)$$

On illustre la méthode sur un exemple 2D similaire à celui de la Figure 3.24, avec une structure constituée de 51×51 particules. Les interactions sont modélisées par des potentiels harmoniques :

$$V_{ij}(\mathbf{z}_i, \mathbf{z}_j) = \frac{1}{2} k_{ij} (r_{ij} - r_{ij}^0)^2 \quad (3.64)$$

avec $r_{ij} = |\mathbf{z}_j - \mathbf{z}_i + \mathbf{x}_j - \mathbf{x}_i|$ la longueur actuelle de la liaison et k_{ij} , r_{ij}^0 des paramètres connus. On considère que $r_{ij} = r_{ij}^0$ quand $\mathbf{z}_i = \mathbf{z}_j = \mathbf{0}$. On se restreint ici aux interactions avec les premiers voisins, avec des paramètres de raideur égaux à 10. Les paramètres du modèle continu correspondant, issus de l'homogénéisation, sont $E = 13,65$, $\nu = 0,32$ et $G = 4,14$.

Dans une première application, on suppose que la structure est fixée sur son bord et soumise à une force ponctuelle unitaire appliquée à la particule centrale P . Dans le modèle couplé Arlequin, la zone où le modèle particulaire est conservé consiste en 13×13 particules centrées autour de la particule P . La région de couplage Ω_o a une largeur égale à 2ℓ et on choisit une fonction de pondération α_c linéaire dans Ω_o . La discrétisation choisie dans Ω_o pour la solution du modèle continu et le champ de multiplicateurs de Lagrange est liée au réseau particulaire (éléments $Q4$ de taille ℓ). On montre sur la Figure 3.37 la solution Arlequin dans la configuration déformée.

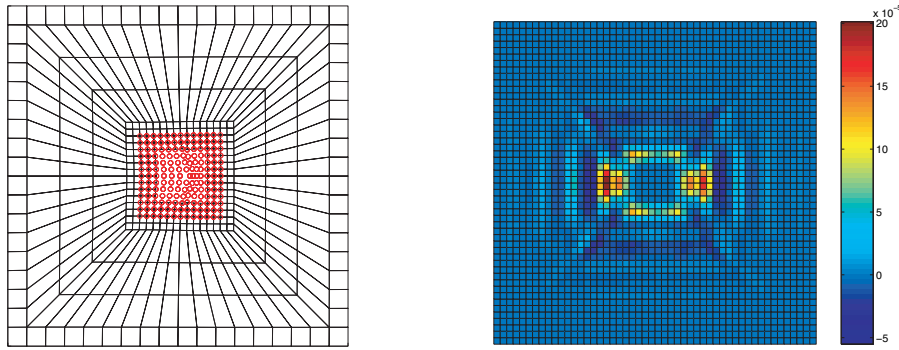


Figure 3.37: Problème Arlequin dans la configuration déformée (gauche) et distribution du résidu (droite)

On s'intéresse au déplacement horizontal de la particule P sur laquelle le chargement est appliqué; cette quantité d'intérêt peut se réécrire :

$$Q(\mathbf{z}) = \mathbf{z}_P \cdot \mathbf{s} \quad (3.65)$$

où $\mathbf{s}$ est le vecteur unitaire $(1, 0)$. Les valeurs exacte et approchée de la quantité d'intérêt sont alors $Q(\mathbf{z}) = 0,1017$ et $Q(\mathbf{z}_0) = 0,0941$, avec une erreur relative de 7%. Pour estimer le terme d'ordre élevé Δ dans (3.54), on calcule la solution adjointe exacte $\tilde{\mathbf{z}} \in \mathcal{Y}$. On a alors :

$$\Delta = Q(\mathbf{z}) - Q(\mathbf{z}_0) - \mathcal{R}(\mathbf{z}_0; \tilde{\mathbf{z}}) = 0,1017 - 0,0941 - 0,0076 = -7,71 \cdot 10^{-8} \quad (3.66)$$

ce qui montre que ce terme est très petit (moins de 0,001% d'erreur). La distribution spatiale de l'estimateur η_{est} , obtenu avec $\tilde{\mathbf{z}}_0 = \tilde{\mathbf{z}}$, est représentée sur la Figure 3.37; la contribution sur chaque zone carrée est obtenue par une moyenne des contributions des quatre particules voisines. On observe que la majorité de l'erreur se concentre au niveau de la zone de couplage Ω_o .

On regarde à présent l'influence des paramètres du couplage Arlequin (discrétisation des multiplicateurs de Lagrange, largeur de la zone Ω_o) sur l'erreur sur Q . Dans une configuration 1, on conserve la même discrétisation que précédemment dans Ω_o mais la largeur de Ω_o est augmentée à 2ℓ . Dans une configuration 2, on garde cette largeur de 2ℓ pour la zone Ω_o mais on prend une discrétisation deux fois plus grossière (éléments $Q4$ de taille 2ℓ). Les deux configurations sont représentées sur la Figure 3.38, et les résultats sur l'erreur sur Q sont donnés dans le Tableau 3.3.

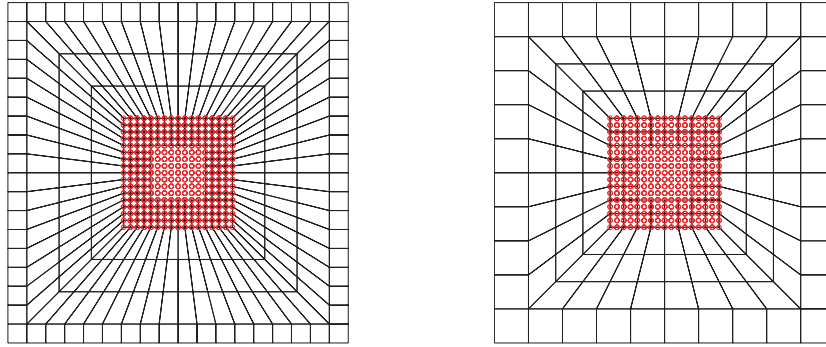


Figure 3.38: Configurations 1 (gauche) et 2 (droite) considérées pour observer l'influence des paramètres de couplage

$Q(\mathbf{y})$	$Q(\mathbf{y}_0)$ (% erreur)	$\mathcal{R}(\mathbf{z}_0; \tilde{\mathbf{z}})$	Δ
0,1017	0,0948(7%)	0,0069	$-7,15 \cdot 10^{-8}$
0,1017	0,0957(6%)	0,0060	$-2,71 \cdot 10^{-8}$

Tableau 3.3: Résultats obtenus sur Q pour les deux configurations de couplage testées

On observe que la taille et la discrétisation de la région de couplage Ω_o jouent un rôle non-négligeable sur la précision de la quantité Q .

On étudie maintenant l'influence de l'approximation $\tilde{\mathbf{z}}_0$ de la solution adjointe $\tilde{\mathbf{z}}$, i.e. la taille de la région dans laquelle le modèle fin est conservé pour résoudre le problème adjoint. On considère trois configurations de couplage pour le problème adjoint, illustrées sur la Figure 3.39. La dernière configuration correspond au même

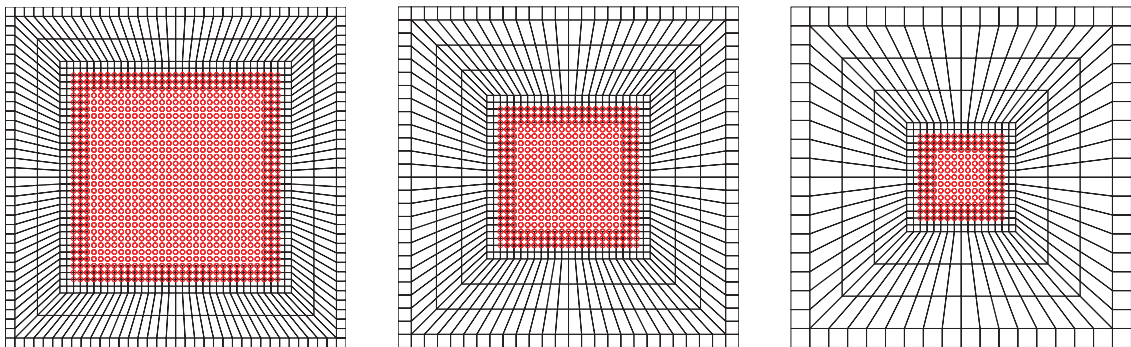


Figure 3.39: Configurations de couplage 1 (gauche), 2 (centre) et 3 (droite) utilisées pour résoudre le problème adjoint

couplage que celui utilisé pour obtenir $\mathbf{z}_0$. On calcule alors les estimateurs d'erreur sur Q obtenus avec ces trois configurations :

$$\eta_{\text{est},1} = 0,0073 \quad ; \quad \eta_{\text{est},2} = 0,0066 \quad ; \quad \eta_{\text{est},3} = 0,0052 \quad (3.67)$$

En comparant ces valeurs avec celle fournie par la solution adjointe de référence $\tilde{\mathbf{z}}$ ($\eta_{\text{est}} = 0,0076$), on remarque que l'estimation d'erreur est précise pourvu qu'on élargisse un peu (une couche d'éléments en pratique) la zone Ω_d pour résoudre le problème adjoint ; conserver une zone Ω_d semblable à celle utilisée pour le problème primal donne un estimateur assez peu pertinent.

Enfin, on regarde l'influence de la taille du domaine Ω_d utilisée dans le problème couplé primal sur l'erreur de réduction vis-à-vis de Q . On considère deux configurations du problème Arlequin et montrons les contributions de l'estimateur correspondantes sur la Figure 3.40. Dans les deux cas, on remarque que les contributions sont majoritaires au niveau de la zone de couplage Ω_o . Dans la première configuration, l'erreur relative sur Q est de 4% tandis qu'elle est de 2% dans la seconde configuration. Cela confirme que l'erreur varie grandement en fonction de la taille de la région Ω_d , et justifie l'emploi d'un algorithme adaptatif pour contrôler l'erreur.

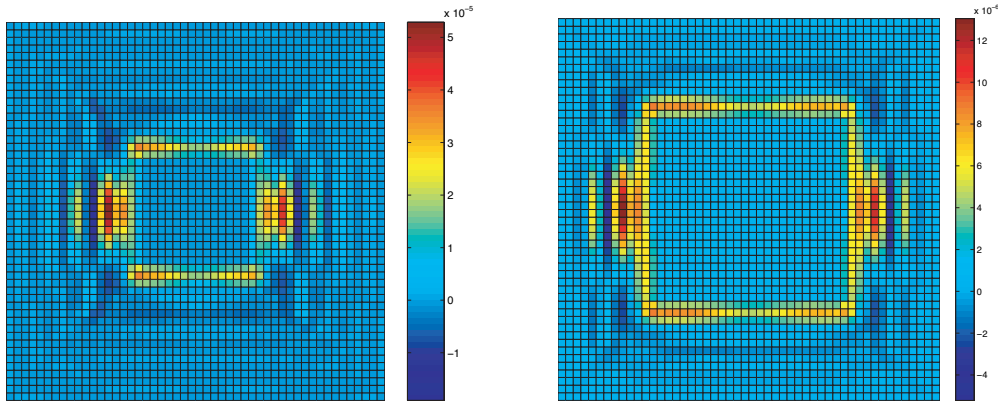


Figure 3.40: Contributions spatiales de l'estimateur d'erreur pour les configurations 1 (gauche) et 2 (droite) étudiées

Sur une seconde application, on montre les performances de l'algorithme adaptatif. On reprend la structure précédente et on suppose qu'elle est soumise à une force ponctuelle verticale $F = 1$ appliquée à la particule P_1 située au milieu du bord supérieur (Figure 3.41). La base de la structure est encastree et les deux faces latérales sont libres d'effort. La fonction de pondération α_c est choisie constante ($\alpha_c(\mathbf{x}) = 0,5$) sur Ω_o . Le maillage de la région Ω_c est constitué d'éléments Q4 réguliers de taille 4ℓ . La quantité d'intérêt considérée est l'étirement de la liaison entre les particules P_1 et P_2 (voir Figure 3.41) ; cette quantité n'étant pas linéaire par rapport à $\mathbf{z}$, on considère par simplicité sa forme linéarisée vis-à-vis de la configuration de référence :

$$Q(\mathbf{z}) = \frac{r_{12} - r_{12}^0}{r_{12}^0} \approx \mathbf{s}_{12} \cdot \frac{\mathbf{z}_{P_2} - \mathbf{z}_{P_1}}{r_{12}^0} \quad (3.68)$$

où r_{12} (resp. r_{12}^0) est la longueur de la liaison étirée (resp. à l'équilibre) entre les deux particules, et $\mathbf{s}_{12} = (\mathbf{x}_{P_2} - \mathbf{x}_{P_1})/r_{12}^0$ est le vecteur unitaire dans la direction

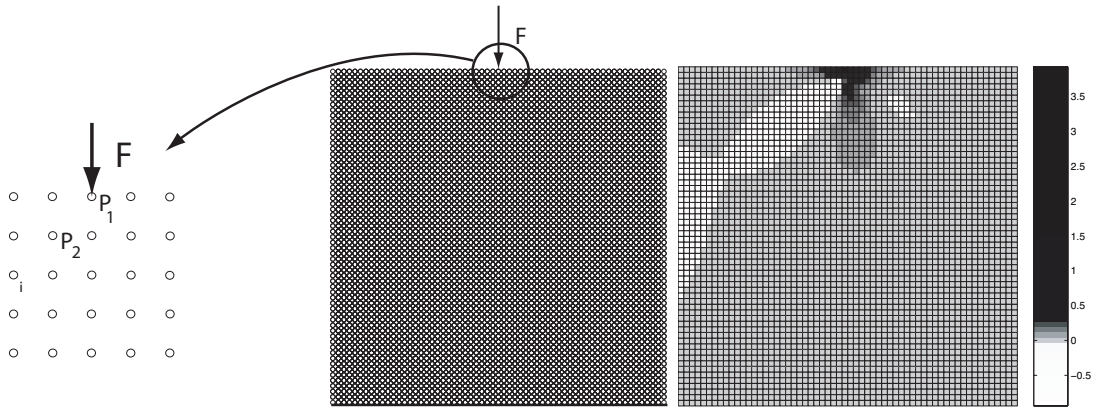


Figure 3.41: Problème de référence avec région d'intérêt (gauche) et composante de déformation ϵ_{yy} de la solution adjointe $\tilde{\mathbf{z}}$ (droite)

de la liaison. La solution adjointe correspondante (composante de déformation) est montrée sur la Figure 3.41.

On met alors en œuvre l'algorithme d'adaptation décrit sur la Figure 3.36, avec $\gamma_a = 0,3$ et $\gamma_{tol} = 1\%$. A partir d'une configuration initiale du couplage Arlequin dans laquelle la zone d'application du modèle particulaire seul est une zone de 5×5 particules autour de la particule P_1 , et pour laquelle l'erreur relative sur Q est de 6,3%, on montre les configurations obtenues aux itérations 1 à 6 de l'algorithme. Les erreurs relatives correspondantes sont respectivement 3,0%, 2,5%, 2,1%, 1,9%, 1,8%, et 1,5%. La tolérance d'erreur γ_{tol} est d'autre part atteinte après 8 itérations.

2.3 Application au couplage transport/diffusion

Dans [Chamoin et Desvillettes 2013], la procédure d'estimation/adaptation est appliquée aux modèles de neutronique (ou transfert radiatif) qui sont utilisés dans l'industrie nucléaire (ou en astrophysique) et sont en pratique très coûteux à résoudre. Pour ce genre de modèles, il est cependant possible de montrer que sous certaines conditions, notamment lorsque le libre parcours moyen (i.e. la distance moyenne parcourue par une particule entre deux collisions successives) est petit par rapport à la taille du domaine d'étude, le modèle de neutronique peut être correctement approché par une équation de diffusion. L'approche multiéchelle consiste donc à coupler les deux modèles.

On considère ici l'équation linéaire de Boltzmann, ou de transport [Duderstadt et Martin 1979], qui décrit l'évolution d'une population de particules (neutrons ici) dans un domaine occupé par un milieu en interaction avec ces particules (noyaux d'uranium par exemple); elle rend donc compte des phénomènes physiques associés au transport de neutrons tels que les collisions, les absorptions ou les émissions. La grandeur solution de cette équation est la densité de neutrons dans l'espace des phases, notée $u(\mathbf{x}, \mathbf{n}, t)$, et définie comme le nombre probable de neutrons à la position $\mathbf{x}$ avec une direction de vitesse $\mathbf{n}$ à l'instant t (on considère que toutes

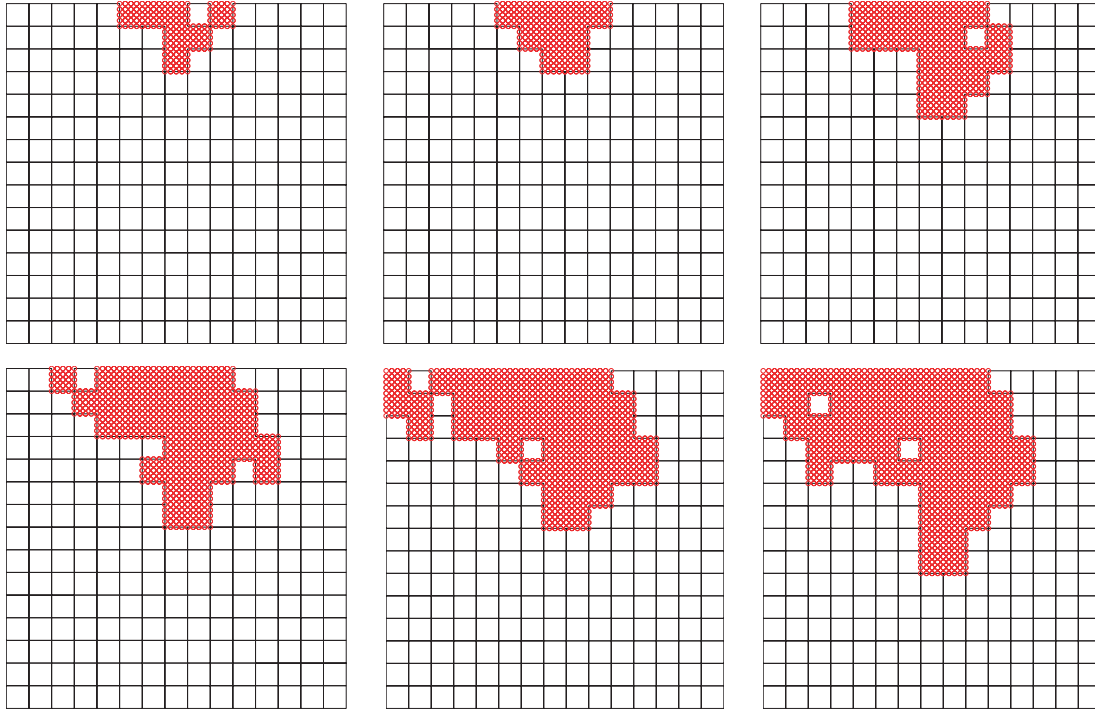


Figure 3.42: Séquence de configurations Arlequin obtenues aux itérations 1, 2, 3, 4, 5 et 6 (de gauche à droite et de haut en bas)

les particules ont la même vitesse v). A cause de sa grande dimensionnalité, la discrétisation de l'équation du transport dans les applications pratiques implique un nombre gigantesque de degrés de liberté, et une simulation précise du modèle de transport est un véritable challenge numérique.

On se restreint à l'écriture 1D de l'équation de transport, valable pour des domaines d'étude en forme de bande 2D ou 3D infinie, de largeur L le long de l'axe transversal $\mathbf{x}$ et invariants par une translation perpendiculaire à cet axe (voir [Larsen et Keller 1974]). En introduisant le paramètre classique de mise à l'échelle parabolique ε (nombre de Knudsen), défini comme le rapport entre le libre parcours moyen λ et une longueur caractéristique du problème, l'équation 1D vérifiée par $u_\varepsilon(x, \mu, t)$ s'écrit :

$$\varepsilon \frac{\partial u_\varepsilon}{\partial t} + \mu v \frac{\partial u_\varepsilon}{\partial x} + \frac{1}{\varepsilon} \Sigma v \mathcal{S} u_\varepsilon = 0 \quad \forall (x, \mu, t) \in]0, L[\times [-1, 1] \times [0, T] \quad (3.69)$$

avec μ le cosinus de l'angle entre l'axe $\mathbf{x}$ et la direction de vitesse $\mathbf{n}$, et Σ (resp. $\mathcal{S}$) la section efficace (resp. l'opérateur) de dispersion. Dans le cas du transport avec dispersion isotrope supposé ici, l'opérateur de dispersion $\mathcal{S}$ s'écrit :

$$\mathcal{S}u = u - \frac{1}{2} \int_{-1}^1 u(x, \mu^*, t) d\mu^* \quad (3.70)$$

L'équation (3.69) est associée aux conditions limites entrantes et initiales suivantes :

$$\begin{cases} u_\varepsilon(0, \mu, t) = u_0(\mu, t), & \forall \mu \geq 0 \\ u_\varepsilon(L, \mu, t) = u_L(\mu, t), & \forall \mu \leq 0 \end{cases} ; \quad u_\varepsilon(x, \mu, 0) = u_I(x) \quad (3.71)$$

Le problème peut alors être écrit sous forme faible :

$$\text{Trouver } u_\varepsilon \in L^\infty \text{ tel que } B(u_\varepsilon, w_\varepsilon) = F(w_\varepsilon) \quad \forall w_\varepsilon \in \mathcal{W} \quad (3.72)$$

avec

$$\begin{aligned} B(u_\varepsilon, w_\varepsilon) &= \int_0^T \int_0^L \int_{-1}^1 u_\varepsilon(x, \mu, t) \left\{ -\varepsilon \partial_t - \mu v \partial_x + \frac{1}{\varepsilon} \Sigma v \mathcal{S} \right\} w_\varepsilon(x, \mu, t) dx d\mu dt \\ F(w_\varepsilon) &= \varepsilon \int_{-1}^1 \int_0^L u_I(x) w_\varepsilon(x, \mu, 0) dx d\mu + \int_0^T \int_0^1 \mu w_\varepsilon(0, \mu, t) u_0(\mu, t) d\mu dt \\ &\quad - \int_0^T \int_{-1}^0 \mu w_\varepsilon(L, \mu, t) u_L(\mu, t) d\mu dt \\ \mathcal{W} &= \{w(x, \mu, t) \in \mathcal{D}([0, L] \times [-1, 1] \times [0, T]); \\ &\quad w(0, \mu, t) = 0 \quad \forall \mu \leq 0, w(L, \mu, t) = 0 \quad \forall \mu \geq 0\} \end{aligned}$$

$\mathcal{D}$ représente l'espace des fonctions C^∞ à support compact, ce qui implique que les fonctions $w \in \mathcal{W}$ sont telles que $w(x, \mu, T) = 0$.

Lorsque le libre parcours moyen des neutrons est petit devant L , i.e $\varepsilon \ll 1$, les phénomènes de dispersion dominent les effets d'écoulement et l'équation de transport (3.69) peut être approximée par une équation de diffusion macroscopique dont les coefficients sont dérivés à partir de ceux de l'équation de transport ; c'est l'approximation diffusive [Duderstadt et Martin 1979, Larsen et Keller 1974, Salvarini et Vasquez 2005, Dautray et Lions 1987], dans laquelle l'inconnue est la densité totale de neutrons dans l'espace physique $\rho(x, t)$. Elle s'écrit dans notre cas :

$$\frac{\partial \hat{u}}{\partial t} - \frac{\partial}{\partial x} \left[\frac{2}{3\Sigma(x)} \frac{\partial \hat{u}}{\partial x} \right] = 0 \quad (3.73)$$

Des conditions limites de Dirichlet (convergence d'ordre 1) ou de Robin (convergence d'ordre 2) sont associées à (3.73). On considère dans la suite ce second cas qui fait intervenir une longueur caractéristique d appelée longueur d'extrapolation [Degond et Mas-Gallic 1987] et reliée au libre parcours moyen λ par un coefficient Λ ($d = \Lambda\lambda$). Pour le modèle 1D considéré, la valeur de Λ s'écrit :

$$\Lambda = \frac{\sqrt{3}}{2} \int_0^1 \mu^2 H(\mu) d\mu \approx 0,7104 \quad (3.74)$$

où H est la *fonction de Chandrasekhar* [Chandrasekhar 1950].

Le modèle de diffusion est utilisé lorsqu'il est valide, i.e. dans les zones opaques, menant à une réduction de modèle et des calculs plus abordables. Le modèle approché résultant est donc fait de deux modèles concurrents couplés [Pomraning et Foglesong 1979, Tidriri 2001] : (i) un modèle à l'échelle mésoscopique gouverné par l'équation du transport (3.69) et appliqué dans les régions critiques où une solution fine est nécessaire, ou dans lesquelles un modèle macroscopique n'est pas valide ; (ii) un modèle macroscopique gouverné par l'équation de diffusion (3.73) dans le reste du

domaine.

Dans la suite, on considère que le domaine $X =]0, L[$ est divisé en deux zones $X_t =]0, a[$ et $X_d =]a, L[$ ($0 < a < L$), la théorie de diffusion donnant une approximation pauvre de la densité de neutrons dans X_t . Le problème couplé associé consiste alors à trouver (u_ε, u) , avec $u_\varepsilon(x, \mu, t)$ (resp. $u(x, t)$) défini dans X_t (resp. X_d), tel que :

$$\begin{aligned}
 \varepsilon \frac{\partial u_\varepsilon}{\partial t} + \mu v \frac{\partial u_\varepsilon}{\partial x} &= -\frac{1}{\varepsilon} \sigma \mathcal{S} u_\varepsilon \quad \text{dans } X_t \times [-1, 1] \times [0, T] \\
 u_\varepsilon(x, \mu, 0) &= u_I(x) \quad \forall (x, \mu) \in X_t \times [-1, 1] \\
 u_\varepsilon(0, \mu, t) &= u_0(\mu, t) \quad \forall \mu \geq 0 \\
 u_\varepsilon(a, \mu, t) &= u(a, t) \quad \forall \mu \leq 0 \\
 \hline
 \frac{\partial u}{\partial t} - \frac{\partial}{\partial x} \left(\frac{2}{3\Sigma(x)} \frac{\partial u}{\partial x} \right) &= 0 \quad \text{dans } X_d \times [0, T] \\
 u(x, 0) &= u_I(x) \quad \forall x \in X_d \\
 u(a, t) - \frac{\Lambda(a)\varepsilon}{\Sigma(a)} \frac{\partial u}{\partial x}(a, t) &= 2 \int_0^1 u_\varepsilon(a, \mu', t) \mu' d\mu' \\
 u(L, t) + \frac{\Lambda(L)\varepsilon}{\Sigma(L)} \frac{\partial u}{\partial x}(L, t) &= 2 \int_{-1}^0 u_L(\mu', t) |\mu'| d\mu'
 \end{aligned} \tag{3.75}$$

Ce problème couplé peut se réécrire sous la forme faible :

$$\text{Trouver } u_0 \in L^\infty \text{ tel que } B_0(u_0, w_0) = F(w_0) \quad \forall w_0 \in \mathcal{W} \tag{3.76}$$

avec $u_0(x, \mu, t) = u_\varepsilon(x, \mu, t)$ dans X_t et $u_0(x, \mu, t) = u(x, t)$ dans X_d .

On étudie une quantité d'intérêt $Q(u_\varepsilon)$, la fonctionnelle Q pouvant être écrite sous la forme générique :

$$Q(u_\varepsilon) = \int_0^T \int_{-1}^1 \int_0^L u_\varepsilon(x, \mu, t) g(x, \mu, t) dx d\mu dt \tag{3.77}$$

et on s'intéresse à l'erreur de réduction de modèle ainsi qu'à l'erreur de discrétisation sur Q . En notant u_0^h l'approximation numérique de u_0 , l'erreur totale sur Q peut s'écrire :

$$\Delta Q = Q(u_\varepsilon) - Q(u_0^h) = [Q(u_\varepsilon) - Q(u_0)] + [Q(u_0) - Q(u_0^h)] = \Delta Q_{red} + \Delta Q_{dis} \tag{3.78}$$

où ΔQ_{red} (resp. ΔQ_{dis}) est l'erreur sur Q due à la réduction de modèle (resp. due à la discrétisation).

Soit $\tilde{u}_0^h$ la solution approchée du problème adjoint obtenue par un couplage multi-échelle; en pratique, l'interface entre les deux modèles concurrents est placée en $x = \tilde{a} := a + \Delta a$ ($\Delta a > 0$) et la discrétisation spatio-temporelle est choisie plus fine que celle utilisée pour calculer u_0^h . Une estimation de l'erreur totale sur Q s'écrit alors :

$$\Delta Q \approx \mathcal{R}(u_0^h, \tilde{u}_0^h) \tag{3.79}$$

Pour estimer l'erreur de discrétisation seule, on peut considérer comme modèle de référence le modèle mathématique couplé (3.76). On obtient donc :

$$\Delta Q_{dis} \approx \mathcal{R}_0(u_0^h, \tilde{u}_0^h) := F(\tilde{u}_0^h) - B_0(u_0^h, \tilde{u}_0^h) \quad (3.80)$$

Une estimation de l'erreur locale de réduction de modèle peut ensuite être dérivée par différence en utilisant (3.78).

On illustre la méthodologie sur un modèle de transport 1D simplifié, appelé modèle de Carleman linéarisé; d'autres modèles sont considérés dans [Chamoin et Desvilletes 2013]. Le modèle de Carleman linéarisé, dans lequel $\mu = \pm 1$, simule la collision entre deux populations de neutrons évoluant avec des directions inverses. Les densités dans l'espace des phases associées sont notées $u^+(x, t)$ et $u^-(x, t)$, respectivement. Le problème mis à l'échelle, en considérant une vitesse constante $v = 1$ et sur le domaine spatio-temporel $]0, L[\times]0, T]$ avec $L = 5$ et $T = 4$, consiste à trouver $u_\varepsilon^+(x, t)$ et $u_\varepsilon^-(x, t)$ tels que :

$$\begin{aligned} \varepsilon \frac{\partial u_\varepsilon^+}{\partial t} + \frac{\partial u_\varepsilon^+}{\partial x} &= \frac{\Sigma(x)}{2\varepsilon} (u_\varepsilon^- - u_\varepsilon^+) \quad ; \quad \varepsilon \frac{\partial u_\varepsilon^-}{\partial t} - \frac{\partial u_\varepsilon^-}{\partial x} = \frac{\Sigma(x)}{2\varepsilon} (u_\varepsilon^+ - u_\varepsilon^-) \\ u_\varepsilon^+(0, t) &= u_0(t) \quad ; \quad u_\varepsilon^-(L, t) = u_L(t) \quad ; \quad u_\varepsilon^\pm(x, 0) = u_I(x) \end{aligned} \quad (3.81)$$

On prend des conditions initiales nulles ($u_I(x) = 0$) et on prescrit pour tout $t \in [0, T]$ les conditions limites entrantes suivantes :

$$u_0(\mu, t) = 1 \quad \forall \mu \geq 0 \quad ; \quad u_L(\mu, t) = 0 \quad \forall \mu \leq 0 \quad (3.82)$$

Le paramètre d'échelle est fixé à $\varepsilon = 5.10^{-2}$, et on considère une évolution linéaire par morceaux pour $\Sigma(x)$:

- $\Sigma(x) = \varepsilon$ pour $x \in [0, 2]$;
- $\Sigma(x)$ linéaire de ε à 1 pour $x \in [2, 3]$;
- $\Sigma(x) = 1$ pour $x \in [3, 5]$.

De ce fait, $\Sigma/\varepsilon = 1$ dans la partie transparente du domaine, et on s'attend à ce que le modèle de diffusion soit valable dans la partie droite du domaine.

Les deux quantités d'intérêt considérées sont la densité et le flux en un point $(x_Q, t_Q) \in]0, L[\times]0, T]$, s'écrivant respectivement :

$$Q_1(u_\varepsilon) = \int_{-1}^1 u_\varepsilon(x_Q, \mu, t_Q) d\mu \quad ; \quad Q_2(u_\varepsilon) = \frac{1}{\varepsilon} \int_{-1}^1 \mu u_\varepsilon(x_Q, \mu, t_Q) d\mu \quad (3.83)$$

avec $g_1(x, \mu, t) = \delta_{x_Q}(x) \delta_{t_Q}(t)$ et $g_2(x, \mu, t) = \frac{\mu}{\varepsilon} \delta_{x_Q}(x) \delta_{t_Q}(t)$ les extracteurs associés. Dans la suite, on prend $x_Q = 1, 5$ et $t_Q = 3, 8$.

On représente sur la Figure 3.43 la solution quasi-exacte du problème (3.81), calculée avec une discrétisation très fine. Les valeurs associées (prises comme référence) des quantités d'intérêt sont $Q_1 = 1, 8398$ et $Q_2 = 1, 2814$.

La limite diffusive pour le modèle (3.81) s'écrit :

$$\frac{\partial u}{\partial t} - \frac{\partial}{\partial x} \left[\frac{1}{\Sigma(x)} \frac{\partial u}{\partial x} \right] = 0 \quad (3.84)$$

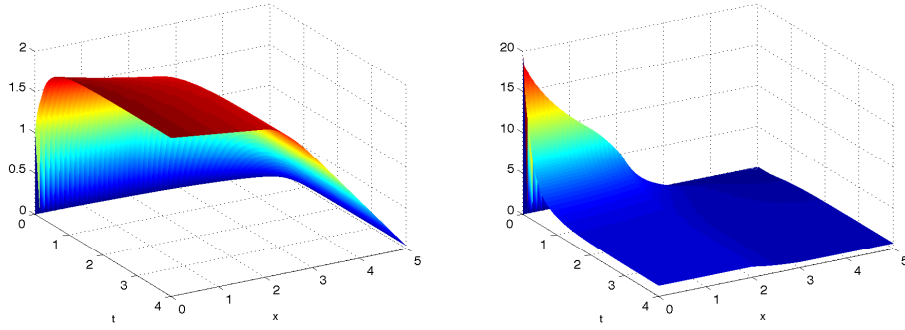


Figure 3.43: Solution quasi-exacte du modèle de Carleman linéarisé dans l'espace spatio-temporel : densité ρ_ε (gauche) et flux j_ε (droite)

et la longueur d'extrapolation associée est $\Lambda = 1$. La résolution numérique du modèle couplé transport-diffusion est menée en utilisant un schéma de différences finies dans X_t et la MEF avec un schéma temporel d'Euler explicite dans X_d . Les pas d'espace et de temps choisis sont $\Delta x = 0,1$ et $\Delta t = 10^{-5}$; ils permettent de respecter les conditions CFL ($\Delta t \leq \varepsilon \Delta x$) et de stabilité. Pour le problème adjoint, on prend $\Delta \tilde{x} = 0,05$ et $\Delta \tilde{t} = 2 \cdot 10^{-6}$. On montre sur la Figure 3.44 la solution approchée obtenue avec trois positions différentes de l'interface i.e. $a = 1,0$, $a = 2,0$ et $a = 3,0$; les valeurs approchées associées des quantités d'intérêt Q_1 et Q_2 sont données dans le Tableau 3.4. On observe clairement que la solution est imprécise quand l'interface est placée dans ou près de la zone transparente.

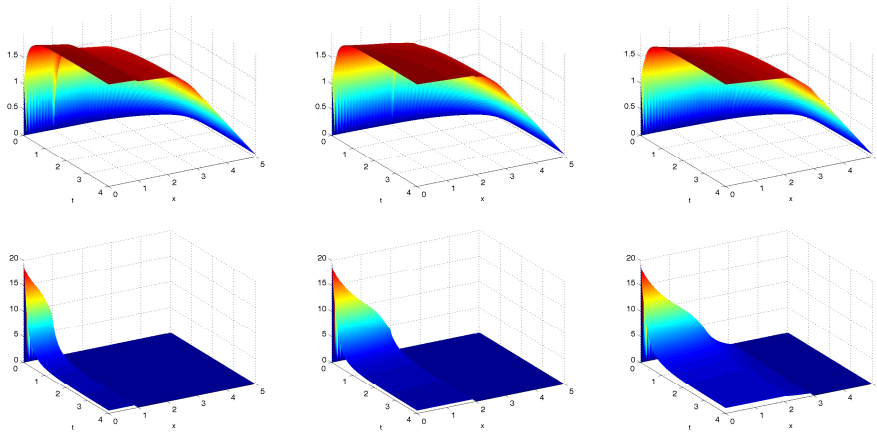
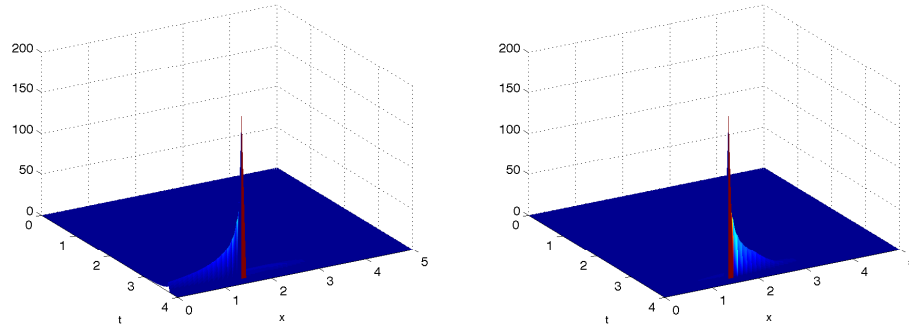


Figure 3.44: Solution approchée du modèle de Carleman linéarisé dans l'espace spatio-temporel, pour $a = 1,0$ (gauche), $a = 2,0$ (centre) et $a = 3,0$ (droite) : densité ρ_0 (haut), flux j_0 (bas)

On utilise maintenant la stratégie d'estimation d'erreur pour adapter le modèle couplé jusqu'à un niveau d'erreur donné. On montre sur la Figure 3.45 (resp. Figure 3.46) les solutions approchées des problèmes adjoints associés à Q_1 (resp. Q_2), pour une interface située en $\tilde{a} = 2,5$. On observe que ces solutions sont très

a	Q_1	erreur sur Q_1 (%)	Q_2	erreur sur Q_2 (%)
1,0	1,8967	3,09	0,0000	100,00
2,0	1,9295	4,88	0,5621	56,13
3,0	1,8693	1,60	1,0435	18,57

Tableau 3.4: Valeurs de Q_1 et Q_2 en fonction de la position a de l'interfaceFigure 3.45: Solution adjointe approchée, dans le domaine spatio-temporel, pour Q_1 : $\tilde{u}_0^-$ (gauche), $\tilde{u}_0^+$ (droite)

localisées en espace et en temps, et ont de forts gradients proche de la zone de définition des quantités d'intérêt.

Avant d'estimer l'erreur de réduction de modèle, on vérifie que l'erreur de discrétisation reste négligeable. Dans le cas $a = 2,0$, on calcule pour Q_1 l'estimateur d'erreur de discrétisation. La représentation de cet estimateur dans le domaine spatio-temporel est montré sur la Figure 3.47 ; on observe que l'erreur de discrétisation est concentrée proche du point (x_Q, t_Q) et s'étend légèrement autour de ce point. La valeur relative de l'estimation d'erreur de discrétisation pour Q_1 est 0,61%, ce qui est négligeable par rapport aux valeurs d'erreur de réduction de modèle mises en évidence dans la suite.

En considérant toujours $a = 2,0$, on calcule maintenant l'estimateur d'erreur totale associé à Q_1 pour plusieurs positions $\tilde{a}$ de l'interface dans le problème adjoint couplé. Les résultats sont reportés dans le Tableau 3.5. On observe que l'estimateur

$\tilde{a}$	2,0	2,2	2,4	2,6	2,8	3,0
estimateur (%)	1,27	4,29	4,87	5,03	5,03	5,03

Tableau 3.5: Valeurs de l'estimateur d'erreur pour Q_1 , en fonction de la position $\tilde{a}$ de l'interface dans le problème adjoint

converge très rapidement, et qu'étendre le domaine d'application du modèle de transport, pour le problème couplé adjoint, au-delà de $\tilde{a} = 2,6$ ne modifie pas la

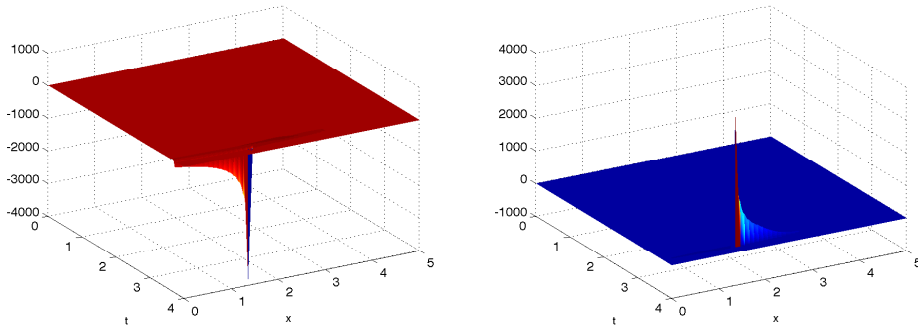


Figure 3.46: Solution adjointe approchée, dans le domaine spatio-temporel, pour Q_2 : $\tilde{u}_0^-$ (gauche), $\tilde{u}_0^+$ (droite)

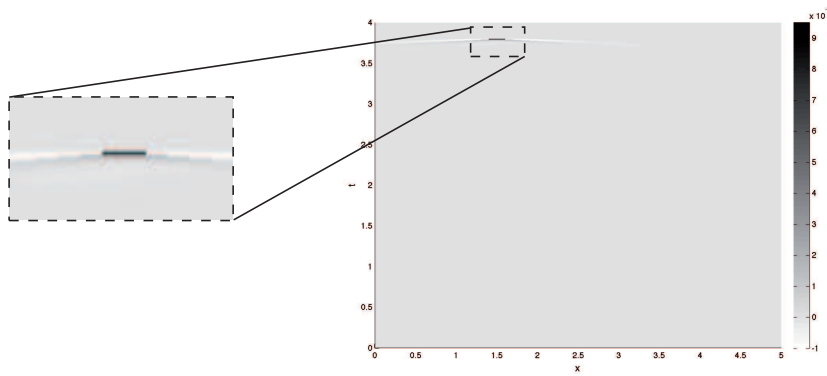


Figure 3.47: Contributions en espace et temps de l'erreur de discrétisation pour Q_1

valeur de l'estimateur. Dans la suite, on prend $\tilde{a} = a + 0,5$.

On met alors en place l'algorithme d'adaptation de modèle pour le contrôle de Q_1 et Q_2 . On commence dans les deux cas par une position initiale $a = 1,6$ de l'interface, et on spécifie la tolérance $\gamma_{\text{tol}} = 3\%$ et le paramètre de raffinement $\gamma_a = 0,2$. On reporte dans le Tableau 3.6 les valeurs estimées de Q_1 et Q_2 au cours des itérations de l'algorithme adaptatif. On observe que la méthode générale de contrôle d'erreur locale est efficace pour les modèles de transport ; elle sera prochainement appliquée à des cas 3D plus réalistes rencontrés dans l'ingénierie nucléaire.

2.4 Application au couplage stochastique/déterministe

Les travaux de thèse de C. Zaccardi [Zaccardi *et al.* 2013] se sont focalisés sur la réduction des modèles continus stochastiques classiquement utilisés pour prendre en compte explicitement les incertitudes sur les données et les variabilités de modèle. L'approche multiéchelle qui a été suivie consiste à coupler le modèle stochastique avec un modèle déterministe (obtenu par homogénéisation stochastique). Le contrôle

itération	a	Q_1	estimateur (%)	itération	a	Q_2	estimateur (%)
0	1,6	1,9767	7,69	0	1,6	0,2734	80,81
1	1,8	1,9426	5,81	1	2,0	0,5621	59,11
2	2,0	1,9295	4,99	2	2,4	0,7311	43,29
3	2,1	1,8877	2,73	3	2,9	1,0396	19,86
				4	3,2	1,1794	8,22
				5	3,5	1,2319	4,19
				6	3,6	1,2492	2,91

Tableau 3.6: Evolution du modèle couplé et estimation d'erreur au cours des itérations pour Q_1 (gauche) et Q_2 (droite)

de la procédure de couplage, issue de [Cottureau *et al.* 2011] et basée sur le cadre Arlequin, a été menée avec les outils exposés dans la Section 2.1. De plus, les diverses sources d'erreur mises en jeu (réduction de modèle, discrétisation spatiale, approximation de Monte-Carlo) ont été estimées pour mener une stratégie adaptative efficace.

2.4.1 Estimation de l'erreur totale

Le monomodèle de référence est un problème stochastique elliptique scalaire défini dans un milieu continu. Plus spécifiquement, on considère le problème de Poisson stochastique représentant une large gamme d'applications physiques (équations de la chaleur, de membrane, de Darcy, ...). Le problème est défini sur un domaine ouvert borné Ω de frontière $\partial\Omega$, séparé en Γ_D et Γ_N tel que $\partial\Omega = \overline{\Gamma_D} \cup \overline{\Gamma_N}$. En considérant l'espace probabiliste $(\Theta, \mathcal{F}, P)$, le problème consiste à trouver $u(x, \theta)$ tel que (p.p. : presque partout, p.s. : presque sûrement) :

$$\nabla \cdot (K(x, \theta) \nabla u(x, \theta)) + f(x) = 0 \quad \text{p.p. dans } \Omega, \text{ et p.s.} \quad (3.85)$$

avec

$$\begin{cases} u(x, \theta) = 0 & \text{p.p. sur } \Gamma_D \text{ et p.s.} \\ K(x, \theta) \nabla u \cdot \mathbf{n} = g(x) & \text{p.p. sur } \Gamma_N \text{ et p.s.} \end{cases} \quad (3.86)$$

Le paramètre de comportement $K(x, \theta)$ correspond à un champ stochastique (Figure 3.48) borné et uniformément coercif [Babuška *et al.* 2004], i.e. vérifiant :

$$0 < K_{min} \leq K(x, \theta) \leq K_{max} < \infty, \quad \text{p.p. dans } \Omega \text{ et p.s.}$$

On note $\mathcal{U}$ l'espace fonctionnel des champs cinématiquement admissibles :

$$\mathcal{W} = \{v \in \mathcal{L}^2(\Theta, \mathcal{H}^1(\Omega)), v = 0 \text{ sur } \Gamma_D, \text{ p.s.}\} \quad (3.87)$$

La formulation faible du problème (3.85)-(3.86) s'écrit alors :

$$\text{Trouver } u \in \mathcal{U} \text{ tel que } a(u, v) = l(v) \quad \forall v \in \mathcal{U} \quad (3.88)$$

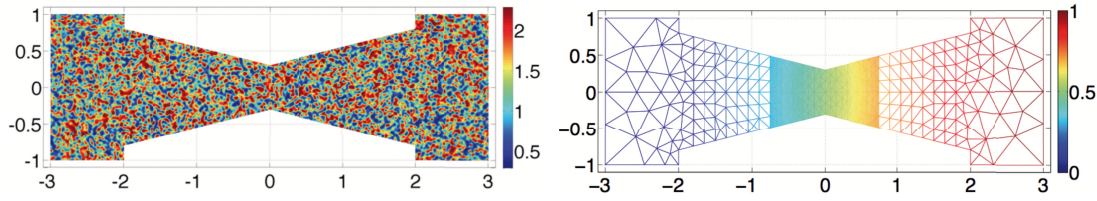


Figure 3.48: Exemple d'une réalisation du champ $K(x, \theta)$ (gauche) et solution en déplacement $(u_d, E_\Theta[u_s])$ (droite)

avec

$$\begin{aligned} a(u, v) &= E_\Theta \left[\int_{\Omega} K(x, \theta) \nabla u(x, \theta) \cdot \nabla v(x, \theta) d\Omega \right] \\ l(v) &= \int_{\Omega} f(x) E_\Theta [v(x, \theta)] d\Omega + \int_{\Gamma_N} g(x) E_\Theta [v(x, \theta)] d\Gamma \end{aligned} \quad (3.89)$$

$E_\Theta [\bullet]$ est l'espérance mathématique.

La résolution de (3.88) par les approches classiques (discrétisation des dimensions spatiale et stochastique) induit généralement un coût de calcul prohibitif. Néanmoins, on suppose que l'objectif de la simulation est l'évaluation d'une quantité d'intérêt définie dans une région locale $\Omega_s \subset \Omega$. Pour reproduire fidèlement les effets stochastiques sur cette quantité, on décide de garder la description stochastique dans la zone Ω_s seule. Le champ stochastique $K(x, \theta)$ est alors changé en un paramètre déterministe noté K_d , obtenu par homogénéisation stochastique [Bourgeat et Piatnitski 2004], dans le reste de la structure. Les modèles stochastique et déterministe, définis respectivement dans les régions Ω_s et Ω_d , sont couplés par la méthode Arlequin ; les deux modèles sont superposés dans une zone de couplage Ω_o et l'énergie est distribuée entre les deux modèles avec des fonctions poids $(\alpha_s(x), \alpha_d(x))$ [Ben Dhia 1998 - 2008, Cottureau et al. 2011]. Le problème Arlequin consiste alors à trouver (u_d, u_s, λ) dans $\mathcal{U}_d \times \mathcal{U}_s \times \mathcal{X}$ tel que :

$$\begin{cases} a_d(u_d, v_d) + C(\lambda, v_d) &= l_d(v_d) & \forall v_d \in \mathcal{U}_d \\ a_s(u_s, v_s) - C(\lambda, v_s) &= l_s(v_s) & \forall v_s \in \mathcal{U}_s \\ C(\mu, u_d - u_s) &= 0 & \forall \mu \in \mathcal{X} \end{cases} \quad (3.90)$$

avec $\mathcal{U}_d = \{v \in \mathcal{H}^1(\Omega_d), v(x) = 0 \text{ sur } \Gamma_D\}$, $\mathcal{U}_s = \mathcal{L}^2(\Theta, \mathcal{H}^1(\Omega_s))$ et les définitions suivantes pour les travaux interne et externe :

$$\begin{aligned} a_d(u, v) &= \int_{\Omega_d} \alpha_d(x) K_d \nabla u(x) \cdot \nabla v(x) d\Omega \\ a_s(u, v) &= E_\Theta \left[\int_{\Omega_s} \alpha_s(x) K(x, \theta) \nabla u(x, \theta) \cdot \nabla v(x, \theta) d\Omega \right] \\ l_d(v) &= \int_{\Omega_d} \alpha_d(x) f(x) v(x) d\Omega + \int_{\Gamma_N} g(x) v(x) d\Gamma \\ l_s(v) &= E_\Theta \left[\int_{\Omega_s} \alpha_s(x) f(x) v(x, \theta) d\Omega \right] \end{aligned} \quad (3.91)$$

L'espace des multiplicateurs de Lagrange (ou espace médiateur) $\mathcal{W}_c$ est construit sous la forme [Cottureau *et al.* 2011] :

$$\mathcal{W}_c = \{\psi + \theta_c \mathbb{I}_{\Omega_c} \mid \psi \in \mathcal{H}^1(\Omega_c), \theta_c \in \mathcal{L}^2(\Theta, \mathbb{R})\} \quad (3.92)$$

avec θ_c une variable aléatoire et $\mathbb{I}_{\Omega_c}$ la fonction caractéristique de Ω_c . Cet espace, qui est le plus petit assurant la stabilité de la formulation Arlequin (3.90), impose implicitement que le champ $E_\Theta[u_s]$ soit égal au champ u_d , presque partout dans Ω_c , et que la variabilité de la quantité moyenne spatiale $\int_{\Omega_c} (E_\Theta[u_s] - u_s) d\Omega$ s'annule. L'opérateur de couplage C est quant à lui défini par :

$$C(u, v) = E_\Theta \left[\int_{\Omega_c} (\kappa_0 uv + \kappa_1 \nabla u \cdot \nabla v) d\Omega \right] \quad (3.93)$$

où κ_0 et κ_1 sont deux paramètres d'échelle.

A partir de (u_d, u_s) , on dérive le champ continu $u_0 \in \mathcal{U}$:

$$u_0 = \begin{cases} u_d & \text{dans } \Omega_d \setminus \Omega_s \\ \alpha_d u_d + \alpha_s u_s & \text{dans } \Omega_s \end{cases} \quad (3.94)$$

Cette solution Arlequin u_0 est une approximation de la solution de référence $u \in \mathcal{U}$. Son évaluation par des techniques numériques introduit deux niveaux supplémentaires d'approximation : une approximation en espace reliée à l'introduction d'un maillage EF, et une autre liée à la dimension stochastique avec l'utilisation de la technique de Monte-Carlo. On note (u_d^a, u_s^a) la solution numérique associée et $u_0^a \in \mathcal{U}$ le champ reconstruit avec (3.94). On montre sur la Figure 3.49, pour le problème 2D précédemment décrit (Figure 3.48), une comparaison entre les solutions exacte et Arlequin le long de l'axe de symétrie $\mathbf{x}$. Dans cette figure : le champ $u_d(x)$ est représenté en ligne continue noire ; les moyennes de u_s et $|\nabla u_s|$ sont représentées en tirets noirs ; un exemple de réalisation est représenté par la ligne continue rouge ; l'intervalle de confiance à 90% sur u_s ou $|\nabla u_s|$ est représenté par la zone jaune ; la solution de référence est représentée en bleu (espérance par une ligne continue et intervalle de confiance à 90% par des tirets).

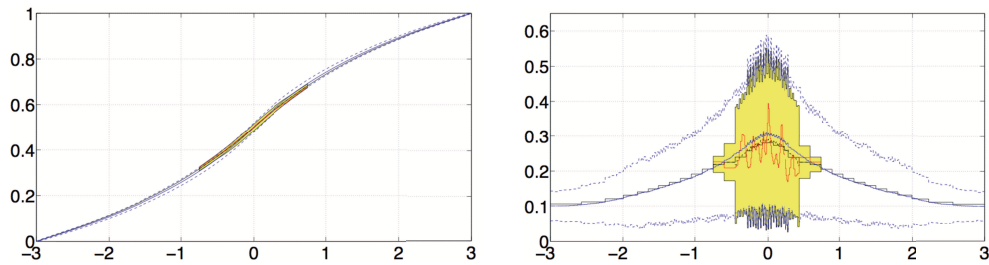


Figure 3.49: Solutions en déplacement (gauche) et en gradient de déplacement (droite) le long de l'axe horizontal

La quantité d'intérêt Q considérée est mise sous la forme globale :

$$Q(u) = E_\Theta \left[\int_{\Omega} u f_{\Sigma} d\Omega + \int_{\Gamma_N} u g_{\Sigma} d\Gamma + \int_{\Omega} \nabla u \cdot \mathbf{p}_{\Sigma} d\Omega \right] \quad (3.95)$$

et on cherche à quantifier l'erreur $Q(u) - Q(u_0^a)$. On introduit alors le problème adjoint correspondant (voir (3.51)), de solution $\tilde{u} \in \mathcal{U}$. Cette solution n'étant pas abordable en pratique, on l'approxime avec un problème couplé Arlequin similaire à (3.90), avec une région Ω_s plus large et des espaces discrétisés plus riches que ceux utilisés pour u_0^a ; la solution obtenue est notée $(\tilde{u}_d^a, \tilde{u}_s^a)$ et on en reconstruit une solution $\tilde{u}_0^a \in \mathcal{U}$ en suivant (3.94) et avec une interpolation dans l'espace stochastique. L'erreur totale sur Q est donc estimée par :

$$\Delta Q = Q(u) - Q(u_0^a) \approx \mathcal{R}(u_0^a, \tilde{u}_0^a) \quad (3.96)$$

avec $\mathcal{R}$ la fonctionnelle des résidus associée à (3.88).

Une difficulté dans le calcul de l'estimateur (3.96), qui implique le produit $K \nabla u_0^a \cdot \nabla \tilde{u}_0^a$, réside dans le fait que les solutions u_0^a et $\tilde{u}_0^a$ sont obtenues pour des tirages de Monte-Carlo différents. Une approche a donc été mise en place dans [Zacardi et al. 2013] pour générer, à partir des réalisations utilisées dans la résolution de (3.90), de nouvelles réalisations de $u_s^a(x, \theta)$ qui correspondent à celles de $K(x, \theta)$ utilisées pour la résolution du problème adjoint, et qui conservent les statistiques (1er et 2ème ordre) du champ u_s^a à disposition.

On suppose pour cela que $K(x, \theta)$ et $u_s^a(x, \theta)$ sont des champs aléatoires gaussiens. On note Ξ et U les vecteurs aléatoires associés à $K(x, \theta)$ et $u_s^a(x, \theta)$ après discrétisation en espace. On génère alors de nouvelles réalisations $\hat{U}$ de U telles que :

$$\begin{cases} \mathbb{E}_\Theta[\hat{U}] = \mathbb{E}_\Theta[U] := \bar{U} \\ \mathbb{E}_\Theta[(\hat{U} - \bar{U})(\hat{U} - \bar{U})^T] = \mathbb{E}_\Theta[(U - \bar{U})(U - \bar{U})^T] := \text{Cov}_U \\ \mathbb{E}_\Theta[(\hat{U} - \bar{U})(\Xi - \bar{\Xi})^T] = \mathbb{E}_\Theta[(U - \bar{U})(\Xi - \bar{\Xi})^T] := \text{Cov}_{U\Xi} \end{cases} \quad (3.97)$$

où $\bar{\bullet}$ et $\text{Cov}_\bullet$ désignent l'espérance et la matrice de covariance (ou d'auto-covariance). Le modèle stochastique de K étant donné, la matrice Cov_Ξ est connue. Toutes les autres espérances et matrices de covariance sont estimées avec des estimateurs statistiques standards et les réalisations originelles de la solution u_0^a . En introduisant un vecteur aléatoire unitaire centré gaussien Υ avec des composantes décorrélées ($\mathbb{E}_\Theta[v_i v_j] = \delta_{ij}$) et indépendantes de Ξ , on montre que le vecteur aléatoire :

$$\hat{U} = \bar{U} + \text{Cov}_{U\Xi} \text{Cov}_\Xi^{-1} (\Xi - \bar{\Xi}) + \sqrt{\text{Cov}_U - \text{Cov}_{U\Xi} \text{Cov}_\Xi^{-1} \text{Cov}_{U\Xi}^T} \Upsilon \quad (3.98)$$

vérifie les statistiques (1er et 2ème ordre) de (3.97).

En pratique, ni $K(x, \theta)$ ni $u_s^a(x, \theta)$ ne sont des champs aléatoires gaussiens; une approche isoprobabiliste est alors suivie pour transformer ces champs en champs gaussiens. Après construction du vecteur gaussien $\hat{U}$, l'approche inverse peut être utilisée pour transformer les nouvelles réalisations de U en la distribution marginale initiale du 1er ordre.

2.4.2 Séparation des sources d'erreur

L'erreur générée entre $Q(u)$ et $Q(u_0^a)$ peut avoir plusieurs sources. Si on décompose la construction de la solution approchée u_0^a en différentes étapes, on peut distinguer (voir l'illustration 1D de la Figure 3.50) :

- l'utilisation d'un modèle réduit généré par un couplage Arlequin ;
- la discrétisation spatiale des modèles continus ;
- l'approximation des moments statistiques (avec la technique de Monte-Carlo par exemple).

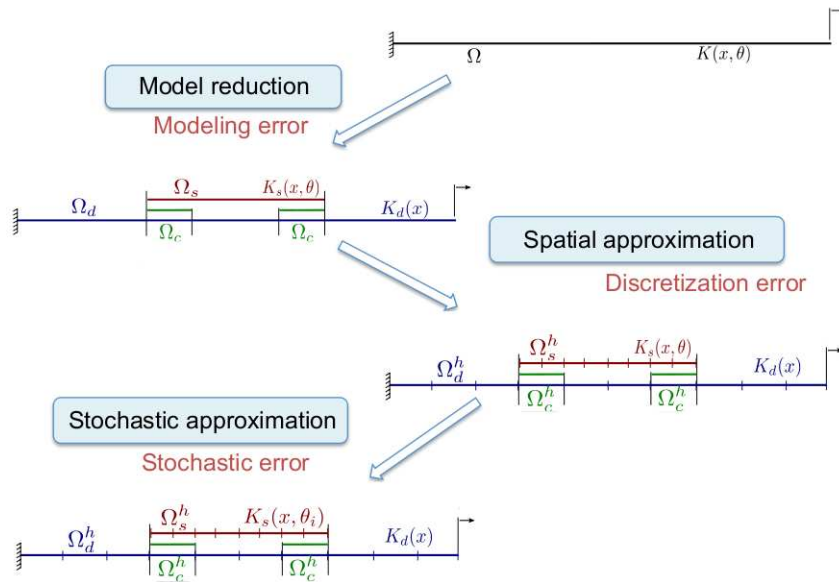


Figure 3.50: Etapes d'approximation pour la résolution du problème stochastique

Chaque étape est associée à des paramètres propres : le modèle Arlequin est décrit par la taille du patch dans lequel le modèle fin est conservé, et celle de la zone de couplage ; la discrétisation spatiale est paramétrée par la taille d'élément sur chaque domaine (Ω_d , Ω_s et Ω_o) ; enfin, si on utilise la technique de Monte-Carlo dans la dimension stochastique, la troisième étape est fonction du nombre de tirages utilisés pour décrire les propriétés matérielles. On propose une méthodologie pour estimer l'erreur induite par chaque étape.

Pour séparer les sources d'erreur, on considère une cascade de problèmes de référence intermédiaires. Outre le problème initial (3.88) (solution u , résidu $\mathcal{R}$), le problème Arlequin continu (3.90) (solution u_0 , résidu $\mathcal{R}_0$) et le problème Arlequin totalement discrétisé finalement résolu (solution u_0^h), on introduit le problème Arlequin couplant un modèle discrétisé déterministe et un modèle stochastique discrétisé spatialement uniquement. On note u_0^h (resp. $\mathcal{R}_0^h$) la solution (resp. le résidu) associée à ce dernier problème. En notant que :

$$\begin{aligned} \Delta Q &= Q(u) - Q(u_0^a) = [Q(u) - Q(u_0)] + [Q(u_0) - Q(u_0^h)] + [Q(u_0^h) - Q(u_0^a)] \\ &= \Delta Q_m + \Delta Q_h + \Delta Q_\theta \end{aligned} \quad (3.99)$$

avec ΔQ_m , ΔQ_h et ΔQ_θ désignant respectivement les erreurs de réduction de modèle, de discrétisation spatiale et de discrétisation stochastique, on dérive les indicateurs

d'erreur suivants :

$$\begin{aligned}\Delta Q_\theta &\approx \eta_\theta = \mathcal{R}_0^h(u_0^a, \tilde{u}_0^a) \\ \Delta Q_h &\approx \eta_h = \mathcal{R}_0(u_0^a, \tilde{u}_0^a) - \mathcal{R}_0^h(u_0^a, \tilde{u}_0^a) \\ \Delta Q_m &\approx \eta_m = \mathcal{R}(u_0^a, \tilde{u}_0^a) - \mathcal{R}_0(u_0^a, \tilde{u}_0^a)\end{aligned}\tag{3.100}$$

A noter que ces indicateurs sont calculés simplement par post-traitement des solutions u_0^a et $\tilde{u}_0^a$ à disposition, ce qui limite les coûts de calcul additionnels.

Le processus d'adaptation, basé sur un algorithme glouton, consiste donc à estimer dans un premier temps l'erreur totale ΔQ . Si cette erreur est plus grande que la tolérance fixée γ_{tol} , les contributions venant des différentes sources sont mesurées avec (3.100) puis la source d'erreur dominante est identifiée (en module) et le paramètre correspondant est modifié.

On applique la stratégie à l'exemple 2D considéré dans cette section (Figure 3.48). La structure, inscrite dans la boîte $[-3, 3] \times [-1, 1]$, est encastrée à sa gauche et soumise à un déplacement de composante horizontale imposée et égale à 1 sur son bord droit ; les bords latéraux sont libres (presque sûrement). Le matériau a des propriétés de comportement aléatoires modélisées par le champ $K(x, \theta)$ de bornes 0,3542 et 2,1938, avec une moyenne géométrique $1/E_\Theta [1/K(x, \theta)]$ égale à 1,0, un écart-type σ_K égal à 0,2 et une corrélation exponentielle avec une longueur de corrélation $L_{\text{cor}} = 0,05$ dans chaque direction. L'estimation numérique du résidu utilise une discrétisation spatiale avec 33792 éléments T3 ainsi que $MC = 5000$ tirages de Monte-Carlo pour représenter le champ stochastique K . La méthode Arlequin est quant à elle utilisée avec un patch (région Ω_s) centré.

La quantité d'intérêt considérée est la composante sur $\mathbf{x}$ du gradient de u dans une zone locale critique $\Omega_Q \subset \Omega$ (voir Figure 3.51) :

$$Q(u) = E_\Theta \left[\int_{\Omega_Q} \nabla u \cdot \mathbf{i}_x d\Omega \right]\tag{3.101}$$

Le problème adjoint correspondant, dont le chargement consiste en une précharge $\mathbf{p}_\Sigma = \mathbf{x}$ dans Ω_Q , est donné sur la Figure 3.51. Il est discrétisé spatialement avec 8448 éléments T3 pour le modèle déterministe et 6144 éléments T3 pour le modèle stochastique dans un patch de demi-longueur $L_s = 1,8$. De plus, on utilise 5000 tirages de Monte-Carlo.

La Figure 3.52 montre les valeurs relatives des divers indicateurs d'erreur pour les quatre premières itérations de l'algorithme itératif. On représente les indicateurs pour l'erreur totale (η_r , en blanc), l'erreur de réduction de modèle (η_m , en gris clair), l'erreur de discrétisation spatiale (η_h , en gris foncé) et l'erreur de discrétisation stochastique (η_θ , en noir). On remarque que les sources d'erreur principales sont dues à l'utilisation d'un nombre faible de tirages de Monte-Carlo et à l'utilisation de la méthode Arlequin. En adaptant les paramètres correspondants, l'estimateur global d'erreur relatif η_r diminue de 21% à 5%.

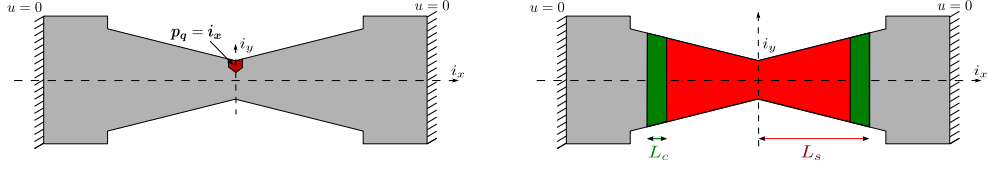


Figure 3.51: Définition du problème adjoint : chargement dans la zone locale Ω_Q (gauche) et couplage Arlequin (droite)

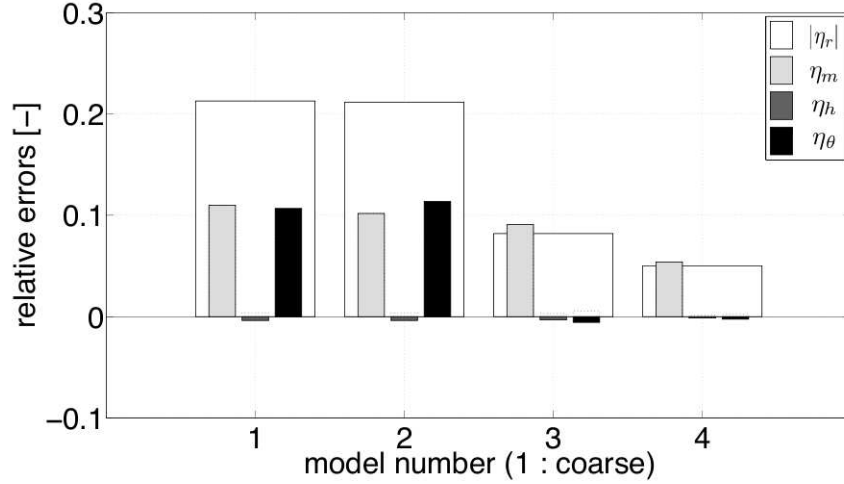


Figure 3.52: Evolution des divers indicateurs relatifs d'erreur durant le processus adaptatif

2.5 Procédure optimale d'adaptation

Une nouvelle approche d'adaptation de modèle, alternative à celle exposée dans la Section 2.1.2, a été introduite et étudiée dans [Ben Dhia *et al.* 2011]. Elle est basée sur les idées générales du contrôle optimal (voir Section 3) et cherche à trouver, par paramétrage de la zone d'application du modèle fin, la configuration de couplage optimale pour l'estimation précise d'une quantité d'intérêt; elle est ainsi dans la lignée des techniques d'optimisation de forme [Allaire 2002]. Les études numériques réalisées montrent que la nouvelle approche d'adaptation a une convergence meilleure que celle de l'approche classique basée sur un post-traitement de l'estimateur de l'erreur de réduction de modèle calculé *a posteriori*.

On expose cette nouvelle approche en reprenant le cadre de couplage déterministe particulaire/continu défini dans la Section 2.2. On suppose que la taille et la forme de la région pleinement particulaire $\Omega_d \setminus \Omega_o$ sont définies au moyen de paramètres de contrôle; ces paramètres peuvent être introduits de différentes façons, et on montre quelques exemples sur la Figure 3.53. Dans la suite, les paramètres de contrôle sont regroupés dans un vecteur $\mathbf{m}$ et la solution Arlequin obtenue avec ces paramètres est notée $\mathbf{z}(\mathbf{m})$. Pour un ensemble donné de paramètres $\mathbf{m} \in \mathcal{T}$, le problème Arlequin (3.15) consiste à trouver $(\mathbf{u}(\mathbf{m}), \mathbf{z}(\mathbf{m}), \boldsymbol{\lambda}(\mathbf{m})) \in \mathcal{U} \times \mathcal{W} \times \mathcal{X}$ tel que :

$$a_{0m}((\mathbf{u}(\mathbf{m}), \mathbf{z}(\mathbf{m}), \boldsymbol{\lambda}(\mathbf{m})); (\mathbf{v}, \mathbf{w}, \boldsymbol{\mu})) = l(\mathbf{w}) \quad \forall (\mathbf{v}, \mathbf{w}, \boldsymbol{\mu}) \in \mathcal{U} \times \mathcal{W} \times \mathcal{X} \quad (3.102)$$

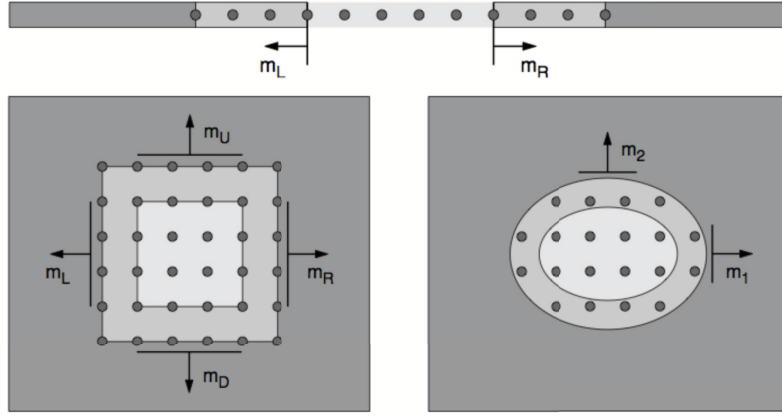


Figure 3.53: Exemples de paramètres utilisés pour définir la région particulière Ω_d : pour une chaîne 1D (haut), pour une structure 2D (bas)

où on a indiqué la dépendance explicite aux paramètres $\mathbf{m}$ de la solution $(\mathbf{u}, \mathbf{z}, \boldsymbol{\lambda})$ et de la forme semi-linéaire a_0 . Par simplicité de notation, ce problème couplé est réécrit sous la forme compacte :

$$\text{Trouver } \mathbf{z}_0(\mathbf{m}) \in \mathcal{V} \text{ tel que } a_{0\mathbf{m}}(\mathbf{z}_0(\mathbf{m}); \mathbf{w}) = l(\mathbf{w}) \quad \forall \mathbf{w} \in \mathcal{V} \quad (3.103)$$

Du point de vue du contrôle optimal, le processus adaptatif consiste à trouver l'ensemble $\mathbf{m}^*$ de paramètres qui minimise l'erreur sur la quantité d'intérêt Q considérée, jusqu'à une tolérance donnée. Le problème s'écrit :

$$\mathbf{m}^* = \operatorname{argmin}_{\mathbf{m} \in \mathcal{T}} \left[\frac{1}{2} (Q(\mathbf{z}) - Q(\mathbf{z}_0(\mathbf{m})))^2 \right] = \operatorname{argmin}_{\mathbf{m} \in \mathcal{T}} J(\mathbf{z}_0(\mathbf{m})) \quad (3.104)$$

où J représente la fonction coût à minimiser. On utilise alors un algorithme itératif de gradient, tel que la valeur approchée de $\mathbf{m}^*$ à l'itération $k + 1$ s'écrit :

$$\mathbf{m}^{(k+1)} = \mathbf{m}^{(k)} - \eta \nabla_{\mathbf{m}} J(\mathbf{z}_0(\mathbf{m}^{(k)})) \quad (3.105)$$

avec η un paramètre de l'algorithme (pas de descente). Le calcul de $\nabla_{\mathbf{m}} J$ est fait par la méthode de l'état adjoint exposée dans la Section 3.1 du Chapitre 1, en définissant le lagrangien associé à (3.104) :

$$\mathcal{L}(\mathbf{z}_0, \tilde{\boldsymbol{\lambda}}, \mathbf{m}) = J(\mathbf{z}_0) - [a_{0\mathbf{m}}(\mathbf{z}_0; \tilde{\boldsymbol{\lambda}}) - l(\tilde{\boldsymbol{\lambda}})] \quad (3.106)$$

La recherche du point-selle définit le problème adjoint :

$$\text{Trouver } \tilde{\boldsymbol{\lambda}} \in \mathcal{V} \text{ tel que : } a'_{0\mathbf{m}, \mathbf{z}_0}(\mathbf{z}_0; \mathbf{w}, \tilde{\boldsymbol{\lambda}}) = J'_{\mathbf{z}_0}(\mathbf{z}_0; \mathbf{w}) \quad \forall \mathbf{w} \in \mathcal{V} \quad (3.107)$$

et on obtient directement, pour une solution $\mathbf{z}_0(\mathbf{m})$ vérifiant (3.103) :

$$\nabla_{\mathbf{m}} J(\mathbf{z}_0(\mathbf{m})) \cdot \delta \mathbf{m} = -a'_{0\mathbf{m}, \mathbf{m}}(\mathbf{z}_0(\mathbf{m}); \tilde{\boldsymbol{\lambda}}, \delta \mathbf{m}) \quad (3.108)$$

cette expression étant valable pour toute direction de recherche $\delta \mathbf{m}$. L'évaluation du second membre de (3.108) est en pratique menée avec un schéma aux différences finies, en définissant les sous-régions $\Omega_{d,\Delta m_i}$ et $\Omega_{0,\Delta m_i}$ de Ω qui changent de statut lorsque le paramètre m_i est incrémenté de Δm_i (voir Figure 3.54).

$$\Omega_{d,\Delta m_i} = (\Omega_d \setminus \Omega_o)_{m_i + \Delta m_i} - (\Omega_d \setminus \Omega_o)_{m_i} \quad ; \quad \Omega_{0,\Delta m_i} = (\Omega_d)_{m_i + \Delta m_i} - (\Omega_d)_{m_i} \quad (3.109)$$

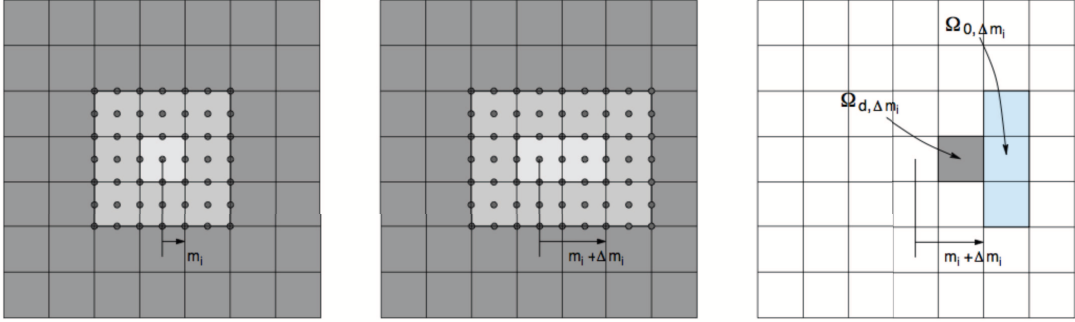


Figure 3.54: Définition des régions $\Omega_{d,\Delta m_i}$ et $\Omega_{0,\Delta m_i}$

Le nouvel algorithme adaptatif basé sur l'approche du contrôle optimal est décrit sur la Figure 3.55.

1. Initialiser $k = 0$
2. Incrémenter $k \rightarrow k + 1$
3. Résoudre le problème couplé primal (3.103) avec $\mathbf{m}^{(k)}$ pour trouver $\mathbf{z}_0(\mathbf{m}^{(k)})$
4. Résoudre le problème couplé adjoint (3.107) avec $\mathbf{m}^{(k)}$ pour trouver $\tilde{\boldsymbol{\lambda}}(\mathbf{m}^{(k)})$
5. Calculer le gradient $\nabla_{\mathbf{m}} J(\mathbf{z}_0(\mathbf{m}^{(k)}))$
6. Mettre à jour $\mathbf{m}^{(k)}$
7. Calculer $\alpha_k = |\nabla_{\mathbf{m}} J(\mathbf{z}_0(\mathbf{m}^{(k)}))|$. Si $\alpha_k \leq \epsilon$, alors stop. Sinon, aller à l'étape 2

Figure 3.55: Algorithme adaptatif basé sur l'approche du contrôle optimal

On compare à présent la procédure d'adaptation présentée dans la Section 2.1.2 (appelée procédure classique) et celle présentée dans cette section (appelée procédure optimale). On considère tout d'abord l'exemple 1D de la Figure 3.56. La structure est une chaîne faite de $n = 89$ particules distantes de $\ell = 0,1$ à l'équilibre. Leurs interactions avec les plus proches voisins sont modélisées par des potentiels harmoniques avec une constante de raideur uniforme $k = 100$. La chaîne est fixée à ses extrémités et soumise à une force unité $f_c = 1$ appliquée à la particule centrale $p_c = (n - 1)/2$. On suppose aussi que la chaîne comprend un défaut local autour de la particule p_d

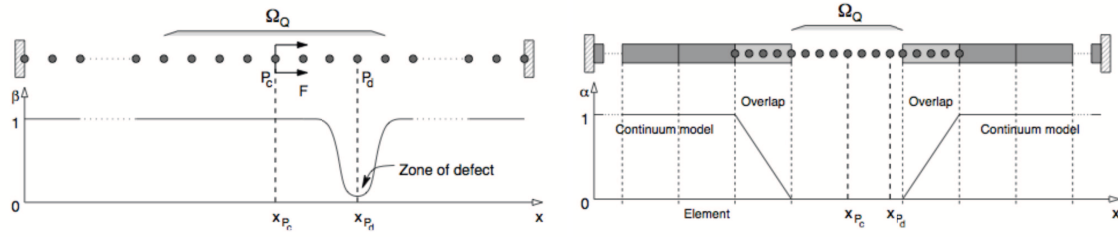


Figure 3.56: Chaîne de particules (gauche) et couplage Arlequin (droite)

(différente de p_c pour générer une adaptation de modèle non-symétrique).

La région d'intérêt est définie par les $n_Q = 9$ particules centrées autour de p_c (Figure 3.56), et on s'intéresse au déplacement moyen de ces n_Q particules :

$$Q(\mathbf{z}) = \frac{1}{n_Q} \sum_{i=1}^{n_Q} z_i \quad (3.110)$$

La valeur exacte de Q , égale à 0,34176, est approximée par un couplage Arlequin (Figure 3.56). La région $\Omega_d \setminus \Omega_o$ coïncide initialement avec la région d'intérêt, et la région de couplage Ω_o s'étend sur cinq particules de chaque côté de $\Omega_d \setminus \Omega_o$. La solution $\mathbf{z}_0$ correspondante, ainsi que la solution adjointe associée à Q , sont montrées sur la Figure 3.57.

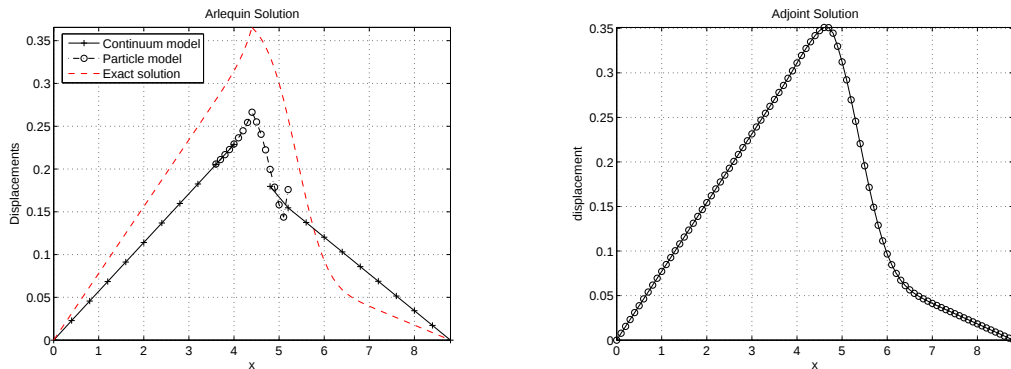


Figure 3.57: Solution Arlequin du problème primal (gauche) et solution exacte du problème adjoint (droite)

On introduit deux paramètres de contrôle, m_L et m_R , comme illustré sur la Figure 3.53; ils représentent les distances d_L et d_R entre la particule p_c et les bords de Ω_c :

$$d_L = (m_L + 1) \times 4\ell \quad ; \quad d_R = (m_R + 1) \times 4\ell \quad (3.111)$$

Les paramètres m_L et m_R sont initialisés à zéro.

En utilisant l'algorithme classique et en imposant une incrémentation unitaire de m_L ou m_R à chaque itération, on obtient les résultats montrés sur la Figure 3.58 et dans le Tableau 3.7

On fournit dans le Tableau 3.8 les valeurs de m_L , m_R et des gradients $\partial_{m_L} J(\mathbf{z}_0(\mathbf{m}))$ et $\partial_{m_R} J(\mathbf{z}_0(\mathbf{m}))$ correspondants calculés à chaque itération avec la

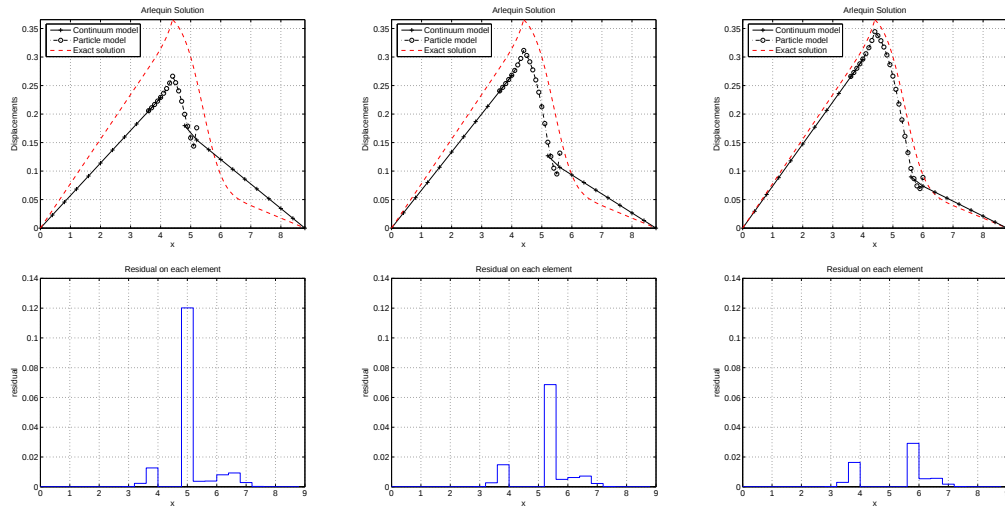


Figure 3.58: Solution Arlequin (haut) et contributions du résidu (bas) obtenues avec la procédure classique d'adaptation aux itérations 0, 1 et 2.

itération	m_L	m_R	valeur de $Q(\mathbf{z}_0)$	erreur (%)
0	0	0	0,23876	30,14
1	0	1	0,28572	16,40
2	0	2	0,31991	6,39
3	0	3	0,33573	1,76
4	1	3	0,33639	1,57
5	1	4	0,34099	0,23
6	1	5	0,34166	$2,82 \cdot 10^{-2}$
7	2	5	0,34170	$1,77 \cdot 10^{-2}$
8	2	6	0,34176	$1,21 \cdot 10^{-3}$

Tableau 3.7: Evolution de la valeur de Q et des paramètres (m_L, m_R) au cours des itérations de la procédure classique d'adaptation

procédure d'adaptation optimale. On observe que les valeurs de m_L et m_R diffèrent légèrement entre les deux approches et que l'approche optimale donne toujours, à itération fixée, une erreur égale ou inférieure à celle donnée par l'approche classique.

On considère à présent la structure 2D représentée sur la Figure 3.59, constituée de $n = 1309$ particules organisées en un réseau régulier de longueur caractéristique $\ell = 0,1$. Les particules interagissent avec leurs premiers et seconds voisins au moyen de potentiels harmoniques de raideur respective $k = 10$ et $k = 5$. La structure est soumise à une densité d'effort vertical $f = 0,01$ appliquée sur chaque particule du bord supérieur droit ; l'action collective de ces forces est représentée par la force F (Figure 3.59).

La quantité d'intérêt considérée est l'élongation de la liaison entre les particules P_1 et P_2 situées dans la région critique de la structure (voir Figure 3.59). Dans la

itération	m_L	m_R	valeur de $Q(\mathbf{z}_0)$	erreur (%)	$ \partial_{m_L} J(\mathbf{z}_0(\mathbf{m})) $	$ \partial_{m_R} J(\mathbf{z}_0(\mathbf{m})) $
0	0	0	0,23876	30,14	$2,98 \cdot 10^{-5}$	0,0015
1	0	1	0,28572	16,40	$2,24 \cdot 10^{-5}$	0,0012
2	0	2	0,31991	6,39	$1,07 \cdot 10^{-5}$	$2,36 \cdot 10^{-4}$
3	0	3	0,33573	1,76	$3,24 \cdot 10^{-6}$	$1,27 \cdot 10^{-5}$
4	0	4	0,34031	0,42	$7,98 \cdot 10^{-7}$	$1,62 \cdot 10^{-7}$
5	1	4	0,34099	0,23	$2,70 \cdot 10^{-8}$	$8,65 \cdot 10^{-8}$
6	1	5	0,34166	$2,82 \cdot 10^{-2}$	$3,39 \cdot 10^{-9}$	$4,43 \cdot 10^{-9}$
7	1	6	0,34172	$1,17 \cdot 10^{-2}$	$1,40 \cdot 10^{-9}$	$1,17 \cdot 10^{-10}$
8	2	6	0,34176	$1,21 \cdot 10^{-3}$	$4,41 \cdot 10^{-12}$	$1,21 \cdot 10^{-11}$

Tableau 3.8: Evolution de la valeur de Q , des paramètres (m_L, m_R) et des gradients de J correspondants au cours des itérations de la procédure optimale d'adaptation

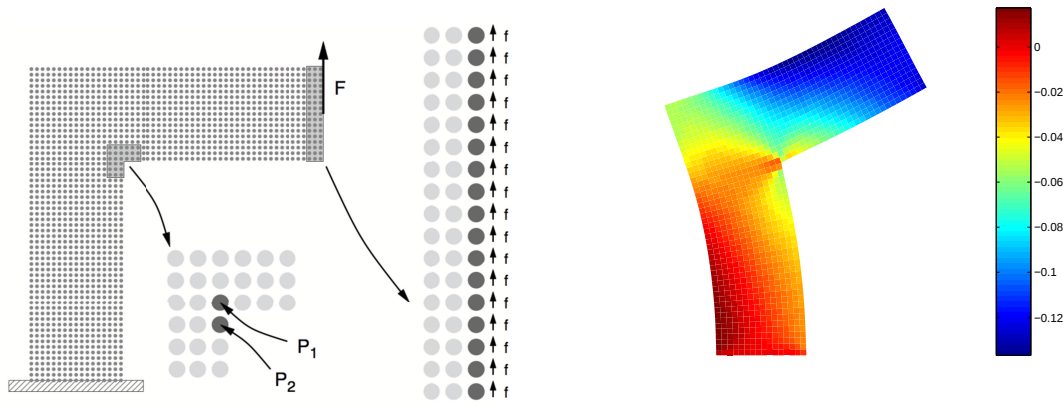


Figure 3.59: Exemple 2D considéré (gauche) et déformation (composante xx) associée (droite)

configuration initiale du couplage Arlequin, le modèle fin particulière est conservé dans le coin et dans la partie supérieure droite de la structure ; le milieu continu placé dans le reste de la structure est discrétisé avec des éléments Q4 de taille 5ℓ . La valeur exacte de la quantité d'intérêt est 1,4599 tandis que celle donnée par la méthode Arlequin est 1,0295.

A chaque itération des deux approches d'adaptation de modèle (classique ou optimale), on enrichit un élément fini à la fois. Pour l'approche optimale, le nombre de paramètres augmente évidemment avec le nombre d'itérations (voir Figure 3.60).

On montre sur la Figure 3.61 les configurations de couplage obtenues aux itérations 0,1 et 3 de la procédure classique d'adaptation. La Figure 3.62 montre quant à elle les configurations obtenues avec l'approche optimale pour les itérations 0, 1 et 3, ainsi que la position du gradient maximal de J par rapport aux paramètres de contrôle. Les valeurs de ce gradient maximal sont respectivement $5,1 \times 10^{-3}$,

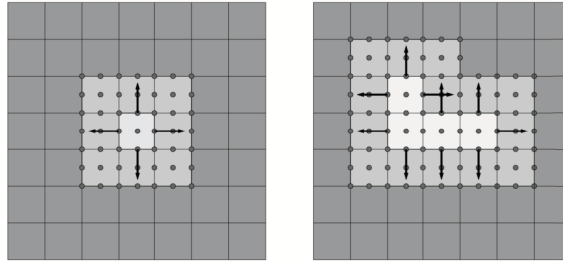


Figure 3.60: Evolution du nombre de paramètres utilisés dans la méthode optimale d'adaptation : 4 paramètres dans la configuration de gauche, 10 paramètres dans la configuration de droite

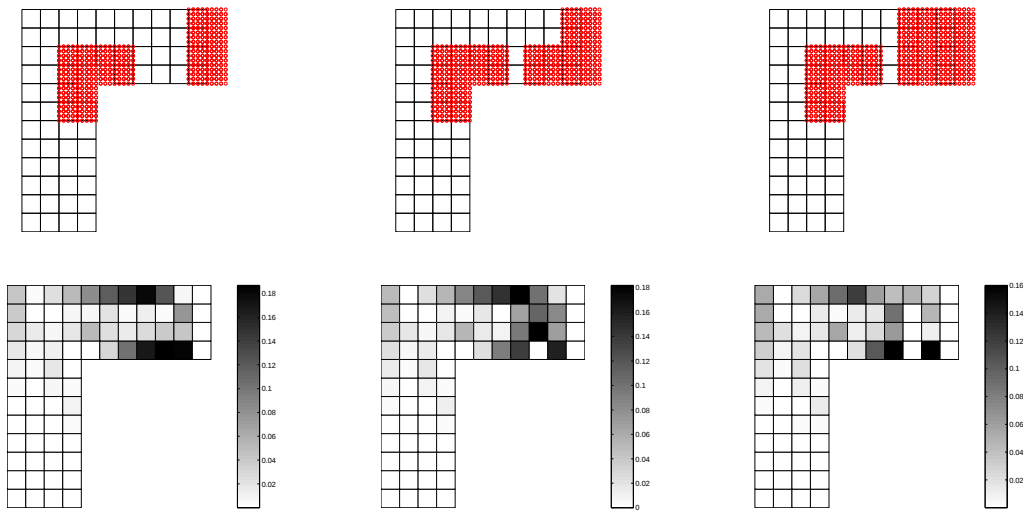


Figure 3.61: Configuration Arlequin et contributions du résidu correspondantes par élément aux itérations 0, 1 et 3 avec l'approche classique

$3,2 \times 10^{-3}$, $3,4 \times 10^{-3}$ et $1,8 \times 10^{-3}$. On compare sur la Figure 3.63 la convergence de la quantité d'intérêt Q en fonction du nombre d'itérations pour les deux procédures d'adaptation. On remarque que l'approche optimale produit toujours des erreurs inférieures à celles que produit l'approche classique et permet de converger plus vite vers la valeur exacte de Q .

On voit donc que la nouvelle procédure d'adaptation est performante pour trouver la configuration de couplage optimale. Néanmoins, dû au fait que cette procédure n'est pas associée à un estimateur d'erreur, le critère d'arrêt utilisé peut s'avérer insuffisant.

3 Contrôle des modèles PGD

Une autre catégorie de modèles complexes est celle des modèles régis par plusieurs paramètres dont l'influence doit être étudiée par simulation ; on peut citer pour exemple les problèmes stochastiques ou d'optimisation. Pour de tels modèles multi-paramétrés, les méthodes d'approximation classiques (utilisation de

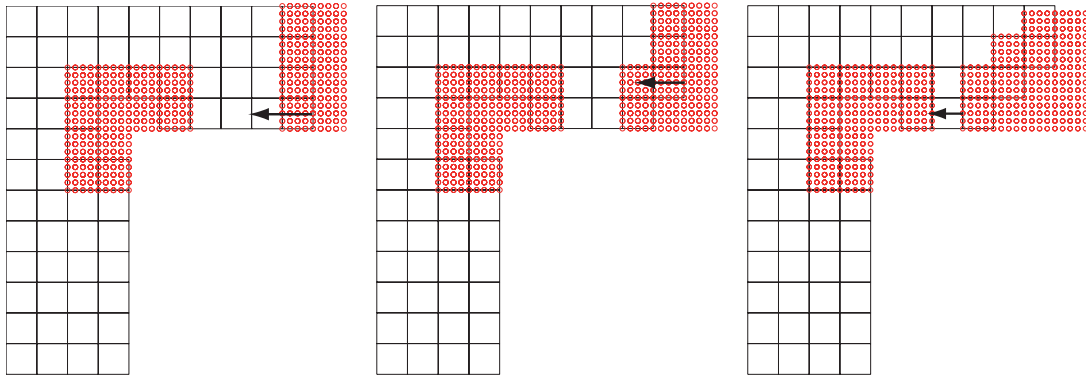


Figure 3.62: Configuration Arlequin et gradient maximal aux itérations 0, 1 et 3 avec l'approche optimale

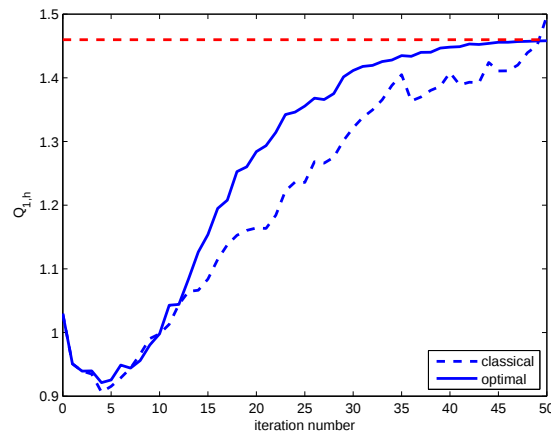


Figure 3.63: Comparaison de la convergence de Q en fonction du nombre d'itérations avec les procédures d'adaptation classique et optimale

grilles dans chaque dimension) mènent à une croissance exponentielle du nombre d'inconnues vis-à-vis du nombre de paramètres. Ce phénomène, connu sous le nom de *malédiction de la dimensionnalité*, peut rapidement être rédhibitoire pour la simulation numérique; des approches alternatives sont donc nécessaires dans ce contexte.

Les techniques de réduction de modèle (ou plus précisément de complexité) utilisant des bases réduites sont un outil approprié pour aborder les modèles multi-paramétrés; elles ont par conséquent été l'objet d'un intérêt grandissant durant la dernière décennie. En particulier, une technique attractive appelée *Proper Generalized Decomposition* (PGD) a émergé récemment et est actuellement le sujet de divers travaux de recherche [Chinesta *et al.* 2010]. La technique, qui reprend les idées initialement développées dans [Ladevèze 1999] pour résoudre les problèmes non-linéaires dépendant du temps, est basée sur la séparation de variables dans une approche de résolution spectrale. Les atouts majeurs sont qu'aucune information sur la solution n'est nécessaire (contrairement à la POD) et que le coût de calcul croît linéairement par rapport au nombre de paramètres. Les fonctions de base PGD (ou modes) sont calculées à la volée, une fois pour toutes et dans un processus *offline*, en

résolvant un ensemble de problèmes mono-paramétrés avec les techniques classiques. L'approximation PGD obtenue, qui dépend explicitement de tous les paramètres du modèle, peut alors être utilisée dans un processus *online* d'analyse ou d'optimisation.

Les performances de la PGD ont été montrées dans beaucoup d'applications exhibant des variations de chargement, de conditions limites ou initiales, de paramètres matériau, de géométrie, ... prises en compte à l'aide de coordonnées additionnelles dans le modèle [Nouy 2007 - 2010, Chinesta *et al.* 2010]. Cependant, une difficulté majeure pour le transfert et l'utilisation intensive des modèles réduits PGD dans l'industrie est le contrôle de leur fiabilité. En effet, certifier la précision de la solution PGD est fondamental pour réaliser un dimensionnement robuste. Le contrôle de la solution PGD nécessite de maîtriser le nombre de modes qui sont calculés (troncature), mais aussi les méthodes numériques qui sont employées dans ces calculs.

Il existe très peu de travaux ayant abordé le contrôle des approximations PGD. Des résultats de base sur l'estimation d'erreur *a priori* pour les représentations par variables séparées sont donnés dans [Ladevèze 1999], tandis qu'un travail pionnier dédié à l'adaptativité peut être trouvé dans [Ammar *et al.* 2010]. Une première approche robuste pour la vérification de la PGD, utilisant le concept d'ERC, a été proposée dans [Ladevèze et Chamoin 2011]. Elle s'applique aux problèmes linéaires elliptiques ou paraboliques dépendant de paramètres, et fournit des modèles réduits PGD garantis vis-à-vis de l'erreur sur des quantités d'intérêt. De plus, l'approche suivie permet d'estimer les contributions des diverses sources d'erreur (discrétisations en espace et temps, troncature de la décomposition PGD, etc.), ce qui constitue une information pertinente pour améliorer efficacement la qualité de la solution PGD. Les performances ont été montrées dans [Ladevèze et Chamoin 2011, Chamoin et Ladevèze 2012, Ladevèze et Chamoin 2012] sur diverses expériences numériques impliquant un modèle de thermique transitoire avec des paramètres matériau fluctuants ; elles sont retranscrites ici sous une forme condensée.

3.1 Présentation de la PGD

On considère un problème de thermique transitoire défini dans un domaine ouvert borné $\Omega \subset R^d$ ($d = 1, 2, 3$), de frontière $\partial\Omega$, sur un intervalle de temps $\mathcal{I} = [0, T]$. Une température nulle est appliquée sur le bord $\partial_1\Omega \neq \emptyset$ de $\partial\Omega$ et le chargement thermique, dépendant du temps, consiste en : (i) un flux thermique imposé $r_d(\mathbf{x}, t)$ sur le bord $\partial_2\Omega \subset \partial\Omega$ complémentaire de $\partial_1\Omega$; (ii) un terme source $f_d(\mathbf{x}, t)$ dans Ω . Les conditions initiales sont prises nulles.

Le matériau qui compose Ω est supposé isotrope mais hétérogène et en partie inconnu. De ce fait, le coefficient de diffusion k et la capacité thermique c dépendent de la variable spatiale $\mathbf{x}$ mais aussi d'un ensemble de N paramètres $\mathbf{p} = [p_1, p_2, \dots, p_N]$ appartenant à un domaine borné $\Theta = \Theta_1 \times \Theta_2 \times \dots \times \Theta_N$ donné.

On note (u, φ) la paire température/flux solution du problème. En définissant $\mathcal{U} = H_0^1(\Omega) = \{v \in H^1(\Omega), v|_{\partial_1\Omega} = 0\}$, la formulation faible en espace du problème

s'écrit pour tout $(t, \mathbf{p}) \in \mathcal{I} \times \Theta$:

$$\text{Trouver } u(\mathbf{x}, t, \mathbf{p}) \in \mathcal{U} \text{ tel que } a(u, v) = l(v) \quad \forall v \in \mathcal{U} \quad (3.112)$$

avec $u|_{t=0^+} = 0$. La forme bilinéaire $a(\bullet, \bullet)$ et la forme linéaire $l(\bullet)$ sont définies par :

$$a(u, v) = \int_{\Omega} \left\{ c \frac{\partial u}{\partial t} v + k \nabla u \cdot \nabla v \right\} d\Omega \quad ; \quad l(v) = \int_{\Omega} f_d v d\Omega - \int_{\partial_2 \Omega} r_d v dS \quad (3.113)$$

On introduit à présent les espaces fonctionnels $\mathcal{T} = L^2(\mathcal{I})$, $\mathcal{P}_i = L^2(\Theta_i)$, et $L^2(\mathcal{I}, \Theta; \mathcal{U}) = \mathcal{U} \otimes \mathcal{T} \otimes_{n=1}^N \mathcal{P}_n$. La formulation faible complète du problème consiste alors à chercher $u \in L^2(\mathcal{I}, \Theta; \mathcal{U})$, avec $\frac{\partial u}{\partial t} \in L^2(\mathcal{I}, \Theta; L^2(\Omega))$, tel que :

$$B(u, v) = L(v) \quad \forall v \in L^2(\mathcal{I}, \Theta; \mathcal{U}) \quad (3.114)$$

avec

$$B(u, v) = \int_{\Theta} \left[\int_{\mathcal{I}} a(u, v) dt + \int_{\Omega} cu(\mathbf{x}, 0^+)v(\mathbf{x}, 0^+)d\Omega \right] d\mathbf{p} \quad ; \quad L(v) = \int_{\Theta} \int_{\mathcal{I}} l(v) dt d\mathbf{p} \quad (3.115)$$

La résolution approchée de (3.114), à partir de la MEF en espace associée à un schéma d'intégration en temps et une grille donnée dans l'espace des paramètres Θ , peut être très coûteuse lorsque le nombre de paramètres augmente.

La technique alternative PGD [Chinesta *et al.* 2011] consiste à construire *a priori* une représentation à variables séparées de la solution u de (3.114). La solution PGD approchée est cherchée sous la forme :

$$u(\mathbf{x}, t, \mathbf{p}) \approx u_m(\mathbf{x}, t, \mathbf{p}) \equiv \sum_{i=1}^m \psi_i(\mathbf{x}) \lambda_i(t) \Gamma_i(\mathbf{p}) \quad \text{avec } \Gamma_i(\mathbf{p}) = \prod_{n=1}^N \gamma_{i,n}(p_n) \quad (3.116)$$

m est l'ordre (i.e. le nombre de modes) de la représentation, tandis que les fonctions d'espace $\psi_i(\mathbf{x})$, du temps $\lambda_i(t)$, et des paramètres $\gamma_{i,n}(p_n)$ appartiennent respectivement à $\mathcal{U}$, $\mathcal{T}$ et $\mathcal{P}_n$.

La construction des modes ne nécessite aucune connaissance particulière sur u , elle est faite à la volée lors de la résolution. On donne ci-après une version classique de la construction des modes, appelée *Galerkin progressive* et inspirée des algorithmes classiques de point fixe.

On suppose qu'une approximation PGD d'ordre $s - 1$ a été calculée. L'approximation PGD d'ordre s est alors définie par :

$$u_s(\mathbf{x}, t, \mathbf{p}) = u_{s-1}(\mathbf{x}, t, \mathbf{p}) + \psi(\mathbf{x}) \lambda(t) \Gamma(\mathbf{p}) \quad \text{avec } \Gamma(\mathbf{p}) = \prod_{n=1}^N \gamma_n(p_n) \quad (3.117)$$

ψ , λ , et γ_n ($n = 1, \dots, N$) sont des fonctions *a priori* inconnues appartenant respectivement aux sous-espaces discrétisés $\mathcal{U}_h$, $\mathcal{T}_h$, et $\mathcal{P}_{nh}$; on suppose que $\mathcal{U}_h$ et $\mathcal{T}_h$ respectent les contraintes cinématiques et les conditions initiales, respectivement. Partant d'une initialisation $\psi^{(0)}(\mathbf{x}) \lambda^{(0)}(t) \Gamma^{(0)}(\mathbf{p})$ pour le mode s , on construit une nouvelle représentation modale $\psi^{(1)}(\mathbf{x}) \lambda^{(1)}(t) \Gamma^{(1)}(\mathbf{p})$ par une approche de Galerkin qui mène à la sous-itération suivante :

→ déterminer $\lambda^{(1)} \in \mathcal{T}_h$ tel que :

$$B(u_{s-1} + \psi^{(0)} \lambda^{(1)} \Gamma^{(0)}, \psi^{(0)} \lambda^* \Gamma^{(0)}) = L(\psi^{(0)} \lambda^* \Gamma^{(0)}) \quad \forall \lambda^* \in \mathcal{T}_h \quad (3.118)$$

→ pour $n_0 = 1, \dots, N$, déterminer $\gamma_{n_0}^{(1)} \in \mathcal{P}_{n_0 h}$ tel que :

$$B(u_{s-1} + \psi^{(0)} \lambda^{(1)} \gamma_{n_0}^{(1)} \Gamma_{/n_0}^{(1,0)}, \psi^{(0)} \lambda^{(1)} \gamma^* \Gamma_{/n_0}^{(1,0)}) = L(\psi^{(0)} \lambda^{(1)} \gamma^* \Gamma_{/n_0}^{(1,0)}) \quad \forall \gamma^* \in \mathcal{P}_{n_0 h} \quad (3.119)$$

$$\text{avec } \Gamma_{/n_0}^{(1,0)} = \prod_{n=1}^{n_0-1} \gamma_n^{(1)} \times \prod_{n=n_0+1}^N \gamma_n^{(0)} ;$$

→ déterminer $\psi^{(1)} \in \mathcal{U}_h$ tel que :

$$B(u_{s-1} + \psi^{(1)} \lambda^{(1)} \Gamma^{(1)}, \psi^* \lambda^{(1)} \Gamma^{(1)}) = L(\psi^* \lambda^{(1)} \Gamma^{(1)}) \quad \forall \psi^* \in \mathcal{U}_h \quad (3.120)$$

La sous-itération consiste en la résolution d'une série de problèmes simples : l'équation différentielle ordinaire en temps issue de (3.118) est résolue par un schéma d'Euler explicite (pas de temps Δt), le problème en espace issu de (3.120) est traité par la MEF (taille caractéristique des éléments h), tandis que la résolution du problème issu de (3.119) est explicite. Quelques sous-itérations sont faites en pratique. De plus, la fonction temporelle $\lambda^{(j)}(t)$ et les fonctions des paramètres $\gamma_n^{(j)}(p_n)$ sont normalisées à chaque sous-itération. Il faut noter qu'une sous-itération se termine avec le problème en espace (3.120), ce qui a une importance pour la technique exposée par la suite. Des optimisations, telles que la mise à jour de l'ensemble $\{\lambda_i\}$ (resp. $\{\Gamma_i\}$) des fonctions du temps (resp. des paramètres) ou l'orthogonalisation des modes spatiaux, sont possibles ; nous ne les considérons pas par la suite même si le cadre de contrôle d'erreur défini dans cette section permet de le faire simplement.

3.2 Estimation d'erreur locale par l'ERC

La stratégie de vérification proposée utilise le concept d'ERC. Soit $(\hat{u}, \hat{\varphi})$ une solution admissible du problème i.e. vérifiant, en plus des conditions initiales, les contraintes cinématiques et les équations d'équilibre pour tout $(t, \mathbf{p}) \in \mathcal{I} \times \Theta$. La mesure de l'ERC dans le domaine spatio-temporel, qui dépend de $\mathbf{p}$, s'écrit alors :

$$e_{ERC}^2(\mathbf{p}) = \frac{1}{2} \int_{\mathcal{I}} \int_{\Omega} \frac{1}{k} [\hat{\varphi} - k \nabla \hat{u}] \cdot [\hat{\varphi} - k \nabla \hat{u}] d\Omega dt \equiv \frac{1}{2} ||| \hat{\varphi} - k \nabla \hat{u} |||_{k^{-1}}^2 \quad (3.121)$$

avec $||| \bullet |||_{k^{-1}}$ la norme énergétique dans le domaine spatio-temporel, et l'extension du théorème de Prager-Synge est :

$$||| \varphi - \hat{\varphi}^m |||_{k^{-1}}^2 + \frac{1}{2} \int_{\Omega} c(u_{ex} - \hat{u})^2 d\Omega = \frac{1}{2} e_{ERC}^2 \quad (3.122)$$

avec $\hat{\varphi}^m = \frac{1}{2} [\hat{\varphi} + k \nabla \hat{u}]$.

A nouveau, le point technique est la construction d'une solution admissible ; on explique ci-dessous comment une telle solution, notée $(\hat{u}_m, \hat{\varphi}_m)$, peut être obtenue par post-traitement de toute l'information disponible à partir du calcul de la solution

PGD u_m .

Construire le champ admissible cinématique $\hat{u}_m(\mathbf{x}, t, \mathbf{p})$ est simple, et on le prend par la suite égal à $\hat{u}_m(\mathbf{x}, t, \mathbf{p})$. Obtenir $\hat{\varphi}_m(\mathbf{x}, t, \mathbf{p})$ est plus technique ; pour utiliser les outils classiques qui permettent de calculer des flux équilibrés (en particulier la condition de prolongement, voir Section 3.1 du Chapitre 2), on doit d'abord construire un champ $\varphi_m(\mathbf{x}, t, \mathbf{p})$ qui satisfait l'équilibre EF pour tout $(t, \mathbf{p}) \in \mathcal{I} \times \Theta$:

$$\int_{\Omega} \varphi_m \cdot \nabla v d\Omega = \int_{\Omega} (f_d - c \frac{\partial \hat{u}_m}{\partial t}) v d\Omega - \int_{\partial_2 \Omega} r_d v dS \quad \forall v \in \mathcal{U}_h \quad (3.123)$$

On suppose tout d'abord que le chargement extérieur peut être écrit sous une forme radiale :

$$(f_d(\mathbf{x}, t), r_d(\mathbf{x}, t)) = \sum_{j=1}^J \alpha_j(t) (f_d^j(\mathbf{x}), r_d^j(\mathbf{x})) \quad (3.124)$$

On calcule alors, pour chaque couple (f_d^j, r_d^j) , un champ $\varphi_d^j(\mathbf{x})$ vérifiant l'équilibre EF :

$$\int_{\Omega} \varphi_d^j \cdot \nabla v d\Omega = \int_{\Omega} f_d^j v d\Omega - \int_{\partial_2 \Omega} r_d^j v dS \quad \forall v \in \mathcal{U}_h \quad (3.125)$$

Ce calcul est en pratique réalisé avec la MEF en déplacement. Il vient, en introduisant $\varphi_d = \sum_{j=1}^J \alpha_j(t) \varphi_d^j(\mathbf{x})$ dans (3.123), que φ_m doit alors vérifier pour tout $(t, \mathbf{p}) \in \mathcal{I} \times \Theta$:

$$\int_{\Omega} (\varphi_m - \varphi_d) \cdot \nabla v d\Omega = - \int_{\Omega} c \frac{\partial \hat{u}_m}{\partial t} v d\Omega = - \sum_{i=1}^m c \dot{\lambda}_i \Gamma_i \int_{\Omega} \psi_i v d\Omega \quad \forall v \in \mathcal{U}_h \quad (3.126)$$

D'autre part, à la fin des sous-itérations pour calculer chaque mode PGD $m_0 \in [1, m]$, la condition (3.120) donne :

$$B(u_{m_0}, \psi^* \lambda_{m_0} \Gamma_{m_0}) = L(\psi^* \lambda_{m_0} \Gamma_{m_0}) \quad \forall \psi^* \in \mathcal{U}_h \quad (3.127)$$

Cette dernière relation peut se reformuler sous la forme :

$$\int_{\Omega} \mathbf{H}_{m_0} \cdot \nabla \psi^* d\Omega = \int_{\Omega} \sum_{i=1}^{m_0} [G_{m_0, i} \psi_i] \psi^* d\Omega \quad \forall \psi^* \in \mathcal{U}_h \quad (3.128)$$

avec

$$\mathbf{H}_{m_0} \equiv \int_{\Theta} \int_{\mathcal{I}} \lambda_{m_0} \Gamma_{m_0} (\varphi_d - k \nabla u_{m_0}) dt d\mathbf{p} \quad ; \quad G_{m_0, i} \equiv \int_{\Theta} \int_{\mathcal{I}} c \lambda_{m_0} \Gamma_{m_0} \dot{\lambda}_i \Gamma_i dt d\mathbf{p} \quad (3.129)$$

Par conséquent, pour tout $m_0 \in [1, m]$, le terme $\mathbf{H}_{m_0}$ équilibre $\sum_{i=1}^{m_0} G_{m_0, i} \psi_i$ au sens EF. Une simple inversion génère ainsi des termes de la forme $\sum_{j=1}^m R_{ij} \mathbf{H}_j$ qui équilibrent au sens EF chaque fonction ψ_i ($i = 1, \dots, m$). Un champ φ_m satisfaisant l'équilibre EF (3.123) (ou (3.126)) est donc :

$$\varphi_m = \varphi_d - \sum_{i=1}^m \sum_{j=1}^m c \dot{\lambda}_i \Gamma_i R_{ij} \mathbf{H}_j \quad (3.130)$$

A partir de φ_m , les techniques usuelles peuvent ensuite être utilisées pour calculer un flux $\hat{\varphi}_m$ qui vérifie l'équilibre de façon stricte :

$$\int_{\Omega} \hat{\varphi}_m \cdot \nabla v d\Omega = \int_{\Omega} (f_d - c \frac{\partial \hat{u}_m}{\partial t}) v d\Omega - \int_{\partial_2 \Omega} r_d v dS \quad \forall v \in \mathcal{U} \quad (3.131)$$

Ce flux s'écrit sous la forme $\hat{\varphi}_m = \hat{\varphi}_d - \sum_{i=1}^m \sum_{j=1}^m c \dot{\lambda}_i \Gamma_i R_{ij} \hat{\mathbf{H}}_j$ où $\hat{\varphi}_d$ et $\hat{\mathbf{H}}_j$ sont calculés en résolvant des problèmes locaux sur chaque élément ou patch d'éléments.

On peut remarquer que dans le cas d'un problème stationnaire, les termes $\mathbf{H}_{m_0} = \int_{\Theta} \Gamma_{m_0} (\varphi_d - k \nabla u_{m_0}) d\mathbf{p}$ ($m_0 = 1, \dots, m$) sont auto-équilibrés au sens des éléments finis. De ce fait, φ_m et $\hat{\varphi}_m$ peuvent être définis par :

$$\varphi_m(\mathbf{x}, \mathbf{p}) = \varphi_d(\mathbf{x}) + \sum_{m_0=1}^m \beta_{m_0}(\mathbf{p}) \mathbf{H}_{m_0}(\mathbf{x}) \quad ; \quad \hat{\varphi}_m(\mathbf{x}, \mathbf{p}) = \hat{\varphi}_d(\mathbf{x}) + \sum_{m_0=1}^m \beta_{m_0}(\mathbf{p}) \hat{\mathbf{H}}_{m_0}(\mathbf{x}) \quad (3.132)$$

où les coefficients β_{m_0} , fonctions de $\mathbf{p}$, sont obtenus explicitement en minimisant $\int_{\Theta} e_{ERC}^2(\mathbf{p}) d\mathbf{p}$.

A partir de l'estimateur d'erreur globale $e_{ERC}^2(\mathbf{p})$ défini précédemment, il est ensuite aisé de construire un estimateur d'erreur locale à partir de techniques similaires à celles détaillées dans le Chapitre 2. Soit $Q(u(\mathbf{p}))$ une quantité d'intérêt définie par les extracteurs $(\varphi_{\Sigma}, \mathbf{f}_{\Sigma})$:

$$Q(u(\mathbf{p})) = \int_{\mathcal{I}} \int_{\Omega} \{ \nabla u(\mathbf{p}) \cdot \varphi_{\Sigma} + u(\mathbf{p}) \cdot \mathbf{f}_{\Sigma} \} d\Omega dt \quad (3.133)$$

On introduit le problème adjoint associé, et on calcule une solution PGD approchée (resp. admissible) $(\tilde{u}_{\tilde{m}}, \tilde{\varphi}_{\tilde{m}})$ (resp. $(\hat{\tilde{u}}_{\tilde{m}}, \hat{\tilde{\varphi}}_{\tilde{m}})$) pour ce problème. En pratique, la résolution PGD du problème adjoint est faite avec un ordre $\tilde{m}$ potentiellement différent de m , ainsi qu'en introduisant des fonctions d'enrichissement locales (cf. Section 3.2 du Chapitre 2). On obtient la majoration :

$$|Q(u(\mathbf{p})) - Q(u_m(\mathbf{p})) - Q_{corr}(\mathbf{p})| \leq e_{ERC}(\mathbf{p}) \cdot \tilde{e}_{ERC}(\mathbf{p}) \quad (3.134)$$

où $Q_{corr}(\mathbf{p})$ est un terme de correction calculable, et $e_{ERC}(\mathbf{p})$ (resp. $\tilde{e}_{ERC}(\mathbf{p})$) est l'erreur en relation de comportement du problème de référence (resp. adjoint). Ainsi, des bornes garanties sur l'erreur locale $Q(u(\mathbf{p})) - Q(u_m(\mathbf{p}))$ (ou directement sur $Q(u(\mathbf{p}))$) peuvent être obtenues pour toute valeur $\mathbf{p}$ des paramètres du modèle.

3.3 Séparation des sources et adaptation

L'erreur sur Q vient de deux sources : (i) la troncature à un ordre m de la représentation PGD (3.116) ; (ii) la discrétisation spatio-temporelle utilisée pour calculer numériquement les modes. On peut écrire :

$$Q(u) - Q(u_m) = [Q(u) - Q(u^{h, \Delta t})] + [Q(u^{h, \Delta t}) - Q(u_m)] = \Delta Q_{dis} + \Delta Q_{trunc} \quad (3.135)$$

avec ΔQ_{tronc} (resp. ΔQ_{dis}) la part d'erreur sur Q venant de la troncature PGD (resp. de la discrétisation). $u^{h,\Delta t}$ est la solution du problème issu de la discrétisation de (3.114) et vu comme problème de référence pour définir l'erreur ΔQ_{tronc} .

Pour contrôler efficacement le processus de calcul avec la PGD, on introduit un indicateur d'erreur pour chacune des sources d'erreur. L'évaluation de ΔQ_{tronc} est faite à partir de l'ERC en prenant comme référence le problème discrétisé (de solution $u^{h,\Delta t}$) qui peut se mettre sous la forme :

$$\mathbf{U}_h^1 = \mathbf{0} \quad ; \quad \mathbb{M}(\mathbf{p}) \frac{\mathbf{U}_h^{k+1} - \mathbf{U}_h^k}{\Delta t} + \mathbb{K}(\mathbf{p}) \mathbf{U}_h^k = \mathbf{F}_h^k \quad \forall k \geq 1 \quad (3.136)$$

avec $\mathbf{U}_h^k$ le vecteur des inconnues nodales de u au piquet de temps k . Une solution admissible est alors reconstruite au sens du problème discrétisé par post-traitement direct de l'information disponible. L'évaluation de ΔQ_{dis} est ensuite obtenue par différence entre celles de $Q(u) - Q(u_m)$ et de ΔQ_{tronc} .

Dès lors, une stratégie adaptative simple consiste à évaluer, après le calcul de chaque mode PGD, les différentes sources d'erreur sur Q et à adapter en fonction de la source dominante : si $|\Delta Q_{dis}|$ est dominant, on définit une discrétisation plus fine (jusqu'à obtenir $|\Delta Q_{dis}| < |\Delta Q_{tronc}|$) ; si $|\Delta Q_{tronc}|$ est dominant, on calcule le mode PGD suivant sans modifier la discrétisation.

On considère pour exemple la structure 2D représentée sur la Figure 3.64. Il s'agit d'une section droite Ω percée de deux trous rectangulaires dans lesquels circule un fluide. Par symétrie, on ne considère qu'un quart de la structure. Un flux $r_d(t) = -1$ est imposé sur les bords du trou tandis qu'un terme source $f_d(x, y) = 200xy$ est imposé dans Ω . Une température nulle est imposée sur le reste de $\partial\Omega$.

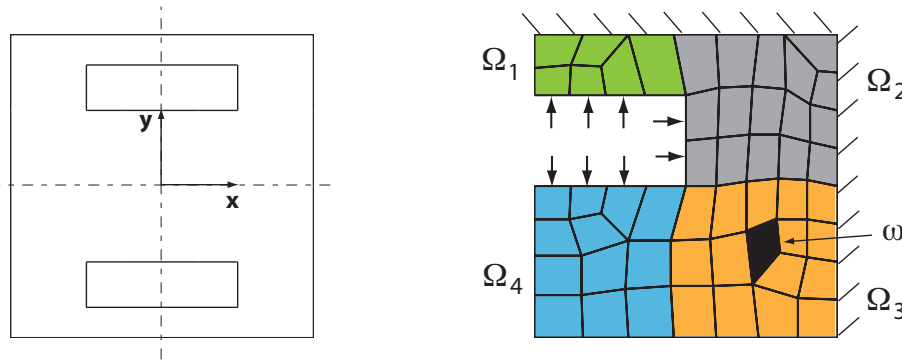


Figure 3.64: Représentation du problème multi-paramétré 2D

On considère que le coefficient de diffusion k fluctue mais reste homogène par morceaux, i.e. est homogène dans chacun des quatre sous-domaines Ω_i ($i = 1, 2, 3, 4$) définis sur la Figure 3.64 et tels que $\Omega_i \cap \Omega_j = \emptyset$, $\Omega_1 \cup \Omega_2 \cup \Omega_3 \cup \Omega_4 = \Omega$. La capacité thermique c est quant à elle supposée homogène dans tout le domaine Ω . Les deux coefficients matériau sont donc définis par 5 paramètres $(\theta_1, \theta_2, \theta_3, \theta_4, \theta_5)$ tels que :

$$k(\mathbf{x}, \theta_i) = 1 + \sum_{i=1}^4 g_i I_{\Omega_i}(\mathbf{x}) \theta_i \quad c(\mathbf{x}, \theta_5) = 1 + 0.2 \theta_5 \quad (3.137)$$

avec $[g_1, g_2, g_3, g_4] = [0, 1; 0, 1; 0, 2; 0, 05]$, et $I_{\Omega_i}(\mathbf{x})$ désignant la fonction indicatrice du sous-domaine Ω_i .

La solution résultante $u(\mathbf{x}, t, \mathbf{p})$, avec $\mathbf{p} = [\theta_1, \theta_2, \theta_3, \theta_4, \theta_5]$, est cherchée par la technique PGD, avec une discrétisation initiale faite de 50 éléments Q4 en espace et 1000 pas de temps.

Dans un premier temps, on considère le cas $\theta_i = 0$ ($i = 1, \dots, 5$); les seules variables du problème sont donc $\mathbf{x}$ et t . On représente sur la Figure 3.65 les trois premiers modes PGD (ψ_i, λ_i) calculés. L'évolution de l'estimateur (relatif) d'erreur

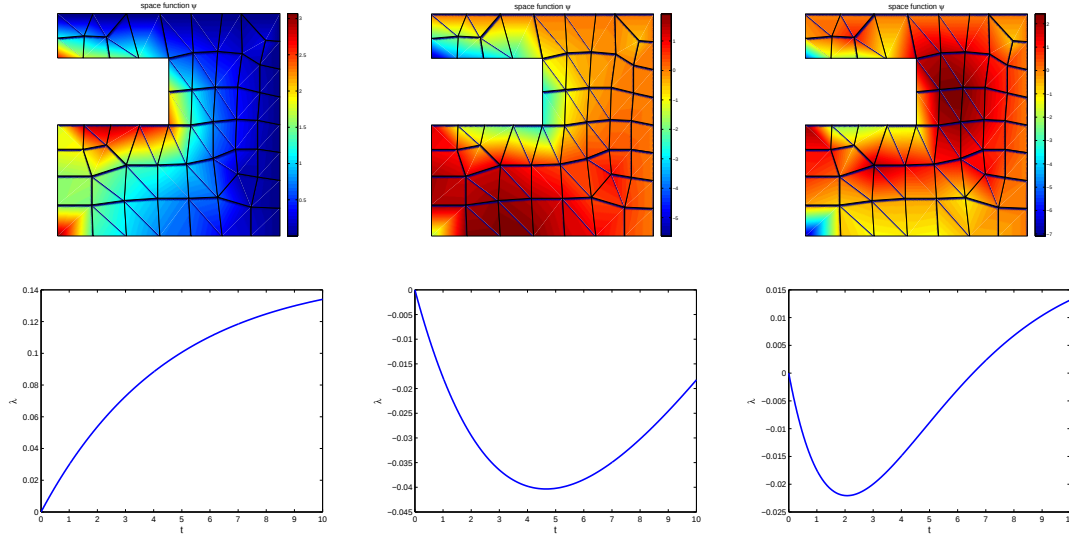


Figure 3.65: Fonctions spatiales ψ_i (haut) et temporelles λ_i (bas) pour $i = 1$ (gauche), $i = 2$ (centre) et $i = 3$ (droite)

globale $e_{ERC}/|||\varphi_{\tilde{m}}|||_{k-1}$, en fonction du nombre m de modes PGD calculés, est représentée sur la Figure 3.66. On observe qu'après 5 modes, l'erreur atteint une valeur asymptotique qui correspond à l'erreur de discrétisation.

La quantité d'intérêt étudiée est la moyenne de température dans une zone locale $\omega \subset \Omega$ (Figure 3.64) à l'instant final T :

$$Q(u) = \frac{1}{|\omega|} \int_{\omega} u|_T d\Omega \quad (3.138)$$

L'erreur locale relative $e_{CRE} \cdot \tilde{e}_{CRE} / (Q(u_m) + Q_{corr})$, ainsi que les indicateurs relatifs sur l'erreur de discrétisation et de troncature sont représentés sur la Figure 3.66 en fonction de m . On observe qu'à partir du mode 3, l'erreur de discrétisation devient dominante et nécessite un raffinement de la discrétisation.

On considère à présent que les paramètres θ_i sont des variables stochastiques normales centrées réduites (tronquées), avec $\theta_i \in [-2, 2]$ ($i = 1, \dots, 5$). On étudie deux quantités d'intérêt : (i) l'espérance mathématique de la moyenne de température dans la zone locale $\omega \subset \Omega$ (Figure 3.64) à l'instant final T ; (ii) la valeur maximale

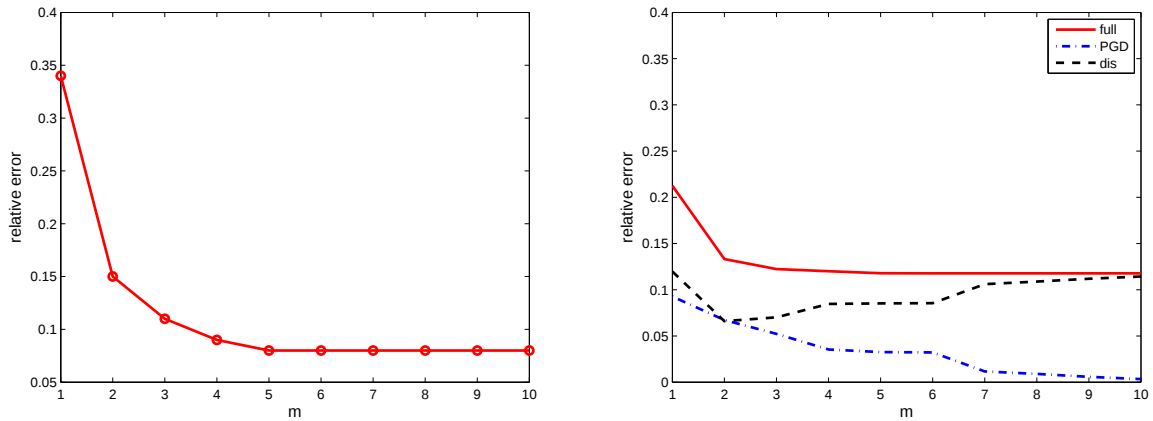


Figure 3.66: Estimateur d'erreur globale relatif (gauche) et estimateur d'erreur locale relatif avec indicateurs sur les sources d'erreur (droite), en fonction du nombre m de modes PGD

de cette moyenne :

$$Q(\mathbf{p}) = \frac{1}{|\omega|} \int_{\omega} u(\mathbf{p})|_T d\Omega \quad ; \quad Q_1 = E_{\Theta} [Q(\mathbf{p})] \quad ; \quad Q_2 = \sup_{\mathbf{p} \in \Theta} [Q(\mathbf{p})] \quad (3.139)$$

L'erreur locale relative $\int_{\Theta} e_{CRE} \cdot \tilde{e}_{CRE} d\mathbf{p} / (Q_1(u_m) + Q_{1,corr})$ (resp. $\sup_{\Theta} e_{CRE} \cdot \tilde{e}_{CRE} / (Q_2(u_m) + Q_{2,corr})$) associée à Q_1 (resp. Q_2), ainsi que les indicateurs relatifs sur l'erreur de discrétisation et de troncature sont représentés sur la Figure 3.67 en fonction de m .

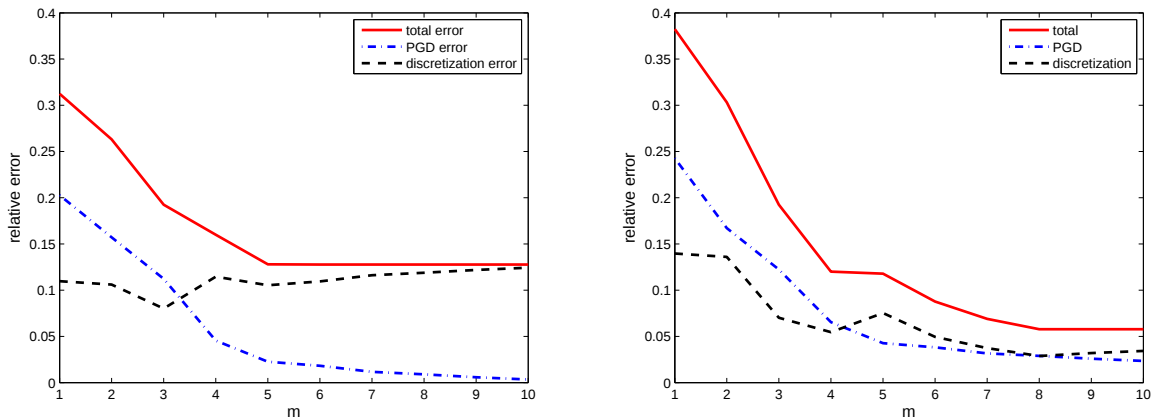


Figure 3.67: Estimateur d'erreur locale relatif et indicateurs sur les sources d'erreur en fonction du nombre m de modes PGD : pour Q_1 (gauche), pour Q_2 (droite)

Pour ce même cas-test, on met en œuvre la procédure adaptative (algorithme glouton) issue de l'évaluation des sources d'erreur. La convergence de l'erreur locale relative obtenue pour Q_1 et Q_2 est montrée sur la Figure 3.68. Les évolutions verticales des courbes indiquent un raffinement de la discrétisation.

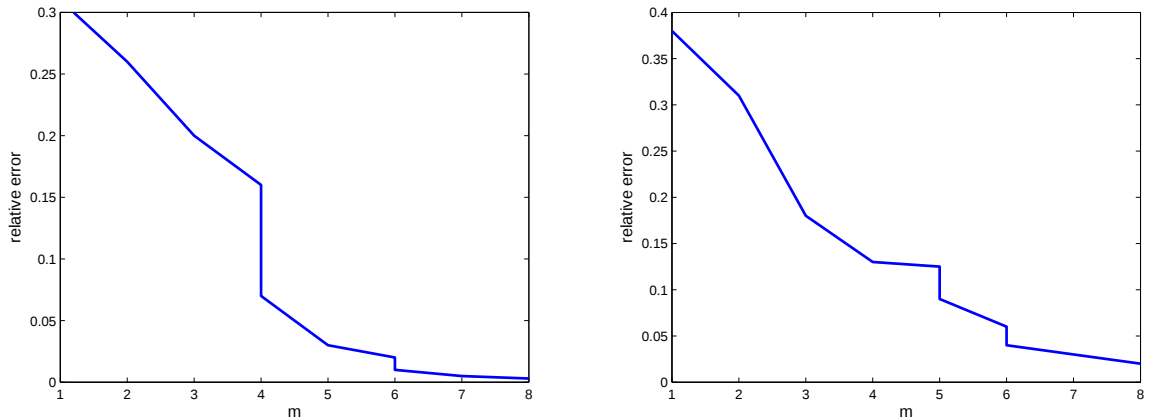


Figure 3.68: Convergence de l'erreur relative en fonction du nombre m de modes PGD dans la stratégie adaptative : pour Q_1 (gauche), pour Q_2 (droite)

On observe que la procédure permet de diminuer efficacement l'erreur sur Q_1 ou Q_2 . Jusqu'au mode 4, la discrétisation grossière initiale est suffisante pour calculer la solution PGD, et l'erreur locale est diminuée en calculant un nouveau mode. Au-delà de $m = 4$, les modes PGD représentent des détails plus fins de la solution qui nécessitent une modification de la discrétisation. Ceci est illustré sur la Figure 3.69 où sont représentés les maillages raffinés (structure *quadtree*) utilisés pour le calcul du huitième mode PGD, pour la stratégie d'adaptation relative à Q_1 ou Q_2 .

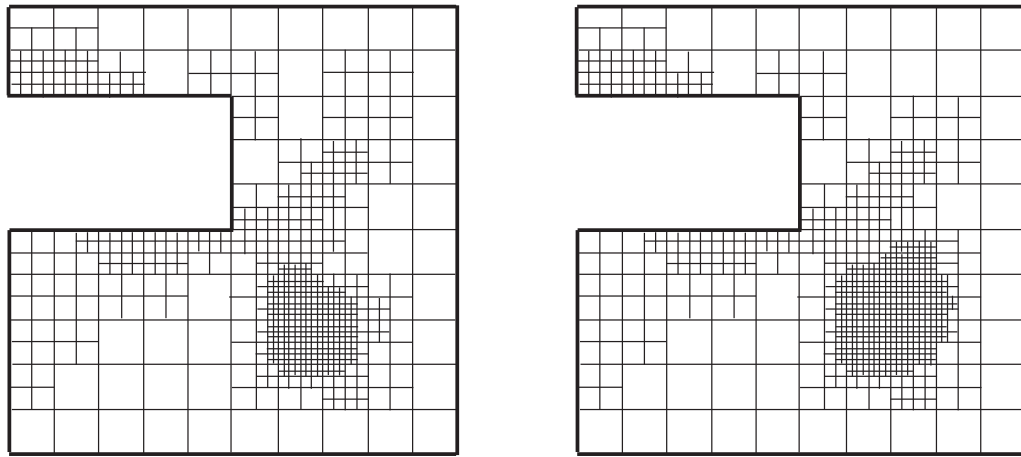


Figure 3.69: Maillages raffinés utilisés pour le calcul du mode PGD $m = 8$ après application de la stratégie adaptative : pour Q_1 (gauche), pour Q_2 (droite)

4 Conclusions partielles et perspectives

Les travaux exposés dans ce chapitre ont montré qu'il était possible de contrôler efficacement les modèles réduits et de certifier leur fiabilité, ce qui est une condition indispensable pour leur utilisation comme outil de dimensionnement dans le monde

industriel. Dans le cas de couplage de modèles, les outils proposés permettent d'une part de faire dialoguer de façon pertinente les modèles et de positionner idéalement l'interface de couplage. Dans le cas d'une représentation modale de la solution, il est possible de savoir le nombre de modes à considérer et la façon de les calculer pour une précision donnée. La thématique de réduction de modèle étant assez jeune, de nombreux travaux restent cependant à être menés sur le contrôle de l'erreur de réduction.

Pour le couplage de modèles proprement dit, la mise en place de méthodes propres de couplage est un axe de recherche important pour les problèmes complexes tels que ceux de dynamique et/ou stochastiques incorporant des échelles variées. L'élaboration de techniques de couplage non-intrusives, permettant notamment de coupler des codes de calcul d'une façon aisée, est aussi un enjeu majeur. L'estimation d'erreur et le contrôle adaptatif du couplage de modèles nécessitent d'être appliqués ou étendus à des cas de couplage autres que ceux présentés dans ce chapitre ; on peut citer notamment le couplage entre des modèles plaque/coque et des modèles 3D qui intéresse fortement le milieu industriel (travail en cours avec P.T. Mendonça). De plus, un paramétrage de la zone de couplage à base de *level-set* et apportant de la flexibilité dans l'adaptation des modèles couplés est à l'étude.

Les travaux menés sur le contrôle de la PGD sont actuellement poursuivis à travers deux thèses au LMT-Cachan. Dans la première (thèse de A. Courard, en partenariat avec EADS Astrium), on s'intéresse aux applications avec variation de géométrie pour la construction d'abaques virtuels garantis exploitant la PGD ; dans la seconde (thèse de P.E. Allier), on traite de problèmes à très grand nombre de paramètres (modèles de composites), une des problématiques étant l'intégration numérique dans les différents espaces paramétriques qui peut s'avérer coûteuse. Une voie de développement prometteuse concerne l'utilisation de l'ERC (cadre d'erreur globale ou locale) pour piloter la méthode PGD, dès la construction des modes ; ceci permet en particulier d'éviter le choix arbitraire d'une norme à minimiser et s'étend aux problèmes non-linéaires. Enfin, des outils de contrôle basés sur l'ERC pourraient être développés pour la POD, pour l'application locale de la PGD dans une structure (modèles cinétiques, optimisation topologique) ou pour d'autres méthodes avancées de réduction et d'approximation comme la TVRC en dynamique vibratoire.

Contrôle de l'erreur locale de modèle

Sommaire

1	Recalage de modèle par l'ERC modifiée	134
2	Vers la validation en temps réel	137
2.1	Construction d'un modèle réduit de processus d'usinage	138
2.2	Contrôle en temps réel par recalage du modèle réduit	141
2.3	Application	143
3	Recalage dédié à des quantités d'intérêt	146
3.1	Définition des nouvelles fonctions coût	147
3.2	Sensibilité vis-à-vis des mesures	152
3.3	Application	154
4	Conclusions partielles et perspectives	156

Nous présentons dans ce chapitre quelques travaux récents, et essentiellement exploratoires, sur le recalage (ou identification selon le degré de connaissance des paramètres) des modèles vis-à-vis de données expérimentales. La validation proprement dite, qui nécessite des analyses complémentaires avec un modèle recalé, n'est pas abordée ici. De plus, l'interaction entre la vérification du modèle numérique et le recalage, i.e. l'impact de l'erreur de discrétisation sur les résultats de recalage (voir [Deraemaeker 2001, Becker et Vexler 2004]), n'est pas considérée en supposant que cette première erreur est contrôlée.

On considère un modèle mathématique linéaire décrit par un ensemble $\mathbf{p}$ de paramètres liés au comportement de la structure. Le problème associé, écrit sous forme faible, consiste à trouver $\mathbf{u}(\mathbf{p}) \in \mathcal{U}$ tel que :

$$A(\mathbf{p}, \mathbf{u}, \mathbf{v}) = a(\mathbf{p}, \mathbf{u}, \mathbf{v}) - l(\mathbf{v}) = 0 \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (4.1)$$

On suppose que la valeur de $\mathbf{p}$ n'est pas parfaitement connue et doit être recalée à partir d'un ensemble d'informations expérimentales locales. Ces dernières informations sont par la suite représentées par un ensemble de quantités observables $\mathbf{s}_{obs}$ obtenues par mesure de la réponse de la structure en des points sélectionnés.

Améliorer la connaissance sur un modèle à partir des données mesurées est une procédure couramment utilisée; un exemple classique est l'identification de défauts par contrôle non-destructif [Andrieux et Bui 2006]. Cette procédure, qui conduit à la résolution d'un problème inverse [Bonnet et Constantinescu 2005] généralement mal posé, consiste à minimiser l'écart entre les prédictions numériques et les mesures expérimentales [Caillaud et Pilvin 1994]. Elle est formulée comme la minimisation d'une fonction coût $\mathcal{F}$ dans un espace de paramètres $\mathcal{P}$:

$$\mathbf{p}_{sol} = \underset{\mathbf{p} \in \mathcal{P}}{\operatorname{argmin}} \mathcal{F}(\mathbf{p}) \quad (4.2)$$

Des techniques spécifiques de régularisation [Tikonov et Arsenin 1977] avec ajout d'information *a priori* sur $\mathbf{p}$ sont généralement utilisées pour stabiliser le processus d'inversion et réduire la sensibilité.

Nous utilisons ici une méthode de recalage basée sur l'ERC [Bonnet et Abdallah 1994, Ladevèze et Chouaki 1999]. Après avoir rappelé les bases de cette méthode, nous présentons un travail où elle est associée à un modèle réduit (construit par la PGD) afin de mener le recalage en temps réel. Nous développons ensuite une adaptation de la méthode de recalage visant à focaliser celui-ci sur la prédiction de quantités d'intérêt.

1 Recalage de modèle par l'ERC modifiée

On présente ici les idées de base de la méthode de recalage considérée dans ce chapitre et utilisant le concept d'ERC. Ce concept a été initialement introduit pour le recalage de modèles dans [Ladevèze *et al.* 1994, Chouaki *et al.* 1996, Ladevèze et Chouaki 1999] pour les modèles dynamiques de vibrations forcées,

avant d'être utilisé dans de nombreuses applications pratiques comme les données et comportements incertains [Faverjon *et al.* 2009] ou les données corrompues [Allix *et al.* 2005]. Particulièrement adapté au recalage du comportement dans les modèles, il présente les avantages de pouvoir localiser les sources d'erreur en espace et d'être robuste vis-à-vis des données bruitées. De plus, la stratégie de recalage associée, qui se base sur une localisation des zones les plus erronées et un recalage hiérarchique, introduit naturellement un processus de régularisation.

On se restreint au cas de l'élasticité linéaire, les formes a et l étant définies comme dans (2.5), et l'opérateur de Hooke $\mathbf{K}$ (potentiellement hétérogène) dépendant de $\mathbf{p}$. La mesure définie par l'ERC s'écrit donc pour tout $\mathbf{p} \in \mathcal{P}$ et tout couple admissible $(\mathbf{v}, \pi) \in \mathcal{U} \times \mathcal{S}$:

$$e_{ERC}^2(\mathbf{p}, \mathbf{v}, \pi) = \frac{1}{2} \|\pi - \mathbf{K}(\mathbf{p})\varepsilon(\mathbf{v})\|_{K^{-1}}^2 \quad (4.3)$$

Dans les premières formulations de recalage utilisant l'ERC, les données expérimentales $\mathbf{s}_{obs}$ étaient incluses comme conditions limites dans la notion d'admissibilité, définissant les nouveaux espaces $\overline{\mathcal{U}}$ et $\overline{\mathcal{S}}$ de champs admissibles et menant à la minimisation de la fonction coût suivante :

$$\mathcal{F}(\mathbf{p}) = \min_{(\mathbf{v}, \pi) \in \overline{\mathcal{U}} \times \overline{\mathcal{S}}} e_{ERC}^2(\mathbf{p}, \mathbf{v}, \pi) \quad (4.4)$$

Cependant, les applications ont montré que cette formulation, qui donne un poids prépondérant aux mesures, était mal adaptée au cas de données bruitées ou corrompues et nécessitait une régularisation.

Dans une version plus récente de la formulation, les conditions limites liées aux données expérimentales ont été relaxées au même titre que le comportement du matériau. Celles-ci sont donc à présent imposées par pénalisation dans une nouvelle mesure de l'ERC appelée erreur en relation de comportement modifiée (ERCm), définie par :

$$e_{ERCm}^2(\mathbf{p}, \mathbf{v}, \pi) = e_{ERC}^2(\mathbf{p}, \mathbf{v}, \pi) + \frac{1}{2} \frac{r}{1-r} \|\mathbf{s}(\mathbf{v}) - \mathbf{s}_{obs}\|_{L^2}^2 \quad (4.5)$$

et associée aux espaces admissibles initiaux $\mathcal{U}$ et $\mathcal{S}$. Les deux termes composant e_{ERCm}^2 peuvent être respectivement interprétés comme des évaluations de l'erreur de modèle et de l'erreur de mesure, $r \in [0, 1]$ étant un paramètre scalaire qui permet de moduler l'influence de ces deux termes. Une caractéristique majeure de l'ERCm est qu'elle permet de déterminer : (i) les régions où les paramètres du modèle sont les plus erronés ; (ii) les éventuels capteurs défectueux. La philosophie de l'ERCm est donc de privilégier les informations théoriques et expérimentales fiables (équilibre, position des capteurs, ...) par rapport aux autres informations (relation de comportement, mesure des capteurs). La nouvelle fonction coût à minimiser dans le processus de recalage est :

$$\mathcal{F}(\mathbf{p}) = \min_{(\mathbf{v}, \pi) \in \mathcal{U} \times \mathcal{S}} e_{ERCm}^2(\mathbf{p}, \mathbf{v}, \pi) \quad (4.6)$$

Par simplicité de notations, la stratégie de minimisation associée à (4.6) est par la suite expliquée à partir de la version discrétisée du problème. L'ERCm discrétisée,

notée e_{ERCm}^h , s'écrit sous la forme :

$$e_{ERCm}^{h2}(\mathbf{p}, \mathbf{U}, \mathbf{V}) = \frac{1}{2}(\mathbf{U} - \mathbf{V})^T \mathbb{K}(\mathbf{p})(\mathbf{U} - \mathbf{V}) + \frac{1}{2} \frac{r}{1-r} (\mathbf{U} - \mathbf{U}_{obs})^T \mathbb{G}_u(\mathbf{U} - \mathbf{U}_{obs}) \quad (4.7)$$

Le vecteur $\mathbf{U} \in \mathcal{U}^h$ est cinématiquement admissible au sens où il vérifie les contraintes cinématiques au sens des EF, tandis que $\mathbf{V} \in \mathcal{S}^h$ est statiquement admissible au sens où il vérifie les équations d'équilibre au sens des EF, i.e. un système de la forme $\mathbb{K}(\mathbf{p})\mathbf{V} = \mathbf{F}$ avec $\mathbb{K}(\mathbf{p})$ la matrice de rigidité globale du problème. $\mathbb{G}_u(\mathbf{p})$ est une matrice de mise à l'échelle qui est en pratique choisie pour que les termes d'erreur de modèle et de mesure puissent être comparés.

La première étape de minimisation impliquée dans (4.6), i.e. le calcul de $\mathcal{F}^h(\mathbf{p})$ tel que :

$$\mathcal{F}^h(\mathbf{p}) = \min_{(\mathbf{U}, \mathbf{V}) \in \mathcal{U}^h \times \mathcal{S}^h} e_{ERCm}^{h2}(\mathbf{p}, \mathbf{U}, \mathbf{V}) \quad (4.8)$$

est appelée *étape de localisation*. Ce problème de minimisation sous contrainte est résolu en recherchant, à $\mathbf{p}$ fixé, le point-selle du lagrangien suivant :

$$\mathcal{L}(\mathbf{p}, \mathbf{U}, \mathbf{V}, \mathbf{\Lambda}) = e_{ERCm}^{h2}(\mathbf{p}, \mathbf{U}, \mathbf{V}) - \mathbf{\Lambda}^T (\mathbb{K}(\mathbf{p})\mathbf{V} - \mathbf{F}) \quad (4.9)$$

Il mène au système d'équations :

$$\begin{aligned} \mathbb{K}(\mathbf{p})(\mathbf{U} - \mathbf{V}) + \frac{r}{1-r} \mathbb{G}_u(\mathbf{U} - \mathbf{U}_{obs}) &= \mathbf{0} \\ \mathbb{K}(\mathbf{p})(\mathbf{U} - \mathbf{V}) + \mathbb{K}(\mathbf{p})\mathbf{\Lambda} &= \mathbf{0} \\ \mathbb{K}(\mathbf{p})\mathbf{V} - \mathbf{F} &= \mathbf{0} \end{aligned} \quad (4.10)$$

la seconde équation correspondant à la version discrétisée d'un problème adjoint. On en déduit que la solution $\mathbf{U}$ vérifie :

$$\overline{\mathbb{K}}(\mathbf{p})\mathbf{U} = \overline{\mathbf{F}} \quad \text{avec} \quad \begin{cases} \overline{\mathbb{K}}(\mathbf{p}) &= \mathbb{K}(\mathbf{p}) + \frac{r}{1-r} \mathbb{G}_u \\ \overline{\mathbf{F}} &= \frac{r}{1-r} \mathbb{G}_u \mathbf{U}_{obs} + \mathbf{F} \end{cases} \quad (4.11)$$

La séparation en contributions spatiales (associées aux différents éléments p_i de $\mathbf{p}$) du terme d'erreur de modèle $e_{ERCm}^{h2}(\mathbf{p}, \mathbf{U}, \mathbf{V})$ obtenu permet alors de localiser les paramètres qui contribuent le plus à l'erreur. La séparation du terme d'erreur de mesure permet quant à lui de détecter les capteurs défectueux.

La seconde étape de minimisation impliquée dans (4.6), i.e. le calcul de $\mathbf{p}_{sol}$ tel que :

$$\mathbf{p}_{sol} = \underset{\mathbf{p} \in \mathcal{P}}{\operatorname{argmin}} \mathcal{F}^h(\mathbf{p}) \quad (4.12)$$

est appelée *étape de correction*. En pratique, $\mathcal{F}^h(\mathbf{p})$ est minimisé uniquement par rapport aux paramètres p_i sélectionnés dans l'étape précédente. Cette minimisation non-linéaire nécessite l'emploi d'un algorithme itératif (méthode de gradient à pas optimal par exemple).

Le gradient de $\mathcal{F}^h(\mathbf{p})$ en fonction de p_i peut être calculé explicitement à partir de

l'information disponible. En effet, pour le couple $(\mathbf{U}, \mathbf{V})$ solution du problème de minimisation sous contraintes (4.8), et avec $\mathbf{\Lambda} = \mathbf{V} - \mathbf{U}$, on a (cf. (1.22)) :

$$\begin{aligned} d_{\mathbf{p}}\mathcal{F}^h(\mathbf{p}; \delta\mathbf{p}) &= \mathcal{L}'_{\mathbf{p}}(\mathbf{p}, \mathbf{U}, \mathbf{V}, \mathbf{\Lambda}; \delta\mathbf{p}) \\ &= \frac{1}{2}(\mathbf{U} - \mathbf{V})^T d_{\mathbf{p}}\mathbb{K}(\mathbf{p}; \delta\mathbf{p})(\mathbf{U} + \mathbf{V}) \end{aligned} \quad (4.13)$$

Bien entendu, ce gradient s'annule lorsque $\mathbf{U} = \mathbf{V}$.

Le processus itératif de recalage incluant les étapes de localisation et correction n'est pas mené jusqu'à convergence ; il est en pratique stoppé lorsque la fonction coût $\mathcal{F}^h$ atteint une valeur seuil donnée (ou un plateau). Ce recalage hiérarchique régularise naturellement la formulation inverse. D'autre part, pour une valeur seuil donnée, il existe généralement un domaine de valeurs possibles pour $\mathbf{p}$.

2 Vers la validation en temps réel

Durant la dernière décennie, les progrès faits en modélisation, analyse numérique et capacités de calcul ont fait émerger l'idée de simulation en temps réel. Cette idée est aujourd'hui portée par le concept de *Dynamic Data Driven Application System* (DDDAS) [DDD 2006], qui définit une nouvelle façon d'appréhender les systèmes dans laquelle les simulations et les mesures *in situ* interagissent profondément et continûment pour former un processus de contrôle et de pilotage (Figure 4.1). Ce concept de contrôle en temps réel, en rupture avec la simulation traditionnelle, est en plein développement dans divers secteurs de pointe comme l'usinage ou la chirurgie (voir [Feng et al. 2009, Fuentes et al. 2009]).

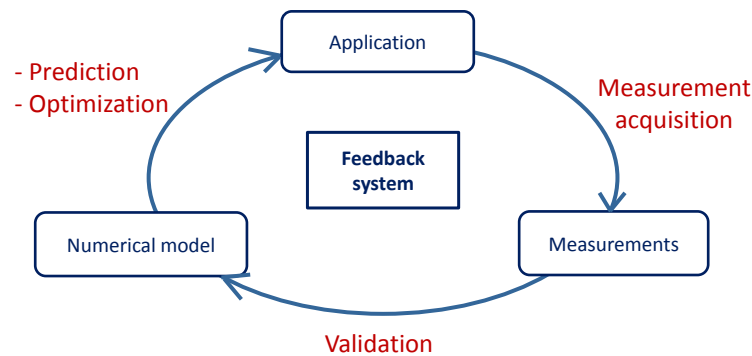


Figure 4.1: Schématisation du concept DDDAS formant un système avec retour d'information

Le recalage de modèle en temps réel se basant sur la notion d'ERCm a été considéré dans [Bouclier et al. 2013], pour une application d'usinage. Dans ce cadre, le concept DDDAS est très bénéfique car l'encastrement de la pièce sur la machine-outil, la position de l'outil, les forces de coupe et la modification de la géométrie sont autant de paramètres qui mènent à des déformations de la pièce variant avec le

temps. Comme ces déformations doivent être prises en compte afin d'obtenir la géométrie finale désirée, une solution apparaît avec le concept DDDAS pour contrôler le système d'usinage et ses évolutions : un modèle numérique simplifié de la pièce et de ses connections avec le bâti peut être introduit et recalé en temps réel à partir de mesures *in situ*, afin de prédire les déformations et corriger la trajectoire de l'outil si besoin. Un point important est que le modèle numérique n'a pas besoin d'être proche de la réalité pendant toute la durée de l'usinage ; il suffit qu'une fois calibré à l'instant t , il puisse refléter correctement la réponse du système réel à cet instant t .

Un challenge majeur est de pouvoir recalé le modèle en temps réel, i.e. pendant l'intervalle de temps entre deux mesures consécutives. Pour cela, une alternative aux calculs coûteux des méthodes itératives de recalage peut être apportée par les méthodes de réduction de modèle. Un processus de recalage utilisant l'ERCm sur un modèle réduit construit par la PGD a donc été proposé dans [Bouclier *et al.* 2013], à partir d'un modèle linéaire et en supposant la position des mesures connue. L'idée de base est d'établir, dans une phase *offline* préliminaire, un méta-modèle PGD à la fois léger et dépendant explicitement des paramètres variables du système. La minimisation de la fonction coût dans la phase *online* de recalage est alors analytique et mène à des calculs itératifs simples et rapides.

2.1 Construction d'un modèle réduit de processus d'usinage

On considère une structure élastique Ω similaire à celle de la Section 1.1.1 du Chapitre 2 et soumise à un chargement quasi-statique (Figure 4.2). Outre les champs donnés $\mathbf{u}_d$, $\mathbf{F}_d$ et $\mathbf{f}_d$ appliqués respectivement sur $\partial_1\Omega$, $\partial_2\Omega$ et Ω , on suppose que la structure Ω (pièce à usiner) est reliée au bâti (E) (machine-outil) par l'intermédiaire de n_k connexions. Celles-ci sont modélisées par des appuis élastiques de raideur $\mathbb{k}_l$ ($l = 1, \dots, n_k$) et définis sur $\Gamma_l \subset \partial\Omega$, avec $\Gamma = \bigcup_{l=1}^{n_k} \Gamma_l$; on suppose que les frontières $\partial_1\Omega$, $\partial_2\Omega$, et Γ sont non-recouvrantes et telles que $\partial_1\Omega \cup \partial_2\Omega \cup \Gamma = \partial\Omega$.

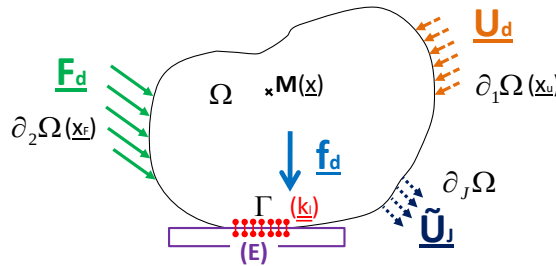


Figure 4.2: Description schématique du problème de référence

La position d'un point M dans la structure est notée $\mathbf{x}$, tandis que les vecteurs $\mathbf{x}_u$ et $\mathbf{x}_F$ sont introduits pour caractériser la position des frontières $\partial_1\Omega$ et $\partial_2\Omega$, respectivement, durant le processus d'usinage. On se restreint par simplicité au cas où les chargements $\mathbf{u}_d$, $\mathbf{F}_d$ et $\mathbf{f}_d$ sont connus *a priori*. L'objectif est alors de calibrer en temps réel les caractéristiques $\mathbb{k}_l$ (tenseurs du second ordre) des liaisons avec le bâti, à partir de mesures expérimentales i.e. les déplacements $\tilde{\mathbf{u}}_J$ de la frontière

$\partial_f \Omega \subset \partial \Omega$. On espère que les variations de $\mathbf{p}$ et les comportements complexes des connexions sur Γ peuvent être prises en compte, à tout instant, en modifiant les comportements linéaires supposés sur Γ .

La calibration du modèle doit être faite tout au long du processus d'usinage, i.e. pour toutes les positions $\mathbf{x}_u$ et $\mathbf{x}_F$ atteintes durant le processus. Par conséquent, les paramètres variables du problème multidimensionnel sont la position spatiale $\mathbf{x}$, les raideurs $(\mathbb{k}_l)_{l=1..n_k}$ des supports élastiques, et les positions $\mathbf{x}_u$ et $\mathbf{x}_F$ des conditions externes. Les espaces associés sont notés Ω , $K_l = [k_{lmin}, k_{lmax}]^{d^2}$ ($l = 1, \dots, n_k$), X_u , et X_F , respectivement. Les solutions en déplacement et contrainte (en tout point de $\Omega \times (\bigotimes_{l=1}^{n_k} K_l) \times X_u \times X_F$) doivent vérifier, en plus de (2.1), (2.2) et (2.3), les relations de comportement sur Γ :

$$\sigma \mathbf{n} = \mathbb{k}_l \mathbf{u} \quad \text{sur } \Gamma_l \quad \forall l = 1 \dots n_k \quad (4.14)$$

La formulation faible en espace, pour tout $(\mathbb{k}_l)_{l=1..n_k}, \mathbf{x}_u, \mathbf{x}_F \in (\bigotimes_{l=1}^{n_k} K_l) \times X_u \times X_F$, est alors similaire à (4.1) avec :

$$a(\mathbf{u}, \mathbf{v}) = \int_{\Omega} \varepsilon(\mathbf{u}) : \mathbf{K} \varepsilon(\mathbf{v}) d\Omega + \sum_{l=1}^{n_k} \int_{\Gamma_l} \mathbb{k}_l \mathbf{u} \cdot \mathbf{v} d\Gamma_l \quad (4.15)$$

En définissant les espaces fonctionnels κ_l , $\mathcal{X}_u$ et $\mathcal{X}_F$ associés respectivement à K_l , X_u et X_F , la formulation faible complète du problème consiste à trouver $\mathbf{u} \in \mathcal{U} \otimes (\bigotimes_{l=1}^{n_k} \kappa_l) \otimes \mathcal{X}_u \otimes \mathcal{X}_F$ tel que :

$$B(\mathbf{u}, \mathbf{v}) = L(\mathbf{v}) \quad \forall \mathbf{v} \in \mathcal{U}_0 \otimes \left(\bigotimes_{l=1}^{n_k} \kappa_l \right) \otimes \mathcal{X}_u \otimes \mathcal{X}_F \quad (4.16)$$

avec

$$\begin{aligned} B(\mathbf{u}, \mathbf{v}) &= \int_{X_F} \int_{X_u} \int_{\bigotimes_{l=1}^{n_k} K_l} a(\mathbf{u}, \mathbf{v}) (d\mathbb{k}_l)_{l=1..n_k} dX_u dX_F \\ L(\mathbf{v}) &= \int_{X_F} \int_{X_u} \int_{\bigotimes_{l=1}^{n_k} K_l} l(\mathbf{v}) (d\mathbb{k}_l)_{l=1..n_k} dX_u dX_F \end{aligned} \quad (4.17)$$

On construit un méta-modèle PGD de (4.16), dépendant explicitement des paramètres $\mathbf{x}$, $(\mathbb{k}_l)_{l=1..n_k}$, $\mathbf{x}_u$ et $\mathbf{x}_F$, avec la même technique que celle exposée dans la Section 3 du Chapitre 3. L'approximation PGD de $\mathbf{u}$, à variables séparées et à l'ordre n , s'écrit :

$$\begin{aligned} \mathbf{u}(\mathbf{x}, (\mathbb{k}_l)_{l=1..n_k}, \mathbf{x}_u, \mathbf{x}_F) &\approx \mathbf{u}_n(\mathbf{x}, (\mathbb{k}_l)_{l=1..n_k}, \mathbf{x}_u, \mathbf{x}_F) \\ &= \sum_{i=1}^n \psi_i(\mathbf{x}) \left(\prod_{l=1}^{n_k} \lambda_{li}(\mathbb{k}_l) \right) \phi_{ui}(\mathbf{x}_u) \phi_{Fi}(\mathbf{x}_F) \end{aligned} \quad (4.18)$$

Les fonctions d'espace ψ_n peuvent être orthonormées pour améliorer l'approximation.

Afin d'illustrer simplement l'implémentation numérique des outils présentés dans cette section, on choisit tout d'abord un exemple académique qui est une version 1D

du problème décrit sur la Figure 4.2. Il s'agit d'une barre élastique de caractéristiques E (module de Young), S (section) et L (longueur) soumise à une force de traction F en $x = X \in [0, L]$ (Figure 4.3). La barre est fixée au bâti par des supports élastiques de raideur k_1 (resp. k_2) en $x = 0$ (resp. $x = L$). Le problème est alors à 4 dimensions, avec les paramètres variables x , k_1 , k_2 , X et les espaces correspondants $[0, L]$, $K_1 = [k_{1min}, k_{1max}]$, $K_2 = [k_{2min}, k_{2max}]$, X_F .

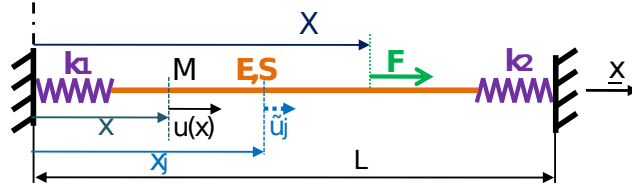


Figure 4.3: Problème de traction/compression 1D

La représentation PGD du déplacement 1D vient naturellement :

$$u(x, k_1, k_2, X) \approx u_n(x, k_1, k_2, X) = \sum_{i=1}^n \psi_i(x) \lambda_{1i}(k_1) \lambda_{2i}(k_2) \phi_i(X) \quad (4.19)$$

En pratique, les fonctions ψ (respectivement λ_1 , λ_2 et ϕ) sont approximées numériquement par les fonctions notées ψ^{hx} (resp. $\lambda_1^{hk_1}$, $\lambda_2^{hk_2}$ et ϕ^{hX}) se mettant sous la forme $\psi^{hx}(x) = \Phi^T(x)\psi$ (resp. $\lambda_1^{hk_1}(k_1) = \mathbf{N}_1^T(k_1)\lambda_1$, $\lambda_2^{hk_2}(k_2) = \mathbf{N}_2^T(k_2)\lambda_2$ et $\phi^{hX}(X) = \bar{\Phi}^T(X)\phi$) où :

- $\Phi(x)$ (resp. $\mathbf{N}_1(k_1)$, $\mathbf{N}_2(k_2)$ et $\bar{\Phi}^T(X)$) est le vecteur colonne des fonctions de forme associées au maillage spatial (resp. aux maillages des domaines des paramètres k_1 , k_2 , et X) ;
- $\psi, \lambda_1, \lambda_2$ et ϕ sont les vecteurs des inconnues nodales.

La version discrétisée des problèmes liés au calcul des modes PGD par la méthode de Galerkin progressive (voir Section 3.1 du Chapitre 3) fait apparaître des matrices de structure usuellement rencontrées dans les codes EF :

- des matrices de raideur classiques pour le paramètre spatial x , à la fois pour la barre élastique et les supports élastiques en $x = 0$ et $x = L$:

$$\mathbb{K} = \int_0^L ES \Phi^T(x) \Phi'(x) dx \quad ; \quad \mathbb{K}_0 = \Phi^T(0) \Phi(0) \quad ; \quad \mathbb{K}_L = \Phi^T(L) \Phi(L) \quad (4.20)$$

- des matrices de masse avec une distribution de masse unitaire pour X , k_1 et k_2 :

$$\begin{aligned} \mathbb{M}_{\bar{\Phi}} &= \int_{X_{0L}} \bar{\Phi}^T(X) \bar{\Phi}(X) dX \quad ; \quad \mathbb{M}_{\Phi\bar{\Phi}} = \int_{X_{0L}} \Phi^T(X) \bar{\Phi}(X) dX \\ \mathbb{M}_{1N} &= \int_{K_1} \mathbf{N}_1^T(k_1) \mathbf{N}_1(k_1) dk_1 \quad ; \quad \mathbb{M}_{2N} = \int_{K_2} \mathbf{N}_2^T(k_2) \mathbf{N}_2(k_2) dk_2 \end{aligned} \quad (4.21)$$

- des matrices de masse avec une distribution de masse linéaire pour les paramètres structuraux :

$$\mathbb{N}_{1N} = \int_{K_1} \mathbf{N}_1^T(k_1) \mathbf{N}_1(k_1) dk_1 \quad ; \quad \mathbb{N}_{2N} = \int_{K_2} \mathbf{N}_2^T(k_2) \mathbf{N}_2(k_2) dk_2 \quad (4.22)$$

– des chargements distribués par vecteur unitaire pour k_1 et k_2 :

$$\int_{K_1} \mathbf{N}_1^T(k_1) dk_1 \quad ; \quad \int_{K_2} \mathbf{N}_2^T(k_2) dk_2 \quad (4.23)$$

La solution PGD (4.19) est cinématiquement admissible, les fonctions ψ^{hx} vérifiant les contraintes cinématiques.

2.2 Contrôle en temps réel par recalage du modèle réduit

L'objectif est à présent de calibrer les raideurs $\mathbb{k}_l$ à partir des mesures expérimentales des déplacements $\mathbf{u}_{obs}^{(j)}$ en $\mathbf{x}_j$ ($j = 1, \dots, n_{obs}$), $\mathbf{x}_j$ étant la coordonnée du point de mesure et n_{obs} le nombre de mesures considérées. La méthode de recalage utilisée est basée sur l'ERCm (Section 1) ; la mesure d'ERC globale s'écrit ici :

$$\begin{aligned} e_{ERC}^2((\mathbb{k}_l)_{l=1..n_k}, \mathbf{v}, \pi) &= \frac{1}{2} \int_{\Omega} (\pi - \mathbf{K}\varepsilon(\mathbf{v})) : \mathbf{K}^{-1}(\pi - \mathbf{K}\varepsilon(\mathbf{v})) d\Omega \\ &+ \frac{1}{2} \sum_{l=1}^{n_k} \int_{\Gamma_l} (\pi \mathbf{n} - \mathbb{k}_l \mathbf{v}) \cdot \mathbb{k}_l^{-1}(\pi \mathbf{n} - \mathbb{k}_l \mathbf{v}) d\Gamma_l \end{aligned} \quad (4.24)$$

La fonctionnelle ERCm à considérer dans la fonction coût (4.6) est alors :

$$e_{ERCm}^2((\mathbb{k}_l)_{l=1..n_k}, \mathbf{v}, \pi) = e_{ERC}^2((\mathbb{k}_l)_{l=1..n_k}, \mathbf{v}, \pi) + \frac{1}{2} \frac{r}{1-r} \sum_{j=1}^{n_{obs}} \|\mathbf{u}(\mathbf{x}_j) - \mathbf{u}_{obs}^{(j)}\|_{L^2}^2 \quad (4.25)$$

Comme décrit précédemment, le problème de calibration qui consiste à trouver le jeu de paramètres $(\mathbb{k}_l)_{l=1..n_k}$ minimisant e_{ERCm}^2 est résolu à partir d'une séquence d'itérations avec deux principales étapes :

1. dans l'étape de localisation, les régions où le modèle est le plus erroné sont identifiées à partir d'une localisation du terme d'erreur de modèle e_{ERC}^2 ;
2. dans l'étape de correction, les paramètres structuraux sont corrigés (par minimisation non-linéaire) pour réduire autant que possible e_{ERCm}^2 .

Sans ajout de technique particulière, la stratégie de recalage précédente est coûteuse car l'étape de correction nécessite la résolution de nombreux problèmes linéaires avec réévaluation de la fonction coût ; les champs $\mathbf{U}$ et $\mathbf{V}$ (4.7) doivent être recalculés à chaque itération de la méthode itérative employée dans cette seconde étape. Ce point peut être grandement amélioré en utilisant un méta-modèle PGD dans lequel la solution dépend analytiquement des paramètres structuraux.

L'adaptation de la méthode de recalage aux modèles réduits PGD apparaît donc à deux niveaux. Le premier est lié à la construction *a priori* de l'écart entre le modèle réduit et les données expérimentales en utilisant l'ERCm, qui nécessite d'associer un méta-modèle PGD au champ cinématique $\mathbf{U}$ du problème (4.11) impliquant les mesures. Le second concerne le calcul analytique des différents termes utilisés dans la stratégie itérative pour réaliser l'étape de correction très rapidement et atteindre la contrainte de temps réel.

2.2.1 Etape de localisation

On explique ici la méthode utilisée pour construire la mesure d'ERCm dans le cas d'un modèle PGD. En reprenant pour illustration le problème 1D de la Figure 4.3, la matrice de raideur globale $\mathbb{K}_T$, la matrice $\mathbb{G}_u$ et le vecteur des mesures $\mathbf{U}_{obs}$ dans (4.7) s'écrivent respectivement :

$$\mathbb{K}_T = \mathbb{K} + k_1 \mathbb{K}_0 + k_2 \mathbb{K}_L \quad ; \quad \mathbb{G}_u = \begin{bmatrix} \cdots & 0 & 0 \\ 0 & \alpha_j & 0 \\ 0 & 0 & \cdots \end{bmatrix} \quad ; \quad \mathbf{U}_{obs} = \begin{pmatrix} \cdots \\ u_{obs}^{(j)} \\ \cdots \end{pmatrix} \quad (4.26)$$

où α_j est un coefficient de pondération ; on choisit dans la suite de donner la même importance à chaque mesure i.e. $\alpha_j = \alpha = \frac{ES}{L}$ ($j = 1, \dots, n_{obs}$). On prend de plus $r = 0,5$ (même poids pour les erreurs de modèle et de mesure).

On dérive alors une représentation PGD de $\mathbf{U}$ et $\mathbf{V}$ de la forme suivante :

$$\begin{aligned} \mathbf{U}(k_1, k_2, X) &\approx \mathbf{U}_n = \sum_{i=1}^n \lambda_{1i}^u(k_1) \lambda_{2i}^u(k_2) \phi_i^u(X) \Psi_i^u \\ \mathbf{V}(k_1, k_2, X) &\approx \mathbf{V}_n = \sum_{i=1}^n \lambda_{1i}^v(k_1) \lambda_{2i}^v(k_2) \phi_i^v(X) \Psi_i^v \end{aligned} \quad (4.27)$$

Le méta-modèle PGD associé à $\mathbf{V}$ est directement issu de (4.19). Le méta-modèle PGD associé à $\mathbf{U}$, qui implique le terme de mesure, est quant à lui construit en utilisant la linéarité du problème (4.11). En notant $\mathbf{U}_0$ et $\mathbf{U}_j$ ($j = 1, \dots, n_{obs}$) les solutions des systèmes suivants :

$$(\mathbb{K}_t + \frac{r}{1-r} \mathbb{G}_u) \mathbf{U}_0 = \mathbf{F} \quad ; \quad (\mathbb{K}_t + \frac{r}{1-r} \mathbb{G}_u) \mathbf{U}_j = \frac{r}{1-r} \mathbb{G}_u \mathbf{1}_j \quad (4.28)$$

avec $\mathbf{1}_j$ le vecteur dont chaque composante est nulle exceptée celle située à la position j qui vaut 1, la solution $\mathbf{U}$ est en effet obtenue par combinaison linéaire :

$$\mathbf{U} = \mathbf{U}_0 + \sum_{j=1}^{n_{obs}} u_{obs}^{(j)} \mathbf{U}_j \quad (4.29)$$

Par conséquent, en supposant connues les positions où l'information expérimentale est donnée, on peut obtenir le méta-modèle PGD de $\mathbf{U}$ *a priori* à partir du méta-modèle PGD de $\mathbf{U}_0$ et des n_{obs} méta-modèles PGD de $\mathbf{U}_j$. Seule la combinaison linéaire avec les valeurs mesurées $u_{obs}^{(j)}$, dont le coût de calcul est très faible, doit être réalisée en temps réel pendant la calibration. Néanmoins, il est nécessaire de construire, dans une phase *offline*, $n_{obs} + 2$ méta-modèles PGD ce qui implique une quantité d'information non négligeable à stocker.

2.2.2 Etape de correction

L'étape de minimisation de la fonction coût $\mathcal{F}^h$ (cf. (4.12) et (4.13)) par rapport à $\mathbf{p}$ (k_1 et k_2 pour le problème 1D) profite pleinement de la construction des méta-modèles PGD ; elle peut être menée très rapidement en utilisant la dépendance

explicite de $\mathcal{F}^h$ par rapport à $\mathbf{p}$. Dans l'algorithme de minimisation non-linéaire (méthode de gradient à pas optimal avec règle d'Armijo ici), chacun des termes qui apparaît peut en effet être exprimé analytiquement, sans résolution à chaque itération des systèmes (4.11) et (4.10) ; la solution pour un couple (k_1, k_2) donné consiste simplement en l'évaluation des fonctions λ_{1i}^u et λ_{1i}^v (respectivement λ_{2i}^u et λ_{2i}^v) en k_1 (resp. k_2). On montre sur les Figures 4.4 et 4.5 l'évolution de la fonction coût $\mathcal{F}^h(k_1, k_2)$ dans le cas de mesures non-bruitées et bruitées, respectivement.

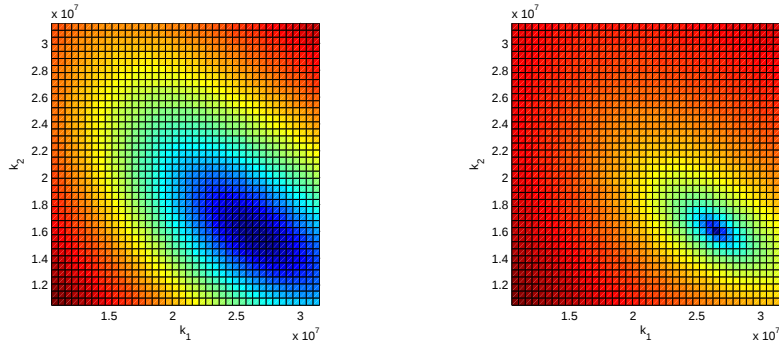


Figure 4.4: Forme de la fonction coût associée au méta-modèle PGD avec des mesures non-bruitées : modèle avec 8 modes (gauche) et avec 48 modes (droite)

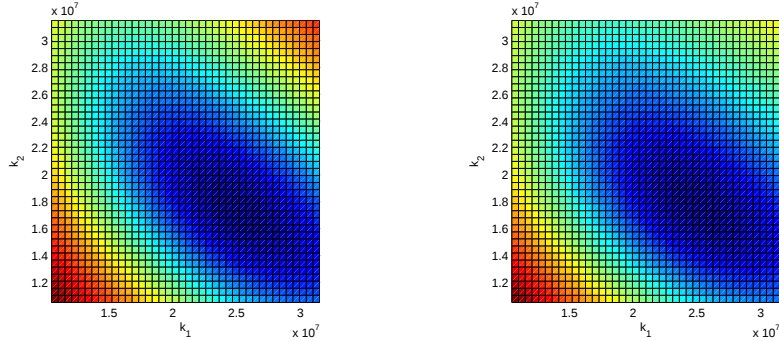


Figure 4.5: Forme de la fonction coût associée au méta-modèle PGD avec des mesures bruitées : modèle avec 8 modes (gauche) et avec 48 modes (droite)

La Figure 4.6 illustre quant à elle la rapidité de la méthode. Elle montre le temps de calcul nécessaire pour la procédure d'optimisation (minimisation non-linéaire), en fonction de la discrétisation spatiale, lorsqu'on utilise le modèle de référence ou un modèle réduit PGD avec 10 modes. On voit clairement que lorsque le nombre d'éléments du maillage augmente, l'utilisation de méta-modèles PGD permet de réduire considérablement le temps de calcul lié au recalage.

2.3 Application

La méthode combinant les outils ERCm et PGD est dans cette section appliquée à une opération de fraisage permettant de réaliser un surfacage sur une pièce (Figure 4.7).

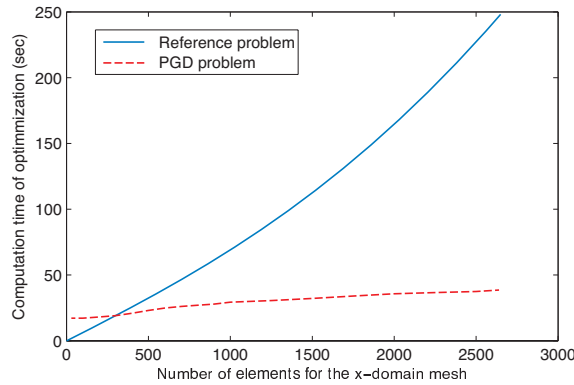


Figure 4.6: Evolution du temps de calcul pour le processus d'optimisation en fonction du nombre d'éléments du maillage spatial

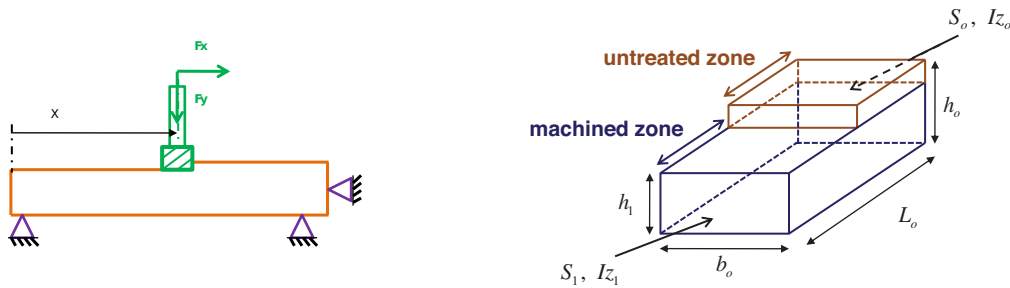


Figure 4.7: Schématisation de l'opération de fraisage sur la pièce

La structure considérée est une poutre en acier (module de Young $E_0 = 210$ GPa, coefficient de Poisson $\nu_0 = 0,3$) de longueur $L_0 = 300$ mm et de section droite rectangulaire (largeur $b_0 = 20$ mm, hauteur $h_0 = 10$ mm avant usinage).

Le modèle numérique de référence utilisé permet de considérer un cas idéal d'usinage. L'opération de fraisage, modélisée par une force verticale de 100 N le long de l'axe vertical, réalise un surfacage de 2 mm de profondeur ; la poutre désirée après usinage est donc de hauteur $h_1 = 8$ mm. Les propriétés géométriques et de rigidité en flexion de la poutre varient durant l'usinage (Figure 4.7) : la section droite (resp. le moment quadratique) vaut S_0 (resp. I_{z_0}) dans la zone saine et S_1 (resp. I_{z_1}) dans la zone usinée. La poutre est fixée au bâti par des brides dont le comportement est supposé élastique et très rigide par rapport à la structure ; la raideur des supports est $k_0 = 300.10^2$ N.m¹ soit environ cent fois supérieure à celle de la poutre en flexion.

Ce modèle de référence, associé à un modèle de flexion de Timoshenko, est représenté sur la Figure 4.8 ; il fournit les mesures simulées.

Le modèle utilisé pour le recalage est quant à lui composé d'une poutre en flexion avec des propriétés géométriques constantes (prises égales à celles avant usinage) et soumise à une force verticale d'amplitude F et de position X variable durant l'usinage (Figure 4.8). La poutre est fixée sur ses bords par deux appuis

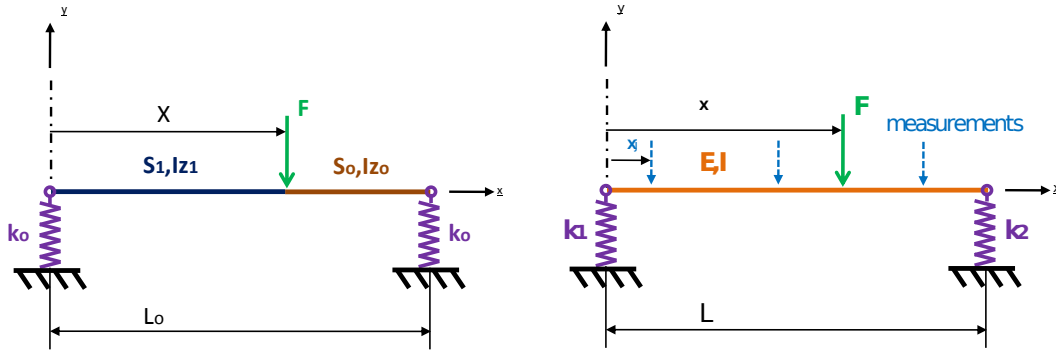


Figure 4.8: Modèle utilisé pour simuler l'usinage idéal (gauche) et modèle utilisé pour le recalage (droite)

élastiques de raideur respective k_1 et k_2 . La valeur de F étant supposée constante et connue, l'objectif est de recalibrer les valeurs de k_1 et k_2 en temps réel à partir de mesures expérimentales qui consistent en des déplacements ponctuels de la pièce réelle aux points x_j . Pour les résultats numériques, on considère 13 points de mesure répartis uniformément le long de la poutre.

Il est important de noter que le but ici n'est pas d'identifier précisément les paramètres de comportement qui affectent la raideur de la structure durant l'usinage, mais de mettre en place un modèle qui reproduit précisément le phénomène de flexion et la réponse en déplacement vertical au niveau de l'outil. Jouer seulement sur les raideurs k_1 et k_2 des supports élastiques pour atteindre cet objectif est un choix, même si ces deux paramètres ne sont pas liés à l'enlèvement de matière.

A chaque instant (i.e. pour chaque position X de l'outil), après recalage du modèle, on peut utiliser celui-ci pour corriger si besoin le déplacement vertical de l'outil afin de garantir l'usinage désiré ; ce protocole de contrôle est schématisé sur la Figure 4.9. Pour respecter la contrainte de temps réel, le recalage suivi de la correction du déplacement de l'outil doivent être faits suffisamment rapidement, i.e. dans l'intervalle de temps entre deux acquisitions de mesure successives. En pratique, en considérant un taux d'avance de l'outil de 2 mm/s et une longueur de poutre $L_0 = 300$ mm, cet intervalle est d'environ 8 s en considérant 19 acquisitions ; le challenge ici est de faire le recalage en 1 s.

Le méta-modèle PGD introduit dépend des paramètres x , k_1 , k_2 et X . Pour le calculer, on choisit des maillages avec 30 éléments pour chaque dimension, les domaines K_1 et K_2 étant pris tels que $K_1 = K_2 = [k_{min}, k_{max}]$ avec $k_{min} = 0,5k_0$ et $k_{max} = 1,5k_0$. On conserve 5 modes dans la représentation PGD utilisée pour le recalage.

La figure 4.10 montre les surfaces usinées obtenues en utilisant ou non la procédure de contrôle en temps réel. On comprend clairement ici l'origine du défaut observé : la force verticale de l'outil fait fléchir la poutre et le phénomène est d'autant plus marqué qu'on est loin des supports (surface en forme de dôme) ; le léger décalage vers la droite est dû au retrait de matière (impliquant une diminution de

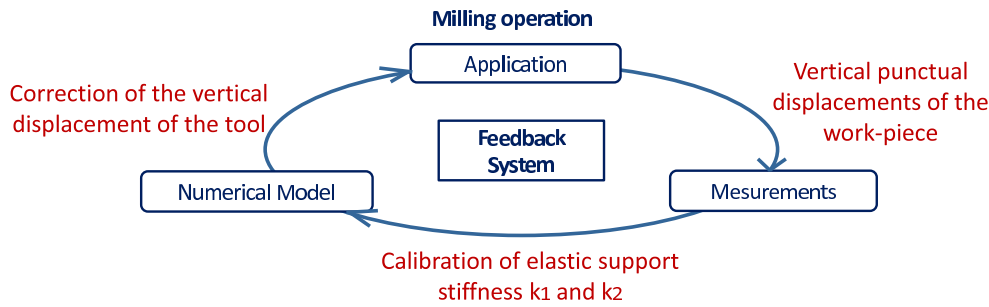


Figure 4.9: Procédure implémentée pour contrôler l'usinage

rigidité) qui commence à gauche. Ce défaut d'usinage diminue fortement quand la procédure en temps réel de recalage et de correction de la position de l'outil est menée ; le défaut est alors divisé par trois en amplitude.

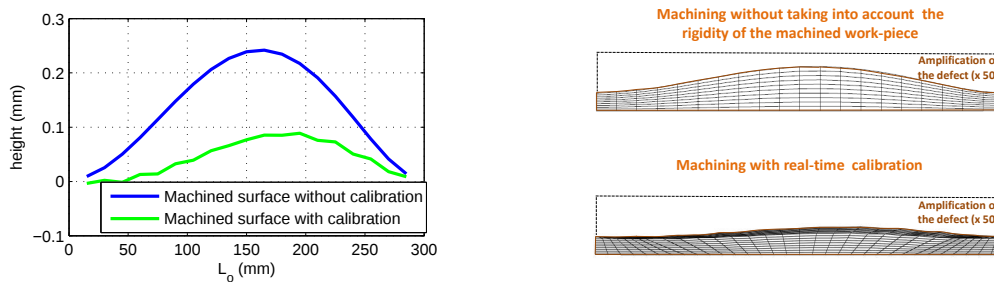


Figure 4.10: Comparaison des structures usinées avec ou sans recalage en temps réel : représentation des surfaces usinées (gauche) et formes du défaut sur la poutre (droite)

3 Recalage dédié à des quantités d'intérêt

Dans la stratégie de recalage par l'ERCm présentée dans la Section 1, le modèle est recalé de façon globale i.e. à partir d'une mesure d'erreur définie sur l'ensemble du domaine spatial Ω et avec l'ensemble $\mathbf{p}$ des paramètres. Néanmoins, lorsque l'objectif d'une simulation est de prédire la valeur d'une quantité d'intérêt Q , cette stratégie n'est pas optimale ; si Q est peu sensible vis-à-vis de certains paramètres, il est probablement inutile de recaler avec précision ces paramètres. L'objectif de cette partie est donc de proposer une approche de recalage de modèle dédiée et partielle, qui permette d'obtenir à un coût raisonnable un modèle précis pour la prédiction d'une quantité d'intérêt donnée. En d'autres termes, à partir d'une quantité d'informations expérimentales donnée, on souhaite que la nouvelle approche puisse exploiter au mieux cette quantité d'informations en vue de l'objectif de la simulation.

Quelques travaux ont déjà abordé le recalage de modèle avec une vision *goal-oriented*. Dans [Becker et Vexler 2004, Johansson et al. 2007] par exemple, un

problème d'optimisation couplé avec une méthode duale a été proposé pour estimer, vis-à-vis d'une quantité d'intérêt, la sensibilité par rapport aux mesures (bruit, incertitudes) et l'erreur de discrétisation affectant la procédure de recalage des paramètres. Dans [Johansson *et al.* 2011], la méthode d'estimation d'erreur *a posteriori* a été étendue pour combiner à la fois le recalage et la simulation finale (ces deux processus étant possiblement gouvernés par différentes équations d'état, et couplés uniquement par l'intermédiaire des paramètres) ; pour une tolérance donnée sur une quantité d'intérêt, dépendant de la solution de la simulation finale, trois sources d'erreur ont été contrôlées : l'erreur de modèle, l'erreur de discrétisation polluant le problème de recalage, et l'erreur de discrétisation liée à la simulation finale.

Dans cette section, le contexte d'étude est celui du recalage sur une structure réelle de grande taille, qui est aussi la structure sur laquelle la simulation finale sera réalisée (avec un autre chargement par exemple). L'objectif de réalisation d'un protocole de recalage optimal par rapport à une quantité d'intérêt Q signifie : (i) recalculer seulement les paramètres pertinents du modèle, i.e. ceux qui affectent Q ; (ii) utiliser des informations expérimentales pertinentes pour lesquelles Q est sensible. Une approche naturelle consisterait à faire une analyse de sensibilité de Q vis-à-vis des paramètres du modèle, comme détaillé dans la Section 3.2 du Chapitre 1. Cependant, les informations obtenues sur les gradients ne sont pas facilement exploitables dans une démarche de recalage, i.e. les choix de paramètres à recalculer et de critère d'arrêt utilisant ces informations sont discutables.

Une procédure de recalage de modèle dédiée aux quantités d'intérêt a été proposée dans [Chamoin *et al.* 2013]. Elle conserve la philosophie de l'ERCm introduite précédemment mais se distingue de la méthode classique par l'introduction de nouvelles fonctions coût spécifiques à minimiser. Une analyse de sensibilité vis-à-vis des mesures est également introduite dans [Chamoin *et al.* 2013] ; elle permet de définir un protocole expérimental optimal (en terme de position des capteurs de mesure et de quantités mesurées) dans l'optique de la prédiction d'une quantité d'intérêt avec le modèle recalé. Les performances des outils proposés sont illustrées ici sur des cas simples d'élasticité linéaire et de thermique instationnaire.

3.1 Définition des nouvelles fonctions coût

3.1.1 Cas d'une quantité d'intérêt non mesurée

On suppose dans cette section que la zone spatiale de définition de Q n'est pas un point de mesure. Dans ce cas, en gardant la philosophie et la flexibilité de l'outil ERCm composé d'un terme d'erreur de modèle et d'un terme d'erreur de mesure, on introduit une nouvelle fonction coût notée $\mathcal{F}_Q(\mathbf{p})$ et définie par :

$$\mathcal{F}_Q(\mathbf{p}) = \min_{q \in \mathbb{R}} \left[\frac{1}{2} |q - Q_{mod}(\mathbf{p})|^2 + \frac{1}{2} \frac{r}{1-r} |q - Q_{obs}(\mathbf{p})|^2 \right] \quad (4.30)$$

Le premier terme de (4.30), lié à une erreur de modèle, implique une valeur Q_{mod} de la quantité d'intérêt définie à partir du modèle mathématique uniquement, i.e.

l'équation d'état (4.1) :

$$Q_{mod}(\mathbf{p}) = Q(\mathbf{u}_1(\mathbf{p})) \quad ; \quad A(\mathbf{p}, \mathbf{u}_1, \mathbf{v}) = 0 \quad \forall \mathbf{v} \in \mathcal{U}_0 \quad (4.31)$$

Le second terme de (4.30), lié à une erreur de mesure, implique une valeur Q_{obs} de la quantité d'intérêt définie par "extrapolation" des données expérimentales $\mathbf{s}_{obs}$:

$$Q_{obs}(\mathbf{p}) = Q(\mathbf{u}_2(\mathbf{p})) \quad ; \quad (\mathbf{u}_2, \sigma_2) = \underset{(\mathbf{v}, \pi) \in \mathcal{U} \times \mathcal{S}}{\operatorname{argmin}} e_{ERCm}^2(\mathbf{p}, \mathbf{v}, \pi) \quad (4.32)$$

La détermination de $\mathbf{u}_2$ fait donc appel à $\mathbf{s}_{obs}$ et mène, après discrétisation, à la résolution d'un système de la forme $\overline{\mathbb{K}}(\mathbf{p})\mathbf{U}_2 = \overline{\mathbf{F}}$ (cf. (4.11)).

Dès lors, la première étape dans le processus itératif de recalage consiste à calculer $\mathcal{F}_Q(\mathbf{p})$. Ce calcul étant une minimisation sous contrainte à $\mathbf{p}$ fixé, on introduit le lagrangien (semi-discrétisé) suivant :

$$\begin{aligned} \mathcal{L}_Q(\mathbf{p}, q, \mathbf{u}_1, \mathbf{U}_2, \boldsymbol{\mu}, \boldsymbol{\Lambda}) = & \frac{1}{2}|q - Q(\mathbf{u}_1)|^2 + \frac{1}{2} \frac{r}{1-r} |q - Q(\mathbf{U}_2)|^2 \\ & - A(\mathbf{p}, \mathbf{u}_1, \boldsymbol{\mu}) - \boldsymbol{\Lambda}^T (\overline{\mathbb{K}}(\mathbf{p})\mathbf{U}_2 - \overline{\mathbf{F}}) \end{aligned} \quad (4.33)$$

La recherche du point-selle associé aboutit à vérifier, outre les conditions issues de (4.31) (resp. (4.32)) pour $\mathbf{u}_1$ (resp. $\mathbf{U}_2$) :

$$\begin{aligned} (q - Q(\mathbf{u}_1)) + \frac{r}{1-r} (q - Q(\mathbf{U}_2)) &= 0 \\ Q(\delta \mathbf{u}_1) (q - Q(\mathbf{u}_1)) + A(\mathbf{p}, \delta \mathbf{u}_1, \boldsymbol{\mu}) &= 0 \quad \forall \delta \mathbf{u}_1 \in \mathcal{U}_0 \\ \frac{r}{1-r} Q(\delta \mathbf{U}_2) (q - Q(\mathbf{U}_2)) + \boldsymbol{\Lambda}^T \overline{\mathbb{K}}(\mathbf{p}) \delta \mathbf{U}_2 &= 0 \quad \forall \delta \mathbf{U}_2 \in \mathcal{U}_0^h \end{aligned} \quad (4.34)$$

les deux dernières équations correspondant à des problèmes adjoints. On obtient alors :

$$q = (1-r)Q(\mathbf{u}_1) + rQ(\mathbf{U}_2) \quad ; \quad \mathcal{F}_Q^h(\mathbf{p}) = \frac{1}{2}r [Q(\mathbf{u}_1) - Q(\mathbf{U}_2)]^2 \quad (4.35)$$

ainsi qu'une reformulation des deux problèmes adjoints :

$$\begin{aligned} A(\mathbf{p}, \delta \mathbf{u}_1, \boldsymbol{\mu}) &= \beta Q(\delta \mathbf{u}_1) \quad \forall \delta \mathbf{u}_1 \in \mathcal{U}_0 \\ \boldsymbol{\Lambda}^T \overline{\mathbb{K}}(\mathbf{p}) \delta \mathbf{U}_2 &= -\beta Q(\delta \mathbf{U}_2) \quad \forall \delta \mathbf{U}_2 \in \mathcal{U}_0^h \end{aligned} \quad (4.36)$$

avec $\beta = r[Q(\mathbf{u}_1) - Q(\mathbf{U}_2)]$, et $\mathbf{u}_1$ (resp. $\mathbf{U}_2$) vérifiant (4.31) (resp. (4.32)).

Dans une seconde étape, le gradient de $\mathcal{F}_Q^h(\mathbf{p})$ est calculé par la méthode de l'état adjoint, à partir des champs $(q, \mathbf{u}_1, \mathbf{U}_2, \boldsymbol{\mu}, \boldsymbol{\Lambda})$ obtenus précédemment :

$$\begin{aligned} d_{\mathbf{p}} \mathcal{F}_Q^h(\mathbf{p}; \delta \mathbf{p}) &= \mathcal{L}'_{Q,\mathbf{p}}(\mathbf{p}, q, \mathbf{u}_1, \mathbf{U}_2, \boldsymbol{\mu}, \boldsymbol{\Lambda}; \delta \mathbf{p}) \\ &= -a'_{\mathbf{p}}(\mathbf{p}, \mathbf{u}_1, \boldsymbol{\mu}; \delta \mathbf{p}) - \boldsymbol{\Lambda}^T d_{\mathbf{p}} \overline{\mathbb{K}}(\mathbf{p}; \delta \mathbf{p}) \mathbf{U}_2 \end{aligned} \quad (4.37)$$

Les éléments de $\mathbf{p}$ associés aux plus fortes composantes du gradient de $\mathcal{F}_Q^h(\mathbf{p})$ sont alors sélectionnés pour mener une minimisation non-linéaire hiérarchique, et les

itérations sont poursuivies jusqu'à atteindre une valeur seuil pour $\mathcal{F}_Q^h(\mathbf{p})$.

On illustre la méthode sur le treillis de poutre représenté sur la Figure 4.11. Il est constitué de 10 barres dont les modules de Young respectifs E_i sont des paramètres mal connus qu'il faut recaler à partir des déplacements mesurés aux nœuds du treillis. Une force ponctuelle verticale, parfaitement connue, est imposée au point P , et on suppose que l'objectif de la simulation finale est le déplacement vertical du point P i.e. $Q(\mathbf{u}) = u_y(P)$.

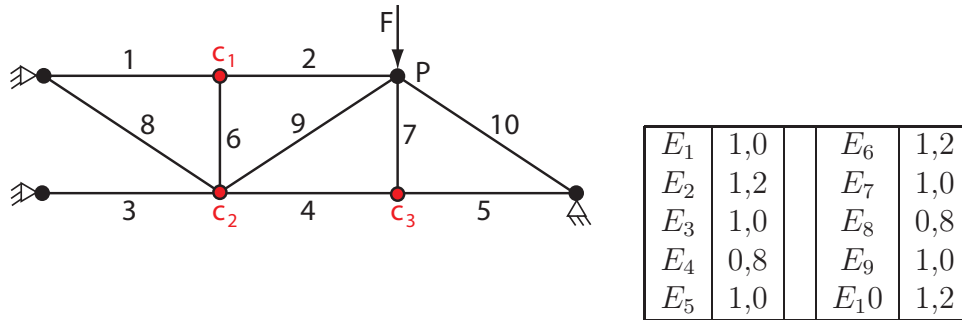


Figure 4.11: Treillis considéré (gauche) et valeurs de référence pour les paramètres des barres (droite)

Les mesures prépondérantes consistent en des déplacements nodaux aux points c_1 , c_2 et c_3 de la structure; leur valeur est simulée en considérant un modèle de référence dans lequel les valeurs des paramètres de raideur E_i sont celles indiquées sur la Figure 4.11. Le modèle initial à recaler est quant à lui obtenu en perturbant les paramètres des barres de 50%, i.e. $E_i^* = 1,5 \times E_i$ ($i = 1, \dots, 10$). Le recalage est mené avec un algorithme de gradient à pas fixe.

La Figure 4.12 montre les valeurs de Q prédites par le modèle après chaque itération du recalage, dans le cas d'un recalage classique (ERCm globale) ou d'un recalage dédié à Q avec la même quantité d'informations expérimentales. On observe que la nouvelle stratégie de recalage fournit, après seulement quelques itérations, un modèle de qualité pour prédire Q . Sur la Figure 4.13, on montre les rapports E_i^*/E_i (paramètres recalés sur paramètres de référence) à la fin de l'itération 10 du recalage, à la fois pour le recalage classique et celui dédié à la quantité d'intérêt. On voit clairement que dans le second cas, la stratégie de recalage privilégie un sous-ensemble de paramètres i.e. ceux les plus influents sur la valeur prédite de Q .

3.1.2 Cas d'une quantité d'intérêt mesurée

Lorsqu'une mesure de la quantité d'intérêt est disponible, il est naturel de l'utiliser directement pour définir Q_{obs} ; dès lors, le processus de régularisation apporté par l'ERCm dans la fonction coût de la Section 3.1.1 n'est plus disponible et une autre fonction coût doit être construite. On propose ici une fonction coût $\mathcal{F}_Q(\mathbf{p})$ basée sur une version locale de l'ERCm i.e. une mesure d'erreur en relation de comportement

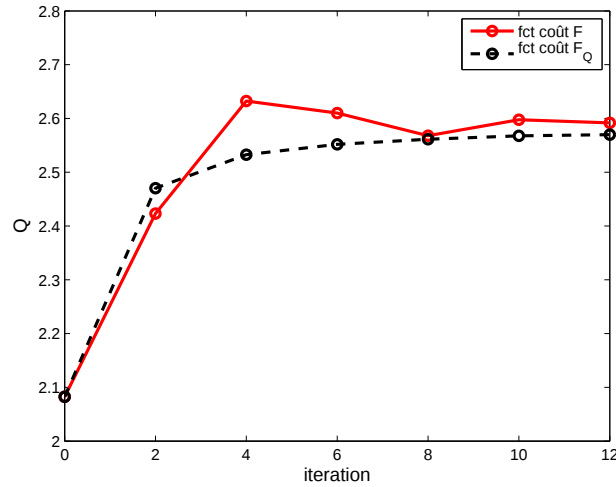


Figure 4.12: Comparaison des valeurs de Q prédites par le modèle après chaque itération du recalage

au point (ou dans la zone) où la quantité d'intérêt est définie. On introduit la quantité Q^* duale de Q , au sens de l'énergie, et on note la relation de comportement locale liant Q^* à Q sous la forme $Q^* = k(\mathbf{p})Q$. On considère alors la fonction coût suivante :

$$\mathcal{F}_Q(\mathbf{p}) = \min_{(\mathbf{v}, \pi) \in \mathcal{U} \times \mathcal{S}} e_{ERCm,loc}^2(\mathbf{p}, \mathbf{v}, \pi) \quad (4.38)$$

avec

$$e_{ERCm,loc}^2(\mathbf{p}, \mathbf{v}, \pi) = \frac{1}{2} |Q^*(\pi) - k(\mathbf{p})Q(\mathbf{v})|^2 + \frac{1}{2} \frac{r}{1-r} |Q(\mathbf{v}) - Q_{obs}|^2 \quad (4.39)$$

Les deux termes composant $e_{ERCm,loc}^2$ représentent des erreurs locales de modèle et de mesure, respectivement. Par ailleurs, en introduisant la solution adjointe $\tilde{\mathbf{u}}$ (au sens de (1.10)) liée à Q , le premier terme peut aussi s'écrire :

$$\frac{1}{2} \left(\int_{\Omega} (\pi - \mathbf{K}(\mathbf{p})\varepsilon(\mathbf{v})) : \varepsilon(\tilde{\mathbf{u}}) d\Omega \right)^2 \quad (4.40)$$

En suivant une démarche similaire à celle présentée dans la Section 1, on écrit la version discrétisée de (4.39) :

$$e_{ERCm,loc}^{h2}(\mathbf{p}, \mathbf{U}, \mathbf{V}) = \frac{1}{2} |Q^*(\mathbf{V}) - k(\mathbf{p})Q(\mathbf{U})|^2 + \frac{1}{2} \frac{r}{1-r} |Q(\mathbf{U}) - Q_{obs}|^2 \quad (4.41)$$

avec $\mathbb{K}(\mathbf{p})\mathbf{V} = \mathbf{F}$ (équilibre), et le calcul de $\mathcal{F}_Q^h(\mathbf{p})$ est fait par recherche du point-selle du lagrangien suivant :

$$\mathcal{L}_Q(\mathbf{p}, \mathbf{U}, \mathbf{V}, \Lambda) = e_{ERCm,loc}^{h2}(\mathbf{p}, \mathbf{U}, \mathbf{V}) - \Lambda^T (\mathbb{K}(\mathbf{p})\mathbf{V} - \mathbf{F}) \quad (4.42)$$

On aboutit alors au système d'équations :

$$\begin{aligned} -k(\mathbf{p}) (Q^*(\mathbf{V}) - k(\mathbf{p})Q(\mathbf{U})) + \frac{r}{1-r} (Q(\mathbf{U}) - Q_{obs}) &= 0 \\ Q^*(\delta\mathbf{V}) (Q^*(\mathbf{V}) - k(\mathbf{p})Q(\mathbf{U})) - \Lambda^T \mathbb{K}(\mathbf{p})\delta\mathbf{V} &= 0 \quad \forall \delta\mathbf{V} \in \mathcal{S}_0^h \\ \mathbb{K}(\mathbf{p})\mathbf{V} - \mathbf{F} &= 0 \end{aligned} \quad (4.43)$$

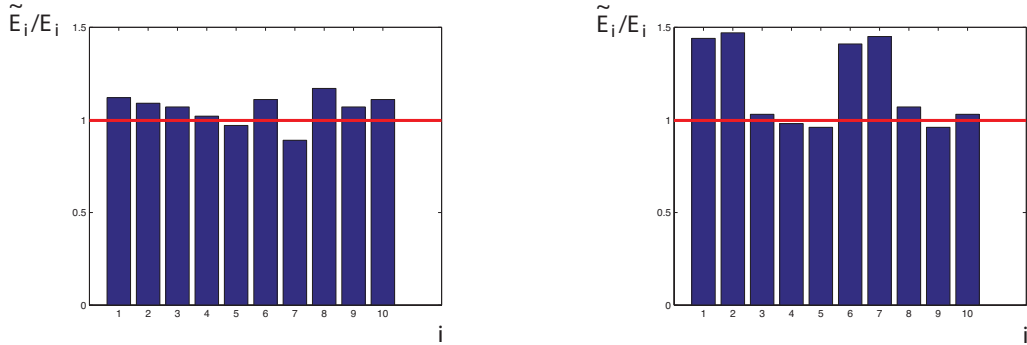


Figure 4.13: Rapports E_i^*/E_i après l'itération 10 du processus de recalage, pour un recalage classique (gauche) ou dédié à Q (droite)

La seconde équation dans (4.43) est la version discrétisée d'un problème adjoint ; la combinaison des deux autres équations donne le système à vérifier par $\mathbf{U}$.

L'évaluation du gradient de $\mathcal{F}_Q^h(\mathbf{p})$ est ensuite directe, avec la méthode de l'état adjoint :

$$\mathbf{d}_{\mathbf{p}}\mathcal{F}^h(\mathbf{p}; \delta\mathbf{p}) = \mathcal{L}'_{\mathbf{p}}(\mathbf{p}, \mathbf{U}, \mathbf{V}, \mathbf{\Lambda}; \delta\mathbf{p}) \quad (4.44)$$

avec $(\mathbf{U}, \mathbf{V}, \mathbf{\Lambda})$ solution de (4.43). Là encore, on sélectionne les éléments de $\mathbf{p}$ associés aux plus fortes composantes du gradient, afin de mener la minimisation de $\mathcal{F}_Q^h(\mathbf{p})$ uniquement par rapport à ces éléments. Les itérations du processus de recalage sont poursuivies jusqu'à atteindre une valeur seuil.

On illustre la méthode sur le même treillis que celui de la section 3.1.1, avec les mêmes valeurs des paramètres de référence E_i , le même chargement et la même quantité d'intérêt, mais en plaçant le capteur de mesure prépondérant au point P (Figure 4.14).

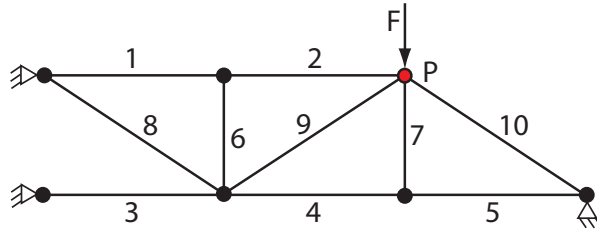


Figure 4.14: Nouvelle configuration du treillis

La Figure 4.15 montre d'une part les valeurs de Q prédites par le modèle après chaque itération du recalage, dans le cas d'un recalage classique ou d'un recalage dédié à Q , et d'autre part les rapports E_i^*/E_i à la fin de l'itération 10 du recalage dédié à Q . On observe à nouveau que la stratégie de recalage mise en place dans cette

section est plus performante que la stratégie classique pour fournir, avec la même quantité d'informations expérimentales, un modèle précis en vue de la prédiction de Q .

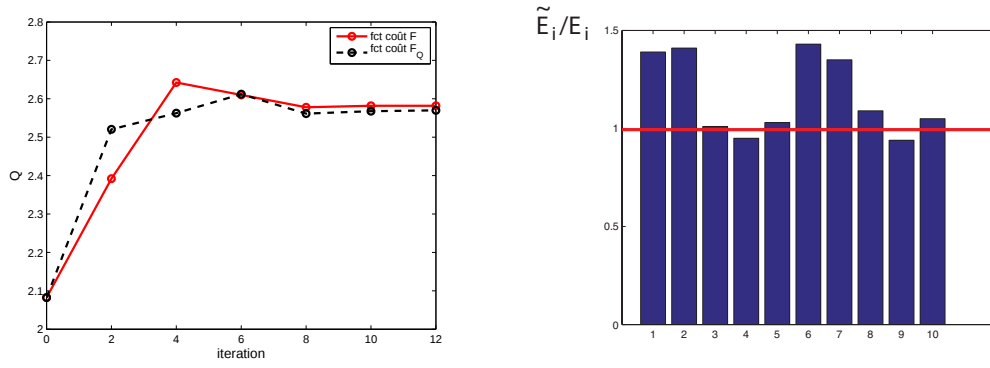


Figure 4.15: Comparaison des valeurs de Q prédites par le modèle après chaque itération du recalage (gauche), et rapports E_i^*/E_i après l'itération 10 du processus de recalage dédié à Q (droite)

3.2 Sensibilité vis-à-vis des mesures

L'objectif de cette section est d'optimiser le protocole expérimental servant au recalage de modèle vis-à-vis de la prédiction d'une quantité d'intérêt. Il y a en pratique deux voies d'optimisation, liées au choix du positionnement des outils de mesure (capteurs) d'une part et au choix des quantités mesurées d'autre part. La démarche naturelle est donc d'estimer la sensibilité de la valeur prédite de Q vis-à-vis des mesures considérées dans le recalage. Cette sensibilité est liée à la relation implicite entre une variation des mesures $\mathbf{s}_{obs}$ et la prédiction pour Q , qui implique le processus de recalage.

Le recalage par l'ERCm présenté dans la Section 1 est, à mesures $\mathbf{s}_{obs}$ données, un problème de minimisation sous contrainte dont la solution correspond au point-selle global du lagrangien suivant :

$$\mathcal{L}(\mathbf{q}, \mathbf{v}, \pi, \boldsymbol{\mu}) = e_{ERCm}^2(\mathbf{q}, \mathbf{v}, \pi) - \left[\int_{\Omega} \pi : \varepsilon(\boldsymbol{\mu}) d\Omega - \int_{\Omega} \mathbf{f}_d \cdot \boldsymbol{\mu} d\Omega - \int_{\partial_2 \Omega} \mathbf{F}_d \cdot \boldsymbol{\mu} dS \right] \quad (4.45)$$

Pour mener l'analyse de sensibilité, on écrit explicitement par la suite la dépendance des grandeurs vis-à-vis des mesures $\mathbf{s}_{obs}$; on note de plus $\mathbf{z} = (\mathbf{p}, \mathbf{u}, \sigma, \boldsymbol{\lambda})$ la solution du problème de recalage à mesures $\mathbf{s}_{obs}$ fixées. Dès lors, la quantité d'intérêt peut s'écrire $Q(\mathbf{z}(\mathbf{s}_{obs}))$ où $\mathbf{z}(\mathbf{s}_{obs})$ est solution du problème de recalage i.e. :

$$\mathcal{L}'_{\mathbf{z}}(\mathbf{z}, \mathbf{s}_{obs}; \delta \mathbf{z}) = 0 \quad \forall \delta \mathbf{z} \quad (4.46)$$

Afin de calculer le gradient de Q par rapport à $\mathbf{s}_{obs}$, on utilise à nouveau la méthode de l'état adjoint ; on introduit pour cela le lagrangien $\mathcal{H}$ défini par :

$$\mathcal{H}(\mathbf{z}, \mathbf{s}_{obs}, \mathbf{y}) = Q(\mathbf{z}) - \mathcal{L}'_{\mathbf{z}}(\mathbf{z}, \mathbf{s}_{obs}; \mathbf{y}) \quad (4.47)$$

Le calcul de la solution adjointe $\mathbf{y}$ issue de la recherche du point-selle de $\mathcal{H}$ et vérifiant :

$$\mathcal{L}''_{\mathbf{z}\mathbf{z}}(\mathbf{z}, \mathbf{s}_{obs}; \mathbf{y}, \delta\mathbf{z}) = Q(\delta\mathbf{z}) \quad \forall \delta\mathbf{z} \quad (4.48)$$

permet d'obtenir le résultat de sensibilité :

$$\begin{aligned} d_{\mathbf{s}_{obs}} Q(\mathbf{z}(\mathbf{s}_{obs}); \delta\mathbf{s}_{obs}) &= \mathcal{H}'_{\mathbf{s}_{obs}}(\mathbf{z}, \mathbf{s}_{obs}, \mathbf{y}; \delta\mathbf{s}_{obs}) \\ &= -\mathcal{L}''_{\mathbf{z}\mathbf{s}_{obs}}(\mathbf{z}, \mathbf{s}_{obs}; \mathbf{y}, \delta\mathbf{s}_{obs}) \end{aligned} \quad (4.49)$$

avec $\mathbf{z}$ (resp. $\mathbf{y}$) vérifiant (4.46) (resp. (4.48)). A noter que la solution adjointe $\mathbf{y}$, à mesures $\mathbf{s}_{obs}$ données, permet aussi d'aborder la sensibilité de Q par rapport à la discrétisation utilisée pour résoudre les problèmes associés au processus de recalage, cette discrétisation pouvant avoir une forte influence sur la valeur des paramètres $\mathbf{p}$ recalés (voir [Becker et Vexler 2004]).

La version discrétisée de l'étude de sensibilité précédente montre que l'effort de calcul est raisonnable. L'écriture discrétisée de (4.45) est similaire à (4.9), et la recherche de son point-selle mène à $\mathbf{p}$ tel que (cf. (4.13)) :

$$B(\mathbf{p}, \mathbf{U}_{obs}; \delta\mathbf{p}) = \frac{1}{2}(\mathbf{U} - \mathbf{V})^T d_{\mathbf{p}} \mathbb{K}(\mathbf{p}; \delta\mathbf{p})(\mathbf{U} + \mathbf{V}) = 0 \quad \forall \delta\mathbf{p} \in \mathcal{P}_0 \quad (4.50)$$

avec $\mathbf{U}$ et $\mathbf{V}$ dépendant de $\mathbf{p}$ et vérifiant (4.10) et (4.11). Après introduction du lagrangien discrétisé :

$$\mathcal{H}(\mathbf{p}, \mathbf{U}_{obs}, \mathbf{q}) = Q(\mathbf{U}(\mathbf{p})) - B(\mathbf{p}, \mathbf{U}_{obs}; \mathbf{q}) \quad (4.51)$$

et résolution du problème adjoint associé :

$$B'_{\mathbf{p}}(\mathbf{p}, \mathbf{U}_{obs}; \mathbf{q}, \delta\mathbf{p}) = Q'_{\mathbf{p}}(\mathbf{U}(\mathbf{p}); \delta\mathbf{p}) \quad \forall \delta\mathbf{p} \in \mathcal{P}_0 \quad (4.52)$$

on obtient directement :

$$d_{\mathbf{U}_{obs}} Q(\mathbf{U}(\mathbf{p}, \mathbf{U}_{obs}); \delta\mathbf{U}_{obs}) = \mathcal{H}'_{\mathbf{U}_{obs}}(\mathbf{p}, \mathbf{U}_{obs}, \mathbf{q}; \delta\mathbf{U}_{obs}) = -B'_{\mathbf{U}_{obs}}(\mathbf{p}, \mathbf{U}_{obs}; \mathbf{q}, \delta\mathbf{U}_{obs}) \quad (4.53)$$

avec $\mathbf{U}$ (resp. $\mathbf{q}$) vérifiant (4.11) (resp. (4.52)).

Le gradient $d_{\mathbf{U}_{obs}} Q$ peut être calculé à chaque itération du processus de recalage à partir de l'information disponible et après résolution de (4.52). L'optimisation du protocole expérimental (position des capteurs, grandeurs mesurées) consiste alors à chercher les extrema de $|s_{obs}^{(i)} d_{s_{obs}^{(i)}} Q|$ ou $|U_{obs}^{(i)} d_{U_{obs}^{(i)}} Q|$ ($i = 1, \dots, n_{obs}$).

En reprenant pour exemple le cas de recalage sur le treillis de la Figure 4.11, on trouve un indice de sensibilité pour Q de 7, 1 (resp. 4, 9 et 12, 4) pour la mesure au point c_1 (resp. c_2 et c_3) à l'itération 0 du recalage. On voit ainsi que le capteur placé en c_3 est celui qui donne initialement l'information la plus pertinente pour le recalage dédié à Q ; si seulement deux points de mesure étaient envisagés, il serait préférable de choisir les points c_1 et c_3 .

3.3 Application

On considère un cas-test académique de thermique instationnaire semblable à ceux utilisés en thermique des bâtiments. Dans ce domaine, les modèles traités sont généralement constitués de zones (chambres) et de parois (murs); on considère ici l'échange thermique à travers une paroi de largeur L (Figure 4.16), sur un intervalle d'observation de 24 h. Les températures dans chaque zone sont supposées homogènes et notées $T_{int}(t)$ (température intérieure) et $T_{ext}(t)$ (température extérieure), tandis que la température dans la paroi est notée $\theta(x, t)$.

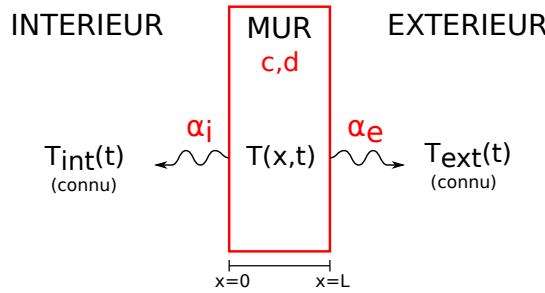


Figure 4.16: Configuration à une paroi étudiée

Les équations du problème sont :

- conditions initiales :

$$T(x, t = 0) = T_0(x) \quad \forall x \in [0, L] \quad (4.54)$$

- équations d'équilibre :

$$\begin{aligned} c \frac{\partial T}{\partial t} + \frac{\partial q}{\partial x} &= 0 \quad \forall x \in]0, L[, \quad \forall t \in [0, T] \\ q + q_i &= 0 \quad \text{en } x = 0, \quad \forall t \in [0, T] \\ -q + q_e &= 0 \quad \text{en } x = L, \quad \forall t \in [0, T] \end{aligned} \quad (4.55)$$

- relations de comportement :

$$\begin{aligned} q &= -d \frac{\partial T}{\partial x} \quad \forall x \in]0, L[, \quad \forall t \in [0, T] \\ q_i &= \alpha_i (T|_{x=0} - T_{int}) \quad \text{en } x = 0, \quad \forall t \in [0, T] \\ q_e &= \alpha_e (T|_{x=L} - T_{ext}) \quad \text{en } x = L, \quad \forall t \in [0, T] \end{aligned} \quad (4.56)$$

avec c (resp. d) la capacité thermique (resp. la conductivité) dans la paroi, et α_i et α_e les coefficients d'échange. Pour les applications numériques, on prend :

$$T_0(x) = 20^\circ C \quad ; \quad T_{int}(t) = 20^\circ C \quad ; \quad T_{ext}(t) = -5 \sin(2\pi t/86400)^\circ C \quad (4.57)$$

Le paramètre $d = 1,0$ est le seul supposé connu. L'ensemble des paramètres à recalculer est donc $\mathbf{p} = [c, \alpha_i, \alpha_e]$; une valeur de référence de ces paramètres, utilisée pour simuler les mesures expérimentales, est $\mathbf{p}_{ref} = [10^6; 4, 3; 17, 7]$.

La quantité d'intérêt étudiée est $T(x = 0, t)$, et le recalage est mené par mesure au cours du temps de $T(x = L/5, t)$. Contrairement au cas statique étudié précédemment, le calcul de la fonction coût implique ici une solution adjointe évoluant en temps inversé et nécessite donc de résoudre un système de la forme :

$$\begin{bmatrix} \mathbb{C} & 0 \\ 0 & -\mathbb{C} \end{bmatrix} \begin{pmatrix} \frac{\partial \mathbf{U}}{\partial t} \\ \frac{\partial \mathbf{\Lambda}}{\partial t} \end{pmatrix} + \begin{bmatrix} \mathbb{K} & -\mathbb{K} \\ \frac{r}{1-r} \mathbb{\Sigma}^T \mathbb{G} \mathbb{\Sigma} & \mathbb{K} \end{bmatrix} \begin{pmatrix} \mathbf{U} \\ \mathbf{\Lambda} \end{pmatrix} = \begin{pmatrix} \mathbf{F} \\ \frac{r}{1-r} \mathbb{\Sigma}^T \mathbb{G} \mathbf{U}_{obs} \end{pmatrix} \quad (4.58)$$

où $\mathbb{\Sigma}$ est un opérateur d'extraction aux points de mesure. La difficulté de la résolution de (4.58) vient du fait qu'il y a des conditions initiales (resp. finales) sur $\mathbf{U}$ (resp. $\mathbf{\Lambda}$) ; des techniques spécifiques peuvent être introduites pour minimiser le coût de calcul associé [Nguyen *et al.* 2008].

A l'itération 0 du processus de recalage, on prend l'ensemble de paramètres $\mathbf{p}_0 = 0,8 \mathbf{p}_{ref}$ (erreur de 20%). Les résultats de recalage obtenus, après 20 itérations, avec l'ERCm classique d'une part et la méthode dédiée à Q (cf. Section 3.1.1) d'autre part sont donnés dans le Tableau 4.1 ; deux valeurs du coefficient de pondération r sont considérées, donnant plus ou moins d'importance aux mesures.

	c	α_i	α_e
$r = 0,5$	0,91	0,95	0,81
$r = 0,9$	0,84	0,97	0,81

	c	α_i	α_e
$r = 0,5$	0,80	1,00	0,81
$r = 0,9$	0,82	1,01	0,80

Tableau 4.1: Valeur adimensionnées (par rapport à $\mathbf{p}_{ref}$) des paramètres recalés avec la méthode de recalage classique (gauche) et celle dédiée à Q (droite)

On observe que le recalage dédié à Q se focalise uniquement sur le paramètre α_1 . On montre sur la Figure 4.17, pour $r = 0,5$ et $r = 0,9$, l'évolution temporelle de la température externe $T_{ext}(t)$, de la valeur de référence de la quantité d'intérêt, et des valeurs $Q(\mathbf{u}_1)$ et $Q(\mathbf{u}_2)$ impliquées dans le processus de recalage dédié à Q ; $Q(\mathbf{u}_1)$ correspond à la valeur de Q prédite par le modèle recalé.

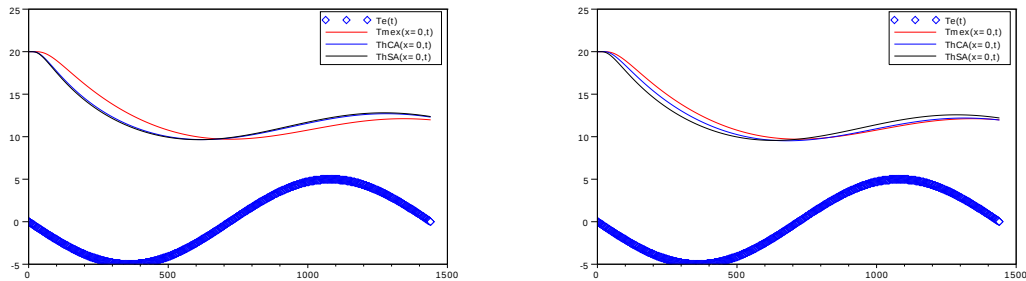


Figure 4.17: Evolution temporelle des différentes températures pour $r = 0,5$ (gauche) et $r = 0,9$ (droite)

On constate donc que le recalage avec la fonction coût dédiée à la quantité d'intérêt Q permet d'obtenir un modèle très précis pour la prédiction de Q .

4 Conclusions partielles et perspectives

On a montré dans ce chapitre des premières études visant à orienter les méthodes de recalage vers des applications cibles, ici le calcul en temps réel et la prédiction de quantités d'intérêt dimensionnantes. Le regroupement des deux sujets traités ici, i.e. l'incorporation d'un recalage dédié à une quantité d'intérêt dans un processus en temps réel semble être une voie pertinente à suivre pour diverses applications pratiques. Beaucoup de travail reste néanmoins à accomplir pour rendre les outils matures et pleinement exploitables en contexte industriel. Notamment, des extensions aux problèmes de dynamique et/ou non-linéaires doivent être faites. De plus, la robustesse des outils vis-à-vis des mesures réelles bruitées doit être vérifiée numériquement.

Une question majeure reste celle des incertitudes de mesure ou de modèle, dont nous avons assez peu parlé ici. Dans le cas de faibles méconnaissances, l'approche proposée par Ladevèze [[Ladevèze et al. 2004](#)] semble représenter un bon compromis entre efficacité et coût de calcul pour les problèmes de Mécanique des Structures. Néanmoins, le cas des fortes méconnaissances est un sujet encore largement ouvert, et une voie peut être l'approche bayésienne qui formule par défaut l'identification dans un cadre stochastique [[Tarantola 2005](#)].

Enfin, le recalage intelligent des modèles de simulation appelle inévitablement à un dialogue accru entre l'expérimental et la simulation. Comme indiqué dans le protocole de DDDAS, une alchimie doit être trouvée pour tirer partie au maximum des possibilités offertes par ce type de dialogue. Comme un essai ne voit pas tout du modèle, une perspective envisagée à court terme est de définir des formes d'éprouvette optimisées, en les considérant comme des structures à part entière, permettant de maximiser la pertinence de l'information expérimentale obtenue afin de prédire efficacement des quantités d'intérêt. Cette vision s'applique également à l'identification de paramètres matériaux, à partir de mesures de champs par exemple, en considérant ces paramètres comme des quantités d'intérêt et en ne gardant que l'information pertinente parmi la grande quantité de données expérimentales recueillies.

Conclusion générale

On a présenté dans ce travail des outils permettant de contrôler toute la chaîne de virtualisation, depuis l'expérience jusqu'à la simulation finale, en vue de la prédiction fiable d'une quantité d'intérêt. Ces outils permettent donc de construire de bons modèles mathématiques et numériques, i.e. des modèles utiles pour l'ingénieur dans une démarche ciblée de conception robuste. Tout au long du manuscrit, on a vu apparaître la notion d'adjoint qui est une notion clé pour l'étude des modèles vis-à-vis de quantités d'intérêt. Un autre concept qui est également apparu souvent est celui d'erreur en relation de comportement, pleinement adapté aux modèles de Mécanique des Structures.

D'un point de vue général, les points-clé importants sur lesquels les futurs travaux de recherche doivent se pencher dans le domaine de la V&V ont été mentionnés dans le manuscrit : non-linéaire, quantification des incertitudes, dialogue essais-calculs, . . . Leur traitement permettra progressivement de passer de la simulation pour comprendre ou prédire à la simulation pour agir.

D'un point de vue plus personnel, je souhaite poursuivre mes travaux sur le contrôle des modèles par rapport à des quantités d'intérêt ; ceci est une activité de recherche importante et porteuse actuellement. La vision *goal-oriented* est très développée à présent dans le cadre de la vérification, même si de nombreux verrous scientifiques persistent ; elle reste à l'être pour le recalage ou la validation des modèles car elle semble offrir des perspectives prometteuses. Ce dernier point constitue un axe de recherche sur lequel je souhaite particulièrement m'investir ces prochaines années.

Bibliographie

[DDD 2006]

January 2006 dddas workshop report. *http://www.dddas.org*, 2006.

[SBE 2006]

Simulation Based Engineering Science - Revolutionizing Engineering Science through Simulation. Report of the National Science Foundation Blue Ribbon Panel, USA, 2006.

[Alfrey 1944]

T. Alfrey. Non-homogeneous stresses in viscoelastic media. *Q. Appl. Math.* II, 2 : 113–119, 1944.

[Allaire 2002]

G. Allaire. Shape optimization by the homogenization method. *Applied Mathematical Sciences*, 146, 2002.

[Allix et al. 2005]

O. Allix, P. Feissel et H.M. Nguyen. Identification strategy in the presence of corrupted measurements. *Engineering Computations*, 22(5-6) : 487–504, 2005.

[Ammar et al. 2010]

A. Ammar, F. Chinesta, P. Diez et A. Huerta. An error estimator for separated representations of highly multidimensional models. *Computer Methods in Applied Mechanics and Engineering*, 199 : 1872–1880, 2010.

[Andrieux et Bui 2006]

S. Andrieux et H.D. Bui. Ecart à la réciprocité et identification de fissures en thermoélasticité isotrope transitoire. *Comptes Rendus Mécanique*, 334(4) : 225–229, 2006.

[Arndt et Luskin 2007]

M. Arndt et M. Luskin. Goal-oriented atomistic-continuum adaptivity for the quasicontinuum approximation. *International Journal for Multiscale Computational Engineering*, 5 : 407–415, 2007.

[Aubin 1967]

J.P. Aubin. Behavior of the error of the approximate solutions of boundary value problems for linear elliptic operators by galerkin and finite difference methods. *Ann Scuola Norm Sup Pisa*, 3 : 599–637, 1967.

[Aubry *et al.* 1999]

D. Aubry, D. Lucas et B. Tie. Adaptive strategy for transient/coupled problems applications to thermoelasticity and elastodynamics. *Computer Methods in Applied Mechanics and Engineering*, 176(1) : 41–50, 1999.

[Babuška 1973]

I. Babuška. The finite element method with lagrangian multipliers. *Numer Math*, 20 : 179–192, 1973.

[Babuška *et al.* 2004]

I. Babuška, R. Tempone et G.E. Zouraris. Galerkin finite element approximations of stochastic elliptic partial differential equations. *SIAM Journal on Numerical Analysis*, 42(2) : 800–825, 2004.

[Babuška et Miller 1984]

I. Babuška et A. Miller. The post-processing approach in the finite element method - part 2 : the calculation of stress intensity factors. *International Journal for Numerical Methods in Engineering*, 20 : 1111–1129, 1984.

[Babuška et Oden 2005]

I. Babuška et J.T. Oden. The reliability of computer predictions : can they be trusted? *International Journal of Numerical Analysis and Modeling*, 1 : 1–18, 2005.

[Babuška et Strouboulis 2001]

I. Babuška et T. Strouboulis. *The finite element method and its reliability*. Oxford university press, 2001.

[Bauman *et al.* 2008]

P.T. Bauman, H. Ben Dhia, N. Elkhodja, J.T. Oden et S. Prudhomme. On the application of the arlequin method to the coupling of particle and continuum models. *Computational Mechanics*, 42 : 511–530, 2008.

[Bauman *et al.* 2009]

P.T. Bauman, J.T. Oden et S. Prudhomme. Adaptive multiscale modeling of polymeric materials : Arlequin coupling and goals algorithms. *Computer Methods in Applied Mechanics and Engineering*, 198 : 799–818, 2009.

[Bazant 1978]

Z.P. Bazant. Spurious reflection of elastic waves in nonuniform finite element grids. *Computer Methods in Applied Mechanics and Engineering*, 16 : 91–100, 1978.

[Becker et Rannacher 2001]

R. Becker et R. Rannacher. An optimal control approach to a posteriori error estimation in finite element methods. *A. Iserles (Ed.), Acta Numerica, Cambridge University Press*, 10 : 1–120, 2001.

[Becker et Vexler 2004]

R. Becker et B. Vexler. A posteriori error estimation for finite element discretization of parameter identification problems. *Numerische Mathematik*, 96 : 435–459, 2004.

[Ben Dhia 1998]

H. Ben Dhia. Multiscale mechanical problems : the arlequin method. *Comptes Rendus de l'Académie des Sciences, Paris, Série II B*, 326(12) : 899–904, 1998.

[Ben Dhia 2008]

H. Ben Dhia. Further insights by theoretical investigations of the multiscale arlequin method. *International Journal for Multiscale Computational Engineering*, 6(3) : 215–232, 2008.

[Ben Dhia et al. 2011]

H. Ben Dhia, L. Chamoin, J.T. Oden et S. Prudhomme. A new adaptive modeling strategy based on optimal control for atomic-to-continuum coupling simulations. *Computer Methods in Applied Mechanics and Engineering*, 200(37-40) : 2675–2696, 2011.

[Ben Dhia et Rateau 2001]

H. Ben Dhia et G. Rateau. Analyse mathématique de la méthode arlequin mixte. *Comptes Rendus de l'Académie des Sciences, Paris, Série I*, 332 : 649–654, 2001.

[Ben Dhia et Rateau 2005]

H. Ben Dhia et G. Rateau. The arlequin method as a flexible engineering design tool. *International Journal for Numerical Methods in Engineering*, 62(11) : 1442–1462, 2005.

[Blanc et al. 2007]

X. Blanc, C. Le Bris et P.L. Lions. Atomistic to continuum limits for computational materials science. *ESAIM : Mathematical Modelling and Numerical Analysis*, 41(2) : 391–426, 2007.

[Bonnet et Abdallah 1994]

M. Bonnet et J. Ben Abdallah. Structural parameter identification using nonlinear gaussian inversion. *Inverse Problems in Engineering Mechanics*, pages 235–242, 1994.

[Bonnet et Constantinescu 2005]

M. Bonnet et A. Constantinescu. Inverse problems in elasticity. *Inverse Problems*, 21 : R1–R50, 2005.

[Bordas et Duflot 2007]

S. Bordas et M. Duflot. Derivative recovery and a posteriori error estimate for extended finite elements. *Computer Methods in Applied Mechanics and Engineering*, 196(35-36) : 3381–3399, 2007.

[Bouclier et al. 2013]

R. Bouclier, F. Louf et L. Chamoin. Real-time validation of mechanical models coupling pgd and constitutive relation error. *Computational Mechanics (online)*, 2013.

[Bourgeat et Piatnitski 2004]

A. Bourgeat et A. Piatnitski. Approximations of effective coefficients in stochastic homogenization. *Annales de l'Institut Henri Poincaré (B) Probability and Statistics*, 40(2) : 153–165, 2004.

[Broughton et al. 1999]

J.Q. Broughton, F.F. Abraham, N. Bernstein et E. Kaxiras. Concurrent coupling of length scales : methodology and application. *Physical Review B*, 60(4) : 2391–2403, 1999.

[Cailletaud et Pilvin 1994]

G. Cailletaud et P. Pilvin. Identification and inverse problems related to material behaviour. In H.D. Bui, M. Tanaka, and al., editors, *Inverse Problems in Engineering Mechanics*, pages 79–86, 1994.

[Cao et Kelly 2003]

T. Cao et D.W. Kelly. Pointwise and local error estimates for the quantities of interest in two-dimensional elasticity. *Computers and Mathematics with Applications*, 46(1) : 69–79, 2003.

[Chamoin et Desvillettes 2013]

L. Chamoin et L. Desvillettes. Control of modeling errors in the coupling of linear transport and diffusion models. *Computer Methods in Applied Mechanics and Engineering*, 261-262 : 83–95, 2013.

[Chamoin et al. 2012]

L. Chamoin, E. Florentin, S. Pavot et V. Visseq. Robust goal-oriented error estimation based on the constitutive relation error for stochastic problems. *Computers and Structures*, 106-107 : 189–195, 2012.

[Chamoin et Ladevèze 2007]

L. Chamoin et P. Ladevèze. Bounds on history-dependent or independent local quantities in viscoelasticity problems solved by approximate methods. *International Journal for Numerical Methods in Engineering*, 71(12) : 1387–1411, 2007.

[Chamoin et Ladevèze 2008]

L. Chamoin et P. Ladevèze. A non-intrusive approach of goal-oriented error estimation for evolution problems solved by the finite element method. *Computer Methods in Applied Mechanics and Engineering*, 197(9-12) : 994–1014, 2008.

[Chamoin et Ladevèze 2009]

L. Chamoin et P. Ladevèze. Strict and practical bounds through a non-intrusive and goal-oriented error estimation method for linear viscoelasticity problems. *Journal of Finite Elements in Analysis and Design*, 45 : 251–262, 2009.

[Chamoin et Ladevèze 2012]

L. Chamoin et P. Ladevèze. Robust control of pgd-based numerical simulations. *European Journal of Computational Mechanics*, 21(3-6) : 195–207, 2012.

[Chamoin et al. 2013]

L. Chamoin, P. Ladevèze et J. Waeytens. Goal-oriented updating of mechanical models using the adjoint framework. (*submitted*), 2013.

[Chamoin et al. 2008]

L. Chamoin, J.T. Oden et S. Prudhomme. A stochastic coupling method for atomic-to-continuum monte-carlo simulations. *Computer Methods in Applied Mechanics and Engineering*, 197(43-44) : 3530–3546, 2008.

[Chamoin et al. 2010]

L. Chamoin, S. Prudhomme, H. Ben Dhia et J.T. Oden. Ghost forces and spurious effects in atomic-to-continuum coupling methods by the arlequin approach. *International Journal for Numerical Methods in Engineering*, 83 : 1081–1113, 2010.

[Chandrasekhar 1950]

S. Chandrasekhar. *Radiative Transfer*. Oxford university press, London, 1950.

[Chinesta et al. 2010]

F. Chinesta, A. Ammar et E. Cueto. Recent advances and new challenges in the use of the proper generalized decomposition for solving multidimensional models. *Archives of Computational Methods in Engineering*, 17(4) : 327–350, 2010.

[Chinesta et al. 2011]

F. Chinesta, P. Ladevèze et E. Cueto. A short review on model reduction based on proper generalized decomposition. *Archives of Computational Methods in Engineering*, 18 : 395–404, 2011.

[Chouaki et al. 1996]

A. Chouaki, P. Ladevèze et L. Proslier. An updating of structural dynamic model with damping. *Inverse Problems in Engineering : Theory and Practice*, pages 335–342, 1996.

[Cirak et Ramm 1998]

F. Cirak et E. Ramm. A posteriori error estimation and adaptivity for linear elasticity using the reciprocal theorem. *Computer Methods in Applied Mechanics and Engineering*, 156 : 351–362, 1998.

[Cochez-Dhondt et Nicaise 2008]

S. Cochez-Dhondt et S. Nicaise. Equilibrated error estimators for discontinuous galerkin methods. *Numerical Methods for Partial Differential Equations*, 24(5) : 1236–1252, 2008.

[Cottureau et al. 2010]

R. Cottureau, L. Chamoin et P. Diez. Strict error bounds for linear and nonlinear solid mechanics problems using a patch-based flux-free method. *Mécanique & Industries*, 11 : 249–254, 2010.

[Cottureau et al. 2011]

R. Cottureau, D. Clouteau, H. Ben Dhia et C. Zaccardi. A stochastic-deterministic coupling method for continuum mechanics. *Computer Methods in Applied Mechanics and Engineering*, 200 : 3280–3288, 2011.

[Curtin et Miller 2003]

W.A. Curtin et R.E. Miller. Atomistic/continuum coupling in computational material science. *Modeling and Simulation in Materials Science and Engineering*, 11 : R33–R68, 2003.

[Dahan et Predeleanu 1980]

M. Dahan et M. Predeleanu. Solutions fondamentales pour un milieu élastique à isotropie transverse. *ZAMP*, 31 : 413–424, 1980.

[Dautray et Lions 1987]

R. Dautray et J.L. Lions. *Analyse Mathématique et Calcul Numérique pour les Sciences et les Techniques*. Masson, 1987.

[Degond et Mas-Gallic 1987]

P. Degond et S. Mas-Gallic. Existence of solutions and diffusion approximation for a model fokker-planck equation. *Transport Theory and Statistical Physics*, 16(4-6) : 589–636, 1987.

[Deraemaeker 2001]

A. Deraemaeker. *Sur la maîtrise des modèles en dynamique des structures à partir de résultats d'essais*. Thèse de doctorat, 2001.

[Duderstadt et Martin 1979]

J.J. Duderstadt et W.R. Martin. *Transport Theory*. Wiley, New York, 1979.

[E et al. 2006]

W. E, J. Lu et J.Z. Yang. Uniform accuracy of the quasicontinuum method. *Physical Review B*, 74(21) : 214115 1–13, 2006.

[Epanechnikov 1969]

V.I. Epanechnikov. Non-parametric estimation of a multivariate probability density. *Theory of Probability and Its Applications*, 14 : 153–158, 1969.

[Eriksson et al. 1995]

K. Eriksson, D. Estep, P. Hansbo et C. Johnson. Introduction to adaptive methods for partial differential equations. *Acta Numerica*, A. Iserles ed, Cambridge University Press, pages 105–159, 1995.

[Estep 2004]

D. Estep. A short course on duality, adjoint operators, green's functions, and a posteriori error analysis. 2004.

[Faverjon et al. 2009]

B. Faverjon, P. Ladevèze et F. Louf. Validation of stochastic linear structural dynamics models. *Computers & Structures*, 87(13-14) : 829–837, 2009.

[Feng et al. 2009]

Y. Feng, D. Fuentes, A. Hawkins, J. Bass et M.N. Rylander. Model-based optimization and real-time control for laser treatment of heterogeneous soft tissues. *Computer Methods in Applied Mechanics and Engineering*, 198(21-26) : 1742–1750, 2009.

[Fish 2006]

J. Fish. Bridging the scales in nano engineering and science. *Journal of Nanoparticle Research*, 8(6) : 577–594, 2006.

[Florentin et al. 2002]

E. Florentin, L. Gallimard et J.P. Pelle. Evaluation of the local quality of stresses in 3d finite element analysis. *Computer Methods in Applied Mechanics and Engineering*, 191 : 4441–4457, 2002.

[Fuentes et al. 2009]

D. Fuentes, J.T. Oden, K.R. Diller, J. Hazle, A. Elliott, A. Shetty et R.J. Stafford. Computational modeling and real-time control of patient-specific laser treatment cancer. *Annals of Biomedical Engineering*, 37(4) : 763–782, 2009.

[Gallimard et Panetier 2006]

L. Gallimard et J. Panetier. Error estimation of stress intensity factors for mixed-mode cracks. *International Journal for Numerical Methods in Engineering*, 68(3) : 299–316, 2006.

[Gartland 1984]

E.C. Gartland. Computable pointwise error bounds and the ritz method in one dimension. *SIAM Journal on Numerical Analysis*, 21(1) : 84–1000, 1984.

[Ghanem et Pelissetti 2002]

R. Ghanem et M. Pelissetti. Error estimation for the validation of model-based predictions. *Proceedings of the 5th World Congress on Computational Mechanics*, 2002.

[Ghanem et Spanos 1991]

R. Ghanem et P. Spanos. *Stochastic Finite Element : a spectral approach*. Springer, 1991.

[Giles et Süli 2002]

M.B. Giles et E. Süli. Adjoint methods for pdes : a posteriori error analysis and postprocessing by duality. *Acta Numerica, Cambridge University Press*, pages 145–236, 2002.

[Graff 1975]

K.F. Graff. *Wave Motion in Elastic Solids*. Dover, New York, 1975.

[Greenberg 1948]

H.J. Greenberg. The determination of upper and lower bounds for the solution of dirichlet problem. *Journal of Mathematical Physics*, 27 : 161–182, 1948.

[Guidault et Belytschko 2007]

P.A. Guidault et T. Belytschko. On the l^2 and the h^1 couplings for an overlapping domain decomposition method using lagrange multipliers. *International Journal for Numerical Methods in Engineering*, 70(3) : 322–350, 2007.

[Johansson et al. 2011]

H. Johansson, F. Larsson et K. Runesson. Application-specific error control for parameter identification problems. *International Journal for Numerical Methods in Biomedical Engineering*, 27 : 608–618, 2011.

[Johansson et al. 2007]

H. Johansson, K. Runesson et F. Larsson. Parameter identification with sensitivity assessment and error computation. *GAMM-Mitt*, 30(2) : 430–457, 2007.

[Knap et Ortiz 2001]

J. Knap et M. Ortiz. An analysis of the quasicontinuum method. *Journal of the Mechanics and Physics of Solids*, 49 : 1899–1923, 2001.

[Ladevèze 1999]

P. Ladevèze. *Nonlinear Computational Structural Mechanics - New Approaches and Non-Incremental Methods of Calculation*. Springer NY, 1999.

[Ladevèze 2006]

P. Ladevèze. Upper error bounds on calculated outputs of interest for linear and nonlinear structural problems. *Comptes Rendus Académie des Sciences - Mécanique, Paris*, 334 : 399–407, 2006.

[Ladevèze 2008]

P. Ladevèze. Strict upper error bounds for computed outputs of interest in computational structural mechanics. *Computational Mechanics*, 42(2) : 271–286, 2008.

[Ladevèze 2003]

P. Ladevèze. Validation et vérification de modèles stochastiques en environnement incertain par la méthode de l’erreur en relation de comportement. *Rapport interne n° 258, LMT-Cachan*, juin 2003.

[Ladevèze et al. 2012]

P. Ladevèze, B. Blaysat et E. Florentin. Strict upper bounds of the error in calculated outputs of interest for plasticity problems. *Computer Methods in Applied Mechanics and Engineering*, 245-246 : 194–205, 2012.

[Ladevèze et Chamoin 2010]

P. Ladevèze et L. Chamoin. Calculation of strict error bounds for finite element approximations of nonlinear pointwise quantities of interest. *International Journal for Numerical Methods in Engineering*, 84 : 1638–1664, 2010.

[Ladevèze et Chamoin 2011]

P. Ladevèze et L. Chamoin. On the verification of model reduction methods based on the proper generalized decomposition. *Computer Methods in Applied Mechanics and Engineering*, 200 : 2032–2047, 2011.

[Ladevèze et Chamoin 2012]

P. Ladevèze et L. Chamoin. *Toward guaranteed PGD-reduced models*. Bytes and Science, G. Zavarise & D.P. Boso (Eds.), CIMNE, 2012.

[Ladevèze et al. 2010]

P. Ladevèze, L. Chamoin et E. Florentin. A new non-intrusive technique for the construction of admissible stress fields in model verification. *Computer Methods in Applied Mechanics and Engineering*, 199(9-12) : 766–777, 2010.

[Ladevèze et Chouaki 1999]

P. Ladevèze et A. Chouaki. Application of a posteriori error estimation for structural model updating. *Inverse Problems*, 15(1) : 49–58, 1999.

[Ladevèze et Florentin 2006]

P. Ladevèze et E. Florentin. Verification of stochastic models in uncertain environments using the constitutive relation error method. *Computer Methods in Applied Mechanics and Engineering*, 196 : 225–234, 2006.

[Ladevèze et Leguillon 1983]

P. Ladevèze et D. Leguillon. Error estimate procedure in the finite element method and application. *SIAM Journal of Numerical Analysis*, 20(3) : 485–509, 1983.

[Ladevèze et Moës 1998]

P. Ladevèze et N. Moës. A new a posteriori error estimation for nonlinear time-dependent finite element analysis. *Computer Methods in Applied Mechanics and Engineering*, 157 : 45–68, 1998.

[Ladevèze et al. 1994]

P. Ladevèze, D. Nedjar et M. Reynier. Updating of finite element models using vibration tests. *AIAA Journal*, 32(7) : 1485–1491, 1994.

[Ladevèze et Pelle 2004]

P. Ladevèze et J.P. Pelle. *Mastering Calculations in Linear and Nonlinear Mechanics*. Springer NY, 2004.

[Ladevèze et al. 2013]

P. Ladevèze, F. Pled et L. Chamoin. New bounding techniques for goal-oriented error estimation applied to linear problems. *International Journal for Numerical Methods in Engineering*, 93(13) : 1345–1380, 2013.

[Ladevèze et al. 2004]

P. Ladevèze, G. Puel, A. Deraemaeker et T. Romeuf. Sur une théorie des méconnaissances en calcul des structures. *Revue Européenne des Eléments Finis*, 13(5-7) : 571–582, 2004.

[Ladevèze et al. 1999]

P. Ladevèze, P. Rougeot, P. Blanchard et J.P. Moreau. Local error estimators for finite element linear analysis. *Computer Methods in Applied Mechanics and Engineering*, 176 : 231–246, 1999.

[Larsen et Keller 1974]

E. Larsen et J.B. Keller. Asymptotic solutions of neutron transport problems. *J. Math. Phys.*, 15 : 75–81, 1974.

[Lions 1971]

J.L. Lions. *Optimal control of systems governed by partial differential equations*. Springer, 1971.

[Liu et al. 2004]

W.K. Liu, E.G. Karpov et H.S. Park. An introduction to computational nanomechanics and materials. *Computer Methods in Applied Mechanics and Engineering*, 193 : 1529–1578, 2004.

[Love 1944]

A.E.H. Love. A treatise on the mathematical theory of elasticity. *Dover, New-York*, pages 186–213, 1944.

[Marchais et al. 2013]

J. Marchais, C. Rey et L. Chamoin. Geometrically consistent approximations of the energy in the quasi-continuum framework. *Journal of the Mechanics and Physics of Solids (submitted)*, 2013.

[Matthies et Bucher 1999]

H.G. Matthies et C.G. Bucher. Finite element for stochastic media problems. *Computer Methods in Applied Mechanics and Engineering*, 168 : 3–17, 1999.

[Miller et Tadmor 2002]

R.E. Miller et E.B. Tadmor. The quasicontinuum method : overview, applications and current directions. *Journal of Computer-Aided Materials Design*, 9 : 203–239, 2002.

[Mindlin 1936]

R.D. Mindlin. Force at a point in the interior of a semi-infinite solid. *Journal of Physics*, 7 : 195–202, 1936.

[Moës et al. 2006]

N. Moës, E. Béchet et E. Tourbier. Imposing dirichlet boundary conditions in the extended finite element method. *International Journal for Numerical Methods in Engineering*, 67 : 1641–1669, 2006.

[Moës et al. 1999]

N. Moës, J. Dolbow et T. Belytschko. A finite element method for crack growth without remeshing. *International Journal for Numerical Methods in Engineering*, 46 : 131–150, 1999.

[Morin et al. 2001]

P. Morin, R.H. Nochetto et K.G. Siebert. Local problems on stars : A posteriori error estimators, convergence, and performance. *Mathematics of Computation*, 72(243) : 1067–1097, 2001.

[Nassiopoulos et Bourquin 2010]

A. Nassiopoulos et F. Bourquin. Fast three-dimensional temperature reconstruction. *Computer Methods in Applied Mechanics and Engineering*, 199 : 3169–3178, 2010.

[Nguyen et al. 2008]

H.M. Nguyen, O. Allix et P. Feissel. A robust identification strategy for rate-dependent models in dynamics. *Inverse Problems*, 24 : 065006, 2008.

[Nitsche 1968]

J. Nitsche. Ein kriterium für die quasi-optimalität des ritzschen verfahrens. *Numer Math*, 11 : 346–348, 1968.

[Nouy 2007]

A. Nouy. A generalized spectral decomposition technique to solve a class of linear stochastic partial differential equations. *Computer Methods in Applied Mechanics and Engineering*, 196(45-48) : 4521–4537, 2007.

[Nouy 2010]

A. Nouy. A priori model reduction through proper generalized decomposition for solving time dependent partial differential equations. *Computer Methods in Applied Mechanics and Engineering*, 199 : 1603–1626, 2010.

[Oberkampf et al. 2003]

W. Oberkampf, T. Trucano et C. Hirsh. *Verification, Validation and Predictive Capability in Computational Engineering and Physics*. Technical report, Sandia 2003-3769, 2003.

[Oden et al. 2005]

J.T. Oden, I. Babuška, F. Nobile, Y. Feng et R. Tempone. Theory and methodology for estimation and control of error due to modeling, approximation, and uncertainty. *Computer Methods in Applied Mechanics and Engineering*, 194(2-5) : 195–204, 2005.

[Oden et Prudhomme 1999]

J.T. Oden et S. Prudhomme. Goal-oriented error estimation and adaptativity for the finite element method. *TICAM repport, Austin (USA)*, 99, 1999.

[Oden et Prudhomme 2002]

J.T. Oden et S. Prudhomme. Estimation of modeling error in computational mechanics. *Journal of Computational Physics*, 182 : 496–515, 2002.

[Oden et al. 2010]

J.T. Oden, S. Prudhomme, P.T. Bauman et L. Chamoin. *Estimation and control of modeling error : a general approach to multiscale modeling*. Bridging the Scales in Science and Engineering, J. Fish (editor), Oxford University Press, 2010.

[Oden et al. 2006]

J.T. Oden, S. Prudhomme, A. Romkes et P.T. Bauman. Multi-scale modeling of

physical phenomena : Adaptive control of models. *SIAM Journal on Scientific Computing*, 28 : 2359–2389, 2006.

[Oden et Vemaganti 2000]

J.T. Oden et K. Vemaganti. Estimation of local modeling error and goal-oriented modeling of heterogeneous materials. part i : Error estimates and adaptive algorithms. *Journal of Computational Physics*, 164 : 22–47, 2000.

[Oden et al. 1999]

J.T. Oden, K. Vemaganti et N. Moës. Hierarchical modeling of heterogeneous solids. *Computer Methods in Applied Mechanics and Engineering*, 172(1-4) : 3–25, 1999.

[Oden et Zohdi 1997]

J.T. Oden et T.I. Zohdi. Analysis and adaptive modeling of highly heterogeneous elastic structures. *Computer Methods in Applied Mechanics and Engineering*, 148 : 367–391, 1997.

[Ortner et Zhang 2011]

C. Ortner et L. Zhang. Construction and sharp consistency estimates for atomistic-continuum coupling methods with general interfaces : a 2d model problem. *Arxiv preprint arXiv :1110.0168*, 2011.

[Panetier et al. 2010]

J. Panetier, P. Ladevèze et L. Chamoin. Strict and effective bounds in goal-oriented error estimation applied to fracture mechanics problems solved with the xfem. *International Journal for Numerical Methods in Engineering*, 81(6) : 671–700, 2010.

[Panetier et al. 2009]

J. Panetier, P. Ladevèze et F. Louf. Strict bounds for computed stress intensity factors. *Computers and Structures*, 871(15-16) : 1015–1021, 2009.

[Paraschivoiu et al. 1997]

M. Paraschivoiu, J. Peraire et A.T. Patera. A posteriori finite element bounds for linear functional outputs of elliptic partial differential equations. *Computer Methods in Applied Mechanics and Engineering*, 150 : 289–312, 1997.

[Pares et al. 2006]

N. Pares, P. Diez et A. Huerta. Subdomain-based flux-free a posteriori error estimators. *Computer Methods in Applied Mechanics and Engineering*, 195(4-6) : 297–323, 2006.

[Pitkaranta et al. 2009]

J. Pitkaranta, I. Babuska et B. Szabo. The problem of verification with reference to the girkmann problem. *IACM Expressions*, 24 : 14–15, 2009.

[Pled et al. 2011]

F. Pled, L. Chamoin et P. Ladevèze. On the techniques for constructing admissible stress fields in model verification : performances on engineering examples.

International Journal for Numerical Methods in Engineering, 88(5) : 409–441, 2011.

[Pled *et al.* 2012]

F. Pled, L. Chamoin et P. Ladevèze. An enhanced method with local energy minimization for the robust a posteriori construction of equilibrated stress fields in finite element analyses. *Computational Mechanics*, 49(3) : 357–378, 2012.

[Pomraning et Foglesong 1979]

G.C. Pomraning et C.M. Foglesong. Transport-diffusion interfaces in radiative transfer. *J. Computational Phys.*, 32 : 420–436, 1979.

[Prager et Synge 1947]

W. Prager et J.L. Synge. Approximation in elasticity based on the concept of functions spaces. *Quarterly of Applied Mathematics*, 5 : 261–269, 1947.

[Prudhomme *et al.* 2006]

S. Prudhomme, P.T. Bauman et J.T. Oden. Error control for molecular statics problems. *International Journal for Multiscale Computational Engineering*, 4 : 647–662, 2006.

[Prudhomme *et al.* 2012]

S. Prudhomme, R. Bouclier, L. Chamoin et J.T. Oden. Analysis of an averaging operator for atomic-to-continuum coupling methods by the arlequin approach. *Lecture Notes in Computational Sciences and Engineering*, 82 : 365–396, 2012.

[Prudhomme *et al.* 2009]

S. Prudhomme, L. Chamoin, H. Ben Dhia et P.T. Bauman. An adaptive strategy for the control of modeling error in two-dimensional atomic-to-continuum coupling simulations. *Computer Methods in Applied Mechanics and Engineering*, 198(21–26) : 1887–1901, 2009.

[Prudhomme *et al.* 2008]

S. Prudhomme, H. Ben Dhia, P.T. Bauman, N. Elkhodja et J.T. Oden. Computational analysis of modeling error for the coupling of particle and continuum models by the arlequin method. *Computer Methods in Applied Mechanics and Engineering*, 197 : 3399–3409, 2008.

[Prudhomme *et al.* 2004]

S. Prudhomme, F. Nobile, L. Chamoin et J.T. Oden. Analysis of a subdomain-based error estimator for finite element approximations of elliptic problems. *Numerical Methods for Partial Differential Equations*, 20(2) : 165–192, 2004.

[Prudhomme et Oden 1999]

S. Prudhomme et J.T. Oden. On goal-oriented error estimation for elliptic problems : application to the control of pointwise errors. *Computer Methods in Applied Mechanics and Engineering*, 176 : 313–331, 1999.

[Rannacher et Suttmeier 1997]

R. Rannacher et F.T. Suttmeier. A feedback approach to error control in finite element methods : application to linear elasticity. *Computational Mechanics*, 19 : 434–446, 1997.

[Roache 1998]

P.J. Roache. *Verification and Validation in Computational Science and Engineering*. Hermosa publishers, 1998.

[Romkes et Oden 2004]

A. Romkes et J.T. Oden. Adaptive modeling of wave propagation in heterogeneous elastic solids. *Computer Methods in Applied Mechanics and Engineering*, 193 : 539–559, 2004.

[Salvarini et Vasquez 2005]

F. Salvarini et J.L. Vasquez. The diffusive limit of carleman-type equations. *Non-linearity*, 9 : 67–80, 2005.

[Schueller 1997]

G. Schueller. A state-of-the-art report on computational stochastic mechanics. *Probabilistic Engineering Mechanics*, 12 : 197–321, 1997.

[Shapeev 2011]

A.V. Shapeev. Consistent energy-based atomistic/continuum coupling for two-body potentials in 1d and 2d. *Multiscale Modeling and Simulation*, 9 : 905–932, 2011.

[Shenoy *et al.* 1999]

V.B. Shenoy, R.E. Miller, E.B. Tadmor, D. Rodney, R. Phillips et M. Ortiz. An adaptive finite element approach to atomic-scale mechanics : the quasicontinuum method. *Journal of the Mechanics and Physics of Solids*, 47 : 611–642, 1999.

[Shimokawa *et al.* 2004]

T. Shimokawa, J.J. Mortensen, J. Schiotz et K.W. Jacobsen. Matching conditions in the quasicontinuum method : Removal of the error introduced at the interface between the coarse-grained and fully atomistic region. *Physical Review B*, 69(21) : 214104 1–10, 2004.

[Sneddon et Hill 1964]

I.N. Sneddon et R. Hill. *Progress in solid mechanics*. North Holland Publishing Company, Amsterdam, 1964.

[Stein 2003]

E. Stein. *Error controlled Adaptive Finite Elements in Solid Mechanics*. J. Wiley, 2003.

[Stengel 1994]

R.F. Stengel. *Optimal Control and Estimation*. Dover Books on Advanced Mathematics, 1994.

[Stone et Babuska 1998]

T.J. Stone et I. Babuska. A numerical method with a posteriori error estimation for determining the path taken by a propagating crack. *Computer Methods in Applied Mechanics and Engineering*, 160 : 245–271, 1998.

[Strouboulis et al. 2000]

T. Strouboulis, I. Babuška et K. Copps. The design and analysis of the generalized finite element method. *Computer Methods in Applied Mechanics and Engineering*, 181(1-3) : 43–69, 2000.

[Sudret et al. 2004]

B. Sudret, M. Berveiller et M. Lemaire. Eléments finis stochastiques en élasticité linéaire. *Comptes Rendus de l'Académie des Sciences, Mécanique*, 332 : 531–537, 2004.

[Tadmor et al. 1996]

E.B. Tadmor, M. Ortiz et R. Philips. Quasicontinuum analysis of defects in solids. *Philosophical Magazine*, A73 : 1529–1563, 1996.

[Tarantola 2005]

A. Tarantola. *Inverse problem theory and model parameter estimation*. SIAM, 2005.

[Tikonov et Arsenin 1977]

A.N. Tikonov et Y. Arsenin. *Solutions to ill-posed problems*. Wintson-Widley, New York, 1977.

[Tidriri 2001]

M.D. Tidriri. New models for the solution of intermediate regimes in transport theory and radiative transfer : existence theory, positivity, asymptotic analysis, and approximations. *Journal of Statistical Physics*, 104(1-2) : 291–325, 2001.

[Vemaganti et Oden 2001]

K. Vemaganti et J.T. Oden. Estimation of local modeling error and goal-oriented modeling of heterogeneous materials ; part 2 : A computational environment for adaptive modeling of heterogeneous elastic solids. *Computer Methods in Applied Mechanics and Engineering*, 190 : 6089–6124, 2001.

[Vijayakumar et Cormack 1987]

S. Vijayakumar et E.C. Cormack. Green's functions for the biharmonic equation : bonded elastic media. *SIAM Journal of Applied Mathematics*, 47(5), 1987.

[Vohralik 2011]

M. Vohralik. Guaranteed and fully robust a posteriori error estimates for conforming discretizations of diffusion problems with discontinuous coefficients. *Journal of Scientific Computing*, 46 : 397–438, 2011.

[Waeytens *et al.* 2012]

J. Waeytens, L. Chamoin et P. Ladevèze. Guaranteed error bounds on point-wise quantities of interest for transient viscodynamics problems. *Computational Mechanics*, 49(3) : 291–307, 2012.

[Wagner et Liu 2003]

G.J. Wagner et W.K. Liu. Coupling of atomistic and continuum simulations using a bridging scale decomposition. *Journal of Computational Physics*, 190(1) : 249–274, 2003.

[Washizu 1953]

K. Washizu. Bounds for solutions of boundary value problems in elasticity. *Journal of Mathematical Physics*, 32 : 117–128, 1953.

[Xiao et Belytschko 2004]

S.P. Xiao et T. Belytschko. A bridging domain method for coupling continua with molecular dynamics. *Computer Methods in Applied Mechanics and Engineering*, 193 : 1645–1669, 2004.

[Zaccardi *et al.* 2013]

C. Zaccardi, L. Chamoin, R. Cottureau et H. Ben Dhia. Error estimation and model adaptation for a stochastic-deterministic coupling method based on the arlequin framework. *International Journal for Numerical Methods in Engineering (online)*, 2013.