



Sur une théorie des méconnaissances en dynamique des structures

Guillaume Puel

► To cite this version:

Guillaume Puel. Sur une théorie des méconnaissances en dynamique des structures. Mécanique [physics.med-ph]. École normale supérieure de Cachan - ENS Cachan, 2004. Français. NNT: . tel-00008493

HAL Id: tel-00008493

<https://theses.hal.science/tel-00008493>

Submitted on 14 Feb 2005

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

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

**THÈSE DE DOCTORAT
DE L'ÉCOLE NORMALE SUPÉRIEURE DE CACHAN**

Spécialité :
MÉCANIQUE - GÉNIE MÉCANIQUE - GÉNIE CIVIL

Présentée à l'École Normale Supérieure de Cachan par

Guillaume PUEL

pour obtenir le grade de

**DOCTEUR
DE L'ÉCOLE NORMALE SUPÉRIEURE DE CACHAN**

Sujet de la thèse :

Sur une théorie des méconnaissances en dynamique des structures

Thèse soutenue le 9 décembre 2004 devant le jury composé de :

M. REYNIER	Directrice de l'ENSAM	Présidente
F. HEMEZ	Los Alamos National Laboratory (USA)	Rapporteur
L. JÉZÉQUEL	Professeur, École Centrale de Lyon	Rapporteur
M. BONNET	Directeur de recherche, École Polytechnique	Examineur
S. COGAN	Chargé de recherche, U. de Franche-Comté	Examineur
E. LOUAAS	Centre National d'Études Spatiales	Examineur
T. ROMEUF	EADS Space Transportation	Examineur
P. LADEVÈZE	Directeur du LMT-Cachan, ENS Cachan	Directeur de thèse

Laboratoire de Mécanique et Technologie
(ENS Cachan/CNRS UMR8535/Université Paris 6)
61 avenue du Président Wilson, 94235 CACHAN CEDEX (France)

*J'ai des tas d'idées brillantes et nouvelles,
mais les brillantes ne sont pas nouvelles,
et les nouvelles ne sont pas brillantes. . .*

Marcel Achard

*Citer les pensées des autres, c'est regretter
de ne pas les avoir trouvées soi-même.*

Sacha Guitry

À ma famille
À Hélène

C'est de nouveau à l'aide de mes dix doigts agiles que je peux consigner ici quelques remerciements adressés à tous ceux qui m'ont supporté (dans tous les sens du terme !) au cours de ces trois dernières années.

Cette thèse a été réalisée au Laboratoire de Mécanique et Technologie de Cachan, sous la baguette de Pierre Ladevèze, que je tiens à remercier pour ne pas m'avoir laissé trop souvent patauger dans la méconnaissance et avoir su constituer le bon dosage entre l'indépendance et la direction dont j'avais besoin dans ce travail de longue haleine.

Je tiens également à exprimer toute ma gratitude à M^{me} Marie Reynier pour l'honneur qu'elle m'a fait en assurant la présidence de mon jury, ainsi qu'à MM. François Hemez et Louis Jézéquel pour avoir accepté la tâche souvent longue et difficile de rapporteur de ce mémoire de thèse. Je remercie enfin les autres participants de ce jury, MM. Marc Bonnet, Scott Cogan, Éric Louaas et Thierry Romeuf, de s'être plongés eux aussi dans le monde des méconnaissances.

Mes remerciements s'adressent maintenant tout particulièrement aux joyeux drilles du bien nommé Bureau Sympa : Mathilde et son sens aigu de la décoration et de la bonne humeur, Jay et son esprit fort, toujours à proposer divers défis intellectuello-sportifs, David et son éternel bon sens (en particulier $\LaTeX$ ien), ainsi que sa planification rigoureuse de l'espace ; merci d'avoir créé un environnement convivial de travail, mais aussi de détente (souvent musicale), indispensable à la sauvegarde de ma santé mentale au cours de certaines périodes plus difficiles que d'autres...

J'en profite pour jeter un clin d'œil amusé aux habitants d'une certaine mezzanine qui, s'ils ont bien essayé, n'ont pas réussi à obtenir la popularité et le rayonnement culturel du Bureau Sympa, et ce, malgré l'instauration récente de la parité ; j'ai nommé, des plus anciens aux plus jeunes : Doud, Cloups, Sandra et Karine, chez lesquels je n'ai pas été dépaycé par l'ambiance.

Je garderai aussi de très bons souvenirs des membres et ex-membres des quelques équipes de travail non officielles auxquelles j'ai pris part : le groupe d'étude Comportement Dynamique d'un Corps Sphérique sous Interactions Multiples, initié par Boubou et Francis, le projet LMTV-Achat, dont les auteurs tiennent encore à garder l'anonymat, ainsi que la toute récente UTR (V)VARUM NICHT, qui fort de la multidisciplinarité de ceux qui la constituent (Mathilde encore, Béa, Olive, Pierrot, Guigui, et moi-même donc), ne demande qu'à prendre son envol...

Je salue aussi les divers camarades de la promo 97 qui ont accompagné toutes ces années mécaniciennes : tout d'abord Alex mon éternel compagnon de vadrouille et RV mon binôme de TP préféré ;-), mais aussi Willy, Braj, Manu, Toto, Tophe, Sami, Patto, Benoît, Bébert, pour ne citer que ceux que je n'ai pas encore eu l'occasion d'évoquer et qui ont plutôt gravité autour de Cachan...

Après ceux de ma génération, je ne peux m'empêcher de faire référence à tous les anciens, encore présents ou non, qui ont fait du LMT ce qu'il est aujourd'hui : Arno, Floflo, Rouchy, Laurent, Fabien, Juju, Pierre, Delphine, François, Séb, Gilles, Marc, Claude, PAB, sans oublier notre modèle à tous (quoique) Frisou... et je souhaite bon courage à tous les petits jeunes qui n'ont pas encore atteint le stade libératoire de la soutenance de thèse, et dont la liste serait trop longue à dresser ici.

Quelques derniers remerciements en vrac : à Guy, Maïté et Marie-T pour m'avoir fait découvrir le monde merveilleux des vibrations, à Yves pour ses conseils avisés, aux potes de Taverny pour leur amitié constante depuis tant d'années, aux membres de l'Opéra Chœur Ouvert et de La Croche Chœur pour la note musicale qu'ils m'ont apportée, à Georges Pérec pour m'avoir appris le sens secret de la recherche, et à Peter Jackson pour avoir orné chacune de mes trois années de thèse d'une perle cinématographique...

Enfin, je lance, en particulier à tous ceux que j'ai oubliés sur cette page, et qui ont contribué par leur présence ou leur intérêt à ce travail, un grand :

MERCI !

TABLE DES MATIÈRES

Table des matières	i
Table des figures	v
Liste des tableaux	vii
Introduction	1
1 État de l'art des méthodes stochastiques en calcul des structures	9
1.1 Petit historique succinct des probabilités	10
1.2 Choix des paramètres stochastiques	11
1.2.1 Types de paramètres stochastiques	11
1.2.2 Choix des lois de probabilité	13
1.2.3 Sélection des paramètres les plus pertinents	15
1.3 Méthode de Monte Carlo	16
1.3.1 Avantages et inconvénients	16
1.3.2 Améliorations de la méthode	17
1.4 Méthode des Éléments Finis Stochastiques	18
1.4.1 Méthode des Éléments Finis Stochastiques par Perturbation	18
1.4.2 Méthode Spectrale des Éléments Finis Stochastiques	20
1.5 Autres méthodes stochastiques	23
1.5.1 Méthodes fiabilistes	23
1.5.2 Méthodes non-paramétriques	25
1.6 Méthodes stochastiques et validation de modèles	26
2 État de l'art des méthodes non stochastiques	27
2.1 Motivation des méthodes non stochastiques	28
2.2 Théorie généralisée de l'information	29
2.2.1 Description	29
2.2.2 Exemple de la modélisation probabiliste	30
2.3 La théorie des intervalles	32
2.3.1 Principes	32
2.3.2 Résolution de systèmes matriciels	33
2.3.3 Applications en calcul des structures	34
2.4 Les ensembles flous	35
2.4.1 Principes et définitions	35
2.4.2 Applications	36
2.5 Ensembles aléatoires selon Dempster-Shafer	38
2.5.1 Définitions et propriétés	38
2.5.2 Applications	42

2.5.3	Conclusion	45
2.6	Les modèles convexes d'incertitude	46
2.6.1	Cadre mathématique associé	46
2.6.2	Applications dans le domaine de la fiabilité	49
2.6.3	Applications dans le domaine de la validation de modèles	49
3	Description envisagée de la réalité	53
3.1	Réalité considérée	54
3.1.1	Famille de structures réelles	54
3.1.2	Données expérimentales retenues	54
3.2	Simulation stochastique d'une famille de structures	56
3.2.1	Principe	56
3.2.2	Structure étudiée	56
3.2.3	Obtention des valeurs caractéristiques	58
4	Définition du concept de méconnaissance	61
4.1	Principes de la modélisation des méconnaissances	62
4.1.1	Définition des méconnaissances de base	62
4.1.2	Choix des lois de probabilité	63
4.2	Exemple : modélisation des dispersions matériaux	64
4.2.1	Justification du choix d'une loi normale	65
4.2.2	Conversion de la loi normale en termes de méconnaissances	66
4.3	Concept de probabilité d'intervalle	66
4.3.1	Intervalle standards et probabilité d'intervalle	66
4.3.2	Interprétation	69
4.3.3	Comparaison avec les méthodes non stochastiques	71
4.4	Obtention pratique et manipulation des méconnaissances de base	72
5	Propagation des méconnaissances de base	77
5.1	Principe	78
5.2	Calcul théorique des méconnaissances effectives	80
5.2.1	Méconnaissance effective sur une pulsation propre	80
5.2.2	Méconnaissance effective sur la valeur en un DDL d'un mode	81
5.2.3	Méconnaissance effective sur la projection d'un mode	82
5.3	Calcul de méconnaissances effectives sur un exemple simple	83
5.3.1	Premier choix de modèle de méconnaissances de base	83
5.3.2	Deuxième choix de modèle de méconnaissances de base	88
5.3.3	Troisième choix de modèle de méconnaissances de base	89
5.4	Prise en compte d'une forte méconnaissance	91
5.4.1	Principe et mise en œuvre	91
5.4.2	Application sur un exemple simple	93
6	Détermination des méconnaissances de base à partir de données expérimentales	95
6.1	Point de vue adopté	96

6.1.1	Données expérimentales	96
6.1.2	Confrontation entre le modèle et les données expérimentales . . .	97
6.1.3	Détermination des méconnaissances de base représentatives . . .	97
6.2	Réduction des méconnaissances de base : première formulation	99
6.2.1	Description du principe associé au procédé de réduction	99
6.2.2	Prise en compte des cas les plus défavorables	100
6.3	Réduction des méconnaissances de base sur un exemple académique . . .	103
6.3.1	Structure étudiée	103
6.3.2	Mise en œuvre du procédé de réduction	105
6.3.3	Cas d'une forte méconnaissance	111
6.4	Réduction des méconnaissances de base : seconde formulation	113
6.4.1	Notion de représentativité d'une mesure expérimentale	113
6.4.2	Illustration	114
6.5	Vers une stratégie générale de réduction des méconnaissances	117
6.5.1	Définition du modèle initial de méconnaissances de base	118
6.5.2	Détermination de l'ordre des réductions successives	118
6.5.3	Processus de réduction des méconnaissances de base	119
6.5.4	Interprétation des résultats	120
7	La théorie des méconnaissances au sein de la démarche de validation de mo-	
	dèles	121
7.1	Recalage de modèles dynamiques	122
7.1.1	Concept d'erreur en relation de comportement	122
7.1.2	Prise en compte des données expérimentales	123
7.1.3	Implémentation de la méthode	125
7.2	Validation d'un modèle simplifié de liaison	127
7.2.1	Présentation de la structure d'étude	127
7.2.2	Recalage du modèle théorique à l'aide des données expérimentales	128
7.2.3	Réduction des méconnaissances de base	129
8	Application de la théorie des méconnaissances à un cas industriel	133
8.1	Structure industrielle étudiée	134
8.1.1	Présentation du SYLDA5	134
8.1.2	Description des mesures expérimentales effectuées	136
8.1.3	Modèle associé	136
8.2	Recalage du SYLDA5	137
8.2.1	Itérations du recalage	137
8.2.2	Bilan du recalage	138
8.3	Réduction des méconnaissances de base du SYLDA5	140
8.3.1	Données du problème	140
8.3.2	Processus de réduction des méconnaissances de base	141
	Conclusion	145

A	Calcul de la méconnaissance effective sur un mode propre	147
A.1	Expression de la sensibilité au premier ordre d'un mode propre	148
A.2	Calcul pratique dans le cas de structures à grand nombre de DDL	150
A.3	Méconnaissance effective sur la projection d'un mode propre	151
	Bibliographie	153

TABLE DES FIGURES

1.1	Densités de probabilité pour une loi normale et une loi log-normale	15
1.2	Principe des méthodes fiabilistes	24
2.1	Comparaison des résultats obtenus sur un exemple simple par une représentation par intervalles aléatoires et une modélisation probabiliste, d'après [Oberkampf <i>et al.</i> 2001]	45
2.2	Panorama des principales méthodes non stochastiques d'après [Joslyn et Booker 2004]	47
3.1	Valeurs expérimentales à P % de la quantité d'intérêt α_{exp}	55
3.2	Treillis plan à quatre barres	57
3.3	Modes propres associés au modèle déterministe du treillis à quatre barres	57
3.4	Répartitions des valeurs de $\omega_{2\text{exp}}^2$ obtenues pour le treillis à quatre barres à l'aide de la méthode de Monte Carlo (50000 tirages)	59
4.1	Méconnaissances de base associées à une loi normale centrée	67
4.2	Intervalle standard et probabilité d'intervalle	69
4.3	Cas général de modélisation des méconnaissances de base	70
4.4	Méconnaissances de base modélisées par une loi uniforme	71
4.5	Densités de probabilité associées à une loi normale non centrée	73
4.6	Interprétation de la somme de deux variables aléatoires	75
5.1	Calcul des méconnaissances effectives sur la quantité d'intérêt $\Delta\alpha_{\text{mod}}$	79
5.2	Répartitions des bornes de $\omega_{2\text{mod}}^2$ obtenues par la méthode de Monte Carlo (50000 tirages)	84
5.3	Répartitions des bornes de $\phi_{24\text{mod}}$ obtenues par la méthode de Monte Carlo (50000 tirages)	85
5.4	Densités de probabilité des bornes de $\phi_{24\text{mod}}$ calculées à l'aide des fonctions caractéristiques	85
5.5	Répartitions des bornes de $\omega_{3\text{mod}}^2$ obtenues par la méthode de Monte Carlo (50000 tirages) pour le treillis à quatre barres dans le cas d'un modèle de sous-structure très simplifié par rapport à la réalité	90
6.1	Valeurs expérimentales à P % de la quantité d'intérêt $\Delta\alpha_{\text{exp}}$	96
6.2	Comparaison entre les méconnaissances effectives et les valeurs expérimentales à P %	97
6.3	Principe de la réduction des méconnaissances de base à l'aide d'informations expérimentales pertinentes	101
6.4	Prise en compte des cas les plus défavorables dans le processus de réduction des méconnaissances de base	102
6.5	Treillis plan à six barres	103

6.6	Modes propres du treillis à six barres	104
6.7	Énergies de déformation modales par groupe de barres pour les modes du treillis à six barres	107
6.8	Modes propres de flexion de la plaque orthotrope	116
7.1	Procédure de recalage	126
7.2	Problème de poutres reliées par une liaison	127
7.3	Première localisation dans le processus de recalage	128
7.4	Seconde localisation dans le processus de recalage	129
7.5	Définition des sous-structures du problème	130
8.1	Principe du support SYLDA5	134
8.2	Constitution du SYLDA5	135
8.3	Protocole d'expérimentation du SYLDA5	135
8.4	Premiers modes propres théoriques de l'ensemble d'étude du SYLDA5	137
8.5	Premières étapes de localisation lors du recalage du SYLDA5	139
8.6	Détail des groupes de sous-structures utilisés dans le processus de réduction des méconnaissances de base du SYLDA5	141
8.7	Énergies de déformation modales par groupe pour les huit premiers modes du SYLDA5	142

LISTE DES TABLEAUX

3.1	Caractéristiques du modèle déterministe du treillis à quatre barres	58
3.2	Détail des pulsations propres et modes propres associés au modèle déterministe du treillis à quatre barres	58
3.3	Valeurs à 70 % et 99 % obtenues avec le modèle stochastique du treillis à quatre barres	59
5.1	Valeurs à 70 % et 99 % obtenues avec le modèle de méconnaissances pour le treillis à quatre barres	86
5.2	Comparaison des valeurs à 70 % et 99 % obtenues pour les pulsations propres du treillis à quatre barres, respectivement pour le modèle avec méconnaissances et pour une simulation de Monte Carlo	87
5.3	Comparaison des valeurs à 70 % et 99 % obtenues pour les valeurs en un DDL des modes propres du treillis à quatre barres, respectivement pour le modèle avec méconnaissances et pour une simulation de Monte Carlo	87
5.4	Valeurs à 70 % et 99 % obtenues avec un modèle de méconnaissance « groupé » pour le treillis à quatre barres	89
5.5	Méconnaissances de base du treillis à quatre barres dans le cas d'un modèle de sous-structure très simplifié par rapport à la réalité	90
5.6	Modélisation des méconnaissances de base du treillis à quatre barres dans le cas d'une forte méconnaissance	94
5.7	Méconnaissances effectives pour le treillis à quatre barres et une forte méconnaissance sur la barre b_{13}	94
6.1	Caractéristiques du treillis à six barres	105
6.2	Caractéristiques relatives des lois de probabilité permettant la simulation des raideurs « expérimentales » des six barres du treillis	105
6.3	Modèle initial de méconnaissances de base pour le treillis à six barres	106
6.4	Résultats après une première réduction des méconnaissances de base du treillis à six barres	108
6.5	Résultats finaux après une seconde réduction des méconnaissances de base du treillis à six barres	109
6.6	Comparaison modèle-expérimental pour les pulsations propres du treillis à six barres	109
6.7	Comparaison modèle-expérimental pour les modes propres du treillis à six barres	109
6.8	Comparaison des valeurs à P % pour deux familles de structures simulées par des répartitions de raideur de lois de probabilité de natures différentes	110
6.9	Caractéristiques relatives des lois de probabilité utilisées dans la simulation stochastique du treillis à six barres dans le cas d'une forte méconnaissance	111

6.10	Modèle initial de méconnaissances de base pour le treillis à six barres dans le cas d'une forte méconnaissance	112
6.11	Résultats de la réduction des méconnaissances de base du treillis à six barres avec forte méconnaissance, sans prise en compte de cette dernière .	112
6.12	Résultats de la réduction des méconnaissances de base du treillis à six barres avec forte méconnaissance, avec prise en compte de cette dernière .	113
6.13	Caractéristiques relatives des lois de probabilité permettant la simulation des valeurs « expérimentales » des coefficients matériaux de la plaque orthotrope	114
6.14	Modèle initial de méconnaissances de base pour la plaque orthotrope . . .	115
6.15	Résultats de la réduction de la méconnaissance de base sur la plaque orthotrope suivant l'information expérimentale utilisée	115
7.1	Caractéristiques du problème de poutres reliées par une liaison	127
7.2	Modèle initial de méconnaissances de base dans le cas des poutres reliées par une liaison	130
7.3	Résultats de la réduction des méconnaissances de base dans le cas des deux poutres reliées par une liaison	130
8.1	Erreur en relation de comportement et erreur en fréquence pour les dix premiers modes propres du SYLDA5	138
8.2	Corrections successives des paramètres du modèle	139
8.3	Modèle initial de méconnaissances de base pour le SYLDA5	141
8.4	Résultats de la réduction des méconnaissances de base du SYLDA5 . . .	142
8.5	Résultats de la réduction des méconnaissances de base du SYLDA5, avec prise en compte de la forte méconnaissance sur le modèle associé au sol .	143
8.6	Comparaison modèle-expérimental pour la première pulsation propre du SYLDA5	144

Dans le domaine de la mécanique comme dans bien d'autres domaines, la simulation prend une place de plus en plus importante ; même si les méthodes numériques permettent de calculer les solutions à des problèmes de plus en plus complexes, on est amené à s'interroger sur la validité des résultats obtenus à l'aide du modèle utilisé. C'est pourquoi **la quantification de la qualité d'un modèle reste encore aujourd'hui une question majeure**, tout particulièrement en dynamique basses fréquences des structures. Deux domaines ont considérablement progressé grâce aux nombreuses recherches entreprises depuis vingt-cinq ans.

Le premier concerne **la quantification de la qualité du modèle numérique utilisé par rapport à la solution « analytique »** (non nécessairement connue) du problème de référence continu que l'on s'est posé ; trois grandes classes de méthodes coexistent, utilisant des concepts bien distincts et introduisant les estimateurs suivants :

- un estimateur construit sur les résidus d'équilibre [Babuska et Rheinboldt 1979] ;
- un estimateur basé sur l'erreur en relation de comportement [Ladevèze et Leguillon 1983, Ladevèze et Oden 1998, Ladevèze et Pelle 2001] ;
- un estimateur utilisant le lissage des contraintes [Zienkiewicz et Zhu 1987].

Bien que cette première approche soit nécessaire, on ne peut garantir que le modèle utilisé soit représentatif de la réalité du phénomène mécanique étudié : **dans le domaine de la confrontation avec une référence expérimentale, de nombreuses méthodes ont été développées pour le recalage des matrices de rigidité, de masse et d'amortissement dans les modèles dynamiques, à partir de résultats d'essais en vibrations libres ou forcées**. On trouvera dans [Mottershead et Friswell 1993] un état de l'art des diverses techniques de recalage existantes ; les plus utilisées sont les méthodes dites paramétriques pour lesquelles les corrections des matrices élémentaires du modèle Éléments Finis sont liées à la variation de certains paramètres physiques du modèle [Golinval et Collignon 1996, Moine *et al.* 1997, Pascual Gimenez *et al.* 1998, Humbert *et al.* 1999, Balmès 2000] ; on citera plus particulièrement la méthode développée au sein du LMT-Cachan qui repose sur le principe de l'erreur en relation de comportement modifiée [Ladevèze 1993, Ladevèze et Chouaki 1999], récemment étendu au domaine de l'acoustique dans [Découvreur *et al.* 2004].

Toutefois, **un modèle numérique, même bien discrétisé et recalé, peut ne pas rendre compte correctement de certains phénomènes**. De façon générale, on distingue :

- des sources d'erreurs sur les charges et leur cumul, ce qui, dans les problèmes de pilotage en aérospatiale par exemple, constitue un point crucial ;
- des sources d'erreurs structurales : ainsi, il peut survenir des dispersions dans les caractéristiques des matériaux utilisés dans l'élaboration de la structure, ou bien encore la modélisation de certaines parties, comme les liaisons, peut être simplifiée, par exemple en adoptant en première approximation un modèle linéaire.

Pour tenir compte de ces incertitudes, il est nécessaire d'utiliser des méthodes non déterministes.

Méthodes stochastiques et validation de modèles

Les premières approches non déterministes sont basées sur l'usage des probabilités, et ceci pour une raison historique : apparu il y a plus de trois siècles, les concepts n'ont cessé de s'affiner, et les moyens numériques aidant, on a fini par arriver à la création d'une discipline à part entière, traitant du calcul stochastique en mécanique comme d'un moyen de modéliser toutes les incertitudes du système étudié ; un état de l'art a été donné dans [Schuëller *et al.* 1997], puis mis à jour dans [Schuëller 2001].

De façon générale, **ces méthodes stochastiques consistent à étudier les effets des incertitudes affectant les paramètres du modèle sur la variabilité de la sortie**. Ces paramètres incertains peuvent être :

- des variables aléatoires, à valeurs dans $\mathbb{R}$, caractérisées chacune par une densité de probabilité ou une fonction de répartition, et les unes vis-à-vis des autres par des fonctions de corrélation ;
- des champs stochastiques, c'est-à-dire des fonctions dépendant à la fois de l'aléa et de l'espace, caractérisés chacun par la donnée d'une moyenne et d'une variance dépendant de l'espace, et d'une fonction de covariance entre deux points de l'espace.

Bien qu'ils soient de manipulation plus complexe, les champs stochastiques représentent de façon plus réaliste les paramètres mécaniques incertains qui montrent une dépendance spatiale non négligeable. Il est toutefois possible de décomposer un champ stochastique sur une base d'espace déterministe, les coordonnées dans cette base étant des variables aléatoires non corrélées ; il s'agit de la décomposition de Karhunen-Loève [Loève 1977]. L'étude des variables aléatoires ainsi obtenues se révèle alors plus aisée, même si en pratique on doit souvent se contenter d'une décomposition approchée, tronquée après quelques termes.

Pour calculer la réponse aléatoire d'une structure modélisée par des paramètres incertains, une première technique s'appuie sur le savoir-faire présent dans l'étude déterministe du problème. Ainsi, **la méthode de Monte Carlo permet de générer des réalisations (ou tirages) des paramètres aléatoires définissant le modèle selon les lois de probabilité et les fonctions de corrélation supposées** ; on est alors en présence, pour chaque tirage des différents paramètres aléatoires, d'une structure déterministe pour laquelle un calcul déterministe de la réponse peut être mené bien souvent au travers d'une méthode de discrétisation spatiale de type Éléments Finis. Il reste enfin à étudier la statistique de ces réponses pour pouvoir caractériser la variabilité de la réponse du système stochastique.

Dans la simulation selon la technique de Monte Carlo, la méthode des Éléments Finis n'est en fait qu'un moyen de mener le calcul déterministe, et **c'est dans le but de mêler plus profondément celle-ci avec le modèle stochastique que s'est développée la méthode des Éléments Finis Stochastiques** (*Stochastic Finite Elements Method* ou *SFEM* en anglais), dont les deux variantes principales sont répertoriées ci-dessous :

- **la méthode des Éléments Finis Stochastiques par Perturbation** (*Perturbation SFEM* en anglais), qui utilise, avec une discrétisation spatiale de type Éléments Finis, une décomposition en série de Taylor afin d'exprimer des relations linéaires sur la base d'une méthode de perturbation entre les caractéristiques de la réponse aléatoire et les paramètres incertains du modèle [Shinozuka et Yamazaki 1988] ;

- **la méthode Spectrale des Éléments Finis Stochastiques** (*Spectral SFEM* en anglais), qui utilise une projection sur le « chaos polynomial » de la partie aléatoire de la réponse cherchée, en association avec une décomposition de Karhunen-Loève des paramètres incertains du modèle, le tout associé à une représentation spatiale par Éléments Finis [[Ghanem et Spanos 1991](#)].

Même si elles propagent ainsi les incertitudes sur les paramètres, et permettent une comparaison entre la variabilité de la réponse du modèle et celle effectivement constatée dans la réalité, **peu de méthodes s’attachent spécifiquement à la validation en tant que telle de modèles mécaniques stochastiques**. C’est en fait bel et bien au travers de cette confrontation que commence à s’esquisser une démarche de validation.

Ainsi, on peut imaginer réaliser à partir de cette comparaison une sorte de propagation inverse des incertitudes, comme ce qui est présenté dans [[Hemez et al. 2004](#)] où les auteurs entreprennent de « calibrer » les paramètres incertains du modèle à l’aide d’une fonction coût définissant une probabilité *a posteriori* d’avoir les paramètres corrects, eu égard aux mesures expérimentales effectuées ; on a alors défini à travers cette fonction coût une quantification plus ou moins satisfaisante de la qualité du modèle stochastique vis-à-vis de la référence expérimentale, et donné le moyen d’obtenir, par minimisation de cette fonction coût, les « meilleurs » paramètres stochastiques du modèle.

Cette technique est à rapprocher de la pléthore des méthodes de calibration statistique [[Beven et Binley 1992](#), [Kennedy et O’Hagan 2001](#)], mais aussi du domaine des méthodes dites inverses [[Tarantola 1987](#), [Menke 1984](#)], qui entrent toutes pleinement dans le cadre stochastique, mais où la mécanique est dans la grande majorité des cas quelque peu absente. Pour combler cette lacune, une nouvelle voie est en cours d’exploration au LMT-Cachan, en ce qui concerne l’idée de vouloir recalibrer les paramètres incertains d’un modèle stochastique à l’aide de données expérimentales : la démarche se situe dans le prolongement des travaux déjà menés dans le cadre déterministe à l’aide du concept de l’erreur en relation de comportement modifiée ; pour plus de précisions, on pourra consulter [[Ladevèze et al. 2005](#)].

Méthodes non stochastiques

Malgré le rôle prédominant joué par les méthodes stochastiques dans la modélisation des incertitudes, certains auteurs se sont peu à peu montrés prudents avec l’utilisation d’une telle démarche qui peut être restrictive : en effet, la volonté d’associer à tout problème de quantification des incertitudes une approche probabiliste peut être exagérée dans des cas où le manque d’informations est tel que supposer une modélisation de cette sorte revient déjà à déformer le problème initial de représentation des incertitudes.

En fait, l’incertitude sur un paramètre peut se définir selon deux schémas différents, comme ce qui est décrit dans [[Klir 1994](#)] :

- **un schéma conflictuel** où des valeurs alternatives et distinctes du paramètre peuvent se produire ; c’est ce que permettent déjà de faire les méthodes stochastiques classiques ;

- **un schéma de non-spécificité** relatif à un ensemble de valeurs alternatives toutes possibles de façon égale ; une illustration en est la réalisation d'une mesure à l'aide d'un appareil, à laquelle est associé un intervalle de confiance ; ce cas de figure ne peut être représenté par la théorie classique des probabilités.

Bien entendu, il arrive que ces deux schémas aient lieu simultanément, et **dès lors que le second cas de figure se produit, une modélisation stochastique entraîne nécessairement une déformation du problème traité**, inacceptable dans certains cas, d'où le recours à des méthodes non stochastiques.

Une première façon naturelle de traiter le schéma de non-spécificité précédemment cité est d'utiliser une méthode d'intervalles déterministes. L'incertitude sur un paramètre est tout simplement modélisée par le fait que ce paramètre se trouve quelque part entre deux bornes données, et ce avec une complète certitude. Il est alors possible de propager cet intervalle au travers d'un modèle quelconque par la seule application des règles d'arithmétique propres aux intervalles. **Cette façon de procéder constitue la méthode des intervalles** (*Interval Analysis* en anglais), introduite dans [Moore 1966], et **qui a été récemment associée à la méthode des Éléments Finis** [Dessombz *et al.* 2001].

Dans les années soixante, une nouvelle théorie des ensembles qui permet de traiter le schéma de non-spécificité a été définie par Zadeh : il s'agit de **la théorie des ensembles flous** (*Fuzzy Sets* en anglais), introduite dans [Zadeh 1965], qui étend la théorie classique des probabilités en relâchant, à travers le concept de participation à un ensemble, la propriété primordiale de la définition d'un ensemble classique, à savoir qu'un élément ne peut qu'appartenir, ou ne pas appartenir, à cet ensemble. Cette théorie a permis d'intégrer des concepts purement linguistiques dans des raisonnements faisant intervenir des modèles mathématiques ; avec une telle approche, on peut alors modéliser de façon relativement simple le jugement d'un expert portant sur des concepts subjectifs. Les outils de la logique floue (*Fuzzy Logic* en anglais) permettent ensuite de propager ces ensembles flous à travers un modèle quelconque. **Dans le domaine de la mécanique, le concept a connu récemment un certain succès, allant jusqu'à la création d'une méthode des Éléments Finis Flous** (*Fuzzy FEM* en anglais), introduite dans [Rao et Sawyer 1995].

À peu près à la même époque s'est développée une nouvelle théorie mathématique définissant ce que l'on pourrait qualifier d'**intervalles aléatoires** (*random intervals* en anglais), qui est due aux travaux successifs de Dempster dès [Dempster 1968], puis de Shafer qui l'intègre quelques années plus tard dans un cadre plus général au sein de la « **théorie des preuves** » (de l'anglais *Evidence Theory*) [Shafer 1976]. En schématisant un peu, on peut dire que la notion de probabilité est remplacée par deux quantités :

- **la croyance** (ou *Belief* en anglais) d'un événement, qui est la somme des « preuves » en faveur de l'occurrence de cet événement ;
- **la plausibilité** (ou *Plausibility* en anglais) d'un événement, qui est le complément de la somme de toutes les « preuves » qui vont à l'encontre de l'occurrence de cet événement.

Ces deux quantités sont alors souvent nommées respectivement probabilités supérieure et inférieure, et permettent ainsi de traduire la notion de non-spécificité.

La **modélisation possibiliste** est une théorie analogue, qui introduit elle aussi deux mesures complémentaires de l'incertitude : la possibilité et la nécessité [Dubois et Prade 1986]. De même que la théorie des preuves, elle n'a que peu d'applications pratiques en mécanique.

Il est bon de souligner que **toutes ces approches ont été réunifiées et hiérarchisées dans la théorie généralisée de l'information** (*Generalized Information Theory*, ou *GIT*, en anglais) introduite dans [Klir 1991], qui considère que toute quantification de l'incertitude est forcément liée à l'information que l'on peut porter à notre connaissance, et c'est d'ailleurs sous cet angle que nous reprendrons dans le détail toutes ces méthodes dans la seconde partie de ce mémoire.

Plus récemment, une généralisation originale du concept d'intervalle déterministe a été menée par Ben-Haim avec la définition des **modèles convexes d'incertitude** (*Convex Models of Uncertainty* en anglais) [Ben-Haim et Elishakoff 1990]. À tout paramètre incertain, on peut associer un ensemble convexe susceptible ou non de le contenir ; ceci rend possible l'élaboration de « modèles de lacunes d'informations » (*information-gap models* en anglais) qui ont conduit certains auteurs à appliquer ce concept dans le domaine de la validation. Dans [Ben-Haim et Hemez 2004], il est montré à l'aide de cette approche que la fidélité du modèle vis-à-vis des données expérimentales se fait nécessairement au détriment de la robustesse vis-à-vis des incertitudes, et que l'obtention d'un compromis entre ces deux critères s'avère indispensable.

Vers une théorie des méconnaissances

De tout ce qui précède résultent deux points principaux concernant la prise en compte des incertitudes dans un modèle mécanique, et la façon de quantifier la qualité de cette prise en compte par rapport à la réalité :

- **l'essentiel des efforts s'est concentré sur les méthodes stochastiques et leur application en mécanique**, avec la mise au point de méthodes efficaces d'étude de modèles mettant en jeu des variables aléatoires ou des champs stochastiques, mais **par contre, il n'existe quasiment aucun moyen basé sur la mécanique de juger quantitativement de la validité de ces modèles ; nous précisons ce constat dans la première partie de ce mémoire ;**
- **les méthodes stochastiques ont pour inconvénient de ne prendre en compte qu'un seul type d'incertitude**, ce qui est très restrictif et peut conduire à des aberrations dans le cas où les informations sur le système à modéliser sont très rares ; malheureusement, **les méthodes non stochastiques sont encore peu employées dans le domaine de la mécanique**, même si une progression sensible commence à être observée, en particulier avec les ensembles flous et les modèles convexes d'incertitudes ; **nous présentons une description rapide de toutes ces méthodes, ainsi que leurs éventuelles applications en mécanique, dans la seconde partie de ce mémoire.**

C'est dans cet état des lieux que l'on peut placer une nouvelle théorie dont les concepts et les résultats sont détaillés dans ce mémoire : la théorie des méconnaissances introduite dans [Ladevèze 2002 - 2003a] à la suite de plusieurs travaux préliminaires conduits au sein du Département « Mécanique des structures lanceurs » du groupe EADS. Cette théorie, qui concerne pour le moment l'étude des réponses modales linéaires en basses fréquences, repose sur **la vision d'une réalité expérimentale, représentée par une famille de structures semblables, mais non rigoureusement identiques, à partir de laquelle sont extraites des quantités caractéristiques que nous décrivons dans la troisième partie de ce mémoire.**

Le concept de méconnaissance repose sur deux idées de base, de nature différente, permettant de traiter d'une façon originale les incertitudes structurales.

La première idée est de nature mécanique : elle consiste à globaliser les incertitudes à l'échelle des sous-structures constitutives de la structure étudiée. On ne cherche donc pas, contrairement aux méthodes stochastiques par exemple, à caractériser précisément les dispersions des paramètres matériaux localement, mais seulement à dégager l'effet de ces dispersions à une échelle supérieure. C'est ainsi que la définition des méconnaissances dites de base (celles attachées aux différentes sous-structures) est liée aux énergies de déformation par sous-structure - vu que pour l'instant, on ne s'intéresse qu'à des incertitudes de type rigidité structurale.

La seconde idée concerne plutôt le traitement mathématique du concept de méconnaissance : elle consiste à s'affranchir de la modélisation purement probabiliste, que l'on serait tenté d'introduire, en définissant la méconnaissance de base sur une sous-structure comme une variable située quelque part dans un intervalle dont les bornes sont des variables aléatoires, au lieu de la définir elle-même directement comme une variable aléatoire. Une telle modélisation permet donc de pallier les inconvénients des approches purement stochastiques dans des cas où la masse d'informations disponibles est si faible qu'une modélisation à l'aide de lois de probabilité serait abusive.

C'est dans la quatrième partie de ce mémoire que nous indiquons en particulier quelles lois de probabilité on peut choisir pour la modélisation des méconnaissances de base ; en particulier, la question est de traduire effectivement les lois de probabilité qui portent sur les bornes des intervalles de méconnaissances de base. Pour cela, nous étudions l'opération inverse, qui permet de transcrire une loi de probabilité portant sur une variable aléatoire en termes de méconnaissances de base, en définissant quelques outils mathématiques concernant l'introduction d'une « probabilité d'intervalle ». Une comparaison avec les méthodes non stochastiques de représentation des incertitudes nous permet enfin de constater la pertinence de la modélisation retenue.

Ainsi défini, le concept de méconnaissance a pour but de chiffrer simplement la comparaison que l'on veut effectuer entre le modèle et le réel. Cette comparaison doit se réaliser sur des quantités accessibles à l'expérimentation, aussi **la première application du concept de méconnaissance doit être de pouvoir transmettre l'incertitude qu'il représente au niveau des grandeurs observées : ceci constitue le thème de notre cinquième partie.**

Ainsi, si l'on considère une quantité d'intérêt définie sur l'ensemble de la structure (une fréquence propre par exemple), **il est possible à partir de la donnée du modèle de méconnaissances de base de calculer un intervalle d'appartenance de cette quantité d'intérêt, dont les bornes sont probabilistes** ; à l'aide d'hypothèses bien précises, il s'agit en fait de traiter le problème de propagation des intervalles de méconnaissances de base au travers du modèle mécanique, à ceci près que ces intervalles ont des bornes qui sont des variables aléatoires. C'est encore grâce au concept de probabilité d'intervalle que l'on peut parvenir à réaliser en pratique cette transmission et à obtenir des valeurs directement comparables : celles-ci sont issues des intervalles de méconnaissances dites effectives, qui ont une certaine probabilité P de contenir la quantité d'intérêt associée au modèle. **Le concept de méconnaissance rend donc possible la manipulation d'un modèle pourtant imparfait par rapport à la réalité considérée.**

La propagation des méconnaissances repose sur une approximation au premier ordre de la variation des quantités d'intérêt modales, ce qui limite l'ordre de grandeur des méconnaissances de base employées ; c'est pourquoi une version permettant de prendre en compte une « forte » méconnaissance sur une sous-structure donnée est également introduite. **L'obtention des méconnaissances effectives sur diverses grandeurs modales est précisée dans la cinquième partie. C'est à l'aide de ces quantités que s'effectue la comparaison avec la réalité** : on peut associer à cette dernière des valeurs caractéristiques qui encadrent une certaine proportion P des grandeurs expérimentales d'intérêt, et comparer ces quantités aux méconnaissances effectives déduites du modèle avec méconnaissances. **C'est ce que nous faisons dans cette cinquième partie sur un exemple académique** avec diverses possibilités de modélisations probabilistes des méconnaissances de base ; le cas d'une forte méconnaissance est également traité.

Le problème de cette comparaison est exposé dans la sixième partie de ce mémoire : il s'agit en effet de savoir comment déterminer les méconnaissances de base du modèle qui soient représentatives de la dispersion observée sur les quantités d'intérêt expérimentales mesurées ; ce point est abordé sous la forme d'une réduction des méconnaissances de base, par l'apport d'informations expérimentales pertinentes, à partir d'une certaine description majorante fixée *a priori*. Par cette idée, on se rapproche quelque peu de la théorie généralisée de l'information [Klir 1991] qui considère l'apport d'information comme la source de réduction des incertitudes.

Ainsi, **plus on a d'informations expérimentales, plus on est susceptible de réduire les méconnaissances de base** : cette réduction s'effectue sous-structure par sous-structure, après sélection de l'information expérimentale la plus pertinente, et en envisageant les cas les plus défavorables en ce qui concerne les contributions des autres sous-structures dans le calcul des méconnaissances effectives. La notion de représentativité d'un essai expérimental est également abordée. Nous détaillons aussi la mise en œuvre de cette technique sur un exemple académique à des fins d'illustration, et concluons sur les premières bases de l'élaboration d'une véritable stratégie de réduction des méconnaissances de base à l'aide de l'information expérimentale disponible.

Une fois résolu le problème de détermination des méconnaissances de base, **nous avons voulu montrer, dans une septième partie, où se situe la théorie des méconnaissances dans la démarche générale de validation d'un modèle** ; sur l'exemple élémentaire d'un modèle de liaison très simplifié par rapport à la réalité, nous voyons que le problème classique de recalage, traité selon la méthode de l'erreur en relation de comportement modifiée, conduit à une erreur résiduelle, qui peut être décrite efficacement à l'aide du concept de méconnaissance.

Dans une huitième et dernière partie, nous appliquons notre stratégie de réduction des méconnaissances de base à un exemple de structure réelle, à savoir le support de satellites SYLDA5 présent dans le lanceur Ariane 5. Nous pouvons ainsi conclure sur la qualité des différents modèles de sous-structures utilisés dans cet exemple.

Tous les calculs sont réalisés à l'aide du logiciel *Matlab* [[MathWorks](#)], et de la boîte d'outils associée *Structural Dynamics Toolbox* [[SDT](#)] permettant de traiter par la méthode des Éléments Finis des problèmes courants de vibrations des structures.

Nous concluons enfin ce travail sur les premiers enseignements, ainsi que les problématiques importantes, engendrés par cette nouvelle méthode qu'est la théorie des méconnaissances.

État de l'art des méthodes stochastiques en calcul des structures

Le but principal du calcul des probabilités, comme son nom même l'indique, est de calculer les probabilités des phénomènes complexes en fonction des probabilités, supposées connues, de phénomènes plus simples.

Émile Borel

Sommaire

1.1	Petit historique succinct des probabilités	10
1.2	Choix des paramètres stochastiques	11
1.2.1	Types de paramètres stochastiques	11
1.2.2	Choix des lois de probabilité	13
1.2.3	Sélection des paramètres les plus pertinents	15
1.3	Méthode de Monte Carlo	16
1.3.1	Avantages et inconvénients	16
1.3.2	Améliorations de la méthode	17
1.4	Méthode des Éléments Finis Stochastiques	18
1.4.1	Méthode des Éléments Finis Stochastiques par Perturbation	18
1.4.2	Méthode Spectrale des Éléments Finis Stochastiques	20
1.5	Autres méthodes stochastiques	23
1.5.1	Méthodes fiabilistes	23
1.5.2	Méthodes non-paramétriques	25
1.6	Méthodes stochastiques et validation de modèles	26

Ce premier chapitre dresse un état de l'art, non exhaustif, des différentes méthodes stochastiques, prenant en compte les incertitudes sous la forme de variables aléatoires et de champs stochastiques dans les modèles de mécanique, dans le calcul des structures ; il ne sera donc pas fait état ici des méthodes probabilistes ou statistiques concernant la modélisation microscopique des matériaux et de leur comportement. Nous étudions ensuite les concepts et les applications que peuvent apporter ces méthodes au domaine de la validation de modèles.

1.1 Petit historique succinct des probabilités

Les méthodes stochastiques sont les méthodes non déterministes les plus utilisées ; la raison essentielle de ce succès vient tout simplement du fait qu'elles reposent sur la théorie des probabilités, dont l'utilisation s'impose naturellement lorsqu'il s'agit de modéliser les effets du hasard, et donc par extension les incertitudes que l'on veut introduire dans les modèles. L'« omniprésence » des probabilités est due quant à elle à la longue maturation historique des concepts qui leur sont liés.

Les toutes premières notions de probabilité datent du seizième siècle et concernent les jeux de hasard (en particulier, les jeux de dés) ; les concepts mathématiques associés commencent à se formaliser dès le début du dix-huitième siècle, à travers les travaux de [Bernoulli](#), de [Bayes](#) et de [Laplace](#).

Très vite pour les mathématiciens de l'époque, la question se pose de pouvoir attribuer des lois de probabilité aux événements qui les entourent : trois visions des probabilités initiées dans les premiers moments sont restées pertinentes, et servent encore aujourd'hui dans le choix d'une loi de probabilité.

- **une première technique se base sur des considérations de symétrie** ; ainsi, on estime pour le lancer d'un dé équilibré que la probabilité d'obtenir une face donnée est la même pour toutes les faces du dé ;
- **une seconde méthode est fondée sur la fréquence d'apparition de l'événement E** au cours d'« expériences » répétées, et on suppose alors que la probabilité d'occurrence de l'événement E est :

$$P(E) = \lim_{n \rightarrow +\infty} \frac{N(E)}{n} \quad (1.1)$$

avec $N(E)$ le nombre de réalisations de l'événement E en n expériences ;

- **une troisième vision, qualifiée de subjectiviste**, consiste à définir la probabilité P d'un événement comme un pari que l'on est prêt à tenir : cela revient à miser « P contre un » sur l'occurrence de cet événement, et « $1 - P$ contre un » sur la non-réalisation de cet événement, sachant que l'on dispose de suffisamment d'informations pour que ce pari soit honnête. Vue sous cet angle, **la probabilité devient un degré de croyance, et s'enrichit de l'utilisation du théorème de Bayes** [[Bayes 1763](#)] qui permet de calculer des probabilités conditionnelles : ainsi, la probabilité $P(E_1 \mid E_2)$ qu'un événement E_1 se produise si l'on a observé la réalisation d'un événement E_2 s'exprime en fonction des probabilités d'occurrence des événements

E_1 et E_2 respectivement, ainsi que de la probabilité $P(E_2 | E_1)$ d'avoir E_2 sachant que l'on a observé E_1 :

$$P(E_1 | E_2) = \frac{P(E_2 | E_1)P(E_1)}{P(E_2)} \quad (1.2)$$

On verra plus tard l'importance de ce théorème dont **une interprétation possible est sa capacité à relier deux probabilités concernant l'événement E_1** : l'une, $P(E_1)$, antérieure et l'autre, $P(E_1 | E_2)$ postérieure à la connaissance de l'événement E_2 , indiquant ainsi comment intégrer la connaissance de l'événement E_2 . L'intérêt de cette approche subjectiviste se fait alors d'autant plus sentir que les données requises pour une étude de fréquence d'apparition sont rares dans le problème traité.

Après une maturation progressive des concepts jusqu'à l'intervention de Kolmogorov qui fonde définitivement dans [Kolmogorov 1956] la théorie des probabilités dans les termes de la théorie de la mesure, des outils se sont développés de façon importante au cours de la seconde moitié du vingtième siècle avec la forte émergence des moyens informatiques. **Dans le domaine de la mécanique, on est peu à peu arrivé à la création d'une discipline à part entière**, traitant du calcul stochastique en mécanique comme du moyen de modéliser précisément les incertitudes du système étudié, dont nous allons maintenant dresser un panorama non exhaustif ; on trouvera plus de précisions dans un état de l'art dressé dans [Schuëller et al. 1997], et complété dans [Schuëller 2001]. Enfin, pour plus de détails sur la théorie des probabilités elle-même, on pourra par exemple consulter [Grimette et Stirzaker 1991].

1.2 Choix des paramètres stochastiques

Comme nous l'avons déjà évoqué dans l'introduction de ce mémoire, les méthodes stochastiques consistent à étudier les effets des incertitudes affectant les paramètres incertains du modèle sur la variabilité de la sortie. Pour cela, il faut déjà préciser comment on va définir les paramètres incertains en termes stochastiques.

1.2.1 Types de paramètres stochastiques

Les types de paramètres incertains que l'on peut définir dans des modèles stochastiques sont les suivants :

- **les variables aléatoires**, à valeurs dans $\mathbb{R}$, généralement caractérisées chacune par une densité de probabilité ou une fonction de répartition, et les unes vis-à-vis des autres par des fonctions de corrélation (idéalement, une description complète de ces variables serait caractérisée par la donnée d'une loi conjointe de probabilité, mais ceci est souvent difficile à établir en pratique) ;
- **les champs stochastiques**, c'est-à-dire des fonctions, à valeurs dans $\mathbb{R}$, dépendant à la fois de l'aléa θ et de l'espace $\underline{x}$; ces champs sont généralement caractérisés chacun par la donnée d'une moyenne et d'une variance dépendant de l'espace, et d'une fonction de covariance entre deux points de l'espace ;

- **les processus stochastiques**, c'est-à-dire des fonctions dépendant à la fois de l'aléa θ et du temps t ; bien sûr, on se rend compte que ces processus sont au temps ce que les champs stochastiques sont à l'espace, ce qui implique des possibilités de traitement analogues.

Bien qu'ils soient de manipulation plus complexe, les champs stochastiques représentent de manière bien plus réaliste des paramètres mécaniques incertains pour lesquels la prise en compte de la variabilité spatiale est nécessaire, comme par exemple les rigidités structurales ; en effet, modéliser le module d'Young d'une poutre en acier par une variable aléatoire seulement revient bien à supposer que deux poutres en acier quelconques n'auront pas forcément le même module, mais par contre que le module de chaque poutre est constant sur toute sa longueur, et cette hypothèse implicite d'homogénéité de la rigidité peut se révéler bien trop restrictive suivant les applications.

Décomposition de Karhunen-Loève

Il est toutefois possible de décomposer un champ stochastique $E(\underline{x}, \theta)$ sur une base d'espace déterministe, les coordonnées dans cette base étant des variables aléatoires ; la plus populaire de ce type de méthodes est la décomposition de Karhunen-Loève [Loève 1977] :

$$E(\underline{x}, \theta) = \overline{E}(\underline{x}) + \sum_{i=1}^{\infty} \xi_i(\theta) \sqrt{\lambda_i} \varphi_i(\underline{x}) \quad (1.3)$$

où λ_i et $\varphi_i(\underline{x})$ sont respectivement les valeurs propres et les vecteurs propres du noyau de l'opérateur de covariance défini comme :

$$C_E(\underline{x}_1, \underline{x}_2) = \overline{(E(\underline{x}_1, \theta) - \overline{E}(\underline{x}_1))(E(\underline{x}_2, \theta) - \overline{E}(\underline{x}_2))} \quad (1.4)$$

avec $\overline{Y}(\underline{x})$ l'espérance de $Y(\underline{x}, \theta)$; les λ_i et $\varphi_i(\underline{x})$ sont en fait solutions de l'équation intégrale, souvent résolue numériquement :

$$\int_D C_E(\underline{x}_1, \underline{x}_2) \varphi_i(\underline{x}_1) d\underline{x}_1 = \lambda_i \varphi_i(\underline{x}_2) \quad (1.5)$$

où D désigne le domaine spatial sur lequel le champ stochastique $E(\underline{x}, \theta)$ est défini. On peut montrer que les valeurs propres λ_i sont des quantités décroissantes quand i augmente, et que, plus le champ à décomposer est proche du bruit blanc, plus on a besoin de termes dans la décomposition, ce qui explique qu'en pratique, comme les champs stochastiques utilisés sont assez « lisses », peu de termes sont gardés dans la décomposition de Karhunen-Loève.

Les variables aléatoires $\xi_i(\theta)$ sont quant à elles non corrélées, d'espérances nulles et d'écarts types unitaires :

$$\overline{\xi_i(\theta)} = 0 \quad \text{et} \quad \overline{\xi_i(\theta) \xi_j(\theta)} = \delta_{ij} \quad (1.6)$$

Ces variables aléatoires sont en fait les moyennes du champ $E(\underline{x}, \theta)$ pondérées par les fonctions de l'espace $\varphi_i(\underline{x})$:

$$\xi_i(\theta) = \frac{1}{\sqrt{\lambda_i}} \int_D E(\underline{x}, \theta) \varphi_i(\underline{x}) d\underline{x} \quad (1.7)$$

D'autres méthodes de décomposition existent également : on peut citer par exemple les décompositions en séries de Fourier [Grigoriu 1993], les décompositions orthogonales [Zhang et Ellingwood 1994], ou encore la décomposition sur le chaos polynomial [Ghanem et Spanos 1991], sur laquelle nous reviendrons plus en détail dans le paragraphe 1.4.2. **Dans tous les cas, l'analyse du champ stochastique se résume alors à l'étude d'un certain nombre de variables aléatoires indépendantes, dont la manipulation se révèle plus aisée, même s'il se pose toujours la question de l'effet de la troncature de ces décompositions.**

1.2.2 Choix des lois de probabilité

Un autre point à débattre concerne l'attribution des lois de probabilité aux paramètres incertains. Trois premières techniques proviennent directement des trois visions des probabilités citées précédemment dans le paragraphe 1.1 : les considérations de symétrie, l'étude des fréquences d'apparition, et l'approche subjectiviste. Toutefois, d'autres façons de faire, détaillées ci-dessous, sont encore envisageables.

Loi de probabilité normale

Il est en général très courant de choisir pour une variable aléatoire $Y(\theta)$ une loi de probabilité dite normale, dont la densité de probabilité s'écrit sous la forme :

$$p(y) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(y - \mu)^2}{2\sigma^2}\right) \quad (1.8)$$

où $\mu = \overline{Y(\theta)}$ est l'espérance (ou moyenne) et $\sigma = \sqrt{\overline{(Y(\theta) - \overline{Y(\theta)})^2}}$ l'écart type de la variable aléatoire $Y(\theta)$ ainsi définie.

L'importance de cette loi est due au résultat du théorème central limite, dont l'énoncé est le suivant :

Si l'on considère n variables aléatoires indépendantes $Y_1(\theta), \dots, Y_n(\theta)$ d'espérances respectives μ_k et d'écarts types respectifs σ_k , alors la loi de répartition de :

$$S_n(\theta) = \sum_{k=1}^n Y_k(\theta) \quad (1.9)$$

tend lorsque $n \rightarrow \infty$ vers une loi normale d'espérance μ et d'écart type σ :

$$\mu = \sum_{k=1}^n \mu_k \quad \text{et} \quad \sigma = \sqrt{\sum_{k=1}^n \sigma_k^2}$$

Ce théorème peut s'énoncer sous la forme des conditions dites de Borel :

« Si les causes de variations aléatoires d'une grandeur expérimentale $Y(\theta)$ sont nombreuses et si aucune d'elles n'est prépondérante, alors on peut considérer que les différentes valeurs y_k obtenues en répétant la mesure dans des conditions (pratiquement) identiques se répartissent suivant une loi normale de Gauss. »

Ceci tend à justifier pourquoi les paramètres structuraux tels que la masse ou la rigidité sont souvent modélisés par des variables aléatoires de loi normale. Ceci étant, cette loi pose rigoureusement un problème théorique : vu que son domaine de définition est l'ensemble des réels tout entier, des réalisations non positives peuvent se produire, ce qui n'a pas de sens physique pour des paramètres tels que la masse ou la raideur. Deux solutions sont alors envisageables :

- si, avec la loi normale utilisée, la probabilité d'occurrence d'une valeur négative est suffisamment faible (typiquement inférieure à un ordre de grandeur de 0,1%), on peut faire en sorte de ne pas prendre en compte ces valeurs aberrantes, quitte à renormaliser ensuite la densité de probabilité de la loi ;
- une loi de probabilité, proche de la loi normale, mais dont le domaine de définition est réduit à $\mathbb{R}_+^*$, peut être utilisée : il s'agit de la loi log-normale, dont la densité de probabilité est la suivante :

$$p(y) = \frac{1}{y\sqrt{2\pi a^2}} \exp\left(-\frac{(\ln y - b)^2}{4a^2}\right) \quad (1.10)$$

les paramètres a et b se reliant à l'espérance μ et l'écart type σ de $Y(\omega)$ par les expressions :

$$a = \sqrt{\ln\left(\frac{\mu^2 + \sigma^2}{\mu^2}\right)} \quad \text{et} \quad b = \ln\left(\frac{\mu^2}{\sqrt{\mu^2 + \sigma^2}}\right) \quad (1.11)$$

On peut voir sur la figure 1.1 la comparaison entre les deux lois, l'une normale, l'autre log-normale, de mêmes espérance $\mu = 1$ et écart type $\sigma = 0,2$. On peut reprocher à la loi log-normale la forme non symétrique de sa densité de probabilité autour de la valeur moyenne.

Principe du maximum d'entropie

Une dernière démarche de détermination des lois de probabilité des paramètres du modèle repose sur le principe du maximum d'entropie, qui donne la « meilleure » loi de probabilité parmi toutes celles qui vérifient un ensemble de contraintes.

La notion d'entropie d'une variable aléatoire, introduite dans [Shannon 1948], permet de caractériser la quantité d'information manquante à propos de cette variable aléatoire comme suit :

$$S(Y(\omega)) = - \int_{D_y} p(y) \ln(p(y)) dy \quad (1.12)$$

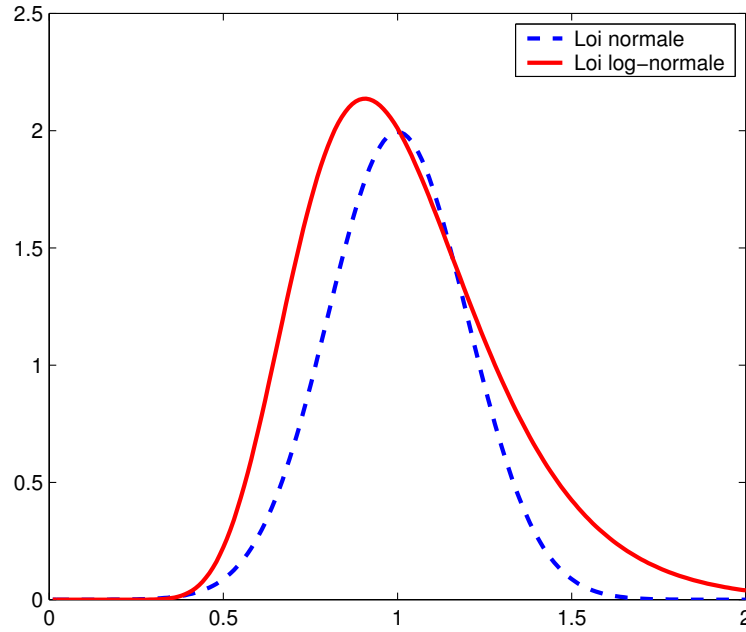


FIG. 1.1 – Densités de probabilité pour une loi normale et une loi log-normale

À partir de là, la méthode du maximum d'entropie, définie dans [Jaynes 1957a - b], permet de choisir la loi de probabilité la plus « honnête », les contraintes à respecter étant données (généralement la donnée des n premiers moments statistiques) : **il s'agit donc de la loi où toute l'information disponible est utilisée, et une incertitude maximale, en ce qui concerne ce que l'on ne connaît pas, est supposée** ; pour plus de précisions sur cette méthode, on pourra consulter [Kapur et Kesavan 1992].

Un exemple simple d'utilisation de ce principe est le suivant : on cherche à déterminer une loi de probabilité sur l'intervalle $[a; b]$, $a \neq b$; la maximisation de l'entropie $S(Y(\omega))$ sous la contrainte de normalisation $\int_a^b p(y)dy = 1$ implique nécessairement :

$$p(y) = \frac{1}{b - a} \quad (1.13)$$

ce qui correspond à une loi uniforme. De manière intuitive, il semble logique en effet *a posteriori* que cette loi permette de maximiser l'incertitude sur un intervalle $[a; b]$, en affirmant qu'aucun point de cet intervalle n'a de raison d'être plus « probable » qu'un autre.

1.2.3 Sélection des paramètres les plus pertinents

On peut également citer une autre famille de méthodes très répandues qui ont leur intérêt dans le processus de validation de modèles : **ces méthodes, issues de la communauté des statisticiens, permettent de déterminer la part de l'effet de l'incertitude de chaque variable du modèle sur l'incertitude finale de la réponse**. On obtient alors une réduction de la liste des variables incertaines du modèle en déterminant celles qui sont les

plus influentes dans la réponse du système. Sans entrer plus dans le détail, on peut citer différentes méthodes comme l'analyse de sensibilité [Saltelli *et al.* 2004], généralement la plus répandue, ou bien encore les études de criblage [Dean et Lewis 2004, Chipman *et al.* 1997].

Nous allons maintenant évoquer comment les méthodes stochastiques permettent de transmettre l'incertitude modélisée par les paramètres stochastiques pour étudier la réponse du modèle.

1.3 Méthode de Monte Carlo

Une première démarche, très répandue, est l'utilisation de la méthode de Monte Carlo, détaillée dans de nombreux ouvrages tels que [Hammersley et Handscomb 1967], [Rubinstein 1981] ou [Fishman 1996], qui est une approche statistique, car l'évaluation des caractéristiques stochastiques de la réponse du système passe par le calcul d'un grand nombre de problèmes déterministes.

En effet, l'utilisation de cette méthode permet de générer des réalisations des paramètres aléatoires définissant le modèle, aussi appelées tirages, qui tiennent compte des lois de probabilité respectives et des fonctions de corrélation entre les différentes variables aléatoires entrant en jeu ; **on obtient ainsi, pour chaque tirage des différents paramètres aléatoires, une structure pour laquelle un calcul déterministe de la réponse peut être mené.** Une étude statistique de ce jeu de réponses nous permet alors de déterminer une moyenne et un écart type, ou encore une probabilité d'occurrence d'un critère mécanique, pour des systèmes qui initialement n'avaient généralement pas de solution analytique.

1.3.1 Avantages et inconvénients

Le principal avantage de la méthode de Monte Carlo est de permettre de mener successivement plusieurs calculs déterministes une fois que les jeux de paramètres ont été tirés ; pour peu que le problème déterministe soit traitable par un code de calcul classique ou spécialement dédié, on peut facilement étudier la réponse du modèle stochastique traité. Néanmoins cet avantage constitue également le principal inconvénient de la méthode car **un nombre suffisamment grand de tirages doit être effectué pour que l'étude statistique de la réponse converge.**

Un résultat de convergence en dimension un peut être donné par une application du théorème central limite (1.9) : on cherche à estimer l'espérance d'une variable aléatoire $Y(\theta)$; la méthode de Monte Carlo pour n tirages nous donne alors une estimation de cette espérance, à l'aide de n variables aléatoires $\{Y_k(\theta)\}_{k=1\dots n}$ de même loi de probabilité que $Y(\theta)$. On montre alors que l'écart

$$\epsilon_n = \overline{Y(\theta)} - \frac{1}{n} \sum_{k=1}^n Y_k(\theta) \quad (1.14)$$

entre l'espérance de $Y(\theta)$ et son estimé suit une loi normale centrée d'écart type $\frac{\sigma}{\sqrt{n}}$, où σ est l'écart type de la variable aléatoire $Y(\theta)$. Autrement dit, **la vitesse de convergence de la méthode de Monte Carlo est donc de l'ordre de $\frac{1}{\sqrt{n}}$, où n est le nombre de tirages.**

Le nombre minimal de tirages requis pour un certain niveau de convergence peut alors devenir considérable suivant l'objet de l'étude statistique, mais aussi suivant le nombre de paramètres aléatoires du problème. C'est pourquoi de nombreuses techniques numériques ont été mises au point pour contrebalancer cette exigence, dont quelques-unes sont présentées ci-dessous.

1.3.2 Améliorations de la méthode

Parallélisation

La structure même de la méthode de Monte Carlo permet une parallélisation aisée des opérations : il suffit de mener les résolutions des systèmes déterministes sur plusieurs processeurs (ou ordinateurs), ce qui réduit ainsi assez fortement les temps de calcul ; on trouvera dans [Johnson *et al.* 1997] et [Papadrakakis et Papadopoulos 1999] quelques exemples récents d'implémentations.

Réduction de variance

Les techniques dites de réduction de variance (faisant partie de ce qu'on appelle également *Importance sampling*) **permettent d'accélérer la convergence de la méthode en augmentant la densité des réalisations dans les régions d'intérêt**, à savoir celles qui contribuent le plus à l'estimation statistique désirée ; en effet, on a vu précédemment que l'écart entre l'espérance cherchée et l'estimation par Monte Carlo suivait une loi normale centrée d'écart type $\frac{\sigma}{\sqrt{n}}$, où σ est l'écart type - c'est-à-dire la racine carrée de la variance - de la variable aléatoire étudiée, d'où l'idée de vouloir « réduire » cette variance.

La principale limitation de ce genre de méthode est liée au nombre de variables aléatoires indépendantes : au-delà de dix, le coût de la technique de réduction de variance devient prohibitif par rapport au coût d'une méthode de Monte Carlo classique. On trouvera plus de détails sur ces techniques dans [Thompson 1992].

Latin Hypercube sampling

La technique dite du *Latin Hypercube sampling*, introduite dans [McKay *et al.* 1979], est très employée ; **sa propriété de « stratification » permet de réduire assez nettement le nombre N de tirages requis** : le domaine de définition de chaque variable aléatoire est divisé en N intervalles d'égales probabilités, et on tire aléatoirement une valeur sur chacun de ces intervalles ; ensuite, les N valeurs ainsi tirées pour la première variable aléatoire sont appariées aléatoirement avec les N valeurs tirées pour la seconde variable aléatoire, formant ainsi N couples qui sont alors associés aléatoirement avec les N valeurs tirées pour la troisième valeur aléatoire, et ainsi de suite pour obtenir N p -uplets, où p est le nombre de variables aléatoires du modèle, qui constituent les N tirages finalement utilisés pour caractériser la réponse.

En pratique, ce nombre N de tirages est bien plus faible que celui requis avec une méthode de Monte Carlo classique pour un même niveau de convergence. Diverses améliorations ont été menées depuis, par exemple dans [Helton et Davis 2003].

Méta-modèles

Les techniques précédentes s'attachent à améliorer la méthode de Monte Carlo elle-même ; **il est bien sûr possible de réduire le temps de calcul requis par les résolutions déterministes en remplaçant le modèle « complet » initial par un modèle approché de substitution qui soit plus rapide à mettre en œuvre** ; on parle souvent de méthodes de surface de réponse (*Response Surface Methods*, ou *RSM* en anglais) [Veneziano et al. 1983, Faravelli 1989], ou encore de méta-modèles, dont le principe consiste à établir une surface approchée de la réponse du modèle.

Pour cela, la réponse du modèle déterministe complet est calculée pour différents jeux de paramètres, de façon à obtenir suffisamment de points de réponse pour interpoler ces derniers par une certaine forme de modèle de régression ; on trouvera quelques exemples récents de cette façon de procéder dans [Schultze et al. 2001, Hemez et al. 2001, Boucard et Champany 2003]. Bien entendu, une étape préliminaire indispensable dans l'élaboration de ces méta-modèles consiste à sélectionner les paramètres du modèle les plus influents dans la réponse, ce à quoi s'attachent les méthodes décrites dans le paragraphe 1.2.3, car comme pour les techniques précédentes, la principale limitation réside dans le nombre de variables aléatoires indépendantes du modèle.

Dans le domaine de la mécanique, les calculs déterministes qui sont réalisés pour les jeux de paramètres tirés font intervenir généralement une discrétisation spatiale du problème, de type Éléments Finis ; on est donc en présence, quoique de façon indirecte, d'une méthode que l'on pourrait qualifier d'Éléments Finis en stochastique.

1.4 Méthode des Éléments Finis Stochastiques

C'est dans l'idée de mêler plus intimement la méthode des Éléments Finis, telle qu'elle existe pour les problèmes déterministes, avec les outils probabilistes que s'est développée la méthode des Éléments Finis Stochastiques (*Stochastic Finite Elements Method* ou *SFEM* en anglais). L'idée est de considérer l'aléa comme une dimension supplémentaire du problème traité, tout en employant la discrétisation spatiale proposée par le concept des Éléments Finis. On peut alors traiter la dimension associée à l'aléa de deux façons générales, ce qui nous conduit aux deux grandes orientations suivantes.

1.4.1 Méthode des Éléments Finis Stochastiques par Perturbation

La première méthode utilise une méthode de perturbation entre les caractéristiques de la réponse aléatoire et les paramètres incertains du modèle : il s'agit de la méthode des Éléments Finis par Perturbation (ou *Perturbation SFEM* en anglais) [Baecher et Ingra 1981].

La première étape concerne la discrétisation spatiale des champs stochastiques utilisés ; pour cela, de nombreuses techniques existent :

- **la discrétisation au milieu de l'élément** (*mid-point method* en anglais) est la façon la plus simple de procéder : elle consiste à associer une variable aléatoire à chaque élément du maillage, en supposant que celle-ci est la valeur du champ stochastique au milieu de l'élément [Shinozuka et Yamazaki 1988] ; les moments statistiques de ces variables aléatoires sont alors obtenus en considérant ceux du champ stochastique aux différents points milieux du maillage ;
- **la discrétisation par moyenne locale** (*local average method* en anglais) consiste cette fois-ci à prendre comme variable aléatoire associée à chaque élément la moyenne du champ stochastique sur cet élément [Vanmarcke et Grigoriu 1983] ; de nombreuses techniques numériques existent alors pour calculer les moments statistiques de ces variables aléatoires en fonction des caractéristiques du champ stochastique initial ;
- **la discrétisation par intégrales pondérées** (*weighted-integral method* en anglais) permet de transformer le champ stochastique initial en un vecteur de variables aléatoires nodales calculées lors de l'intégration numérique requise sur chaque élément pour le calcul de la rigidité élémentaire [Takada 1990a - b, Deodatis 1991a - b] ; dans [Takada 1992], il a été montré que cette discrétisation donnait une très bonne approximation de la variance de la réponse, tandis que la discrétisation au milieu de l'élément avait tendance à la surestimer, et la discrétisation par moyenne locale à la sous-estimer ;
- d'une façon plus indirecte peut être également envisagée une discrétisation basée sur la décomposition de Karhunen-Loève (1.3) décrite dans le paragraphe 1.2.1 ; il suffit pour cela de projeter les fonctions de l'espace obtenues dans la décomposition sur la base de discrétisation Éléments Finis.

La seconde étape consiste à approximer les fonctions des variables aléatoires par leur développement en série de Taylor autour de leur valeur moyenne, à l'ordre un ou deux, en supposant que les variables aléatoires varient peu autour de leurs valeurs moyennes : par exemple, si le problème discrétisé revient à résoudre un système de la forme $KU = F$, des développements à l'ordre deux par rapport aux n variables aléatoires ϵ_k nous donnent :

$$K = K^0 + \sum_{i=1}^n K_i^I \epsilon_i + \sum_{i=1}^n \sum_{j=1}^n K_{ij}^{II} \epsilon_i \epsilon_j \quad (1.15)$$

$$U = U^0 + \sum_{i=1}^n U_i^I \epsilon_i + \sum_{i=1}^n \sum_{j=1}^n U_{ij}^{II} \epsilon_i \epsilon_j \quad (1.16)$$

où K^0 , K_i^I et K_{ij}^{II} désignent respectivement la moyenne, et les dérivées premières et secondes de la matrice de rigidité K par rapport aux variables aléatoires ; U^0 , U_i^I et U_{ij}^{II} sont les quantités analogues définies pour le champ de déplacement solution U que l'on recherche.

La résolution se fait alors successivement ordre par ordre : voici les expressions pour la détermination des ordres zéro, un et deux, avec $i, j = 1 \dots n$:

$$\mathbf{K}^0 U^0 = F \quad (1.17a)$$

$$\mathbf{K}^0 U_i^I = -\mathbf{K}_i^I U^0 \quad (1.17b)$$

$$\mathbf{K}^0 U_{ij}^{II} = -\mathbf{K}_i^I U_j^I - \mathbf{K}_j^I U_i^I - \mathbf{K}_{ij}^{II} U^0 \quad (1.17c)$$

De nombreuses applications à des problèmes linéaires et non-linéaires ont été réalisées, en statique comme en dynamique, avec de bons résultats lorsque les paramètres incertains fluctuent dans des bandes étroites ; citons par exemple [Liu *et al.* 1986a - b - 1988, Shinozuka et Yamazaki 1988] pour les premiers travaux, ainsi que plus récemment [Elishakoff *et al.* 1995, Muscolino *et al.* 2000].

Dans [Van den Nieuwenhof 2003, Van den Nieuwenhof et Coyette 2003] a été développée une approche modale originale qui permet de calculer directement la dispersion des Fonctions de Réponse en Fréquence (*FRF*) pour des paramètres incertains donnés ; pour cela, l'étude de la variabilité du problème aux valeurs propres a été menée dans le cas d'incertitudes concernant aussi bien les rigidités que la géométrie. Numériquement, une méthode mixte de perturbation couplée à des tirages de Monte Carlo permet de traiter, avec un faible coût, la variabilité des *FRF*, dont l'obtention est pourtant issue d'un problème fortement non-linéaire au niveau des pics de résonances.

1.4.2 Méthode Spectrale des Éléments Finis Stochastiques

La seconde méthode utilise une projection sur le « chaos polynomial » de la partie aléatoire de la réponse cherchée, en association avec une décomposition de Karhunen-Loève des paramètres incertains du modèle, le tout épaulé par une représentation spatiale par Éléments Finis : il s'agit de la méthode Spectrale des Éléments Finis Stochastiques (*Spectral SFEM* en anglais).

La toute première étape de la méthode consiste à décomposer les champs stochastiques du modèle sur la base d'espace proposée par la décomposition de Karhunen-Loève décrite dans le paragraphe 1.2.1 : par exemple, si l'on considère que c'est le module d'Young qui est représenté par un champ $E(\underline{x}, \theta)$, on peut écrire :

$$E(\underline{x}, \theta) = \overline{E}(\underline{x}) + \sum_{i=1}^{\infty} \xi_i(\theta) \sqrt{\lambda_i} \varphi_i(\underline{x}) \quad (1.18)$$

où $\xi_i(\theta)$ sont des variables aléatoires non corrélées d'espérances nulles et d'écarts types unitaires, formant ainsi une base orthonormée de l'aléa. Si le champ stochastique $E(\underline{x}, \theta)$ suit une loi normale, alors les variables aléatoires $\xi_i(\theta)$ sont elles aussi normales ; c'est l'hypothèse que l'on suppose dans la suite.

La seconde étape consiste à utiliser cette base orthonormée de variables aléatoires $\xi_i(\theta)$ pour décomposer le vecteur $\hat{U}(\theta)$ contenant les inconnues nodales, provenant de la discrétisation Éléments Finis du champ de déplacement solution $U(\underline{x}, \theta)$ recherché.

On peut montrer que le vecteur de variables aléatoires inconnues $\widehat{U}(\theta)$ peut s'écrire sous la forme d'une somme de polynômes des variables aléatoires normales $\xi_i(\theta)$:

$$\widehat{U}(\theta) = a_0 \Gamma_0 + \sum_{i_1=1}^{\infty} a_{i_1} \Gamma_1(\xi_{i_1}(\theta)) + \sum_{i_1=1}^{\infty} \sum_{i_2=1}^{i_1} a_{i_1 i_2} \Gamma_2(\xi_{i_1}(\theta), \xi_{i_2}(\theta)) + \dots \quad (1.19)$$

où $\Gamma_p(\xi_{i_1}(\theta), \dots, \xi_{i_p}(\theta))$ désigne le chaos polynomial d'ordre p , qui engendre un sous-espace de l'aléa appelé chaos homogène d'ordre p [Wiener 1938] ; on trouvera dans [Ghanem et Spanos 1991] le mode d'obtention de ces polynômes. En pratique, on réordonne les termes sous la forme d'une simple somme, que l'on tronque après le $(P + 1)$ ième terme :

$$\widehat{U}(\theta) = \sum_{j=0}^P U_j \Psi_j(\xi_1(\theta), \dots, \xi_N(\theta)) = \sum_{j=0}^P U_j \psi_j(\theta) \quad (1.20)$$

si l'on ne considère qu'un chaos polynomial de dimension N , c'est-à-dire si seules les N premières variables aléatoires $\xi_i(\theta)$ sont conservées. La relation qui donne le nombre de termes $P + 1$ de la décomposition (1.20) en fonction de l'ordre p et la dimension N du chaos est la suivante :

$$P + 1 = 1 + \sum_{j=1}^p \frac{1}{j!} \prod_{k=0}^{j-1} (N + k) \quad (1.21)$$

Voici le détail de ces $P + 1 = 10$ polynômes pour un ordre deux et une dimension trois ($p = 2$ et $N = 3$) :

$$\{\Psi_j\} = \{1, \xi_1(\theta), \xi_2(\theta), \xi_3(\theta), \xi_1(\theta)\xi_2(\theta), \xi_1(\theta)\xi_3(\theta), \xi_2(\theta)\xi_3(\theta), \xi_1(\theta)^2 - 1, \xi_2(\theta)^2 - 1, \xi_3(\theta)^2 - 1\} \quad (1.22)$$

Ces polynômes vérifient :

$$\overline{\psi_j(\theta)} = 0 \quad \forall j \geq 1 \quad \text{et} \quad \overline{\psi_j(\theta)\psi_k(\theta)} = \overline{\psi_j(\theta)^2} \delta_{jk} \quad (1.23)$$

La détermination du champ déplacement solution revient donc à déterminer les vecteurs U_j de la décomposition (1.20).

La discrétisation Éléments Finis du problème conduit au calcul d'une matrice de rigidité globale issue de l'assemblage de matrices élémentaires de la forme :

$$\mathbf{K}_E(\theta) = \int_E \mathbf{B}^T \mathbf{C}(\underline{x}, \theta) \mathbf{B} dE \quad (1.24)$$

où $\mathbf{B}$ est la matrice gradient exprimant les déformations en fonction du vecteur des inconnues nodales, et $\mathbf{C}(\underline{x}, \theta)$ est la matrice de relation de comportement, faisant intervenir dans notre cas le champ $E(\underline{x}, \theta)$. La décomposition de Karhunen-Loève (1.18) de ce champ, tronquée à partir du N ième terme, nous donne formellement :

$$\mathbf{K}_E(\theta) = \sum_{i=0}^N \xi_i(\theta) \mathbf{K}_{Ei} \quad (1.25)$$

où $\mathbf{K}_{E,i,i>1}$ correspond à l'intégration spatiale sur E de la quantité $\sqrt{\lambda_i}\varphi_i(\underline{x})$ issue de la décomposition, et où l'on a posé pour plus de commodité $\mathbf{K}_{E,0} = \overline{\mathbf{K}_E}$ et $\xi_0(\theta) = 1$.

Le problème Éléments Finis revient alors à résoudre un système de la forme $\mathbf{K}U = F$, ce qui nous amène, en utilisant les décompositions (1.25) et (1.20), à résoudre l'équation suivante :

$$\sum_{j=0}^P \left(\sum_{i=0}^N \xi_i(\theta) \mathbf{K}_i \right) \psi_j(\theta) U_j = F \quad (1.26)$$

que l'on multiplie par $\psi_k(\theta)$, ce qui nous donne, en prenant l'espérance :

$$\sum_{j=0}^P \sum_{i=0}^N \overline{\xi_i(\theta) \psi_j(\theta) \psi_k(\theta) \mathbf{K}_i} U_j = \overline{F \psi_k(\theta)}, k = 0, \dots, P \quad (1.27)$$

Au bilan, on doit résoudre le système par blocs suivant :

$$\begin{pmatrix} \widehat{\mathbf{K}}_{00} & \widehat{\mathbf{K}}_{01} & \cdots & \widehat{\mathbf{K}}_{0P} \\ \widehat{\mathbf{K}}_{10} & \widehat{\mathbf{K}}_{11} & \cdots & \widehat{\mathbf{K}}_{1P} \\ \vdots & \vdots & \ddots & \vdots \\ \widehat{\mathbf{K}}_{P0} & \widehat{\mathbf{K}}_{P1} & \cdots & \widehat{\mathbf{K}}_{PP} \end{pmatrix} \begin{pmatrix} U_0 \\ U_1 \\ \vdots \\ U_P \end{pmatrix} = \begin{pmatrix} F \\ 0 \\ \vdots \\ 0 \end{pmatrix} \quad (1.28)$$

où chaque bloc est défini par :

$$\widehat{\mathbf{K}}_{kj} = \sum_{i=0}^N \overline{\xi_i(\theta) \psi_j(\theta) \psi_k(\theta) \mathbf{K}_i} \quad \forall j, k = 0 \dots P \quad (1.29)$$

La matrice finale étant d'ordre $N(P+1)$, on a tout intérêt à ne pas l'assembler explicitement ; dans [Ghanem et Kruger 1996], deux algorithmes de résolutions sont proposés : le premier repose sur un schéma itératif basé sur l'algorithme du Gradient Conjugué (pour plus de précisions, voir [Axelsson 1994]), tandis que le second repose sur une hiérarchisation des systèmes à résoudre, rendue possible par le caractère lui-même hiérarchique de la technique du chaos polynomial.

Des méthodes hybrides ont également été développées dans [Ghanem 1999] : elles permettent de lier la *SFEM* à la méthode de Monte Carlo et ses améliorations décrites dans le paragraphe 1.3.2, telles que l'*Importance sampling* ou le *Latin Hypercube sampling*.

Dans le cas où le champ stochastique $E(\underline{x}, \theta)$ ne suit pas une loi de probabilité normale, on n'utilise plus la décomposition de Karhunen-Loève (1.18) : à la place, on applique directement la décomposition sur le chaos polynomial, à savoir :

$$E(\underline{x}, \theta) = \sum_{i=0}^P e_i(\underline{x}) \psi_i(\theta) \quad (1.30)$$

où les vecteurs de la base d'espace $e_i(\underline{x})$ sont donnés par :

$$e_i(\underline{x}) = \frac{\psi_i(\theta) E(\underline{x}, \theta)}{\overline{\psi_i(\theta)^2}} \quad (1.31)$$

Le raisonnement reste alors le même, les blocs matriciels s'exprimant maintenant comme suit :

$$\widehat{\mathbf{K}}_{kj} = \sum_{i=0}^P \overline{\psi_i(\theta) \psi_j(\theta) \psi_k(\theta)} \mathbf{K}_i \quad \forall j, k = 0 \dots P \quad (1.32)$$

avec $\mathbf{K}_i$ résultant de l'intégration spatiale de $e_i(\underline{x})$.

De la même façon, si le problème à traiter est non-linéaire, la matrice de rigidité n'est plus une fonction linéaire des variables aléatoires $\xi_i(\theta)$; on la décompose alors sur le chaos polynomial comme dans le cas de champs non normaux.

1.5 Autres méthodes stochastiques

1.5.1 Méthodes fiabilistes

Ces méthodes s'attachent à déterminer la probabilité de défaillance d'une structure dans des cas où une simulation directe de Monte Carlo serait très difficile à mettre en œuvre. Pour cela, on associe au problème :

- la donnée de n variables aléatoires X_i qui sont les paramètres incertains du problème, pour lesquels la densité conjointe de probabilité est $p(x_1, \dots, x_n)$;
- un ou plusieurs scénarios de défaillance caractérisés chacun par une fonction de performance $G(x_1, \dots, x_n)$, dont les valeurs positives constituent les réalisations du succès tandis que les valeurs négatives traduisent la défaillance.

Avec ces notations, la probabilité de défaillance P_f d'une structure se calcule alors comme :

$$P_f = \int_{G(x_1, \dots, x_n) \leq 0} p(x_1, \dots, x_n) dx_1 \dots dx_n \quad (1.33)$$

L'évaluation directe de cette probabilité de défaillance par une méthode de Monte Carlo se révèle alors impossible, car pour estimer correctement des probabilités P_f de l'ordre de 10^{-m} , il faudrait effectuer de l'ordre de 10^{m+2} , voire 10^{m+3} tirages, or on s'intéresse à des événements de très faible occurrence, pour lesquels on a généralement $m \geq 4$. De plus, cette technique nécessite également la connaissance de la densité conjointe de probabilité $p(x_1, \dots, x_n)$, qui n'est pas toujours simple à établir.

Classiquement, la probabilité de défaillance est estimée par l'intermédiaire du calcul de l'indice de fiabilité β de Hasofer-Lind [[Hasofer et Lind 1974](#)].

La première étape consiste à appliquer une transformation dite isoprobabiliste T permettant d'agir dans un espace de variables aléatoires $(U_1, \dots, U_n) = T(X_1, \dots, X_n)$ qui vérifient chacune une loi de probabilité normale centrée d'écart type unitaire ; l'état-limite est alors caractérisé par $H(u_1, \dots, u_n) = 0$ avec H tel que :

$$H(u_1, \dots, u_n) = G(T^{-1}(u_1, \dots, u_n)) \quad (1.34)$$

Plusieurs transformations sont possibles selon la quantité d'information disponible à propos des variables X_i ; en particulier, il n'est pas nécessaire de connaître la densité conjointe de probabilité, car certaines techniques approchées ne requièrent que la connaissance des lois de probabilité marginales et des fonctions de corrélations entre les différentes variables aléatoires.

La seconde étape consiste à calculer l'indice de fiabilité β , qui, dans l'espace construit, est défini géométriquement comme le minimum de la distance de l'origine de l'espace des variables normées à la fonction d'état-limite $H(u_1, \dots, u_n)$, ce qui revient à :

$$\beta = \min_{\substack{(u_1, \dots, u_n) \\ H(u_1, \dots, u_n)=0}} \sqrt{\sum_{i=1}^n u_i^2} \quad (1.35)$$

Le point solution de ce problème de minimisation est appelé point de défaillance le plus probable et noté P^* .

La probabilité de défaillance se calcule alors de plusieurs façons :

- la plus simple consiste à utiliser une approximation du premier ordre (*First Order Reliability Method* ou *FORM* en anglais) en remplaçant la fonction d'état-limite par son hyperplan tangent au point P^* :

$$P_f = \Phi(-\beta) = 1 - \Phi(\beta) \quad (1.36)$$

où Φ est la fonction de répartition de la loi normale centrée de variance unitaire ; la figure 1.2 illustre cette situation ;

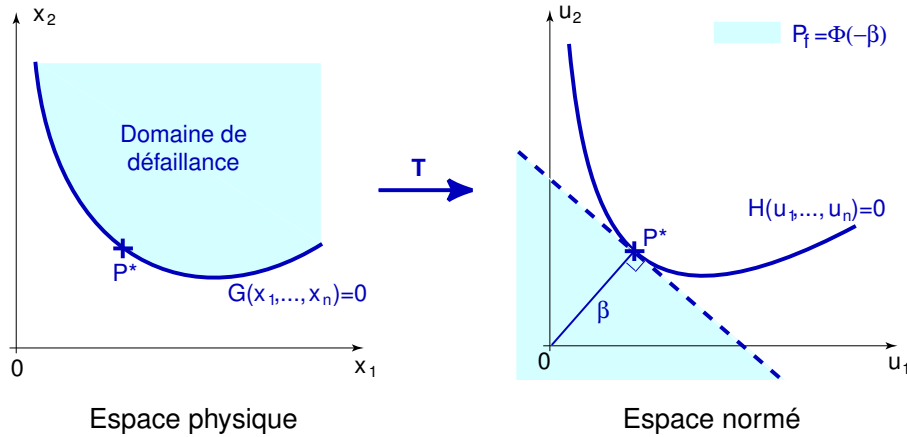


FIG. 1.2 – Principe des méthodes fiabilistes : exemple de l'approximation FORM

- une technique plus précise consiste à approcher la géométrie de la fonction d'état-limite au voisinage du point P^* par une courbe plus riche (*Second Order Reliability Method* ou *SORM* en anglais) :

$$P_f = \Phi(-\beta) \prod_{i=1}^{n-1} \frac{1}{\sqrt{1 + \beta \kappa_i}} \quad (1.37)$$

- où les κ_i sont les courbures principales de la fonction d'état-limite au point P^* ;
- une simulation de Monte Carlo améliorée selon les méthodes du paragraphe 1.3.2 peut être employée pour réaliser l'essentiel des tirages au voisinage du point P^* ; on peut aussi envisager d'utiliser une approximation à l'aide des surfaces de réponse, couplée avec un plan d'expérience [Lemaire 1997].

Pour plus de précisions sur les méthodes fiabilistes, on pourra consulter par exemple [Ditlevsen et Madsen 1996].

1.5.2 Méthodes non-paramétriques

Dans toutes les méthodes stochastiques que nous venons de passer en revue, le choix des paramètres incertains du modèle, à savoir le type de données (variables aléatoires ou champs stochastiques) et les lois de probabilité associées, se révèle crucial pour l'étude de la variabilité de la réponse. Dans ce cas, on suppose que le modèle est parfaitement connu, les incertitudes ne portant en effet que sur les paramètres de ce modèle, or pour des systèmes mécaniques suffisamment complexes, cette hypothèse peut être excessive ; certaines zones, comme les liaisons par exemple, mettent en jeu des phénomènes qui souvent ne sont pas pris en compte dans le détail, et sont donc représentées par des modèles simplifiés, pour lequel le choix des paramètres incertains se révèle un peu artificiel. **Il faut dans ces cas pouvoir prendre en compte des incertitudes de modélisation.**

Peu de travaux considèrent ce type d'incertitudes ; citons par exemple [Soize 1986] pour une première approche concernant le « flou structural », défini comme l'ensemble des sous-systèmes mécaniques secondaires reliés à la structure maîtresse, et qui ne sont pas modélisables classiquement, la finalité étant d'introduire une modélisation probabiliste globale de ce flou structural : ce concept n'a pas pour but premier d'étudier la variabilité de la réponse du système, mais vise à améliorer la prévision des *FRF* déterministes, initialement dans le domaine des « Moyennes Fréquences », puis dans la quantification des effets de l'amortissement dans le domaine des Basses Fréquences [Soize 1994]. Les lois de probabilité étant appliquées à des quantités globales telles que des impédances de frontière entre le flou structural et la structure maîtresse, cette dernière approche est qualifiée de non-paramétrique.

Plus récemment dans [Soize 2000], une autre méthode non-paramétrique a été proposée pour prendre en compte les incertitudes de modélisation, cette fois-ci avec pour but réel d'étudier la variabilité des *FRF*. Ici, la méthode est non-paramétrique car les paramètres incertains ne sont pas modélisés directement par des variables aléatoires ou des champs stochastiques : **les incertitudes sont prises en compte de façon globale en modélisant directement les matrices du modèle dynamique par des matrices aléatoires, formées à partir du principe du maximum d'entropie** décrit dans le paragraphe 1.2.3. Puisque les niveaux de dispersion peuvent varier fortement au sein d'une structure complexe, cette approche non-paramétrique a été couplée dans [Chebli 2002] avec la méthode de sous-structuration dynamique de Craig-Bampton [Craig et Bampton 1968] pour pouvoir l'appliquer sur chaque sous-domaine où le niveau d'incertitude peut être considéré comme homogène ; les résultats de cette méthode sont alors confrontés avec des mesures expérimentales sur deux plaques boulonnées.

1.6 Méthodes stochastiques et validation de modèles

Même si elles propagent ainsi les incertitudes sur les paramètres, et permettent une comparaison entre la variabilité de la réponse du modèle et celle effectivement constatée dans la réalité, **peu de méthodes s'attachent spécifiquement à la validation de modèles mécaniques probabilistes en tant que telle**. C'est donc bien au travers de cette comparaison que peut commencer à s'esquisser une démarche de validation.

On imagine ainsi réaliser à partir de cette comparaison une sorte de propagation inverse des incertitudes, comme ce qui est présenté dans [Hemez *et al.* 2004] où les auteurs entreprennent de « calibrer » les paramètres incertains du modèle à l'aide d'une fonction coût définissant une probabilité *a posteriori* que les paramètres du modèle soient corrects, eu égard aux mesures expérimentales effectuées ; cette probabilité *a posteriori* peut s'obtenir aisément à l'aide du théorème de Bayes (1.2) car elle est proportionnelle au produit de la probabilité *a priori* d'avoir les bons paramètres du modèle par celle d'avoir les valeurs expérimentales mesurées avec les paramètres choisis *a priori*. On a alors défini à travers cette fonction coût une quantification plus ou moins satisfaisante de la qualité du modèle stochastique vis-à-vis de la référence expérimentale, et donné le moyen d'obtenir, par minimisation de cette fonction coût, les « meilleurs » paramètres stochastiques.

Cette technique est à rapprocher de la pléthore des méthodes de calibration statistique [Beven et Binley 1992, Kennedy et O'Hagan 2001], mais aussi du domaine des méthodes dites inverses [Tarantola 1987, Menke 1984], qui entrent toutes pleinement dans le cadre stochastique, mais où **la mécanique est dans la grande majorité des cas quelque peu absente**. En effet, une comparaison entre la résolution de problèmes inverses par la théorie de l'estimation [Tarantola 1987] et la méthode de recalage à l'aide de l'erreur en relation de comportement modifiée [Ladevèze et Chouaki 1999] est menée dans [Deraemaeker *et al.* 2004] ; les auteurs montrent que, même si les formulations dans les deux cas sont similaires, les deux méthodes sont de nature profondément différente, puisque le choix des normes est d'origine statistique dans le cadre de la théorie de l'estimation, au lieu de présenter une justification mécanique comme dans le cas de l'erreur en relation de comportement, et c'est pourquoi de telles théories mathématiques peuvent difficilement constituer des estimateurs « mécaniques » de la qualité d'un modèle. Les auteurs proposent alors une extension du concept de l'erreur en relation de comportement modifiée dans le cas de données incertaines.

Plus récemment, l'idée précédente est encore approfondie pour envisager le cas d'un modèle stochastique, afin de pouvoir recalculer, à l'aide de données expérimentales, les paramètres incertains de ce modèle : la démarche se situe dans le prolongement des travaux déjà menés dans le cadre déterministe à l'aide du concept de l'erreur en relation de comportement modifiée ; pour cela, une extension au cadre stochastique, associée à une reconstruction des données expérimentales à l'aide du modèle stochastique, permet l'élaboration d'un indicateur d'erreur nul si le modèle stochastique est « exact », c'est-à-dire si les données expérimentales s'identifient aux valeurs prédites par le modèle. Pour plus de précisions, on pourra consulter [Ladevèze 2003b, Ladevèze *et al.* 2005, Caignot 2004].

État de l'art des méthodes non stochastiques

Les savants sont des gens qui, sur la route des choses inconnues, s'embourbent un peu plus loin que les autres.

Alphonse Karr

Sommaire

2.1	Motivation des méthodes non stochastiques	28
2.2	Théorie généralisée de l'information	29
2.2.1	Description	29
2.2.2	Exemple de la modélisation probabiliste	30
2.3	La théorie des intervalles	32
2.3.1	Principes	32
2.3.2	Résolution de systèmes matriciels	33
2.3.3	Applications en calcul des structures	34
2.4	Les ensembles flous	35
2.4.1	Principes et définitions	35
2.4.2	Applications	36
2.5	Ensembles aléatoires selon Dempster-Shafer	38
2.5.1	Définitions et propriétés	38
2.5.2	Applications	42
2.5.3	Conclusion	45
2.6	Les modèles convexes d'incertitude	46
2.6.1	Cadre mathématique associé	46
2.6.2	Applications dans le domaine de la fiabilité	49
2.6.3	Applications dans le domaine de la validation de modèles	49

Dans ce second chapitre, nous essayons de donner un état de l'art, non exhaustif, des différentes méthodes modélisant les incertitudes par des techniques n'employant pas la théorie des probabilités. Ces approches, qui soit étendent à un niveau supérieur, soit prennent le contre-pied total de ces méthodes stochastiques sont classiquement regroupées dans ce que l'on appelle la théorie généralisée de l'information [Klir 1991].

2.1 Motivation des méthodes non stochastiques

On ne peut nier, ainsi qu'on l'a montré dans le chapitre précédent, le rôle prédominant joué par les méthodes stochastiques dans la modélisation des incertitudes. Cependant, certains auteurs se sont peu à peu montrés prudents avec l'utilisation de telles approches qui peuvent être restrictives, et les concepts apportés par les mathématiques se multipliant, se sont mis à considérer d'autres voies. En effet, **associer une vision probabiliste à tout problème de quantification des incertitudes peut être exagéré dans certains cas** où le manque d'informations est tel que supposer une telle vision revient déjà à déformer le problème initial de modélisation des incertitudes.

Ainsi, dans [Worden *et al.* 2003], les auteurs ont mis en évidence que les méthodes stochastiques pouvaient être mises en défaut sur un simple exemple de non-linéarité forte, à savoir un système masse-ressort à un degré de liberté pour lequel la relation de comportement du ressort exprime l'effort comme une fonction cubique du déplacement (oscillateur de Duffing).

On peut montrer que, soumis à une excitation sinusoïdale, ce système peut avoir soit une réponse à faible amplitude, soit une réponse à forte amplitude, et ce pour une même valeur de la fréquence de l'excitation ; l'une ou l'autre réponse s'établit suivant les conditions initiales en position et vitesse de l'oscillateur. Ce comportement rend l'approximation de la réponse par un méta-modèle, dont le principe a été exposé dans le paragraphe 1.3.2, tout bonnement impossible, car la « surface » de réponse présente un caractère fractal, deux tirages de Monte Carlo de valeurs très proches pouvant donner deux réponses d'amplitudes très éloignées.

De même, l'utilisation des méthodes *FORM/SORM* décrites dans le paragraphe 1.5.1 se révèle totalement impossible à cause du phénomène de bifurcation présent dans l'oscillateur. En fait, le comportement très non-linéaire de ce dernier rend totalement caduc l'emploi d'une modélisation stochastique des incertitudes : il faut donc dans ce cas se tourner vers d'autres schématisations des incertitudes.

En fait, comme nous l'avons déjà évoqué dans l'introduction de ce mémoire, on peut considérer que l'incertitude à propos d'une grandeur physique se manifeste selon deux schémas différents, comme décrit dans [Klir 1994] :

- **un schéma conflictuel** où des valeurs alternatives et distinctes du paramètre peuvent se produire ; l'illustration classique de ce cas de figure est le lancer de dés, où l'occurrence d'une valeur exclut toutes les autres, et la théorie des probabilités est tout à fait adaptée à l'étude de tels phénomènes ;

- **un schéma de non-spécificité** relatif à un ensemble de valeurs alternatives toutes possibles de façon égale ; une illustration en est la réalisation d'une mesure à l'aide d'un appareil, à laquelle est associé un intervalle de confiance ; même si l'on arrivait à obtenir par d'autres moyens une valeur précise de la quantité mesurée, ceci ne rentrerait pas en conflit avec l'annonce initiale que la quantité recherchée est quelque part dans un certain intervalle, une sorte de région minimale d'ignorance pour ainsi dire ; dans ce schéma, la notion d'intervalle devient primordiale, et ne peut être représentée par la théorie classique des probabilités.

Bien entendu, ces deux schémas peuvent avoir lieu simultanément, et dès lors que le second cas de figure se produit, une modélisation probabiliste entraîne nécessairement une déformation du problème traité, qui peut être inacceptable dans certains cas, d'où le recours à des méthodes non stochastiques, dont nous allons maintenant dresser un panorama. Dans ce qui suit, bien que les formulations initiales concernent des supports ensemblistes, nous allons à chaque fois étudier les restrictions à des quantités réelles, qui sont les objets effectivement manipulés dans les applications réelles.

2.2 Théorie généralisée de l'information

2.2.1 Description

Depuis une quarantaine d'années s'est développée l'idée que le concept d'incertitude avait pour pendant la notion d'information : en effet, il est tout à fait compréhensible que si l'on reçoit de l'information à propos d'un système, une part de l'incertitude initiale concernant ce système est levée ; ainsi, **on peut considérer que le concept de l'information représente une réduction potentielle de l'incertitude.**

Cette idée a pris naissance à partir des travaux de Shannon [Shannon 1948, Shannon et Weaver 1964] qui définissent l'information comme une mesure statistique d'une loi de probabilité : ainsi en va-t-il de la notion d'entropie d'une variable aléatoire qui s'exprime selon la relation (1.12). Parallèlement, à partir des années soixante, ont commencé à se développer toutes sortes de formalismes différents, soit liés à la théorie des probabilités, soit orientés dans une toute autre direction ; devant cette prolifération de méthodologies et de terminologies, la notion d'une théorie généralisée de l'information (Generalized Information Theory ou GIT en anglais) a été introduite dans [Klir 1991], de façon à réunifier les concepts et à comparer les différentes visions proposées.

Dans ce formalisme sont définies les opérations suivantes :

- la fonction complément $c : [0; 1] \rightarrow [0; 1]$ ayant les caractéristiques suivantes :

$$c(0) = 1 \quad \text{et} \quad c(1) = 0 \quad (2.1a)$$

$$x \leq y \Rightarrow c(x) \geq c(y) \quad (2.1b)$$

- les normes $\sqcap : [0; 1]^2 \rightarrow [0; 1]$ et conormes $\sqcup : [0; 1]^2 \rightarrow [0; 1]$ qui sont des fonctions associatives et commutatives vérifiant les propriétés suivantes :

$$\forall x \leq y \quad \forall z \leq w, \quad x \sqcap z \leq y \sqcap w \quad \text{et} \quad x \sqcup z \leq y \sqcup w \quad (2.2)$$

et dont les éléments neutres sont respectivement un et zéro :

$$\forall x, \quad 1 \sqcap x = x \sqcap 1 = x \quad \text{et} \quad 0 \sqcup x = x \sqcup 0 = x \quad (2.3)$$

Il s'agit alors de voir dans ce cadre comment considérer l'apport d'information nouvelle. Supposons qu'un certain niveau d'information, noté A , est connu, ou supposé, à propos d'un système donné, et qu'un degré d'information, noté B , supplémentaire et indépendant du précédent, peut être obtenu. Deux façons différentes de combiner ces deux états d'information coexistent :

- **soit les deux informations A et B correspondent à des solutions distinctes**, c'est-à-dire qu'elles correspondent à des conditions qui doivent être satisfaites simultanément ; par exemple, en termes de la théorie classique des ensembles, cela revient à considérer leur intersection :

$$C = A \cap B \quad (2.4)$$

dans ce cas, l'association des deux états d'information A et B conduit à une réduction de l'incertitude sur le système étudié ;

- **soit les deux informations A et B ne sont que des indications de possibilité**, c'est-à-dire qu'elles indiquent simplement que l'on ne peut avoir simultanément les informations « contraires » à A et B , ce qui s'écrit en termes de la théorie classique des ensembles comme :

$$C = (A^C \cap B^C)^C = A \cup B \quad (2.5)$$

ce qui revient à considérer l'union des deux états d'information A et B , et donc à augmenter l'incertitude sur le système étudié.

Ces deux opérations d'union et d'intersection dans le cadre de la théorie classique des ensembles se traduit de façon générale à travers les concepts de norme et conorme respectivement, dans toutes les méthodes qui s'inscrivent dans la théorie généralisée de l'information.

2.2.2 Prise en compte de l'information nouvelle : exemple de la modélisation probabiliste

En ce qui concerne la prise en compte de l'information, une question qui se pose également concerne la transmission de l'incertitude sur une variable d'un certain espace X à travers un modèle représenté par une certaine fonction f , déterministe, à valeurs dans un espace Y . On l'illustre ici dans le cas d'une modélisation de l'incertitude à l'aide des outils probabilistes.

Supposons que l'on connaisse la densité de probabilité $p_X(x)$, avec x élément quelconque de l'espace X ; la connaissance de la fonction f nous permet d'étendre cette connaissance de l'espace X sur l'espace Y , en définissant la densité de probabilité associée p_Y comme :

$$p_Y(y) = \int_{f(x)=y} p_X(x) dx \quad (2.6)$$

Ceci est en fait un cas très particulier de l'application du théorème de Bayes (1.2) écrit à l'aide des densités de probabilité marginales p_X et p_Y , ainsi que de la densité de probabilité conjointe p_{XY} sur $X \times Y$:

$$p_Y(y) = \int_X p_{XY}(x, y) dx = \int_X p_Y(y | x) p_X(x) dx \quad (2.7)$$

où le fait d'apporter la relation déterministe $y = f(x)$ revient à avoir comme densité de probabilité conditionnelle :

$$p_Y(y | x) = \begin{cases} 1 & \text{si } y = f(x) \\ 0 & \text{sinon} \end{cases} \quad (2.8)$$

De manière plus générale, **le théorème de Bayes permet, dans une modélisation probabiliste de l'incertitude, de mettre à jour une connaissance *a priori* avec de l'information nouvelle.** En effet, la relation (1.2) relie deux types de probabilité concernant l'événement E_1 :

$$P(E_1 | E_2) = \frac{P(E_2 | E_1) P(E_1)}{P(E_2)} \quad (2.9)$$

On a d'une part $P(E_1)$ qui représente la probabilité *a priori* d'avoir l'événement E_1 , indépendamment de toute source d'information supplémentaire, et d'autre part une probabilité *a posteriori* $P(E_1 | E_2)$ concernant E_1 après l'apport d'information représenté par l'événement E_2 .

Si l'on enlève le dénominateur de la relation précédente, on exprime une proportionnalité entre les deux probabilités concernant l'événement E_1 :

$$P(E_1 | E_2) \propto P(E_2 | E_1) P(E_1) \quad (2.10)$$

Si l'on observe effectivement l'événement $E_2 = e_2$, l'équation précédente devient :

$$P(E_1, E_2 = e_2) \propto P(E_2 = e_2 | E_1) P(E_1) \quad (2.11)$$

Cependant, la quantité $P(E_2 = e_2 | E_1)$ ne peut plus s'interpréter comme une probabilité, mais comme la vraisemblance (*likelihood* en anglais) d'avoir E_1 vu que l'on a effectivement observé $E_2 = e_2$; on la note $L(E_1, e_2)$, ce qui aboutit à une « nouvelle » écriture du théorème de Bayes :

$$P(E_1, E_2 = e_2) \propto L(E_1, e_2) P(E_1) \quad (2.12)$$

Comme illustration, prenons l'exemple de la probabilité de défaillance chez une famille de pièces mécaniques similaires. On pose que $P(E_1 = e_1)$ est la probabilité de défaillance *a priori* d'un élément de cette famille, provenant de l'expérience ou des hypothèses concernant cette dernière. Imaginons que l'on réalise des essais afin d'améliorer la connaissance de la bonne tenue de ces pièces ; ce sont les résultats $E_2 = e_2$ de ces essais qui forment la vraisemblance $L(E_1 = e_1; e_2)$, qui nous permet de déterminer une nouvelle probabilité de défaillance $P(E_1 = e_1, E_2 = e_2)$, *a posteriori*, intégrant à la fois la connaissance initiale et celle apportée par les essais.

Bien sûr, tout ce qui précède peut être réécrit en termes de densités de probabilité :

$$p_X(x|y) = \frac{p_Y(y|x)p_X(x)}{\int_X p_Y(y|x)p_X(x)dx} \quad (2.13)$$

puisque, comme on l'a vu plus haut, $p_Y(y) = \int_X p_Y(y|x)p_X(x)dx$.

La vraisemblance devient alors une fonction de x , vu que l'on a effectivement mesuré $y = y_0$:

$$p_X(x, y_0) \propto L(x, y_0)P_X(x) \quad (2.14)$$

Nous aurons l'occasion d'étudier plus loin d'autres applications du théorème de Bayes.

2.3 La théorie des intervalles

2.3.1 Principes

Utiliser un intervalle englobant tous les possibles permet de répondre à la question de la non-spécificité, en décrivant des incertitudes qui ne correspondent pas à de la variabilité pure ; la quantité $x \in \mathbb{R}$ associée est alors bornée par deux valeurs, respectivement supérieure I^+ et inférieure I^- . L'intervalle $[I]$ d'incertitude peut être représenté par sa fonction caractéristique χ_I telle que :

$$\begin{aligned} \chi_I : \mathbb{R} &\rightarrow \{0, 1\} \\ x &\mapsto \chi_I(x) = \begin{cases} 1 & \text{si } I^- \leq x \leq I^+ \\ 0 & \text{sinon} \end{cases} \end{aligned} \quad (2.15)$$

Dans [Moore 1966 - 1979] a été définie une arithmétique qui permet de propager à travers tout modèle les intervalles d'incertitude : de façon générale, on peut définir une opération donnée $*$ $\in \{+, -, \times, \div, \min, \max\}$ entre deux intervalles $[I]$ et $[J]$ comme suit :

$$[I] * [J] = \{x * y \mid x \in [I], y \in [J]\} \quad (2.16)$$

Ainsi, nous avons pour la somme de deux intervalles :

$$[I] + [J] = [I^-; I^+] + [J^-; J^+] = \{x + y \mid x \in I, y \in J\} = [I^- + J^-; I^+ + J^+] \quad (2.17)$$

mais il faut souligner que, en général, on a $[I] * [J] \neq [I^- * J^-; I^+ * J^+]$; en effet, par exemple, la multiplication de deux intervalles nous donne dans le cas général :

$$\begin{aligned} [I] \times [J] &= [I^-; I^+] \times [J^-; J^+] = \{xy \mid x \in [I], y \in [J]\} \\ &= [\min(I^-J^-, I^-J^+, I^+J^-, I^+J^+); \max(I^-J^-, I^-J^+, I^+J^-, I^+J^+)] \end{aligned} \quad (2.18)$$

Ainsi, toute fonction à n variables $f : (x_1, \dots, x_n) \mapsto \mathbb{R}$ peut être appliquée à n intervalles, et l'image est alors définie par :

$$f([I_1], \dots, [I_n]) = \{f(x_1, \dots, x_n) \mid x_i \in [I_i], i = 1, \dots, n\} \quad (2.19)$$

Dans le cas général, cet ensemble n'est pas un intervalle ; on peut néanmoins lui associer le plus petit intervalle contenant $f([I_1], \dots, [I_n])$, noté $\square f([I_1], \dots, [I_n])$.

Propagation des intervalles

La propagation d'intervalles montre alors un inconvénient majeur, parfois rédhibitoire : elle peut être très pessimiste. Si, par exemple, on choisit une fonction $f(x) = 2x - x$ appliquée à un intervalle $[I] = [I^-; I^+]$, on s'attend bien sûr à retrouver $[I]$ comme image. Toutefois, l'application de la règle de sommation des intervalles nous conduit à $f([I]) = [2I^- - I^+; 2I^+ - I^-] \supset [I]$: le fait d'avoir « oublié » la dépendance des occurrences de x nous conduit à une surestimation de l'intervalle image. La seule situation n'engendrant pas ce problème se produit quand chaque variable n'intervient qu'une seule fois dans la fonction : il est donc nécessaire d'utiliser des expressions factorisées au maximum, en raison de la propriété dite de « sous-distributivité » :

$$[I] \times ([J] + [K]) \subset [I] \times [J] + [I] \times [K] \quad (2.20)$$

Diverses améliorations de l'arithmétique des intervalles ont été proposées afin de pouvoir prendre facilement en compte les dépendances entre les différentes variables considérées : on peut citer par exemple l'arithmétique affine [Comba et Stolfi 1993]. Des informations régulièrement actualisées sur l'arithmétique des intervalles peuvent également être trouvées dans [IC].

Pour quantifier l'incertitude représentée par un intervalle $[I]$, on peut lui associer le diamètre de ce dernier, ou encore, en en prenant le logarithme de base deux :

$$S([I]) = \log_2(I^+ - I^-) \quad (2.21)$$

dont la définition est à rapprocher de la notion d'entropie d'une variable aléatoire, qui s'exprime selon la relation (1.12).

Cette définition permet à nouveau de constater le phénomène pessimiste de surestimation :

$$\begin{aligned} S([I] + [J]) &= \log_2((I^+ - I^-) + (J^+ - J^-)) \\ &\geq \log_2(I^+ - I^-) + \log_2(J^+ - J^-) = S([I]) + S([J]) \end{aligned} \quad (2.22)$$

à savoir que l'incertitude de la somme $[I] + [J]$ est plus grande que la somme des incertitudes de $[I]$ et de $[J]$. Ainsi, sans précaution, l'incertitude selon la théorie des intervalles croît quand on la propage dans un modèle, ce qui peut rendre l'analyse inutile, quand les niveaux d'incertitude propagés ont tendance à être exagérés, et totalement invraisemblables vis-à-vis des phénomènes réels.

2.3.2 Résolution de systèmes matriciels

Les problèmes mécaniques mènent généralement, après une discrétisation appropriée, à la résolution de systèmes linéaires ; voyons comment il est possible d'utiliser une modélisation par intervalles en introduisant les objets suivants.

On définit un vecteur intervalle $[I]$ comme un vecteur dont les composantes sont des intervalles :

$$[I] = ([I_1] \cdots [I_n])^T \quad (2.23)$$

On peut interpréter « géométriquement » cette définition en remarquant qu'un tel vecteur intervalle est en fait un parallélépipède à n dimensions, dont les côtés sont colinéaires aux axes de la base canonique de $\mathbb{R}^n$; c'est pourquoi les vecteurs intervalles sont souvent qualifiés de « pavés ».

De la même façon, on peut introduire la notion de matrice intervalle $[\mathbf{I}]$, qui est une matrice dont les composantes sont des intervalles. La plupart des définitions associées aux intervalles (centre, rayon, ...) peuvent être appliquées aux pavés et aux matrices intervalles.

Considérons la forme générale des systèmes linéaires que l'on peut être amené à résoudre en mécanique :

$$\mathbf{K}\underline{U} = \underline{F} \quad (2.24)$$

où, typiquement, $\mathbf{K}$ est associé à la raideur de la structure, $\underline{F}$ aux efforts exercés sur cette dernière, et $\underline{U}$ aux déplacements correspondants.

Suivant ce qui est connu et inconnu dans l'équation précédente, de nombreux problèmes sont envisageables, comme décrit dans [Chen et Ward 1997]. Le cas le plus courant est celui où l'on cherche les déplacements de la structure, avec la raideur et les efforts définis respectivement comme une matrice intervalle $[\mathbf{K}]$ et un pavé $[\underline{F}]$. La solution théorique de ce problème, définie par :

$$\Sigma_{\exists\exists}([\mathbf{K}], [\underline{F}]) = \{\underline{x} \in \mathbb{R}^n \mid \exists \mathbf{K} \in [\mathbf{K}], \exists \underline{F} \in [\underline{F}] \quad \mathbf{K}\underline{x} = \underline{F}\} \quad (2.25)$$

est donnée de façon exacte dans [Oettli et Prager 1964]. Comme généralement cette solution n'est pas un pavé et est difficile à calculer telle quelle, de nombreuses méthodes numériques existent pour déterminer le plus petit intervalle $\square\Sigma_{\exists\exists}([\mathbf{K}], [\underline{F}])$ la contenant [Ning et Kearfott 1997].

2.3.3 Applications en calcul des structures

Avec le développement des outils mathématiques et numériques évoqués juste avant, il est devenu possible récemment de résoudre des problèmes mécaniques dont les incertitudes sont modélisées par des intervalles.

En particulier, dans [Dessombz 2000, Dessombz et al. 2001], a été introduite une formulation par intervalles appliquée à un problème de dynamique des structures discrétisé par la méthode des Éléments Finis ; pour cela, on fait l'hypothèse que les paramètres varient peu et sont spatialement peu corrélés, ce qui revient à considérer que l'on peut écrire les variations de la matrice de raideur comme suit :

$$\mathbf{K} = \mathbf{K}_0 + \sum_{i=1}^n [\varepsilon_i] \mathbf{K}_i \quad (2.26)$$

où les quantités $[\varepsilon_i]$ sont des intervalles indépendants, généralement pris égaux à $[-1; 1]$; on peut effectuer la même décomposition pour le vecteur $\underline{F}$ des efforts. Cette écriture n'est pas sans rappeler celle associée à la méthode des Éléments Finis Stochastiques, à ceci près que les matrices $\mathbf{K}_i$ étaient multipliées par des variables aléatoires indépendantes.

Les auteurs, ayant constaté que des techniques approchées, telles qu'une méthode de perturbation ou une projection sur une base de polynômes orthogonaux de type Hermite, ne permettent pas d'obtenir une enveloppe de la solution, introduisent alors une méthode de résolution spécifique : celle-ci s'inspire de la méthode d'inclusion de Rump [Rump 1983] qui est un algorithme itératif basé sur le théorème du point fixe.

L'utilisation de cette méthode rend alors possible la détermination des *FRF* enveloppes, grâce à un mode de calcul spécifique qui permet de ne pas trop surestimer les intervalles résultants. De plus, le coût numérique est faible, comparé à une simulation directe de Monte Carlo. Des applications dans [Dessombz 2000] permettent de constater la prise en compte de phénomènes tels que l'apparition de résonances dues à un désaccordage d'une partie d'un système à symétrie cyclique.

2.4 Les ensembles flous

2.4.1 Principes et définitions

C'est en partie pour résoudre le problème d'amplification des incertitudes modélisées par des intervalles qu'a été introduite une généralisation du concept classique d'ensemble avec la définition des ensembles flous dans [Zadeh 1965]. Pour cela, **la notion de fonction caractéristique (2.15) d'un ensemble est remplacée par une fonction, pouvant prendre n'importe quelle valeur dans $[0; 1]$, appelée fonction de participation** (*membership function* en anglais) associée à l'ensemble flou $\tilde{A}$:

$$\mu_{\tilde{A}} : \Omega \rightarrow [0; 1] \quad (2.27)$$

On remarquera que, avec cette définition, la fonction caractéristique χ_A est une fonction de participation particulière : comme cette dernière est telle que $\chi_A(x) = 1$ si et seulement si $x \in A$, et $\chi_A(x) = 0$ si et seulement si $x \notin A$, il est logique d'interpréter la valeur en x de la fonction de participation $\mu_{\tilde{A}}(x)$ comme le « degré d'appartenance », compris entre zéro et un, de x à l'ensemble flou $\tilde{A}$.

Cette définition permet alors d'établir une logique floue (*fuzzy logic* en anglais), associée à des « degrés » de vérité qui, assignés à des propositions, vont de zéro (faux) à un (vrai) avec toutes les gradations possibles. Ce type de logique autorise la prise en compte des jugements humains, ou des concepts « linguistiques », qui doivent être intégrés dans des raisonnements mathématiques. Par exemple, supposons que nous sommes intéressés par l'ensemble flou $\tilde{A}$, dont la caractéristique est de contenir tous les entiers « moyens », c'est-à-dire « ni trop petits, ni trop grands », parmi ceux compris entre zéro et dix. Nul doute que l'on mettrait cinq sans hésiter dans cet ensemble, mais qu'en serait-il de sept par exemple ? L'incertitude concernant ce que l'on entend par « moyen » est contenue dans la notion d'ensemble flou : ce dernier ne peut avoir de bornes « rigides » : ainsi l'entier sept peut avoir une certaine participation à l'ensemble flou $\tilde{A}$, mais aussi une participation *a priori* non nulle dans l'ensemble flou $\tilde{A}^C$ complémentaire de $\tilde{A}$, c'est-à-dire dans l'ensemble des entiers compris entre zéro et dix « qui ne sont pas moyens ».

L'un des principaux aspects proposés par Zadeh est l'extension aux ensembles flous des opérations d'union, intersection et complément issues de la théorie classique des ensembles : ces extensions sont en fait des cas particuliers des normes et conormes prévues par la *GIT*; d'autres choix sont en effet possibles et tout à fait envisageables :

$$\text{intersection : } \mu_{\tilde{A} \cap \tilde{B}}(x) = \mu_{\tilde{A}}(x) \sqcap \mu_{\tilde{B}}(x) = \min(\mu_{\tilde{A}}(x), \mu_{\tilde{B}}(x)) \quad (2.28a)$$

$$\text{union : } \mu_{\tilde{A} \cup \tilde{B}}(x) = \mu_{\tilde{A}}(x) \sqcup \mu_{\tilde{B}}(x) = \max(\mu_{\tilde{A}}(x), \mu_{\tilde{B}}(x)) \quad (2.28b)$$

$$\text{complément : } \mu_{\tilde{A}^c}(x) = c(\mu_{\tilde{A}}(x)) = 1 - \mu_{\tilde{A}}(x) \quad (2.28c)$$

Bien sûr, on peut interpréter ces opérations entre ensembles flous en termes de logique floue : la conjonction « et » entre deux propositions est une intersection entre deux ensembles tandis que la disjonction « ou » est une union, et la négation est associée à l'opération de complément.

Intervalles flous

Après avoir défini de manière générale les ensembles flous, il s'agit maintenant d'étudier leur restriction à l'ensemble des réels $\mathbb{R}$, afin d'envisager des applications plus pratiques. Ainsi, une quantité floue est un ensemble flou $\tilde{I}$ tel que $\mu_{\tilde{I}} : \mathbb{R} \rightarrow [0; 1]$. Cette quantité floue est un intervalle flou si elle est à la fois :

$$- \text{ normale, ie si } \max_{x \in \mathbb{R}} \mu_{\tilde{I}}(x) = 1 \quad (2.29)$$

$$- \text{ convexe, ie si } \forall x, y \in \mathbb{R} \quad \forall z \in [x; y] \quad \mu_{\tilde{I}}(z) \geq \min(\mu_{\tilde{I}}(x), \mu_{\tilde{I}}(y)) \quad (2.30)$$

Associé à cet intervalle flou, on définit le support comme l'intervalle (au sens classique du terme) $I = \{x \mid \mu_{\tilde{I}}(x) > 0\}$.

Si la fonction de participation $\mu_{\tilde{I}}(x)$ est telle qu'il existe un seul $x \in \mathbb{R}$ tel que $\mu_{\tilde{I}}(x) = 1$, alors l'intervalle flou est qualifié de nombre flou. Ces dénominations sont liées au fait qu'elles généralisent directement la notion d'intervalle et de nombre réel, ces deux cas correspondant à une fonction de participation de type fonction caractéristique (2.15).

Enfin, si $\tilde{I}_1, \dots, \tilde{I}_n$ sont des intervalles flous, alors le produit cartésien $\tilde{I}_1 \times \dots \times \tilde{I}_n$ est un ensemble flou dont la fonction de participation est :

$$\mu_{\tilde{I}_1 \times \dots \times \tilde{I}_n}(x_1, \dots, x_n) = \min(\mu_{\tilde{I}_1}(x_1), \dots, \mu_{\tilde{I}_n}(x_n)) \quad (2.31)$$

2.4.2 Applications

Propagation au travers d'un modèle

Maintenant que les ensembles flous ont été définis, nous allons préciser comment l'on peut transmettre la description de ces derniers au travers d'un modèle. Supposons que

celui-ci se résume à une fonction $f : \mathbb{R}^n \rightarrow \mathbb{R}$: l'ensemble flou $\tilde{I}_1 \times \dots \times \tilde{I}_n \subset \mathbb{R}^n$ se propage alors de la façon suivante :

$$\mu(z) = \begin{cases} \sup_{\substack{(x_1, \dots, x_n) \\ f(x_1, \dots, x_n) = z}} \min(\mu_{\tilde{I}_1}(x_1), \dots, \mu_{\tilde{I}_n}(x_n)) \\ 0 & \text{si } z \text{ n'a pas d'antécédent par } f \end{cases} \quad (2.32)$$

Ce calcul n'étant pas facilement réalisable sur le plan numérique pour un ensemble flou quelconque $\tilde{A}$, on utilise en pratique une représentation équivalente en termes d'ensembles classiques ${}^\alpha A$, appelés α -coupures (α -cuts en anglais), définis comme :

$${}^\alpha A = \{\underline{x} \mid \mu_{\tilde{A}}(\underline{x}) > \alpha\} \quad (2.33)$$

pour diverses valeurs de α comprises dans $[0; 1]$. Inversement, la donnée de ces α -coupures est suffisante pour reconstruire la fonction de participation de l'ensemble flou :

$$\mu_{\tilde{A}}(\underline{x}) = \max_{\alpha} \min(\alpha, \chi^{\alpha A}) \quad (2.34)$$

Avec cette description en termes d'intervalles, il est facile d'obtenir les α -coupures ${}^\alpha Z$ de l'ensemble flou $\tilde{Z}$ obtenu par propagation :

$${}^\alpha Z = \left[\min_{x_1 \in {}^\alpha I_1, \dots, x_n \in {}^\alpha I_n} f(x_1, \dots, x_n); \max_{x_1 \in {}^\alpha I_1, \dots, x_n \in {}^\alpha I_n} f(x_1, \dots, x_n) \right] \quad (2.35)$$

Dans le cas particulier où la fonction f est de classe C^1 et strictement monotone, la méthode des sommets (*Vertex Method* en anglais) permet de réduire considérablement l'effort de calcul en affirmant que les bornes des α -coupures ${}^\alpha Z$ sont directement les images de certaines bornes des α -coupures ${}^\alpha I_i$ [Dong et Shah 1987].

Quoique plus facile à mener, le calcul à l'aide des α -coupures a tendance à donner des évaluations surestimées, car il est basé sur la propagation d'intervalles déterministes dont on avait vu dans le paragraphe 2.3 la tendance au « pessimisme ». C'est pourquoi diverses méthodes ont été proposées pour résoudre cette difficulté [Möller et al. 2000, Hanss 2002].

Apport d'informations nouvelles

Supposons que la connaissance à propos d'un système se traduise par la donnée de deux ensembles flous $\tilde{A}$ et $\tilde{B}$; **la loi de combinaison floue suivante permet d'intégrer les deux informations précédentes sous la forme d'un nouvel ensemble flou $\tilde{C}$ dont la fonction de participation est donnée par :**

$$\mu_{\tilde{C}}(\underline{x}) = \min(\mu_{\tilde{A}}(\underline{x}), \mu_{\tilde{B}}(\underline{x})) \quad (2.36)$$

Notons que cette loi est valide également quand les deux ensembles flous ne sont pas définis sur le même espace.

Dans la situation courante où les deux ensembles flous sont définis sur $\mathbb{R}^n$, et donc *a fortiori* sur le même espace, **cette loi correspond à la définition de la fonction de participation d'un ensemble flou intersection :**

$$\mu_{\tilde{C}}(\underline{x}) = \mu_{\tilde{A} \cap \tilde{B}}(\underline{x}) \quad (2.37)$$

ce qui représente donc l'extension naturelle de la relation (2.4) concernant des ensembles standards.

Un autre aspect de l'apport d'information à l'aide des ensembles flous provient du fait que, comme on l'a déjà évoqué précédemment, ces derniers permettent, à travers la logique floue, de manipuler des jugements humains ou des concepts « linguistiques ». **Un domaine très étudié concerne la combinaison d'une connaissance « mathématique » avec un avis d'expert :** ceci peut se faire à l'aide du théorème de Bayes, dont on avait déjà évoqué l'intérêt dans le paragraphe 2.2.2, grâce auquel la connaissance *a priori* d'une loi de probabilité concernant le système étudié peut être mise à jour à l'aide de la fonction de vraisemblance, qui, dans ce cas, s'écrit sous la forme d'un ensemble flou correspondant au jugement des experts [Booker et Singpurwalla 2003].

Applications en calcul des structures

Les applications dans le domaine de la mécanique se sont développées récemment et couvrent désormais de nombreux aspects : de l'identification de modèles [Hanss *et al.* 2002] à la dynamique des structures [Lallemant *et al.* 1999, Massa *et al.* 2004].

Plus précisément, **on peut également souligner la création d'une méthode des Éléments Finis Flous** (*Fuzzy FEM*, ou *FFEM*, en anglais) dans [Rao et Sawyer 1995]. Pour cela, les paramètres incertains sont modélisés sous la forme d'intervalles flous à l'aide de fonctions de participation, et les matrices élémentaires sont obtenues pour différentes α -coupures de ces intervalles flous ; la réponse est ensuite calculée à l'aide de l'arithmétique des intervalles.

Pour lutter contre la tendance au « pessimisme » de cette dernière, on peut encore une fois utiliser des méthodes adaptées permettant de résoudre cette difficulté ; on citera par exemple [Moens et Vandepitte 2002] où les auteurs peuvent dériver des *FRF* « floues », à partir des paramètres incertains du modèle.

2.5 Ensembles aléatoires selon Dempster-Shafer

2.5.1 Définitions et propriétés

Ce sont les travaux de Dempster, dans les années soixante, sur la propagation des lois de probabilité par des fonctions à valeurs dans des ensembles de dimension supérieure à un, qui sont à l'origine de cette théorie.

En effet, considérons un ensemble Z sur lequel existe une distribution de probabilité $P(z) \quad \forall z \in Z$. Soit une fonction $G : Z \rightarrow 2^X$, où 2^X désigne l'ensemble des

sous-ensembles d'un ensemble X donné : $2^X = \{A \mid A \subset X\}$. La propagation de la distribution de probabilité de Z sur l'ensemble X se fait selon la relation :

$$P(Y) = \sum_{G(z) \subset Y} P(z) \quad \forall Y \subset X \quad (2.38)$$

qui est analogue à l'égalité (2.6) utilisée dans le cas d'une fonction à valeurs dans $\mathbb{R}$. Cependant, avec la définition précédente, $P(Y)$ n'est pas une mesure de probabilité sur X , car on peut montrer que l'on n'a pas forcément la relation sur les ensembles classiques $P(A \cup B) = P(A) + P(B) - P(A \cap B)$. Il faut, pour résoudre cette difficulté, introduire de nouveaux concepts.

Ensembles aléatoires

Si, encore une fois, on considère une fonction G entre un espace Z , sur lequel existe une distribution de probabilité, et un ensemble X de dimension finie, on peut associer à chaque élément z_i de Z un sous-ensemble $A_i = G(z_i)$ de X et attribuer à ce dernier une quantité $m(A_i)$ telle que :

$$m(A_i) = P(z_i) \quad \text{et} \quad \sum_{i=1}^{|X|} m(A_i) = 1 \quad (2.39)$$

L'ensemble des sous-ensembles A_i de X avec les fonctions $m(A_i)$ associées est qualifié d'ensemble aléatoire sur X (*random set* en anglais), noté $\mathcal{A}(A_i, m(A_i))$; les sous-ensembles A_i sont appelés éléments focaux (*focal elements* en anglais), auxquels sont attribuées les quantités $m(A_i)$ qualifiées d'affectations probabilistes de base (*basic probabilistic assignments* en anglais).

La quantité $m(A_i)$ est une sorte de pondération quantifiant le fait que toute l'information dont on dispose permet d'affirmer qu'un élément de X appartient à A_i exactement ; en particulier, on ne peut rien affirmer à propos des sous-ensembles de A_i : en effet, si c'était le cas, c'est-à-dire si l'on était capable de dire qu'un élément donné de X était susceptible d'appartenir à un sous-ensemble B de A_i , il faudrait attribuer une affectation probabiliste $m(B)$ à ce sous-ensemble. Ceci revient à dire que chaque élément focal A_i doit être traité comme un objet en tant que tel, sans chercher pour le moment à tirer des conclusions sur sa nature d'ensemble : en quelque sorte, $m(A_i)$ désigne la probabilité d'occurrence de l'objet A_i .

Quand on associe à X une structure $\mathcal{A}(A_i, m(A_i))$ d'ensemble aléatoire, **il n'est pas possible de déterminer la probabilité $P(B)$ qu'un élément de X appartienne à un sous-ensemble B donné de X ;** par contre, il est possible de l'encadrer par deux bornes [Dempster 1968] :

$$\sum_{A_i \subset B} m(A_i) \leq P(B) \leq \sum_{A_i \cap B \neq \emptyset} m(A_i) \quad (2.40)$$

En fait, la borne inférieure est constituée de la somme des affectations probabilistes des éléments focaux « en faveur » de l'occurrence de l'événement B , tandis que la borne supérieure est concernée par le complément de la somme des affectations probabilistes de tous les éléments focaux qui vont « à l'encontre » de l'occurrence de cet événement.

Croyance et plausibilité

Dans [Shafer 1976], ces deux bornes sont qualifiées respectivement de croyance et plausibilité (*Belief* et *Plausibility* en anglais), et notées $\text{Bel}(B)$ et $\text{Pl}(B)$:

$$\text{Bel}(B) = \sum_{A_i \subset B} m(A_i) = 1 - \text{Pl}(B^C) \quad (2.41a)$$

$$\text{Pl}(B) = \sum_{A_i \cap B \neq \emptyset} m(A_i) = 1 - \text{Bel}(B^C) \quad (2.41b)$$

Ces fonctions ont les propriétés suivantes, qualifiées respectivement de super- et sous-additivité :

$$\text{Bel}(B \cup C) \geq \text{Bel}(B) + \text{Bel}(C) - \text{Bel}(B \cap C) \quad (2.42a)$$

$$\text{Pl}(B \cap C) \leq \text{Bel}(B) + \text{Bel}(C) - \text{Bel}(B \cup C) \quad (2.42b)$$

Avec l'introduction de ces fonctions, nous entrons dans l'étude des ensembles aléatoires selon le formalisme dit de Dempster-Shafer.

Dans le cas particulier où chaque élément focal n'est en fait réduit qu'à un singleton :

$$\forall A_i \quad \exists ! x_i \mid A_i = \{x_i\} \quad (2.43)$$

l'écart entre la croyance et la plausibilité s'annule, et ces deux mesures nous permettent de retrouver la mesure classique de probabilité :

$$\forall B \subset X \quad \text{Bel}(B) = \text{Pl}(B) = P(B) \quad (2.44)$$

Ensembles aléatoires à plusieurs dimensions

Supposons que l'ensemble X sur lequel est défini un ensemble aléatoire $\mathcal{A}$ est de dimension supérieure à un. De la même façon que l'on peut obtenir des densités de probabilité marginales à partir de la densité de probabilité conjointe, on peut déterminer des éléments focaux marginaux, qui sont les projections des éléments focaux A_i de l'ensemble aléatoire $\mathcal{A}$ sur les « axes » x_j de X . On associe à chacune de ces projections, notées A_{i,X_j} , l'affectation probabiliste de base de l'élément focal A_i dont elles sont issues.

Plus précisément, pour que les familles d'éléments focaux marginaux constituent chacune un ensemble aléatoire sur la projection de X suivant l'une de ses dimensions, il faut regrouper toutes les projections qui donnent un même élément focal, de la façon suivante :

$$m(A_{i,X_j}) = \sum_{\substack{k \\ A_{k,X_j} = A_{i,X_j}}} m(A_k) \quad (2.45)$$

Ainsi, on obtient pour chaque dimension de X un ensemble aléatoire $\mathcal{I}_{X_j}$ défini par les éléments focaux marginaux A_{i,X_j} auxquels sont associées les affectations $m(A_{i,X_j})$.

Par contre, il n'est pas toujours possible de faire l'opération inverse, de la même façon que la donnée des densités de probabilité marginales ne peut permettre la reconstruction de la densité de probabilité conjointe si l'on ne dispose pas de plus d'informations.

Intervalle aléatoires

Nous allons étudier maintenant la restriction d'un ensemble aléatoire à l'ensemble des réels : cette restriction, appelée intervalle aléatoire (*random interval* en anglais), est composée d'éléments focaux I_i qui sont des intervalles déterministes de $\mathbb{R}$. Dans ce cas, de nouveaux outils peuvent être introduits qui permettent une représentation plus aisée de ces objets.

Une boîte de probabilité (*probability box* en anglais, ou encore *p-box*) [Williamson et Downs 1990] est un objet $\mathcal{B} = \{\underline{B}, \overline{B}\}$ où $\underline{B}, \overline{B} : \mathbb{R} \rightarrow [0; 1]$ sont des fonctions monotones telles que $\underline{B} \leq \overline{B}$ et :

$$\lim_{x \rightarrow -\infty} \underline{B}(x) = \lim_{x \rightarrow -\infty} \overline{B}(x) = 0 \quad (2.46a)$$

$$\lim_{x \rightarrow \infty} \underline{B}(x) = \lim_{x \rightarrow \infty} \overline{B}(x) = 1 \quad (2.46b)$$

On peut interpréter ces deux fonctions $\underline{B}$ et $\overline{B}$ comme les bornes inférieure et supérieure de fonctions de répartition, qui définissent pour ainsi dire une classe de mesures de probabilité : en effet, la donnée de ces deux fonctions à travers $\mathcal{B}$ permet d'identifier l'ensemble de toutes les fonctions $\{F \mid \underline{B} \leq F \leq \overline{B}\}$ telles que F est la fonction de répartition associée à une certaine mesure de probabilité sur $\mathbb{R}$.

Ainsi, si $\mathcal{I}$ est un intervalle aléatoire, alors :

$$\mathcal{B}(\mathcal{I}) = \{\text{BEL}, \text{PL}\} \quad (2.47)$$

est une boîte de probabilité, où $\text{BEL}, \text{PL} : \mathbb{R} \rightarrow [0; 1]$ sont les fonctions de répartition de croyance et de plausibilité respectivement, définies par :

$$\text{BEL}(x) = \text{Bel}([-\infty; x]) \quad \text{et} \quad \text{PL}(x) = \text{Pl}([-\infty; x]) \quad (2.48)$$

Un autre outil de description d'un intervalle aléatoire $\mathcal{I}$ est sa trace $r_{\mathcal{I}}$ définie par :

$$r_{\mathcal{I}}(x) = \text{Pl}(\{x\}) \quad (2.49)$$

Cette fonction est en fait, vis-à-vis des mesures de croyance et de plausibilité, ce que la densité de probabilité est vis-à-vis de la mesure de probabilité. Dans le cas d'une boîte de probabilité dérivée de (2.47), on a simplement :

$$r_{\mathcal{I}} = \overline{B} - \underline{B} \quad (2.50)$$

Toutes ces notions sont bien sûr généralisables pour une dimension supérieure à un, avec $\mathbb{R}^n$ plutôt que $\mathbb{R}$, toutes les fonctions précédemment définies faisant alors intervenir plusieurs variables.

Quantification de l'incertitude

L'incertitude représentée par un ensemble aléatoire est quelque peu difficile à quantifier, vu que ce dernier intègre à la fois les deux schémas de représentation de l'incertitude, à savoir le conflit et la non-spécificité.

Dans cette perspective, un premier bon « candidat » revient à généraliser la relation (2.21) caractérisant la non-spécificité d'un intervalle déterministe :

$$U_N(\mathcal{I}) = \sum_{I_i \in \mathcal{I}} m(I_i) \log_2(I_i^+ - I_i^-) \quad (2.51)$$

où I_i^+ et I_i^- désignent les bornes supérieure et inférieure de l'intervalle déterministe I_i .

Bien que cette mesure permette à la fois de tenir compte du diamètre des éléments focaux, et du « poids » probabiliste qui leur est associé, elle peut ne pas refléter toutes les caractéristiques associées au schéma de conflit, qui est bien décrit par la notion d'entropie (1.12), d'où par exemple l'expression suivante :

$$U_C(\mathcal{I}) = - \sum_{I_i \in \mathcal{I}} m(I_i) \log_2 \left(\sum_{k=1}^n m(I_k) \frac{I_{ik}^+ - I_{ik}^-}{I_i^+ - I_i^-} \right) \quad (2.52)$$

avec $I_{ik} = I_i \cap I_k$; cette quantité représente une mesure d'entropie basée sur les éléments focaux.

Malgré l'utilité de ces mesures dans des cas bien précis, une expression plus simple permet à la fois de prendre en compte le schéma conflictuel et la non-spécificité : il s'agit de la plus grande entropie de toutes celles associées à des mesures de probabilité cohérentes avec la donnée de l'intervalle aléatoire, ce qui s'écrit :

$$U_{NS}(\mathcal{I}) = \max_{\substack{F \\ B \leq F \leq \overline{B}}} - \int_{\mathbb{R}} \frac{dF}{dx} \log_2 \left(\frac{dF}{dx} \right) dx \quad (2.53)$$

où les F sont des fonctions de répartition comprises dans la boîte de probabilité associée à $\mathcal{I}$.

2.5.2 Applications

Propagation au travers d'un modèle

Il s'agit de voir ici comment la connaissance d'un ensemble aléatoire sur $\mathbb{R}^n$ peut être propagée à l'aide d'une fonction $f : X \rightarrow Z$ dans un autre espace Z . Supposons que $\mathcal{A}$ est un ensemble aléatoire défini par ses éléments focaux A_i et les affectations associées $m(A_i)$. Conformément à la définition des affectations probabilistes, ces dernières sont liées aux éléments focaux :

$$R_i = f(A_i) = \{z = f(\underline{x}) \mid \underline{x} \in A_i\} \quad (2.54)$$

images sur Z des éléments focaux de $\mathcal{A}$, de la façon suivante :

$$m(R_i) = \sum_{f(A_k) = R_i} m(A_k) \quad (2.55)$$

Cette relation permet de tenir compte du cas où plusieurs éléments focaux de $\mathcal{A}$ ont la même image R_i ; au bilan, on a défini sur Z un ensemble aléatoire $\mathcal{R}$ caractérisé par les R_i et $m(R_i)$.

Cette propagation se réalise aussi facilement dans le cas de fonctions $G : X \rightarrow 2^Z$; dans cette situation, chaque élément focal image R_i est l'union de tous les sous-ensembles de Z images par G des différents points de l'élément focal A_i :

$$R_i = \bigcup_{x \in A_i} G(x) \quad (2.56)$$

Apport d'informations nouvelles

Supposons que l'on connaisse deux ensembles aléatoires $\mathcal{A}$ et $\mathcal{B}$ sur X ; la question à résoudre consiste à savoir comment associer ces deux connaissances indépendantes de X , afin de pouvoir déterminer un ensemble aléatoire $\mathcal{C}$ résultant.

Pour cela, Dempster a proposé la loi de combinaison suivante [Dempster 1968], dans laquelle les éléments focaux C_i associés à $\mathcal{C}$ sont issus des intersections des éléments focaux de $\mathcal{A}$ et $\mathcal{B}$:

$$m(C_k) = \frac{1}{K} \sum_{\substack{i,j \\ A_i \cap B_j = C_k}} m(A_i)m(B_j) \quad \text{avec} \quad K = 1 - \sum_{\substack{i,j \\ A_i \cap B_j = \emptyset}} m(A_i)m(B_j) \quad (2.57)$$

Les affectations probabilistes de $\mathcal{C}$ sont en fait les produits des affectations probabilistes associées aux éléments focaux concernés par l'intersection, mais normalisés pour traduire le fait qu'une partie de l'information initiale peut se « perdre » dans des intersections vides. Notons également que l'idée de l'intersection rappelle la façon dont on prendrait en compte simultanément les deux informations dans la théorie classique des ensembles, selon la relation (2.4).

Cette loi est intéressante dans la mesure où si le premier ensemble aléatoire se résume en fait à une description probabiliste, et si le second ensemble est purement déterministe, on retrouve exactement le théorème de Bayes ; on rappelle que ce dernier permet de « réviser » une probabilité *a priori* grâce à l'apport d'une information déterministe. Vue sous cet angle, **la loi de combinaison de Dempster est en fait une généralisation du théorème de Bayes dans le cas de l'apport d'une information nouvelle qui n'est pas déterministe.**

Ensembles aléatoires et ensembles flous

Un grand nombre des applications réalisées à l'aide des ensembles aléatoires se rapproche de celles obtenues à l'aide des ensembles flous. En effet, si un ensemble aléatoire $\mathcal{A}$ défini sur X est dit consonant, c'est-à-dire si ses éléments focaux vérifient :

$$A_i \subset A_{i+1} \quad (2.58)$$

alors pour chaque singleton $\underline{x}$ de X , la plausibilité $\text{Pl}(\{\underline{x}\})$ est équivalente à la fonction de participation d'un ensemble flou $\tilde{A}$, dont les α -coupures associées sont les éléments focaux de $\mathcal{A}$:

$$A_i = {}^{\alpha_i}A = \{x \mid \mu(x) \geq \alpha_i\} \quad \text{avec} \quad m(A_i) = \alpha_i - \alpha_{i+1} \quad (2.59)$$

où $1 = \alpha_1 > \alpha_2 > \dots > \alpha_{n-1} > \alpha_n = 0$.

Les ensembles aléatoires sont donc des structures plus générales que les ensembles flous.

Exemple académique

On veut juste montrer sur un exemple très simple, présenté dans [Oberkampff *et al.* 2001], comment utiliser les ensembles aléatoires. Supposons que la réponse y d'un système est reliée à la donnée de deux entrées a et b selon la relation $y = a^b$, et qu'il y a défaillance de ce système lorsque $y > 12$.

La connaissance des deux entrées est traduite en termes d'intervalles aléatoires $\mathcal{A}$ et $\mathcal{B}$ d'éléments focaux :

$$\begin{array}{llll} A_1 = [1; 1, 25] & A_2 = [1; 1, 5] & A_3 = [1; 1, 75] & A_4 = [1; 2] \\ B_1 = [2, 5; 3] & B_2 = [2, 5; 3, 5] & B_3 = [2; 3, 5] & B_4 = [2; 4] \end{array}$$

dont les affectations probabilistes sont toutes égales :

$$m(A_i) = m(B_j) = 0,25 \quad \forall i, j = 1, \dots, 4$$

La propagation de l'ensemble aléatoire $\mathcal{A} \times \mathcal{B}$ se fait avec la fonction $f(a, b) = a^b$, et la croyance associée à l'ensemble $[y; +\infty[$ sur Y , pour y fixé, est calculée. On déduit de ce calcul que le risque d'avoir $y > 12$ est d'au moins 10^{-1} , vu que la croyance est un minorant de la probabilité vraie.

Si l'on n'avait pas voulu utiliser l'information telle qu'elle était disponible, et ramener le problème dans un cadre plus classique, tel que celui des probabilités, on aurait dû assigner des densités de probabilité arbitraires pour la répartition des éléments a et b dans les différents ensembles focaux A_i et B_i respectivement. Le principe du maximum d'entropie, décrit dans le paragraphe 1.2.2, nous permet d'opter pour une loi uniforme, afin de minimiser les effets de l'hypothèse d'une loi de probabilité sur les éléments focaux en jouant la « carte » de l'incertitude maximale.

Les densités de probabilité résultantes pour a et b sont ensuite propagées à l'aide de la fonction $f(a, b) = a^b$, et on obtient une probabilité d'avoir $y > 12$ égale à 10^{-3} , soit un facteur cent dans le sens de l'insécurité vis-à-vis de la solution avec les intervalles aléatoires. On a représenté sur la figure 2.1 respectivement les fonctions $y \mapsto \text{Bel}([y; +\infty[)$ et $y \mapsto P(Y > y) = 1 - F(y)$ avec F la fonction de répartition de Y ; on constate alors que les deux représentations de l'incertitude sur y sont très différentes.

Encore une fois, on peut constater que le fait d'imposer une modélisation stochastique, alors que les informations sont insuffisantes pour valider une telle hypothèse, a conduit à des résultats très différents de ceux que l'on obtient en se tenant rigoureusement à la connaissance du système.

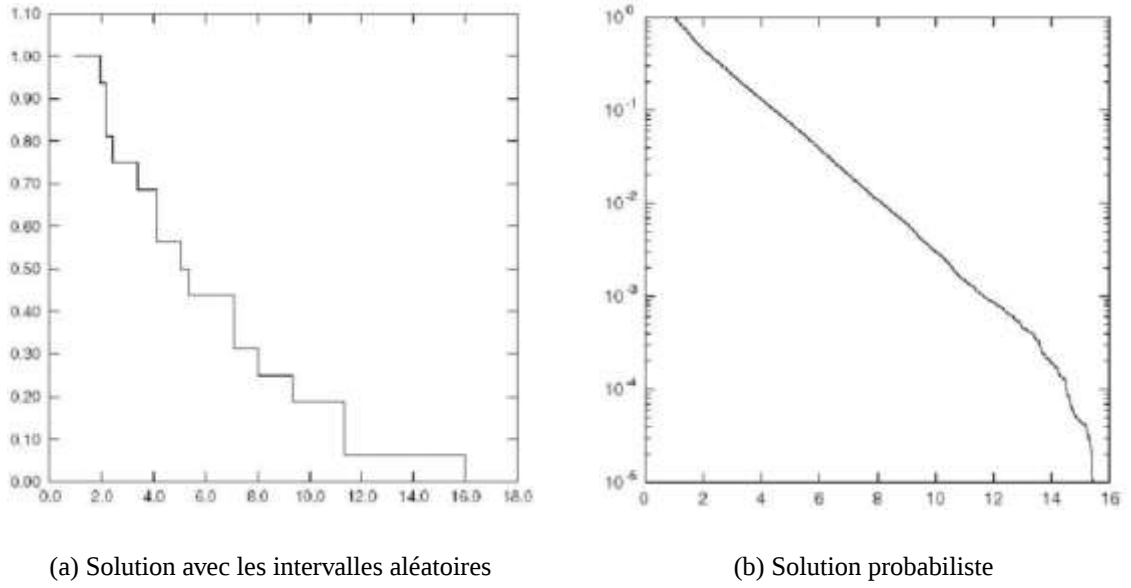


FIG. 2.1 – Comparaison des résultats obtenus sur un exemple simple par une représentation par intervalles aléatoires et une modélisation probabiliste, d’après [Oberkamp *et al.* 2001]

2.5.3 Conclusion

Toutes les méthodes non stochastiques que l’on vient de présenter sont une réelle alternative aux méthodes reposant sur une modélisation stochastique des incertitudes ; même si l’on n’a pas totalement balayé la totalité des méthodes existantes, on a pu juger des idées et capacités de ces techniques. En particulier, on peut encore évoquer la théorie possibiliste, dont un très rapide survol est donné ci-dessous.

Théorie possibiliste

Cette théorie [Dubois et Prade 1986] introduit la notion de distribution de possibilité $\Pi(x)$ (*possibility distribution* en anglais) qui est en fait la fonction de participation d’un ensemble flou. À partir de cette définition, il est possible d’associer pour tout ensemble A deux mesures analogues à la croyance et la plausibilité associées aux ensembles aléatoires : il s’agit de la nécessité *Nec* et la possibilité *Pos* (respectivement *Necessity* et *Possibility* en anglais), définis par :

$$\text{Nec}(A) = \max_{x \in A} \Pi(x) \quad (2.60a)$$

$$\text{Pos}(A) = 1 - \max_{x \in A} \Pi(x) \quad (2.60b)$$

Quoique analogues aux deux mesures associées aux ensembles aléatoires, la nécessité et la possibilité ont des propriétés différentes, que nous ne détaillerons pas ici.

Bilan

Afin de conclure sur toutes ces méthodes, on a représenté sur la figure 2.2 un panorama des différentes modélisations non stochastiques, ainsi que les liens qui les unissent. Bien que le point de départ soit constitué par les ensembles aléatoires de Dempster-Shafer, il faut noter l'existence de théories mathématiques bien plus générales qui englobent tous ces objets, et dont la présentation dépasse le cadre de notre présente étude ; on peut citer la théorie des « probabilités imprécises » (*imprecise probabilities* en anglais) pour laquelle on trouvera plus de détails dans [Walley 1990] ou encore [IPP]. Les boîtes de probabilité, décrites précédemment, sont un mode de représentation courant associé à cette théorie, et dans ce cadre, un ensemble aléatoire de Dempster-Shafer constitue une approximation discrète possible d'une certaine boîte de probabilité.

Validation de modèles

L'utilisation dans le domaine de la mécanique des méthodes non stochastiques connaît un succès croissant depuis quelques années ; toutefois il n'existe pas véritablement d'applications concernant la validation de modèles. Les seules utilisations de données expérimentales se résument généralement à des comparaisons, et éventuellement à des processus d'identification, comme dans [Hanss et al. 2002] où les auteurs identifient un modèle de liaison défini par des paramètres de raideur et d'amortissement flous.

2.6 Les modèles convexes d'incertitude

2.6.1 Cadre mathématique associé

Dans [Ben-Haim et Elishakoff 1990] a été introduite une généralisation originale du concept d'intervalle déterministe, qui permet en outre de constituer un cadre novateur à toutes les théories de prise en compte de l'incertitude que nous avons évoquées plus haut ; elle s'est développée dans le domaine de la fiabilité des structures à partir du constat, évoqué dans le paragraphe 2.1, qu'une modélisation purement probabiliste des incertitudes pouvait poser problème, et s'avérer abusive dans le cas d'un fort manque d'informations.

Le cadre des modèles convexes d'incertitude permet de préciser la capacité d'intégrer de l'information nouvelle dans le cas où cette dernière n'est pas déterministe, d'une façon différente de ce que pouvaient décrire les ensembles flous ou aléatoires : **l'incertitude présente dans l'information nouvelle est ici vue comme la « lacune » d'information entre ce qui est connu et ce qui devrait l'être pour rendre possible une prise de décision fiable**, d'où l'autre dénomination des modèles convexes d'incertitude : les modèles de lacune d'information (*info-gap models* en anglais).

Les modèles convexes utilisés sont des ensembles caractérisés par deux quantités différentes :

- un réel $\alpha \geq 0$;
- un élément $\underline{u}$ de l'espace sur lequel le modèle convexe est défini.

Ceci est fait de la façon suivante :

$$\mathcal{U}(\alpha, \underline{u}) = \{ \underline{u} \mid \|\underline{u} - \underline{u}\|^2 \leq \alpha^2 \} \quad (2.61)$$

où $\|\cdot\|$ est une certaine norme définie sur l'espace des éléments u .

L'exemple concret suivant permet d'illustrer cette définition : si l'on considère les forces incertaines $\underline{u}(t)$ s'exerçant pendant un séisme sur un immeuble, une force donnée mesurée s'écarte d'une certaine quantité incertaine de $\underline{u}(t)$, qui est une excitation nominale ou typique connue par expérience ou issue d'hypothèses. L'ensemble de toutes les forces dont l'énergie cumulée mise en œuvre est bornée par α^2 constitue le modèle convexe d'incertitude suivant :

$$\mathcal{U}(\alpha, \underline{u}) = \left\{ \underline{u}(t) \mid \int_0^{+\infty} \|\underline{u} - \underline{u}\|^2 dt \leq \alpha^2 \right\} \quad (2.62)$$

L'incertitude se traduit à deux niveaux différents au sein de ce type de modèle.

Pour une valeur de α fixée, l'ensemble $\mathcal{U}(\alpha, \underline{u})$ représente le degré de variabilité incertaine de la force $\underline{u}(t)$: plus α est grand, plus la variation envisageable de $\underline{u}(t)$ par rapport à $\underline{u}(t)$ l'est également ; c'est pourquoi α est généralement qualifié de paramètre d'incertitude. De plus, la valeur de α est généralement inconnue, ce qui constitue le second niveau de description de l'incertitude. $\underline{u}(t)$ est quant à lui nommé centre de $\mathcal{U}(\alpha, \underline{u})$.

Les modèles convexes d'incertitude présentent différentes propriétés utiles :

- **une propriété d'« imbrication »**, proche de la notion de consonance introduite précédemment à propos des ensembles aléatoires : $\alpha \leq \beta$ implique que :

$$\mathcal{U}(\alpha, \underline{u}) \subset \mathcal{U}(\beta, \underline{u}) \quad (2.63)$$

- **une propriété de « contraction »** : le modèle $\mathcal{U}(0, \underline{0})$ est un singleton contenant son centre :

$$\mathcal{U}(0, \underline{0}) = \{\underline{0}\} \quad (2.64)$$

- **une propriété de translation** :

$$\mathcal{U}(\alpha, \underline{u}) = \underline{u} + \mathcal{U}(\alpha, 0) \quad (2.65)$$

où $\underline{u} + \mathcal{U}$ signifie que $\underline{u}$ est ajouté à chaque élément de $\mathcal{U}$;

- **une propriété d'expansion linéaire** :

$$\mathcal{U}(\beta, \underline{0}) = \frac{\beta}{\alpha} \mathcal{U}(\alpha, \underline{0}) \quad \forall \alpha \neq 0 \quad (2.66)$$

où $\beta \mathcal{U}$ signifie que β multiplie chaque élément de $\mathcal{U}$.

2.6.2 Applications dans le domaine de la fiabilité

Les premières applications basées sur les modèles convexes d'incertitude ont permis de mesurer la fiabilité d'un système comme la taille de la plus grande lacune d'information n'entraînant pas de défaillance ; ainsi définie, **la fiabilité d'un système est quantifiée comme le degré d'immunité vis-à-vis de l'incertitude, elle est d'autant plus grande que sa vulnérabilité face à l'« imprévu » est petite** [Ben-Haim 1995].

Ainsi, si $\underline{q}$ représente l'ensemble des paramètres de conception d'un système, et $R(\underline{q}, \underline{u})$ est une fonction de performance, semblable à celles définies dans les méthodes fiabilistes décrites dans le paragraphe 1.5.1, cette fonction, qui fait intervenir à la fois les paramètres de conception $\underline{q}$ et la réponse incertaine $\underline{u}$ du modèle, doit être plus grande qu'un certain seuil r_c pour qu'il n'y ait pas de défaillance. On définit alors la robustesse $\hat{\alpha}(\underline{q}, r_c)$ du système comme la plus grande valeur du paramètre d'incertitude α telle qu'il n'y ait pas de défaillance :

$$\hat{\alpha}(\underline{q}, r_c) = \arg \max_{\alpha} \left\{ \min_{\underline{u} \in \mathcal{U}(\alpha, \underline{u})} R(\underline{q}, \underline{u}) \geq r_c \right\} \quad (2.67)$$

Le choix des paramètres de conception $\underline{q}$ consiste alors à faire en sorte d'avoir :

$$\hat{\alpha}(\underline{q}, r_c) \geq \tilde{\alpha}_d \quad (2.68)$$

où $\tilde{\alpha}_d$ est une valeur seuil qui peut être déterminée de bien des façons [Ben-Haim 1999].

Divers domaines d'applications sont envisageables ; on trouvera en particulier dans [Vinot *et al.* 2002] des exemples d'études dans le domaine de la dynamique des structures.

Bien que les modèles convexes soient les plus utilisés, il est possible de généraliser la plupart des concepts précédents dans le cas de modèles linéaires seulement, voire pour des modèles ne vérifiant que la propriété d'imbrication et une forme affaiblie de la propriété de translation [Ben-Haim 2000]. Il est alors possible de traiter le problème de la fiabilité de systèmes dans des cas de vibration forcée, vis-à-vis de leur fréquence d'excitation.

Bien que l'approche par modèles de lacune d'information ne nécessite pas l'introduction de lois de probabilité, il est tout à fait possible de combiner des hypothèses de nature stochastique avec les ensembles définis précédemment [Ben-Haim 1999]. En fait, de façon plus générale, **on peut interpréter les modèles de lacune d'information comme un cadre d'application de toutes les représentations de l'incertitude existantes** : comme précisé dans [Hemez 2004], on peut associer à un modèle probabiliste de variabilité un ensemble convexe dont le paramètre d'incertitude α contrôle directement les termes de la matrice de covariance ; on peut aussi envisager de considérer des fonctions de participation, utilisées par exemple pour traduire un jugement d'expert, en les paramétrant à l'aide de α .

2.6.3 Applications dans le domaine de la validation de modèles

Dans [Ben-Haim et Hemez 2004], les auteurs utilisent la description de l'incertitude à l'aide des modèles de lacune d'information pour montrer des résultats intéressants la validation des modèles.

Le constat initial est le suivant : lors de la calibration d'un modèle grâce à des mesures expérimentales, il n'est pas sain de viser la fidélité par rapport aux données à tout prix ; en effet, on peut montrer que des modèles optimaux, au sens qu'ils minimisent les écarts avec les données expérimentales, n'ont aucune robustesse vis-à-vis des incertitudes, c'est-à-dire que de légères variations dans les paramètres du modèle, ou dans la forme mathématique du modèle, peuvent conduire à des prédictions très peu fidèles [Ben-Haim 2001].

On représente de façon générale un modèle donné par :

$$\underline{y} = M(\underline{p}, \underline{q}) \quad (2.69)$$

où $\underline{y}$ constitue la réponse mesurable du modèle, et $\underline{p}$ et $\underline{q}$ les entrées, représentant respectivement la configuration d'étude et la constitution mathématique du modèle ; plus précisément, tandis que $\underline{p}$ caractérise des réglages de l'ordre du protocole expérimental, $\underline{q}$ est associée aux paramètres du modèle, mais aussi à la forme fonctionnelle de ce modèle, voire au fait que plusieurs types de modèles différents pourraient être utilisés. C'est pourquoi on décide d'associer à l'ensemble des modèles possibles un modèle de lacune d'information de la forme :

$$\mathcal{U}(\alpha, \bar{\underline{q}}) = \{ M(\underline{p}, \underline{q}) \mid \|\underline{q} - \bar{\underline{q}}\| \leq \alpha \} \quad (2.70)$$

Les auteurs [Ben-Haim et Hemez 2004] préconisent alors l'introduction de nouveaux critères d'évaluation d'un modèle, à savoir la robustesse et la dispersion de prédiction. Précisons ces trois notions :

- **la fidélité vis-à-vis des données** est associée simplement à la distance R entre la réponse $\underline{y}$ du modèle choisi et la réponse $\underline{y}_{\text{exp}}$ expérimentale, selon une certaine norme :

$$R = \|\underline{y}_{\text{exp}} - \underline{y}\| \quad (2.71)$$

plus cette distance est petite, plus le modèle est fidèle vis-à-vis des données expérimentales ;

- **la robustesse** α est en fait l'étendue de la dispersion des paramètres $\underline{q}$ qui conduisent à une erreur R inférieure à un certain seuil R_M ; ainsi toutes les prédictions menées avec des paramètres $\underline{q}$ à l'intérieur de $\mathcal{U}(\alpha, \bar{\underline{q}})$ permettent d'obtenir le niveau de fidélité désiré ;
- **la dispersion de prédiction** $\Delta \underline{y}$ est l'étendue des réponses que l'on peut obtenir quand on utilise chaque modèle de $\mathcal{U}(\alpha, \bar{\underline{q}})$:

$$\Delta \underline{y} = \max_{\underline{p} \in \mathcal{U}(\alpha, \bar{\underline{q}})} M(\underline{p}_k, \underline{q}) - \min_{\underline{p} \in \mathcal{U}(\alpha, \bar{\underline{q}})} M(\underline{p}_k, \underline{q}) \quad (2.72)$$

où $\underline{p}_k$ est une configuration de test donnée. L'importance de cette dispersion est liée au fait que, plus les prédictions issues des différents modèles de $\mathcal{U}(\alpha, \bar{\underline{q}})$ sont proches les unes des autres, plus on peut avoir confiance dans ces prédictions.

Plutôt donc que de minimiser directement l'écart entre la réponse du modèle et la mesure effectuée, une stratégie analogue à la relation (2.67) utilisée dans l'étude de la fiabilité des structures est proposée :

$$\hat{\alpha} = \arg \max_{\alpha \geq 0} \left\{ \min_{M \in \mathcal{U}(\alpha, \bar{q})} R \leq R_M \right\} \quad (2.73)$$

Dans ce cas, $\hat{\alpha}$ représente la plus grande quantité d'incertitude concernant la définition du modèle assurant une fidélité suffisante vis-à-vis des quantités mesurées. Notons encore une fois que l'utilisation des modèles de lacune d'information permet si nécessaire d'introduire toute forme de modélisation de l'incertitude (probabiliste, floue, ...).

À partir de cette représentation, **on peut mettre en évidence tous les compromis qu'il est nécessaire d'avoir à l'esprit lors de la calibration d'un modèle :**

$$\frac{\partial \Delta y}{\partial \alpha} \geq 0 \quad \frac{\partial \alpha}{\partial R} \geq 0 \quad \frac{\partial \Delta y}{\partial R} \geq 0 \quad (2.74)$$

Plus précisément :

- la dispersion de prédiction augmente quand la robustesse augmente ;
- la robustesse diminue quand la fidélité augmente, ce qu'on a déjà évoqué plus haut ;
- la dispersion de prédiction diminue quand la fidélité augmente.

Ainsi, bien qu'il soit préférable d'avoir des modèles robustes, ceci a tendance à conduire à des prédictions plus dispersées, et implique que l'obtention de prédictions plus précises est conditionnée par les hypothèses sur lesquelles le modèle est basé.

Dans un registre proche de la question de la validation de modèles, une procédure robuste de sélection des emplacements de capteurs a pu être validée dans [Vinot et Cogan 2004] en utilisant des modèles de lacune d'information afin de modéliser de façon neutre les incertitudes concernant la structure instrumentée, et l'effet de ces dernières sur le placement optimal des capteurs.

Description envisagée de la réalité

La science cherche le mouvement perpétuel. Elle l'a trouvé : c'est elle-même.

Victor Hugo

Sommaire

3.1	Réalité considérée	54
3.1.1	Famille de structures réelles	54
3.1.2	Données expérimentales retenues	54
3.2	Simulation stochastique d'une famille de structures	56
3.2.1	Principe	56
3.2.2	Structure étudiée	56
3.2.3	Obtention des valeurs caractéristiques	58

Avant d'aborder les principes de la théorie des méconnaissances, nous décrivons dans ce troisième chapitre comment est considérée la « réalité » que l'on souhaite étudier, et plus particulièrement les quantités représentatives qui peuvent en être extraites. Nous présentons ensuite un exemple classique de simulation du réel à l'aide d'une modélisation stochastique par la technique de Monte Carlo ; à partir d'un certain nombre de structures simulées, nous calculons les quantités représentatives qui nous serviront de bases de comparaison avec les résultats de la théorie des méconnaissances dans le chapitre 5.

3.1 Réalité considérée

3.1.1 Famille de structures réelles

La réalité telle que nous la considérerons dans la suite de ce mémoire est constituée d'une famille de structures semblables, mais non rigoureusement identiques ; en effet, diverses sources d'incertitudes peuvent provenir de l'élaboration même de ces structures, entraînant des dispersions notables, notamment au niveau des caractéristiques des matériaux utilisés ainsi que des spécifications géométriques. Dans ce travail, **on se limite à des dispersions concernant seulement les rigidités structurales**, vu que la théorie des méconnaissances que nous allons introduire ne permet de prendre en compte que des incertitudes de ce type pour le moment.

Bien entendu, nous ne pouvons avoir accès directement à la connaissance de ces incertitudes au travers de cette famille de structures réelles : nous devons passer par l'intermédiaire de quantités mesurables expérimentalement. Dans le domaine de l'analyse modale, il s'agit généralement de déterminer les pulsations propres et les formes propres de chaque structure, qui sont des quantités d'intérêt définies sur l'ensemble de celle-ci. De manière plus synthétique, on désigne dans la suite une quantité d'intérêt quelconque à l'aide de la notation α : par exemple, on peut se préoccuper en particulier de la troisième pulsation propre chez les structures étudiées.

Quand on réalise des expérimentations sur l'ensemble des structures de la famille, on obtient un ensemble de pulsations propres et formes propres mesurées ; pour reprendre la notation que l'on vient juste d'introduire, **on dispose pour une quantité d'intérêt α donnée d'un ensemble de valeurs expérimentales α_{exp} issues des différentes structures testées.**

3.1.2 Données expérimentales retenues

En vue d'une comparaison ultérieure avec les résultats fournis par la théorie des méconnaissances, **on souhaite extraire de l'ensemble des valeurs α_{exp} quelques quantités représentatives de la dispersion de la quantité d'intérêt étudiée.** Typiquement, on associe alors à la famille de structures réelles deux quantités $\alpha_{\text{exp}}^-(P)$ et $\alpha_{\text{exp}}^+(P)$ qui, pour une valeur P donnée, encadrent une proportion P des quantités mesurées α_{exp} .

Pour éviter toute ambiguïté et obtenir ainsi un couple unique de quantités pour chaque valeur de P , on impose également que ces deux quantités soient celles qui définissent le plus petit intervalle contenant une proportion P de valeurs mesurées α_{exp} . On a illustré le principe de ces « valeurs expérimentales à P % » sur la figure 3.1.

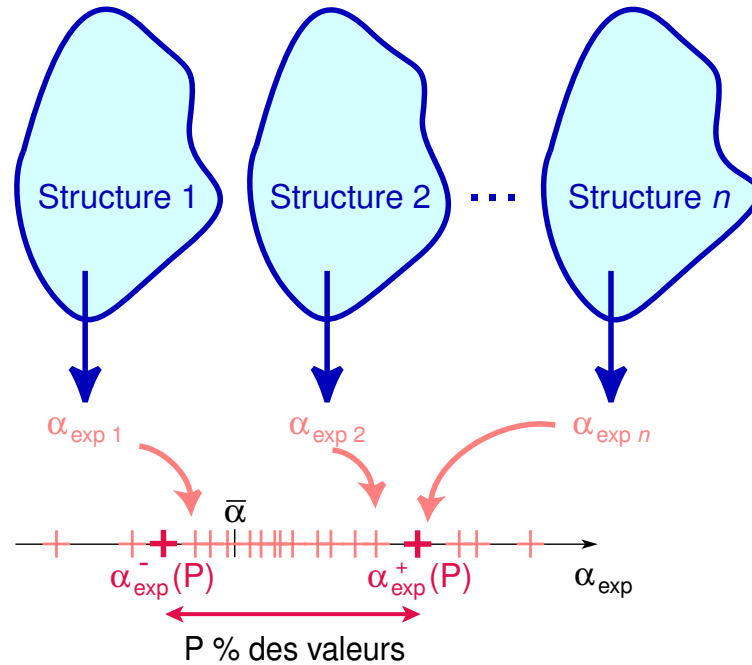


FIG. 3.1 – Valeurs expérimentales à P % de la quantité d'intérêt α_{exp}

En pratique, on se contente de certaines valeurs bien précises de P :

- les « **valeurs à 99 %** » $\alpha_{\text{exp}}^-(0, 99)$ et $\alpha_{\text{exp}}^+(0, 99)$, correspondant à $P = 0, 99$, permettent d'apprécier l'étendue presque complète des valeurs que la quantité d'intérêt expérimentale est susceptible de prendre ; le fait de choisir 99 % plutôt que 100 % permet de filtrer des cas aberrants ou très rares qui peuvent survenir, et que l'on estime non représentatifs de la famille de structures étudiée ;
- les « **valeurs à 70 %** » $\alpha_{\text{exp}}^-(0, 70)$ et $\alpha_{\text{exp}}^+(0, 70)$, correspondant à $P = 0, 70$, permettent d'avoir un ordre de grandeur, plus ou moins juste suivant la répartition des valeurs mesurées, de l'écart type de la quantité d'intérêt expérimentale.

Il est en fait inutile de vouloir conserver une description plus précise de la réalité dans la mesure où le nombre de structures réelles effectivement testées est souvent relativement faible, étant donné le coût généralement important de la mise en œuvre de mesures expérimentales.

Ceci est d'ailleurs une des raisons pour lesquelles nous avons décidé dans ce travail de considérer des exemples académiques d'étude simulés numériquement, dans le but de représenter une certaine « réalité » expérimentale ; en contrôlant parfaitement les dispersions présentées par notre famille de structures « réelles » simulées, nous disposerons de résultats qui pourront être confrontés ultérieurement aux prédictions de la théorie des méconnaissances. Nous aurons par contre l'occasion dans le chapitre 8 d'utiliser de véritables données expérimentales issues d'une structure industrielle réelle.

3.2 Simulation stochastique d'une famille de structures

3.2.1 Principe

Nous introduisons ici un exemple académique d'étude dont les résultats pourront servir de « référence » dans la suite de notre travail. En l'occurrence, on fait l'hypothèse d'une réalité parfaitement connue, qui peut être simplement décrite à l'aide d'une méthode stochastique, dont l'utilisation constitue un moyen courant de simuler une réalité « expérimentale ». Afin d'obtenir les dispersions présentes dans la famille considérée de structures semblables, nous décidons de représenter les incertitudes à l'aide de paramètres stochastiques qui, dans le cadre de ce travail, se limitent aux rigidités structurales.

En pratique, nous avons besoin de définir la nature des paramètres stochastiques intervenant dans la simulation ; nous avons déjà eu l'occasion de discuter de ce problème dans le paragraphe 1.2.1 et nous avons montré que, par l'intermédiaire de la décomposition de Karhunen-Loève (1.3), tous les paramètres stochastiques peuvent être représentés à l'aide de variables aléatoires, dont la manipulation est plus aisée.

À partir de là, une façon simple et directe de simuler la famille des structures semblables consiste à utiliser une technique de Monte Carlo qui permet d'obtenir des réalisations des différentes variables aléatoires intervenant dans le problème, selon les lois de probabilité choisies. Chaque structure de la famille est alors caractérisée par un jeu donné de paramètres obtenus par tirage, et présente des pulsations propres et formes propres de vibrations que l'on peut obtenir classiquement comme on le ferait avec un modèle déterministe.

Pour chaque quantité d'intérêt, on a donc un ensemble de valeurs obtenues pour les différentes structures considérées, et il est alors possible de déterminer les valeurs à P % $\alpha_{\text{exp}}^-(P)$ et $\alpha_{\text{exp}}^+(P)$, comme ce qui serait effectué sur les résultats issus d'une famille de structures réelles.

3.2.2 Structure étudiée

L'exemple d'étude est un treillis plan composé de quatre barres rectilignes, reliées par des rotules, comme représenté sur la figure 3.2. On suppose que les liaisons avec le bâti aux nœuds 1 et 2 sont parfaites, et que les barres ne travaillent qu'en traction-compression. Les caractéristiques de ces barres dont le matériau est supposé homogène et isotrope sont données dans le tableau 3.1 : elles permettent de définir le modèle théorique déterministe qui sert de base au modèle stochastique.

Le problème aux valeurs propres associé au modèle déterministe permet de déterminer quatre modes propres de vibrations, représentés sur la figure 3.3 ; les composantes de ces modes $\bar{\phi}_i$ (normalisés par rapport à leur maximum), ainsi que les carrés des pulsations propres $\bar{\omega}_i^2$ associées, sont consignés dans le tableau 3.2 ; on a posé $\bar{\phi}_i = (x_3 \ y_3 \ x_4 \ y_4)^T$, où $\underline{U}_j = x_j \underline{x} + y_j \underline{y}$ est le déplacement du nœud j .

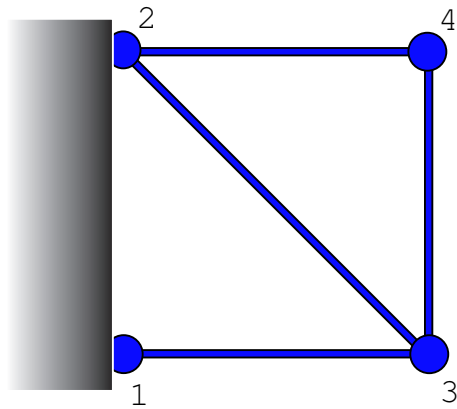


FIG. 3.2 – Treillis plan à quatre barres

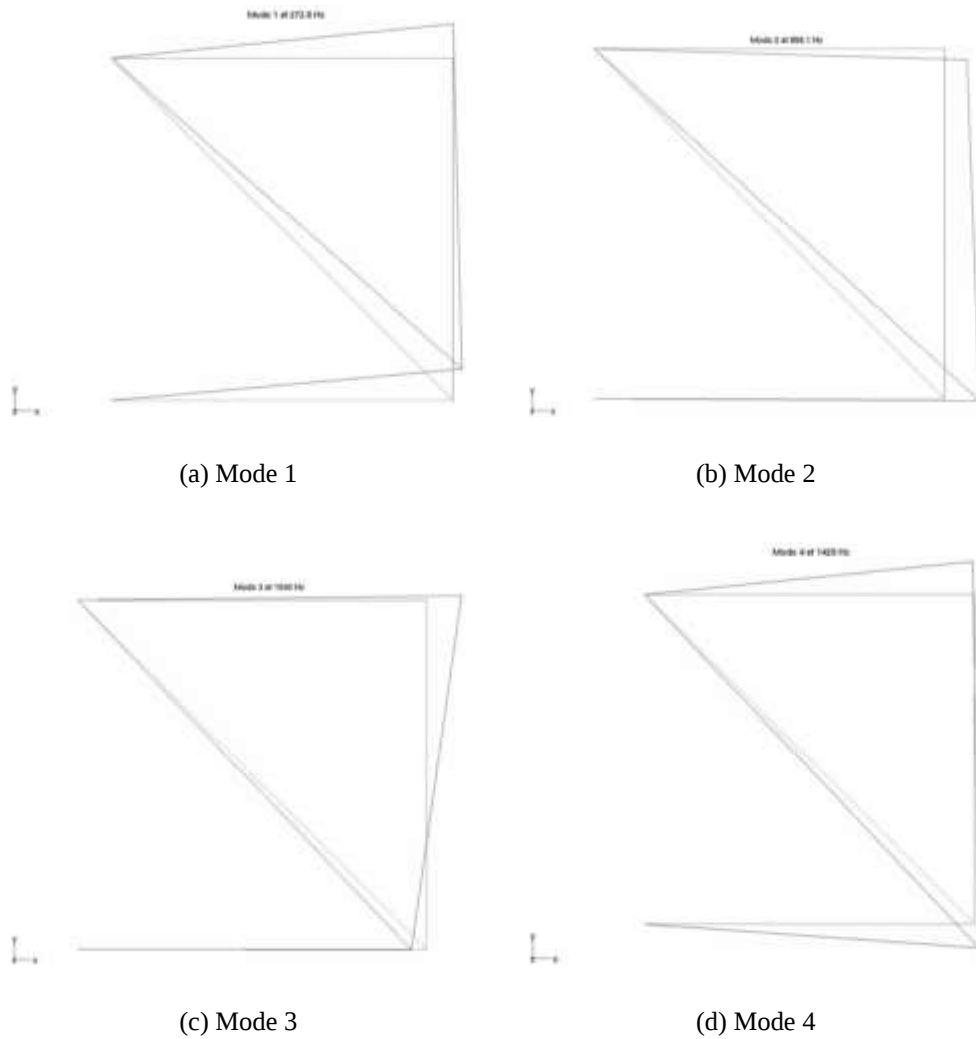


FIG. 3.3 – Modes propres associés au modèle déterministe du treillis à quatre barres

Barres	Section	Longueur	Matériau	Module d'Young	Masse vol.
{1-3,2-4,3-4}	10^{-4}m^2	1m	aluminium	$\bar{E}_{\{b13,b24,b34\}} = 72\text{GPa}$	2700kg/m^3
2-3	10^{-4}m^2	$\sqrt{2}\text{m}$	aluminium	$\bar{E}_{b23} = 72\text{GPa}$	2700kg/m^3

TAB. 3.1 – Caractéristiques du modèle déterministe du treillis à quatre barres

Quantités	$i = 1$	$i = 2$	$i = 3$	$i = 4$
Pulsations $\bar{\omega}_i^2$	$3,02 \cdot 10^6$	$2,93 \cdot 10^7$	$4,54 \cdot 10^7$	$8,33 \cdot 10^7$
Modes $\bar{\phi}_i$	$\begin{pmatrix} 0,26 \\ 0,91 \\ 0,01 \\ 1,00 \end{pmatrix}$	$\begin{pmatrix} 1,00 \\ -0,07 \\ 0,69 \\ -0,30 \end{pmatrix}$	$\begin{pmatrix} -0,48 \\ -0,01 \\ 1,00 \\ 0,14 \end{pmatrix}$	$\begin{pmatrix} 0,13 \\ -0,71 \\ -0,07 \\ 1,00 \end{pmatrix}$

TAB. 3.2 – Détail des pulsations propres et modes propres associés au modèle déterministe du treillis à quatre barres

3.2.3 Obtention des valeurs caractéristiques

Le modèle stochastique est obtenu en introduisant, à partir du modèle théorique déterministe précédent, des dispersions de raideurs caractérisées sur chaque barre par une variable aléatoire donnée. Le choix de la loi de probabilité est guidé selon les réflexions menées dans le paragraphe 1.2.2 : si l'on suppose que les sources d'incertitude proviennent uniquement des coefficients matériaux, on peut décider de choisir une loi de probabilité normale pour chacune de ces variables aléatoires ; ceci revient bien sûr à considérer dans ce cas que le comportement de ces barres est seulement lié à la donnée de leur module d'Young.

On décide par exemple d'utiliser pour la simulation de chaque barre une loi normale centrée, mais bornée (et renormalisée) pour éviter l'occurrence de réalisations non physiques, comme on a déjà eu l'occasion de l'évoquer auparavant ; par convention, on peut poser que les bornes que l'on introduit sont les valeurs qui encadrent 99 % des réalisations de la variable aléatoire selon la loi normale non bornée. Ici, on choisit d'appliquer sur la variable aléatoire de chaque barre une loi normale bornée d'étendue ± 5 % de la valeur théorique de la raideur de la barre considérée.

Il suffit alors d'utiliser une simulation de Monte Carlo afin d'obtenir plusieurs jeux de paramètres de raideurs, chaque jeu correspondant à une structure donnée de la famille considérée. On effectue alors le calcul des pulsations propres et des formes propres de chacune des structures simulées, et on obtient donc pour chaque quantité d'intérêt une certaine dispersion de valeurs. Par exemple, on a représenté sur la figure 3.4 la répartition des occurrences des carrés des secondes pulsations propres obtenues avec une méthode de Monte Carlo mettant en œuvre 50000 tirages afin d'obtenir une bonne convergence des valeurs caractéristiques estimées.

On peut ainsi déterminer de ces répartitions les valeurs à P % qui constituent une description suffisante de la famille de structures simulées ; on a répertorié dans le tableau 3.3 les valeurs correspondant aux carrés des pulsations propres et des valeurs en un DDL des modes, pour des proportions $P = 0,70$ et $P = 0,99$.

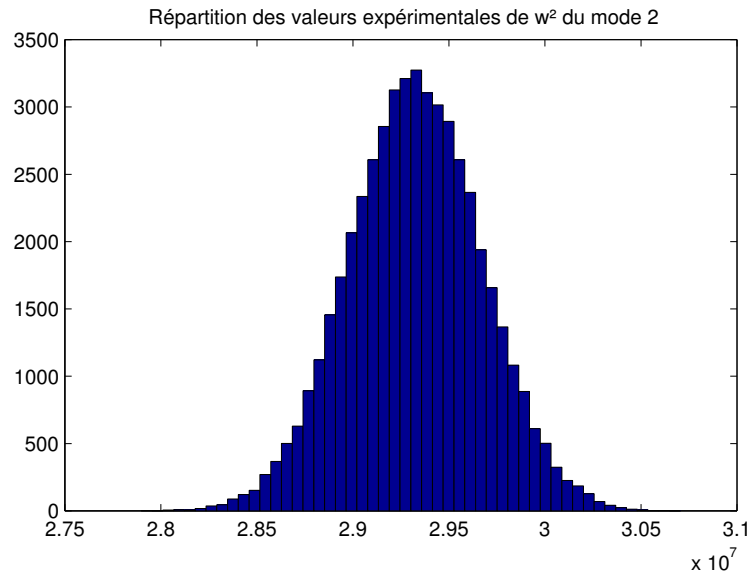


FIG. 3.4 – Répartitions des valeurs de $\omega_{2\text{exp}}^2$ obtenues pour le treillis à quatre barres à l'aide de la méthode de Monte Carlo (50000 tirages)

Valeurs de P	Modes i	Pulsations propres $[\omega_{i\text{exp}}^{2-}; \omega_{i\text{exp}}^{2+}]$	DDL k	Valeurs en un DDL des modes $[\phi_{ki\text{exp}}^{-}; \phi_{ki\text{exp}}^{+}]$
0,70	1	$[2,97; 3,06] \cdot 10^5$	2	$[0,90; 0,91]$
	2	$[2,90; 2,97] \cdot 10^6$	3	$[0,64; 0,74]$
	3	$[4,47; 4,61] \cdot 10^6$	1	$[-0,50; -0,45]$
	4	$[8,18; 8,49] \cdot 10^6$	2	$[-0,72; -0,71]$
0,99	1	$[2,91; 3,13] \cdot 10^5$	2	$[0,90; 0,91]$
	2	$[2,84; 3,02] \cdot 10^6$	3	$[0,58; 0,83]$
	3	$[4,37; 4,72] \cdot 10^6$	1	$[-0,54; -0,43]$
	4	$[7,95; 8,71] \cdot 10^6$	2	$[-0,72; -0,71]$

TAB. 3.3 – Valeurs à 70 % et 99 % obtenues avec le modèle stochastique du treillis à quatre barres

Nous aurons l'occasion de comparer ces résultats avec ceux obtenus à l'aide de différents modèles de méconnaissances présentés dans le chapitre 5.

Définition du concept de méconnaissance

Définir, c'est entourer d'un mur de mots un terrain vague d'idées.

Samuel Butler

Sommaire

4.1	Principes de la modélisation des méconnaissances	62
4.1.1	Définition des méconnaissances de base	62
4.1.2	Choix des lois de probabilité	63
4.2	Exemple : modélisation des dispersions matériaux	64
4.2.1	Justification du choix d'une loi normale	65
4.2.2	Conversion de la loi normale en termes de méconnaissances	66
4.3	Concept de probabilité d'intervalle	66
4.3.1	Intervalles standards et probabilité d'intervalle	66
4.3.2	Interprétation	69
4.3.3	Comparaison avec les méthodes non stochastiques	71
4.4	Obtention pratique et manipulation des méconnaissances de base	72

Nous introduisons dans ce quatrième chapitre le concept de méconnaissance, qui a pour but la manipulation de modèles imparfaits par rapport à la réalité ; pour cela, nous détaillons en particulier les idées intervenant dans la modélisation des méconnaissances de base, ainsi que les outils mathématiques permettant une description plus aisée de celles-ci. Quelques comparaisons avec les différentes méthodes non stochastiques étudiées dans le chapitre précédent sont aussi menées.

4.1 Principes de la modélisation des méconnaissances

Le concept de méconnaissance a pour but de chiffrer simplement la comparaison que l'on effectue entre le modèle et le réel, et ainsi d'être capable de manipuler des modèles qui ne sont qu'une représentation imparfaite de la réalité. La théorie se place dans le cadre d'une structure modélisée comme un assemblage de sous-structures, sachant que les liaisons peuvent être assimilées à des sous-structures particulières. **C'est à l'échelle de ces sous-structures que l'on choisit de quantifier les incertitudes ;** on globalise ainsi facilement toutes les sources d'erreurs qui se produisent à divers degrés plus ou moins locaux de la sous-structure : hétérogénéités, inclusions plastiques, défauts mineurs localisés, etc. Pour le moment, toutefois, on ne considère que des incertitudes concernant les rigidités structurales, mais d'autres types sont parfaitement envisageables par la suite à l'aide du concept de méconnaissance.

4.1.1 Définition des méconnaissances de base

La théorie des méconnaissances repose sur l'utilisation d'un modèle théorique déterministe donné, auquel on ajoute un « modèle de méconnaissances » ; celui-ci consiste à considérer que, à chaque sous-structure E d'une structure réelle Ω , on peut associer une quantité m évoluant quelque part à l'intérieur d'un intervalle dont les bornes m_E^+ et m_E^- sont définies formellement comme suit :

$$(1 - m_E^-) \bar{K}_E \leq K_E \leq (1 + m_E^+) \bar{K}_E \quad (4.1)$$

où K_E et $\bar{K}_E$ désignent la matrice de rigidité associée à la sous-structure E , respectivement pour la structure réelle et pour le modèle théorique déterministe.

Cette association ne se fait pas directement à travers la quantité m , dont on ne sait rien si ce n'est qu'elle appartient avec certitude à l'intervalle $[-m_E^-; m_E^+]$; c'est donc au travers de ces deux bornes m_E^- et m_E^+ , qui sont définies de façon à être forcément positives, qu'est véhiculée la notion de méconnaissance. **Ces deux quantités relatives à la sous-structure E forment ce que l'on appelle la méconnaissance de base sur la sous-structure E .**

L'inégalité matricielle (4.1) s'interprète au sens des valeurs propres, mais pour traduire en pratique cette inégalité, on choisit de faire intervenir les énergies de déformation :

$$(1 - m_E^-) \bar{e}_E(\underline{U}) \leq e_E(\underline{U}) \leq (1 + m_E^+) \bar{e}_E(\underline{U}) \quad (4.2)$$

où :

- la quantité $e_E(\underline{U}) = \frac{1}{2} \underline{U}^T \mathbf{K}_E \underline{U}$ désigne l'énergie de déformation dans le champ de déplacement $\underline{U}$ d'une structure réelle appartenant à la famille de structures semblables étudiées ;
- la quantité $\bar{e}_E(\underline{U}) = \frac{1}{2} \underline{U}^T \bar{\mathbf{K}}_E \underline{U}$ représente l'énergie de déformation dans le champ de déplacement $\underline{U}$ du modèle théorique déterministe considéré.

Cette inégalité doit être vérifiée quel que soit le champ de déplacement $\underline{U}$; ainsi formulée, elle est totalement équivalente à la définition de la relation (4.1).

Pour chaque sous-structure, la quantité m est située à l'intérieur de l'intervalle de méconnaissance de base $[-m_E^-; m_E^+]$ sans que l'on puisse être plus précis. Sans plus d'effort, on est ramené à une modélisation de l'incertitude à l'aide d'intervalles, ce qui n'est pas sans rappeler la théorie des intervalles dont on a parlé dans le paragraphe 2.3. On avait vu que le principal inconvénient de ce type de modélisation était sa tendance à propager des incertitudes de façon surestimée.

C'est pourquoi les bornes m_E^- et m_E^+ ne sont pas forcément introduites comme des quantités déterministes : afin de limiter l'effet de surestimation, **on définit ces bornes comme des variables aléatoires caractérisées par des lois de probabilité données.** Pour plus de clarté, on distinguera la variable aléatoire α d'une réalisation donnée α .

4.1.2 Choix des lois de probabilité

Les méconnaissances de base peuvent suivre diverses lois de probabilité ; la nature de ces lois doit être définie *a priori*, en fonction d'hypothèses formulées sur les sources de dispersion qui se produisent dans la sous-structure concernée.

Une hypothèse simplificatrice consiste à supposer que l'on peut avoir des réalisations équiprobables des méconnaissances de base m_E^- et m_E^+ , comprises dans un certain domaine borné :

$$m_E^- \in [0; \bar{m}_E^-] \quad \text{et} \quad p^-(m_E^-) = Cte \quad (4.3a)$$

$$m_E^+ \in [0; \bar{m}_E^+] \quad \text{et} \quad p^+(m_E^+) = Cte \quad (4.3b)$$

où p^- et p^+ sont les densités de probabilité respectives des variables aléatoires m_E^- et m_E^+ , fonctions des valeurs $\bar{m}_E^-$ et $\bar{m}_E^+$.

Cette idée revient à suivre les recommandations du maximum d'entropie (1.12) qui permet de choisir une loi uniforme de probabilité quand les seules contraintes données sont d'avoir des réalisations bornées ; c'est en effet le cas des méconnaissances de base dont les réalisations sont positives, et de plus les réalisations m_E^- sont nécessairement limitées en amplitude, car chacune d'entre elles constitue une borne inférieure de l'énergie de déformation, forcément positive. Les réalisations m_E^+ , elles aussi positives, n'ont quant à elles pas de raison d'être majorées, mais il est logique de le supposer tout de même, par analogie avec m_E^- . Nous reparlerons plus précisément du cas de la loi uniforme dans le paragraphe 4.3.2.

Imposer des lois de probabilité aux bornes d'un intervalle, au lieu de les appliquer directement à des paramètres du modèle, pose des problèmes de description des phénomènes : en effet, **comment peut-on traduire la probabilité d'appartenance de l'énergie de déformation à un intervalle défini à partir d'une probabilité d'occurrence pour chacune de ses deux bornes ?** C'est pour cette raison que l'on est amené à introduire quelques outils mathématiques afin de définir une notion que l'on pourrait qualifier de « probabilité d'intervalle », dont la description est donnée dans le paragraphe 4.3. Ces outils sont en fait issus du raisonnement inverse : comment transcrire une loi de probabilité, servant à décrire l'incertitude sur un paramètre donné du modèle, sous la forme d'une méconnaissance de base ? C'est cette question que nous traiterons en premier, avec le cas particulier de la loi normale décrit dans le paragraphe 4.2. Enfin, une comparaison avec les méthodes non stochastiques de représentation des incertitudes, dans le paragraphe 4.3.3, nous permettra de constater la pertinence des outils introduits dans le cadre de la théorie des méconnaissances.

Revenons un instant sur l'idée de définir les bornes des intervalles de méconnaissance comme des variables aléatoires : on peut toutefois être en présence de certains cas extrêmes où la modélisation est (volontairement ou non) très imparfaite : c'est le cas, par exemple, d'une liaison au comportement non-linéaire qui est représentée par un modèle linéaire. Dans cette situation, le manque d'information est tel que l'on peut décider de ne pas associer de loi de probabilité aux méconnaissances de base : ceci équivaut alors à affirmer que la méconnaissance m est située quelque part dans un intervalle (déterministe) noté $[-\bar{m}_E^-; \bar{m}_E^+]$. On est ainsi revenu à une description de l'incertitude par intervalles telle que déjà définie auparavant dans le paragraphe 2.3.

4.2 Exemple : modélisation des dispersions matériaux à l'aide des méconnaissances

Dans le cas où l'on suppose que les seules sources d'erreurs sur la sous-structure E sont les incertitudes concernant le matériau, on peut raisonnablement considérer que l'on peut associer à la méconnaissance m , représentant l'énergie de déformation de la sous-structure E , une loi normale centrée de densité de probabilité $p(m)$:

$$p(m) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{m^2}{2\sigma^2}\right) \quad (4.4)$$

L'écart type σ de la loi normale peut être paramétré à l'aide de deux valeurs notées $\bar{m}_E^+$ et $\bar{m}_E^-$ en posant par exemple que l'on doit avoir :

$$\int_{-\bar{m}_E^-}^{\bar{m}_E^+} p(m) dm = 0,99 \quad (4.5)$$

4.2.1 Justification du choix d'une loi normale

Le choix d'une telle loi de probabilité est motivé par une application du théorème central limite (1.9) évoquée dans le paragraphe 1.2.2 sous la forme des conditions de Borel.

Ainsi, prenons l'exemple d'une poutre droite de longueur L dont le module d'Young est modélisé par un champ stochastique quelconque $E(x, \theta)$. On peut réécrire ce champ à l'aide de la décomposition de Karhunen-Loève (1.3) présentée dans le paragraphe 1.2.1 :

$$E(x, \theta) = \overline{E}(x) + \sum_{i=1}^{\infty} \xi_i(\theta) \sqrt{\lambda_i} \varphi_i(x) \quad (4.6)$$

où les quantités $\xi_i(\theta)$ sont des variables aléatoires non corrélées d'espérances nulles et d'écarts types unitaires.

L'énergie de déformation e_d de cette poutre dans un champ de déplacement quelconque est alors proportionnelle à la moyenne spatiale du champ stochastique le long de la poutre :

$$e_d \propto \int_0^L E(x, \theta) dx \quad (4.7)$$

En utilisant la décomposition (4.6), l'énergie de déformation peut alors s'écrire, en posant $\lambda_0 = 1$, $\xi_0(\theta) = 1$ et $\varphi_0(x) = \overline{E}(x)$, comme :

$$e_d \propto \int_0^L \sum_{i=0}^{\infty} \xi_i(\theta) \sqrt{\lambda_i} \varphi_i(x) dx \quad (4.8)$$

L'énergie de déformation est donc proportionnelle à une somme infinie de variables aléatoires $\eta(\theta)$:

$$e_d \propto \sum_{i=0}^{\infty} \eta(\theta) \quad (4.9)$$

Suivant la nature de ces variables $\eta(\theta)$, le choix d'une loi normale pour l'évolution de l'énergie de déformation est plus ou moins valide ; il l'est d'autant plus que le nombre de variables de « poids » équivalents est grand. Ainsi, en définissant les méconnaissances à une échelle structurale, on peut supposer des lois de probabilité simples pour les méconnaissances de base associées aux énergies de déformation par sous-structure.

On a vu néanmoins dans le paragraphe 1.2.2 que l'utilisation d'une loi de probabilité normale pouvait conduire à des réalisations non physiques : n'étant pas à support borné, elle rend possible l'occurrence d'énergies négatives ($-m_E^- < -1$). Parmi les solutions proposées pour résoudre cette difficulté, on choisit d'éliminer toutes les réalisations non physiques pour lesquelles on suppose tout de même qu'elles sont caractérisées par une fréquence d'apparition très faible ; ceci revient à modifier la loi de probabilité en imposant que la densité de probabilité doit être nulle par exemple en dehors de l'intervalle $[-\overline{m}_E^-; \overline{m}_E^+]$, quitte à la renormaliser ensuite.

4.2.2 Conversion de la loi normale en termes de méconnaissances

Maintenant que l'on a supposé que la quantité m , associée à l'énergie de déformation de la sous-structure E , vérifiait une loi de probabilité normale, il s'agit de préciser comment traduire cette dernière en termes de méconnaissances de base. Pour cela, il suffit de déterminer la probabilité d'observer m dans un intervalle $[-m_E^-; m_E^+]$ donné :

$$P(-m_E^- \leq m \leq m_E^+) = \int_{-m_E^-}^{m_E^+} p(m) dm \quad (4.10)$$

Puisque les méconnaissances de base m_E^+ et m_E^- sont définies dans (4.2) de part et d'autre du modèle théorique, l'événement précédent peut être décrit par deux événements indépendants :

- soit $m \in [0; m_E^+]$, ce qui signifie que l'on observe l'événement ($m_E^- = 0, m_E^+ \geq 0$) avec une probabilité $P^+(m_E^+)$:

$$P^+(m_E^+) = \int_0^{m_E^+} p(m) dm \quad (4.11)$$

- soit $m \in [-m_E^-; 0]$, c'est-à-dire que l'on constate l'événement ($m_E^- \geq 0, m_E^+ = 0$) avec une probabilité $P^-(m_E^-)$:

$$P^-(m_E^-) = \int_{-m_E^-}^0 p(m) dm \quad (4.12)$$

On a alors $P^+(\infty) + P^-(\infty) = 1$, et même $P^+(\infty) = P^-(\infty) = \frac{1}{2}$ dans le cas particulier d'une loi normale centrée ; cette situation est représentée sur la figure 4.1.

Ce cas de figure montre comment on peut modéliser une incertitude représentée par une loi de probabilité classique à l'aide de notre concept de méconnaissance. Toutefois, comme l'utilisation des deux expressions P^+ et P^- n'est guère commode en pratique, divers outils mathématiques décrits dans le paragraphe suivant sont introduits pour en simplifier l'usage.

4.3 Concept de probabilité d'intervalle

4.3.1 Intervalles standards et probabilité d'intervalle

Le but recherché est de remplacer le couple de quantités P^+ et P^- par une seule quantité stochastique. Considérons pour cela une famille d'intervalles de méconnaissances de base $[-m_E^-; m_E^+]$ de même longueur $L = m_E^+ + m_E^-$.

L'intervalle $[-m_E^-; m_E^+]$ est dit standard si pour une longueur L donnée, la probabilité d'avoir m contenu dans $[-m_E^-; m_E^+]$ est maximale sur l'ensemble des intervalles de longueur L :

$$I(L) = \arg \max_{\substack{[-m_E^-; m_E^+] \\ m_E^+ + m_E^- = L}} P^+(m_E^+) + P^-(m_E^-) \quad (4.13)$$

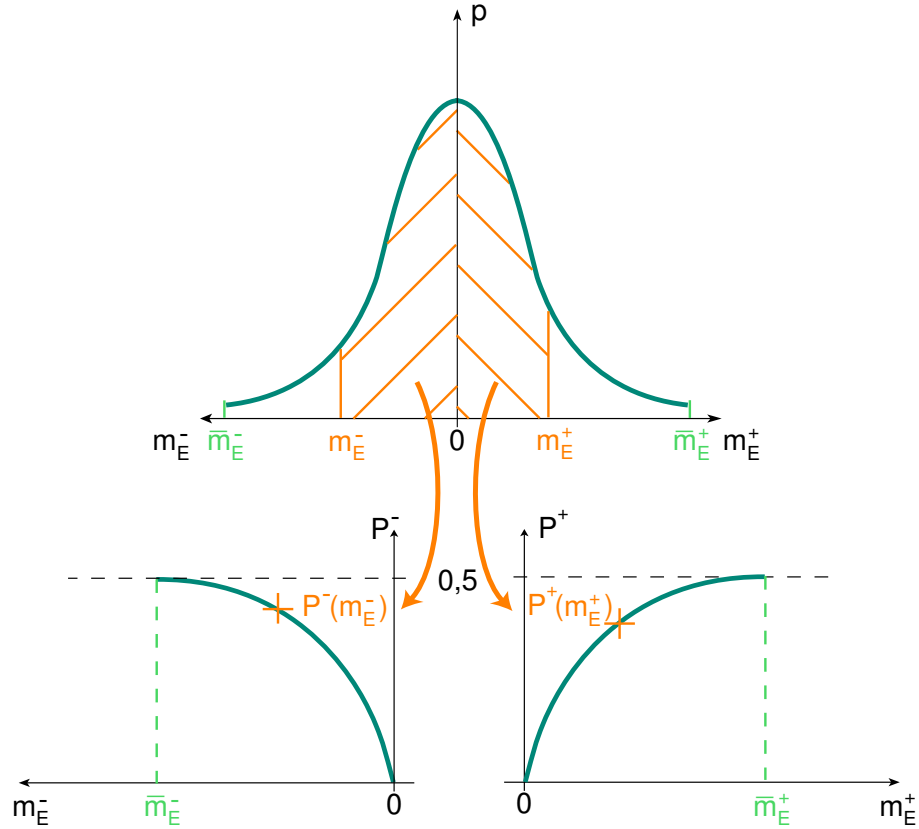


FIG. 4.1 – Méconnaissances de base associées à une loi normale centrée

À la suite de cette définition, on peut introduire la notion de probabilité d'intervalle, en posant que pour une longueur L donnée, cette probabilité $P(L)$ est la probabilité d'avoir m appartenant à l'intervalle standard $I(L)$:

$$P(L) = P(m \in I(L)) = \max_{\substack{[-m_E^-; m_E^+] \\ m_E^+ + m_E^- = L}} P^+(m_E^+) + P^-(m_E^-) \quad (4.14)$$

Quelques propriétés à propos des intervalles standards peuvent alors être démontrées.

Tout d'abord, les bornes d'un intervalle standard $I(L) = [-m_E^-; m_E^+]$, pour une longueur L donnée, vérifient nécessairement :

$$p(m_E^+) = p(-m_E^-) \quad (4.15)$$

En effet, ces bornes vérifient la relation suivante, issue de la recherche du maximum de $P^+(m_E^+) + P^-(m_E^-)$ sous la contrainte $m_E^- = L - m_E^+$:

$$\frac{\partial P^+}{\partial m_E^+}(m_E^+) - \frac{\partial P^-}{\partial m_E^+}(L - m_E^+) = 0 \quad (4.16)$$

avec :

$$\frac{\partial P^+}{\partial m_E^+}(m_E^+) = \frac{\partial}{\partial m_E^+} \int_0^{m_E^+} p(m) dm = p(m_E^+) \quad (4.17)$$

$$\begin{aligned} \frac{\partial P^-}{\partial m_E^+}(L - m_E^+) &= \frac{\partial}{\partial m_E^+} \int_{-(L-m_E^+)}^0 p(m) dm = \frac{\partial}{\partial m_E^+} \int_0^{L-m_E^+} p(-m) dm \\ &= p(-(L - m_E^+)) = p(-m_E^-) \end{aligned} \quad (4.18)$$

ce qui prouve la propriété énoncée.

Dans la situation courante où les fonctions $m_E^+ \mapsto p(m_E^+)$ et $m_E^- \mapsto p(-m_E^-)$ sont strictement décroissantes, l'ensemble des intervalles standards vérifie une relation d'inclusion :

$$\forall L \geq L' \quad I(L') \subset I(L) \quad (4.19)$$

Pour démontrer cette relation, prenons $L \geq L'$ donnés, et les intervalles standards associés $I(L) = [-m_E^-(L); m_E^+(L)]$ et $I(L') = [-m_E^-(L'); m_E^+(L')]$ respectivement.

D'après la proposition que l'on vient juste de démontrer, on peut écrire les relations suivantes :

$$p(m_E^+(L)) = p(-m_E^-(L)) \quad \text{et} \quad p(m_E^+(L')) = p(-m_E^-(L')) \quad (4.20)$$

Raisonnons maintenant par l'absurde : on suppose que $m_E^+(L) < m_E^+(L')$, ce qui entraîne nécessairement d'une part que :

$$p(m_E^+(L)) > p(m_E^+(L')) \quad (4.21)$$

D'autre part, l'hypothèse implique aussi que :

$$-m_E^-(L) = m_E^+(L) - L < m_E^+(L') - L < m_E^+(L') - L' = -m_E^-(L') \quad (4.22)$$

et donc que :

$$p(-m_E^-(L)) < p(-m_E^-(L')) \quad (4.23)$$

En combinant les équations (4.21) et (4.23) avec la relation (4.20), on arrive au résultat suivant :

$$p(m_E^+(L)) > p(m_E^+(L')) = p(-m_E^-(L')) > p(-m_E^-(L)) = p(m_E^+(L)) \quad (4.24)$$

ce qui est bien sûr contradictoire, donc l'hypothèse de départ est fausse, ce qui démontre que $m_E^+(L') \leq m_E^+(L)$. Par un raisonnement analogue, on peut montrer également que $-m_E^-(L) \leq -m_E^-(L')$, ce qui prouve l'inclusion.

4.3.2 Interprétation

Une interprétation possible de ces définitions est la suivante : si l'on cherche un intervalle de méconnaissance de base tel que la quantité m ait une probabilité P donnée d'être à l'intérieur, on considère l'intervalle standard $I(L_P)$ tel que la probabilité d'intervalle associée $P(L_P)$ est égale à P . En fait, cet intervalle $I(L_P)$ est le plus petit intervalle $[-m_E^-; m_E^+]$ tel que $P^+(m_E^+) + P^-(m_E^-) = P$:

$$I(L_P) = \arg \min_{\substack{[-m_E^-, m_E^+] \\ P^+(m_E^+) + P^-(m_E^-) = P}} m_E^+ + m_E^- \quad (4.25)$$

Ceci est une conséquence de la propriété (4.19) qui n'est rigoureusement vraie que si les fonctions $m_E^+ \mapsto p(m_E^+)$ et $m_E^- \mapsto p(-m_E^-)$ sont strictement décroissantes ; ainsi, si l'on est en présence d'une loi normale non centrée, cette interprétation ne reste valide que pour des longueurs $L \geq 2|\bar{m}_E^+ - \bar{m}_E^-|$. La figure 4.2 montre une illustration de cette situation : on a représenté l'intervalle standard ainsi que la probabilité d'intervalle égale à l'aire hachurée.

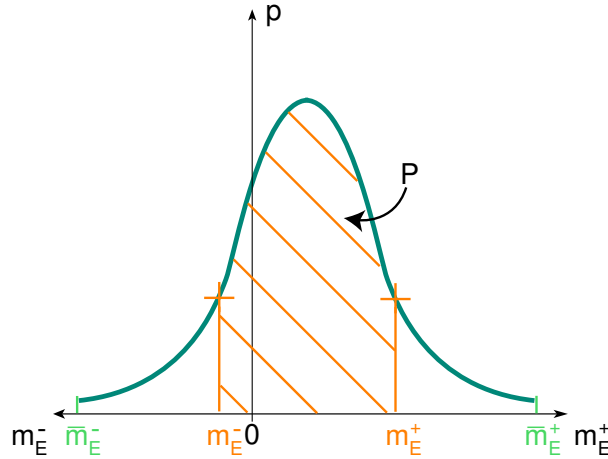


FIG. 4.2 – Intervalle standard et probabilité d'intervalle pour une loi normale non centrée

Cette interprétation reste bien évidemment vraie même s'il n'y a pas de loi de probabilité $p(m)$ définissable directement sur la quantité m ; on détermine simplement les intervalles standards et les probabilités d'intervalle associées à l'aide des quantités $P^+(m_E^+)$ et $P^-(m_E^-)$. De plus, certains cas ne seraient pas « traduisibles » sous la forme d'une densité de probabilité d'une variable aléatoire ; il suffit en effet d'envisager des quantités $P^+(m_E^+)$ et $P^-(m_E^-)$ telles qu'il existe un domaine non réduit à un singleton où ces quantités sont nulles :

$$\exists \mathcal{M}_E^- \text{ tel que } \forall m_E^- \in [0; \mathcal{M}_E^-] \quad P^-(m_E^-) = 0 \quad (4.26a)$$

$$\exists \mathcal{M}_E^+ \text{ tel que } \forall m_E^+ \in [0; \mathcal{M}_E^+] \quad P^+(m_E^+) = 0 \quad (4.26b)$$

En d'autres termes, il existe un intervalle minimal $[-\mathcal{M}_E^-; \mathcal{M}_E^+]$ à propos duquel on ne sait rien, puisqu'aucune probabilité ne peut lui être associée, et qui constitue l'intersection commune de tous les intervalles standards. En cela, cet intervalle d'« ignorance maximale » répond bien au schéma de non-spécificité, évoqué dans le paragraphe 2.1, de description des incertitudes. Ainsi, dans le cas général, on se trouve en présence de la situation décrite par la figure 4.3.

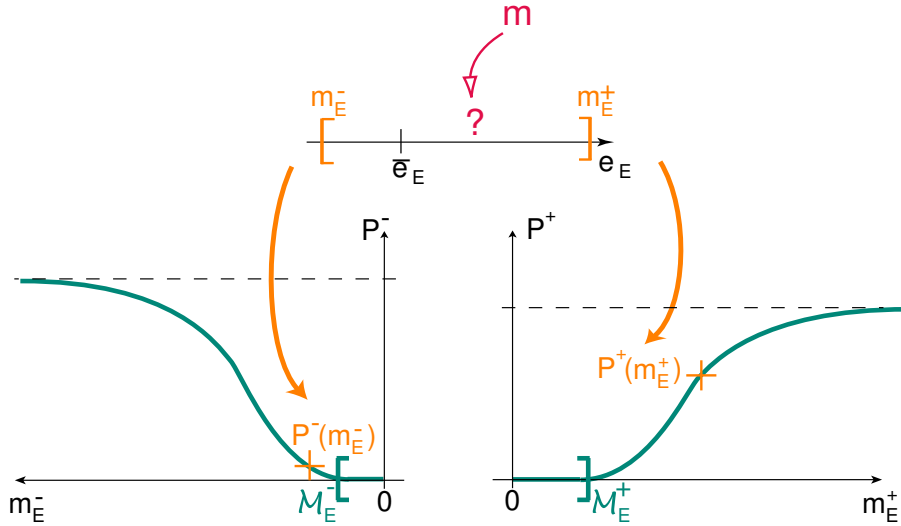


FIG. 4.3 – Cas général de modélisation des méconnaissances de base

Après cette généralisation, revenons un instant sur le cas particulier de la loi uniforme, juste évoqué dans le paragraphe 4.1.2. Dans cette situation, et avec les notations de (4.3), les probabilités respectives d'occurrence de la méconnaissance sont :

$$P^-(m_E^-) = \frac{m_E^-}{\overline{m}_E^- + \overline{m}_E^+} \quad (4.27a)$$

$$P^+(m_E^+) = \frac{m_E^+}{\overline{m}_E^- + \overline{m}_E^+} \quad (4.27b)$$

ce qui permet d'assurer $P^+(\infty) + P^-(\infty) = 1$. Ces probabilités sont en fait les mêmes que celles que l'on obtiendrait si l'on supposait une répartition uniforme directement sur la quantité m , dans le domaine $[-\overline{m}_E^-; \overline{m}_E^+]$:

$$p(m) = \frac{1}{\overline{m}_E^+ + \overline{m}_E^-} \quad (4.28)$$

Cette équivalence, illustrée sur la figure 4.4, montre bien l'efficacité de la représentation à l'aide du concept de méconnaissance.

Notons également que dans le cas uniforme, il n'y a pas unicité de l'intervalle standard pour une longueur donnée L . Ce qui importe est de pouvoir considérer une famille d'intervalles standards définie par une seule variable stochastique ; il en résulte une simplification notable dans l'analyse.

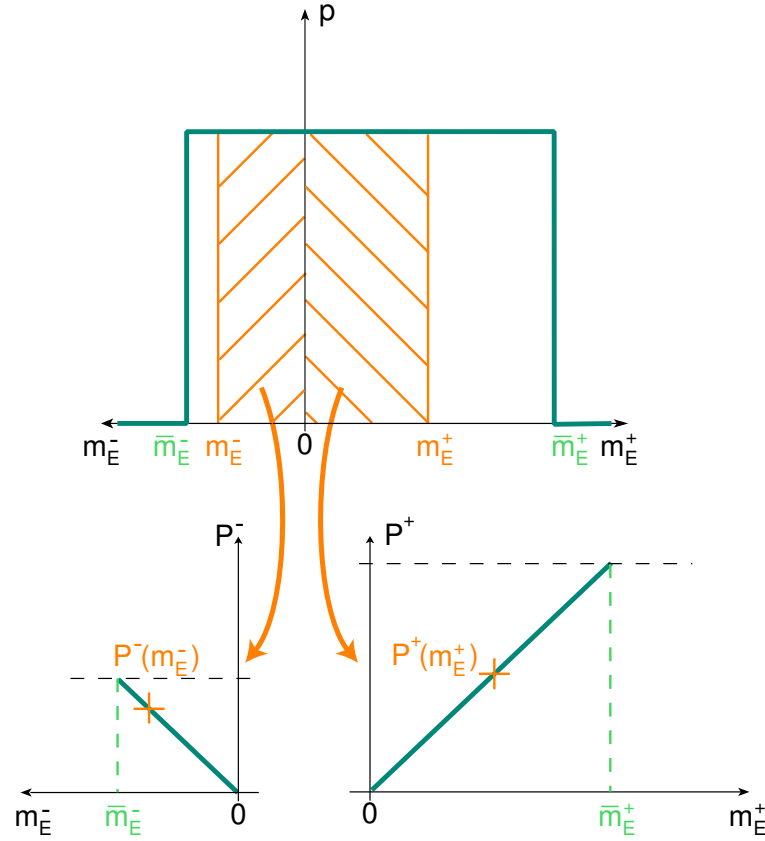


FIG. 4.4 – Méconnaissances de base modélisées par une loi uniforme

4.3.3 Comparaison avec les méthodes non stochastiques de description des incertitudes

Le concept de méconnaissance, qui met en jeu des intervalles dont les bornes sont des variables aléatoires, n'est pas sans rappeler certaines méthodes non stochastiques décrites dans le chapitre 2.

En particulier, les ensembles aléatoires selon Dempster-Shafer, présentés dans le paragraphe 2.5, correspondent à ce schéma. Pour simplifier la comparaison, on ne regarde pour le moment que le cas où l'on est en présence, dans notre théorie, d'intervalles du type $[0; m_E^+]$, de probabilités d'occurrence $P^+(m_E^+)$. Le nombre d'intervalles de ce type n'est pas fini, aussi décide-t-on de « discrétiser » cette donnée, à des fins de comparaison avec les ensembles aléatoires. Par exemple, avec quatre points tels que :

$$P^+(m_i^+) = \frac{i}{4} \quad \forall i = 1, \dots, 4 \quad (4.29)$$

on peut définir un intervalle aléatoire $\mathcal{I}$ caractérisé par les éléments focaux et affectations probabilistes suivants :

$$I_i = [0; m_i^+] \quad \text{et} \quad m(I_i) = \frac{1}{4} \quad \forall i = 1, \dots, 4 \quad (4.30)$$

La probabilité d'avoir m compris dans un intervalle $J \subset \mathbb{R}^+$ est donnée par la relation de Dempster (2.40) :

$$\sum_{I_i \subset J} m(I_i) \leq P(J) \leq \sum_{I_i \cap J \neq \emptyset} m(I_i) \quad (4.31)$$

Prenons le cas particulier où $J = I_k$, avec k donné ; on obtient alors :

$$\frac{k}{4} = \sum_{I_i \subset I_k} m(I_i) \leq P(I_k) \leq \sum_{I_i \cap I_k \neq \emptyset} m(I_i) = 1 \quad (4.32)$$

La probabilité d'être dans l'intervalle I_k est donc d'au moins $\frac{k}{4} = P^+(m_k^+)$: la relation de Dempster nous permet donc de constater simplement que la probabilité $P^+(m_E^+)$ est en fait le meilleur minorant de la probabilité de trouver la méconnaissance m effectivement dans l'intervalle $[0; m_E^+]$. En associant les deux événements indépendants, cela revient également à dire que la probabilité $P^+(m_E^+) + P^-(m_E^-)$ est le meilleur minorant de la probabilité d'avoir m dans $[-m_E^-; m_E^+]$.

Ceci est cohérent avec la façon dont nous avons traduit la densité de probabilité normale $p(m)$ en termes de méconnaissances de base : en posant que :

$$P^+(m_E^+) = \int_0^{m_E^+} p(m) dm \quad (4.33)$$

et en considérant ensuite que la probabilité d'avoir m dans $[0; m_E^+]$ était effectivement $P^+(m_E^+)$, on a fait l'hypothèse implicite de considérer le cas le plus défavorable, à savoir de choisir le meilleur minorant de la probabilité vraie, ce qui joue dans le sens de la sécurité.

Notons enfin que comme l'intervalle aléatoire $\mathcal{I}$ équivalent à notre façon de modéliser les méconnaissances de base est consonant, c'est-à-dire tel que $I_i \subset I_{i+1}$, il peut être représenté par un intervalle flou $\tilde{I}$ dont les caractéristiques sont liées à celles de $\mathcal{I}$, comme précisé dans le paragraphe 2.5.2. Ainsi, les méthodes associées au traitement des ensembles flous pourraient également être utilisées dans notre théorie.

4.4 Obtention pratique et manipulation des méconnaissances de base

Principe

Si la notion de probabilité d'intervalle est commode pour traduire les méconnaissances de base, l'obtention pratique de ces dernières se fait nécessairement au travers des deux probabilités $P^+(m_E^+)$ et $P^-(m_E^-)$. En effet, on suppose que l'on est en présence de deux événements indépendants : m appartient à un intervalle qui est soit du type $[0; m_E^+]$, soit du type $[-m_E^-; 0]$. C'est donc ainsi que doivent être obtenues les méconnaissances.

Reprenons le cas de la loi normale : les fonctions de répartition F^- et F^+ des variables stochastiques m_E^- et m_E^+ sont liées aux probabilités $P^+(m_E^+)$ et $P^-(m_E^-)$:

$$F^-(m_E^-) = P(\textcolor{brown}{m}_E^- \leq m_E^-) = P^-(m_E^-) + P^+(\infty) \quad (4.34a)$$

$$F^+(m_E^+) = P(\textcolor{brown}{m}_E^+ \leq m_E^+) = P^+(m_E^+) + P^-(\infty) \quad (4.34b)$$

En effet, comme $P^-(\infty)$ est la probabilité d'avoir un intervalle du type $[-m_E^-; 0]$, c'est aussi la probabilité d'avoir $\textcolor{brown}{m}_E^+$ égal à zéro, et c'est ce terme qui permet d'avoir une fonction de répartition normalisée, c'est-à-dire telle que $F^+(\infty) = 1$.

En termes de densités de probabilité, c'est à l'aide de distributions de Dirac δ qu'on peut les exprimer :

$$p^-(m_E^-) = \left\langle \frac{1}{\sqrt{2\pi}\sigma^2} \exp\left(-\frac{m_E^{-2}}{2\sigma^2}\right) \right\rangle_- + P^+(\infty)\delta(m_E^-) \quad (4.35a)$$

$$p^+(m_E^+) = \left\langle \frac{1}{\sqrt{2\pi}\sigma^2} \exp\left(-\frac{m_E^{+2}}{2\sigma^2}\right) \right\rangle_+ + P^-(\infty)\delta(m_E^+) \quad (4.35b)$$

où $x \mapsto \langle x \rangle_-$ et $x \mapsto \langle x \rangle_+$ sont les fonctions caractéristiques de $\mathbb{R}^-$ et $\mathbb{R}^+$ respectivement. On a représenté sur la figure 4.5 les densités de probabilité pour une loi normale non centrée ; toutes les deux sont bien normalisées dans la mesure où, à chaque fois, la somme de l'aire hachurée et du poids de la distribution de Dirac est égale à un.

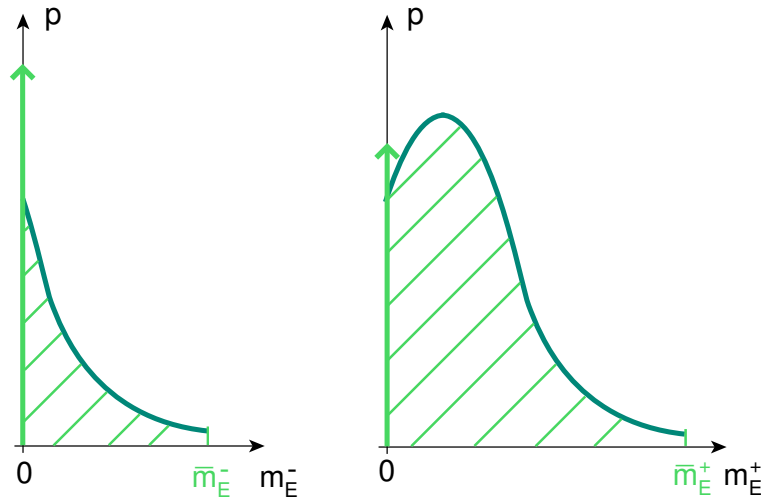


FIG. 4.5 – Densités de probabilité associées à une loi normale non centrée

En anticipant un peu sur le chapitre suivant, où il sera question de la propagation des méconnaissances au travers du modèle de la structure étudiée, nous indiquons plus bas comment sont obtenues en pratique les méconnaissances, et comment on peut envisager leur manipulation ultérieure.

Méthode de Monte Carlo

La première façon consiste à utiliser une technique de Monte Carlo avec les densités de probabilité retenues, puis à assurer la propagation des intervalles de méconnaissances tirage par tirage.

Si l'on prend l'exemple de la loi normale, obtenir des tirages selon les lois (4.35) se fait comme suit :

- on réalise un tirage m à partir de la loi de probabilité « complète », c'est-à-dire la densité de probabilité normale initiale (4.4) ;
- en fonction du signe de m , on est en présence de l'un des deux événements décrits précédemment : si m est positif, on forme l'intervalle $[-m_E^-; m_E^+] = [0; m]$; dans le cas contraire, c'est l'intervalle $[-m_E^-; m_E^+] = [-m; 0]$ que l'on obtient ;
- les réalisations sont alors propagées à travers le modèle selon le principe exposé dans le paragraphe 5.1 ; une étude statistique permet ensuite de déduire les quantités que l'on recherche.

Le cas de la loi de probabilité uniforme se traite de la même façon.

En ce qui concerne la situation extrême où aucune loi de probabilité n'est supposée, le démarche est encore plus simple, puisqu'il suffit de supposer pour la sous-structure concernée l'intervalle déterministe $[-\bar{m}_E^-; \bar{m}_E^+]$, et de le combiner avec les tirages de méconnaissances de base des autres sous-structures.

Fonctions caractéristiques

À toute variable aléatoire X , on peut associer une fonction Φ appelée fonction caractéristique et définie par :

$$\begin{aligned} \Phi : \mathbb{R} &\rightarrow \mathbb{C} \\ u &\mapsto \Phi(u) = \int_{-\infty}^{+\infty} p(x) \exp(iux) dx \end{aligned} \quad (4.36)$$

quand la variable aléatoire est caractérisée par sa densité de probabilité $p(x)$; la fonction caractéristique de X est en fait liée à la transformée de Fourier de sa densité de probabilité $p(x)$. La dénomination de fonction caractéristique, à ne pas confondre avec celle d'un intervalle, est tout de même appropriée, puisque deux variables aléatoires X et Y ont même loi de probabilité si et seulement si elles ont même fonction caractéristique.

Cette quantité s'avère très utile dans le calcul de combinaisons linéaires de variables aléatoires indépendantes, associé, comme on le verra dans le chapitre suivant, à la propagation des méconnaissances de base à travers le modèle.

Si Z est la somme de deux variables aléatoires X et Y indépendantes, de densités de probabilité respectives $p_X(x)$ et $p_Y(y)$, la fonction de répartition de Z est donnée par :

$$F(z) = \int_{-\infty}^{+\infty} \int_{-\infty}^{z-x} p_X(x) p_Y(y) dy dx = \int_{-\infty}^{+\infty} \int_{-\infty}^z p_X(x) p_Y(z' - x) dz' dx \quad (4.37)$$

Il s'agit en fait de l'aire de la partie du plan telle que $x + y < z$, après avoir exprimé la dépendance de x et y vis-à-vis de z ; ceci est illustré sur la figure 4.6.

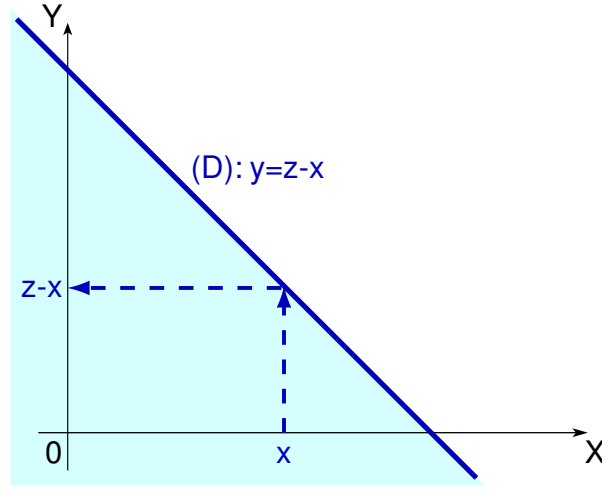


FIG. 4.6 – Interprétation de la somme de deux variables aléatoires

Comme la densité de probabilité est la dérivée de la fonction de répartition, Z a pour densité de probabilité :

$$p_Z(z) = \int_{-\infty}^{+\infty} p_X(x)p_Y(z-x)dx \quad (4.38)$$

ou encore :

$$p_Z = p_X * p_Y = p_Y * p_X \quad (4.39)$$

où $*$ désigne le produit de convolution.

L'utilisation des fonctions caractéristiques se révèle alors très pratique, puisque la transformée de Fourier du produit de convolution est le produit des transformées de Fourier :

$$\Phi_{X+Y} = \Phi_X \Phi_Y \quad (4.40)$$

De plus, avec la propriété $\Phi_{aX}(u) = \Phi_X(au)$, où a est une constante scalaire donnée, on déduit très facilement l'expression de la fonction caractéristique d'une combinaison linéaire de variables aléatoires X_i :

$$\Phi_{\sum_{i=1}^n a_i X_i}(u) = \prod_{i=1}^n \Phi_{X_i}(a_i u) \quad (4.41)$$

Comme on le verra dans le paragraphe 5.1, la propagation des méconnaissances de base à travers le modèle revient en fait à prendre des combinaisons linéaires de ces dernières ; l'emploi des fonctions caractéristiques permet d'obtenir les densités de probabilité correspondantes avec plus d'efficacité que la technique de Monte Carlo, qui, dans sa version de base, nécessite un nombre minimal de tirages pour avoir une convergence satisfaisante.

En pratique, les fonctions caractéristiques sont obtenues à l'aide de la transformée de Fourier rapide (*FFT*) implantée dans Matlab ; il faut bien entendu mettre en œuvre un pas de discrétisation suffisamment fin pour éviter les difficultés liées au phénomène de repliement de spectre. Ces problèmes bien identifiés dans le domaine du traitement du signal sont faciles à résoudre.

Propagation des méconnaissances de base

La connaissance progresse en intégrant en elle l'incertitude, non en l'exorcisant.

Edgar Morin

Sommaire

5.1 Principe	78
5.2 Calcul théorique des méconnaissances effectives	80
5.2.1 Méconnaissance effective sur une pulsation propre	80
5.2.2 Méconnaissance effective sur la valeur en un DDL d'un mode	81
5.2.3 Méconnaissance effective sur la projection d'un mode	82
5.3 Calcul de méconnaissances effectives sur un exemple simple	83
5.3.1 Premier choix de modèle de méconnaissances de base	83
5.3.2 Deuxième choix de modèle de méconnaissances de base	88
5.3.3 Troisième choix de modèle de méconnaissances de base	89
5.4 Prise en compte d'une forte méconnaissance	91
5.4.1 Principe et mise en œuvre	91
5.4.2 Application sur un exemple simple	93

L'objet de ce cinquième chapitre est le calcul de la méconnaissance sur une quantité d'intérêt définie sur l'ensemble de la structure, telle qu'une pulsation propre ou une forme propre de vibration. C'est à partir de la donnée du modèle de méconnaissances de base que l'on est capable de déterminer deux valeurs encadrant la quantité d'intérêt, par propagation des intervalles définis par les méconnaissances de base. Le calcul de ces valeurs est réalisé simplement à l'aide d'approximations au premier ordre, quand on est en présence de faibles valeurs de méconnaissances de base sur les différentes sous-structures. Néanmoins, le cas d'une « forte » méconnaissance est également étudié.

5.1 Principe

Le but de la démarche est d'établir une procédure permettant de calculer des valeurs caractéristiques de la dispersion d'une quantité d'intérêt quelconque, à partir d'un modèle de méconnaissances de base donné ; autrement dit, on souhaite grâce au concept de méconnaissance obtenir des informations appréciables sur la famille de structures que l'on étudie.

Le point de départ est la donnée du modèle de méconnaissances de base $(m_E^-, m_E^+)_{E \in \Omega}$ sur les différentes sous-structures E de la structure étudiée Ω . Supposons que l'on étudie une certaine quantité d'intérêt α concernant l'ensemble de la structure. On pose :

$$\Delta\alpha_{\text{mod}} = \alpha_{\text{mod}} - \bar{\alpha} \quad (5.1)$$

où $\bar{\alpha}$ est la valeur de la quantité d'intérêt pour le modèle théorique déterministe, et α_{mod} celle associée au modèle avec méconnaissances. Compte tenu de la façon dont les méconnaissances de base ont été définies, ce n'est pas directement cette quantité $\Delta\alpha_{\text{mod}}$ que l'on souhaite déterminer.

À celle-ci, on peut associer grâce au modèle de méconnaissances deux bornes $\Delta\alpha_{\text{mod}}^+$ et $\Delta\alpha_{\text{mod}}^-$ qui l'encadrent : il s'agit simplement, comme on le montrera dans le paragraphe 5.2, de la propagation au travers du modèle des intervalles de méconnaissances de base $([-m_E^-; m_E^+])_{E \in \Omega}$.

Comme on connaît les lois de probabilité définissant les méconnaissances de base, on est capable de déterminer les lois que vérifient ces bornes $\Delta\alpha_{\text{mod}}^+$ et $\Delta\alpha_{\text{mod}}^-$: que l'on ait adopté pour cela une technique de Monte Carlo ou un calcul numérique des fonctions caractéristiques, comme précisé dans le paragraphe 4.4, on obtient, indirectement dans le premier cas, et directement dans le second, les densités de probabilité des bornes $\Delta\alpha_{\text{mod}}^+$ et $\Delta\alpha_{\text{mod}}^-$ encadrant la quantité $\Delta\alpha_{\text{mod}}$, et donc les probabilités respectives d'occurrence associées aux deux bornes $P_\alpha^-(\Delta\alpha_{\text{mod}}^-)$ et $P_\alpha^+(\Delta\alpha_{\text{mod}}^+)$.

De la même façon que pour les méconnaissances de base dans le paragraphe 4.3, on associe pour plus de commodité à ces deux quantités P_α^- et P_α^+ la probabilité d'intervalle $P_\alpha(L)$ qui caractérise la probabilité d'avoir la quantité $\Delta\alpha_{\text{mod}}$ calculée par le modèle au

sein de l'intervalle standard $I_{\Delta\alpha_{\text{mod}}}(L)$ de longueur L donnée :

$$P(\Delta\alpha_{\text{mod}} \in I_{\Delta\alpha_{\text{mod}}}(L)) = P_{\alpha}(L) \quad \forall L \quad (5.2)$$

Encore une fois, cela revient à dire que, pour une probabilité P donnée, on considère l'intervalle standard $I_{\Delta\alpha_{\text{mod}}}(L_P)$ qui est tel que $P(\Delta\alpha_{\text{mod}} \in I_{\Delta\alpha_{\text{mod}}}(L_P)) = P$; la figure 5.1 résume l'ensemble du principe de la propagation des méconnaissances de base.

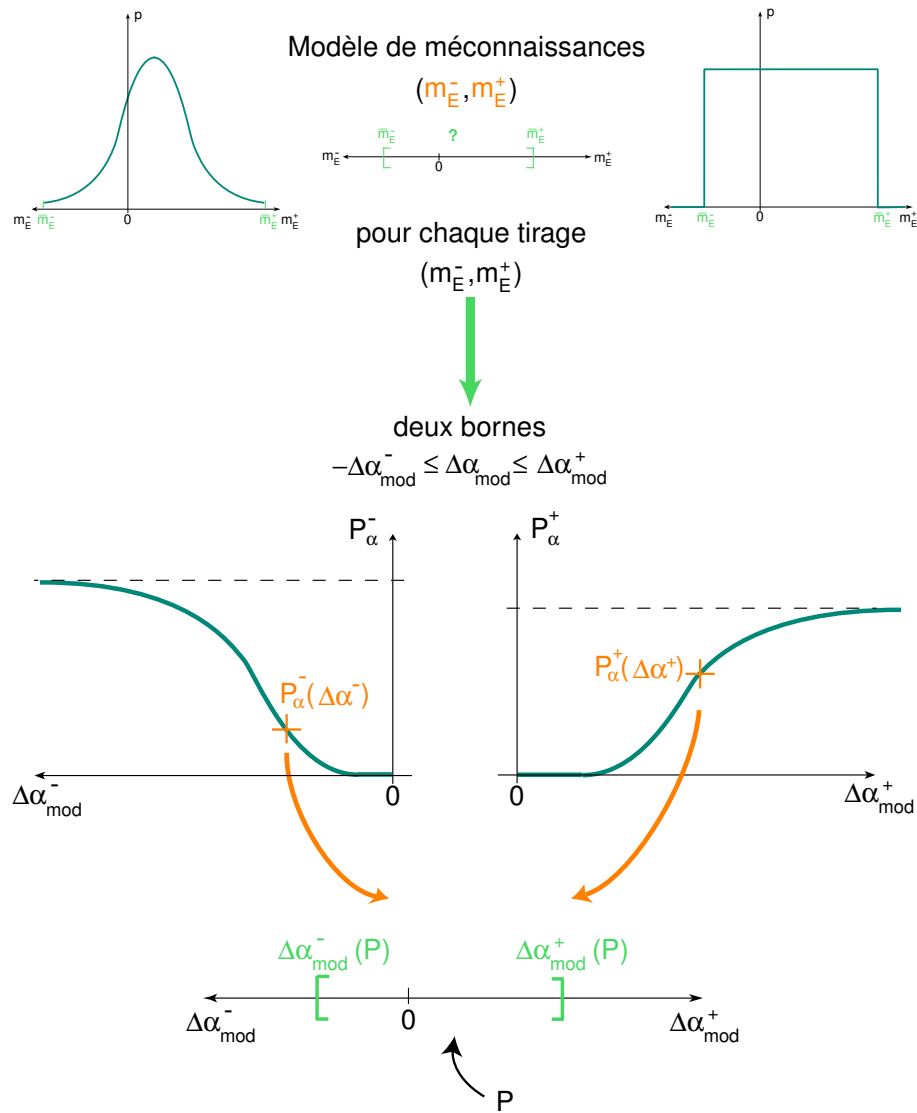


FIG. 5.1 – Calcul des méconnaissances effectives sur la quantité d'intérêt $\Delta\alpha_{\text{mod}}$

Les deux bornes de cet intervalle $I_{\Delta\alpha_{\text{mod}}}(L_P)$, notées $\Delta\alpha_{\text{mod}}^-(P)$ et $\Delta\alpha_{\text{mod}}^+(P)$, constituent ce que l'on appelle la *méconnaissance effective* sur la quantité d'intérêt α .

5.2 Calcul théorique des méconnaissances effectives

Le calcul de méconnaissances effectives concerne des quantités α définies sur l'ensemble de la structure : comme notre propos se limite à l'étude des vibrations libres d'un système mécanique, les quantités prises en considération sont des pulsations propres ω_i et des formes propres de vibration $\underline{\phi}_i$, qui sont respectivement les valeurs propres et les vecteurs propres du problème aux valeurs propres suivant :

$$(\mathbf{K} - \omega_i^2 \mathbf{M}) \underline{\phi}_i = \underline{0} \quad (5.3)$$

De la même façon que nous avons déjà écrit l'énergie de déformation du modèle théorique déterministe sous la forme $\bar{e}_E$, toutes les quantités qui sont surmontées d'un trait horizontal, $\bar{\phi}_i$ par exemple, concernent ce modèle déterministe. Toutefois, cette notation ne veut en aucun cas signifier que ce modèle est un modèle moyen.

On rappelle enfin que le couple de valeurs $\Delta\alpha_{\text{mod}}^-(P)$ et $\Delta\alpha_{\text{mod}}^+(P)$ constitue ce qu'on appelle la *méconnaissance effective* liée à la quantité $\Delta\alpha_{\text{mod}}$; nous allons maintenant voir comment obtenir ces valeurs pour les quantités précisées ci-dessus.

5.2.1 Méconnaissance effective sur une pulsation propre

La première quantité qui nous intéresse est la pulsation propre ω_i ; si les modes $\bar{\phi}_i$ du modèle théorique déterministe (associés respectivement aux pulsations propres $\bar{\omega}_i$) sont normalisés par rapport à la matrice de masse, la différence entre le carré d'une pulsation propre du modèle avec méconnaissances et du modèle déterministe $\Delta\omega_{i\text{mod}}^2 = \omega_{i\text{mod}}^2 - \bar{\omega}_i^2$ s'exprime avec une approximation du premier ordre comme suit :

$$\begin{aligned} \Delta\omega_{i\text{mod}}^2 &= \underline{\phi}_i^T \mathbf{K} \underline{\phi}_i - \bar{\phi}_i^T \bar{\mathbf{K}} \bar{\phi}_i \\ &\simeq \bar{\phi}_i^T (\mathbf{K} - \bar{\mathbf{K}}) \bar{\phi}_i = 2 \sum_{E \in \Omega} \left(e_E(\bar{\phi}_i) - \bar{e}_E(\bar{\phi}_i) \right) \end{aligned} \quad (5.4)$$

La relation (4.2) nous permet alors de propager les intervalles de méconnaissances $([-m_E^-; m_E^+])_{E \in \Omega}$, où pour chaque sous-structure E , (m_E^-, m_E^+) est une réalisation donnée des méconnaissances de base selon les lois de probabilité choisies ; cette propagation se fait de la façon suivante :

$$-\Delta\omega_{i\text{mod}}^{2-} \leq \Delta\omega_{i\text{mod}}^2 \leq \Delta\omega_{i\text{mod}}^{2+} \quad (5.5)$$

avec :

$$\Delta\omega_{i\text{mod}}^{2-} = 2 \sum_{E \in \Omega} m_E^- \bar{e}_E(\bar{\phi}_i) \quad (5.6a)$$

$$\Delta\omega_{i\text{mod}}^{2+} = 2 \sum_{E \in \Omega} m_E^+ \bar{e}_E(\bar{\phi}_i) \quad (5.6b)$$

Connaissant les lois de probabilité sur les méconnaissances de base, on est capable d'obtenir les dispersions des bornes $\Delta\omega_{i\text{mod}}^{2-}$ et $\Delta\omega_{i\text{mod}}^{2+}$ qui encadrent la quantité d'intérêt $\Delta\omega_{i\text{mod}}^2$. On peut donc déterminer, pour une valeur de probabilité P donnée, les deux bornes $\Delta\omega_{i\text{mod}}^{2-}(P)$ et $\Delta\omega_{i\text{mod}}^{2+}(P)$ de l'intervalle standard $I_{\Delta\omega_{i\text{mod}}^2}(L_P)$ associé, ou autrement dit la méconnaissance effective sur le carré de la fréquence propre ω_i^2 .

Après les pulsations propres, les modes propres sont des quantités d'intérêt pour lesquelles un effort supplémentaire doit être fourni ; en effet, ce sont des quantités vectorielles, or les méconnaissances effectives sont des scalaires ; c'est pourquoi on cherche à extraire des grandeurs scalaires de ces vecteurs. Deux possibilités s'offrent à nous :

- soit on considère la valeur du mode propre en un degré de liberté donné ;
- soit on projette ce mode propre dans une direction donnée.

Ces deux possibilités sont successivement étudiées dans les deux paragraphes suivants.

5.2.2 Méconnaissance effective sur la valeur en un degré de liberté d'un mode propre

On s'intéresse ici à la différence de valeur en un degré de liberté (DDL) particulier entre un mode propre du modèle avec méconnaissances et le mode propre correspondant du modèle déterministe, c'est-à-dire à la quantité $\Delta\phi_{ki\text{mod}} = \phi_{ki\text{mod}} - \bar{\phi}_{ki}$, où k est l'indice de la k -ième ligne du i -ème mode propre.

Pour des réalisations $(m_E^-, m_E^+)_{E \in \Omega}$ de faibles valeurs, on peut approcher au premier ordre cette variation par :

$$\Delta\phi_{ki\text{mod}} \simeq \underline{U}^T \Delta \mathbf{K} \bar{\phi}_i = \sum_{E \in \Omega} \underline{U}^T (\mathbf{K}_E - \bar{\mathbf{K}}_E) \bar{\phi}_i \quad (5.7)$$

où $\underline{U}$ est un vecteur donné dont nous détaillons le calcul dans l'annexe A ; notons également que dans cette expression, nous avons cette fois-ci normalisé les modes par rapport à leur maximum.

En utilisant l'égalité suivante faisant intervenir deux énergies de déformation différentes $e_E(\underline{U} + \bar{\phi}_i)$ et $e_E(\underline{U} - \bar{\phi}_i)$:

$$\underline{U}^T \mathbf{K}_E \bar{\phi}_i = \frac{1}{2} e_E(\underline{U} + \bar{\phi}_i) - \frac{1}{2} e_E(\underline{U} - \bar{\phi}_i) \quad (5.8)$$

et la relation fondamentale (4.2), on peut propager les intervalles de méconnaissances de base de la façon suivante :

$$-\Delta\phi_{ki\text{mod}}^- \leq \Delta\phi_{ki\text{mod}} \leq \Delta\phi_{ki\text{mod}}^+ \quad (5.9)$$

avec :

$$\Delta\phi_{ki\text{ mod}}^- = \frac{1}{2} \sum_{E \in \Omega} \left\{ m_E^- \bar{e}_E(\underline{U} + \bar{\phi}_i) + m_E^+ \bar{e}_E(\underline{U} - \bar{\phi}_i) \right\} \quad (5.10a)$$

$$\Delta\phi_{ki\text{ mod}}^+ = \frac{1}{2} \sum_{E \in \Omega} \left\{ m_E^+ \bar{e}_E(\underline{U} + \bar{\phi}_i) + m_E^- \bar{e}_E(\underline{U} - \bar{\phi}_i) \right\} \quad (5.10b)$$

On peut alors déterminer, pour une valeur de probabilité P donnée, les deux bornes $\Delta\phi_{ki\text{ mod}}^-(P)$ et $\Delta\phi_{ki\text{ mod}}^+(P)$ de l'intervalle standard $I_{\Delta\phi_{ki\text{ mod}}}(L_P)$ associé, ou autrement dit la méconnaissance effective sur la valeur au DDL k de la forme propre $\underline{\phi}_i$.

5.2.3 Méconnaissance effective sur la projection selon une direction d'un mode propre

Le but est ici de pouvoir prendre en compte la totalité du mode propre plutôt que sa valeur en un degré de liberté donné. Néanmoins, comme les méconnaissances effectives sont des quantités scalaires, c'est la projection selon une direction $\underline{N}$ donnée que l'on considère.

On exprime la variation $\Delta\phi_{i\text{ mod}}^N = (\phi_{i\text{ mod}} - \bar{\phi}_i) \cdot \underline{N}$ comme :

$$\Delta\phi_{i\text{ mod}}^N = \sum_{\substack{k=1 \\ k \neq i}}^r [\Delta\phi_{ki}] \bar{\phi}_k \cdot \underline{N} \quad (5.11)$$

où r est le nombre de modes propres retenus, et :

$$\Delta\phi_{ki} = \frac{1}{\bar{\omega}_k^2 - \bar{\omega}_i^2} \sum_{E \in \Omega} \left\{ \bar{\phi}_k^T \Delta \mathbf{K}_E \bar{\phi}_i \right\} \quad (5.12)$$

Les détails de l'obtention de cette relation sont précisés dans le paragraphe A.3 ; pour cela, on a en particulier supposé que les valeurs de méconnaissances étaient suffisamment petites pour utiliser une approximation au premier ordre, et que les modes étaient à nouveau normalisés par rapport à la matrice de masse. En utilisant l'égalité suivante :

$$\bar{\phi}_k^T \mathbf{K}_E \bar{\phi}_i = \frac{1}{2} e_E(\bar{\phi}_k + \bar{\phi}_i) - \frac{1}{2} e_E(\bar{\phi}_k^T - \bar{\phi}_i) \quad (5.13)$$

et la relation fondamentale (4.2), on peut propager les intervalles de méconnaissances de base de la façon suivante :

$$-\Delta\phi_{i\text{ mod}}^{N-} \leq \Delta\phi_{i\text{ mod}}^N \leq \Delta\phi_{i\text{ mod}}^{N+} \quad (5.14)$$

avec :

$$\Delta\phi_{i\text{mod}}^{N-} = \sum_{E \in \Omega} \sum_{\substack{k=1 \\ k \neq i}}^r \frac{1}{2(\bar{\omega}_k^2 - \bar{\omega}_i^2)} \left\{ m_E^- \bar{e}_E(\bar{\phi}_i + \bar{\phi}_k) + m_E^+ \bar{e}_E(\bar{\phi}_i - \bar{\phi}_k) \right\} \bar{\phi}_k \cdot \underline{N} \quad (5.15a)$$

$$\Delta\phi_{i\text{mod}}^{N+} = \sum_{E \in \Omega} \sum_{\substack{k=1 \\ k \neq i}}^r \frac{1}{2(\bar{\omega}_k^2 - \bar{\omega}_i^2)} \left\{ m_E^+ \bar{e}_E(\bar{\phi}_i + \bar{\phi}_k) + m_E^- \bar{e}_E(\bar{\phi}_i - \bar{\phi}_k) \right\} \bar{\phi}_k \cdot \underline{N} \quad (5.15b)$$

On peut alors déterminer, pour une valeur de probabilité P donnée, les deux bornes $\Delta\phi_{i\text{mod}}^{N-}(P)$ et $\Delta\phi_{i\text{mod}}^{N+}(P)$ de l'intervalle standard $I_{\Delta\phi_{i\text{mod}}^{N-}}(L_P)$ associé, ou autrement dit la méconnaissance effective sur la projection de la forme propre $\bar{\phi}_i$ sur la direction $\underline{N}$.

5.3 Calcul de méconnaissances effectives sur un exemple simple

La mise en œuvre de la méthode de calcul des méconnaissances effectives est tout d'abord réalisée sur un exemple académique simple : il s'agit du treillis plan à quatre barres que nous avons eu l'occasion de décrire dans le paragraphe 3.2.2. Le modèle Éléments Finis associé, dont les caractéristiques sont précisées dans le tableau 3.1, constitue le modèle théorique déterministe associé au modèle de méconnaissances de base choisi ; nous présentons par la suite plusieurs possibilités de modélisations de ces méconnaissances.

5.3.1 Premier choix de modèle de méconnaissances de base

On considère ici que les seules sources d'erreur sont liées aux incertitudes concernant le matériau, ce qui oriente notre choix vers une loi normale, centrée autour du modèle déterministe théorique, comme nous l'avons déjà justifié dans le paragraphe 4.2.

Plus précisément, pour chaque barre, la loi normale est caractérisée par deux quantités $\bar{m}_E^-$ et $\bar{m}_E^+$, qui sont les valeurs relatives au-delà desquelles la densité de probabilité initiale (4.4) est mise à zéro, évitant ainsi d'obtenir des réalisations non physiques. Ici, on choisit d'imposer pour chaque barre :

$$\bar{m}_E^- = 0,05 \quad \text{et} \quad \bar{m}_E^+ = 0,05$$

On verra dans le chapitre 6 comment on peut déterminer les méconnaissances de base les plus représentatives d'une réalité expérimentale donnée.

Calcul des méconnaissances effectives sur les pulsations propres et les modes propres

En ce qui concerne les pulsations propres, les intervalles de méconnaissances de base sont propagés à l'aide de la relation (5.6) sous la forme de bornes $\Delta\omega_{i\text{mod}}^{2-}$ et $\Delta\omega_{i\text{mod}}^{2+}$ encadrant la quantité $\Delta\omega_{i\text{mod}}^2$, la transmission des lois de probabilité se faisant soit par la méthode de Monte Carlo, soit par le calcul des fonctions caractéristiques, comme on l'a précisé dans le paragraphe 4.4. On a représenté par exemple sur la figure 5.2 la répartition des bornes $\Delta\omega_{i\text{mod}}^{2-}$ et $\Delta\omega_{i\text{mod}}^{2+}$ obtenues avec la technique de Monte Carlo pour le deuxième mode.

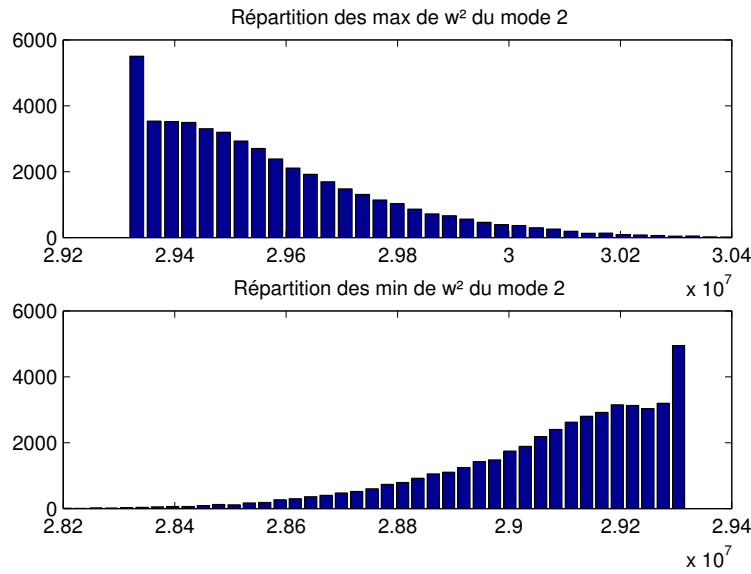


FIG. 5.2 – Répartitions des valeurs de $\bar{\omega}_2^2 + \Delta\omega_{2\text{mod}}^{2+}$ et $\omega_2^2 - \Delta\omega_{2\text{mod}}^{2-}$ obtenues par la méthode de Monte Carlo (50000 tirages)

De la même façon, on trouve sur la figure 5.3 la répartition des bornes $\Delta\phi_{ki\text{mod}}^-$ et $\Delta\phi_{ki\text{mod}}^+$ pour le quatrième mode obtenues par la propagation des intervalles de méconnaissances de base selon la relation (5.10); on peut comparer cette répartition avec les densités de probabilité de ces mêmes quantités calculées à l'aide des fonctions caractéristiques et représentées sur la figure 5.4.

Sur les répartitions des bornes sur les pulsations propres, on peut constater une accumulation des réalisations autour de la valeur théorique $\bar{\omega}_i^2$: ceci est dû à la façon dont on modélise la loi de probabilité normale au travers du concept de méconnaissance, qui ne revient à propager que des intervalles du type $[-m_E^-; 0]$ et $[0; m_E^+]$. Une explication similaire peut également être proposée pour la forme particulière « à deux bosses » des répartitions des valeurs en un DDL d'un mode propre.

À partir de ces distributions, on peut déterminer, pour chaque pulsation propre, et pour une valeur de probabilité P donnée, les deux bornes $\Delta\omega_{i\text{mod}}^{2-}(P)$ et $\Delta\omega_{i\text{mod}}^{2+}(P)$ de l'intervalle standard $I_{\Delta\omega_{i\text{mod}}^2}(L_P)$ associé, ou autrement dit la méconnaissance effective sur le carré de la fréquence propre ω_i^2 . De même, on détermine la méconnaissance effective concernant les formes propres. Comme on le verra plus loin, les phénomènes

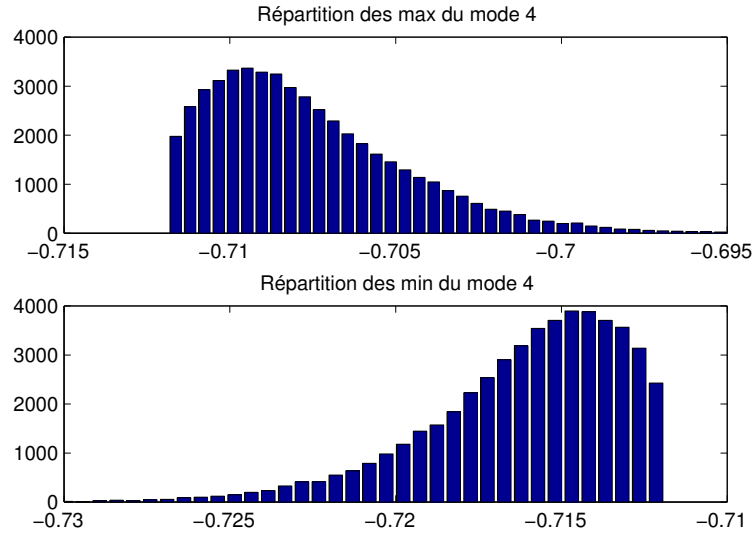


FIG. 5.3 – Répartitions des valeurs de $\bar{\phi}_{24} + \Delta\phi_{24\text{mod}}^+$ et $\bar{\phi}_{24} - \Delta\phi_{24\text{mod}}^-$ obtenues par la méthode de Monte Carlo (50000 tirages)

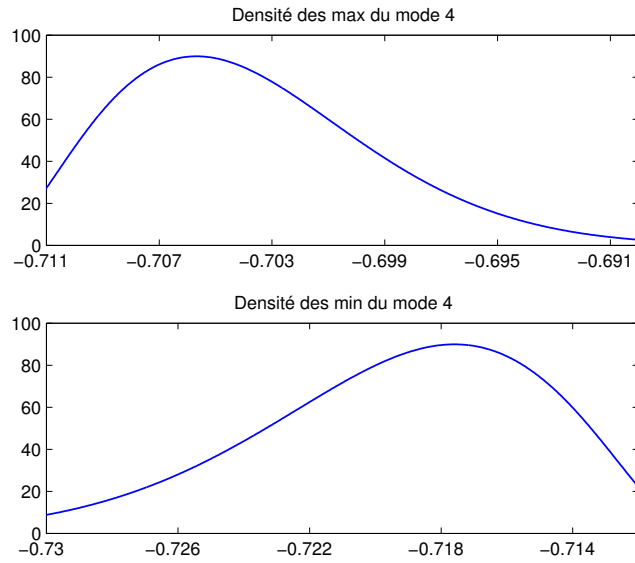


FIG. 5.4 – Densités de probabilité de $\bar{\phi}_{24} + \Delta\phi_{24\text{mod}}^+$ et $\bar{\phi}_{24} - \Delta\phi_{24\text{mod}}^-$ calculées à l'aide des fonctions caractéristiques

d'accumulation constatés dans les distributions des bornes peuvent avoir une influence non négligeable dans la détermination de ces valeurs.

On a récapitulé dans le tableau 5.1 les intervalles de méconnaissances effectives obtenus pour les quatre pulsations propres et les valeurs en un DDL des quatre modes propres, et pour des probabilités $P = 0,70$ et $P = 0,99$. En ce qui concerne les valeurs en un DDL des modes, comme ces derniers sont normalisés par rapport à leur maximum, le DDL considéré dans l'estimation des méconnaissances effectives correspond à la plus grande composante du mode en valeur absolue (après bien sûr la composante normalisée à un) ;

on pourrait bien entendu retenir un tout autre DDL, mais ce choix permet de s'affranchir au maximum des erreurs numériques en ne travaillant pas sur des quantités trop petites.

Valeurs de P	Modes i	Pulsations propres $[\bar{\omega}_i^2 - \Delta\omega_{i\text{mod}}^2; \bar{\omega}_i^2 + \Delta\omega_{i\text{mod}}^2]$	DDL k	Valeurs en un DDL des modes $[\bar{\phi}_{ki} - \Delta\phi_{ki\text{mod}}^-; \bar{\phi}_{ki} + \Delta\phi_{ki\text{mod}}^+]$
0,70	1	$[2, 98; 3, 05] \cdot 10^5$	2	$[0, 90; 0, 92]$
	2	$[2, 90; 2, 96] \cdot 10^6$	3	$[0, 64; 0, 73]$
	3	$[4, 50; 4, 59] \cdot 10^6$	1	$[-0, 50; -0, 46]$
	4	$[8, 24; 8, 42] \cdot 10^6$	2	$[-0, 72; -0, 71]$
0,99	1	$[2, 92; 3, 12] \cdot 10^5$	2	$[0, 89; 0, 93]$
	2	$[2, 85; 3, 02] \cdot 10^6$	3	$[0, 57; 0, 81]$
	3	$[4, 38; 4, 71] \cdot 10^6$	1	$[-0, 53; -0, 42]$
	4	$[7, 97; 8, 68] \cdot 10^6$	2	$[-0, 73; -0, 70]$

TAB. 5.1 – Valeurs à 70 % et 99 % obtenues avec le modèle de méconnaissances pour le treillis à quatre barres

Comparaison avec une simulation stochastique

Pour pouvoir juger de la qualité des résultats précédents, nous les comparons avec ceux issus de la modélisation stochastique détaillée dans le paragraphe 3.2. En effet, dans cet exemple, la raideur de chaque barre est modélisée par une variable aléatoire, dont la loi de probabilité est la même que celle qui est à l'origine de la loi retenue pour le modèle avec méconnaissances utilisé, à savoir une loi normale centrée « tronquée à $\pm 5\%$ ». On rappelle que cette simulation stochastique à l'aide de la technique de Monte Carlo est en quelque sorte représentative d'une réalité « expérimentale » parfaitement connue.

Il est alors possible de confronter les répartitions « expérimentales » de chaque quantité d'intérêt avec celles obtenues par propagation des méconnaissances de base et de leurs lois de probabilité associées. Par exemple, on peut comparer les répartitions des valeurs de la seconde pulsation propre du treillis, respectivement pour la simulation stochastique sur la figure 3.4 et pour le modèle avec méconnaissances sur la figure 5.2.

Néanmoins, le véritable point de comparaison concerne plutôt la confrontation des différentes valeurs à $P\%$ qui sont les quantités que l'on estime représentatives des dispersions des quantités d'intérêt ; ceci est effectué pour les valeurs à 70 % et 99 % dans les tableaux 5.2 et 5.3. **On se rend alors compte de la relativement bonne correspondance entre les valeurs prévues par le modèle avec méconnaissances et celles obtenues à l'aide de la simulation stochastique qui nous sert ici de référence** ; en effet, les écarts relatifs entre les valeurs correspondantes sont généralement bien inférieurs à 1 %.

On peut également constater que **les valeurs à 99 % sont généralement plus précises que celles à 70 %**, ceci étant vraisemblablement lié à l'accumulation des réalisations provoquée par la façon de réaliser les tirages de méconnaissances de base, et constatée sur les figures 5.2 et 5.3 ; en effet, plus la valeur de la probabilité P est faible, plus les méconnaissances effectives sont proches de la valeur associée au modèle théorique déterministe, et donc plus elles sont perturbées par l'influence de cette accumulation.

Valeurs de P	Modes i	Modèle avec méconnaissances $[\bar{\omega}_i^2 - \Delta\omega_{i\text{mod}}^2; \bar{\omega}_i^2 + \Delta\omega_{i\text{mod}}^2]$	Simulation stochastique $[\omega_{i\text{exp}}^2; \omega_{i\text{exp}}^2]$
0,70	1	$[2, 98; 3, 05] \cdot 10^5$	$[2, 97; 3, 06] \cdot 10^5$
	2	$[2, 90; 2, 96] \cdot 10^6$	$[2, 90; 2, 97] \cdot 10^6$
	3	$[4, 50; 4, 59] \cdot 10^6$	$[4, 47; 4, 61] \cdot 10^6$
	4	$[8, 24; 8, 42] \cdot 10^6$	$[8, 18; 8, 49] \cdot 10^6$
0,99	1	$[2, 92; 3, 12] \cdot 10^5$	$[2, 91; 3, 13] \cdot 10^5$
	2	$[2, 85; 3, 02] \cdot 10^6$	$[2, 84; 3, 02] \cdot 10^6$
	3	$[4, 38; 4, 71] \cdot 10^6$	$[4, 37; 4, 72] \cdot 10^6$
	4	$[7, 97; 8, 68] \cdot 10^6$	$[7, 95; 8, 71] \cdot 10^6$

TAB. 5.2 – Comparaison des valeurs à 70 % et 99 % obtenues pour les pulsations propres du treillis à quatre barres, respectivement pour le modèle avec méconnaissances et pour une simulation de Monte Carlo

Valeurs de P	Modes i	DDL k	Modèle avec méconnaissances $[\bar{\phi}_{ki} - \Delta\phi_{ki\text{mod}}; \bar{\phi}_{ki} + \Delta\phi_{ki\text{mod}}]$	Simulation stochastique $[\phi_{ki\text{exp}}; \phi_{ki\text{exp}}]$
0,70	1	2	$[0, 90; 0, 92]$	$[0, 90; 0, 91]$
	2	3	$[0, 64; 0, 73]$	$[0, 64; 0, 74]$
	3	1	$[-0, 50; -0, 46]$	$[-0, 50; -0, 45]$
	4	2	$[-0, 72; -0, 71]$	$[-0, 72; -0, 71]$
0,99	1	2	$[0, 89; 0, 93]$	$[0, 90; 0, 91]$
	2	3	$[0, 57; 0, 81]$	$[0, 58; 0, 83]$
	3	1	$[-0, 53; -0, 42]$	$[-0, 54; -0, 43]$
	4	2	$[-0, 73; -0, 70]$	$[-0, 72; -0, 71]$

TAB. 5.3 – Comparaison des valeurs à 70 % et 99 % obtenues pour les valeurs en un DDL des modes propres du treillis à quatre barres, respectivement pour le modèle avec méconnaissances et pour une simulation de Monte Carlo

Enfin, **les intervalles de méconnaissances effectives sur les modes propres ont tendance à être plus larges que ceux qui concernent les pulsations propres** ; ceci est dû au mode de propagation des intervalles de méconnaissances de base qui font intervenir dans ce cas deux fois l'inégalité fondamentale (4.2). L'effet est tout de même limité par le fait que les méconnaissances de base sont des variables aléatoires ; on se trouverait sinon dans la situation de la propagation d'intervalles déterministes qui, comme on l'a montré dans le paragraphe 2.3, a tendance à surestimer les incertitudes, quelquefois de façon très exagérée.

Quelques remarques sur la mise en œuvre de la théorie des méconnaissances

L'application de la théorie des méconnaissances afin de déterminer la méconnaissance effective sur une quantité d'intérêt donnée est aisée ; les différentes étapes nécessaires pour sa mise en œuvre, ainsi que les estimations grossières des coûts de calcul associés, sont les suivantes :

- **la résolution complète du modèle théorique déterministe**, ce qui, en pratique, se traduit par la résolution d'un problème aux valeurs propres ;

- **une phase de pré-traitement qui se résume au calcul des énergies de déformation modales de ce modèle**, peu coûteux, ainsi que, éventuellement, à celui des énergies intervenant dans l’expression de la méconnaissance effective sur la valeur en un DDL d’un mode propre, qui nécessite des résolutions de problèmes aux valeurs propres ; heureusement, pour les modèles à grand nombre de degrés de liberté, il est possible d’utiliser des bases de réduction pour diminuer de façon très nette le coût de ces résolutions, comme il est précisé dans le paragraphe A.2 ;
- **la phase de calcul des répartitions des bornes sur les quantités d’intérêt**, dont le coût réside principalement dans la façon d’obtenir les méconnaissances de base et de réaliser leur propagation : soit par tirage de Monte Carlo, puis propagation des intervalles ainsi obtenus, soit par l’intermédiaire des fonctions caractéristiques en utilisant une transformée de Fourier rapide ; dans les deux cas, le coût réside dans le nombre de tirages ou le nombre de points de discrétisation, la seconde méthode étant, à précisions égales, la moins onéreuse ;
- **une phase de post-traitement associée à la détermination des méconnaissances effectives** pour des valeurs de probabilité données.

Ainsi, à la simplicité de mise en œuvre de la méthode s’ajoute aussi un coût de calcul total bien plus faible que celui correspondant à la simulation stochastique de Monte Carlo dont les résultats nous ont servi ici de base de comparaison ; en effet, cette technique nécessite la résolution d’un problème aux valeurs propres pour chaque structure simulée par un jeu donné de tirages de valeurs de raideur.

5.3.2 Deuxième choix de modèle de méconnaissances de base

Les quatre barres étant constituées du même matériau, on peut envisager de les considérer toutes ensemble dans le modèle avec méconnaissances, renforçant ainsi le principe de globalisation mis en avant par ce concept ; en pratique, ceci revient à associer une seule méconnaissance de base pour le groupe de barres, et à utiliser la somme des énergies de déformation des barres, au lieu de considérer ces énergies individuellement.

On choisit une loi de probabilité normale centrée pour modéliser la méconnaissance de base $(\mathbf{m}_g^-, \mathbf{m}_g^+)$ relative au groupe $g = \{b13, b24, b34, b23\}$, telle que :

$$\overline{m}_g^- = 0,05 \quad \text{et} \quad \overline{m}_g^+ = 0,05$$

Il s’agit donc d’une loi identique à celle que l’on a appliquée individuellement à chaque barre dans le modèle proposé dans le paragraphe précédent.

Les résultats de la propagation de cette méconnaissance de base sur le groupe g sont mentionnés dans le tableau 5.4. On constate que les intervalles de méconnaissances effectives obtenus sont plus larges que ceux présentés dans le tableau 5.1, qui sont issus de la donnée de méconnaissances de base sur chaque barre – et donc aussi de ceux correspondant à la simulation stochastique, indiqués dans les tableaux 5.2 et 5.3. En effet, les écarts relatifs sont de l’ordre de 0,5 % à 2 % sur les méconnaissances effectives concernant les pulsations propres, et peuvent aller jusqu’à presque 3 % pour les modes propres.

Valeurs de P	Modes i	Pulsations propres $[\bar{\omega}_i^2 - \Delta\omega_{i\text{mod}}^2; \bar{\omega}_i^2 + \Delta\omega_{i\text{mod}}^2]$	DDL k	Valeurs en un DDL des modes $[\bar{\phi}_{ki} - \Delta\phi_{ki\text{mod}}^-; \bar{\phi}_{ki} + \Delta\phi_{ki\text{mod}}^+]$
0,70	1	[2, 96; 3, 08]. 10^5	2	[0, 90; 0, 92]
	2	[2, 88; 2, 99]. 10^6	3	[0, 64; 0, 73]
	3	[4, 45; 4, 63]. 10^6	1	[-0, 50; -0, 46]
	4	[8, 16; 8, 50]. 10^6	2	[-0, 72; -0, 71]
0,99	1	[2, 87; 3, 17]. 10^5	2	[0, 89; 0, 93]
	2	[2, 79; 3, 08]. 10^6	3	[0, 57; 0, 80]
	3	[4, 32; 4, 77]. 10^6	1	[-0, 53; -0, 43]
	4	[7, 92; 8, 75]. 10^6	2	[-0, 73; -0, 70]

TAB. 5.4 – Valeurs à 70 % et 99 % obtenues avec un modèle de méconnaissance « groupé » pour le treillis à quatre barres

Ceci peut s'expliquer par l'hypothèse que l'on pose de façon implicite quand on regroupe les quatre barres dans la description de la méconnaissance de base : chacune des barres a exactement le même comportement que les autres, sa méconnaissance étant identique à celle des quatre autres. On retrouverait les mêmes résultats dans la simulation stochastique si l'on considérait que la dispersion de raideur était la même simultanément sur les barres, chaque tirage de cette raideur étant alors appliqué sur les quatre barres. On conçoit aisément que ceci a tendance à rigidifier ou assouplir unilatéralement le comportement d'un treillis simulé, puisqu'il n'y a pas de possibilité de « compenser » un excès de raideur sur une barre par une diminution sur une autre.

Malgré cette surestimation, les résultats restent convenables, et **le choix de grouper des sous-structures qui sont censées subir les mêmes dispersions est une solution valide**, qui s'inscrit dans le concept de méconnaissance qui a pour but de globaliser les méconnaissances.

5.3.3 Troisième choix de modèle de méconnaissances de base

Pour illustrer un autre cas de modélisation des méconnaissances de base, on reprend l'exemple du treillis à quatre barres. La différence majeure concerne la barre 2-3, pour laquelle on décide de modéliser la méconnaissance de base par un intervalle purement déterministe, afin de traduire que le modèle théorique déterministe de la barre est très simplifié par rapport à la réalité : par exemple, on peut imaginer qu'on modélise de façon linéaire une barre qui, dans la réalité, présente un comportement non-linéaire. Les caractéristiques des méconnaissances de base sont résumées dans le tableau 5.5 ; pour plus de clarté, on a précisé les lois normales de deux façons distinctes, mais équivalentes, d'une part par leur « étendue » caractérisée par les valeurs $(\bar{m}_E^-, \bar{m}_E^+)$, d'autre part par leurs deux premiers moments statistiques (exprimés de façon relative, car divisés par la valeur de raideur du modèle déterministe).

La propagation des méconnaissances de base se fait classiquement selon les expressions (5.6) et (5.10), avec une méconnaissance de base déterministe sur la barre 2-3, telle que :

$$m_{b23}^- = 0,1 \quad \text{et} \quad m_{b23}^+ = 0,2$$

On peut voir sur la figure 5.5 la distribution des bornes de la quantité $\omega_{3\text{mod}}^2$ issue des choix précédents de méconnaissances : on constate que la méconnaissance de base modélisée par un intervalle purement déterministe s'est propagée sous la forme d'un intervalle lui aussi déterministe. **On obtient un intervalle d'ignorance maximale en ce qui concerne la quantité d'intérêt**, inclus dans tous les intervalles de méconnaissances effectives associés à des valeurs P de probabilité ; ainsi, **la quantité d'intérêt est toujours susceptible de se situer dans cet intervalle sans que l'on puisse en quantifier la probabilité**.

Barres	Modèle déterministe	Loi choisie	Domaine relatif	Moments statistiques relatifs (μ : moyenne / σ : écart type)
1-3	$\bar{E}_{b13} = 72\text{GPa}$	normale	$[-0,05; 0,05]$	$\mu = 0 / \sigma = 0,019$
2-4	$\bar{E}_{b24} = 72\text{GPa}$	normale	$[-0,05; 0,05]$	$\mu = 0 / \sigma = 0,019$
3-4	$\bar{E}_{b34} = 72\text{GPa}$	normale	$[-0,05; 0,05]$	$\mu = 0 / \sigma = 0,019$
2-3	$\bar{E}_{b23} = 72\text{GPa}$	pas de loi	$[-0,10; 0,20]$	aucun

NB : Les caractéristiques des lois sont données relativement au module d'Young du modèle déterministe.

TAB. 5.5 – Méconnaissances de base du treillis à quatre barres dans le cas d'un modèle de sous-structure très simplifié par rapport à la réalité

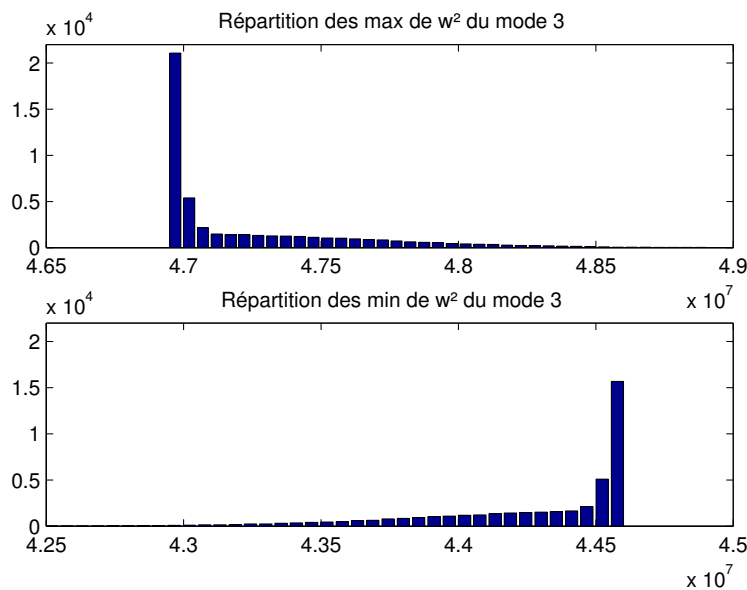


FIG. 5.5 – Répartitions des valeurs de $\bar{\omega}_3^2 + \Delta\omega_{3\text{mod}}^2$ et $\bar{\omega}_3^2 - \Delta\omega_{3\text{mod}}^2$ obtenues par la méthode de Monte Carlo (50000 tirages) pour le treillis à quatre barres dans le cas d'un modèle de sous-structure très simplifié par rapport à la réalité

Le concept de méconnaissance a donc permis d'utiliser un modèle de sous-structure très imparfait par rapport à la réalité afin d'obtenir des valeurs caractéristiques de la dispersion des quantités d'intérêt : il a suffi pour cela de décrire la méconnaissance sur la sous-structure concernée comme un intervalle purement déterministe, ce qui n'aurait pas été directement réalisable avec une simulation stochastique.

5.4 Prise en compte d'une forte méconnaissance

5.4.1 Principe et mise en œuvre

Tous les calculs de méconnaissances effectives réalisés précédemment reposent sur les expressions de la variation d'une pulsation propre et d'un mode propre approximées au premier ordre, ce qui est suffisant pour de « faibles » valeurs de méconnaissances de base. Bien qu'il soit difficile de préciser ce que l'on entend par « faibles » valeurs, tant cette notion dépend de la structure étudiée, on peut toutefois admettre que des valeurs de méconnaissances de l'ordre de 0,4 ou 0,5 (c'est-à-dire des incertitudes de l'ordre de 40 % ou 50 %) entraînent vraisemblablement une certaine prudence quant à la validité de l'approximation utilisée.

On va décrire ici comment résoudre ce problème dans le cas d'une méconnaissance de base pouvant prendre des valeurs particulièrement fortes, sans toutefois nuire trop à la simplicité et à l'efficacité de la méthode initiale. Notons $\mathcal{E}$ la sous-structure dont la méconnaissance de base est susceptible de prendre de telles valeurs. L'idée consiste à être capable pour chaque réalisation de la méconnaissance sur $\mathcal{E}$ de déterminer les pulsations propres et modes propres de la structure « perturbée » par la méconnaissance forte $m_{\mathcal{E}}$, et ce de façon peu coûteuse, mais relativement fine, et ensuite à appliquer, pour les méconnaissances sur les autres sous-structures, la méthode approchée habituelle.

Problème aux valeurs propres associé à la forte méconnaissance et résolution

On cherche à être capable de déterminer, pour toute valeur $m_{\mathcal{E}}$, les pulsations propres $\tilde{\omega}_i(m_{\mathcal{E}})$ et modes propres $\tilde{\phi}_i(m_{\mathcal{E}})$ du modèle associé au modèle perturbé par la méconnaissance sur la sous-structure $\mathcal{E}$ seulement ; on doit pour cela résoudre le problème aux valeurs propres suivant :

$$(\bar{\mathbf{K}} + m_{\mathcal{E}}\Delta\mathbf{K}_{\mathcal{E}})\tilde{\phi}_i = \tilde{\omega}_i^2\mathbf{M}\tilde{\phi}_i \quad (5.16)$$

où $\Delta\mathbf{K}_{\mathcal{E}}$ est la contribution de la sous-structure $\mathcal{E}$ à la matrice de rigidité du modèle déterministe. Précisons que, conformément à la façon dont sont définies les méconnaissances de base, la quantité $m_{\mathcal{E}}$ correspond soit à la borne supérieure $m_{\mathcal{E}}^+$, soit à la borne inférieure $m_{\mathcal{E}}^-$ de méconnaissance de base.

Comme il serait très coûteux de résoudre ce problème pour chaque valeur de $m_{\mathcal{E}}$, on développe une technique qui va permettre d'évaluer de façon approchée les quantités $\tilde{\omega}_i(m_{\mathcal{E}})$ et $\tilde{\phi}_i(m_{\mathcal{E}})$.

La première étape consiste à introduire une base modale de réduction pour diminuer la taille des problèmes aux valeurs propres que l'on doit résoudre. Pour cela, on commence par retenir les n modes $\bar{\phi}_i$ tels que la sous-structure $\mathcal{E}$ apporte une contribution à l'énergie de déformation totale de la structure qui soit parmi les plus fortes constatées : on sélectionne $\bar{\phi}_i$ si :

$$\bar{e}_{\mathcal{E}}(\bar{\phi}_i) \geq k_{\%}\bar{e}(\bar{\phi}_i) \quad (5.17)$$

où, typiquement, on choisit $k_{\%} \simeq \frac{1}{N}$, où N désigne le nombre de modes propres associés au modèle théorique déterministe utilisé.

Pour augmenter l'efficacité de cette base de réduction, on enrichit cette dernière de modes statiques $\underline{\psi}_i$ définis par :

$$\bar{\mathbf{K}}\underline{\psi}_i = \Delta\mathbf{K}_{\mathcal{E}}\bar{\underline{\phi}}_i \quad (5.18)$$

Après élimination des modes statiques colinéaires aux modes initiaux, et après normalisation par rapport à la matrice de masse ($\underline{\psi}_j^T \mathbf{M} \underline{\psi}_i = \delta_{ji}$), on obtient des relations d'orthogonalité entre les n modes initiaux retenus et les $m \leq n$ modes statiques calculés :

$$\forall i, j \quad \underline{\psi}_j^T \mathbf{M} \bar{\underline{\phi}}_i = 0 \quad \text{et} \quad \underline{\psi}_j^T \mathbf{K} \bar{\underline{\phi}}_i = 0 \quad (5.19)$$

La solution du problème aux valeurs propres est alors recherchée sous la forme :

$$\tilde{\underline{\phi}}_i = \mathbf{\Phi} \underline{x} + \mathbf{\Psi} \underline{y} = (\mathbf{\Phi} | \mathbf{\Psi}) \begin{pmatrix} \underline{x} \\ \underline{y} \end{pmatrix} \quad (5.20)$$

où $\mathbf{\Phi} = (\bar{\underline{\phi}}_1 \quad \dots \quad \bar{\underline{\phi}}_n)$ et $\mathbf{\Psi} = (\underline{\psi}_1 \quad \dots \quad \underline{\psi}_m)$, avec $m \leq n$.

En multipliant à gauche par $(\mathbf{\Phi} | \mathbf{\Psi})^T$ les deux termes de l'égalité, le problème se réécrit finalement comme :

$$\left\{ \left(\begin{array}{c|c} \begin{matrix} \omega_1^2 & 0 \\ & \ddots \\ 0 & \omega_n^2 \end{matrix} & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{\Psi}^T \bar{\mathbf{K}} \mathbf{\Psi} \end{array} \right) + m_{\mathcal{E}} \left(\begin{array}{c|c} \mathbf{\Phi}^T \Delta\mathbf{K}_{\mathcal{E}} \mathbf{\Phi} & \mathbf{\Phi}^T \Delta\mathbf{K}_{\mathcal{E}} \mathbf{\Psi} \\ \hline \mathbf{\Psi}^T \Delta\mathbf{K}_{\mathcal{E}} \mathbf{\Phi} & \mathbf{\Psi}^T \Delta\mathbf{K}_{\mathcal{E}} \mathbf{\Psi} \end{array} \right) \right\} \begin{pmatrix} \underline{x} \\ \underline{y} \end{pmatrix} = \tilde{\omega}_i^2 \begin{pmatrix} \underline{x} \\ \underline{y} \end{pmatrix} \quad (5.21)$$

La résolution de ce problème aux valeurs propres de dimension $(n + m) \times (n + m)$ beaucoup plus petite que celle du problème initial conduit donc à une estimation efficace de $\tilde{\omega}_i(m_{\mathcal{E}})$ et $\tilde{\phi}_i(m_{\mathcal{E}})$.

Interpolation des pulsations propres et modes propres

Malgré la diminution substantielle du coût de calcul offerte par la réduction précédente, il serait encore trop onéreux de résoudre le problème aux valeurs propres précédent pour chaque réalisation de $m_{\mathcal{E}}$; on décide donc de ne faire cet effort que pour quelques valeurs particulières, qui permettront dans un second temps de réaliser une approximation par interpolation de $\tilde{\omega}_i(m_{\mathcal{E}})$ et $\tilde{\phi}_i(m_{\mathcal{E}})$.

Typiquement, **on retient une interpolation quadratique pour l'estimation des pulsations propres et modes propres**, c'est-à-dire un ordre de plus que l'hypothèse linéaire de la méthode initiale. Cette interpolation quadratique permet d'obtenir également des valeurs approchées pour les énergies de déformation modales :

$$\tilde{e} = \frac{1}{2} \tilde{\underline{\phi}}_i^T (\bar{\mathbf{K}} + m_{\mathcal{E}} \Delta\mathbf{K}_{\mathcal{E}}) \tilde{\underline{\phi}}_i \quad (5.22)$$

L'interpolation quadratique des modes propres implique une interpolation d'ordre cinq pour l'estimation des énergies de déformation. Typiquement, on se sert du domaine de définition de la loi de probabilité de la méconnaissance de base $(\mathbf{m}_{\mathcal{E}}^-, \mathbf{m}_{\mathcal{E}}^+)$, c'est-à-dire de la valeur des quantités $\overline{m}_{\mathcal{E}}^-$ et $\overline{m}_{\mathcal{E}}^+$; on décide alors de baser l'interpolation sur l'estimation de $\tilde{e}$ en six points régulièrement espacés sur chaque domaine $[-\overline{m}_{\mathcal{E}}^-; 0]$ et $[0; \overline{m}_{\mathcal{E}}^+]$, par exemple :

$$\left\{ -\overline{m}_{\mathcal{E}}^-; -\frac{2\overline{m}_{\mathcal{E}}^-}{3}; -\frac{\overline{m}_{\mathcal{E}}^-}{3}; \frac{\overline{m}_{\mathcal{E}}^+}{3}; \frac{2\overline{m}_{\mathcal{E}}^+}{3}; \overline{m}_{\mathcal{E}}^+ \right\}$$

Les six points peuvent être aussi utilisés pour l'approximation des fréquences propres et modes propres, en réalisant une interpolation quadratique sur chacun des intervalles $[-\overline{m}_{\mathcal{E}}^-; 0]$ et $[0; \overline{m}_{\mathcal{E}}^+]$ pour plus de précision :

$$\left\{ -\overline{m}_{\mathcal{E}}^-; -\frac{2\overline{m}_{\mathcal{E}}^-}{3}; -\frac{\overline{m}_{\mathcal{E}}^-}{3} \right\} \quad \text{et} \quad \left\{ \frac{\overline{m}_{\mathcal{E}}^+}{3}; \frac{2\overline{m}_{\mathcal{E}}^+}{3}; \overline{m}_{\mathcal{E}}^+ \right\}$$

Calcul des méconnaissances effectives

L'intégration des méconnaissances de base sur les autres sous-structures se fait de la même façon qu'auparavant, à ce détail près que c'est à partir du modèle perturbé par la réalisation $\mathbf{m}_{\mathcal{E}}$ que l'on vient ajouter à chaque fois la contribution des autres sous-structures. Ainsi, par propagation des intervalles de méconnaissances de base $([-\overline{m}_{\mathcal{E}}^-; \overline{m}_{\mathcal{E}}^+])$, on obtient par exemple les bornes suivantes pour les pulsations propres :

$$\Delta\omega_{i \text{ mod}}^{2-} = 2 \sum_{E \neq \mathcal{E}} \mathbf{m}_{\mathcal{E}}^- \tilde{e}_E \left(\tilde{\phi}_i(\mathbf{m}_{\mathcal{E}}) \right) - (\tilde{\omega}_i(\mathbf{m}_{\mathcal{E}})^2 - \bar{\omega}_i^2) \quad (5.23a)$$

$$\Delta\omega_{i \text{ mod}}^{2+} = 2 \sum_{E \neq \mathcal{E}} \mathbf{m}_{\mathcal{E}}^+ \tilde{e}_E \left(\tilde{\phi}_i(\mathbf{m}_{\mathcal{E}}) \right) + \tilde{\omega}_i(\mathbf{m}_{\mathcal{E}})^2 - \bar{\omega}_i^2 \quad (5.23b)$$

et des expressions analogues pour les bornes sur les modes propres.

On obtient ainsi comme auparavant les dispersions des bornes $\Delta\omega_{i \text{ mod}}^{2-}$ et $\Delta\omega_{i \text{ mod}}^{2+}$ qui encadrent la quantité $\Delta\omega_{i \text{ mod}}^2$ et donc, pour une valeur de probabilité P donnée, les deux bornes $\Delta\omega_{i \text{ mod}}^{2-}(P)$ et $\Delta\omega_{i \text{ mod}}^{2+}(P)$ de l'intervalle standard $I_{\Delta\omega_{i \text{ mod}}^2}(L_P)$ associé, ou autrement dit la méconnaissance effective sur le carré de la fréquence propre ω_i^2 .

5.4.2 Application sur un exemple simple

Pour pouvoir juger de la méthode introduite pour la prise en compte d'une forte méconnaissance sur une sous-structure, on étudie à nouveau l'exemple, présenté dans le paragraphe 5.3.1, du treillis plan à quatre barres dont les méconnaissances de base sont modélisées par des lois normales à la différence que l'une des barres est maintenant caractérisée par une loi très « dispersée ».

Barres	Modèle déterministe	Loi choisie	Domaine $[-\bar{m}_E^-; \bar{m}_E^+]$	Moments statistiques relatifs (μ : moyenne / σ : écart type)
{2-4,3-4,2-3}	$\bar{E}_{\{b_{24},b_{34},b_{23}\}} = 72\text{GPa}$	normale	$[-0,05; 0,05]$	$\mu = 0 / \sigma = 0,019$
1-3	$\bar{E}_{b_{23}} = 72\text{GPa}$	normale	$[-0,1; 0,7]$	$\mu = 0,3 / \sigma = 0,155$

TAB. 5.6 – Modélisation des méconnaissances de base du treillis à quatre barres dans le cas d’une forte méconnaissance

Les modèles supposés de méconnaissances sont récapitulés dans le tableau 5.6. On a volontairement très exagéré l’écart type de la loi s’appliquant à la barre b_{13} pour rendre le phénomène de « forte » méconnaissance plus visible sur cet exemple très simple.

C’est le premier mode qui est le plus affecté par la forte méconnaissance ; c’est donc celui-ci qui est retenu pour la détermination des méconnaissances effectives. Même si ce n’est pas nécessaire vu le faible nombre de degrés de liberté de la structure, on applique la méthode dans son ensemble avec une base de réduction comprenant les deux premiers modes propres, complétée par un mode statique.

La propagation des méconnaissances de base est réalisée par une méthode de Monte Carlo avec 50000 tirages, et les distributions obtenues permettent de déterminer les méconnaissances effectives sur la première pulsation propre, pour des valeurs de probabilités P de 0,70 et 0,99. On a aussi effectué, à des fins de comparaison, ce calcul avec la méthode initiale et son approximation au premier ordre. De plus, on a également mis en œuvre une simulation de Monte Carlo avec 50000 tirages de valeurs de raideur selon des lois de probabilité en accord avec les lois définissant les méconnaissances de base sur les barres ; les résultats provenant de cette simulation stochastique peuvent être encore une fois considérés comme issus d’un calcul exact sur une réalité « expérimentale » parfaitement connue, et constituent donc une référence pour les comparaisons effectuées.

Les méconnaissances effectives pour les deux méthodes, ainsi que les valeurs à P % pour la référence « expérimentale », obtenues pour la première pulsation propre, sont consignées dans le tableau 5.7. On constate que la méthode initiale surestime de près de 5 % la méconnaissance effective à 99 % sur la pulsation propre, et que la technique qui prend en compte la forte méconnaissance sur la barre 2-3 donne des résultats plus proches de ceux obtenus à l’aide de la simulation stochastique.

Valeurs de P	Méthode initiale	Simulation stochastique $[\omega_{1\text{exp}}^2(P); \omega_{1\text{exp}}^2(P)]$	Méthode avec forte méconnaissance
0,70	$[2,95; 3,31].10^6$	$[3,02; 3,35].10^6$	$[3,13; 3,36].10^6$
0,99	$[2,90; 3,63].10^6$	$[2,92; 3,49].10^6$	$[2,89; 3,50].10^6$

NB : Les méconnaissances effectives sont de la forme $[\bar{\omega}_1^2 - \Delta\omega_{1\text{mod}}^2(P); \bar{\omega}_1^2 + \Delta\omega_{1\text{mod}}^2(P)]$

TAB. 5.7 – Méconnaissances effectives pour le treillis à quatre barres et une forte méconnaissance sur la barre b_{13}

Néanmoins, on peut également remarquer des résultats moins bons en ce qui concerne les bornes inférieures des intervalles de méconnaissances effectives calculés ; une explication possible, déjà évoquée dans le paragraphe 5.3.1, est liée à la façon de modéliser les méconnaissances de base, qui entraîne des résultats moins bons pour les valeurs à 70 % que pour les valeurs à 99 %.

Détermination des méconnaissances de base à partir de données expérimentales

L'incertitude n'est pas dans les choses, mais dans notre tête : l'incertitude est une méconnaissance.

Jacob Bernoulli

Sommaire

6.1 Point de vue adopté	96
6.1.1 Données expérimentales	96
6.1.2 Confrontation entre le modèle et les données expérimentales	97
6.1.3 Détermination des méconnaissances de base représentatives	97
6.2 Réduction des méconnaissances de base : première formulation	99
6.2.1 Description du principe associé au procédé de réduction	99
6.2.2 Prise en compte des cas les plus défavorables	100
6.3 Réduction des méconnaissances de base sur un exemple académique	103
6.3.1 Structure étudiée	103
6.3.2 Mise en œuvre du procédé de réduction	105
6.3.3 Cas d'une forte méconnaissance	111
6.4 Réduction des méconnaissances de base : seconde formulation	113
6.4.1 Notion de représentativité d'une mesure expérimentale	113
6.4.2 Illustration	114
6.5 Vers une stratégie générale de réduction des méconnaissances	117
6.5.1 Définition du modèle initial de méconnaissances de base	118
6.5.2 Détermination de l'ordre des réductions successives	118
6.5.3 Processus de réduction des méconnaissances de base	119
6.5.4 Interprétation des résultats	120

La détermination des méconnaissances de base grâce à l'information expérimentale constituée par les données mesurées sur une famille de structures réelles constitue l'objet de ce sixième chapitre ; la méthode employée est un procédé de réduction des méconnaissances de base avec l'idée fondatrice que les données expérimentales sont une source d'information susceptible de réduire les méconnaissances sur les différentes sous-structures à partir d'un niveau majorant fixé *a priori*. L'application à un exemple académique constitue le prélude d'une stratégie générale de réduction des méconnaissances de base.

6.1 Point de vue adopté

On a vu dans le chapitre précédent comment déterminer à partir des méconnaissances de base les méconnaissances effectives, encadrant une quantité d'intérêt définie sur l'ensemble de la structure, pour différentes valeurs de probabilité P . Pour pouvoir valider le modèle de méconnaissances grâce à ces résultats, il est nécessaire de les confronter avec des données équivalentes provenant de la réalité observée.

6.1.1 Données expérimentales

On a précisé dans le paragraphe 3.1.2 comment définir les quantités représentatives de la dispersion d'une grandeur d'intérêt donnée sur la famille de structures réelles considérée. Ce sont ces quantités que nous allons confronter aux méconnaissances effectives calculées à partir des modèles proposés, à la différence que nous allons redéfinir ces valeurs expérimentales relativement à la quantité d'intérêt associée au modèle théorique déterministe, comme c'est déjà le cas dans la définition des méconnaissances.

Ainsi, on associe à la famille de structures réelles deux valeurs $\Delta\alpha_{\text{exp}}^-(P)$ et $\Delta\alpha_{\text{exp}}^+(P)$ qui, pour une valeur P donnée, encadrent une proportion P de valeurs mesurées de la quantité expérimentale $\Delta\alpha_{\text{exp}} = \alpha_{\text{exp}} - \bar{\alpha}$. Là encore, pour éviter toute ambiguïté, on impose que ces deux valeurs soient celles qui définissent le plus petit intervalle contenant une proportion P des quantités mesurées $\Delta\alpha_{\text{exp}}$, ce qui constitue en fait une condition analogue à celle que l'on applique au calcul des méconnaissances effectives. Le principe des valeurs expérimentales à P % est de nouveau représenté sur la figure 6.1.

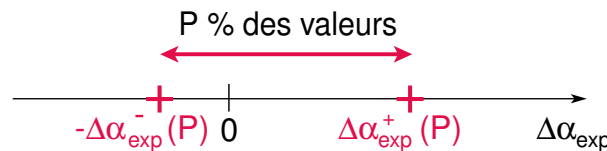


FIG. 6.1 – Valeurs expérimentales à P % de la quantité d'intérêt $\Delta\alpha_{\text{exp}}$

6.1.2 Confrontation entre le modèle avec méconnaissances et les données expérimentales

De façon intuitive, on peut affirmer que l'écart entre les méconnaissances effectives calculées et les valeurs expérimentales à $P\%$ doit traduire la qualité de la représentativité du modèle avec méconnaissances par rapport à la réalité mesurée.

On ajoute toutefois une condition supplémentaire : on souhaite imposer que les méconnaissances de base soient telles que la probabilité d'avoir la quantité d'intérêt expérimentale dans l'intervalle standard associé à une probabilité d'intervalle P donnée soit plus grande que cette dernière, ce qui s'écrit :

$$P(\Delta\alpha_{\text{exp}} \in I_{\Delta\alpha_{\text{mod}}}(L)) \geq P_{\alpha}(L) \quad \forall L \quad (6.1)$$

Compte tenu des remarques évoquées dans le paragraphe 4.3.3, où l'on avait précisé que la probabilité d'intervalle $P_{\alpha}(L)$ était en fait un minorant de la probabilité d'avoir effectivement la quantité d'intérêt associée au modèle dans l'intervalle standard associé à $P_{\alpha}(L)$, le choix de cette condition est logique : en effet, il permet de définir un critère de sécurité, en imposant que la probabilité d'avoir la quantité d'intérêt expérimentale à l'intérieur de l'intervalle standard $I_{\Delta\alpha_{\text{mod}}}(L)$ doit être elle aussi minorée par la probabilité d'intervalle $P_{\alpha}(L)$.

En termes d'intervalles standards, cela revient à dire que l'on doit avoir des relations d'inclusions des valeurs expérimentales à $P\%$ considérées vis-à-vis des méconnaissances effectives calculées, à savoir :

$$\Delta\alpha_{\text{exp}}^{-}(P) \leq \Delta\alpha_{\text{mod}}^{-}(P) \quad \text{et} \quad \Delta\alpha_{\text{exp}}^{+}(P) \leq \Delta\alpha_{\text{mod}}^{+}(P) \quad (6.2)$$

et ce, pour toute valeur de probabilité P donnée ; c'est cette situation qui est représentée sur la figure 6.2.

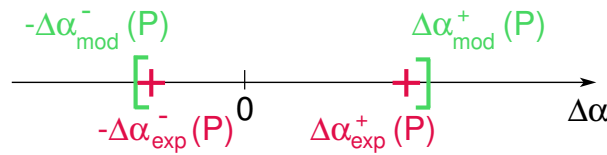


FIG. 6.2 – Comparaison entre les méconnaissances effectives et les valeurs expérimentales à $P\%$

6.1.3 Détermination des méconnaissances de base représentatives des incertitudes de la structure réelle

Cette comparaison entre les méconnaissances effectives et les valeurs expérimentales à $P\%$ doit nous permettre de déterminer les méconnaissances de base les plus représentatives des incertitudes présentes dans les structures réelles.

Pour cela, **une première tentative a consisté à faire en sorte que les méconnaissances effectives soient les plus proches possibles des valeurs expérimentales à P %**, tout en les encadrant, définissant ainsi une sorte de méthode d'identification par résolution d'un problème inverse. Cela revient à supposer dans ce cas que tous les résultats expérimentaux dont on dispose proviennent d'essais qui « voient » tous parfaitement les différentes méconnaissances sur les sous-structures considérées, ce qui n'est guère réaliste en pratique.

On trouvera néanmoins quelques exemples d'application de cette méthode d'identification dans [Puel *et al.* 2003, Barthe *et al.* 2003, Ladevèze *et al.* 2004]. La résolution de ce problème inverse est obtenue par la minimisation d'une fonction coût caractéristique de l'écart entre les valeurs provenant du modèle et de l'expérimental, sous la contrainte d'inclusion des valeurs mesurées par les méconnaissances effectives associées au modèle ; pour cela, la forme de cette fonction coût fait intervenir un certain nombre de quantités d'intérêt α , pour une ou plusieurs valeurs de probabilité P :

$$\mathcal{F}(\bar{m}_E^+, \bar{m}_E^-) = \sum_{\alpha} \left| \frac{\Delta\alpha_{\text{mod}}^- - \Delta\alpha_{\text{exp}}^-}{\Delta\alpha_{\text{exp}}^-} \right| + \sum_{\alpha} \left| \frac{\Delta\alpha_{\text{mod}}^+ - \Delta\alpha_{\text{exp}}^+}{\Delta\alpha_{\text{exp}}^+} \right| \quad (6.3)$$

On suppose bien sûr *a priori* des formes de loi de probabilité (normale, uniforme, ...) pour la méconnaissance de base sur chaque sous-structure, et le problème de détermination des méconnaissances de base se résume alors à trouver pour chaque sous-structure E les valeurs de $(\bar{m}_E^+, \bar{m}_E^-)$ qui minimisent cette fonction coût $\mathcal{F}(\bar{m}_E^+, \bar{m}_E^-)$.

Cette méthode n'a finalement pas été retenue : en effet, d'un point de vue technique, on peut avoir des soucis de précision dans la détermination de méconnaissances de base qui n'ont qu'un faible impact dans les valeurs de méconnaissances effectives calculées, d'autant plus que l'on ne sait pas vraiment comment tirer profit au mieux des données expérimentales disponibles : il est clair qu'utiliser toute l'information expérimentale d'un seul coup n'est guère efficace, d'autant moins que le nombre de sous-structures étudiées est important. Mais même en tentant d'améliorer le processus en sélectionnant les données expérimentales pertinentes, on ne parvient pas vraiment à tirer intelligemment profit de l'information apportée par les données expérimentales, et encore moins de la connaissance que l'on peut avoir *a priori* de la structure.

Enfin, comme on l'a cité plus haut, **une telle méthode suppose que l'essai expérimental traduise parfaitement l'incertitude que l'on cherche à mesurer à l'aide du concept de méconnaissance**, ce qui n'est guère évident en général. Ainsi, pour donner plus de valeur à la notion d'information apportée expérimentalement, et à la question de sa pertinence, on introduit le principe de réduction des méconnaissances de base décrit, en deux temps successifs, dans le paragraphe suivant, puis dans le paragraphe 6.4.

6.2 Réduction des méconnaissances de base : première formulation

L'idée majeure associée au procédé que l'on décrit ici est que **plus on a de données expérimentales, plus on est susceptible de réduire les méconnaissances de base**. Les mesures expérimentales sont autant d'apports d'informations permettant d'accroître de manière intelligente la connaissance disponible à propos de la structure étudiée. Afin de simplifier la présentation du procédé de réduction, **on suppose dans ce paragraphe que les données expérimentales sont parfaitement représentatives du comportement de la structure**, et donc qu'elles permettent de traduire parfaitement les différentes méconnaissances sur les sous-structures. On verra dans un second temps une tentative de prise en compte de la représentativité des mesures expérimentales, dans le paragraphe 6.4.

6.2.1 Description du principe associé au procédé de réduction

Le principe associé à la volonté de réduire les méconnaissances de base sur les différentes sous-structures suppose tout d'abord que **l'on puisse avoir au départ une certaine description, éventuellement grossière, mais nécessairement majorante, de ces diverses méconnaissances**. Ceci peut être obtenu grâce à de la connaissance *a priori* ou à de l'expérience concernant la structure étudiée : ainsi, si l'on connaît bien le matériau, on peut déjà fixer l'ordre de grandeur de l'écart type de la loi normale que l'on est susceptible d'attribuer à la méconnaissance de base de la sous-structure correspondante.

Cette description majorante est caractérisée par la donnée des méconnaissances de base initiales $(m_E^{+0}, m_E^{-0})_{E \in \Omega}$, ou plus exactement des lois de probabilité qui les modélisent ; il est d'ailleurs crucial de vérifier que cette modélisation est admissible, car sinon la réduction ne sera pas possible : pour cela, on réalise un calcul de propagation des méconnaissances de base initiales, et on vérifie que les valeurs expérimentales à P % sont toutes incluses dans les intervalles correspondants de méconnaissances effectives.

Le procédé de réduction consiste à utiliser l'apport d'informations expérimentales pertinentes pour réduire le niveau de méconnaissance sur une sous-structure à la fois. Une façon de juger de la pertinence des mesures effectuées consiste à observer, pour chaque mode $\bar{\phi}_i$, quelle est la contribution de chaque sous-structure dans l'énergie de déformation totale du modèle théorique déterministe, associée au mode en question ; en effet, toutes les expressions qui permettent de calculer les méconnaissances effectives sur une quantité d'intérêt, à savoir les équations (5.6), (5.10) et (5.15), font intervenir plus ou moins directement ces énergies.

Considérons par exemple que c'est la sous-structure E^* qui est fortement concernée par l'apport d'information expérimentale représenté par la quantité d'intérêt α . Il s'agit de déterminer une méconnaissance de base $(m_{E^*}^-, m_{E^*}^+)$ plus « petite », c'est-à-dire dont l'étendue des lois de probabilité utilisées est moins grande que celle des lois qui modélisent la méconnaissance initiale $(m_{E^*}^{-0}, m_{E^*}^{+0})$; pendant ce temps, les méconnaissances $(m_E^{+0}, m_E^{-0})_{E \neq E^*}$ sur les autres sous-structures n'évoluent pas.

Cette volonté de réduire les méconnaissances se traduit plus précisément en termes de probabilités d'intervalle par l'inégalité suivante :

$$P_{E^*}^0(L) \leq P_{E^*}(L) \quad \forall L \quad (6.4)$$

ou encore en termes d'intervalles standards de méconnaissance effective, pour toute valeur P de probabilité, par :

$$\Delta\alpha_{\text{mod}}^-(P) \leq \Delta\alpha_{\text{mod}}^{-0}(P) \quad \text{et} \quad \Delta\alpha_{\text{mod}}^+(P) \leq \Delta\alpha_{\text{mod}}^{+0}(P) \quad (6.5)$$

L'information apportée par les valeurs expérimentales à P % intervient dans le processus de réduction au travers de l'application de la condition (6.1), associée aux quantités $\Delta\alpha_{\text{exp}}^-(P)$ et $\Delta\alpha_{\text{exp}}^+(P)$, qui s'écrit pour tout P :

$$\Delta\alpha_{\text{exp}}^-(P) \leq \Delta\alpha_{\text{mod}}^-(P) \quad \text{et} \quad \Delta\alpha_{\text{exp}}^+(P) \leq \Delta\alpha_{\text{mod}}^+(P) \quad (6.6)$$

On peut résumer notre démarche par la volonté de diminuer au maximum, sous-structure par sous-structure, la dispersion des lois de probabilité associées *a priori* aux méconnaissances de base, sous le contrôle offert par les données expérimentales proposées ; la figure 6.3 illustre l'ensemble du processus de réduction. Écrit en termes de probabilités d'intervalles, ce problème de minimisation devient un problème de maximisation sous la contrainte (6.1) :

$$\max_{P(\Delta\alpha_{\text{exp}} \in I_{\Delta\alpha_{\text{mod}}}(L)) \geq P_{\alpha}(L)} P_{E^*}(L, (\mathbf{m}_{E^*}^-, \mathbf{m}_{E^*}^+)) \quad \forall L \quad (6.7)$$

Cette philosophie est proche des idées de prise en compte de l'information nouvelle telles qu'on les a décrites dans le chapitre 2, et dont la version la plus simple se traduit par l'interprétation du théorème de Bayes donnée dans le paragraphe 2.2.2 : l'apport d'information peut être un moyen efficace pour passer d'un état de connaissance initial *a priori* à un nouvel état de connaissance *a posteriori*.

6.2.2 Prise en compte des cas les plus défavorables

Comme la réduction des méconnaissances de base se fait sous-structure par sous-structure, la seule vérification des contraintes (6.6) n'est pas suffisante, car elle ne permet pas d'être réaliste vis-à-vis de toutes les situations possibles. En effet, le couple de valeurs $\Delta\alpha_{\text{mod}}^+$ et $\Delta\alpha_{\text{mod}}^-$ constituant la méconnaissance effective est issu de contributions de la part de toutes les sous-structures :

$$\Delta\alpha_{\text{mod}}^+ = \Delta\alpha_{E^*}^+ + \sum_{E \neq E^*} \Delta\alpha_E^+ \quad (6.8a)$$

$$\Delta\alpha_{\text{mod}}^- = \Delta\alpha_{E^*}^- + \sum_{E \neq E^*} \Delta\alpha_E^- \quad (6.8b)$$

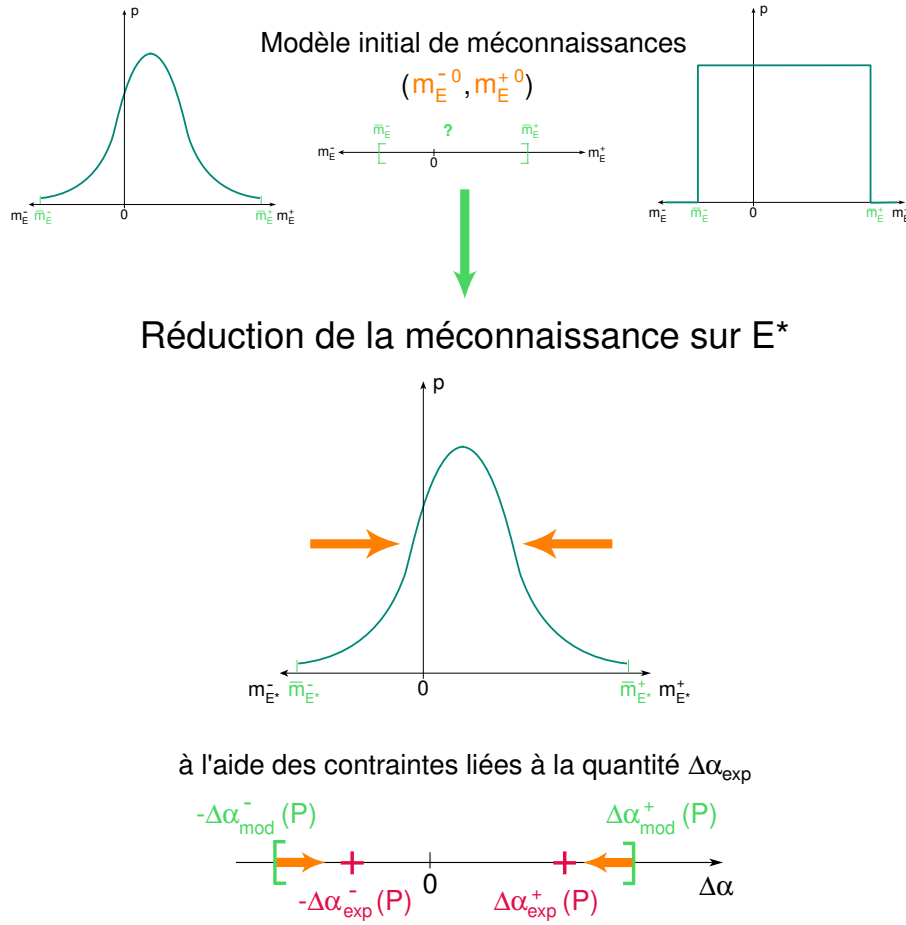


FIG. 6.3 – Principe de la réduction des méconnaissances de base à l'aide d'informations expérimentales pertinentes

Ainsi, si l'on ne s'en tient qu'aux contraintes (6.6), la réduction de la méconnaissance sur E^* peut tout à fait être faussée par des niveaux de méconnaissance initiaux très surestimés sur les autres sous-structures $E \neq E^*$, d'autant plus si ces dernières ont une contribution non négligeable dans l'énergie de déformation associée au mode qui sert à réduire la méconnaissance sur E^* .

Pour éviter cette situation probable, on envisage de tenir compte des cas de figures les plus défavorables en ce qui concerne toutes les autres sous-structures :

$$\Delta\alpha_{\text{mod}}^{+\text{pire}} = \Delta\alpha_{E^*}^{+} + \sum_{E \neq E^*} \Delta\alpha_E^{+\text{pire}} \quad (6.9a)$$

$$\Delta\alpha_{\text{mod}}^{-\text{pire}} = \Delta\alpha_{E^*}^{-} + \sum_{E \neq E^*} \Delta\alpha_E^{-\text{pire}} \quad (6.9b)$$

La définition du « cas pire » lors de la propagation d'un jeu de réalisations de méconnaissances de base est le suivant : tandis que l'intervalle dont la méconnaissance est en cours de réduction est propagé classiquement à l'aide des relations du paragraphe 5.2, on ne retient de chacun des autres intervalles de méconnaissances de base que la valeur, comprise dans l'intervalle, la plus défavorable dans l'évaluation des quantités $\Delta\alpha_{\text{mod}}^{+}$ et $\Delta\alpha_{\text{mod}}^{-}$.

Par exemple, dans le cas de la méconnaissance effective sur le carré d'une pulsation propre, pour chacune des sous-structures $E \neq E^*$, c'est la borne inférieure de l'intervalle de méconnaissance de base qui constitue le cas le plus défavorable pour l'évaluation de la borne supérieure sur la quantité d'intérêt, et vice versa, ce qui s'écrit :

$$\Delta\alpha_E^{+ \text{pire}} = -2m_E^- \bar{e}_E(\bar{\phi}_i) \quad (6.10a)$$

$$\Delta\alpha_E^{- \text{pire}} = -2m_E^+ \bar{e}_E(\bar{\phi}_i) \quad (6.10b)$$

On a également représenté sur la figure 6.4 le principe d'une telle prise en compte pour un tirage donné de méconnaissances de base sur les différentes sous-structures.

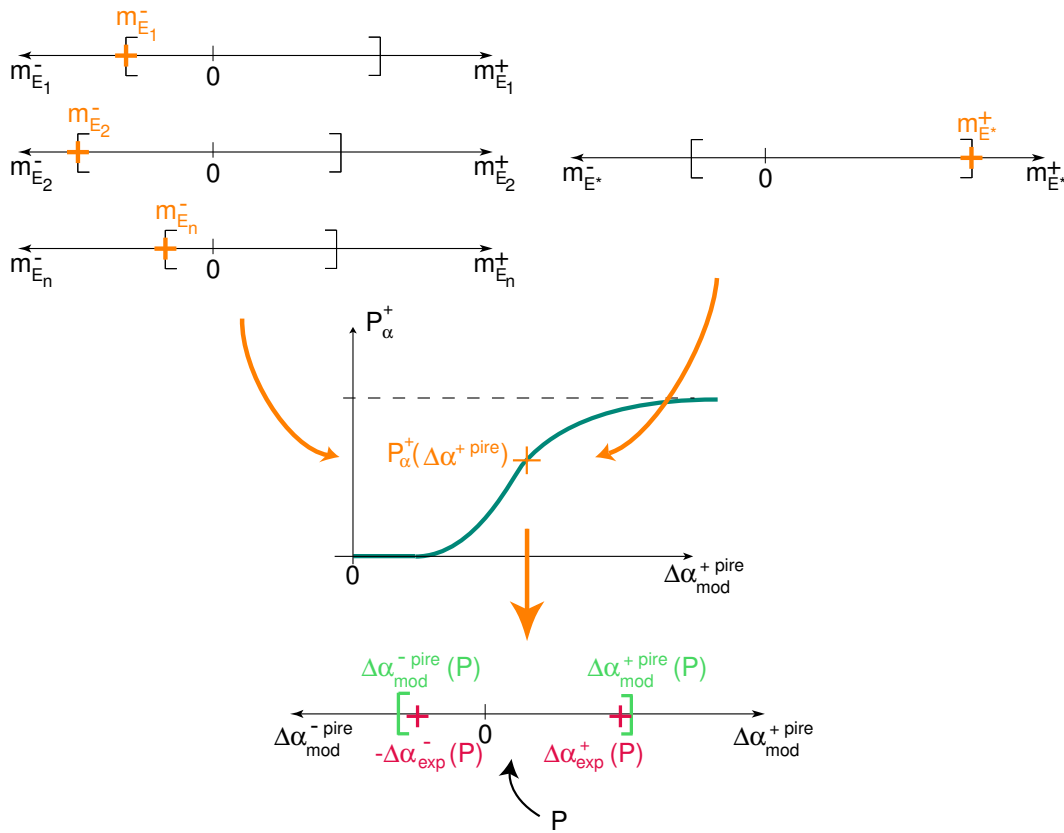


FIG. 6.4 – Prise en compte des cas les plus défavorables dans le processus de réduction des méconnaissances de base

À partir des répartitions des bornes $\Delta\alpha_{\text{mod}}^{+ \text{pire}}$ et $\Delta\alpha_{\text{mod}}^{- \text{pire}}$, on peut déduire une nouvelle probabilité d'intervalle $P_{\alpha^{\text{pire}}}(L)$ prenant en compte l'ensemble des cas défavorables, et en déduire les deux bornes $\Delta\alpha_{\text{mod}}^{+ \text{pire}}(P)$ et $\Delta\alpha_{\text{mod}}^{- \text{pire}}(P)$ de l'intervalle standard $I_{\Delta\alpha_{\text{mod}}^{\text{pire}}}(L_P)$ associé à une valeur de probabilité P donnée. Finalement, de nouvelles contraintes, représentées sur la figure 6.4, doivent être prises en compte :

$$\Delta\alpha_{\text{exp}}^-(P) \leq \Delta\alpha_{\text{mod}}^{- \text{pire}}(P) \quad \text{et} \quad \Delta\alpha_{\text{exp}}^+(P) \leq \Delta\alpha_{\text{mod}}^{+ \text{pire}}(P) \quad (6.11)$$

Ces dernières conditions garantissent une certaine sécurité dans le processus de réduction des méconnaissances de base, quitte à obtenir des valeurs un peu trop conservatives : elles font en sorte que la méconnaissance de base réduite puisse conduire à des résultats raisonnables. Elles permettent surtout d'éviter une erreur grossière de choix de données expérimentales à utiliser, ou encore d'avoir une assez bonne réduction dans le cas où il n'y a pas de donnée expérimentale totalement discriminante pour la réduction sur la sous-structure E^* , ce qui est une situation que l'on peut rencontrer selon les structures étudiées. Dans ce dernier cas, après une première réduction des méconnaissances de chaque sous-structure, une nouvelle itération est possible, et ainsi de suite, comme on aura l'occasion de le voir dans l'exemple suivant.

6.3 Réduction des méconnaissances de base sur un exemple académique

6.3.1 Structure étudiée

La structure analysée dans cette partie est un treillis plan dérivé de celui présenté dans le paragraphe 3.2 : il est composé de six barres rectilignes, constituées de différents matériaux, et reliées par des rotules, comme représenté sur la figure 6.5.

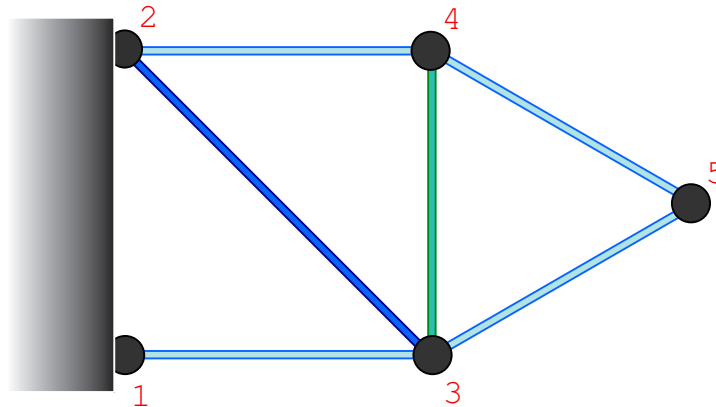


FIG. 6.5 – Treillis plan à six barres

Le modèle théorique déterministe associé fait l'hypothèse de liaisons rotules parfaites, de sollicitations de traction-compression dans toutes les barres et de matériaux homogènes et isotropes. Les caractéristiques de ce modèle sont répertoriées dans le tableau 6.1, classées par groupes de mêmes matériaux. Le problème aux valeurs propres associé au modèle déterministe permet de déterminer six modes propres de vibrations, qui sont représentés sur la figure 6.6.

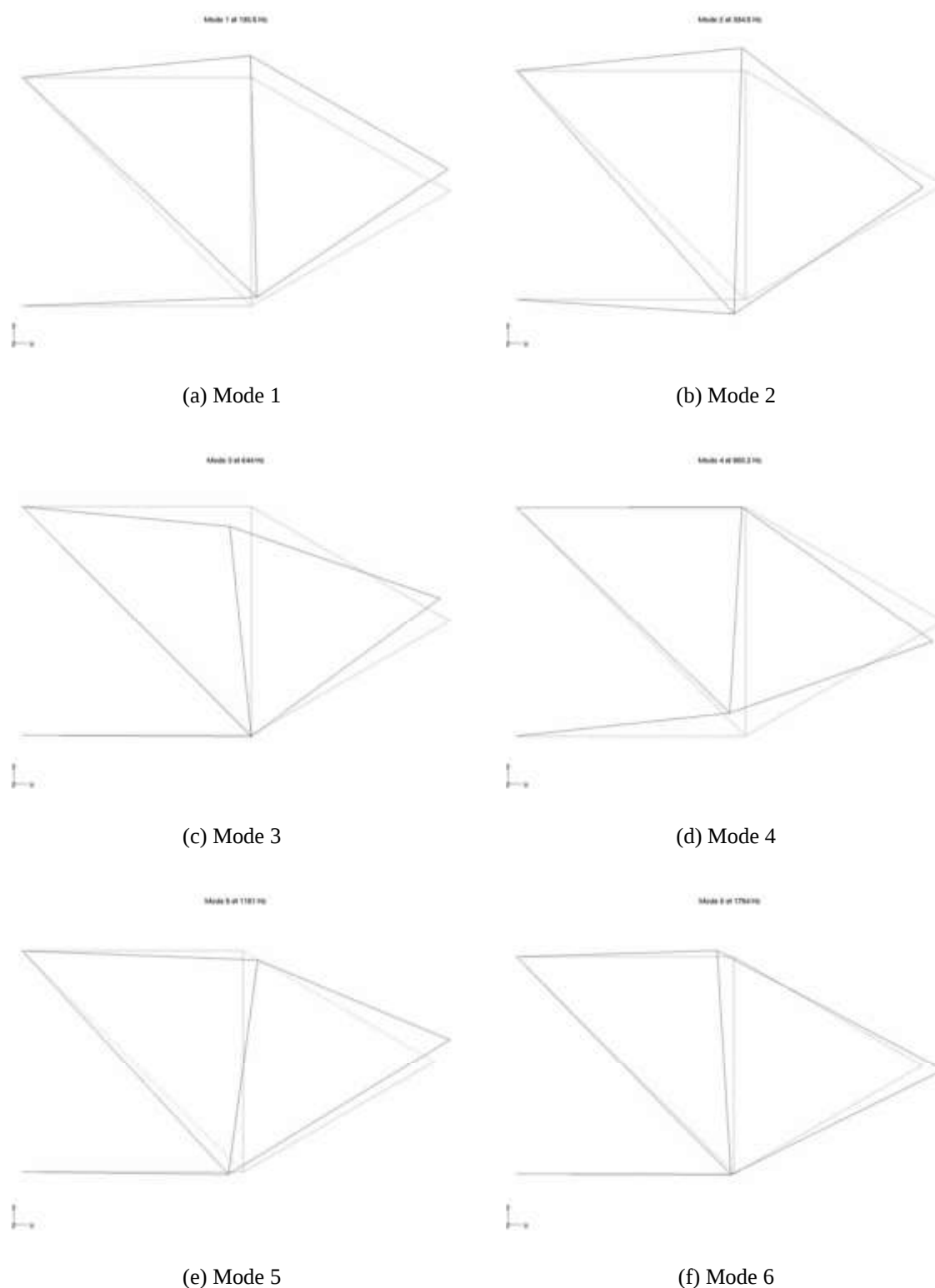


FIG. 6.6 – Modes propres du treillis à six barres

Groupes	Section	Longueur	Matériau	Module d'Young	Masse vol.
$g1$	10^{-4}m	1m	aluminium	$\bar{E}_{g1} = 72\text{GPa}$	2700kg/m^3
$g2 = \{b23\}$	10^{-4}m	$\sqrt{2}\text{m}$	acier	$\bar{E}_{g2} = 210\text{GPa}$	7800kg/m^3
$g3 = \{b34\}$	10^{-4}m	$\sqrt{2}\text{m}$	« X »	$\bar{E}_{g3} = 10\text{GPa}$	1500kg/m^3

NB : Le groupe $g1$ est composé des barres $\{1-3,2-4,3-5,4-5\}$.

TAB. 6.1 – Caractéristiques du treillis à six barres

6.3.2 Mise en œuvre du procédé de réduction

Données expérimentales

L'exemple présenté ici étant purement académique, les données « expérimentales » sont obtenues par une simulation stochastique de la famille des structures considérées, définissant ainsi une réalité supposée parfaitement connue.

Des dispersions de raideur correspondant à des lois de probabilité précises sont introduites dans le modèle théorique déterministe ; elles sont détaillées dans le tableau 6.2, où l'on donne à la fois l'étendue et les deux premiers moments statistiques exprimés relativement aux modules d'Young du modèle déterministe. Les lois choisies pour les barres en acier et en aluminium sont des lois normales, mais bornées pour éviter toute réalisation non physique. On considère que le matériau intitulé « X » est très mal connu ; ainsi, la loi de probabilité choisie est de type uniforme, et d'étendue plus importante que les deux autres, le but étant de tester le processus de réduction avec différents types de lois.

Barres	Modèle déterministe	Loi choisie	Domaine relatif	Moments statistiques relatifs (μ : moyenne / σ : écart type)
$\{1-3,2-4,3-5,4-5\}$	$\bar{E}_{g1} = 72\text{GPa}$	normale	$[-0,05; 0,05]$	$\mu = 0,00 / \sigma = 0,019$
2-3	$\bar{E}_{g2} = 210\text{GPa}$	normale	$[-0,15; 0,05]$	$\mu = -0,05 / \sigma = 0,039$
3-4	$\bar{E}_{g3} = 10\text{GPa}$	uniforme	$[-0,10; 0,20]$	$\mu = 0,05 / \sigma = 0,087$

NB : Les caractéristiques des lois sont données relativement au module d'Young du modèle déterministe.

TAB. 6.2 – Caractéristiques relatives des lois de probabilité permettant la simulation des raideurs « expérimentales » des six barres du treillis

Pour chacun de ces treillis « expérimentaux » simulés, on est capable de calculer les fréquences propres et modes propres, et donc de déterminer la répartition des quantités d'intérêt expérimentales. Le calcul est mené à l'aide d'une technique de Monte Carlo mettant en œuvre 50000 tirages. Les informations « expérimentales » retenues sont les valeurs à $P\%$ $\Delta\omega_{i\text{exp}}^{2+}(P)$ et $\Delta\omega_{i\text{exp}}^{2-}(P)$, ainsi que $\Delta\phi_{ki\text{exp}}^{+}(P)$ et $\Delta\phi_{ki\text{exp}}^{-}(P)$, qui encadrent $P\%$ des valeurs « expérimentales » de pulsations propres et modes propres ; en particulier, on s'intéresse à des probabilités $P = 0,99$ et $P = 0,70$, afin de caractériser respectivement l'étendue et un ordre de grandeur de l'écart type des quantités d'intérêt expérimentales. On pourrait bien sûr envisager une description encore plus riche en utilisant d'autres valeurs de probabilité P .

Choix du modèle initial de méconnaissances

Le modèle initial de méconnaissances qui complète le modèle théorique déterministe est précisé dans le tableau 6.3 ; encore une fois, pour plus de clarté, on précise à la fois le domaine d'étendue de chaque loi, et les deux premiers moments statistiques associés, définis relativement aux valeurs du modèle théorique déterministe. On a décidé de réunir dans un même groupe les quatre barres en aluminium, ce qui va dans le sens de la globalisation effectuée à travers le concept de méconnaissance, et permet de plus de réduire encore les temps de calcul.

Groupes	Loi choisie	Domaine $[-\overline{m}_E^{-0}; \overline{m}_E^{+0}]$	Moments statistiques relatifs (μ : moyenne / σ : écart type)
$g1$	normale	$[-0, 50; 0, 50]$	$\mu = 0, 00 / \sigma = 0, 194$
$g2 = \{b23\}$	normale	$[-0, 50; 0, 50]$	$\mu = 0, 00 / \sigma = 0, 194$
$g3 = \{b34\}$	uniforme	$[-0, 50; 0, 50]$	$\mu = 0, 00 / \sigma = 0, 289$

NB : Le groupe $g1$ est composé des barres $\{1-3, 2-4, 3-5, 4-5\}$.

TAB. 6.3 – Modèle initial ($\overline{m}_E^{-0}$, $\overline{m}_E^{+0}$) de méconnaissances de base pour le treillis à six barres

La description retenue est largement majorante, fait que l'on peut facilement vérifier en calculant les méconnaissances effectives associées, et en constatant qu'elles encadrent très largement les quantités expérimentales. Pour le moment, on a supposé dans le modèle initial de méconnaissances des lois de probabilité du même type que celles que l'on a introduites dans la simulation stochastique des structures prises comme références « expérimentales » ; ceci nous permettra de réaliser plus facilement des comparaisons avec la référence expérimentale après l'application du processus de réduction. On discutera plus loin dans ce paragraphe de l'influence du choix du type de loi de probabilité.

Processus de réduction

Compte tenu du modèle initial de méconnaissances et des données expérimentales, il est important de sélectionner quelles sont les mesures les plus pertinentes pour réaliser les réductions successives sur les différentes sous-structures.

Une méthode efficace consiste à s'appuyer sur le fait que **la sensibilité des méconnaissances effectives vis-à-vis des méconnaissances de base est directement liée aux énergies de déformation modales du modèle théorique déterministe**. En effet, les expressions des bornes $\Delta\alpha_{\text{mod}}^{+}$ et $\Delta\alpha_{\text{mod}}^{-}$ encadrant la quantité $\Delta\alpha_{\text{mod}}$ relative au modèle avec méconnaissances sont des combinaisons linéaires des méconnaissances de base, dont les coefficients sont des énergies $\overline{e}_E$, comme on peut le constater dans les expressions (5.6), (5.10) et (5.15).

Ainsi, **les informations expérimentales les plus pertinentes dans le processus de réduction de la méconnaissance de base sur la structure E^* sont celles pour lesquelles l'énergie de déformation modale se trouve majoritairement dans E^*** . Par exemple, pour ce qui concerne les pulsations propres, nous avons besoin de considérer les énergies $\overline{e}_E(\underline{\phi}_i)$: en effet, la méconnaissance effective sur une pulsation propre est déterminée à

partir des distributions des bornes suivantes :

$$\Delta\omega_{i\text{mod}}^{2-} = 2 \sum_{E \in \Omega} m_E^- \bar{e}_E(\bar{\phi}_i) \quad (6.12a)$$

$$\Delta\omega_{i\text{mod}}^{2+} = 2 \sum_{E \in \Omega} m_E^+ \bar{e}_E(\bar{\phi}_i) \quad (6.12b)$$

et on constate que ces deux valeurs sont des combinaisons linéaires des méconnaissances de base, avec comme coefficients de linéarité les énergies de déformation modales $\bar{e}_E(\bar{\phi}_i)$. Les parts respectives de ces dernières sont représentées sur la figure 6.7.

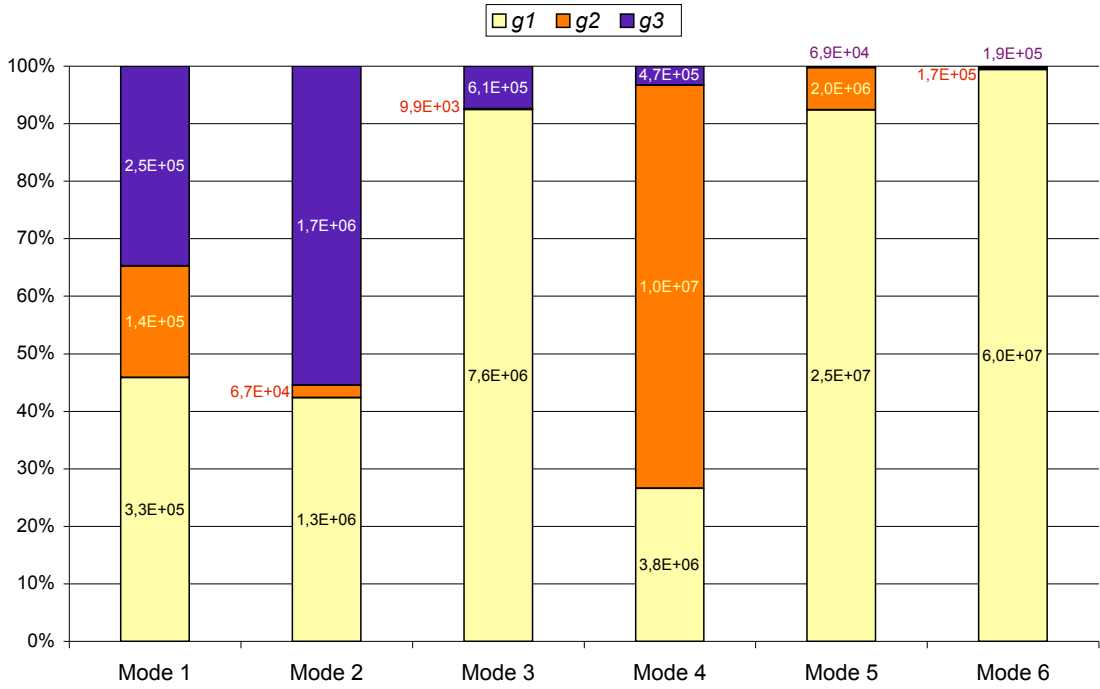


FIG. 6.7 – Énergies de déformation modales par groupe de barres pour les modes du treillis à six barres

L'observation des énergies nous permet de définir un ordre pour les réductions successives, qui sont exposées dans le tableau 6.4. Il est important de souligner qu'il s'agit véritablement d'une prise en compte successive de l'information expérimentale : ainsi, après la réduction de la méconnaissance sur le groupe $g1$, on se sert de l'expression de la nouvelle méconnaissance de base sur ce groupe dans la réduction de la méconnaissance sur le groupe $g2$.

On peut comparer les méconnaissances obtenues après réduction, répertoriées dans le tableau 6.4, directement avec les lois de probabilité définissant la simulation stochastique de la famille de structures étudiées ; ces lois sont indiquées dans le tableau 6.2. En effet, dans ce cas simple, les seules incertitudes « expérimentales » portent sur les valeurs de raideur des barres, et les dispersions relatives qui définissent la simulation stochastique sont les mêmes que celles concernant les énergies de déformation.

De manière générale, les méconnaissances obtenues après réduction sont plus grandes que ce que l'on était en droit d'attendre vis-à-vis des dispersions de raideur expérimentales. Ce n'est par contre pas le cas de la méconnaissance sur le groupe $g1$, qui fait intervenir les quatre barres en aluminium simultanément dans le modèle avec méconnaissances : ceci a donc tendance soit à rigidifier, soit à assouplir le comportement global de la structure, comme on l'a déjà remarqué avec les résultats du tableau 5.4, et conduit ainsi à l'obtention de valeurs légèrement sous-estimées.

Gr.	Données expérimentales	Loi	Méconnaissances après réduction	
			$[-\overline{m}_E^-; \overline{m}_E^+]$	Moments statistiques relatifs
$g1$	Mode 6	normale	$[-0,041; 0,042]$	$\mu = 0,001 / \sigma = 0,016$
$g2$	Mode 4	normale	$[-0,172; 0,075]$	$\mu = -0,049 / \sigma = 0,048$
$g3$	Mode 2	uniforme	$[-0,121; 0,214]$	$\mu = 0,047 / \sigma = 0,097$

NB : Les données expérimentales sont composées pour chaque mode de $\Delta\omega_{i\text{exp}}^{2-}(P)$, $\Delta\omega_{i\text{exp}}^{2+}(P)$, $\Delta\phi_{ki\text{exp}}^-(P)$, $\Delta\phi_{ki\text{exp}}^+(P)$ pour $P = 0,70$ et $P = 0,99$

TAB. 6.4 – Résultats après une première réduction des méconnaissances de base du treillis à six barres

Le procédé peut être itératif : après avoir réduit la méconnaissance sur chaque groupe, on peut recommencer la réduction sur le groupe $g1$, sachant que la contribution des autres groupes a diminué entre-temps ; dans ce cas, on ne constate aucune amélioration. Toutefois, si, par exemple, on avait commencé l'ensemble du processus par la réduction de la méconnaissance sur le groupe $g1$ à l'aide du premier mode, cette dernière aurait été bien moins réduite vu les contributions non négligeables des deux autres groupes ; dans ce cas, après la réduction sur les trois groupes, une deuxième itération nous aurait montré qu'il était encore possible de réduire la méconnaissance sur le groupe $g1$.

Ainsi, les méconnaissances que l'on obtient après le processus de réduction sont encore forcément surestimées, même après plusieurs itérations : en effet, les contraintes (6.11) introduites font intervenir les cas les plus défavorables. Ceci nous avait permis avant tout de ne pas obtenir de résultats erronés à la suite d'un mauvais choix d'information expérimentale, et donc d'obtenir des méconnaissances conduisant à des résultats raisonnables dans d'autres situations expérimentales.

Si l'on considère que l'ensemble des données « expérimentales » que l'on a utilisées décrit la totalité des situations envisageables pour la famille de structures, on peut essayer de réduire encore le niveau des méconnaissances de base obtenues, en imposant cette fois-ci les contraintes (6.6) définies initialement. Les résultats du tableau 6.5 montrent que les méconnaissances peuvent être réduites quelque peu, se rapprochant ainsi de ce qu'on pouvait en attendre compte tenu de la référence « expérimentale » représentée par la simulation stochastique décrite dans le tableau 6.2. Une nouvelle itération ne fait plus évoluer les niveaux de méconnaissances de base : on peut donc estimer avoir tiré profit au maximum de l'information expérimentale disponible.

Gr.	Données expérimentales	Loi	Méconnaissances après réduction	
			$[-\bar{m}_E^-; \bar{m}_E^+]$	Moments statistiques relatifs
$g1$	Mode 6	normale	$[-0,039; 0,041]$	$\mu = 0,001 / \sigma = 0,016$
$g2$	Mode 4	normale	$[-0,158; 0,058]$	$\mu = -0,050 / \sigma = 0,042$
$g3$	Mode 2	uniforme	$[-0,108; 0,203]$	$\mu = 0,048 / \sigma = 0,090$

NB : Les données expérimentales sont composées pour chaque mode de $\Delta\omega_{i\text{exp}}^{2-}(P)$, $\Delta\omega_{i\text{exp}}^{2+}(P)$, $\Delta\phi_{ki\text{exp}}^-(P)$, $\Delta\phi_{ki\text{exp}}^+(P)$ pour $P = 0,70$ et $P = 0,99$

TAB. 6.5 – Résultats finaux après une seconde réduction des méconnaissances de base du treillis à six barres

Capacité de prédiction

Avec les méconnaissances de base que l'on vient juste de déterminer, on peut calculer les méconnaissances effectives pour les trois modes que l'on n'a pas encore utilisés afin d'évaluer la qualité des résultats de notre procédé de réduction en comparant les méconnaissances effectives associées au modèle de méconnaissances réduites avec les valeurs à P % issues de la simulation stochastique prise comme référence « expérimentale ». La confrontation est réalisée sur les valeurs à 70 % et à 99 % dans les tableaux 6.6 et 6.7.

Valeurs de P	Modes i	Expérimental $[\omega_{i\text{exp}}^{2-}; \omega_{i\text{exp}}^{2+}]$	Modèle avec méconnaissances $[\bar{\omega}_i^2 - \Delta\omega_{i\text{mod}}^{2-}; \bar{\omega}_i^2 + \Delta\omega_{i\text{mod}}^{2+}]$
0,70	1	$[1,39; 1,49].10^6$	$[1,38; 1,47].10^6$
	3	$[1,62; 1,67].10^7$	$[1,62; 1,66].10^7$
	5	$[5,41; 5,56].10^7$	$[5,44; 5,56].10^7$
0,99	1	$[1,35; 1,53].10^6$	$[1,35; 1,54].10^6$
	3	$[1,58; 1,71].10^7$	$[1,58; 1,71].10^7$
	5	$[5,30; 5,68].10^7$	$[5,30; 5,70].10^7$

TAB. 6.6 – Comparaison modèle-expérimental pour les pulsations propres du treillis à six barres

Valeurs de P	Modes i	DDL k	Expérimental $[\phi_{ki\text{exp}}^-; \phi_{ki\text{exp}}^+]$	Modèle avec méconnaissances $[\bar{\phi}_{ki} - \Delta\phi_{ki\text{mod}}^-; \bar{\phi}_{ki} + \Delta\phi_{ki\text{mod}}^+]$
0,70	1	4	$[0,89; 0,96]$	$[0,92; 0,99]$
	3	3	$[-0,96; -0,93]$	$[-0,98; -0,92]$
	5	1	$[-0,69; -0,65]$	$[-0,70; -0,65]$
0,99	1	4	$[0,88; 0,99]$	$[0,86; 1,00]$
	3	3	$[-0,98; -0,91]$	$[-0,99; -0,90]$
	5	1	$[-0,72; -0,62]$	$[-0,73; -0,62]$

TAB. 6.7 – Comparaison modèle-expérimental pour les modes propres du treillis à six barres

Les résultats obtenus sont plutôt bons : on a généralement moins de 0,5 % d'erreur relative sur les quantités liées aux pulsations propres, et moins de 1 % en ce qui concerne les modes propres. On note toujours la moins bonne qualité de l'estimation des valeurs

à 70 % par rapport à celle des valeurs à 99 %, et la légère surestimation des méconnaissances effectives sur les modes propres. **Dans l'ensemble, on peut conclure sur la bonne détermination des méconnaissances de base vis-à-vis de la réalité mesurée.**

Quelques commentaires sur le choix des lois de probabilité

La bonne qualité des résultats précédents provient en partie du fait que l'on a choisi les « bonnes » lois de probabilité dans la modélisation des méconnaissances de base, en accord avec les dispersions introduites dans la simulation stochastique.

Pour pouvoir juger de l'influence d'un tel choix, nous commençons par réaliser une autre simulation stochastique avec des dispersions de raideurs simulées par des densités de probabilité uniformes, de mêmes étendues que celles du tableau 6.2. Les valeurs à P % obtenues sont comparées à celles issues de notre simulation initiale dans le tableau 6.8.

$[\omega_{i \text{ exp}}^{2-}(P); \omega_{i \text{ exp}}^{2+}(P)]$	i	Simulation initiale	Cas nouveau
$P = 0,70$	1	$[1, 39; 1, 49].10^6$	$[1, 39; 1, 49].10^6$
	2	$[6, 00; 6, 69].10^6$	$[5, 95; 6, 65].10^6$
	3	$[1, 62; 1, 67].10^7$	$[1, 60; 1, 68].10^7$
	4	$[2, 74; 2, 92].10^7$	$[2, 70; 2, 99].10^7$
	5	$[5, 41; 5, 56].10^7$	$[5, 37; 5, 62].10^7$
	6	$[1, 20; 1, 23].10^7$	$[1, 19; 1, 24].10^7$
$P = 0,99$	1	$[1, 34; 1, 53].10^6$	$[1, 33; 1, 53].10^6$
	2	$[5, 77; 6, 84].10^6$	$[5, 74; 6, 87].10^6$
	3	$[1, 58; 1, 71].10^7$	$[1, 57; 1, 72].10^7$
	4	$[2, 61; 3, 03].10^7$	$[2, 60; 3, 04].10^7$
	5	$[5, 30; 5, 68].10^7$	$[5, 26; 5, 70].10^7$
	6	$[1, 17; 1, 26].10^7$	$[1, 16; 1, 27].10^7$

TAB. 6.8 – Comparaison des valeurs à P % pour deux familles de structures simulées par des répartitions de raideur de lois de probabilité de natures différentes

La confrontation entre ces simulations de deux réalités expérimentales proches mène aux réflexions suivantes : pour pouvoir déterminer à l'aide des données expérimentales disponibles quelles sont les « bonnes » lois de probabilité pour la modélisation des méconnaissances de base, il faudrait pouvoir juger précisément de l'écart résiduel entre les différentes valeurs expérimentales à P % et les méconnaissances effectives correspondantes lors de la réduction sur une sous-structure donnée ; en effet, cet écart résiduel doit être plus petit lorsque l'on a choisi la « bonne » forme de la loi de probabilité que pour toute autre forme de loi.

Malheureusement, les résultats du tableau 6.8 montrent que l'écart relatif entre les valeurs à P % de l'une et l'autre simulation est souvent inférieur à 1 % ; comme, en pratique, les méconnaissances effectives que l'on calcule sont entachées d'une erreur d'un ordre de grandeur à peine inférieur, **il semble difficile de pouvoir conclure avec certitude sur la pertinence du type de loi choisi a priori.**

De plus, notre comparaison dans le tableau 6.8 a été réalisée à l'aide de 50000 structures simulées par tirage de Monte Carlo dans chaque cas d'étude, ce qui nous a permis d'obtenir des valeurs expérimentales à P % très précises à partir de la réalité que l'on voulait simuler à chaque fois ; malgré cela, les écarts sont peu importants. Il est alors clair qu'en pratique, les données expérimentales étant issues d'un nombre peu important de structures réelles (une dizaine tout au plus), les valeurs à P % déduites de ces dernières ne sont vraisemblablement pas exactement représentatives de la « vraie » dispersion expérimentale, les évaluations statistiques risquant d'être faussées par le faible nombre de structures testées. **Il est donc à coup sûr illusoire de vouloir trancher pour une forme de loi de probabilité particulière à partir de la donnée de ces quantités expérimentales.**

Finalement, **le choix le plus raisonnable consiste plutôt à utiliser la connaissance *a priori* de la structure**, en supposant une loi normale quand on fait l'hypothèse de dispersions matériaux uniquement, une loi uniforme si l'on n'est pas sûr des mécanismes d'incertitudes entrant en jeu dans la sous-structure, et enfin pas de loi de probabilité du tout quand on estime que le modèle est très simplifié par rapport à la réalité.

6.3.3 Cas d'une forte méconnaissance

On illustre ici le processus de réduction des méconnaissances de base dans le cas d'une « forte méconnaissance » qui nous amène à utiliser les expressions données dans le paragraphe 5.4.1.

Données expérimentales

Également ici, les données « expérimentales » sont obtenues à l'aide d'une simulation stochastique en introduisant dans le modèle théorique déterministe des dispersions de raideur correspondant à des lois de probabilité précises. Ces dernières sont détaillées dans le tableau 6.9. On a gardé les mêmes types de lois que dans le cas précédent, mais les caractéristiques ont changé ; une très forte dispersion a été supposée pour la barre b_{23} de façon à mettre en évidence le phénomène de la prise en compte d'une forte méconnaissance.

Barres	Modèle déterministe	Loi choisie	Domaine relatif	Moments statistiques relatifs (μ : moyenne / σ : écart type)
{1-3,2-4,3-5,4-5}	$\bar{E}_{g1} = 72\text{GPa}$	normale	$[-0,05; 0,05]$	$\mu = 0,00 / \sigma = 0,019$
2-3	$\bar{E}_{g2} = 210\text{GPa}$	normale	$[-0,20; 0,70]$	$\mu = 0,25 / \sigma = 0,175$
3-4	$\bar{E}_{g3} = 10\text{GPa}$	uniforme	$[-0,15; 0,05]$	$\mu = -0,05 / \sigma = 0,039$

NB : Les caractéristiques des lois sont données relativement au module d'Young du modèle déterministe.

TAB. 6.9 – Caractéristiques relatives des lois de probabilité utilisées dans la simulation stochastique du treillis à six barres dans le cas d'une forte méconnaissance

La simulation stochastique nous permet de déterminer la répartition des quantités d'intérêt associées à la réalité simulée. On garde comme informations « expérimentales » les valeurs $\Delta\omega_{i\text{exp}}^{2+}(0,99)$ et $\Delta\omega_{i\text{exp}}^{2-}(0,99)$ qui encadrent 99 % des valeurs expérimentales des fréquences propres mesurées.

Modèle initial de méconnaissances

Le modèle initial de méconnaissances est précisé dans le tableau 6.10. Les caractéristiques définies permettent d'avoir des méconnaissances effectives qui encadrent toutes les valeurs expérimentales à P % disponibles.

Groupes	Loi choisie	Domaine $[-\overline{m}_E^{-0}; \overline{m}_E^{+0}]$	Moments statistiques relatifs (μ : moyenne / σ : écart type)
$g1$	normale	$[-0, 25; 0, 25]$	$\mu = 0, 00 / \sigma = 0, 097$
$g2 = \{b23\}$	normale	$[-0, 75; 0, 75]$	$\mu = 0, 00 / \sigma = 0, 289$
$g3 = \{b34\}$	uniforme	$[-0, 25; 0, 25]$	$\mu = 0, 00 / \sigma = 0, 097$

NB : Le groupe $g1$ est composé des barres $\{1-3, 2-4, 3-5, 4-5\}$.

TAB. 6.10 – Modèle initial ($\overline{m}_E^{-0}$, $\overline{m}_E^{+0}$) de méconnaissances de base pour le treillis à six barres dans le cas d'une forte méconnaissance

Processus de réduction

Notre critère de sélection des données expérimentales pertinentes est basé sur les contributions des différents groupes dans les énergies de déformation modales de la structure, qui, dans le cas du calcul des méconnaissances effectives en présence d'une forte méconnaissance, dépendent des réalisations de cette dernière ; toutefois, malgré ces variations, la contribution relative de chaque groupe reste à peu près constante, et on continue donc de se baser sur les rapports entre les différentes énergies de déformation du modèle déterministe, détaillées sur la figure 6.7.

Le processus de réduction est tout d'abord mené sans traiter de façon spécifique la forte méconnaissance qui concerne le groupe $g2$; les réductions successives sont mentionnées dans le tableau 6.11. Une seule itération sur tous les groupes a été réalisée.

Gr.	Données expérimentales	Loi	Méconnaissances après réduction $[-\overline{m}_E^{-}; \overline{m}_E^{+}]$	Caractéristiques
$g1$	$(\Delta\omega_{6\text{exp}}^{2+}, \Delta\omega_{6\text{exp}}^{2-})$	normale	$[-0, 032; 0, 034]$	$\mu = 0, 001 / \sigma = 0, 013$
$g2$	$(\Delta\omega_{4\text{exp}}^{2+}, \Delta\omega_{4\text{exp}}^{2-})$	normale	$[-0, 148; 0, 581]$	$\mu = 0, 217 / \sigma = 0, 142$
$g3$	$(\Delta\omega_{2\text{exp}}^{2+}, \Delta\omega_{2\text{exp}}^{2-})$	uniforme	$[-0, 165; 0, 075]$	$\mu = -0, 045 / \sigma = 0, 035$

TAB. 6.11 – Résultats de la réduction des méconnaissances de base du treillis à six barres avec forte méconnaissance, sans prise en compte de cette dernière

La méconnaissance de base sur le groupe $g2$ après réduction se révèle peu précise, comparée à la référence « expérimentale » donnée par le tableau 6.9 ; la réduction des méconnaissances sur les autres groupes n'est pas perturbée de manière sensible. **Quoique sous-estimée, la méconnaissance obtenue sur $g2$ est associée à une loi normale de grand écart type, ce qui suggère a posteriori d'utiliser l'approche adaptée au cas de la forte méconnaissance.**

On prend maintenant en compte la forte méconnaissance sur le groupe g_2 , en utilisant comme base de réduction les quatre premiers modes propres du modèle théorique ; les réductions successives sont mentionnées dans le tableau 6.12. Encore une fois, une seule itération sur tous les groupes a été réalisée. Les résultats sont cette fois-ci plus conformes à ce que présentent les structures « expérimentales » dans le tableau 6.9.

Gr.	Données expérimentales	Loi	Méconnaissances après réduction	
			$[-\overline{m}_E^-; \overline{m}_E^+]$	Caractéristiques
g_1	$(\Delta\omega_{6\text{exp}}^{2+}, \Delta\omega_{6\text{exp}}^{2-})$	normale	$[-0,032; 0,034]$	$\mu = 0,001 / \sigma = 0,013$
g_2	$(\Delta\omega_{4\text{exp}}^{2+}, \Delta\omega_{4\text{exp}}^{2-})$	normale	$[-0,195; 0,703]$	$\mu = 0,254 / \sigma = 0,174$
g_3	$(\Delta\omega_{2\text{exp}}^{2+}, \Delta\omega_{2\text{exp}}^{2-})$	uniforme	$[-0,155; 0,075]$	$\mu = -0,04 / \sigma = 0,033$

TAB. 6.12 – Résultats de la réduction des méconnaissances de base du treillis à six barres avec forte méconnaissance, avec prise en compte de cette dernière

6.4 Réduction des méconnaissances de base : seconde formulation

6.4.1 Notion de représentativité d'une mesure expérimentale

Dans tout ce qui précède, nous avons supposé que les données expérimentales utilisées dans le processus de réduction permettaient de donner une image fidèle des méconnaissances de base sur chaque sous-structure. En pratique, il faudrait pouvoir tenir compte du fait que **l'information apportée par l'essai expérimental est généralement partielle**.

Une solution envisageable est d'introduire lors de la réduction des méconnaissances un coefficient multiplicateur sur la contribution de E^* qui permet de quantifier à quel point l'information expérimentale sélectionnée est représentative du comportement de la structure : ce terme $\rho_{E^*} \in]0; 1]$, appelé *coefficient de représentativité*, est égal à un quand les données expérimentales traduisent parfaitement la mécanique globale de la structure.

Ce problème de représentativité est lié à ce qu'on pourrait appeler le caractère « isotrope » de la méconnaissance de base étudiée. Ainsi, on conçoit aisément que les données expérimentales issues d'un essai de traction sont très adaptées dans la réduction de la méconnaissance concernant un modèle de sous-structure isotrope : dans ce cas, il est légitime de considérer un coefficient ρ_{E^*} égal à un.

Inversement, pour un modèle de sous-structure orthotrope, cet essai directionnel ne peut donner d'informations que dans la direction de sollicitation. L'apport d'informations qui lui est associé est donc très partiel : en effet, l'essai ne permet de voir qu'une « partie » de la méconnaissance de base qui représente les incertitudes sur la sous-structure concernée, et on pondère alors l'essai avec un coefficient ρ_{E^*} strictement inférieur à un.

Au bilan, les quantités qui sont calculées sont les suivantes :

$$\Delta\alpha_{\text{mod}}^{+\text{pire}} = \rho_{E^*} \Delta\alpha_{E^*}^{+} + \sum_{E \neq E^*} \Delta\alpha_E^{+\text{pire}} \quad (6.13a)$$

$$\Delta\alpha_{\text{mod}}^{-\text{pire}} = \rho_{E^*} \Delta\alpha_{E^*}^{-} + \sum_{E \neq E^*} \Delta\alpha_E^{-\text{pire}} \quad (6.13b)$$

La question qui subsiste concerne la façon dont on pourrait déterminer la valeur de ce coefficient pour chaque sous-structure ; pour tenter de répondre à cette question, nous proposons l'illustration suivante.

6.4.2 Illustration

Le cas d'étude que nous considérons ici est le suivant : on analyse les modes propres de flexion d'une plaque carrée dont le matériau est orthotrope afin de déterminer la méconnaissance de base caractérisant les dispersions du matériau choisi ; il s'agit en l'occurrence du nid d'abeille d'aluminium, dont les coefficients, définis de façon classique, sont :

$$\begin{aligned} E_1 &= 0,1 \text{ GPa} & E_2 &= 0,07 \text{ GPa} & \nu_{12} &= 0,4 \\ G_{12} &= 0,044 \text{ GPa} & G_{13} &= 0,186 \text{ GPa} & G_{23} &= 0,09 \text{ GPa} \end{aligned}$$

Le modèle Éléments Finis construit à partir des coefficients précédents constitue notre modèle théorique déterministe ; la plaque est de forme carrée de côté de longueur $l = 1\text{m}$, et d'épaisseur $e = 10\text{cm}$.

Données expérimentales

Les structures « expérimentales » sont simulées à partir du modèle théorique déterministe en ajoutant des dispersions sur les différents coefficients du matériau orthotrope choisi ; ces dispersions sont obtenues à l'aide de tirages de Monte Carlo avec des lois de probabilité normales, dont les caractéristiques sont précisées dans le tableau 6.13.

Coefficients	Loi choisie	Domaine relatif	Moments statistiques relatifs (μ : moyenne / σ : écart type)
E_1	normale	$[-0,10; 0,10]$	$\mu = 0,00 / \sigma = 0,039$
E_2	normale	$[-0,01; 0,01]$	$\mu = 0,00 / \sigma = 0,004$
$\{G_{12}; G_{13}; G_{23}\}$	normale	$[-0,05; 0,05]$	$\mu = 0,00 / \sigma = 0,019$

NB : Les caractéristiques des lois sont données relativement au module d'Young du modèle déterministe.

TAB. 6.13 – Caractéristiques relatives des lois de probabilité permettant la simulation des valeurs « expérimentales » des coefficients matériaux de la plaque orthotrope

À partir de cette famille de structures simulées, on peut déduire la répartition des pulsations propres « expérimentales », et ainsi déterminer les valeurs « expérimentales » associées à des valeurs de probabilité données.

Modèle initial de méconnaissances

Le modèle initial de méconnaissances est précisé dans le tableau 6.14 ; on a décidé de considérer la structure dans son ensemble, vu qu'il n'y a pas de raison de privilégier un découpage particulier. Ceci nous permet de plus de nous affranchir de l'étude des cas défavorables au cours du processus de réduction.

Groupe	Loi choisie	Domaine $[-\overline{m}_E^{-0}; \overline{m}_E^{+0}]$	Moments statistiques relatifs (μ : moyenne / σ : écart type)
Ω	normale	$[-0,25; 0,25]$	$\mu = 0 / \sigma = 0,097$

TAB. 6.14 – Modèle initial ($\overline{m}_E^{-0}, \overline{m}_E^{+0}$) de méconnaissances de base pour la plaque orthotrope

Mise en œuvre du processus de réduction

Étant donné que l'on considère directement la méconnaissance sur l'ensemble de la structure, tous les modes sont susceptibles de constituer une information pertinente afin de réduire le niveau de méconnaissance initial. On réalise donc des réductions indépendantes à partir de chaque couple de valeurs « expérimentales » à 99 % concernant les pulsations propres. On récapitule dans le tableau 6.15 les résultats de la réduction de la méconnaissance sur la plaque orthotrope à l'aide de chacun des neuf premiers modes propres de la structure.

Données expérimentales	Loi	Méconnaissances après réduction $[-\overline{m}_E^{-0}; \overline{m}_E^{+0}]$	Caractéristiques
$(\Delta\omega_{1\text{exp}}^{2+}, \Delta\omega_{1\text{exp}}^{2-})$	normale	$[-0,032; 0,036]$	$\mu = 0,002 / \sigma = 0,013$
$(\Delta\omega_{2\text{exp}}^{2+}, \Delta\omega_{2\text{exp}}^{2-})$	normale	$[-0,037; 0,029]$	$\mu = -0,004 / \sigma = 0,013$
$(\Delta\omega_{3\text{exp}}^{2+}, \Delta\omega_{3\text{exp}}^{2-})$	normale	$[-0,064; 0,070]$	$\mu = -0,003 / \sigma = 0,028$
$(\Delta\omega_{4\text{exp}}^{2+}, \Delta\omega_{4\text{exp}}^{2-})$	normale	$[-0,037; 0,037]$	$\mu = 0,000 / \sigma = 0,014$
$(\Delta\omega_{5\text{exp}}^{2+}, \Delta\omega_{5\text{exp}}^{2-})$	normale	$[-0,064; 0,064]$	$\mu = 0,000 / \sigma = 0,025$
$(\Delta\omega_{6\text{exp}}^{2+}, \Delta\omega_{6\text{exp}}^{2-})$	normale	$[-0,002; 0,003]$	$\mu = 0,001 / \sigma = 0,001$
$(\Delta\omega_{7\text{exp}}^{2+}, \Delta\omega_{7\text{exp}}^{2-})$	normale	$[-0,090; 0,091]$	$\mu = 0,001 / \sigma = 0,035$
$(\Delta\omega_{8\text{exp}}^{2+}, \Delta\omega_{8\text{exp}}^{2-})$	normale	$[-0,048; 0,044]$	$\mu = -0,002 / \sigma = 0,018$
$(\Delta\omega_{9\text{exp}}^{2+}, \Delta\omega_{9\text{exp}}^{2-})$	normale	$[-0,035; 0,036]$	$\mu = 0,001 / \sigma = 0,014$

TAB. 6.15 – Résultats de la réduction de la méconnaissance de base sur la plaque orthotrope suivant l'information expérimentale utilisée

On se rend alors compte de la forte disparité des résultats obtenus : ceci est dû au caractère plus ou moins « orienté » de certains modes, avec pour conséquence que l'un des paramètres du matériau a une plus forte influence que les autres dans la dispersion de la pulsation propre considérée ; ainsi, si l'on observe plus particulièrement les modes 6 et 7 représentés sur la figure 6.8, on se rend compte que ces derniers correspondent à des ondulations quasiment planes, chacune selon l'une des deux directions principales de la plaque ; ils font donc essentiellement intervenir les modules E_2 et E_1 respectivement,

et on retrouve d'ailleurs directement dans le niveau de méconnaissance de base après réduction une estimation relativement bonne de la dispersion introduite pour la simulation stochastique de ces deux coefficients, à savoir 1 % et 10 % respectivement.

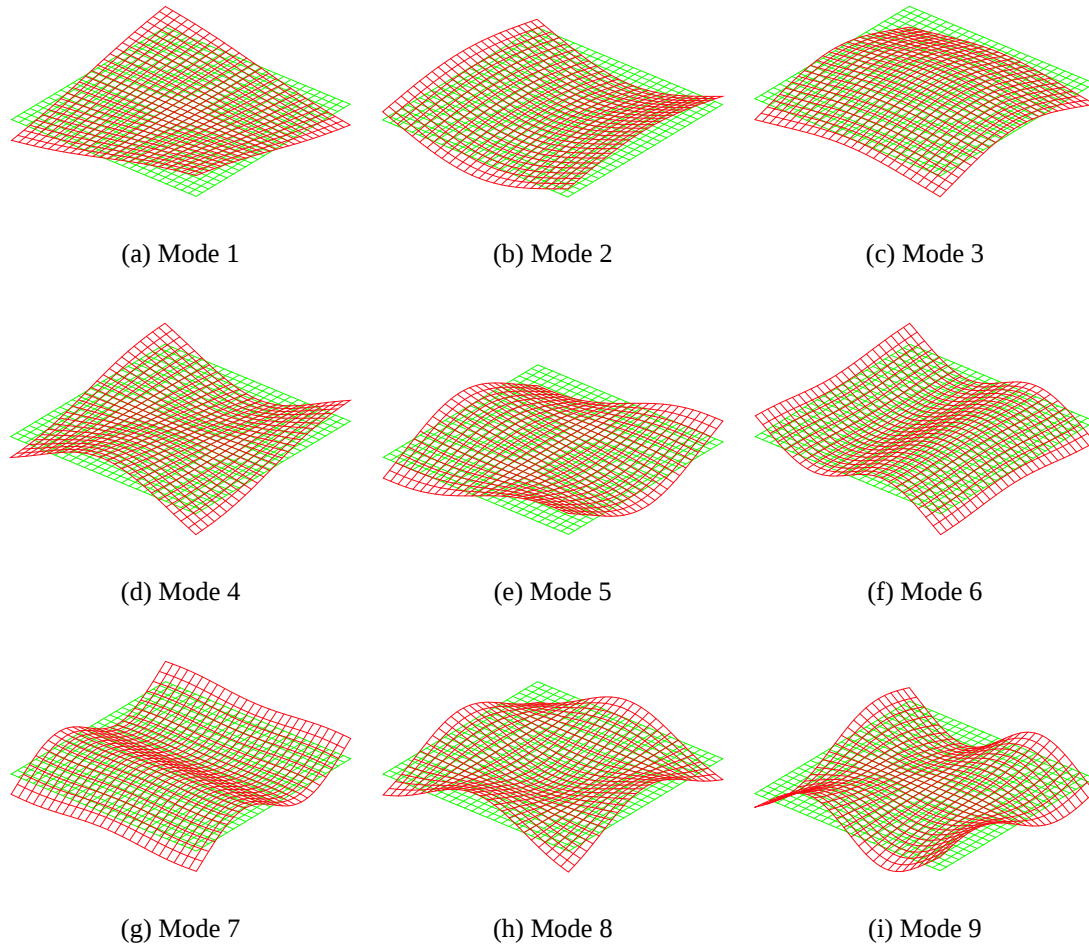


FIG. 6.8 – Modes propres de flexion de la plaque orthotrope

Inversement, pour des modes plus « symétriques », comme les modes 1, 8 et 9 représentés sur la figure 6.8, on obtient à peu près la même description de méconnaissance après réduction, car ils font intervenir les différentes dispersions avec des poids similaires.

Il est donc indispensable de pouvoir quantifier le « degré de représentativité » de l'information expérimentale utilisée dans le processus de réduction de la sous-structure considérée ; c'est au travers de la définition du coefficient de représentativité ρ_E que l'on entend répondre à ce besoin, en caractérisant à quel point le test expérimental utilisé permet une vision juste des dispersions représentées par la méconnaissance de base sur la sous-structure. Encore faut-il s'accorder sur le niveau que l'on veut prendre comme référence parmi les différentes méconnaissances de base obtenues après réduction à l'aide des différents modes expérimentaux utilisés.

Comme on considère que la détermination des méconnaissances les plus représentatives de la réalité étudiée doit toujours être menée dans le sens de la sécurité, c'est le niveau de méconnaissance de base obtenu avec le mode 7 que nous décidons de considérer comme l'indicateur des incertitudes présentes dans la structure étudiée. Nous associons donc à ce mode un coefficient de représentativité égal à un.

Le mode 6, au contraire, ne donne qu'une représentation partielle de cette méconnaissance de base, qui se traduit par une sous-estimation du niveau obtenu après réduction. C'est pourquoi on associe à ce mode un coefficient de représentativité différent de un. Dans notre cas, ce coefficient est égal au rapport entre la méconnaissance obtenue après réduction avec le mode 6, et la méconnaissance prise comme référence, c'est-à-dire une valeur très proche de zéro qui confirme bien le fait que le mode en question traduit un phénomène lié à une direction de sollicitation presque orthogonale à celle du mode 7 qui a conduit à l'estimation de la méconnaissance de base caractéristique de la structure étudiée.

Le même raisonnement est valable pour les modes que l'on n'a pas encore analysés ; aux modes « symétriques » 1, 8 et 9 sont associés des coefficients de représentativité de valeurs sensiblement égales, traduisant bien le fait que ces trois modes ont la même « vision » de l'état de méconnaissance de la structure.

Ainsi, **la notion de coefficient de représentativité permet théoriquement de quantifier à quel point l'essai « voit » la méconnaissance associée à la sous-structure dont on est en train de réduire la méconnaissance**, et cette notion est indispensable lorsque l'on étudie des structures au comportement très fortement non isotrope, comme dans l'exemple que l'on vient de détailler ; cependant, **en pratique, il est pour l'instant difficile de déterminer *a priori* les valeurs de ce coefficient en fonction du type de mode considéré**. Dans l'exemple que nous venons de présenter, il semble même que l'utilisation de ce seul coefficient soit insuffisante, et qu'il faille introduire une description plus riche avec tout un jeu de coefficients qui dépendraient de la forme du mode considéré en relation avec le type de modèle de sous-structure utilisé.

6.5 Vers une stratégie générale de réduction des méconnaissances

Il est maintenant possible de poser les bases d'une stratégie générale de réduction des méconnaissances de base par l'apport d'informations expérimentales pertinentes.

On rappelle que ces dernières sont constituées, dans notre cadre d'étude, par des pulsations propres et des formes propres de vibration, et que l'on peut déterminer des valeurs qui encadrent une certaine proportion P de ces quantités d'intérêt dérivées de la réalité mesurée. Ainsi, **une étape préalable à l'application de notre stratégie de réduction consiste à juger des valeurs de P les plus pertinentes pour la description des données expérimentales**.

Comme on l'avait déjà souligné auparavant, il est souvent difficile d'avoir des mesures effectuées sur un « grand » nombre de structures réelles semblables, compte tenu du coût généralement élevé des campagnes d'essais ; généralement, on se contentera donc de valeurs à 99 %, qui seront souvent en fait des valeurs à 100 % (on se réserve juste la possibilité d'éliminer une, voire deux structures dont les résultats semblent manifestement aberrants par rapport à la dispersion de la majorité des autres), et dans certains cas, on considérera aussi les valeurs à 70 %, qui permettent d'obtenir une estimation (grossière) de l'écart type des mesures.

6.5.1 Définition du modèle initial de méconnaissances de base

Le point de départ de la stratégie est l'élaboration d'un premier modèle de méconnaissances de base pour chaque sous-structure, sachant que la description de ces méconnaissances doit nécessairement être majorante. Par cela, on veut dire que n'importe quel couple de valeurs à P % dérivé des quantités d'intérêt expérimentales doit être inclus dans l'intervalle de méconnaissance effective associé au modèle initial :

$$\Delta\alpha_{\text{exp}}^-(P) \leq \Delta\alpha_{\text{mod}}^{-0}(P) \quad \text{et} \quad \Delta\alpha_{\text{exp}}^+(P) \leq \Delta\alpha_{\text{mod}}^{+0}(P) \quad (6.14)$$

Ce niveau initial de méconnaissance peut bien sûr être déterminé à partir d'une connaissance préalable que l'on détient : ainsi, par exemple, la connaissance de la dispersion des matériaux de la structure considérée permet d'avoir une idée assez juste de la dispersion résultante en termes de méconnaissances sur cette sous-structure. Par contre, le modèle de méconnaissance initial devra légèrement surévaluer cette dispersion estimée, car on a impérativement besoin d'une description majorante, et il est très probable que certaines sources d'incertitude aient été oubliées lors de cette estimation.

Il est important de signaler que, sauf cas exceptionnel, **il n'y a pas de raison d'avoir une forte disparité des niveaux de méconnaissance sur les différentes sous-structures :** c'est seulement si un modèle est issu de simplifications importantes du comportement de la sous-structure associée que l'on peut avoir une valeur de méconnaissance beaucoup plus importante sur cette sous-structure que sur les autres.

Enfin, il faut évoquer à nouveau la question du choix des lois de probabilité des méconnaissances de base : encore une fois, **vu le nombre souvent limité des structures réelles expérimentées, il est difficile de vérifier *a posteriori* le bon choix des lois de probabilité.** Une conclusion raisonnable consiste alors à faire preuve de bon sens, et à choisir une loi normale quand il n'est question *a priori* que de dispersions matériaux dans la sous-structure, et une loi uniforme sinon, voire pas de loi de probabilité du tout, si l'on estime déjà que le modèle de la sous-structure est extrêmement simplifié.

6.5.2 Détermination de l'ordre des réductions successives

Il s'agit maintenant de préciser dans quel ordre on doit mener les réductions, et surtout à partir de quelle information expérimentale.

On a vu précédemment que la sensibilité des méconnaissances effectives vis-à-vis des méconnaissances de base est directement liée aux énergies de déformation modales du modèle théorique déterministe. En effet, les expressions des bornes $\Delta\alpha_{\text{mod}}^+$ et $\Delta\alpha_{\text{mod}}^-$ encadrant la quantité $\Delta\alpha_{\text{mod}}$ relative au modèle avec méconnaissances sont des combinaisons linéaires des méconnaissances de base, dont les coefficients de linéarité sont des énergies associées au modèle théorique $\bar{e}_E$.

Pour cela, on dresse le tableau des énergies de déformation par groupe et par mode. **On considère qu'un mode est une information pertinente pour la réduction de la méconnaissance sur une sous-structure donnée si la part de cette dernière dans l'énergie de déformation totale associée à ce mode est supérieure à une certaine proportion $k\%$ de l'énergie totale** ; typiquement, on choisit $k\% \approx \frac{100}{N}\%$, où N est le nombre de sous-structures pour lesquelles on cherche à déterminer la méconnaissance de base.

À partir de là, la démarche de sélection est la suivante :

- on classe pour chaque sous-structure les modes pertinents par ordre décroissant de leurs contributions relatives ;
- on commence par réduire la méconnaissance sur la sous-structure qui présente la plus forte contribution relative, en utilisant les valeurs expérimentales pertinentes associées ; on considère alors que ce mode ne servira pas pour la réduction de la méconnaissance sur une autre sous-structure ;
- la réduction suivante concerne la sous-structure présentant la deuxième valeur la plus forte de contribution relative ; ceci permet de limiter au maximum lors de chaque réduction le poids de la prise en compte des cas défavorables sur les autres sous-structures ;
- on continue de la même façon avec les autres sous-structures.

Il est possible que certaines sous-structures ne présentent pas de contribution dominante dans l'une des énergies de déformation modales ; dans ce cas, on cherche le mode pour lequel la contribution proposée par la sous-structure est la plus grande. La réduction à l'aide de ce mode devra donc être réalisée impérativement après la réduction de toutes les méconnaissances de sous-structures dont la contribution dans l'énergie associée au mode retenu est plus grande que $k\%$.

6.5.3 Processus de réduction des méconnaissances de base

On suit l'ordre défini dans le paragraphe précédent : dès que l'on a réduit la méconnaissance de base sur une sous-structure, le résultat obtenu est utilisé pour la réduction suivante sur une autre sous-structure, dans la définition des cas défavorables.

Une fois que l'on a réduit la méconnaissance de base sur l'ensemble des sous-structures, on peut recommencer encore une fois l'ensemble des réductions, vu que, par rapport à la première tentative de réduction sur la première sous-structure, le niveau des méconnaissances de base sur les autres sous-structures aura diminué. **On peut ainsi être amené à réaliser plusieurs fois la totalité des réductions, tant qu'il est possible de les réaliser.**

Dès qu'une nouvelle itération sur l'ensemble des sous-structures n'apporte plus de modification dans le niveau obtenu pour les méconnaissances, on peut considérer qu'aucune amélioration n'est possible, en tout cas en gardant l'idée de considérer les cas défavorables à chaque fois. **Si l'on estime que l'ensemble des données expérimentales considérées couvre l'ensemble des situations envisageables pour la famille de structures étudiées, on peut refaire une itération sur toutes les sous-structures, mais cette fois-ci en ne considérant pas les contraintes associées à la prise en compte des cas défavorables.**

Enfin, dans le cadre de ce travail, on choisit pour chaque réduction un coefficient de représentativité égal à un.

6.5.4 Interprétation des résultats

Une fois les résultats obtenus, il faut encore observer les étendues des lois de probabilité des méconnaissances déterminées par le processus de réduction ; en effet, si, sur une sous-structure, on est susceptible d'obtenir des réalisations de la méconnaissance de base qui excèdent des valeurs de l'ordre de 25 % ou 30 %, on peut légitimement s'interroger sur la validité des approximations au premier ordre permettant de propager cette méconnaissance.

Dans ce cas, on doit recommencer tout le processus de réduction, en utilisant la version « forte méconnaissance » de calcul des méconnaissances effectives ; il est important de tout reprendre à zéro, car les fortes valeurs de méconnaissance peuvent aussi modifier les résultats de réduction sur les autres sous-structures, vu que les énergies de déformation associées peuvent dépendre des réalisations de la forte méconnaissance.

On compare ensuite les valeurs de méconnaissances obtenues avec cette prise en compte à celles déterminées initialement, pour constater si l'hypothèse de forte méconnaissance est nécessaire ou non.

À partir de là, **on peut réaliser des prédictions sur des quantités d'intérêt non utilisées dans le processus de réduction ;** on a pu juger sur les quelques exemples précédents de la qualité de ces prédictions, qui nous permettent d'obtenir, à partir d'un modèle pourtant imparfait par rapport à la réalité étudiée, des résultats quantifiés sur diverses quantités d'intérêt.

Les différents niveaux de méconnaissances de base déterminés sont également des indicateurs de la qualité des différents modèles de sous-structures utilisés, comme on le précise d'ailleurs dans le chapitre suivant.

Enfin, **la théorie des méconnaissances permet d'envisager l'élaboration de nouveaux essais expérimentaux permettant de mieux connaître certaines sous-structures :** il suffit de trouver une configuration de test dans laquelle la sous-structure concernée présente une contribution importante dans l'énergie de déformation totale. Comme le niveau de méconnaissance de base est une propriété intrinsèque de la sous-structure, il sera possible d'utiliser les informations obtenues dans d'autres configurations.

La théorie des méconnaissances au sein de la démarche de validation de modèles

Tous les modèles sont faux ; quelques-uns sont utiles.

George Box

Sommaire

7.1 Recalage de modèles dynamiques	122
7.1.1 Concept d'erreur en relation de comportement	122
7.1.2 Prise en compte des données expérimentales	123
7.1.3 Implémentation de la méthode	125
7.2 Validation d'un modèle simplifié de liaison	127
7.2.1 Présentation de la structure d'étude	127
7.2.2 Recalage du modèle théorique à l'aide des données expérimentales	128
7.2.3 Réduction des méconnaissances de base	129

Ce septième chapitre présente l'analyse des résultats de la détermination des méconnaissances de base dans le cas d'un modèle de sous-structure très simplifié par rapport à la référence expérimentale ; la théorie des méconnaissances est placée dans le cadre général de la démarche de validation d'un modèle, et les possibilités associées sont également confrontées aux indications données par une méthode classique de validation de modèles, à savoir le recalage à l'aide du concept de l'erreur en relation de comportement modifiée.

7.1 Recalage de modèles dynamiques avec le concept de l'erreur en relation de comportement modifiée

Dans ce paragraphe, nous présentons rapidement les idées forces du recalage à l'aide du concept de l'erreur en relation de comportement modifiée, initié dès [Ladevèze 1983] et proposé dans sa forme actuelle dans [Ladevèze 1993]. On trouvera également plus de précisions dans [Ladevèze *et al.* 1994a - b, Chouaki 1997, Ladevèze et Chouaki 1999, Deraemaeker 2001].

7.1.1 Concept d'erreur en relation de comportement

On considère une structure Ω pendant un certain intervalle de temps $[0; T]$. Sur le contour $\partial\Omega$ sont imposés des déplacements $\underline{U}_d$ et des efforts $\underline{F}_d$, respectivement sur $\partial_1\Omega$ et $\partial_2\Omega$ tels que $\partial_1\Omega \cup \partial_2\Omega = \partial\Omega$; en tout point de la structure Ω sont définies des forces volumiques $\underline{f}_v$.

Le problème de référence consiste à déterminer, en tout point $\underline{M} \in \Omega$ et à chaque instant $t \in [0; T]$, le champ de déplacement $\underline{U}(\underline{M}, t)$, le champ de contrainte $\underline{\sigma}(\underline{M}, t)$ et la quantité d'accélération $\underline{\Gamma}(\underline{M}, t)$ qui vérifient :

- les conditions initiales à $t = 0$;
- les conditions limites en déplacement ;
- les équations d'équilibre dynamique ;
- les relations de comportement.

La construction d'une erreur en relation de comportement est alors basée sur l'idée de séparer les équations précédentes en une partie fiable et une partie moins fiable :

<u>Équations fiables</u>	<u>Équations moins fiables</u>
conditions initiales	
conditions limites en déplacement	relations de comportement
équations d'équilibre dynamique	

Cette séparation permet de définir deux notions complémentaires :

- **l'admissibilité** : les inconnues du problème sont dites admissibles si elles vérifient les équations fiables, ce que l'on écrit formellement comme :

$$s = (\underline{U}, \underline{\sigma}, \underline{\Gamma}) \in \mathcal{S}_{ad}$$

- **l'erreur en relation de comportement** : cette erreur représente un résidu sur les équations moins fiables du problème, qui est nul si celles-ci sont vérifiées, et positif dans le cas contraire :

$$\begin{aligned}\xi^2(s) &\geq 0 \\ \xi^2(s) &= 0 \iff \text{les relations de comportement sont vérifiées}\end{aligned}$$

où $\xi^2(s)$ représente l'erreur en relation de comportement.

Ces définitions permettent de réécrire le problème de référence sous la forme :

$$\left| \begin{array}{l} \text{Trouver } s = (\underline{U}, \underline{\sigma}, \underline{\Gamma}) \in \mathcal{S}_{\text{ad}} \\ \text{qui minimise } \xi^2(s') \text{ avec } s' \in \mathcal{S}_{\text{ad}} \end{array} \right. \quad (7.1)$$

Pour fixer les idées, si l'on étudie les vibrations forcées, à une pulsation ω donnée, d'une structure dont la relation de comportement est purement élastique, les relations de comportement que l'on considère sont :

$$\underline{\sigma} = \mathbf{K}\epsilon(\underline{U}) \quad (7.2a)$$

$$\underline{\Gamma} = -\rho\omega^2\underline{U} \quad (7.2b)$$

où $\mathbf{K}$ et ρ désignent respectivement le tenseur de Hooke et la masse volumique du matériau constitutif de la structure Ω .

Dans ce cas, l'erreur en relation de comportement associée à ces équations réputées moins fiables s'écrit :

$$\begin{aligned}\xi^2(s') &= \gamma \int_{\Omega} (\underline{\sigma} - \mathbf{K}\epsilon(\underline{U})) \mathbf{K}^{-1} (\underline{\sigma} - \mathbf{K}\epsilon(\underline{U})) \, d\Omega \\ &\quad + (1 - \gamma) \int_{\Omega} (\underline{\Gamma} + \rho\omega^2\underline{U}) \frac{1}{\rho\omega^2} (\underline{\Gamma} + \rho\omega^2\underline{U}) \, d\Omega\end{aligned} \quad (7.3)$$

où γ est un coefficient qui permet de pondérer le poids de chaque résidu associé à une relation de comportement donnée. L'expression précédente peut ensuite s'écrire sous forme discrétisée grâce à la méthode des Éléments Finis.

7.1.2 Prise en compte des données expérimentales

Lorsque l'on est en présence de résultats d'essais, les données expérimentales ne peuvent pas toutes être considérées comme fiables. Il convient donc de redéfinir l'admissibilité en tenant compte de ces nouvelles quantités. À titre d'exemple, on considère une structure, soumise à un effort dont on mesure l'amplitude, et instrumentée de capteurs qui enregistrent des données en sortie (déplacements, accélérations, déformations, ...).

On définit alors comme quantités fiables :

- la pulsation ω ;
- la position et la direction des actionneurs et capteurs.

Les quantités moins fiables sont données par :

- les amplitudes des forces mesurées aux actionneurs ;
- les grandeurs de sortie aux capteurs, constituant un vecteur de dimension finie contenant les valeurs discrètes des mesures.

Compte tenu de l'introduction de ces grandeurs expérimentales moins fiables, le résidu $\xi^2(s)$ doit être augmenté d'une erreur sur les amplitudes mesurées. L'erreur ainsi construite est appelée *erreur en relation de comportement modifiée*.

Comme le traitement d'un signal temporel par *FFT* et l'acquisition expérimentale par balayage sinus procurent des données de type *FRF* sur une bande fréquentielle $[\omega_{\min}, \omega_{\max}]$, nous considérons que les données mesurées sont les déplacements à la pulsation ω , qui sont notés $\underline{U}_\omega$. Le vecteur $\{f\}$ représente l'amplitude de la force mesurée et $\{u\}$ le déplacement au droit des actionneurs. L'erreur modifiée fait intervenir toutes ces quantités moins fiables :

$$e_\omega^2 = \xi_\omega^2 + \frac{r}{1-r} \left\{ \|\Pi \underline{U} - \underline{U}_\omega\|^2 + |\{f\} - \{\tilde{f}\}|^2 + |\{u\} - \{\tilde{u}\}|^2 \right\} \quad (7.4)$$

où r est un coefficient qui permet de quantifier la confiance que l'on place dans les mesures expérimentales, et Π un opérateur de projection sur l'espace des degrés de liberté mesurés.

La qualité d'un modèle par rapport à un ensemble de résultats d'essais est alors caractérisée par le résidu en erreur en relation de comportement modifiée introduit ci-dessus. La construction de celui-ci demande de résoudre le problème suivant :

$$\left| \begin{array}{l} \text{Trouver } s = (\underline{U}, \sigma, \underline{\Gamma}, \{u\}, \{f\}) \in \mathcal{S}_{\text{ad}}^\omega \\ \text{qui minimise } e_\omega^2(s') \text{ avec } s' \in \mathcal{S}_{\text{ad}}^\omega \end{array} \right. \quad (7.5)$$

Notons l'introduction des variables $\{u\}$ et $\{f\}$ représentant le déplacement et la force en chaque point d'excitation du problème considéré.

La solution s du problème de minimisation est ensuite introduite dans l'expression de l'erreur en relation de comportement : la valeur de $\xi_\omega^2(s)$ représente la qualité de vérification des relations de comportement de la solution s . Elle est donc représentative de la qualité du modèle supposé par rapport aux résultats d'essais.

La structure étant subdivisée en plusieurs sous-structures $E \in \Omega$, on peut écrire :

$$\xi_\omega^2(s) = \sum_{E \in \Omega} \xi_{E,\omega}^2(s) \quad (7.6)$$

Le terme de droite représente la contribution de chaque sous-structure à l'erreur.

Afin de pouvoir utiliser simultanément plusieurs mesures obtenues pour différentes valeurs de la pulsation d'excitation dans le cas de vibrations forcées, ou plusieurs modes propres de vibrations dans le cas de vibrations libres, on introduit un facteur de pondération $p(\omega) \geq 0$ tel que :

$$\int_{\omega_{\min}}^{\omega_{\max}} p(\omega) d\omega = 1 \quad (7.7)$$

En pratique, $p(\omega) = 0$ presque partout, sauf aux valeurs de pulsations que l'on souhaite étudier. On est alors en présence des contributions locales suivantes :

$$\xi_{ET}^2(s) = \int_{\omega_{\min}}^{\omega_{\max}} p(\omega) \xi_{E,\omega}^2(s) d\omega \quad (7.8)$$

L'erreur en relation de comportement modifiée donne ainsi une mesure de la qualité d'un modèle par rapport à un ensemble de résultats d'essais au niveau global mais aussi au niveau local. Il s'agit maintenant de préciser l'exploitation de ces deux grandeurs utilisées en vue d'améliorer la qualité d'un modèle.

7.1.3 Implémentation de la méthode

Le problème de référence dépend d'un certain nombre de paramètres sujets à caution, comme les épaisseurs, modules d'Young ou coefficients d'amortissement par exemple ; ces paramètres structuraux sont rangés dans un vecteur $\underline{k}$ et l'espace correspondant est noté $\mathbf{k}$. Le problème consiste à trouver la valeur optimale de $\underline{k}$ dans $\mathbf{k}$, c'est-à-dire celle qui donne la meilleure corrélation du modèle avec les essais. **La démarche proposée s'appuie sur une approche itérative dans laquelle chaque itération est décomposée en deux étapes :**

- une étape de localisation des zones les plus erronées ;
- une étape de correction de ces zones les plus erronées.

Étape de localisation

Cette étape consiste à calculer l'indicateur local ξ_{ET}^2 défini précédemment. La sélection des sous-structures les plus erronées se fait en retenant celles qui vérifient :

$$\xi_{ET}^2 \geq \delta \max_{E \in \Omega} \xi_{ET}^2 \quad (7.9)$$

avec $\delta = 0,8$ par exemple. Soit $\mathbf{Z}$ l'ensemble des sous-structures vérifiant (7.9) et $\mathbf{k}_Z$ l'ensemble des paramètres des sous-structures de $\mathbf{Z}$.

Étape de correction

Le problème s'écrit :

$$\left| \begin{array}{l} \text{Trouver } \underline{k} \in \mathbf{k}_Z \text{ minimisant :} \\ \underline{k} \rightarrow J(\underline{k}) \\ \mathbf{k}_Z \rightarrow \mathbb{R} \end{array} \right. \quad (7.10)$$

où la fonctionnelle $J(\underline{k})$ est donnée par :

$$J(\underline{k}) = \int_{\omega_{\min}}^{\omega_{\max}} p(\omega) e_{\omega}^2 d\omega$$

L'indicateur e_{ω}^2 représente l'erreur en relation de comportement modifiée définie à la pulsation ω à partir d'un modèle défini par les paramètres $\underline{k}$.

L'ensemble k_z représente le sous-espace des paramètres structuraux les plus erronés définis dans l'étape de localisation. Le nombre de ces paramètres peut encore être réduit de deux manières différentes :

- soit par une connaissance de la structure qui permet de mettre en doute des éléments structuraux difficiles à modéliser tels que des éléments de liaison ou des amortissements mal connus *a priori* ;
- soit par une analyse de sensibilité classique dans laquelle des paramètres ayant peu d'influence sur la réponse sont éliminés.

À l'issue de chaque itération, l'erreur est réévaluée et une nouvelle itération est effectuée si le seuil d'erreur requis n'est pas atteint. C'est cette procédure qui assure la régularisation du problème, car aucune régularisation n'est introduite comme il est habituel de le faire dans les problèmes inverses.

L'étape de localisation des zones les plus erronées est cruciale dans le cas de structures complexes dans lesquelles le nombre de paramètres est tellement important qu'une étude de sensibilité sur tous les paramètres du modèle s'avère très coûteuse. La localisation préalable des zones les plus erronées permet de concentrer les efforts de correction sur certaines zones et de s'approcher de manière itérative du modèle offrant la meilleure corrélation avec les mesures, en résolvant à chaque itération un problème de minimisation dans un sous-espace de très faible taille.

On résume sur la figure 7.1 la démarche adoptée pour le recalage d'un modèle dans le domaine fréquentiel à partir de *FRF* ou de données modales. Dans un premier temps, la qualité du modèle ξ_{Tot}^2 est évaluée sur la bande de fréquences d'intérêt. Si l'erreur est jugée trop élevée, on effectue une itération comprenant une étape de localisation basée sur le calcul de ξ_{ET}^2 et une étape de correction qui consiste à minimiser e_T^2 en modifiant les paramètres des zones sélectionnées. Les itérations continuent jusqu'à atteindre le niveau d'erreur requis ξ_0^2 .

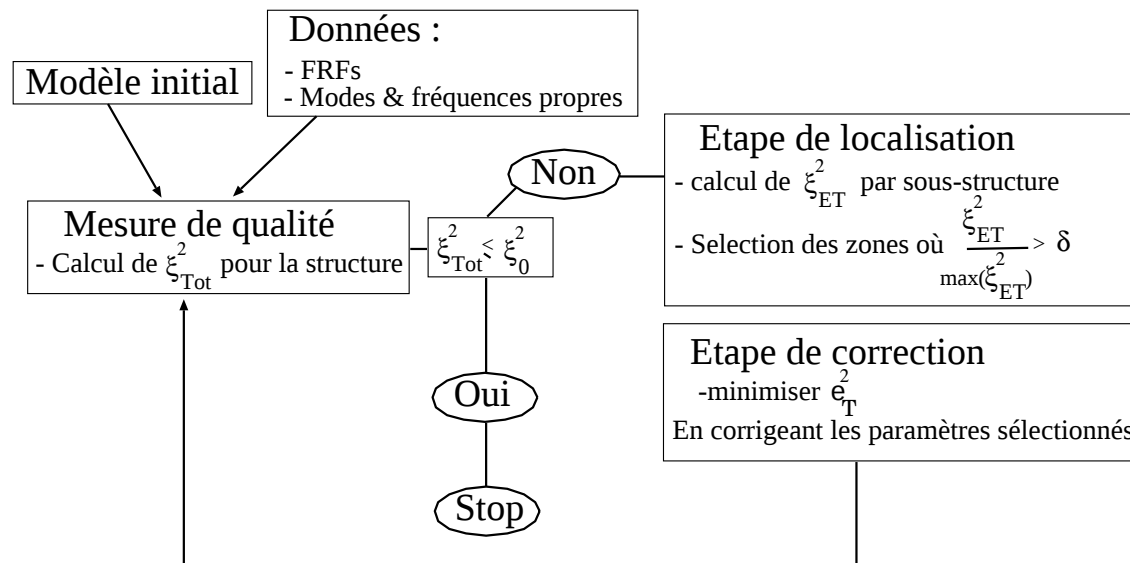


FIG. 7.1 – Procédure de recalage

7.2 Validation d'un modèle simplifié de liaison

7.2.1 Présentation de la structure d'étude

Données expérimentales

Le cas d'étude est constitué de deux poutres droites de même longueur reliées l'une à l'autre à une extrémité par une liaison, et respectivement encastree et libre à l'autre extrémité. La structure « expérimentale » est simulée à l'aide d'un modèle Éléments Finis détaillé dans le tableau 7.1. La liaison présente au milieu de cette structure est modélisée par deux ressorts K et k , représentant respectivement une raideur en translation et une raideur en rotation entre les deux nœuds des extrémités correspondantes des deux poutres ; la structure ainsi simulée est schématisée sur la figure 7.2.

Poutre	Section	Longueur	Matériau	Module d'Young	Masse vol.	Modèle EF
1	$1 \times 1\text{cm}^2$	1m	aluminium	$\bar{E} = 72\text{GPa}$	2700kg/m^3	50 éléments
2	$1 \times 1\text{cm}^2$	1m	aluminium	$\bar{E} = 72\text{GPa}$	2700kg/m^3	50 éléments

TAB. 7.1 – Caractéristiques du problème de poutres reliées par une liaison

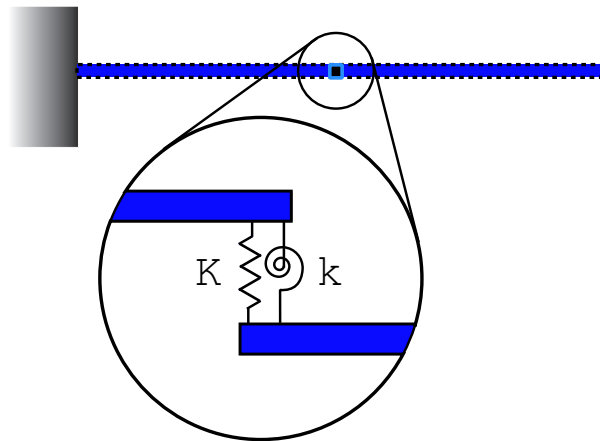


FIG. 7.2 – Problème de poutres reliées par une liaison

Les valeurs de raideurs au niveau de la liaison sont les suivantes :

$$K = 7 \cdot 10^7 \text{N/m} \quad \text{et} \quad k = 1000 \text{N.m/rad}$$

Cette structure simulée nous permet de constituer la représentation d'une réalité donnée, pour laquelle on peut déterminer les pulsations propres et modes propres « expérimentaux » associés.

Modèle théorique

Le modèle théorique que l'on souhaite valider est lui aussi un modèle Éléments Finis qui présente les mêmes caractéristiques matériaux que la structure « expérimentale » simulée, à la différence près que la liaison n'est pas considérée : la simplification effectuée

dans ce modèle consiste donc à supposer que les deux poutres sont reliées par une liaison encastrement parfaite. Néanmoins on garde à l'esprit l'existence de cette liaison au milieu de la structure « réelle ».

7.2.2 Recalage du modèle théorique à l'aide des données expérimentales

Les données expérimentales sont constituées par les huit premiers modes propres et pulsations propres de la structure étudiée. Un premier calcul de l'erreur globale en relation de comportement modifiée nous permet de déterminer une première estimation de la qualité du modèle théorique à l'aide du terme faisant intervenir les résidus sur les relations de comportement : on obtient alors une valeur :

$$\xi_{\text{Tot}}^2 = 4,35 \%$$

Une première étape de localisation est menée afin de déterminer les zones les plus erronées du modèle théorique ; la figure 7.3 permet d'observer une forte contribution à l'erreur des deux éléments de poutre encadrant l'emplacement de la liaison, ce qui est conforme à nos attentes.

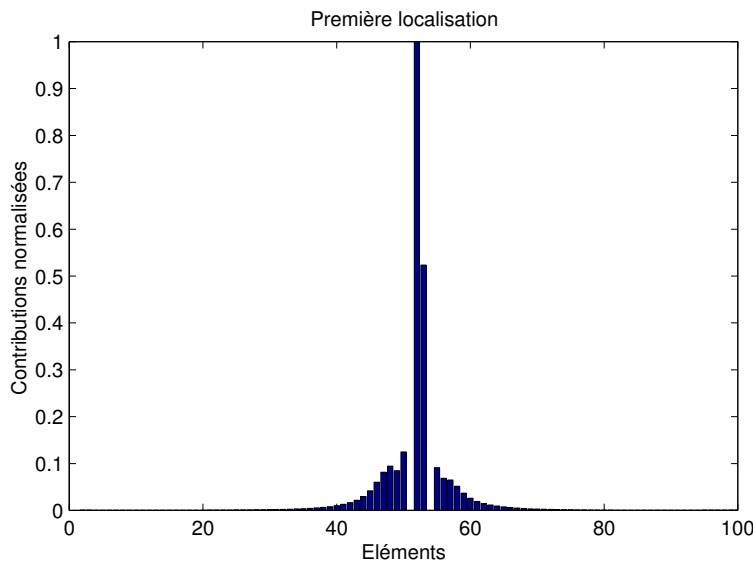


FIG. 7.3 – Première localisation dans le processus de recalage

Une phase de correction est menée sur les paramètres de raideur des deux éléments que l'on vient juste de sélectionner, et la minimisation de l'erreur modifiée entraîne une correction relative :

$$\Delta K_E = -59 \%$$

Une nouvelle estimation de l'erreur globale en relation de comportement modifiée nous conduit alors à :

$$\xi_{\text{Tot}}^2 = 1,26 \%$$

L'étape de localisation associée, représentée sur la figure 7.4, montre nettement les contributions encore importantes des deux éléments encadrant l'emplacement de la liaison.

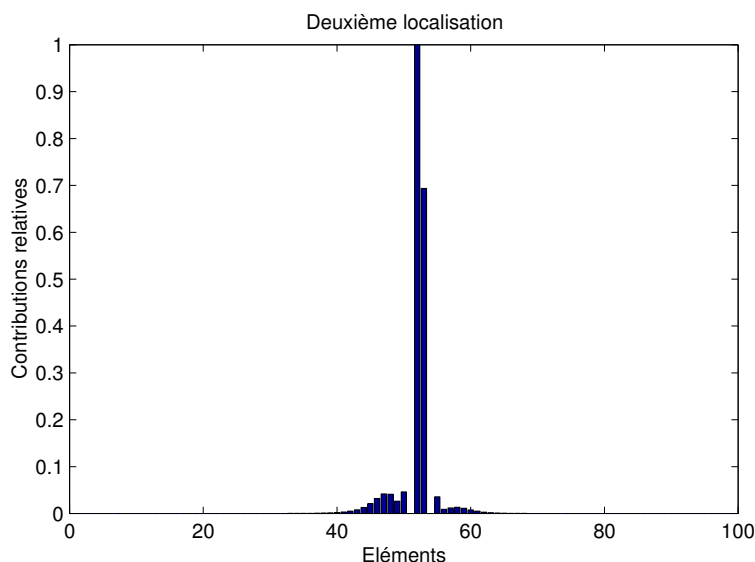


FIG. 7.4 – Seconde localisation dans le processus de recalage

Une nouvelle étape de correction sur les paramètres de ces deux éléments ne nous apportant aucune amélioration, on peut conclure que cette stagnation de l'erreur est significative de l'insuffisance du modèle théorique utilisé à représenter le comportement de la structure « réelle », et les fortes contributions locales au niveau de la liaison nous indiquent naturellement la zone « incorrecte » du modèle.

7.2.3 Réduction des méconnaissances de base

Données expérimentales

Les données « expérimentales » sont constituées par les pulsations propres et modes propres obtenus à partir de la simulation décrite dans le tableau 7.1 ; vu que l'on n'est en présence que d'une seule structure « expérimentale » simulée, on associe à chaque quantité d'intérêt un intervalle de variation, de la forme $[0; \Delta\alpha_{\text{exp}}^+]$ ou $[-\Delta\alpha_{\text{exp}}^-; 0]$, que l'on ne souhaite pas probabiliser.

Modèle initial de méconnaissances

Un découpage intuitif en sous-structures de la structure étudiée nous permet de dégager trois groupes d'éléments : le groupe g_2 , comprenant les deux éléments situés autour de la liaison, est encadré par les groupes g_1 et g_3 , comptant chacun 49 éléments. La situation est représentée sur la figure 7.5. Le modèle théorique déterministe considéré est le modèle recalé dans le paragraphe précédent.

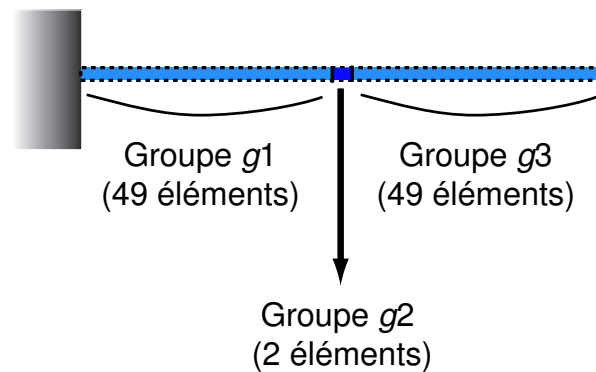


FIG. 7.5 – Définition des sous-structures du problème

Le modèle initial de méconnaissances de base est précisé dans le tableau 7.2 ; on estime que les seules incertitudes sur les poutres sont dues aux dispersions matériaux, d'où l'hypothèse d'une loi normale pour les groupes $g1$ et $g3$; au niveau de la liaison par contre, aucune loi de probabilité n'est supposée, vu que l'on considère que le modèle théorique est très approché.

Groupes	Loi choisie	Domaine $[-\overline{m}_E^{-0}; \overline{m}_E^{+0}]$	Moments statistiques relatifs (μ : moyenne / σ : écart type)
$g1$	normale	$[-0,25; 0,25]$	$\mu = 0,00 / \sigma = 0,097$
$g2$	pas de loi	$[-0,25; 0,25]$	aucun
$g3$	normale	$[-0,25; 0,25]$	$\mu = 0,00 / \sigma = 0,097$

TAB. 7.2 – Modèle initial ($\overline{m}_E^{-0}$, $\overline{m}_E^{+0}$) de méconnaissances de base dans le cas des poutres reliées par une liaison

Mise en œuvre du processus de réduction

L'observation des différentes énergies de déformation modales pour les trois groupes permet de sélectionner l'information expérimentale la plus pertinente pour la réduction de la méconnaissance de base sur chaque groupe. Les réductions successives sont précisées dans le tableau 7.3 ; on remarquera en particulier que l'on a mené la réduction sur les groupes $g1$ et $g3$ simultanément à partir des informations « expérimentales » correspondant au troisième mode.

Gr.	Données expérimentales	Loi	Méconnaissances après réduction $[-\overline{m}_E^{-}; \overline{m}_E^{+}]$	Caractéristiques
$g1$ et $g3$	Mode 3	normale	$[-0,000; 0,000]$	$\mu = 0,000 / \sigma = 0,000$
$g2$	Mode 1	pas de loi	$[-0,000; 0,013]$	aucune

TAB. 7.3 – Résultats de la réduction des méconnaissances de base dans le cas des deux poutres reliées par une liaison

La seule méconnaissance non nulle concerne le groupe g_2 , c'est-à-dire les deux éléments situés de part et d'autre de la liaison entre les poutres ; plus précisément, l'intervalle résultant de la modélisation supposée déterministe des méconnaissances de base sur ce groupe permet de préciser un écart résiduel de 1,3 % vis-à-vis de l'énergie de déformation modale associée au modèle théorique. Il est intéressant de noter la correspondance de cette valeur avec l'erreur en relation de comportement résiduelle obtenue après le processus de recalage présenté dans le paragraphe 7.2.2, à savoir une contribution de 1,26 % concentrée dans les deux éléments du groupe g_2 .

Ainsi, **l'erreur résiduelle après recalage traduit bel et bien une méconnaissance du modèle théorique utilisé vis-à-vis de la réalité simulée** ; il est intéressant de constater la similarité des valeurs obtenues. Bien sûr, **au sein de la démarche générale de validation de modèles, la théorie des méconnaissances va plus loin que les techniques de recalage**, vu qu'elle permet de manipuler un modèle imparfait par rapport à la réalité afin d'obtenir des valeurs prédictives sur des quantités d'intérêt données, et ce simplement en enrichissant le modèle déterministe recalé avec les méconnaissances de base déterminées lors du processus de réduction. La démarche utilisée permet donc d'étendre très facilement le champ d'application du modèle initial.

Application de la théorie des méconnaissances à un cas industriel

Il n'y a rien de plus beau qu'une clef tant que l'on ne sait pas ce qu'elle ouvre.

Maurice Maeterlinck

Sommaire

8.1	Structure industrielle étudiée	134
8.1.1	Présentation du SYLDA5	134
8.1.2	Description des mesures expérimentales effectuées	136
8.1.3	Modèle associé	136
8.2	Recalage du SYLDA5	137
8.2.1	Itérations du recalage	137
8.2.2	Bilan du recalage	138
8.3	Réduction des méconnaissances de base du SYLDA5	140
8.3.1	Données du problème	140
8.3.2	Processus de réduction des méconnaissances de base	141

Dans ce huitième et dernier chapitre est présentée l'application de la théorie des méconnaissances à un véritable cas d'étude industriel : le support de satellites SYLDA5 contenu dans le lanceur Ariane 5 développé par le groupe EADS (*European Aeronautic Defence and Space company*) et étudié de façon détaillée dans [Le Loch 2003]. Le processus de réduction des méconnaissances de base est mené à l'aide de résultats d'essais expérimentaux sur la structure réelle.

8.1 Structure industrielle étudiée

8.1.1 Présentation du SYLDA5

Le SYstème de Lancement Double d'Ariane 5 (SYLDA5) est un support contenu dans la coiffe du lanceur qui permet à ce dernier de mettre en orbite deux satellites lors d'un même lancement : une première charge utile prend place au-dessus du SYLDA5 tandis qu'une deuxième charge utile est à l'intérieur, comme représenté sur la figure 8.1.

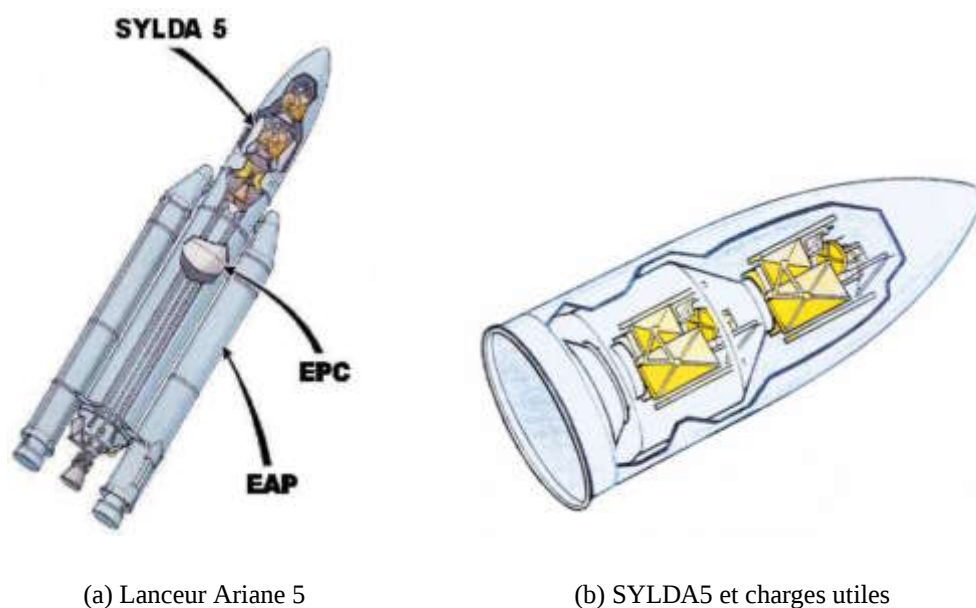
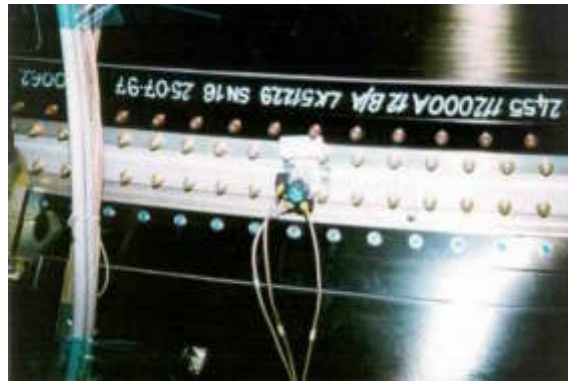


FIG. 8.1 – Principe du support SYLDA5

Cette structure est réalisée par un assemblage de cônes et de cylindres en composite sandwich, dont la peau est faite de stratifiés en carbone/époxy, et l'âme est en « nid d'abeille » (NIDA) d'aluminium. Les différentes parties sont reliées par des liaisons collées, à part la liaison boulonnée du « Système de Séparation du SYLDA5 » (SSS) ; cette liaison, représentée sur la figure 8.2, contient le cordon pyrotechnique permettant de séparer le SYLDA5 du lanceur avant le lancement de la charge utile basse.



(a) Support SYLDA5

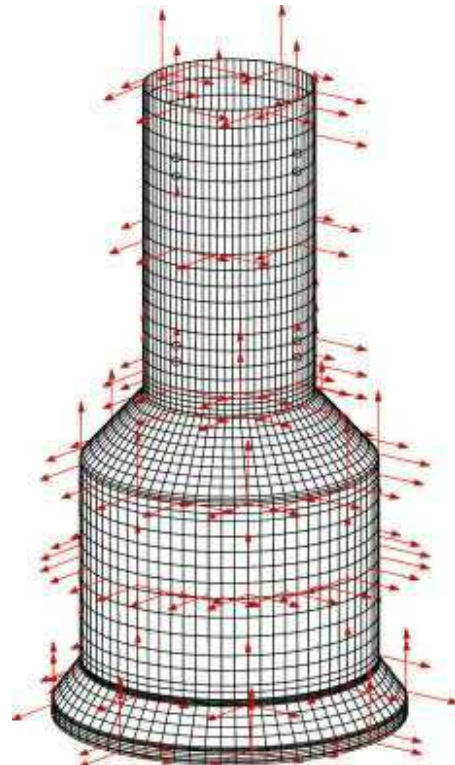


(b) Détail de la liaison SSS

FIG. 8.2 – Constitution du SYLDA5



(a) Configuration de test



(b) Positionnement des capteurs

FIG. 8.3 – Protocole d'expérimentation du SYLDA5

8.1.2 Description des mesures expérimentales effectuées

Une campagne d'essais a été menée en 1998 par IABG pour DASA/Dornier (l'Agence Spatiale allemande) sous contrat avec le Centre National d'Études Spatiales. L'ensemble testé met en jeu, de haut en bas :

- une structure représentant la « charge utile haute » (CUH), composée d'un cylindre nervuré en acier de 3,3 tonnes, surmonté d'une poutre en acier de 1,35 tonne ;
- l' « adaptateur de charge utile » (ACU), qui est un anneau en aluminium réalisant la liaison entre le SYLDA5 et la CUH ;
- le SYLDA5 proprement dit ;
- la liaison avec le sol, composée d'un anneau en aluminium.

Toutes ces structures sont reliées les unes aux autres par des liaisons boulonnées ; la hauteur totale de l'ensemble est d'environ 10 mètres, pour 5 mètres de diamètre.

Deux types d'essais ont été réalisés :

- un balayage en fréquence d'excitation forcée, avec 110 capteurs et 5 pots vibrants ;
- une appropriation modale avec 260 capteurs et 5 pots vibrants.

La configuration de test, ainsi que le positionnement des différents capteurs, sont représentés sur la figure 8.3. Il est important de noter que le sol de la salle d'expérimentation a vibré de façon non négligeable lors des mesures, ce dont on tiendra compte dans le modèle associé.

8.1.3 Modèle associé

Le modèle Éléments Finis associé utilise des éléments plaques pour les parties cylindriques et coniques, et des éléments poutres pour les différentes liaisons collées, ainsi que pour la liaison boulonnée SSS.

Le matériau sandwich constituant le corps du SYLDA5 est modélisé par un matériau orthotrope dont les caractéristiques ont été déterminées par homogénéisation des plis de composite et de l'âme en NIDA d'aluminium. Les nervures de la CUH n'ont pas été reproduites géométriquement dans le modèle ; c'est pourquoi un matériau orthotrope est utilisé dans le modèle afin de représenter les raideurs dans les différentes directions du cylindre.

Comme on l'avait souligné dans le paragraphe précédent, le sol de la salle d'expérimentation s'est déformé de façon sensible lors de la réalisation des mesures. Il a donc été décidé de modéliser d'une façon extrêmement simple le sol à l'aide d'un ressort ayant trois raideurs en rotation et une raideur en translation longitudinale ; de plus, une contrainte de mouvement de corps rigide est ajoutée à tous les nœuds de la base en contact avec le sol.

Du fait du grand nombre de matériaux et de sous-parties en présence, le modèle Éléments Finis est divisé en 32 sous-structures, dont les caractéristiques matériaux sont définies par 64 paramètres différents de masses et raideurs ; il comporte au final 9728 éléments et 27 648 degrés de liberté. C'est pourquoi une base modale de réduction, composée des 40 premiers modes propres et des 265 modes statiques associés aux excitateurs et aux capteurs, est utilisée afin de diminuer les coûts de calcul.

Les premiers modes propres associés à ce modèle sont représentés sur la figure 8.4 ; les modes de flexion sont doubles à cause de l'axisymétrie de la géométrie et des matériaux. On peut ensuite observer des modes d'ovalisation du SYLDA et de la CUH.

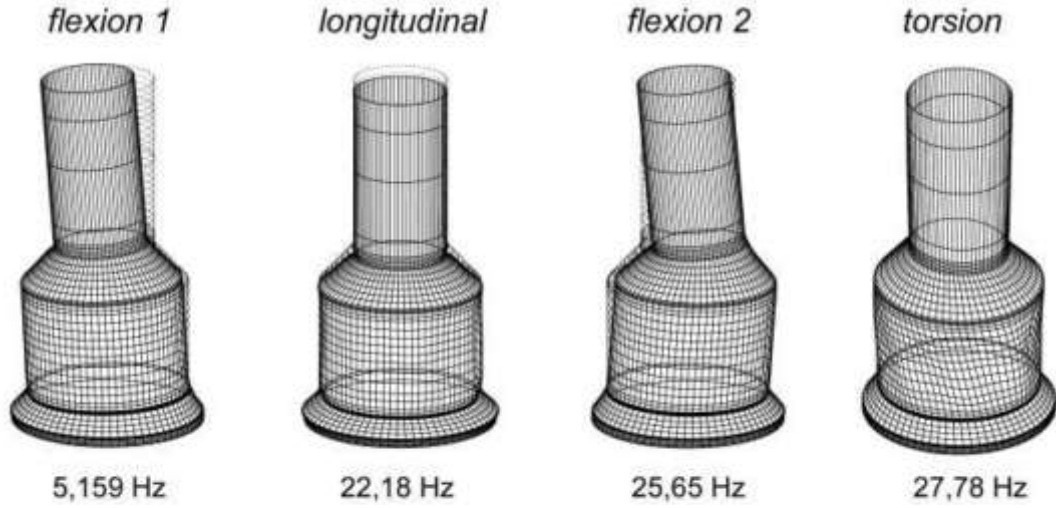


FIG. 8.4 – Premiers modes propres théoriques de l'ensemble d'étude du SYLDA5

8.2 Recalage du SYLDA5

Le recalage résultant de la confrontation entre les prédictions du modèle et les résultats expérimentaux modaux décrits dans le paragraphe précédent est mené à l'aide du concept de l'erreur en relation de comportement modifiée, déjà présenté dans le paragraphe 7.2.2 ; les résultats de ce recalage modal proviennent du travail [Le Loch 2003] auquel on pourra se référer pour plus de détails, ainsi qu'à [Barthe *et al.* 2004].

8.2.1 Itérations du recalage

La qualité initiale du modèle est donnée par l'erreur en relation de comportement totale ξ_{Tot}^2 , égale à 12,4 % ; les valeurs de l'erreur en relation de comportement pour chacun des modes, ainsi que le détail des écarts relatifs en fréquence Δf , sont donnés pour les dix premiers modes dans le tableau 8.1 ; les valeurs du MAC^1 , non précisées ici, sont déjà excellentes (plus de 95 %). Mis à part les modes 7 et 8, très mal représentés, l'erreur initiale sur chaque mode est inférieure à 10 %.

¹Le MAC (*Modal Assurance Criterion* en anglais) est un critère de corrélation modèle-essais très utilisé, car simple à implémenter. Il est défini par :

$$MAC(\bar{\phi}_i, \phi_{j \text{ exp}}) = \frac{(\bar{\phi}_i^T \phi_{j \text{ exp}})^2}{(\bar{\phi}_i^T \bar{\phi}_i)(\phi_{j \text{ exp}}^T \phi_{j \text{ exp}})}$$

Une corrélation parfaite entre deux modes résulte en un MAC égal à 100 %.

Avant d'effectuer la première itération du processus de recalage, les valeurs mesurées en certains points sont retirées : elles proviennent de capteurs présentant une contribution très élevée au niveau du terme d'erreur sur les mesures, et pour lesquels une étude attentive des résultats délivrés a permis de conclure à leur défectuosité.

La première étape de localisation met en évidence des défauts de modélisation au niveau de la liaison entre l'ACU et le SYLDA5, ainsi que sur la CUH, comme on peut le constater sur la figure 8.5. Les paramètres correspondants sont ensuite corrigés ; leurs nouvelles valeurs sont mentionnées dans le tableau 8.2.

It. Modes	Initial(0)		1		2		3		4		5*	
	ξ (%)	Δf (%)	ξ (%)	Δf (%)	ξ (%)	Δf (%)	ξ (%)	Δf (%)	ξ (%)	Δf (%)	ξ (%)	Δf (%)
1	3,45	-3,51	2,65	-2,69	0,24	0,28	2,05	2,05	1,67	1,86	1,26	-1,35
2	3,25	-3,31	2,46	-2,49	0,39	0,37	2,19	2,15	1,54	1,43	1,25	-1,18
3	5,10	-5,24	4,06	-4,14	0,40	-0,40	1,49	1,48	1,49	1,48	0,18	0,18
4	10,29	-10,85	9,78	-10,29	9,03	-9,48	4,09	-4,15	4,28	-4,22	0,61	0,52
5	9,86	-10,39	9,36	-9,83	8,61	-8,99	3,66	-3,75	3,76	-3,98	0,18	0,22
6 ⁺	3,47	3,45	1,91	1,90	0,10	0,10	2,66	2,63	0,00	0,00	0,96	-0,95
7	27,15	23,60	1,27	1,27	1,27	1,27	3,36	3,30	2,71	2,67	3,39	-3,34
8	27,77	24,09	1,92	1,90	1,92	1,90	4,00	3,92	3,36	3,30	4,04	-3,96
9	1,78	-1,79	0,09	-0,09	0,09	-0,09	0,02	0,02	0,10	-0,10	0,04	0,04
10	1,24	-1,24	0,44	0,44	0,44	0,44	0,55	0,55	0,43	0,43	0,50	-0,50
ξ_{Tot}	12,39		4,32		3,69		2,62		2,27		1,66	

* Nouvelle étape de correction à la suite de la prise en compte de masses additionnelles dans le modèle

⁺ Mode permuté avec les modes 4 et 5

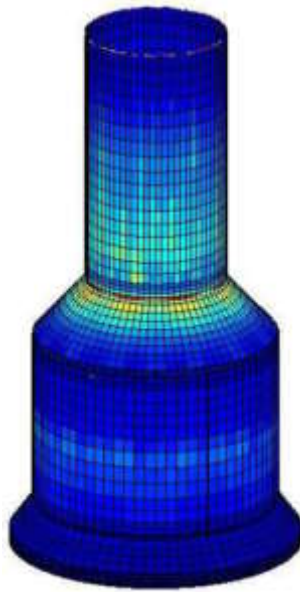
TAB. 8.1 – Erreur en relation de comportement et erreur en fréquence pour les dix premiers modes propres du SYLDA5

Un nouveau calcul de l'erreur totale en relation de comportement peut alors être mené ; il conduit à une erreur résiduelle de 4,32 %, pour laquelle les contributions se situent essentiellement au niveau des paramètres du modèle simplifié du sol, ce qu'on observe sur la droite de la carte de localisation de la figure 8.5 ; ces paramètres sont ensuite corrigés, et le processus itératif de recalage peut encore être mené à plusieurs reprises. Les corrections successives de paramètres sont données dans le tableau 8.2, ainsi que les niveaux de l'erreur pour chaque mode dans le tableau 8.1.

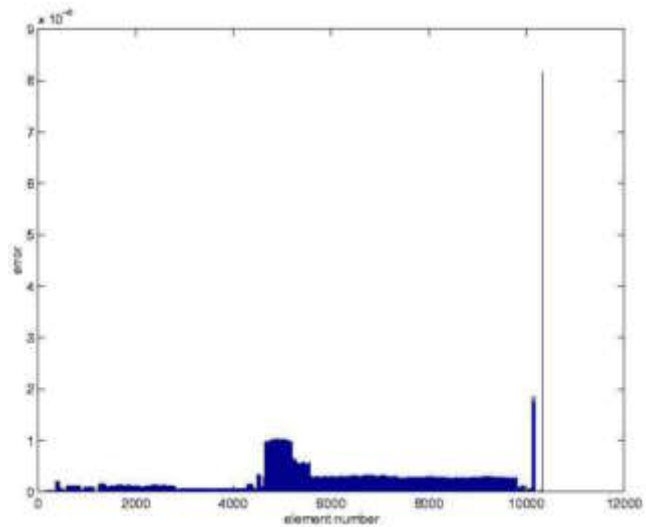
8.2.2 Bilan du recalage

Les différentes étapes du recalage que l'on vient de préciser illustrent bien la démarche de validation de modèles qui est menée par étapes successives, et qui permet de juger des approximations faites lors du choix des modèles constitutifs.

Ainsi, les principaux paramètres qui ont été corrigés correspondent à l'ACU et au sol, mais aussi à la CUH. On avait déjà évoqué pour cette dernière le choix d'une modélisation orthotrope du matériau plutôt que d'une représentation géométrique des nervures ; les valeurs forfaitaires initiales de raideur E_1 pour les éléments coques associés ont été fortement corrigées, notamment lors de la première étape de correction.



(a) Première localisation



(b) Deuxième localisation

FIG. 8.5 – Premières étapes de localisation lors du recalage du SYLDA5

Itérations		0	1	2	3	4	5*
ACU	poutre	1	0		0,76		1,13
	E_1	1			0,57		0,67
	E_2	1			1,56		1,84
	G_{12}	1			1,79		2,93
	lashing	1			0,09		0,05
Sol	T_x	1		0,25			0,41
	R_y	1		0,39		0,41	0,48
	R_z	1		0,40		0,47	0,54
	R_x	1		1,23		1,88	3,75
CUH	plaques E_1	1	6,36			6,51	6,51
	poutre	1	0,84				0,85
	ξ_{Tot} (%)	12,39	4,32	3,69	2,62	2,27	1,66

Les quantités indiquées sont des fractions des valeurs initiales des paramètres

* Nouvelle étape de correction à la suite de la prise en compte de masses additionnelles

TAB. 8.2 – Corrections successives des paramètres du modèle

Les paramètres de raideur pour le modèle du sol ont été corrigés lors de la seconde itération ; là encore, les valeurs initiales de ces derniers avaient été grossièrement estimées.

En ce qui concerne l'ACU, cette sous-structure complexe de liaison a été modélisée avec des éléments coques orthotropes et des éléments poutres, afin de représenter correctement le comportement de la liaison avec des paramètres de raideur équivalents. Une fois encore, il est logique que les paramètres associés aient été corrigés à plusieurs reprises.

Après la quatrième étape de correction, les contributions à l'erreur résiduelle se situent essentiellement au niveau des paramètres de raideur du sol, et l'erreur totale est relativement basse (2,27 %). Néanmoins, après l'obtention de ces résultats, une vérification du modèle a montré l'absence de prise en compte du matériel d'expérimentation utilisé pour l'acquisition des mesures ; des modifications concernant les paramètres de masse de certaines sous-structures ont donc été apportées, et une nouvelle itération a été menée, entraînant la correction de l'ensemble des paramètres précédemment impliqués, correction qui s'est toutefois révélée sensible seulement sur les paramètres de l'ACU.

8.3 Réduction des méconnaissances de base du SYLDA5

Nous déterminons maintenant les méconnaissances de base présentes sur les différentes sous-structures du SYLDA5 grâce aux mesures expérimentales disponibles.

8.3.1 Données du problème

Données expérimentales

Les données expérimentales sont composées des résultats de l'appropriation modale effectuée sur la structure de test dans la configuration décrite sur la figure 8.3. Vu que l'on n'a en fait qu'une seule structure expérimentale, on associe à chaque quantité d'intérêt expérimentale un intervalle de variation, du type $[0; \Delta\alpha_{\text{exp}}^+]$ ou $[-\Delta\alpha_{\text{exp}}^-; 0]$, que l'on ne souhaite pas probabiliser. Ici, seules les pulsations propres sont considérées.

Modèle initial de méconnaissances de base

Le modèle Éléments Finis obtenu après la procédure de recalage décrite dans le paragraphe précédent sert de base à la détermination des méconnaissances de base les plus représentatives de la réalité : il constitue le modèle théorique déterministe.

Afin de se placer dans le cadre de la théorie qui vise à globaliser l'effet des incertitudes à l'échelle des sous-structures, nous regroupons ces dernières en quatre grands groupes distincts, représentés sur la figure 8.6 :

- le groupe $g1$ est associé à la CUH ;
- le groupe $g2$ correspond à l'ACU ;
- le groupe $g3$ est le SYLDA5 proprement dit ;
- le groupe $g4$ est associé au modèle de sol.

En ce qui concerne le modèle initial de méconnaissances de base, on suppose que les seules incertitudes présentes dans les groupes $g1$ et $g3$ sont les dispersions matériaux, d'où le choix de lois de probabilité de type normal ; pour les liaisons, modélisées de façon quelque peu simplifiée, on suppose des lois de type uniforme. Les méconnaissances de base initiales sont résumées dans le tableau 8.3 ; l'étendue des lois a été choisie de telle façon que les méconnaissances effectives encadrent largement les valeurs expérimentales correspondantes.

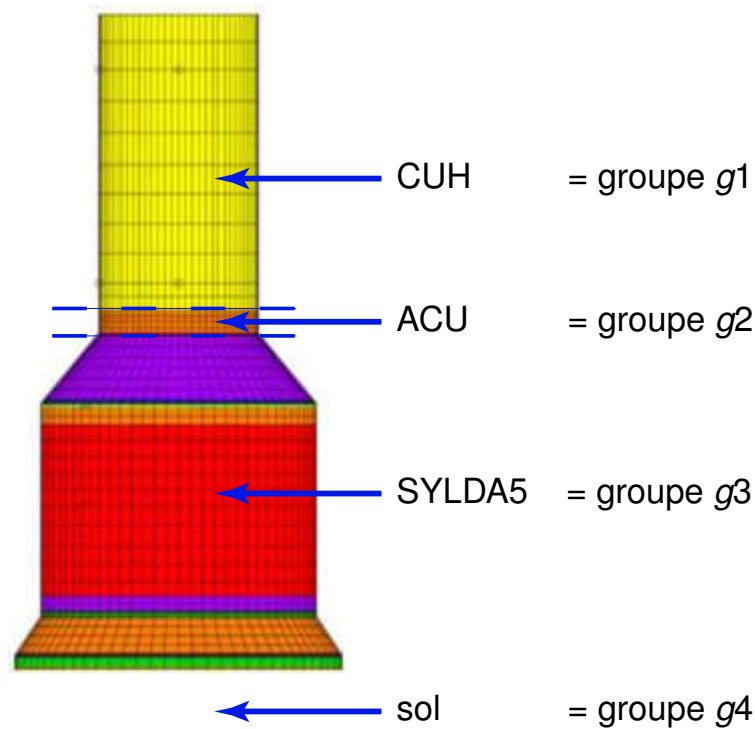


FIG. 8.6 – Détail des groupes de sous-structures utilisés dans le processus de réduction des méconnaissances de base du SYLDA5

Groupes	Loi choisie	Domaine $[-\overline{m}_E^{-0}; \overline{m}_E^{+0}]$	Moments statistiques relatifs (μ : moyenne / σ : écart type)
g_1	normale	$[-0, 20; 0, 20]$	$\mu = 0 / \sigma = 0,078$
g_2	uniforme	$[-0, 20; 0, 20]$	$\mu = 0 / \sigma = 0,115$
g_3	normale	$[-0, 20; 0, 20]$	$\mu = 0 / \sigma = 0,078$
g_4	uniforme	$[-0, 50; 0, 50]$	$\mu = 0 / \sigma = 0,289$

TAB. 8.3 – Modèle initial ($\overline{m}_E^{-0}$, $\overline{m}_E^{+0}$) de méconnaissances de base pour le SYLDA5

8.3.2 Processus de réduction des méconnaissances de base

Mise en œuvre du processus

Conformément à la stratégie précisée dans le paragraphe 6.5, les données expérimentales pertinentes pour la réduction de la méconnaissance de base sur une sous-structure donnée sont telles que l'énergie de déformation associée est caractérisée par une contribution prépondérante de la part de cette sous-structure. Les parts respectives de ces dernières sont représentées sur la figure 8.7.

On précise dans le tableau 8.4 dans quel ordre et avec quelles données est réalisée la réduction, ainsi que les méconnaissances de base obtenues après réduction. Les coefficients de représentativité ρ_E sont pris égaux à 1. Trois itérations ont été menées avant d'obtenir les méconnaissances réduites répertoriées dans le tableau.

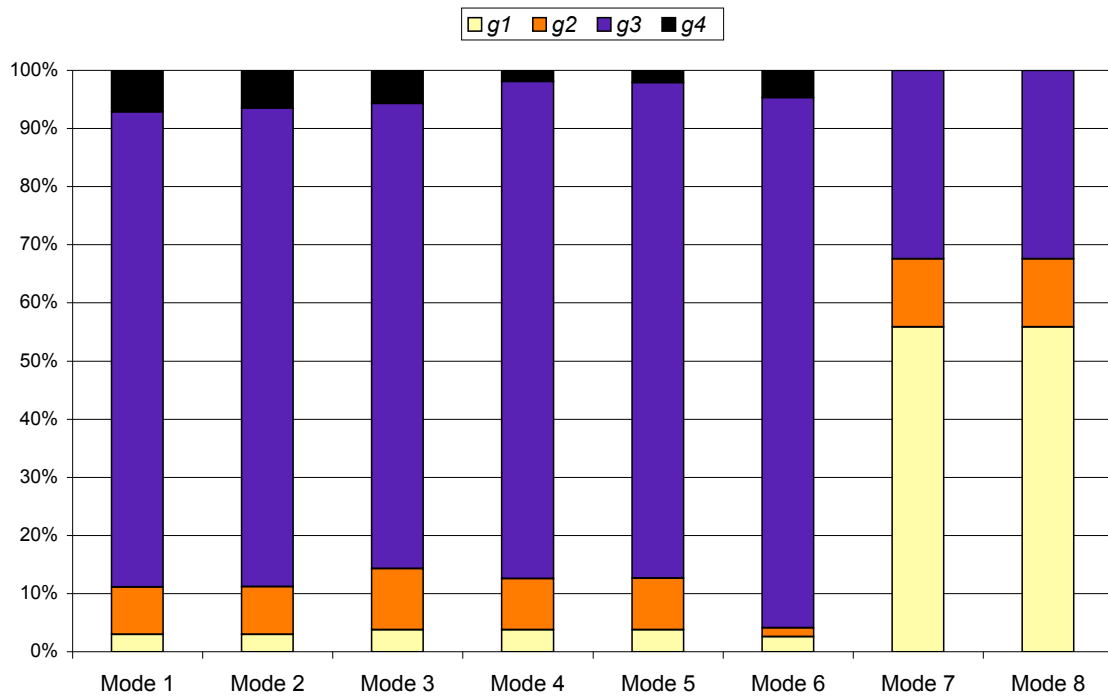


FIG. 8.7 – Énergies de déformation modales par groupe pour les huit premiers modes du SYLDA5

Gr.	Données expérimentales	Loi	Méconnaissances après réduction	
			$[-\overline{m}_E^-; \overline{m}_E^+]$	Caractéristiques
$g3$	$(\Delta\omega_{4\text{exp}}^{2+}, \Delta\omega_{4\text{exp}}^{2-})$	normale	$[-0,016; 0,000]$	$\mu = -0,008 / \sigma = 0,003$
$g1$	$(\Delta\omega_{8\text{exp}}^{2+}, \Delta\omega_{8\text{exp}}^{2-})$	normale	$[0,000; 0,144]$	$\mu = 0,072 / \sigma = 0,028$
$g4$	$(\Delta\omega_{6\text{exp}}^{2+}, \Delta\omega_{6\text{exp}}^{2-})$	uniforme	$[0,000; 0,435]$	$\mu = 0,218 / \sigma = 0,126$
$g2$	$(\Delta\omega_{3\text{exp}}^{2+}, \Delta\omega_{3\text{exp}}^{2-})$	uniforme	$[-0,060; 0,000]$	$\mu = -0,030 / \sigma = 0,012$

TAB. 8.4 – Résultats de la réduction des méconnaissances de base du SYLDA5

Prise en compte d'une forte méconnaissance

Ces résultats montrent en particulier que le groupe $g4$ correspondant au modèle simplifié du sol est caractérisé par une méconnaissance de base très dispersée, ce qui suggère de recommencer à zéro le processus de réduction des méconnaissances à l'aide des expressions tenant compte de la « forte » méconnaissance sur le groupe $g4$, afin de vérifier ou non la validité des calculs précédents. La base de réduction utilise les huit premiers modes propres du modèle déterministe ; les très bonnes valeurs de MAC entre ces modes théoriques et les modes expérimentaux correspondants suggèrent que le choix de cette base de réduction est valide et n'aura pas d'effet majeur sur le processus de réduction.

Les différentes itérations sont précisées dans le tableau 8.5 ; l'ordre des réductions successives n'a pas évolué en raison de la très faible influence de la forte méconnaissance dans l'estimation des contributions des autres groupes dans les énergies de déformation. C'est aussi pour cette raison que les méconnaissances après réduction pour ces groupes sont les mêmes que celles obtenues sans prendre en compte la forte méconnaissance.

Gr.	Données expérimentales	Loi	Méconnaissances après réduction	
			$[-\overline{m}_E^-; \overline{m}_E^+]$	Caractéristiques
$g3$	$(\Delta\omega_{4\text{exp}}^{2+}, \Delta\omega_{4\text{exp}}^{2-})$	normale	$[-0,016; 0,000]$	$\mu = -0,008 / \sigma = 0,003$
$g1$	$(\Delta\omega_{8\text{exp}}^{2+}, \Delta\omega_{8\text{exp}}^{2-})$	normale	$[0,000; 0,144]$	$\mu = 0,072 / \sigma = 0,028$
$g4$	$(\Delta\omega_{6\text{exp}}^{2+}, \Delta\omega_{6\text{exp}}^{2-})$	uniforme	$[0,000; 0,521]$	$\mu = 0,261 / \sigma = 0,150$
$g2$	$(\Delta\omega_{3\text{exp}}^{2+}, \Delta\omega_{3\text{exp}}^{2-})$	uniforme	$[-0,060; 0,000]$	$\mu = -0,030 / \sigma = 0,012$

TAB. 8.5 – Résultats de la réduction des méconnaissances de base du SYLDA5, avec prise en compte de la forte méconnaissance sur le modèle associé au sol

La méconnaissance de base associée au sol est caractérisée par une loi un peu plus dispersée que celle déterminée dans le tableau 8.4 ; la différence est suffisamment sensible pour que l'on puisse estimer que l'hypothèse de forte méconnaissance doit être retenue pour tous les calculs menés à propos du SYLDA5.

Quelques commentaires sur les valeurs de méconnaissances de base obtenues

Les résultats confirment la très bonne qualité du modèle du SYLDA5 proprement dit (groupe $g3$), dont les coefficients associés à la description du matériau avaient été obtenus par homogénéisation : le niveau de méconnaissance de base obtenu (1,6 %) est proche de la valeur finale de l'erreur résiduelle totale (1,66 %), ce qui semble juste vu que l'essentiel de l'énergie de déformation pour chaque mode étudié provient du groupe $g3$. On note également la bonne représentativité du modèle de l'ACU (groupe $g2$), pour lequel le processus de recalage a permis de trouver les meilleures valeurs des paramètres de raideur équivalente.

En ce qui concerne les deux autres groupes, si l'on n'est que peu surpris de la forte valeur de la méconnaissance de base (52,1 %) sur le modèle du sol (groupe $g4$), vu la simplicité de la modélisation utilisée et le faible poids énergétique associé, on peut s'étonner du niveau relativement élevé (14,4 %) de la méconnaissance sur la CUH (groupe $g1$). Cependant, quelques remarques s'imposent : le modèle associé au groupe $g1$ est volontairement simplifié, les nervures n'étant représentées que par leur influence sur la rigidité de la sous-structure, et on constate d'ailleurs que les modes qui font le plus intervenir énergétiquement la CUH, comme les modes 7 et 8 par exemple, sont ceux qui sont encore les plus erronés après les cinq itérations entreprises ; finalement, le niveau obtenu pour le groupe $g1$ semble logique.

Capacité de prédiction

Afin de vérifier la qualité des méconnaissances de base obtenues à l'aide du processus de réduction, on compare dans le tableau 8.6 les valeurs expérimentales à 99 % avec la méconnaissance effective associée à la première pulsation propre de la structure ; en effet, on n'avait pas utilisé l'information associée à cette quantité d'intérêt dans le processus de réduction des méconnaissances.

Les prédictions obtenues avec le modèle de méconnaissances déterminé plus haut sont en bon accord avec les valeurs expérimentales correspondantes, ce qui montre la représentativité du modèle vis-à-vis des données expérimentales disponibles.

i	$\bar{\omega}_i^2$	$[\omega_{i\text{exp}}^2; \omega_{i\text{exp}}^2]$	$[\bar{\omega}_i^2 - \Delta\omega_{i\text{mod}}^2; \bar{\omega}_i^2 + \Delta\omega_{i\text{mod}}^2]$
1	$1,02.10^3$	$[1,02.10^3; 1,05.10^3]$	$[1,01.10^3; 1,06.10^3]$

TAB. 8.6 – Comparaison modèle-expérimental pour la première pulsation propre du SYLDA5

Au final, la théorie des méconnaissances nous a permis de quantifier très facilement, et avec un coût de calcul peu important, la qualité des différents modèles des sous-structures constitutives de la structure industrielle étudiée, mais aussi d'enrichir le modèle théorique initial à l'aide du concept de méconnaissances de base, afin de calculer facilement des prédictions sur des quantités d'intérêt données.

La méthode permet aussi de nous orienter vers les améliorations qu'il faudrait peut-être apporter aux modèles, ainsi que vers les nouvelles expérimentations envisageables afin de mieux comprendre le comportement de la structure.

CONCLUSION

Dans ce mémoire a été exposée dans le détail la théorie des méconnaissances, dont la finalité est de proposer un concept permettant d'utiliser un modèle, aussi imparfait soit-il par rapport à la réalité, afin d'obtenir des informations prédictives sur toutes sortes de quantités d'intérêt.

Un premier point a défini précisément le concept de méconnaissance : ainsi, notre modélisation consiste à enrichir un modèle théorique déterministe de quantités, appelées méconnaissances de base, qui permettent de globaliser à l'échelle des sous-structures l'ensemble des incertitudes présentes dans la réalité étudiée. Divers outils mathématiques facilitant la description des méconnaissances de base ont été également introduits, et différents types de lois de probabilité ont été proposés, puis éprouvés dans la suite, où **notre première contribution a permis de montrer la facilité de la mise en œuvre d'une propagation des méconnaissances de base**, afin de déterminer pour des quantités d'intérêt courantes en dynamique des intervalles de variation dont les bornes sont probabilisées, et ainsi d'obtenir des informations prédictives sur la réalité observée.

Nous nous sommes alors intéressés à la détermination des méconnaissances de base les plus représentatives de la dispersion constatée expérimentalement sur la famille de structures étudiées, ce qui constitue notre seconde contribution. Une technique dite de réduction a été mise en place afin de partir d'une description majorante, définie *a priori*, des méconnaissances de base sur les différentes sous-structures, et de réduire ces dernières à l'aide d'informations expérimentales pertinentes. Les grandes lignes d'une stratégie générale ont été dégagées, afin de pouvoir optimiser la détermination de ces méconnaissances de base, et ont été éprouvées avec succès sur un premier exemple académique.

Nous avons aussi cherché à placer la théorie des méconnaissances dans la démarche générale de validation d'un modèle mécanique ; nous avons ainsi traité l'exemple d'un modèle de liaison très simplifié par rapport à la réalité des phénomènes en appliquant une technique classique de recalage à partir des données expérimentales disponibles, qui utilise le concept de l'erreur en relation de comportement modifiée. Nous avons alors constaté que l'erreur résiduelle présente à la fin du processus de recalage, sans qu'aucune correction de paramètre ne soit encore possible, était parfaitement décrite par le niveau de méconnaissance de base sur le modèle de liaison concerné.

Enfin, **notre dernière contribution a consisté à appliquer notre stratégie de réduction des méconnaissances de base à un exemple de structure réelle**, à savoir le support de satellites SYLDA5 du lanceur Ariane 5 ; nous avons ainsi pu conclure sur les qualités respectives des différents modèles de sous-structures employés.

L'ensemble de ce travail a donc permis de juger l'apport de la théorie des méconnaissances dans le domaine de la validation de modèles ; elle rend en effet possible l'évaluation de la dispersion de quantités d'intérêt définies pour une famille de structures prati-

quement identiques, en associant à un modèle théorique déterministe, pourtant imparfait par rapport à la réalité des phénomènes présents dans les structures étudiées, une description des méconnaissances de base sur l'ensemble des sous-structures constitutives de ces structures.

Dans cet état des lieux très positif subsistent toutefois quelques éléments susceptibles d'être améliorés.

Tout d'abord, **le choix des lois de probabilité pour la modélisation des méconnaissances de base doit être étudié de manière plus approfondie** afin de déterminer s'il est possible de le justifier de façon plus rigoureuse, et également afin d'améliorer la précision des résultats de la propagation de ces méconnaissances à travers le modèle. De plus, comme nous avons montré le positionnement de la théorie des méconnaissances vis-à-vis de la très vaste famille des différentes méthodes mathématiques de modélisation de l'incertitude, il est peut-être envisageable d'approfondir la façon de représenter l'incertitude dans notre domaine du calcul de structures en mécanique, et par là-même proposer d'autres pistes mathématiques pour la modélisation des méconnaissances de base.

Ensuite, **nous avons juste esquissé dans notre sixième partie le rôle du coefficient de représentativité dans le processus de réduction des méconnaissances de base** ; nous avons vu dans un exemple académique traitant des vibrations de flexion d'une plaque orthotrope la nécessité de quantifier à quel point les données expérimentales utilisées dans le processus de réduction permettent une vision juste des méconnaissances de base présentes dans la sous-structure considérée. Nous avons néanmoins constaté combien il est difficile de fixer la valeur de ce coefficient *a priori*, et il serait donc intéressant de mener des études complémentaires dans cette voie.

Enfin, il nous faut citer les très nombreuses perspectives associées à la méthode. Pour le moment, la théorie des méconnaissances ne permet de prendre en compte que des incertitudes concernant les rigidités structurales ; **il faut maintenant envisager son extension à la description des méconnaissances concernant les excitations et l'environnement incertain entourant les structures réelles**. De telles extensions devraient permettre de traiter, à l'aide du concept de méconnaissances, des cas d'études qui relèvent de la conception robuste, autrement dit des problèmes de dimensionnement et d'optimisation en présence d'incertitudes de toutes sortes.

Une autre voie d'études complémentaires est liée à l'évaluation de la pertinence des données expérimentales que nous avons eu l'occasion de mener dans le processus de réduction des méconnaissances. Il serait intéressant d'approfondir cette évaluation, afin de pouvoir déterminer quels essais, complémentaires de ceux dont on dispose déjà, permettraient de réduire au mieux la méconnaissance attachée à un modèle de sous-structure donné. Il serait alors possible d'établir une stratégie du choix des meilleurs essais à effectuer *a priori* sur une structure donnée, permettant ainsi une optimisation importante des coûts liés à l'expérimentation.

Calcul de la méconnaissance effective sur un mode propre : quelques compléments

N'importe quelle idée semble personnelle dès qu'on ne se rappelle plus à qui on l'a empruntée.

Jules Renard

Sommaire

A.1	Expression de la sensibilité au premier ordre d'un mode propre	148
A.2	Calcul pratique dans le cas de structures à grand nombre de DDL . . .	150
A.3	Méconnaissance effective sur la projection d'un mode propre	151

Dans cette annexe, on indique de façon précise comment est calculée la méconnaissance effective sur la valeur en un degré de liberté d'un mode propre : c'est en particulier l'obtention du vecteur $\underline{U}$ traduisant la sensibilité au premier ordre d'un mode propre qui est détaillée. On déduit également de cette dernière un mode de calcul de la méconnaissance effective sur la projection d'un mode propre selon une direction donnée.

On rappelle que toutes les quantités qui sont surmontées d'un trait horizontal, $\overline{\phi}_i$ par exemple, concernent le modèle déterministe, sans pour autant signifier que ce modèle est un modèle moyen.

A.1 Expression de la sensibilité au premier ordre d'un mode propre

Comme on l'a évoqué dans le paragraphe 5.2.2, le calcul de la méconnaissance sur la valeur en un DDL d'un mode propre nécessite de pouvoir exprimer la variation du mode entre le modèle avec méconnaissances et le modèle déterministe.

Pour cela, le point de départ est le problème aux valeurs propres qui permet de déterminer pulsations propres et modes propres : on a, d'une part pour le modèle théorique,

$$(\overline{\mathbf{K}} - \overline{\omega}_i^2 \overline{\mathbf{M}}) \overline{\phi}_i = 0 \quad (\text{A.1})$$

et d'autre part, pour le modèle avec méconnaissances,

$$(\mathbf{K} - \omega_{i \text{ mod}}^2 \mathbf{M}) \phi_{i \text{ mod}} = 0 \quad (\text{A.2})$$

et les deux modèles sont reliés par :

$$\mathbf{K} = \overline{\mathbf{K}} + \Delta \mathbf{K} \quad (\text{A.3a})$$

$$\mathbf{M} = \overline{\mathbf{M}} \quad (\text{A.3b})$$

$$\omega_{i \text{ mod}}^2 = \overline{\omega}_i^2 + \Delta \omega_{i \text{ mod}}^2 \quad (\text{A.3c})$$

$$\phi_{i \text{ mod}} = \overline{\phi}_i + \Delta \phi_{i \text{ mod}} \quad (\text{A.3d})$$

En injectant (A.3) dans (A.2), on obtient :

$$(\overline{\mathbf{K}} + \Delta \mathbf{K} - (\overline{\omega}_i^2 + \Delta \omega_{i \text{ mod}}^2) \overline{\mathbf{M}}) (\overline{\phi}_i + \Delta \phi_{i \text{ mod}}) = 0 \quad (\text{A.4})$$

que l'on peut écrire, après simplification avec l'équation (A.1), et avec une approximation au premier ordre, sous la forme :

$$\overline{\mathbf{Z}}(\overline{\omega}_i) \Delta \phi_{i \text{ mod}} = \underline{\mathbf{B}}(\overline{\omega}_i) \quad (\text{A.5})$$

où :

$$\overline{\mathbf{Z}}(\overline{\omega}_i) = (\overline{\mathbf{K}} - \overline{\omega}_i^2 \overline{\mathbf{M}}) \quad (\text{A.6})$$

$$\underline{\mathbf{B}}(\overline{\omega}_i) = -(\Delta \mathbf{K} - \Delta \omega_{i \text{ mod}}^2 \overline{\mathbf{M}}) \overline{\phi}_i \quad (\text{A.7})$$

Or, de par la définition du problème aux valeurs propres (A.1) pour le modèle déterministe, la matrice $\bar{\mathbf{Z}}$ est singulière en $\bar{\omega}_i$, donc cette équation n'a pas forcément de solution ; en fait, des solutions existent si et seulement si :

$$\bar{\phi}_i^T \underline{B}(\bar{\omega}_i) = 0 \quad (\text{A.8})$$

ce qui se réécrit comme :

$$\bar{\phi}_i^T \underline{\mathbf{M}} \bar{\phi}_i \Delta \omega_{i \text{ mod}}^2 = \bar{\phi}_i^T \Delta \underline{\mathbf{K}} \bar{\phi}_i \quad (\text{A.9})$$

Dans le cas particulier où les modes propres sont normalisés par rapport à la matrice de masse ($\bar{\phi}_j^T \underline{\mathbf{M}} \bar{\phi}_i = \delta_{ji}$), on retrouve l'expression au premier ordre de $\Delta \omega_{i \text{ mod}}^2$, utilisée dans le calcul de la méconnaissance effective sur une pulsation propre, présenté dans le paragraphe 5.2.1.

Dans le cas où la condition A.8 est respectée, les solutions de l'équation (A.5) sont de la forme générale :

$$\Delta \phi_{i \text{ mod}} = \psi_i + a_i \bar{\phi}_i \quad (\text{A.10})$$

où ψ_i est une solution particulière, pour laquelle généralement on impose nulle une coordonnée donnée [Nelson 1976]. On peut donc éliminer la ligne et la colonne correspondantes afin d'inverser facilement une matrice $\bar{\mathbf{Z}}^*(\bar{\omega}_i)$ de taille $(n-1) \times (n-1)$.

La valeur de a_i est alors liée au choix de la condition de normalisation utilisée pour les modes ; ici, on choisit :

$$\max_k \bar{\phi}_{ki} = \bar{\phi}_{k^*i} = 1 \quad (\text{A.11})$$

Cette condition de normalisation, reportée également pour les modes propres du modèle avec méconnaissances, implique que $\Delta \phi_{k^*i \text{ mod}} = 0$, et si l'on impose pour la solution particulière $\psi_{k^*i} = 0$, on aboutit nécessairement à $a_i = 0$, et donc à :

$$\Delta \phi_{i \text{ mod}} = \psi_i \quad (\text{A.12})$$

Si l'on s'intéresse à la k ième coordonnée du mode, on obtient finalement :

$$\Delta \phi_{ki \text{ mod}} = \psi_{ki} = \underline{U}^T \Delta \underline{\mathbf{K}} \bar{\phi}_i \quad (\text{A.13})$$

avec :

$$\underline{U}^T = \underline{\Pi} \left(\underline{\mathbf{P}}^T \bar{\mathbf{Z}}^{*-1} \underline{\mathbf{P}} \right) \left(\frac{1}{\bar{\phi}_i^T \underline{\mathbf{M}} \bar{\phi}_i} \underline{\mathbf{M}} \bar{\phi}_i \bar{\phi}_i^T - \underline{\mathbf{I}}_n \right) \quad (\text{A.14})$$

où $\underline{\mathbf{I}}_n$ est la matrice identité de dimension $n \times n$ et :

$$\underline{\mathbf{P}} = \left(\begin{array}{c|c} \underline{\mathbf{I}}_{k^*-1} & \underline{\mathbf{0}} \\ \hline \underline{\mathbf{0}} & \underline{\mathbf{I}}_{n-k^*} \end{array} \right) \quad \text{et} \quad \underline{\Pi} = (0 \dots 0 | 1 | 0 \dots 0) \quad (\text{A.15})$$

Pour que les trois termes du produit (A.13) soient d'ordres de grandeur comparables, on normalise également le champ $\underline{U}$ en posant $\underline{U}^* = \frac{1}{\hat{U}} \underline{U}$, avec $\hat{U} = \max_j |U_j|$, d'où :

$$\psi_{ki} = \underline{U}^{*T} (\hat{U} \Delta \underline{\mathbf{K}}) \bar{\phi}_i \quad (\text{A.16})$$

C'est finalement ce champ $\underline{U}^*$ qui sert à déterminer la méconnaissance effective sur la valeur au DDL k du i ème mode propre, comme précisé dans le paragraphe 5.2.2.

Dans le cas de modes multiples, la méthode présentée ici n'est plus valable directement, mais de nombreuses références traitent cette situation, comme [Shaw et Jayasuria 1992] ou [Zhang et Wei 1995].

A.2 Calcul pratique de la méconnaissance effective sur la valeur en un degré de liberté d'un mode propre dans le cas de structures à grand nombre de degrés de liberté

Dans le cas de structures industrielles, pour lesquelles le nombre de DDL est souvent très important (de l'ordre de plusieurs milliers), le calcul de champ $\underline{U}$ traduisant la sensibilité au premier ordre d'un mode propre donné s'avère très coûteux s'il est réalisé selon la méthode décrite dans le paragraphe précédent.

Classiquement, on utilise une base de réduction qui permet de traiter un problème de dimension bien plus faible : les modes du modèle avec méconnaissance $\underline{\phi}_{i \text{ mod}}$ sont cherchés sous la forme $\underline{\phi}_{i \text{ mod}} = \mathbf{T} \underline{\phi}_{i \text{ mod}}^R$, avec $\underline{\phi}_{i \text{ mod}}^R$ solution du problème aux valeurs propres :

$$\mathbf{T}^T (\mathbf{K} - \omega_{i \text{ mod}}^2 \mathbf{M}) \mathbf{T} \underline{\phi}_{i \text{ mod}}^R = 0 \quad (\text{A.17})$$

où $\mathbf{T}$ est la matrice de réduction, de dimension $r \times n$ avec $r \ll n$, généralement constante. Une expression similaire peut s'établir pour le modèle théorique déterministe.

On peut alors procéder de la même façon que pour le système de taille complète pour déterminer la variation du i ème mode propre ; on a en effet $\Delta \underline{\phi}_{i \text{ mod}} = \mathbf{T} \Delta \underline{\phi}_{i \text{ mod}}^R$ avec :

$$\Delta \underline{\phi}_{i \text{ mod}}^R = \underline{\psi}_i^R + a_i \bar{\underline{\phi}}_i^R \quad (\text{A.18})$$

et $\underline{\psi}_i^R$ est la solution particulière de :

$$\bar{\mathbf{Z}}^R(\bar{\omega}_i) \Delta \underline{\phi}_{i \text{ mod}}^R = \underline{B}^R(\bar{\omega}_i) \quad (\text{A.19})$$

où :

$$\bar{\mathbf{Z}}^R(\bar{\omega}_i) = \mathbf{T}^T (\bar{\mathbf{K}} - \bar{\omega}_i^2 \bar{\mathbf{M}}) \mathbf{T} \quad (\text{A.20})$$

$$\underline{B}^R(\bar{\omega}_i) = -\mathbf{T}^T (\Delta \mathbf{K} - \Delta \omega_{i \text{ mod}}^2 \bar{\mathbf{M}}) \mathbf{T} \bar{\underline{\phi}}_i^R \quad (\text{A.21})$$

La formule de Nelson décrite précédemment pour le système complet peut tout à fait s'appliquer directement au problème réduit, et permet donc de déterminer facilement la

solution particulière $\underline{\psi}_i^R$.

La question qui reste encore en suspens est le choix de la base $\mathbf{T}$ de réduction ; de manière classique [Fox et Kapoor 1968], ce sont les r premiers modes $\bar{\phi}_i$ du modèle théorique déterministe qui sont retenus. Toutefois, cette option s'avère souvent peu efficace, et il est courant de compléter cette base avec des vecteurs supplémentaires ; de nombreuses possibilités ont d'ailleurs été répertoriées dans [Balmès 1998].

Notre choix, issu de [Bouazzouni et al. 1997], est le suivant : le second membre $\underline{B}(\bar{\omega}_i)$ de l'équation (A.5) traduisant la variation du mode propre $\bar{\phi}_i$ correspond à un effort représentatif de la modification ; il peut donc être intéressant de compléter la base $\mathbf{T}$ par les réponses statiques associées à cet effort, d'où la base de réduction finale suivante :

$$\mathbf{T} = (\bar{\phi}_1 \quad \dots \quad \bar{\phi}_r \quad \mathbf{K}^{-1}\underline{B}(\bar{\omega}_1^2) \quad \dots \quad \mathbf{K}^{-1}\underline{B}(\bar{\omega}_r^2)) \quad (\text{A.22})$$

A.3 Calcul pratique de la méconnaissance effective sur la projection d'un mode propre selon une direction donnée

On précise ici comment obtenir la relation (5.11) permettant de déterminer la méconnaissance effective sur la projection d'un mode propre selon une direction donnée.

En utilisant les relations du paragraphe précédent, avec comme base de réduction :

$$\mathbf{T} = (\bar{\phi}_1 \quad \dots \quad \bar{\phi}_r) \quad (\text{A.23})$$

comme indiqué dans [Fox et Kapoor 1968], on détermine la solution de l'équation (A.19) :

$$\underline{\psi}_i^R = \left(\frac{1}{\bar{\omega}_1^2 - \bar{\omega}_i^2} \bar{\phi}_1 \Delta \mathbf{K} \bar{\phi}_i \quad \dots \quad \frac{1}{\bar{\omega}_p^2 - \bar{\omega}_i^2} \bar{\phi}_p \Delta \mathbf{K} \bar{\phi}_i \right) \quad (\text{A.24})$$

Comme la condition de normalisation implique dans ce cas particulier que $a_i = 0$, on obtient finalement :

$$\Delta \phi_{i \text{ mod}} = \sum_{\substack{k=1 \\ k \neq i}}^r \frac{1}{\bar{\omega}_k^2 - \bar{\omega}_i^2} \sum_{E \in \Omega} \left\{ \bar{\phi}_k^T \Delta \mathbf{K}_E \bar{\phi}_i \right\} \bar{\phi}_k \quad (\text{A.25})$$

BIBLIOGRAPHIE

Axelsson 1994

O. Axelsson. *Iterative Solution Methods*. Cambridge University Press, 1994.

Babuska et Rheinboldt 1979

I. Babuska et W. C. Rheinboldt. Adaptive approaches and reliability estimators in finite element analysis. *Computer Methods in Applied Mechanics and Engineering*, 17/18 : 519–540, 1979.

Baecher et Ingra 1981

G. B. Baecher et T. S. Ingra. Stochastic FEM in settlement predictions. *Journal of the Geotechnical Engineering Division*, 107(4) : 449–463, 1981.

Balmès 1998

E. Balmès. Efficient sensitivity analysis based on finite element model reduction. Dans *Proceedings of the 16th International Modal Analysis Conference (IMAC-XVI)*, pages 1118–1124, Santa Barbara, Californie, février 1998.

Balmès 2000

E. Balmès. Review and evaluation of shape expansion methods. Dans *Proceedings of the 18th International Modal Analysis Conference (IMAC-XVIII)*, San Antonio, Texas, février 2000.

Barthe et al. 2003

D. Barthe, A. Deraemaeker, P. Ladevèze et G. Puel. On a theory of the quantification of the lack of knowledge in structural computation. Dans *Proceedings of the 21st International Modal Analysis Conference (IMAC-XXI)*, Kissimmee, Floride, février 2003.

Barthe et al. 2004

D. Barthe, A. Deraemaeker, P. Ladevèze et S. Le Loch. Validation and updating of industrial models based on the constitutive relation error. *AIAA Journal*, 42(7) : 1427–1434, 2004.

Bayes

T. Bayes. An essay towards solving a problem in the doctrine of chances. *Philosophical Transactions of the Royal Society*, LIII : 376–418, 1763.

Ben-Haim 1995

Y. Ben-Haim. A non-probabilistic measure of reliability of linear systems based on expansion of convex models. *Structural Safety*, 17 : 91–109, 1995.

Ben-Haim 1999

Y. Ben-Haim. Design certification with information-gap uncertainty. *Structural Safety*, 21 : 269–289, 1999.

Ben-Haim 2000

Y. Ben-Haim. Robustness of model-based fault diagnosis : decisions with information-gap models of uncertainty. *International Journal of Systems Science*, 31(11) : 1511–1518, 2000.

Ben-Haim 2001

Y. Ben-Haim. *Information-Gap Decision Theory*. Academic Press, 2001.

Ben-Haim et Elishakoff 1990

Y. Ben-Haim et I. Elishakoff. *Convex Models of Uncertainty in Applied Mechanics*. Elsevier, 1990.

Ben-Haim et Hemez 2004

Y. Ben-Haim et F. M. Hemez. Robustness-to-uncertainty, fidelity-to-data and prediction-looseness of models. Dans *Proceedings of the 22nd International Modal Analysis Conference (IMAC-XXII)*, Dearborn, Michigan, janvier 2004.

Bernoulli

J. Bernoulli. *Ars Conjectandi*. 1713.

Beven et Binley 1992

K. Beven et A. Binley. The future of distributed models : model calibration and uncertainty prediction. *Hydrological Processes*, 6 : 279–298, 1992.

Booker et Singpurwalla 2003

J. M. Booker et N. D. Singpurwalla. Using probability and fuzzy set theory for reliability assessment. Dans *Proceedings of the 21st International Modal Analysis Conference (IMAC-XXI)*, Kissimmee, Floride, février 2003.

Bouazzouni et al. 1997

A. Bouazzouni, G. Lallement et S. Cogan. Selecting a Ritz basis for the reanalysis of the frequency response functions of modified structures. *Journal of Sound and Vibration*, 199(2) : 309–322, 1997.

Boucard et Champaney 2003

P.-A. Boucard et L. Champaney. A suitable computational strategy for the parametric analysis of problems with multiple contact. *International Journal for Numerical Methods in Engineering*, 57(9) : 1259–1281, 2003.

Caignot 2004

A. Caignot. Validation de modèles en dynamique des structures dans le cadre stochastique. Mémoire de DEA, École Normale Supérieure de Cachan, LMT-Cachan, 2004.

Chebli 2002

H. Chebli. *Modélisation des incertitudes aléatoires non homogènes en dynamique des structures pour le domaine des basses fréquences*. Thèse de doctorat, Conservatoire National des Arts et des Métiers - Office National d'Études et de Recherche Aérospatiales, 2002.

Chen et Ward 1997

R. Chen et A. C. Ward. Generalizing interval matrix operations for design. *Journal of Mechanical Design*, 119 : 65–72, 1997.

Chipman et al. 1997

H. Chipman, M. Hamada et C. F. J. Wu. A bayesian variable-selection approach for analyzing designed experiments with complex aliasing. *Technometrics*, 39(4), 1997.

Chouaki 1997

A. Chouaki. *Recalage des modèles dynamiques de structures avec amortissement*. Thèse de doctorat, École Normale Supérieure de Cachan, LMT-Cachan, 1997.

Comba et Stolfi 1993

J. L. D. Comba et J. Stolfi. Affine arithmetic and its applications to computer graphics. Dans *Proceedings of SIBGRAP'93*, octobre 1993.

Craig et Bampton 1968

R. R. Craig et M. C. C. Bampton. Coupling of substructures of dynamic analyses. *AIAA Journal*, 6(7) : 1313–1319, 1968.

Dean et Lewis 2004

A. M. Dean et S. M. Lewis. *Screening*. Springer Verlag, 2004.

Découvreur et al. 2004

V. Découvreur, P. Bouillard, A. Deraemaeker et P. Ladevèze. Updating 2D acoustic models with the constitutive equation error method. *Journal of Sound and Vibration*, 278(4-5) : 773–787, 2004.

Dempster 1968

A. P. Dempster. Upper and lower probabilities generated by a random interval. *Annals of Mathematical Statistics*, 39(3) : 957–966, 1968.

Deodatis 1991a

G. Deodatis. Weighted integral method - I. Stochastic stiffness matrix. *Journal of Engineering Mechanics*, 117(8) : 1851–1864, 1991.

Deodatis 1991b

G. Deodatis. Weighted integral method - II. Response variability and reliability. *Journal of Engineering Mechanics*, 117(8) : 1865–1877, 1991.

Deraemaeker 2001

A. Deraemaeker. *Sur la maîtrise des modèles en dynamique des structures à partir de résultats d'essais*. Thèse de doctorat, École Normale Supérieure de Cachan, LMT-Cachan, 2001.

Deraemaeker et al. 2004

A. Deraemaeker, P. Ladevèze et T. Romeuf. Model validation in the presence of uncertain experimental data. *Engineering Computations*, 21(8) : 808–833, 2004.

Dessombz 2000

O. Dessombz. *Analyse dynamique de structures comportant des paramètres incertains*. Thèse de doctorat, École Centrale de Lyon, 2000.

Dessombz et al. 2001

O. Dessombz, F. Thouverez, J.-P. Laine et L. Jézéquel. Analysis of mechanical systems using interval computations applied to finite element methods. *Journal of Sound and Vibration*, 239(5) : 949–968, 2001.

Ditlevsen et Madsen 1996

O. Ditlevsen et H. O. Madsen. *Structural Reliability Methods*. John Wiley and Sons, 1996.

Dong et Shah 1987

W. Dong et H. C. Shah. Vertex method for computing functions of fuzzy variables. *Fuzzy Sets and Systems*, 24 : 65–78, 1987.

Dubois et Prade 1986

D. Dubois et H. Prade. *Possibility Theory : an Approach to Computerised Processing of Uncertainty*. Plenum Press, 1986.

Elishakoff et al. 1995

I. Elishakoff, Y. J. Ren et M. Shinozuka. Improved finite element method for stochastic problems. *Chaos, Solitons and Fractals*, 5(5) : 833–846, 1995.

Faravelli 1989

L. Faravelli. Response-surface approach for reliability analysis. *Journal of Engineering Mechanics*, 115(12) : 2673–2781, 1989.

Fishman 1996

G. S. Fishman. *Monte Carlo : Concepts, Algorithms and Applications*. Springer Verlag, 1996.

Fox et Kapoor 1968

R. Fox et M. Kapoor. Rate of change of eigenvalues and eigenvectors. *AIAA Journal*, 6 : 2426–2429, 1968.

Ghanem 1999

R. Ghanem. Ingredients for a general purpose stochastic finite elements implementation. *Computer Methods in Applied Mechanics and Engineering*, 168 : 19–34, 1999.

Ghanem et Kruger 1996

R. G. Ghanem et R. M. Kruger. Numerical solution of spectral stochastic finite element systems. *Computer Methods in Applied Mechanics and Engineering*, 129 : 289–303, 1996.

Ghanem et Spanos 1991

R. G. Ghanem et P. D. Spanos. *Stochastic Finite Elements : a Spectral Approach*. Springer, 1991.

Golinval et Collignon 1996

J.-C. Golinval et P. Collignon. Comparison of model updating methods adapted to local error detection. Dans *Proceedings of the 21st International Conference on Noise and Vibration Engineering (ISMA 21)*, Leuven, Belgique, septembre 1996.

Grigoriu 1993

M. Grigoriu. Simulation of nonstationary gaussian processes by random trigonometric polynomials. *Journal of Engineering Mechanics*, 119(2) : 328–343, 1993.

Grimette et Stirzaker 1991

G. R. Grimette et D. R. Stirzaker. *Probability and Random Process*. Oxford Science Publications, 1991.

Hammersley et Handscomb 1967

J. M. Hammersley et D. C. Handscomb. *Monte Carlo Methods*. Fletcher and Son Ltd, 1967.

Hanss 2002

M. Hanss. The transformation method for the simulation and analysis of systems with uncertain parameters. *Fuzzy Sets and Systems*, 130(3) : 1–13, 2002.

Hanss et al. 2002

M. Hanss, S. Oexl et L. Gaul. Identification of a bolted-joint model with fuzzy parameters loaded normal to the contact interface. *Mechanics Research Communications*, 29 : 177–187, 2002.

Hasofer et Lind 1974

A. M. Hasofer et N. C. Lind. Exact and invariant second-moment code format (an exact and invariant first order reliability format). *Journal of Engineering Mechanics*, 100 : 111–121, 1974.

Helton et Davis 2003

J. C. Helton et F. J. Davis. Latin hypercube sampling and the propagation of uncertainty in analyses of complex systems. *Reliability Engineering and System Safety*, 81 : 23–69, 2003.

Hemez 2004

F. M. Hemez. The good, the bad and the ugly of predictive science. Dans *Proceedings of the Fourth International Conference on Sensitivity Analysis of Modeling Output*, Santa Fe, Nouveau Mexique, mars 2004.

Hemez et al. 2001

F. M. Hemez, A. C. Wilson et S. W. Doebling. Design of computer experiments for improving an impact test simulation. Dans *Proceedings of the 19th International Modal Analysis Conference (IMAC-XIX)*, pages 977–985, Kissimmee, Floride, février 2001.

Hemez et al. 2004

F. M. Hemez, S. W. Doebling et M. C. Anderson. A brief tutorial on verification and validation. Dans *Proceedings of the 22th International Modal Analysis Conference (IMAC-XXII)*, Dearborn, Michigan, janvier 2004.

Humbert et al. 1999

L. Humbert, F. Thouverez et L. Jézéquel. Finite element dynamic model updating using modal thermoelastic fields. *Journal of Sound and Vibration*, 228(2) : 397–420, 1999.

IC

IC. Interval Computations. V. Kreinovich, 2004. <http://www.cs.utep.edu/interval-comp>

IPP

IPP. The Imprecise Probability Project. G. de Cooman et P. Walley, 2004. <http://ipp-serv.rug.ac.be>

Jaynes 1957a

E. T. Jaynes. Information theory and statistical mechanics. *Physical Review*, 106(4) : 620–630, 1957.

Jaynes 1957b

E. T. Jaynes. Information theory and statistical mechanics II. *Physical Review*, 108(2) : 171–190, 1957.

Johnson et al. 1997

E. A. Johnson, L. A. Bergman et B. F. Spencer. Parallel implementation of Monte Carlo simulation - comparative studies from stochastic structural dynamics. *Probabilistic Engineering Mechanics*, 12(4) : 208–212, 1997. (Special Issue on Computational Stochastic Mechanics).

Joslyn et Booker 2004

C. Joslyn et J. Booker. *Engineering Design Reliability Handbook*, chapitre Generalized Information Theory for Engineering Modeling and Simulation. CRC Press, éditeurs : E. Nikolaidis, D. M. Ghiocel, S. N. Singhal, 2004.

Kapur et Kesavan 1992

J. N. Kapur et H. K. Kesavan. *Entropy Optimization Principles with Applications*. Academic Press, 1992.

Kennedy et O'Hagan 2001

M. C. Kennedy et A. O'Hagan. Bayesian calibration of computer models (with discussion). *Journal of the Royal Statistical Society (Series B)*, 63 : 425–464, 2001.

Klir 1991

G. Klir. Generalized information theory. *Fuzzy Sets and Systems*, 40 : 127–142, 1991.

Klir 1994

G. J. Klir. *Uncertainty Modelling and Analysis : Theory and Applications*, chapitre The Many Faces of Uncertainty, pages 3–19. Elsevier, éditeurs : B. M. Ayyub, M. M. Gupta, 1994.

Kolmogorov 1956

A. N. Kolmogorov. *Foundations of the Theory of Probability*. Chelsea, 1956.

Ladevèze 1983

P. Ladevèze. Recalage de modélisations des structures complexes. Rapport technique 33.11.01.4, Aérospatiale, Les Mureaux, 1983.

Ladevèze 1993

P. Ladevèze. Erreur en relation de comportement en dynamique : théorie et application au recalage de modèles de structures. Rapport interne 150, LMT-Cachan, 1993.

Ladevèze 2002

P. Ladevèze. Sur une théorie des méconnaissances en calcul des structures. Programme RAF 2001 SY/XS 136 127, EADS Launch Vehicles, avril 2002.

Ladevèze 2003a

P. Ladevèze. Théorie des méconnaissances en dynamique des structures : développement et premiers exemples. Programme RAF 2002 LY22 139 423, EADS Launch Vehicles, mai 2003.

Ladevèze 2003b

P. Ladevèze. Validation et vérification de modèles stochastiques en environnement incertain par la méthode d l'erreur en relation de comportement. Rapport interne 258, LMT-Cachan, 2003.

Ladevèze et Chouaki 1999

P. Ladevèze et A. Chouaki. Application of a posteriori error estimation for structural model updating. *Inverse Problems*, 15 : 49–58, 1999.

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 Oden 1998

P. Ladevèze et J. T. Oden. *Advances in Adaptive Computational Methods in Mechanics*. Elsevier, 1998.

Ladevèze et Pelle 2001

P. Ladevèze et J.-P. Pelle. *La Maîtrise du Calcul en Mécanique linéaire et non-linéaire*. Hermès, 2001.

Ladevèze et al. 1994a

P. Ladevèze, D. Nedjar et M. Reynier. Updating of finite element models using vibration tests. *AIAA Journal*, 32(7) : 1485–1491, 1994.

Ladevèze et al. 1994b

P. Ladevèze, M. Reynier et N. M. M. Maia. *Inverse Problems in Engineering*, chapitre Error on the Constitutive Relation in Dynamics, pages 251–256. Balkema, éditeurs : H. D. Bui, et M. Tanaka, 1994.

Ladevèze et al. 2004

P. Ladevèze, G. Puel et T. Romeuf. On a strategy of reduction of the lack of knowledge (LOK) of a structural dynamics model. Dans *Proceedings of the 22nd International Modal Analysis Conference (IMAC XXII)*, Dearborn, Michigan, janvier 2004.

Ladevèze et al. 2005

P. Ladevèze, G. Puel, A. Deraemaeker et T. Romeuf. Validation of structural dynamics model with uncertainties. *Computer Methods in Applied Mechanics and Engineering*, 2005 (à paraître).

Lallemant et al. 1999

B. Lallemant, A. Cherki, T. Tison et P. Level. Fuzzy modal finite element analysis of structures with imprecise material properties. *Journal of Sound and Vibration*, 220(2) : 353–364, 1999.

Laplace

P. S. Laplace. *Mémoire sur la Probabilité des Causes par les Événements*. 1774.

Le Loch 2003

S. Le Loch. *Modélisation et identification de l'amortissement dans les structures spatiales*. Thèse de doctorat, École Normale Supérieure de Cachan, LMT-Cachan, 2003.

Lemaire 1997

M. Lemaire. Reliability and mechanical design. *Reliability Engineering and System Safety*, 55(2) : 163–170, 1997.

Liu et al. 1986a

W. K. Liu, T. Belytschko et A. Mani. Random field finite elements. *International Journal of Numerical Methods in Engineering*, 23 : 1831–1845, 1986.

Liu et al. 1986b

W. K. Liu, G. Besterfield et A. Mani. Probabilistic finite element methods in nonlinear structural dynamics. *Computer Methods in Applied Mechanics and Engineering*, 57 : 61–81, 1986.

Liu et al. 1988

W. K. Liu, G. Besterfield et T. Belytschko. Transient probabilistic systems. *Computer Methods in Applied Mechanics and Engineering*, 1988 : 27–54, 1988.

Loève 1977

M. Loève. *Probability Theory*. Springer Verlag, 1977. 4ème édition.

Massa et al. 2004

F. Massa, B. Lallemand, T. Tison et P. Level. Fuzzy eigensolutions of mechanical structures. *Engineering Computations*, 21(1) : 66–77, 2004.

MathWorks

MathWorks. Matlab 6.0 (2002). <http://www.mathworks.com>

McKay et al. 1979

M. D. McKay, R. J. Beckman et W. J. Conover. A comparison of three methods for selecting values of input variables in the analysis of output from a computer code. *Technometrics*, 21(2) : 239–245, 1979.

Menke 1984

W. Menke. *Geophysical data analysis : discrete inverse theory*. Academic Press, 1984.

Moens et Vandepitte 2002

D. Moens et D. Vandepitte. Fuzzy finite element method for frequency response function analysis of uncertain structures. *AIAA Journal*, 40 : 126–136, 2002.

Moine et al. 1997

P. Moine, L. Billet et D. Aubry. Two updating methods for dissipative models with non-symmetric matrices. Dans *Proceedings of the 15th International Modal Analysis Conference (IMAC-XV)*, pages 1931–1936, Orlando, Floride, février 1997.

Möller et al. 2000

B. Möller, W. Graf et M. Beer. Fuzzy structural analysis using α -level optimization. *Computing Mechanics*, 26(6) : 547–565, 2000.

Moore 1966

R. M. Moore. *Interval Analysis*. Prentice-Hall, 1966.

Moore 1979

R. M. Moore. *Methods and Applications of Interval Analysis*. Studies in Applied Mathematics (SIAM), 1979.

Mottershead et Friswell 1993

J. E. Mottershead et M. I. Friswell. Model updating in structural dynamics : a survey. *Journal of Sound and Vibration*, 167(2) : 347–375, 1993.

Muscolino et al. 2000

G. Muscolino, G. Ricciardi et N. Impollonia. Improved dynamic analysis of structures with mechanical uncertainties under deterministic input. *Probabilistic Engineering Mechanics*, 15 : 199–212, 2000.

Nelson 1976

R. Nelson. Simplified calculations of eigenvector derivatives. *AIAA Journal*, 14 : 1201–1205, 1976.

Ning et Kearfott 1997

S. Ning et R. B. Kearfott. A comparison of some methods for solving linear interval equations. *SIAM Journal of Numerical Analysis*, 34(4) : 1289–1305, 1997.

Oberkampf et al. 2001

W. L. Oberkampf, J. C. Helton et K. Sentz. Mathematical representation of uncertainty. Dans *Proceedings of the AIAA Non-Deterministic Approaches Forum*, Seattle, Washington, avril 2001.

Oettli et Prager 1964

W. Oettli et W. Prager. Compatibility of approximate solution of linear equations with given error bounds for coefficients and right-hand sides. *Mathematics*, 6 : 405–409, 1964.

Papadrakakis et Papadopoulos 1999

M. Papadrakakis et V. Papadopoulos. Parallel solution methods for stochastic finite element analysis using Monte Carlo simulation. *Computer Methods in Applied Mechanics and Engineering*, 168 : 305–320, 1999.

Pascual Gimenez et al. 1998

R. Pascual Gimenez, J.-C. Golinval et M. Razeto. On the reliability of error localization indicators. Dans *Proceedings of the 23rd International Conference on Noise and Vibration Engineering (ISMA 23)*, Leuven, Belgique, septembre 1998.

Pérec 1991

G. Pérec. *Cantatrix Sopranica*, chapitre Experimental demonstration of the tomatotopic organization of the soprano. Éditions du Seuil, 1991.

Puel et al. 2003

G. Puel, P. Ladevèze, A. Deraemaeker et D. Barthe. Sur une théorie des méconnaissances en calcul des structures. Dans *Actes du Sixième Colloque National en Calcul des Structures*, pages 399–406, Giens, Var, mai 2003.

Rao et Sawyer 1995

S. Rao et P. Sawyer. Fuzzy finite element approach for the analysis of imprecisely defined systems. *AIAA Journal*, 33(12) : 2364–2370, 1995.

Rubinstein 1981

R. Y. Rubinstein. *Simulation and the Monte Carlo Method*. John Wiley and Sons, 1981.

Rump 1983

S. M. Rump. *A New Approach to Scientific Computation*, chapitre Solving algebraic systems with high accuracy, pages 51–120. Academic Press, éditeurs : U. Kulisch, W. Miranker, 1983.

Saltelli et al. 2004

A. Saltelli, S. Tarantola, F. Campolongo et M. Ratto. *Sensitivity Analysis in Practice : A Guide to Assessing Scientific Models*. John Wiley and Sons, 2004.

Schuëller 2001

G. I. Schuëller. Computational stochastic mechanics - recent advances. *Computers and Structures*, 79 : 2225–2234, 2001.

Schuëller et al. 1997

G. I. Schuëller, L. A. Bergman, C. G. Bucher, G. Dasgupta, G. Deodatis, R. G. Ghanem, M. Grigoriu, M. Hoshiya, E. A. Johnson, A. Naess, H. J. Pradlwarter, M. Shinozuka,

K. Sobczyk, P. D. Spanos, B. F. Spencer, A. Sutoh, T. Takada, W. V. Wedig, S. F. Wojtkiewicz, I. Yoshida, B. A. Zeldin et R. Zhang. A state-of-the-art report on computational stochastic mechanics. *Probabilistic Engineering Mechanics*, 12(4) : 197–321, 1997.

Schultze et al. 2001

J. F. Schultze, F. M. Hemez, S. W. Doebling et Hoon Sohn. Statistical based non-linear model updating using feature extraction. Dans *Proceedings of the 19th International Modal Analysis Conference (IMAC-XIX)*, pages 18–26, Kissimmee, Floride, février 2001.

SDT

SDT. Structural Dynamics Toolbox 5.1 (for use with Matlab). E. Balmès et J. Leclère, 2003. <http://www.sdtools.com>

Shafer 1976

G. Shafer. *A Mathematical Theory of Evidence*. Princeton University Press, 1976.

Shannon 1948

C. E. Shannon. A mathematical theory of communication. *Bell System Technology Journal*, 27 : 379–423 ; 623–659, 1948.

Shannon et Weaver 1964

C. E. Shannon et W. Weaver. *Mathematical Theory of Communication*. University of Illinois Press, 1964.

Shaw et Jayasuria 1992

J. Shaw et S. Jayasuria. Modal sensitivities for repeated eigenvalues and eigenvalues derivatives. *AIAA Journal*, 30(3) : 850–852, 1992.

Shinozuka et Yamazaki 1988

M. Shinozuka et F. Yamazaki. *Stochastic Structural Dynamics : Progress in Theory and Applications*, chapitre Stochastic Finite Element Analysis : an Introduction. Elsevier Applied Sciences, éditeurs : S. T. Ariaratnam, G. I. Schuëller, I. Elishakoff, 1988.

Soize 1986

C. Soize. Modélisation probabiliste du flou structural en dynamique linéaire des systèmes mécaniques complexes - I. Éléments théoriques. *La Recherche Aérospatiale*, 5 : 337–362, 1986.

Soize 1994

C. Soize. Vibration damping in low-frequency range due to structural complexity - a model based on the theory of fuzzy structures and model parameters estimation. *Computers and Structures*, 58(5) : 901–915, 1994.

Soize 2000

C. Soize. A nonparametric model of random uncertainties for reduced matrix models in structural dynamics. *Probabilistic Engineering Mechanics*, 15(3) : 277–294, 2000.

Takada 1990a

T. Takada. Weighted integral method in stochastic finite element analysis. *Probabilistic Engineering Mechanics*, 5(3) : 146–156, 1990.

Takada 1990b

T. Takada. Weighted integral method in multi-dimensional stochastic finite element analysis. *Probabilistic Engineering Mechanics*, 5(4) : 158–166, 1990.

Takada 1992

T. Takada. Variability response functions and stochastic field discretization in stochastic finite element methods. Dans *Proceedings of the 6th Specialty Conference on Probabilistic Mechanics and Structural and Geotechnical Reliability*, Denver, Connecticut, juillet 1992.

Tarantola 1987

A. Tarantola. *Inverse Problem Theory*. Elsevier, 1987.

Thompson 1992

S. K. Thompson. *Sampling*. Wiley Interscience, 1992.

Van den Nieuwenhof 2003

B. Van den Nieuwenhof. *Stochastic Finite Elements for Elastodynamics : Random Field and Shape Uncertainty Modelling Using Direct and Modal Perturbation-Based Approaches*. Thèse de doctorat, Université catholique de Louvain - Faculté des sciences appliquées - Unité de Génie Civil et Environnemental, 2003.

Van den Nieuwenhof et Coyette 2003

B. Van den Nieuwenhof et J.-P. Coyette. Modal approaches for the stochastic finite element analysis of structures with material and geometric uncertainties. *Computer Methods in Applied Mechanics and Engineering*, 192 : 3705–3729, 2003.

Vanmarcke et Grigoriu 1983

E. H. Vanmarcke et M. Grigoriu. Stochastic finite element analysis of simple beams. *Journal of Engineering Mechanics*, 109(5) : 1203–1214, 1983.

Veneziano et al. 1983

D. Veneziano, F. Casciati et L. Faravelli. Method of seismic fragility for complicated systems. Dans *Proceedings of the 2nd Specialistic Meeting on Probabilistic Methods in Seismic Risk Assessment for NPP*, Livermore, Californie, 1983.

Vinot et Cogan 2004

P. Vinot et S. Cogan. A robust model-based test planning procedure. Dans *Proceedings of the 22nd International Modal Analysis Conference (IMAC-XXII)*, Dearborn, Michigan, janvier 2004.

Vinot et al. 2002

P. Vinot, S. Cogan et Y. Ben-Haim. Reliability of structural dynamic models based on info-gap models. Dans *Proceedings of the 20th International Modal Analysis Conference (IMAC-XX)*, pages 1057–1063, Los Angeles, Californie, février 2002.

Walley 1990

P. Walley. *Statistical Reasoning with Imprecise Probabilities*. Chapman and Hall, 1990.

Wiener 1938

N. Wiener. The homogeneous chaos. *American Journal of Mathematics*, 60 : 897–936, 1938.

Williamson et Downs 1990

R. C. Williamson et T. Downs. Probabilistic arithmetic I : numerical methods for calculating convolutions and dependency bounds. *International Journal of Approximate Reasoning*, 4 : 89–158, 1990.

Worden et al. 2003

K. Worden, G. Manson, T. M. Lord et M. I. Friswell. Some thoughts on uncertainty quantification and propagation. Dans *Proceedings of the 21st International Modal Analysis Conference (IMAC-XXI)*, Kissimmee, Floride, février 2003.

Zadeh 1965

L. A. Zadeh. Fuzzy sets. *Information and Control*, 8 : 338–353, 1965.

Zhang et Ellingwood 1994

J. Zhang et B. Ellingwood. Orthogonal series expansion of random fields in reliability analysis. *Journal of Engineering Mechanics*, 120(12) : 2660–2677, 1994.

Zhang et Wei 1995

D. Zhang et F. Wei. Computations of eigenvector derivatives with repeated eigenvalues using a complete modal space. *AIAA Journal*, 33(9) : 1749–1753, 1995.

Zienkiewicz et Zhu 1987

O. C. Zienkiewicz et J. Z. Zhu. A simple error estimator and adaptative procedure for practical engineering analysis. *International Journal for Numerical Methods in Engineering*, 24 : 337–357, 1987.

Résumé

Dans le cadre de la quantification de la qualité d'un modèle par rapport à une référence expérimentale, nous présentons une nouvelle théorie basée sur le concept de méconnaissance qui permet de manipuler un modèle qui est forcément une représentation imparfaite de la réalité, et d'obtenir des résultats prédictifs qui tiennent compte du manque de connaissance que l'on a vis-à-vis de la structure réelle étudiée. Cette théorie, qui mêle théorie des intervalles et théorie des probabilités, repose sur la globalisation des diverses sources d'erreur à l'échelle des sous-structures en utilisant une variable interne scalaire contenue dans un intervalle dont les bornes suivent des lois de probabilités.

À partir de ces bornes, si l'on considère une quantité d'intérêt définie sur l'ensemble de la structure (une fréquence propre par exemple), il est possible à partir de la donnée du modèle de méconnaissances de base de calculer un intervalle d'appartenance de cette quantité d'intérêt, dont les bornes sont probabilistes. On définit alors des intervalles de méconnaissances dites effectives, qui ont une certaine probabilité de contenir la quantité d'intérêt associée au modèle.

La quantification de la qualité d'un modèle par rapport à une référence expérimentale se traduit alors par la détermination des méconnaissances de base qui sont les plus représentatives des dispersions constatées. L'idée majeure de cette détermination est de considérer que plus on a d'informations expérimentales, plus on est susceptible de réduire les méconnaissances de base. La démarche dite de réduction des méconnaissances est appliquée avec succès sur des cas d'études académiques ainsi que sur un exemple réel de structure industrielle : le support de satellites SYLDA5 présent dans le lanceur Ariane 5.

Mots clés : méconnaissances, incertitudes, validation de modèles, dynamique des structures

Abstract

Concerning the quantification of the quality of a model in comparison with an experimental reference, we present a new theory based on the concept of Lack Of Knowledge (LOK) combining interval analysis with probabilistic features : the basic principle consists in globalizing the various sources of errors on the substructure level using a scalar internal variable, called the LOK variable, defined over an interval whose upper and lower bounds follow probabilistic laws.

From these basic LOKs, we can derive for the whole structure the effective LOK for a quantity of interest (e.g. an eigenpulsation), resulting in an interval with stochastic bounds that we can compare with experimental values.

The quantification of the quality of a model results in the determination of the basic LOKs the most representative of the experimental data. The process uses additional experimental information to reduce the LOKs and is successfully applied to academic as well as industrial cases.

Keywords : lacks of knowledge, uncertainties, model validation, structural dynamics