



Réduction de modèles thermiques par amalgame modal

Abdelhakim Oulefki

► To cite this version:

Abdelhakim Oulefki. Réduction de modèles thermiques par amalgame modal. Autre. Ecole Nationale des Ponts et Chaussées, 1993. Français. NNT: . tel-00523620

HAL Id: tel-00523620

<https://pastel.hal.science/tel-00523620>

Submitted on 5 Oct 2010

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.

79969

NS 16745

α (4)

THESE DE DOCTORAT

Spécialité: Sciences et Techniques du Bâtiment

Présentée par: Abdelhakim OULEFKI

**Réduction de modèles thermiques par
amalgame modal**

soutenue le 09 Février 1993 devant le jury composé de:

MM:	J. BRANSIER	Rapporteur
	R. DEHAUSSE	Examineur
	P. DUHAMEL	Examineur
	A. NEVEU	Examineur
	D. PETIT	Rapporteur
	J. RILLING	Examineur



03 02

à tous les miens
à Danielle Lega (minou)

Résumé

On présente la méthode d'amalgame modal. Il s'agit d'une approche pour réduire un modèle d'état modal quelconque. La méthode est ici appliquée dans le cadre de la thermique. Le principe repose sur une partition judicieuse de l'espace d'état modal en quelques sous espaces disjoints. La dynamique de chaque sous espace est ensuite approchée au mieux par un pseudo-éléments propre. L'optimalité de la démarche est prouvée au sens d'un critère d'écart quadratique de qualité. Le modèle obtenu conserve des liens formels avec le modèle d'origine. Du point de vue algorithmique, la méthode est automatique: on peut chercher le meilleur modèle réduit respectant une contrainte de précision et/ou de taille. La méthode est performante en temps de calcul. La réduction par amalgame modal est comparée à celles d'autres méthodes. Des exemples de réduction de modèles modaux 1D, 2D et 3D sont donnés.

Abstract

A complete method (the amalgam method) for reducing state models of linear thermal systems without mass transfer is presented here. We apply the approach to thermal problems. The idea set to work is based on partitioning of the model state space into some orthogonal subspaces. The dynamic behavior of each subspace is then approximated by only one pseudo-eigen function. The reduced model obtained with aforementioned method is shown to be optimal with respect to significant error measure. The algorithm runtime is short. Several examples are presented, for which the amalgam method is compared with other approaches.

Remerciements

Le travail de recherche rédigé dans cette thèse a été effectué au sein du Groupe Informatique et Systèmes Energétiques (GISE)¹. Je tiens à souligner le plaisir d'avoir préparé ma thèse au sein de cette équipe dynamique.

Monsieur G. LEFEBVRE, responsable du GISE, m'a accueilli et conseillé tout au long de ces trois dernières années. Il a énormément œuvré pour que je dispose de bonnes conditions (financières et autres) nécessaires au bon déroulement de cette thèse. Je voudrais lui témoigner ici de ma profonde reconnaissance.

Je rends hommage à Monsieur A. NEVEU, mon directeur de thèse, Maître de conférences de l'Ecole des Mines de Paris, pour m'avoir sérieusement suivi dans mon travail. Son expérience en analyse modale mais aussi sa rigueur scientifique m'ont été très utiles à maintes reprises.

Je remercie sincèrement Monsieur J. RILLING, responsable scientifique à l'ENPC et Directeur de recherche au CSTB. Je le remercie pour avoir accepté de participer au jury de ma thèse, mais aussi pour avoir soutenu ma demande d'obtention de bourse auprès de l'ENPC.

Monsieur D. PETIT, Docteur d'état et Maître de conférences de l'université de Provence (Aix-Marseille I) doit être ici gratifié. Outre son acceptation d'être rapporteur de ma thèse, il a manifesté une grande disponibilité lorsque j'ai eu besoin d'appliquer ses résultats de recherche.

Je suis très sensible à l'honneur que me font

Monsieur R. DEHAUSSE professeur de l'Ecole des Mines de Paris

Monsieur P. DUHAMEL professeur de l'IUT Saint Denis

de participer au jury de thèse.

Je n'oublie pas Monsieur J. BRANSIER, Maître de conférences de l'université Pierre & Marie Curie, pour plusieurs raisons. D'abord parce que c'est grâce à lui que j'ai saisi l'intérêt des méthodes numériques dans la modélisation en thermique (cours et stage de DEA). Ensuite pour m'avoir mis en contact avec le GISE pour effectuer cette thèse. Aujourd'hui, il me fait honneur d'être rapporteur de ma thèse et je tiens à lui renouveler toute ma reconnaissance.

Mes remerciements vont aussi

- aux membres du GISE, notamment R. EBERT, B. FLAMENT, K. EL KHOURY et C. DEVILLE CAVELLIN pour leur aide amicale.

- à l'Agence de l'Environnement et de la Maîtrise de l'Energie (ADEME) qui a soutenu financièrement le projet SYMBOL dans lequel s'intègre les travaux de recherche de ma thèse.

¹(Le GISE) est un groupe de recherche commun à l'Ecole des Mines de Paris et l'Ecole Nationale des Ponts et Chaussées.

Table des Matières

1	Le formalisme d'état modal	9
1.1	Présentation générale	9
1.2	Equation de l'énergie	11
1.3	Formulation modale d'un problème thermique	12
1.3.1	Formulation modale continue	12
1.3.2	Formulation modale par discrétisation spatiale	16
1.3.3	Formulation modale pour les régimes forcés	18
2	Réduction de modèles: les différentes approches	21
2.1	Preliminaire	21
2.2	Les méthodes de troncature	23
2.2.1	Présentation	23
2.2.2	Critère de Marshall	24
2.2.3	Critère de Litz	26
2.2.4	Critère de dominance	28
2.2.5	Critère de dominance pondérée	29
2.2.6	Récapitulatif sur la troncature modale	30
2.3	Les méthodes de minimisation	33
2.3.1	Présentation	33

2.3.2	Approche d'agrégation	35
2.3.3	Approche d'Eitelberg	37
2.3.4	Approche mixte réduction-identification	38
2.3.5	Récapitulatif sur les méthodes de minimisation	39
3	Réduction d'un modèle d'état modal par amalgame modal	41
3.1	Preliminaire	41
3.2	Mesure de l'erreur de réduction	44
3.3	Description de la méthode	45
3.3.1	Présentation générale	45
3.3.2	Notations	46
3.3.3	Détermination des modes amalgamés	47
3.3.4	Partition de l'espace des modes propres	51
3.4	Ajustement temporel du modèle réduit	59
3.5	Automatisation de la réduction	60
3.5.1	Normalisation de l'erreur	60
3.5.2	Algorithme d'amalgame	65
3.6	Structure modale du modèle réduit	67
3.7	Estimation du temps de calcul	71
	Exemples de réduction de modèles thermiques	73
4	Paroi tricouche	75
4.1	Description de la paroi	75
4.2	Réduction modale	77
4.2.1	Les sous-espaces d'amalgame	77

4.2.2	Les modes amalgamés	78
4.2.3	Simulation de l'évolution thermique	83
4.3	Erreurs de réduction	85
4.4	Influence du domaine temporel $\mathcal{D}_t$	86
5	Bâtiment bizone	89
5.1	Description du bâtiment	89
5.2	Etude préliminaire	91
5.3	Amalgame modal	97
5.3.1	Les sous-espaces d'amalgame	97
5.3.2	Les modes amalgamés	97
5.3.3	Simulation de l'évolution thermique	98
5.3.4	Réduction par l'agrégation	103
5.4	Erreurs de réduction	104
5.5	Rôle de la pondération sur les sollicitations	106
6	Pont thermique	115
6.1	Modèle de référence	115
6.2	Simulation avec le modèle de référence	118
6.3	Amalgame modal	122
6.4	Méthode mixte réduction-identification	133
7	Vilebrequin	143
7.1	Modèle de référence	143
7.2	Amalgame modal	145
7.3	Simulation de l'évolution thermique	153

8 Conclusion & perspectives	157
Références bibliographiques	163
Annexes	169
A Etude énergétique de la troncature	171
A.1 Dominances des modes propres	171
A.1.1 Critère de dominance	173
A.1.2 Critère de dominance pondérée	173
A.1.3 Erreur de troncature	174
A.1.4 Troncature automatique	175
B Equation de Lyapunov	177
C Formulation modale pour les régimes forcés	179
D Rappels sur la dérivation tensorielle	183
E Mesure de l'erreur	187
F généralités sur la méthode d'amalgame	191
F.1 Détermination des modes amalgamés	191
F.2 Contribution à $\mathcal{M}$ des modes propres	193
F.3 Expression de $\mathcal{M}$ pour les différents modèles	194
F.4 Ajustement temporel du modèle amalgamé	195

Nomenclature

Symboles latins

B ($\tilde{B}$)	Matrice de commande du modèle détaillé (réduit).
$\mathcal{B}$	Opérateur des conditions aux limites.
$C(M)$	Capacité calorifique volumique au point M .
$\mathcal{D}$	Domaine spatial du système thermique.
$\mathcal{D}_t$	Un domaine temporel dans $[0, \infty[$.
d_i	Dominance modale du $i^{\text{ème}}$ élément propre.
$e(M, t)$	Erreur de réduction au point M et à l'instant t .
F ($\tilde{F}$)	Matrice diagonale de valeurs propres du modèle détaillé (réduit) .
h	Coefficient d'échange.
H ($\tilde{H}$)	Matrice des sorties du modèle détaillé (réduit).
H_k	$k^{\text{ème}}$ sous-espace d'amalgame.
$\mathcal{L}$	Opérateur de la chaleur.
M	Variable spatiale du domaine $\mathcal{D}$.
$\mathcal{M}$	Mesure de l'erreur de réduction.
$\overline{\mathcal{M}}$	Mesure normée de l'erreur (ou erreur de réduction).
$\mathcal{M}_k$	Contribution à $\mathcal{M}$ des modes propres du sous-espace H_k .
$\mathcal{M}_{km}$	Contribution à $\mathcal{M}$ du mode propre m de H_k .
N	Ordre du Modèle détaillé à réduire.
n	Ordre du Modèle réduit (égal au nombre de sous-espaces d'amalgames).
n_k	Dimension du sous-espace H_k ou nombre de ses modes propres.
p	Nombre de sollicitations appliquées au système thermique.
$\mathbf{P}$	Tableau des modes du modèle détaillé.
$\mathbf{P}_k$	Tableau des modes du sous-espace d'amalgane H_k .
$\tilde{\mathbf{P}}$	Tableau des pseudo-modes du modèle réduit.
S ($\tilde{S}$)	Matrice statique du modèle détaillé (réduit) .
t	Variable temporelle du domaine $\mathcal{D}_t$.
$T(M, t)$	Température donnée par le modèle détaillé.
$T_d(M, t)$ ($T_g(M, t)$)	Partie dynamique (glissante) de $T(m, t)$.
$\tilde{T}(M, t)$	Température donnée par le modèle réduit.
$U(t)$	Vecteur $[p]$ des sollicitations.
$u^k(i)$	Fonction de projection donnant l'indice dans $\mathbf{P}$ du $i^{\text{ème}}$ mode propre de H_k .
$V_i(M)$	$i^{\text{ème}}$ Fonction propre (ou mode propre).
$\tilde{V}_i(M)$	Mode amalgamé issu du sous-espace H_i .
$x(t)$	Vecteur $[N]$ d'état du modèle détaillé.
$\tilde{x}(t)$	Vecteur $[n]$ d'état du modèle réduit.
y ($\tilde{y}$)	vecteur des sorties $[q]$ du modèle détaillé (réduit).
$\tilde{u}(M)$	produit du vecteur vitesse $\tilde{v}(M)$ par $C(M)$.

Symboles grecs

δ_{ij}	Symbole de Kronecker.
ω_k	Vecteur $[n_k]$ des coefficients de décomposition de $\tilde{V}_k(M)$ sur la base modale de H_k .
π_k	Matrice $[N, n_k]$ de projection de $\mathbf{P}_k$.
ξ_k	Vecteur $[N]$ de projection du mode principal de H_k .
τ_i	Constante de temps associée au $i^{\text{ème}}$ mode propre.
λ_i	valeur propre associée au $i^{\text{ème}}$ mode propre.

INTRODUCTION

La thermique est sans doute l'un des premiers domaines sur lesquels l'homme a commencé à réfléchir. Il a d'abord "trouvé" le moyen de se protéger du froid et celui de faire le feu bien avant d'avoir inventé la roue. Il est vrai que l'objectif était de survivre dans une nature rude. Mais aujourd'hui encore, avec essentiellement l'enjeu économique, les arguments en faveur de la recherche en thermique sont nombreux et témoignent tous de la nécessité de mieux comprendre le comportement dynamique des systèmes thermiques afin de gérer rationnellement l'énergie.

Par ailleurs, le nombre de travaux en modélisation des phénomènes thermiques en régime transitoire est impressionnant. Cela indique aussi bien l'intérêt que la difficulté à mener ce type de recherches.

En réalité à l'heure actuelle, le thermicien dispose du savoir théorique nécessaire à la formulation mathématique rigoureuse de phénomènes thermiques très complexes. En revanche, la résolution formelle du problème mathématique n'est malheureusement accessible que dans des cas académiques ou en adoptant des hypothèses simplificatrices de manière à éliminer les obstacles aujourd'hui insurmontables.

Heureusement que beaucoup de systèmes thermiques peuvent être raisonnablement représentés moyennant des hypothèses telles que la linéarité des échanges thermiques et la constance des paramètres thermophysiques. C'est notamment le cas des systèmes qui ne sont pas le siège de forts gradients thermiques. Ce type de systèmes bénéficie déjà d'un large éventail de méthodes mathématiques.

La discrétisation des relations formelles d'un problème thermique aboutit à des équations algébrique. Naturellement, le nombre d'équations et de paramètres croît vite avec la taille spatiale du système, la complexité des phénomènes thermiques et le degré de finesse souhaité en modélisation. La prise en compte du régime dynamique est particulièrement pénalisante en temps calcul sans compter l'effort initial nécessaire à la mise au point du modèle mathématique.

On peut rajouter à cela les difficultés de nature numérique. En effet, l'étude des grands systèmes thermiques se fait principalement par des algorithmes mathématiques (Runge-Kutta, Adams, Gauss ...) ou de programmation non-linéaire (méthode des gradients, algorithme de Marquardt ...). En pratique, le temps de calcul augmente en n^3 (n désignant la dimension du modèle) et la taille mémoire nécessaire à la machine en n^2 . De plus, à cause des erreurs d'arrondi de la machine, la convergence des algorithmes devient de plus en plus incertaine lorsque la taille du modèle augmente.

Malgré l'évolution permanente de l'électronique et de l'informatique, les remarques précédentes justifient l'intérêt que portent aujourd'hui les thermiciens à la réduction des modèles thermiques. Dans les domaines du contrôle et de la commande, cet intérêt est majeur. Aussi, on assiste depuis environ deux décennies à une éclosion de méthodes de simplification de modèles thermiques. La plupart de ces méthodes peuvent se ranger dans deux grandes

classes:

a) Il y a d'abord les méthodes d'identification qui valident une forme de modèle initialement choisie de petite taille, à partir de mesures souvent réelles. Ces modèles se limitent généralement à un petit nombre d'entrées-sorties du système thermique. Il n'est pas toujours évident de donner un sens physique aux paramètres identifiés, et il ne sont valables que pour les entrées-sorties pour lesquelles ils sont calculés. Cette technique nécessite peu de connaissances précises sur le système.

b) Il y a ensuite les méthodes qui élaborent un modèle réduit à partir d'un modèle de connaissance détaillé. Selon la technique de réduction employée, il est possible de donner une interprétation aux différents paramètres du modèle réduit. Dans certains cas, on peut même donner une analyse riche et pertinente sur la qualité de la réduction à partir du lien mathématique entre le modèle de départ et celui obtenu en fin de la démarche.

En dehors de ces deux grandes classes de méthodes, il existe aussi certaines techniques de simplification de modèles qui restent encore au stade de développement. On peut mentionner

La technique d'homogénéisation dont l'objectif est de représenter un système hétérogène grâce à des paramètres physiques *équivalents*. Le système est ensuite caractérisé par un modèle de faible taille élaboré en considérant que le système est *homogène* avec les paramètres physiques équivalents. Cette technique comporte une phase expérimentale délicate.

La technique de maillage dynamique qui consiste à restreindre la discrétisation fine uniquement aux zones de forts gradients. Le maillage évolue alors dans le temps selon la localisation des zones à forts gradients.

C'est dans le cadre de la catégorie **b)** des méthodes de réduction que s'inscrit le travail présenté dans cette thèse. Plus exactement, nous allons aborder la réduction de modèles de connaissance écrits dans le formalisme d'état modal. Ce formalisme est en effet largement utilisé dans de nombreux domaines, notamment ceux de la commande automatique, du traitement de l'information et du comportement des organismes vivants. De plus, ce formalisme bénéficie d'un large éventail de méthodes développées dans le cadre de la théorie des systèmes et de l'optimisation. La forme canonique des modèles d'état modaux est très adaptée à l'étude et analyse du comportement thermique d'un système linéaire, conduisant ainsi à une démarche de réduction mieux justifiée.

Le plan de ce rapport fait apparaître globalement quatre parties. Le lecteur non habitué au formalisme d'état de façon générale et à l'analyse modale en particulier peut trouver dans la partie I les éléments nécessaires pour une lecture facile du travail présenté. Il ne s'agit pas d'une présentation détaillée du formalisme d'état modal, mais d'un résumé utile. Les différentes approches couramment utilisées pour la réduction de modèles thermiques

sont présentées dans la partie II. A travers cette partie commentée, nous verrons mieux la nécessité du travail entrepris dans cette thèse. Nous présentons ensuite *la méthode de réduction par amalgame modal* dans la partie III. Enfin dans la partie IV, des exemples de réduction de modèles de divers systèmes thermiques (1D, 2D et 3D) permettront de montrer l'intérêt du travail accompli dans cette thèse.

Chapitre 1

Le formalisme d'état modal

1.1 Présentation générale

Le formalisme d'état modal que nous utilisons de manière presque exclusive dans cette thèse va être résumé dans ce chapitre. Il s'agit surtout d'éclairer le lecteur pour lequel la méthode d'analyse modale n'est pas familière. En particulier, il n'est pas question de redémontrer les propriétés fondamentales sur lesquelles repose ce formalisme. Le lecteur soucieux de comprendre les preuves mathématiques des résultats peut se référer à plusieurs travaux antérieurs [60][58][41][39][7][8]. Quant au lecteur qui connaît déjà l'analyse modale, il peut passer directement au chapitre suivant où sont résumés les principales méthodes de réduction de modèles d'état.

Notre travail dans cette thèse repose sur l'hypothèse de posséder un modèle d'état modal suffisamment détaillé. Par *suffisamment détaillé*, nous entendons une *bonne* représentativité du comportement thermique du système par le modèle modal. Cette hypothèse permet de nous affranchir de la manière dont a été construit initialement le modèle, notamment la méthode de discrétisation qui est le plus souvent nécessaire. Néanmoins, les systèmes thermiques auxquels s'appliquent nos résultats doivent vérifier trois conditions: la linéarité, l'invariance (ou la stationnarité *structurelle*) et la réciprocité. Aussi, nous donnons la signification de ces conditions:

Linéarité

Tout échange d'énergie à l'intérieur d'un domaine spatial $\mathcal{D}$ ou avec son milieu environnant, est supposé linéaire ou linéarisé. C'est une hypothèse raisonnable pour de nombreux systèmes qui ne sont pas le siège de gradients thermiques forts. Le bâtiment constitue un bon exemple où cette approximation est admise et donne des

résultats satisfaisants. Elle est vraie pour la conduction et est admise pour la convection et le rayonnement. En l'état actuel, cette hypothèse reste nécessaire pour l'exploitation du principe de superposition [30] utilisé également dans le formalisme d'état modal.

Invariance

Les paramètres thermophysiques intervenant dans les équations mathématiques sont supposés être indépendants du temps et donc a fortiori de la température. Cette restriction est largement acceptée. Un système thermique invariant tend vers un régime asymptotique lorsqu'il est soumis à des excitations constantes dans le temps. Du point de vue formel, l'invariance permet d'obtenir un problème aux valeurs propres [18] dont la résolution ne dépend que de l'espace.

Réciprocité

On dit que les échanges thermiques sont réciproques [39] si à tout instant, l'inversion des températures entre deux points quelconques du système entraîne celle du flux transféré entre ces deux points, le flux *gardant la même valeur absolue*. On vérifie facilement que les transferts thermiques par conduction, convection (linéarisée au moyen d'un coefficient d'échange) et rayonnement sont réciproques. Il n'en est pas ainsi pour le transport d'énergie par transfert de masse (enthalpie). C'est par exemple le cas des systèmes à caloporteurs où les forces de pression sont la cause du transport d'énergie. La réciprocité exige donc que les échanges d'énergie résultent d'un écart de température. Elle entraîne la symétrie des équations d'échange, conduisant à un problème aux valeurs propres [2][39] à solutions réelles.

En vérité, le formalisme d'état modal a été généralisé par El Khoury [39] aux systèmes thermiques non-réciproques. Les solutions du problème aux valeurs propres obtenues sont à valeurs complexes. Bien que les systèmes thermiques non-réciproques ne fassent pas vraiment partie du travail rédigé dans ce rapport, nous les évoquerons parmi nos perspectives visant à généraliser ultérieurement notre travail.

On peut également citer quelques démarches montrant la possibilité d'étendre le formalisme d'état modal pour résoudre des problèmes thermiques non-linéaires. L'idée admise repose sur le fait que les paramètres thermophysiques peuvent être considérés constants par pas de temps. On obtient alors un modèle d'état modal pour chaque pas de temps. Cependant, on montre [43] que dans le cas de faibles non-linéarités, les modèles modaux associés aux différents pas de temps ne sont pas très différents. Ceci permet de réduire le nombre de modèles modaux en considérant des pas de temps assez grands. Neveu [50] aboutit à un résultat similaire en montrant formellement qu'une faible perturbation (par variation des paramètres physiques) appliquée à un modèle modal, ne modifie pas notablement le caractère découplé de ce dernier. Signalons aussi une autre démarche sur le domaine du non linéaire qui s'inspire des techniques dites d'*approche systémique* [32][33]. Elle consiste à subdiviser le système thermique en plusieurs composants pouvant être individuellement décrits par des modèles linéaires. Les composants sont ensuite assemblés par des équations de couplage où sont introduites les non-linéarités.

1.2 Equation de l'énergie

Considérons un système thermique $\mathcal{D}$ dont l'extension spatiale est finie. Ce système peut être composé de sous-domaines solides ou fluides. On notera par $\partial\mathcal{D}$ l'ensemble des points situés aux divers interfaces externes entre $\mathcal{D}$ et son milieu environnant. L'équation de l'énergie qui découle de l'application du premier principe de la thermodynamique [27], et qui régit l'évolution thermique de $\mathcal{D}$ se met [39] sous la forme synthétique suivante

$$\begin{aligned} \forall M \in \mathcal{D}^0 \quad C(M) \frac{\partial T(M, t)}{\partial t} &= \mathcal{L}[T(M, t)] + \varphi(M, t) \\ \forall M \in \partial\mathcal{D} \quad \mathcal{B}[T(M, t)] + \varphi(M, t) &= 0 \end{aligned} \quad (1.1)$$

Dans (1.1), $\mathcal{L}$ et $\mathcal{B}$ sont respectivement [39] l'opérateur spatial généralisé de la chaleur et l'opérateur des conditions aux limites associé à $\mathcal{L}$. Dans le cas linéaire et instationnaire (régime variable), leurs formes explicites sont données par

$$\mathcal{L}(T(M, t)) = \vec{\nabla} \cdot [K(M) \vec{\nabla} T(M, t)] - \vec{u}(M) \cdot \vec{\nabla} T(M, t) + \int_{\mathcal{D}} r(M, M') T(M', t) dM' \quad (1.2)$$

$$\mathcal{B}(T(M, t)) = - [K(M) \vec{\nabla} T(M, t)] \cdot \vec{n}(M) - h(M) T(M, t) \quad (1.3)$$

Dans (1.3), les cas particuliers des conditions aux limites de Dirichlet (Température imposée) et de Neumann (flux imposé) s'obtiennent en faisant tendre $h(M)$ respectivement vers l'infini et zéro. $\varphi(M, t)$ désigne l'ensemble des puissances volumiques imposées au point M et qui interviennent comme des sollicitations en ce point (flux radiatif provenant de l'extérieur, sources, etc).

On peut reconnaître facilement que les trois termes de (1.2) représentent successivement les trois modes de transfert: conduction, transport et rayonnement. $K(M)$ est le tenseur des conductivités thermiques qui devient sphérique [30] dans le cas isotrope. $\vec{u}(M)$ est le produit du vecteur vitesse $\vec{v}(M)$ par la capacité calorifique volumique $C(M)$. Enfin $r(M, M')$ du terme intégral correspond à une fonction de distribution d'échanges à distance (rayonnement). On montre que $r(M, M') = r(M', M)$ ce qui confère le caractère réciproque au terme radiatif.

La présence du terme de transport dans (1.2) rend l'opérateur $\mathcal{L}$ non-autoadjoint [2] [39]. Il possède des propriétés mathématiques remarquables résumées par El Khoury [39]. Le spectre de $\mathcal{L}$ est infini mais dénombrable et appartient au corps des complexes. C'est la raison pour laquelle nous supposons que ce terme est nul et dans ce cas, le spectre de $\mathcal{L}$ devient réel.

Le système (1.1) suffit à caractériser complètement le comportement thermique du système. Il peut être résolu par la voie numérique moyennant une double discrétisation dans l'espace

et dans le temps. A ce sujet, on peut mettre en œuvre diverses méthodes classiques: éléments finis, différences finies, méthode nodale, etc. Si cette discrétisation est fine, la résolution de (1.1) devient pénalisante en temps de calcul. Même si la finalité est d'obtenir l'évolution de la température en un nombre limité de points, cette contrainte ne peut être évitée.

D'autre part, les modèles discrétisés à la fois dans le temps et dans l'espace ne sont pas adaptés à la tâche de réduction. En effet, l'approche utilisée pour exploiter ce type de modèles repose, le plus souvent, sur l'estimation à divers pas de temps, de l'ensemble des températures associées aux nœuds de la discrétisation. Pour diminuer le coût en temps de calcul, on a généralement recours à deux techniques: la première consiste à "grossir" le maillage dans certaines zones du système thermique. La technique dite de condensation de Guyan [37] en est un exemple. La seconde technique consiste à "dilater" le pas de temps lorsque cela est possible.

1.3 Formulation modale d'un problème thermique

1.3.1 Formulation modale continue

Nous allons voir maintenant les étapes qui permettent de transcrire les équations (1.1) dans le formalisme d'état modal. On verra alors apparaître des possibilités de réduction remarquables.

En règle générale, le modélisateur cherche à connaître l'évolution de quelques grandeurs données, intensives ou extensives, en un point, sur une surface ou dans un volume limité du système thermique. On désignera ces grandeurs par le vocable de "sorties" et on les notera y_i $i = 1, 2, \dots$. l'expression de chaque sortie est donnée [39] par

$$y_i(t) = j_i [T(M, t)] + g_i [\varphi(M, t)] \quad (1.4)$$

j_i et g_i sont deux fonctionnelles spatiales à coefficients réels. La seconde fonctionnelle décrit la dépendance directe de la sortie vis-à-vis des sollicitations agissant sur le système thermique. Notons que si la sortie y_i est la température au point M , alors $j_i = 1$ (identité) et $g_i = 0$.

Dans le formalisme d'état modal, la dynamique d'un système thermique s'exprime comme une combinaison linéaire de fonctions spatiales $V_i(M)$ appelées "fonctions propres" ou encore "modes propres". Les coefficients de décomposition x_i , appelés "états des modes" ou "amplitudes des modes", dépendent des valeurs propres λ_i associées aux modes propres $V_i(M)$. Un couple (mode propre, valeur propre associée) est appelé "élément propre". Leur nombre est théoriquement infini mais dénombrable. Mais il existe plusieurs travaux qui montrent que la dynamique d'un système thermique est globalement véhiculée par

seulement quelques éléments propres. Nous reviendrons plus en détail sur ce point essentiel pour notre travail au §2.

Les éléments propres sont solution du problème aux valeurs propres suivant

$$\begin{aligned} \forall M \in \mathcal{D}^0 \quad \mathcal{L}[V_i(M)] &= \lambda_i C(M) V_i(M) \\ \forall M \in \partial \mathcal{D} \quad B[V_i(M)] &= 0 \end{aligned} \quad (1.5)$$

Parmi les propriétés importantes des éléments propres, que nous exploiterons d'ailleurs plusieurs fois dans cette thèse, il y a la relation d'orthonormalité que vérifient les fonctions propres

$$\begin{aligned} \forall i = 1, 2, \dots \quad \forall j = 1, 2, \dots \\ \langle V_i(M), C(M) V_j(M) \rangle = \int_{\mathcal{D}} V_i(M) C(M) V_j(M) dM = \delta_{ij} \end{aligned} \quad (1.6)$$

Il existe aussi une relation d'ordre basée sur le classement des constantes de temps τ_i

$$\tau_1 > \tau_2 > \tau_3 > \dots \quad (1.7)$$

$$\text{avec } \tau_i = -\frac{1}{\lambda_i} \quad (1.8)$$

Le champs de température $T(M, t)$ ne peut se décomposer complètement sur la base des fonctions propres car contrairement à celles-ci, il vérifie des conditions aux limites non homogènes comme le montre (1.1). Il est donc courant de décomposer $T(M, t)$ en deux termes: le terme "dynamique" $T_d(M, t)$ s'écrivant comme une combinaison linéaire des fonctions propres et le terme "glissant"¹ $T_g(M, t)$ ne pouvant généralement pas l'être car non homogène aux limites. On écrit donc

$$\forall M \in \mathcal{D}, \quad \forall t \geq 0 \quad T(M, t) = T_d(M, t) + T_g(M, t) \quad (1.9)$$

La partie dynamique complètement décomposable² sur la base des fonctions propres est donnée par

$$T_d(M, t) = \sum_{i=1}^{\infty} x_i(t) V_i(M) \quad (1.10)$$

¹La température glissante $T_g(M, t)$ est équivalente au régime permanent qu'atteindrait le point M si U gardait une valeur constante égale à celle de l'instant t . C'est aussi la température $T(M, t)$ si le système thermique avait un temps de réponse nul (ou capacité calorifique nulle).

²Les coefficients de décomposition $x_i(t)$ sont les états des modes. On les désigne [43] aussi par "variables douées de mémoire" en raison de l'information qu'elles contiennent sur le passé du système

et vérifie le problème suivant

$$\left| \begin{array}{l} \forall M \in \mathcal{D}^0 \quad C(M) \frac{\partial T_d(M, t)}{\partial t} = \mathcal{L}[T_d(M, t)] - C(M) \frac{\partial T_g(M, t)}{\partial t} \\ \forall M \in \partial \mathcal{D} \quad \mathcal{B}[T_d(M, t)] = 0 \end{array} \right. \quad (1.11)$$

alors que la partie glissante s'obtient à partir de l'expression

$$T_g(M, t) = \sum_{l=1}^p S_l(M) U_l(t) \quad (1.12)$$

Où $S_l(M)$ est le champ statique dans $\mathcal{D}$ lorsque seule la $l^{\text{ème}}$ sollicitation n'est pas nulle et de valeur unitaire. Son calcul fait souvent intervenir des fonctions de Green. Sa définition mathématique peut être consultée dans la référence [39]. $U_l(t)$ est l'évolution temporelle de la $l^{\text{ème}}$ sollicitation appliquée au système thermique (ces quantités se déduisent des puissances volumiques $U(M, t)$).

La partie glissante de la température vérifie le problème suivant

$$\left| \begin{array}{l} \forall M \in \mathcal{D}^0 \quad \mathcal{L}[T_g(M, t)] + \varphi(M, t) = 0 \\ \forall M \in \partial \mathcal{D} \quad \mathcal{B}[T_g(M, t)] + \varphi(M, t) = 0 \end{array} \right. \quad (1.13)$$

Remarque: Le champ glissant $T_g(M, t)$ vérifie une équation différentielle de l'espace, la variable temporelle de (1.13) ne jouant que le rôle de paramètre.

A ce stade, on dispose des équations nécessaires à l'obtention du modèle d'état modal. Pour cela, on effectue deux opérations simples:

- On remplace dans la première relation de (1.11) $T_d(M, t)$ et $T_g(M, t)$ par leurs équivalents obtenus à partir de (1.10) et (1.12) respectivement.
- On multiplie l'équation obtenue par $V_k(M)$ et on procède à une transformation intégrale sur l'espace, faisant ainsi apparaître la relation d'orthonormalité (1.6) afin d'éliminer la variable spatiale.

Ces opérations associées à (1.5) mènent à la relation suivante:

$$\forall i = 1, 2, \dots \quad \frac{dx_i(t)}{dt} = \lambda_i x_i(t) - \sum_{l=1}^p \left(\int_{\mathcal{D}} C(M) S_l(M) V_i(M) dM \right) \dot{U}_l(t) \quad (1.14)$$

qui peut également s'écrire

$$\dot{x}_i(t) = \lambda_i x_i(t) + \sum_{l=1}^p B_{i,l} \dot{U}_l(t) \quad (1.15)$$

$$\text{avec } B_{i,l} = - \int_{\mathcal{D}} C(M) S_l(M) V_i(M) dM \quad (1.16)$$

Dans la pratique, on ne retient qu'un nombre N fini de modes propres. Ce nombre peut varier de quelques modes seulement (typiquement les systèmes inertes) à quelques centaines, voire quelques milliers. Aussi, les N équations du type (1.15) que vérifient les modes propres se mettent sous une forme matricielle compacte

$$\dot{x}(t) = Fx(t) + B\dot{U}(t) \quad (1.17)$$

B est appelée *matrice de commande* et est de dimension $[N, p]$. Elle décrit les effets directs que peut avoir le vecteur dérivé (par rapport au temps) des sollicitations $\dot{U}(t)$ (dimension $[p]$) sur le vecteur d'état $x(t)$ (dimension $[N]$). F est la matrice $[N, N]$ diagonale des valeurs propres du problème (1.11).

Il faut souligner que (1.17) forme un système d'équations différentielles *découplées*. En effet, la détermination de l'état $x_i(t)$ n'est pas liée à celle des autres états. C'est ce que montre (1.14). Le calcul du vecteur d'état dépend des sollicitations agissant sur le système thermique et son expression générale est donnée par

$$x(t) = e^{Ft} \left(x(0) + \int_0^t e^{-F\xi} B \dot{U}(\xi) d\xi \right) \quad (1.18)$$

L'autre intérêt de la formulation (1.17) repose sur le fait que les états des modes n'ont pas la même importance dans (1.10). Il existe divers "critères quantitatifs" qui peuvent renseigner sur l'importance de chaque mode propre par rapport au champs des sorties observée. La réduction de (1.15) peut alors se faire en ne gardant que les états associés aux modes dominants. Nous reviendrons en détail sur ce point au prochain chapitre.

Une démarche mathématique similaire peut être appliquée à l'équation des sorties (1.4). On remplace dans (1.4) $T(M, t)$ par son équivalent issu de (1.9), (1.10) et (1.12) et on obtient la forme

$$y_i(t) = \sum_{j=1}^N H_{i,j} x_j(t) + \sum_{l=1}^p S_{i,l} U_l(t) \quad (1.19)$$

et pour les q sorties on peut former le système matriciel

$$y(t) = Hx(t) + SU(t) \quad (1.20)$$

H est appelée *matrice d'observation* (dimension $[q, N]$) et S *matrice statique* (dimension $[q, p]$).

Les deux équations (1.17) et (1.20) forment le modèle d'état modal. Les quatre matrices F, B, H et S sont appelées *estimateur*. L'estimateur suffit à caractériser complètement le comportement d'un système thermique.

1.3.2 Formulation modale par discrétisation spatiale

Dans la pratique, il est rare que la description formelle précédente puisse être directement exploitée pour répondre à des problèmes thermiques. Sauf pour quelques cas académiques³ simples, le système (1.1) nécessite une discrétisation avant toute simulation.

Nous donnons ci-dessous les grandes lignes de la démarche à suivre et qui peut être facilement justifiée en revoyant les résultats formels donnés dans la section précédente. Pour obtenir le modèle d'état modal discret de (1.1), seule une discrétisation sur la variable spatiale est nécessaire, et on obtient la forme matricielle

$$C \frac{dT(t)}{dt} = AT(t) + EU(t) \quad (1.21)$$

Avec

- $T(t)$ vecteur $[N]$ des températures aux noeuds de discrétisation.
- C matrice $[N, N]$ des capacités calorifiques aux noeuds de discrétisation. Elle est diagonale ($C_{ii} = \rho_i c_{p_i}$) si la méthode de discrétisation adoptée est de type nodale ou différences finies.
- A matrice $[N, N]$ des échanges thermiques internes à $\mathcal{D}$.
- E matrice $[N, p]$ de couplage thermique entre $\mathcal{D}$ et son milieu environnant.

³Nous avons développé un logiciel nommé MurAna (ie:Mur Analytique) pour la caractérisation thermique des structures planes composites. La formulation utilisée est analytique. L'idée consiste à mettre en œuvre une équation transcendante aux valeurs propres qui sera ensuite résolue par une dichotomie. Le logiciel comporte des algorithmes particuliers qui le rendent fiable face au risque d'oubli de valeurs propres au cours de leur exploration. Un module graphique est interfacé à MurAna afin de permettre une visualisation des sorties, modes propres, spectres de réponse ..

- $U(t)$ vecteur $[p]$ des sollicitations appliquées au système thermique.

L'équation des sorties est donnée par

$$y(t) = JT(t) + G_0U(t) \quad (1.22)$$

où J et G_0 sont deux matrices ($[q, N]$ et $[q, p]$ respectivement) appropriées. Dans le cas particulier où $y(t) = T(t)$, alors $J = I_N$ (matrice identité d'ordre N) et G_0 est la matrice nulle.

La diagonalisation de $C^{-1}A$ fournit les valeurs propres et les vecteurs propres associés au problème thermique. Les éléments propres sont tels que

$$C^{-1}A = PFP^{-1} \quad (1.23)$$

Avec

- P matrice $[N, N]$ des modes propres V_i placés en colonnes. Ces modes propres sont l'approximation discrète des fonctions propres $V_i(M)$ de la section précédente.
- F matrice $[N, N]$ diagonale des valeurs propres associées aux modes propres.

La relation d'orthonormalité (1.6) devient

$$\forall i = 1, 2, \dots \quad \forall j = 1, 2, \dots$$

$$V_i^T C V_j = \delta_{ij} \quad (1.24)$$

$$\text{ou bien} \quad P^T C P = I_N \quad (\text{matrice identité d'ordre } N) \quad (1.25)$$

Comme précédemment on peut écrire

$$T(t) = T_d(t) + T_g(t) \quad (1.26)$$

$$\text{avec} \quad T_d(t) = Px(t) \quad (1.27)$$

$$\text{et} \quad T_g(t) = S U(t) \quad (\text{d'après (1.21)}) \quad (1.28)$$

$$S = -A^{-1}E \quad [N, p]$$

Finalement, en remplaçant $T(t)$ de (1.21) par son équivalent donné par (1.26) et (1.27) on obtient

$$\dot{x}(t) = Fx(t) + B\dot{U}(t) \quad (1.29)$$

$$\text{avec } B = P^T C A^{-1} E \quad (1.30)$$

En ce qui concerne l'équation des sorties (1.22), on effectue la même substitution et on arrive à

$$y(t) = Hx(t) + GU(t) \quad (1.31)$$

$$\text{avec } H = JP \quad \text{et} \quad G = G_0 - JA^{-1}E \quad (1.32)$$

Remarque: $G = S$ si $y(t)$ est le champ $T(t)$ de température car dans ce cas, $G_0 = 0$ (matrice nulle) et $J = I$ (matrice identité).

L'équation d'état (1.29) et celle des sorties (1.31) constituent le modèle d'état modal discret. Le prix du passage à la formulation modale est donc principalement celui de la diagonalisation de $C^{-1}A$. Ce prix peut être considérable en termes de temps et moyens de calcul si le nombre de nœuds de discrétisation est grand. Cependant, les travaux récents de Flament [28] sur la synthèse modale montrent comment le modèle modal d'un système thermique de grande dimension peut être déduit à partir du raccordement des modèles modaux des composants élémentaires de ce même système. On est donc amené selon cette nouvelle approche, à diagonaliser plusieurs matrices de faibles dimensions au lieu de faire celle d'une grande matrice. Parmi les avantages obtenus, il y a le gain en temps de calcul.

1.3.3 Formulation modale pour les régimes forcés

Le modèle d'état modal (1.17) et (1.20) fait intervenir le vecteur d'entrée $U(t)$ et sa première dérivée temporelle $\dot{U}(t)$ qu'on peut noter aussi $U^{(1)}(t)$. En fait, on peut généraliser cette formulation aux dérivées $U^{(j)}(t)$ pour $j \geq 1$ [55]. On donne les résultats⁴ suivants, qui sont démontrés dans l'annexe C:

$$\left| \begin{array}{l} \dot{x}_{(j)}(t) = Fx_{(j)}(t) + B_{(j)}U^{(j)}(t) \\ y(t) = Hx_{(j)}(t) + g(U^{(j-1)}(t), U^{(j-2)}(t), \dots, U(t)) \end{array} \right. \quad (1.33)$$

⁴On suppose que les composantes du vecteur $U(t)$ vérifient les conditions de dérivabilité autant de fois que nécessaire.

avec

$$B_{(j)} = F^{-j} B$$

$$\text{et } g\left(U^{(j-1)}(t), U^{(j-2)}(t), \dots, U(t)\right) = S U(t) - \sum_{k=1}^{j-1} H F^{-k} B U^{(k)}(t)$$

Le nouveau vecteur $x_{(j)}(t)$ est lié au vecteur d'état $x(t)$ par l'expression

$$x_{(j)}(t) = x(t) + \sum_{k=1}^{(j-1)} F^{-k} B U^{(k)}(t)$$

L'intérêt de cette formulation généralisée apparaît lorsque l'on veut étudier les régimes établis complexes correspondant à $U^{(m)} = Cte$ pour $m \geq 1$. Si l'on doit réduire un modèle d'état du type (1.33), alors il est utile de ne pas altérer la fonctionnelle $g(\dots U(t))$ lors de la réduction. Autrement dit, au niveau de l'équation des sorties de (1.33), la procédure de réduction doit s'appliquer seulement à la composante dynamique $H x_{(j)}(t)$ de la réponse.

Chapitre 2

Réduction de modèles: les différentes approches

2.1 Préliminaire

La formalisme modal appliqué à la thermique vient d'être présentée dans ses grandes lignes. L'exploitation de ce formalisme pour modéliser un système thermique consiste à mettre en forme le système matriciel regroupant (1.17) et (1.20) dans le cas continu, ou (1.29) et (1.31) dans le cas discret. Quelques avantages offerts par cette formulation ont été cités. Parmi eux, la possibilité de réduire la forme canonique (1.17) et (1.20) nous intéresse particulièrement dans ce chapitre. Mais avant de donner un résumé bibliographique sur ce sujet, posons d'abord le problème de réduction:

Reprenons le modèle d'état modal (1.17) et (1.20), qui est de dimension¹ N

$$\begin{aligned}\dot{x}(t) &= Fx(t) + B\dot{U}(t) \\ y(t) &= Hx(t) + SU(t)\end{aligned}\tag{2.1}$$

Partant de cette description, l'objectif de la réduction est de trouver un modèle de dimension n , avec $n \ll N$, et qu'on note

$$\begin{aligned}\tilde{\dot{x}}(t) &= \tilde{F}\tilde{x}(t) + \tilde{B}\dot{U}(t) \\ \tilde{y}(t) &= \tilde{H}\tilde{x}(t) + \tilde{S}U(t)\end{aligned}\tag{2.2}$$

¹On désigne par dimension du modèle d'état, le nombre de ses modes propres ou encore la dimension du vecteur d'état $x(t)$. Un modèle modal de dimension m a les tableaux suivants: $x[m]$, $F[m, m]$, $B[m, p]$, $U[p]$, $y[q]$, $H[q, m]$ et $S[q, p]$.

De façon générale, les qualités que l'on peut attendre de ce nouveau modèle sont les suivantes:

- ▶ fournir une bonne approximation (au sens d'un critère à préciser et qui dépend de la méthode) du vecteur des sorties $y(t)$, éventuellement pour une classe d'entrées et/ou une plage temporelle données.
- ▶ Atteindre le même régime permanent que celui donné par (2.1).
- ▶ Posséder les mêmes caractéristiques de stabilité que le système (2.1).
- ▶ Etre simple et si possible, relié mathématiquement même de façon approchée, au système (2.1).

A ces qualités, qui pratiquement sont des contraintes, il est nécessaire de rajouter celles de la méthode employée pour opérer la réduction

- ▶ Guider l'utilisateur lors du choix de la dimension " n " du modèle réduit.
- ▶ Ne pas être pénalisante en temps de calcul et/ou espace mémoire. Ces facteurs dépendent du matériel utilisé.

Certaines de ces conditions peuvent être plus déterminantes que d'autres et dans certaines situations, on peut être amené à rajouter d'autres contraintes spécifiques. Cela dépend beaucoup de l'utilisation réservée au système (2.2).

Naturellement, on peut penser que la complexité de la réduction s'accroît lorsque le nombre des contraintes imposées au système (2.2) augmente. En réalité, il existe des situations où une contrainte peut s'avérer particulièrement profitable. On verra ainsi, lorsque nous présenterons au chapitre 3 la *méthode de réduction par amalgame modal*, qu'une contrainte sur l'orthogonalité (relation semblable à (1.6)) de $\tilde{P}$ (matrice des pseudo-modes associés au système (2.2)) facilitera les développements formels.

Les recherches sur l'approximation des grands systèmes par des modèles d'état de dimension réduite a donné lieu à beaucoup de publications. En 1979, Michailesco [48] avait déjà recensé plus de deux cent publications relevant de divers domaines: automatique, mécanique, thermique ... Aujourd'hui ce sujet semble retenir sérieusement l'attention des thermiciens. Le dernier colloque [14] de la SFT (Société Française des Thermiciens) et l'école d'été de Cargèse [36] ont porté une attention particulière à la réduction de modèles thermiques et l'approche modulaire en modélisation.

Les modèles réduits ne sont pas confrontés au problème d'instabilité numérique car le spectre des valeurs propres de l'opérateur $\mathcal{L}$ de la chaleur (équation (1.2)) est à valeurs réelles négatives. Les méthodes de réduction que nous présentons dans ce chapitre produisent des modèles dont le spectre des valeurs propres est un sous ensemble du spectre original. Ainsi, ces méthodes ne comprennent pas de contrainte mathématique sur la stabilité du

modèle réduit. Par ailleurs il n'est pas rare de rencontrer des techniques de réduction qui reposent sur la connaissance approfondie qu'on a du système étudié. Ces techniques n'étant pas généralisables, nous les écartons ici.

Notre domaine d'investigation est donc celui des méthodes de réduction applicables à tout système thermique pouvant être décrit par un modèle du type (2.1). Parmi celles que nous décrirons, il y a celles que nous avons testées et qui nous ont permis d'apporter quelques conclusions. Il y a ensuite celles qui ont retenu notre attention parce qu'elles présentent un intérêt ou s'inscrivent dans la même démarche de réduction que celle que nous préconisons ici. A la lumière de ce chapitre, nous comprendrons mieux les difficultés liées à ce sujet et la démarche qui nous a conduit à la formulation d'une nouvelle méthode: *la réduction par amalgame modal* que nous donnerons au §3.

Nous classons les méthodes de réduction dans ce chapitre en deux groupes:

- Les méthodes de troncature qui consistent à ne conserver que les éléments propres jugés *intéressants* et les états associés. Cela nécessite un critère pour quantifier l'importance de chaque élément propre dans le comportement thermique du système. De manière générale, ces méthodes sont simples à mettre en œuvre.
- Les méthodes qui font appel aux techniques de minimisation. Elles se justifient par l'insuffisance des méthodes de troncature d'une part et les larges possibilités offertes par le domaine de l'optimisation d'autre part. Leur mise en œuvre est beaucoup plus complexe que celle des méthodes de troncature.

2.2 Les méthodes de troncature

2.2.1 Présentation

Toutes les méthodes de troncature partent de la même idée: pour la plupart des systèmes thermiques, seuls quelques modes propres sont importants pour la connaissance de leur comportement thermique relativement à un ensemble donné d'entrées et de sorties. Il est donc envisageable de limiter le modèle détaillé (2.1) à la dynamique des seuls modes importants. Pour cela, on réorganise² (2.1) de façon à faire apparaître deux parties:

- une partie (indice d) associée aux modes importants (ou **dominants**).
- une deuxième partie (indice f) associée aux modes qui ne sont pas jugés importants (ou **faibles**).

$$\begin{pmatrix} \dot{x}_d(t) \\ \dot{x}_f(t) \end{pmatrix} = \begin{pmatrix} F_d & 0 \\ 0 & F_f \end{pmatrix} \begin{pmatrix} x_d(t) \\ x_f(t) \end{pmatrix} + \begin{pmatrix} B_d \\ B_f \end{pmatrix} \dot{U}(t) \quad (2.3)$$

²Il s'agit de permuter des lignes ou des colonnes des différents tableaux du modèle d'état détaillé.

$$y(t) = \begin{pmatrix} H_d & H_f \end{pmatrix} \begin{pmatrix} x_d(t) \\ x_f(t) \end{pmatrix} + S \begin{pmatrix} U(t) \end{pmatrix} \quad (2.4)$$

Le modèle réduit (2.2) s'obtient alors en négligeant toute la partie indicée en f . Ainsi, le modèle réduit (2.2) est ici donné par

$$\tilde{x}(t) = x_d(t) \quad \tilde{F} = F_d \quad \tilde{B} = B_d \quad \text{et} \quad \tilde{H} = H_d$$

On peut imaginer la troncature modale de manière simple en considérant un point d'observation particulier M_o . La température dynamique en ce point est donnée par

$$T_d(M_o, t) = \sum_{i=1}^N x_i(t) V_i(M_o)$$

où $V_i(M_o)$ $i = 1 \dots N$ sont les fonctions propres du système thermique considéré au point M_o . Si l'on ordonne maintenant l'ensemble des modes propres de telle sorte que les n premiers ($V_i^d(M_o)$ $i = 1 \dots n$) soient (en un sens à préciser par un critère) dominants et les autres jugés négligeables ($V_j^f(M_o)$ $j = n + 1 \dots N$) alors $T_d(M_o, t)$ peut aussi s'écrire

$$T_d(M_o, t) = \sum_{i=1}^n x_i^d(t) V_i^d(M_o) + \sum_{j=n+1}^N x_j^f(t) V_j^f(M_o)$$

Cette dernière relation peut être représentée dans un repère dont l'abscisse représente la "direction" importante $\mathbf{P}_d = [V_1^d \dots V_n^d]$ et l'ordonnée la "direction" faible $\mathbf{P}_f = [V_{n+1}^f \dots V_N^f]$. Cette représentation imagée est donnée sur la figure 2.1. La troncature est alors équivalente à faire l'approximation de $T_d(M_o, t)$ par seulement sa projection sur $\mathbf{P}_d$ qui est $\sum_{i=1}^n x_i^d(t) V_i^d(M_o)$. La qualité de la troncature conditionne la distance entre le point M_o et l'axe des abscisses: *la meilleure troncature est obtenue lorsque, à tout instant, la distance entre M_o et l'axe des abscisses est minimale*. Le choix des modes propres de la direction $\mathbf{P}_d$ est donc ici fondamental. Notons enfin que, quelle que soit la position initiale du point M_o , sa trajectoire rejoint l'axe des abscisses au bout d'un temps fini qui ne dépasse pas globalement $4\tau_n$.

De manière générale, les méthodes de troncature ne diffèrent entre elles que par la manière de choisir les éléments propres importants qui conditionnent la partition (2.3). Chacune utilise un "critère" de sélection des modes importants. Aussi, nous résumons ci-dessous quelques critères de troncature modale souvent utilisés.

2.2.2 Critère de Marshall

Cette méthode introduite par Marshall [46], également utilisée par Chidambara [16], Fos-sard [29] et Davison [19] repose sur l'analyse du spectre des valeurs propres contenues dans

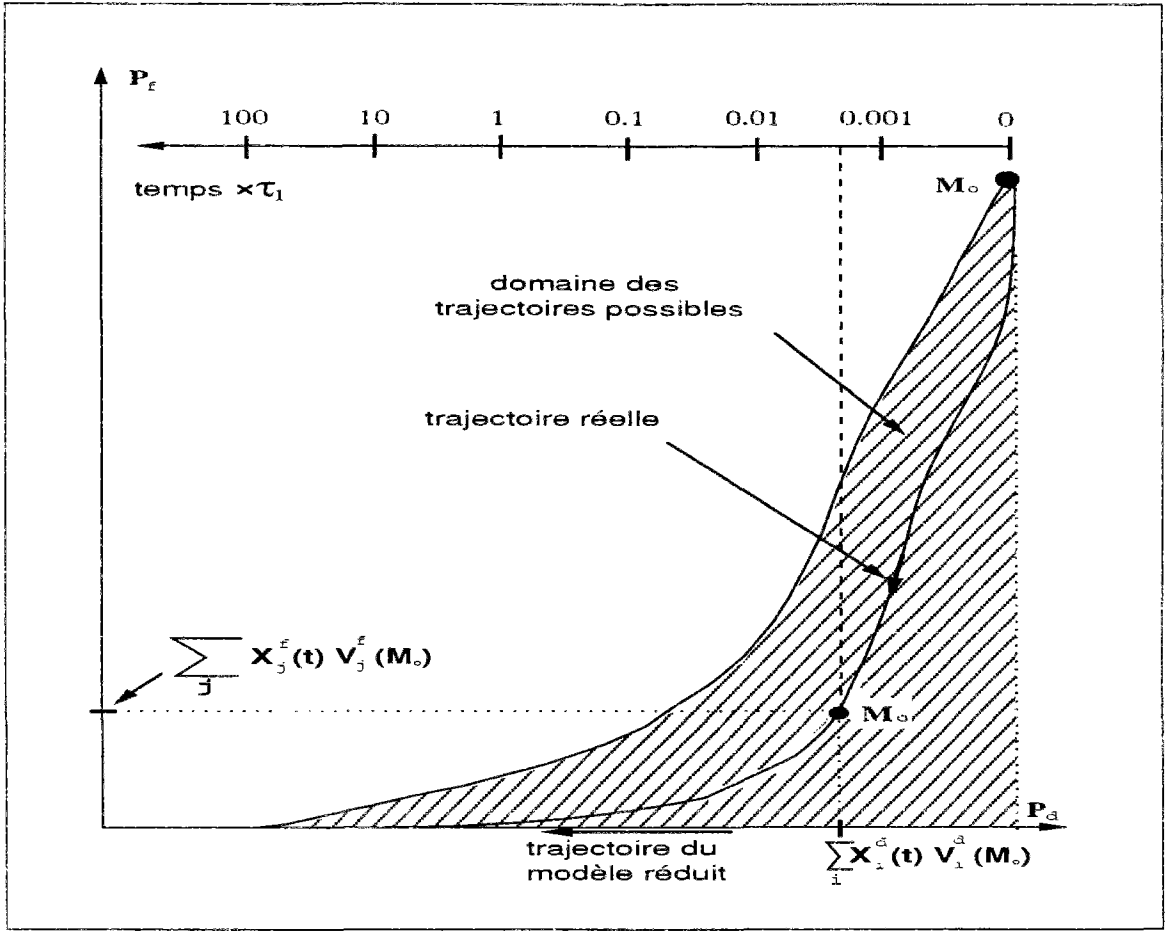


Figure 2.1: Principe de la troncature modale

la matrice F de (2.1). Elle a aussi été testée par de nombreux chercheurs de l'Ecole des Mines de Paris [60][42][41] [8][7][13]. Elle est pratiquée de manière presque systématique en mécanique vibratoire. L'importance d'un élément propre est ici fonction de la valeur de sa constante de temps définie par

$$\tau_i = -\frac{1}{\lambda_i} \quad i = 1 \dots N \quad \text{avec} \quad \lambda_i = F_{ii} \quad (2.5)$$

Les éléments propres sélectionnés sont donc les n premiers de la relation d'ordre (1.7). L'idée vient du fait que le champ de température fait intervenir lors d'une relaxation une série de termes $(a_i \exp(\lambda_i t))$ relatifs aux éléments propres $i = 1, 2 \dots N$. On considère alors que ces quantités subissent un amortissement d'autant plus rapide que τ_i est petit.

La méthode de Marshall est la plus simple possible pour la mise en œuvre pratique. Elle peut être utile si l'on cherche à étudier le régime d'évolution lente [15] d'un système

thermique. Mais elle ne peut être prise comme une démarche systématique de réduction du fait que la dynamique rapide est totalement écartée dans la troncature. En effet, l'importance des modes propres est à la fois contenue dans les constantes de temps et les coefficients a_i , ces derniers étant fonction de la matrice $[B]$ de commande et des amplitudes des modes.

Notons enfin que beaucoup d'auteurs dont Fossard [29], Lefebvre [41], Salgon [58], Marshall [46] et Litz [44] ont, chacun à sa manière, apporté une amélioration de cette méthode en rajoutant à l'équation des sorties de (2.2) un terme supplémentaire qui provient des états rapides négligés. Par exemple, l'approche de Fossard se base sur le fait que les états contenus dans $x_f(t)$ atteignent presque instantanément leurs valeurs asymptotiques, ce qui permet d'écrire d'après (2.3)

$$\dot{x}_f = 0 \Rightarrow x_f(t) \simeq -F_f^{-1} B_f \dot{U}(t)$$

et l'équation *réduite* des sorties s'obtient alors de (2.4)

$$\tilde{y}(t) \simeq H_d x_d(t) + K \dot{U}(t) + S U(t) \quad \text{avec} \quad K = -H_f F_f^{-1} B_f$$

En pratique, il semble que l'introduction du terme de couplage $K \dot{U}(t)$ entre les sorties du système et la dérivée première de $U(t)$ n'a pas d'effet sensible devant le terme $S U(t)$. Ceci est rigoureusement vrai pour des sollicitations constantes. Dans ce qui suit, nous n'allons pas tenir compte de ce terme de couplage.

2.2.3 Critère de Litz

La méthode de Litz [44] est plus performante que celle de Marshall en ce sens qu'à précision égale, elle conduit à retenir un nombre plus faible de modes. Elle est plus sélective. Toutefois, elle s'applique à un couple entrée-sortie. L'étude et l'analyse de cette méthode a fait l'objet de plusieurs travaux [41] [42][15] [13] [8] [58].

Considérons le modèle d'état discret formé de (1.29) et (1.31). La sortie y_i lorsqu'elle est sollicitée par seulement la " $j^{\text{ème}}$ " entrée U_j est notée y_{ij} . En prenant pour U_j l'échelon de Heaviside, y_{ij} est la réponse indicielle associée et qui est donnée par

$$y_{ij} = S_{ij} \left(1 - \sum_{m=1}^N \frac{H_{im} B_{mj}}{S_{ij}} e^{\lambda_m t} \right) \quad (2.6)$$

Les quantités réelles $\left(\alpha_{ij}^m = \frac{H_{im} B_{mj}}{S_{ij}} \quad m = 1 \cdots N \right)$ sont les valeurs des raies associées aux différents modes propres. L'ensemble de ces raies forment le *spectre de réponse indicielle* [41] du couple: entrée " j " / sortie " i ". Pour représenter ce spectre, on a coutume

d'utiliser une échelle de type logarithmique³ sur laquelle on place chaque raie α_{ij}^m verticalement et à l'instant correspondant à la constante de temps τ_m auquel elle correspond. La figure 1.1 illustre l'aspect général d'un spectre de réponse indicielle sur lequel on trace également l'évolution de la sortie "i" soumise à l'action de la "j^{ème}" sollicitation.

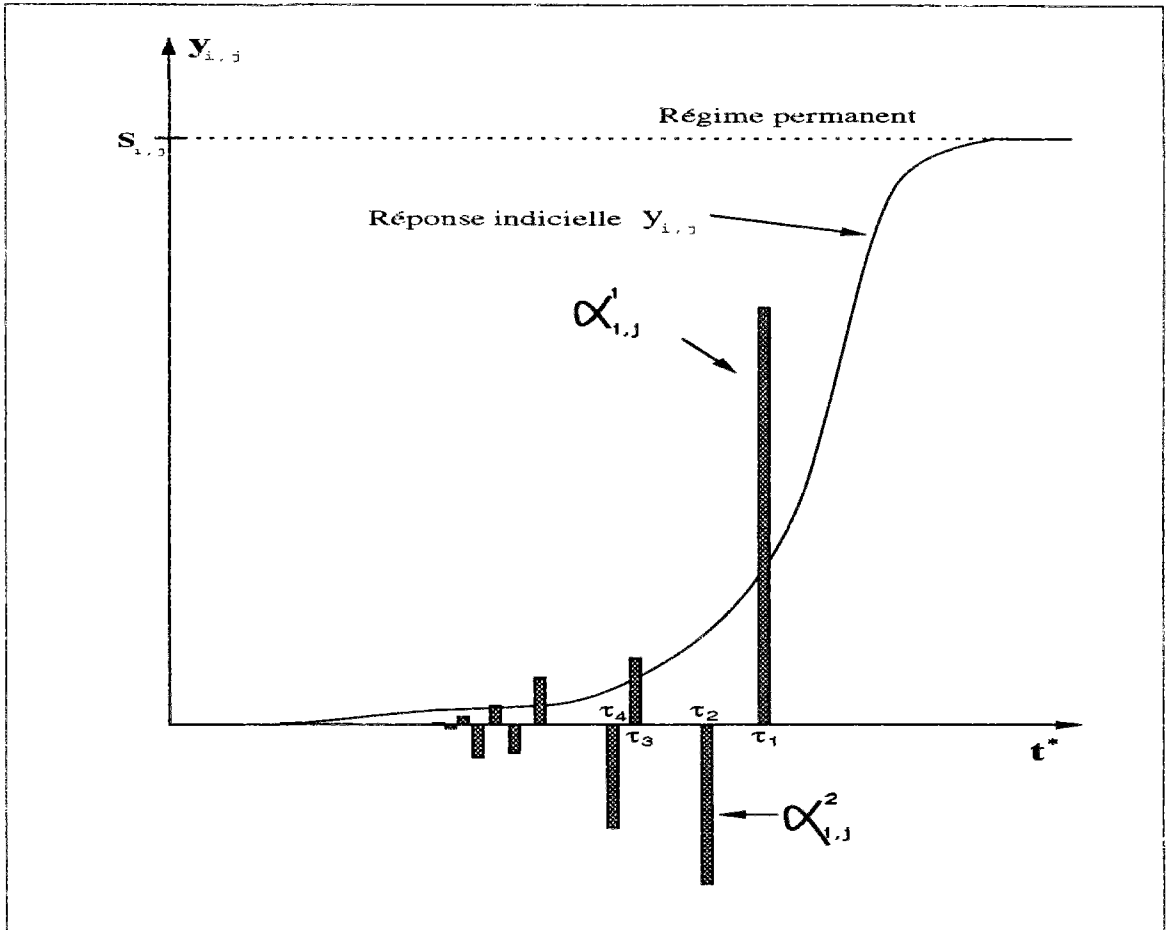


Figure 2.2: Présentation d'un spectre de réponse indicielle (alterné)

La méthode de Litz convient alors de considérer comme importants les n modes ayant les raies les plus grandes. Le signe d'une raie n'est pas pris en compte lors de la sélection, cependant il indique pour chaque mode propre s'il contribue instantanément à la réponse (cas d'une raie positive) ou s'il la retarde (cas d'une raie négative).

Le critère de Litz n'est pas toujours efficace. En effet, il ne réalise pas toujours le minimum de l'écart $|y_{ij} - \tilde{y}_{ij}|$ car la série (2.6) peut être alternée. Par ailleurs, même si les raies

³Nous utilisons une échelle temporelle définie par la variable t^* telle que: $t^* = \frac{\text{Ln}\left(1 + 100 \frac{t}{t_{ref}}\right)}{\text{Ln}(101)}$. Pour $t = t_{ref}$, on obtient $t^* = 1$. La variation du paramètre t_{ref} permet d'avoir une échelle dilatable de telle sorte que les raies puissent être distinguées et notamment celles qui concernent la dynamique rapide.

sont toutes du même signe, il faut aussi tenir compte de l'amortissement exponentiel ($\exp(\lambda_m t)$). Ce critère reste *local* puisqu'il s'applique à un couple entrée-sortie.

remarque: On peut généraliser le critère de Litz pour l'ensemble des entrées et des sorties du système (1.29) et (1.31). Pour chaque mode, il faut alors associer la somme de ses raies relativement à tous les couples entrées-sorties possibles (on peut procéder à une normalisation des nouvelles raies ainsi obtenues). Mais la qualité des résultats peut se dégrader beaucoup même si les raies sont de même signe.

2.2.4 Critère de dominance

Le problème de réduction des modèles d'état a été d'abord une préoccupation des automaticiens [19] [20]. En thermique, il existe des travaux [15] préconisant une démarche similaire à celle des automaticiens pour réduire les modèles d'état en introduisant un critère "énergétique".

L'énergie⁴ associée à un système thermique durant la plage temporelle $\mathcal{D}_t = [t_1, t_2]$ et sollicité par la seule composante $U_j(t)$ du vecteur des entrées $U(t)$ est

$$E_j = \int_{\mathcal{D}_t} \int_{\mathcal{D}} \left(T_d^j(M, t) \right)^2 C(M) \, dM \, dt \quad (2.7)$$

Si l'on considère toutes les composantes du vecteur $U(t)$ de manière indépendante et en utilisant des excitations tests (échelon de Heaviside ou Dirac), alors la *dominance* des modes propres est [annexe A]

$$\forall i = 1 \dots N \quad d_i = \frac{1}{p} \sum_{j=1}^p e_i^j \quad (2.8)$$

avec

$$e_i^j = -\frac{B_{ij}^2 \lambda_i^{2s-1}}{2 E_j} \left(e^{2\lambda_i t_1} - e^{2\lambda_i t_2} \right) \quad (2.9)$$

et l'on a

$$0 \leq d_i \leq 1 \quad \text{et} \quad \sum_{i=1}^N d_i = 1$$

Dans (2.9), $s = 0$ pour l'échelon de Heaviside et $s = 1$ pour le Dirac. La dominance d_i (dite aussi énergie de mode) représente un indicateur quantitatif de l'importance du $i^{\text{ème}}$ mode propre vis-à-vis de l'évolution dynamique de l'intégralité du système thermique

⁴Il convient de noter que le mot énergie doit être interprété ici au sens que lui donnent les automaticiens. Elle est sans rapport avec l'énergie physique. En particulier, elle n'est pas une grandeur extensive et ne peut être exprimée en Joules.

durant l'intervalle $\mathcal{D}_t$. Ce critère de sélection est donc *global*, ce qui est une qualité intéressante. Néanmoins, l'expérimentation numérique de cette approche montre que les éléments propres sélectionnés par (2.8) dépendent de la forme de la sollicitation choisie:

Pour l'échelon, les modes propres dominants représentent globalement la dynamique lente.

Pour l'impulsion, le critère a tendance à imposer trop de modes rapides. Il peut alors se produire une propagation des erreurs pour des instants assez grands (approximativement de $0.1\tau_1$ à $2\tau_1$).

2.2.5 Critère de dominance pondérée

Les modèles tronqués par le critère de dominance (§2.2.4) sont généralement biaisés pour les temps d'évolution rapide dans le cas des sollicitations échelons. L'idée qui nous a conduit à établir ce nouveau critère est la suivante: on espère disposer d'un critère spatio-temporel (donc global) qui permette de sélectionner des modes importants répartis sur toutes les dynamiques. Cela signifie qu'on cherche des modèles réduits présentant des erreurs faibles mais réparties sur une large plage temporelle. On évite ainsi les "grands" biais localisés au début de l'évolution thermique.

Pour cela, on remplace la variable de l'intégrale temporelle de (2.7) par la variable t^* de l'axe logarithmique (celui utilisé pour tracer le spectre de réponse indicelle de la figure 1.3). En effet, sur cet axe, les constantes de temps sont réparties plus régulièrement. la variables t^* se déduit de t par la relation

$$t^* = \frac{\text{Ln} \left(1 + 100 \frac{t}{t_{ref}} \right)}{\text{Ln}(101)} = g(t) \quad (2.10)$$

De manière similaire à (2.7), on écrit

$$E_j = \int_{\mathcal{D}_t} \int_{\mathcal{D}} (T_d(M, t))^2 C(M) dM dt^* \quad \text{avec} \quad dt^* = \left(\frac{\partial g(t)}{\partial t} \right) dt \quad (2.11)$$

La dominance de chaque mode propre au sens de (2.11) est alors [annexe A]

$$\forall i = 1 \dots N \quad d_i = \frac{1}{p} \sum_{j=1}^p e_i^j \quad (2.12)$$

avec

$$\left| \begin{aligned} e_i^j &= \frac{1}{\text{Ln}(101)} B_{ij}^2 \lambda_i^{2s} e^{\frac{2}{\beta \tau_i}} \int_{y_1}^{y_2} \frac{e^{c_i y}}{y} dy \\ \beta &= \frac{100}{t_{ref}}, \quad c_i = -\frac{2}{\beta \tau_i} \quad \text{et} \quad (y_k = 1 + \beta t_k \quad k = 1, 2) \end{aligned} \right. \quad (2.13)$$

et l'on a

$$0 \leq d_i \leq 1 \quad \text{et} \quad \sum_{i=1}^N d_i = 1$$

Comme précédemment, $s = 0$ pour l'échelon et $s = 1$ pour le Dirac. Le calcul effectif des dominances, fait intervenir l'intégrale non classique de l'équation (2.13). Elle peut être facilement résolue par la voie numérique. Pour notre part, nous l'avons tabulée et vérifiée par la méthode d'intégration de Newton-Côtes [18].

L'expérimentation numérique montre que les modèles tronqués par ce critère présentent certes un biais, mais globalement étalé sur une plage de temps assez large, ce qui peut être intéressant si l'erreur en chaque point et à tout instant reste raisonnable.

Remarques:

- Il ne suffit pas de choisir des modes propres ayant des constantes de temps réparties selon toutes les dynamiques pour obtenir un modèle tronqué acceptable. Il faut également que les modes propres vérifient des propriétés spatiales de manière à couvrir l'intégralité du domaine géométrique définissant le système thermique. C'est en ce sens que la dominance est représentative. Michailesco [48] résume bien cette réalité en écrivant: **un mode rapide de forte amplitude n'est pas négligeable par rapport à un mode lent de faible amplitude.**
- Il n'est pas évident de conclure qu'un modèle tronqué obtenu par le critère présenté dans cette section soit "plus précis" que celui obtenu par le critère de dominance de §2.2.4. En effet, seul l'estimation de la surface écart entre $\tilde{T}(M, t)$ et $T(M, t)$ dans les deux cas le prouverait. En revanche, le critère de dominance pondéré permet généralement de *mieux répartir* l'erreur en l'*étalant* sur une plage temporelle assez large.

2.2.6 Récapitulatif sur la troncature modale

Nous venons de décrire la réduction de modèles modaux par la troncature. Quatre critères de sélection des modes propres ont été esquissés: Marshall, Litz, dominance et dominance pondérée. Les caractéristiques de ces critères sont résumées dans le tableau 2.1

Deux principales caractéristiques de la troncature modale méritent d'être citées:

En ce qui concerne l'automatisation de la réduction par les divers critères présentés, elle n'est possible qu'en imposant l'ordre "n" souhaité du modèle réduit. Autrement dit, on ne dispose pas de mesure de l'erreur capable de fixer automatiquement l'ordre "n" minimal du système (2.2). Ceci est évident pour les deux premiers critères. Dans le cas des deux critères énergétiques (dominance et dominance pondérée), on peut penser à procéder de la manière suivante: prendre autant de modes dominants (au sens de (2.8) ou (2.12)) nécessaires, de telle sorte que la somme de leurs dominances soit supérieure à un seuil d_{min}

fixé préalablement (0.98 par exemple). Dans la pratique, ce procédé reste inexploitable: deux modèles réduits de deux systèmes thermiques différents et réalisant la même dominance totale, peuvent être très différents au niveau de leur précision (écart sur les températures).

La réduction des modèles d'état par la troncature modale peut être efficace, et dans la majorité des cas, à condition que le rapport⁵ de la troncature ne soit pas trop élevé. A titre indicatif, à partir d'un modèle détaillé basé sur quelques centaines d'éléments propres, on peut obtenir un modèle tronqué suffisant ne comprenant que quelques dizaines d'éléments propres. On sait qu'en mécanique vibratoire, la troncature par le critère de Marshall décrit en §2.2.2 autorise des réductions fortes (environ 3 à 5 modes). En thermique, l'utilisation de la troncature est souvent délicate car les dominances des modes propres sont plus dissimulées. Il semble que la répartition de l'information thermique sur les modes propres soit plus "dispersée" que ne l'est l'information vibratoire.

Malgré ces insuffisances, les méthodes de troncature ne doivent pas être définitivement écartées. Par leur simplicité, elles restent en effet d'une grande utilité dans les travaux relevant de l'estimation simplifiée (évolution qualitative d'une sortie, besoins de chauffage, déperditions thermiques...).

⁵On désigne ici par rapport de troncature, le rapport entre les dimensions N du modèle détaillé (2.1) d'une part, et n du modèle réduit (2.2) d'autre part.

critère	type	régime permanent	mesure de l'erreur	automatisation	notes
Marshall	global	conservé	non	Dans la pratique on peut fixer seulement l'ordre 'n' du modèle réduit	suffisant pour les dynamiques lentes. Grands biais sur les dynamiques rapides
Litz	local				peut être généralisé à un nombre quelconque d'entrées sorties mais avec dégradation des résultats
Dominance	global		oui		insuffisant selon le type d'entrée utilisée
Dominance pondérée	global				erreur répartie sur une large plage temporelle

Tableau 2.1: Comparaison qualitative des critères de troncature modale

2.3 Les méthodes de minimisation

2.3.1 Présentation

L'insuffisance de la troncature modale dans la réduction des modèles d'état modaux justifie pleinement le recours aux méthodes s'inspirant du domaine de l'optimisation [17] [57]. Les développements mathématiques et la disponibilité d'algorithmes puissants dans le domaine de l'optimisation sont, sans aucun doute, des facteurs encourageants pour aller dans cette voie. Le nombre de travaux ayant contribué à l'adaptation des techniques d'optimisation pour la recherche de modèles réduits est tellement grand qu'il est difficile de les citer tous ici. A ce sujet, on peut se référer utilement aux travaux de Michaillesco [48] et Petit [55].

Nous nous contentons ici d'évoquer les idées communément exploitées dans ces méthodes et d'en décrire quelques unes qui sont proches de notre démarche de réduction.

Dans les méthodes de réduction de modèles d'état par la voie de la minimisation, il est courant de se baser sur une *mesure de l'erreur de réduction*. Dans la plupart des cas, cette mesure repose sur un écart quadratique [5] [17]. La forme générale de la mesure de l'erreur de réduction peut s'écrire

$$\mathcal{M} = \sum_{j=1}^p \int_{\mathcal{D}_t} \left[\left[\Delta z^j(t) \right]^T \quad f \quad \left[\Delta z^j(t) \right] \right] dt \quad (2.14)$$

avec

$$\Delta z^j(t) = z^j(t) - \tilde{z}^j(t)$$

$z^j(t)$ (resp. $\tilde{z}^j(t)$) représente le plus souvent le vecteur des sorties du modèle détaillé (2.1) (resp. réduit (2.2)) lorsque seule la composante $U_j(t)$ du vecteur $U(t)$ des entrées n'est pas nulle. On dit alors que la mesure $\mathcal{M}$ repose sur une *erreur de sortie* [51]. Par contre si $\Delta z^j(t) = x^j(t) - \tilde{x}^j(t)$, on parle de mesure reposant sur une *erreur d'équation* [24] [22] [26]. f est un paramètre qui peut avoir plusieurs sens parmi lesquels:

- f peut être une fonction scalaire de la variable temporelle en vue de privilégier une plage temporelle donnée.
- f peut aussi être une fonction de filtrage sur les entrées. Ce sujet a été étudié par Galiana [31].
- f peut enfin représenter une matrice $[q, q]$ permettant des pondérations sur les sorties du modèle.

Dans le cas d'une erreur de sortie, le vecteur $y(t)$ contenant des températures en un nombre fini de points, la mesure $\mathcal{M}$ devient une énergie (au même sens que (2.7)) à la pondération par $f(t)$ près. L'appellation "énergie de l'écart" (ou énergie de l'erreur) utilisée en théorie des systèmes est bien adaptée. Elle est donnée dans ce cas par

$$\mathcal{M} = \sum_{j=1}^p \sum_{i=1}^q \int_{\mathcal{D}_t} \left(y_i^j(t) - \tilde{y}_i^j(t) \right)^2 f \, dt$$

Lorsqu'on minimise $\mathcal{M}$, on diminue en réalité "l'énergie de l'écart des sorties" entre les modèles détaillé et réduit.

Il existe plusieurs développements [62] [3] [21] [56] [48] [24] [23] [52] bâtis à partir de la mesure (2.14) ou de l'une de ses formes dérivées. Dans le cadre des problèmes liés à la thermique, elle a été notamment utilisée par Petit et Ben Jaafar [55][11], Saulnier [59] et Merour [47]. La mesure $\mathcal{M}$ fait intervenir en pratique les matrices F , B et H du modèle détaillé, ainsi que $\tilde{F}$, $\tilde{B}$ et $\tilde{H}$ du modèle réduit. On considère souvent $\tilde{S} = S$, ce qui permet d'assurer la conservation des régimes statiques au niveau du modèle réduit. Cependant, Decoster, Noldus et Van Cauwenberghe [21] imposent explicitement la contrainte de conservation des régimes statiques, ce qui requiert des algorithmes spécifiques. Le plus souvent, la minimisation de $\mathcal{M}$ se fait par rapport à un tableau $\mathcal{T}$ que l'on recherche. Si l'on fixe par exemple $\tilde{F}$, $\tilde{H}$ et $\tilde{S}$, il est possible de rechercher $\tilde{B}$ en écrivant

$$\frac{\partial \mathcal{M}}{\partial \tilde{B}} = 0$$

La plupart des méthodes nécessitent la mise en œuvre d'algorithmes de programmation non linéaire [17] [18] parmi lesquels les algorithmes des gradients (simplifiés ou conjugués) [3] [62] et de Marquardt [6] sont de loin les plus utilisés. Lorsque la taille du modèle à réduire est grande, les problèmes de convergence des algorithmes de minimisation ne sont pas exclus et les temps de calculs nécessaires se prolongent.

Notons que la mesure de l'erreur $\mathcal{M}$ sous sa forme (2.14) ne peut être retenue dans une démarche visant à automatiser⁶ la réduction. Ce point sera repris et revu en détail au §3.5 où nous proposerons une normalisation de $\mathcal{M}$. Néanmoins, on peut déjà donner une explication qualitative relevant du sens physique. En effet, la mesure $\mathcal{M}$ représente⁷ le carré de l'écart total dans $\mathcal{D}_t$ entre les sorties données par (2.1) et (2.2) et pour toutes les sollicitations prises individuellement. Cette mesure dépend donc de plusieurs facteurs: la constante de temps principale (notamment si $\mathcal{D}_t = [0, \infty[$), le nombre de sorties et le nombre d'entrées. Toutefois, Michaïlesco [48] et Petit [55] ont abordé le problème de l'automatisation de la réduction et les solutions apportées se résument comme suit:

Michaïlesco utilise une forme normalisée $\overline{\mathcal{M}}$ de la mesure de l'erreur. Cette forme normalisée est définie par le rapport entre $\mathcal{M}$ évaluée pour le modèle réduit et la même mesure évaluée pour un modèle statique (lorsqu'on néglige tous les modes)

$$\overline{\mathcal{M}} = \mathcal{M} / \sum_{j=1}^p \int_{\mathcal{D}_t} \left[\begin{bmatrix} \Delta z_{\infty}^j \end{bmatrix}^T \quad f(t) \quad \begin{bmatrix} \Delta z_{\infty}^j \end{bmatrix} \right] dt \quad (2.15)$$

$$\text{où} \quad \Delta z_{\infty}^j = z^j(t) - z^j(t \rightarrow \infty)$$

⁶ Par automatisation, on entend surtout la possibilité de rechercher le plus petit modèle réduit respectant une contrainte d'erreur ($\mathcal{M} < \mathcal{M}_{max}$) imposée préalablement. Cela suppose que nous disposons d'une échelle qualitative permettant de prédire le degré d'efficacité de la réduction effectuée à partir des valeurs de $\mathcal{M}$. De plus, cette échelle doit être valable quel que soit le système thermique modélisé.

⁷ Pour simplifier ici, on peut considérer que le paramètre f est égal à 1.

ce qui permet d'avoir $0 \leq \overline{\mathcal{M}} \leq 1$ avec $\overline{\mathcal{M}} = 0$ pour le modèle détaillé (le plus précis) et $\overline{\mathcal{M}} = 1$ pour le modèle statique (le plus biaisé car ne contenant aucune information sur la dynamique du système). Mais la fixation de l'ordre "n" de (2.2) repose sur une classification par la *contribution énergétique* des modes propres (approche semblable à celle présentée en §2.2.4). On classe alors les modes dans l'ordre décroissant de leurs contributions $\mathcal{C}_i$, et l'ordre "n" est fixé de telle sorte que $\mathcal{C}_n \gg \mathcal{C}_{n+1}$.

Dans les travaux de Petit [55] où l'intégrale temporelle de (2.14) est remplacée par une somme discrète dans le temps, on emploie une technique différente: on incrémente l'ordre "n" du modèle réduit autant de fois que nécessaire pour obtenir $\mathcal{M}(\text{ordre } n) \simeq \mathcal{M}(\text{ordre } n+1)$. Cette technique suppose implicitement une décroissance régulière de $\mathcal{M}$ lorsque l'ordre n augmente.

Le calcul effectif du modèle réduit (2.2) diffère d'une méthode à l'autre. Elles restent néanmoins proches dans leur esprit du fait qu'elles reposent toutes sur la minimisation d'une entité directement liée aux systèmes (2.1) et (2.2). Nous décrirons ci-dessous brièvement trois méthodes qui peuvent donner une idée concise de l'ensemble des travaux liés à la réduction de modèles d'état par la voie de la minimisation.

2.3.2 Approche d'agrégation

Le principe de l'agrégation date de 1956 et a été introduit en économétrie [45] et appliqué en automatique en 1968 [4]. il repose sur une liaison linéaire entre les vecteurs d'état $x(t)$ de (2.1) et $\tilde{x}(t)$ de (2.2) telle que

$$\tilde{x}(t) = L x(t) \quad (2.16)$$

où L est une matrice rectangulaire $[n, N]$. Aoki [4], puis Michaïlesco [48] ont appliqué cette approche aux systèmes dynamiques invariants [43]. En thermique elle a été également étudiée par Petit [55] et Neirac [49]. Résumons l'approche de Michaïlesco qui est également reprise dans [55].

L'obtention du modèle (2.2) recherché (dit *agrégé*) se fait en deux étapes *indépendantes*.

étape 1

On fait d'abord le choix de "n" modes selon un critère de type énergétique. Pour cela, on construit la mesure $\mathcal{M}$ (relation (2.14)) sur la base de l'écart $(y(t) - y(\infty))$ du système (2.1). on montre ensuite que $\mathcal{M}$ se met sous la forme

$$\mathcal{M} = \sum_{i=1}^N \sum_{j=1}^N \mathcal{C}_{ij}$$

Le classement des modes est ensuite effectué suivant leurs *contributions énergétiques*⁸ C_{ii} ($i = 1 \cdots N$) qui sont données dans le cas des sollicitations échelons par

$$C_{ii} = - \frac{\sum_{l=1}^q \sum_{k=1}^p (H_{l,i} B_{i,k})^2}{2 \lambda_i}$$

Les matrices $\tilde{F}$ et $\tilde{B}$ de (2.2) sont enfin obtenues par simple troncature conformément à §2.2 et en prenant comme importants les premiers modes de plus grandes contributions énergétiques.

Remarques:

► Si l'ordre "n" de (2.2) n'est pas fixé a priori, Michailesco propose de le fixer de telle sorte que $C_{nn} \gg C_{(n+1)(n+1)}$.

► Dans le cas où les modes dominants s'avèrent insuffisants pour l'évolution rapide, on rajoute un ou quelques modes rapides.

► La méthode d'agrégation est une méthode de troncature vis-à-vis de l'équation d'état modale.

étape 2

L'étape 1 a permis d'obtenir $\tilde{F}$ et $\tilde{B}$. Il reste donc à chercher la matrice $\tilde{H}$. Pour cela, on exploite la mesure

$$\mathcal{M} = \sum_{j=1}^p \int_{\mathcal{D}_t} [\Delta y^j(t)]^T f [\Delta y^j(t)] dt$$

avec $\Delta y^j(t) = [y^j(t) - y^j(\infty)] - [\tilde{y}^j(t) - \tilde{y}^j(\infty)]$

En pratique, le calcul de $\Delta y^j(t)$ revient à celui de l'erreur de sortie $[y^j(t) - \tilde{y}^j(t)]$ car à l'horizon infini, les modèles (2.1) et (2.2) sont ici équivalents.

En annulant la première dérivée de $\mathcal{M}$ par rapport à $\tilde{H}$

$$\frac{\partial \mathcal{M}}{\partial \tilde{H}} = 0$$

et en prenant des sollicitation échelons unitaires, on obtient

$$\tilde{H} = \begin{pmatrix} H & W & L^T \end{pmatrix} \begin{pmatrix} L & W & L^T \end{pmatrix}^{-1}$$

où L est définie par la relation d'agrégation (2.16) et W est la matrice $[N, N]$ solution de l'équation de Lyapunov [57] [annexe B] suivante

$$F W + W F = - B B^T$$

⁸ Les termes C_{ij} ($i \neq j$) sont dits *contributions croisées*. Lorsque $C_{ij} < 0$ cela est interprété comme une opposition des contributions des modes i et j . Si $C_{ii} \simeq C_{jj}$ $i \neq j$ et $C_{ij} < 0$, alors ces modes (i et j) ne peuvent être retenus dans l'étape 1 même s'ils sont dominants, car ils se compensent au niveau de l'observation.

2.3.3 Approche d'Eitelberg

Cette approche repose sur la minimisation d'un écart quadratique défini à partir d'une erreur d'équation. Elle a été utilisée par plusieurs auteurs [12] [51] [24] [25] [26]. En thermique, elle a été notamment étudiée par Petit et Ben Jaafar [55] [11], Saulnier [59], Mérour[47] et Ben Jaafar [10]. Nous résumons ici l'approche d'Eitelberg qu'on peut également consulter dans [55] et [24] pour ses détails. On peut la présenter en trois étapes.

étape 1

La méthode suppose comme préalable la connaissance des **états importants** du système (2.1). On suppose que ces états importants sont contenus dans $\tilde{x}(t)$ de (2.2). On peut donc définir une matrice R rectangulaire $[n, N]$ telle que⁹

$$\tilde{x}(t) = R x(t)$$

En réalité, la méthode d'Eitelberg opère dans l'espace physique. Le vecteur d'état $x(t)$ contient alors les températures aux nœuds de la discrétisation et non pas les états des modes propres. Exceptionnellement, il faut considérer ici que les modèles détaillé (2.1) et réduit (2.2) ne sont pas sous forme modale. Le choix des composantes de $\tilde{x}(t)$ qui conditionne celui de R revient à sélectionner les nœuds privilégiés pour l'observation. Le cas fréquent est celui où $\tilde{y}(t) = \tilde{x}(t)$, ce qui conduit à fixer simplement

$$\tilde{H} = I_n \quad (\text{matrice identité } [n, n])$$

étape 2

On impose au vecteur d'état de (2.2) de converger à l'horizon infini ($t \rightarrow \infty$) vers celui de (2.1). Pour cela, on écrit

$$[\dot{x}(\infty) = 0 \Rightarrow Fx + B\dot{U} = 0] \quad \text{et} \quad [\tilde{x}(\infty) = 0 \Rightarrow \tilde{F}\tilde{x} + \tilde{B}\dot{U} = 0]$$

et en introduisant la relation $\tilde{x}(t) = Rx(t)$ de l'étape 1, on obtient une relation entre $\tilde{F}$ et $\tilde{B}$ telle que

$$\tilde{B} = \tilde{F}R F^{-1} B$$

A ce stade, il ne reste que la matrice $\tilde{F}$ à déterminer ($\tilde{S} = S$), ce qui fait l'objet la dernière étape.

étape 3

On considère la mesure $\mathcal{M}$ de l'erreur (2.14) basée sur l'écart

$$\Delta z^j(t) = (\dot{\tilde{x}}^j) - (\tilde{F}\tilde{x}^j(t) + \tilde{B}\dot{U}_j(t))$$

⁹De ce point de vue, on peut considérer que la méthode d'Eitelberg est aussi une méthode d'agrégation

définie lorsque seule la $j^{\text{ème}}$ sollicitation n'est pas nulle. La minimisation de la mesure ainsi définie par rapport à l'inconnue $\tilde{F}$

$$\frac{\partial \mathcal{M}}{\partial \tilde{F}} = 0$$

avec l'utilisation de sollicitations échelons mène au résultat

$$\tilde{F} = RFVR^T (RV R^T)^{-1}$$

où V est la solution de l'équation de Lyapunov suivante

$$F V + V F^T = -B B^T$$

Opérant dans l'espace physique, l'approche d'Eitelberg peut permettre la prise en compte de quelques non linéarités telles que les échanges radiatifs.

2.3.4 Approche mixte réduction-identification

Cette approche due à Petit [55] a l'avantage de s'appliquer aussi bien à la réduction d'un modèle détaillé qu'à l'identification de modèles à partir de données expérimentales. Nous résumons ici la version relative à la réduction de modèle détaillé. La méthode peut être scindée en deux étapes.

étape 1

Cette étape concerne la détermination de n , $\tilde{F}$ et $\tilde{H}$. La mesure (2.14) est ici reformulée de façon échantillonnée dans le temps. Si r est le nombre de points de la discrétisation temporelle alors on définit la quantité

$$\mathcal{M}(n, \tilde{F}, \tilde{H}) = \sum_{i=1}^q \sum_{m=1}^r [\Delta y_i(t_m)]^2$$

avec

$$\Delta y_i(t_m) = y_i(t_m) - \tilde{y}_i(t_m)$$

La procédure mise en œuvre permet d'**identifier** $\tilde{F}$ et $\tilde{H}$ et utilise alternativement

- Une minimisation au sens des moindres carrés de $\mathcal{M}$ pour calculer $\tilde{H}$.
- Les gradients conjugués (non linéaire) pour le calcul de $\tilde{F}$.

Cette étape d'identification présente l'avantage de contourner la sélection d'un nombre fini de modes propres qui est délicate. Les sorties $y_i(t_m)$ sont issues du modèle de référence lorsque le système thermique évolue d'un état initial $x(0)$ vers un état asymptotique correspondant à $U = \text{constante}$. De ce fait, le choix de $x(0)$ semble avoir une influence sur les résultats. Il est préférable d'utiliser un état initial $x(0)$ suffisamment représentatif des divers sollicitations. L'ordre n est arrêté de manière à avoir $\mathcal{M}(n+1) \simeq \mathcal{M}(n)$.

étape 2

La matrice $\tilde{B}$ est la seule inconnue qui reste à déterminer. En considérant des sollicitations de type échelon, la minimisation de $\mathcal{M}$ telle qu'elle est donnée par (2.14) (reposant sur l'erreur de sortie) donne

$$\frac{\partial \mathcal{M}}{\partial \tilde{B}} = 0 \quad \Rightarrow \quad \tilde{B} = A_1^{-1} A_2 B$$

où A_1 et A_2 sont deux intégrales de Lyapounov dont la solution [Annexe B] peut être obtenue explicitement car F et $\tilde{F}$ sont diagonales

$$A_1 = \int_{\mathcal{D}_t} e^{\tilde{F}t} \tilde{H}^T f \tilde{H} e^{\tilde{F}t} dt \quad [n, n]$$

$$A_2 = \int_{\mathcal{D}_t} e^{\tilde{F}t} \tilde{H}^T f H e^{Ft} dt \quad [n, N]$$

où f est ici une matrice $[q, q]$ autorisant des pondérations sur les différentes sorties.

Remarque: Les matrices $\tilde{F}$ et $\tilde{H}$ sont obtenues pour des sollicitations $U = \text{constante}$ (étape 1). Il s'en suit que les régimes permanents associés aux sollicitations en rampes peuvent être mal représentés par le modèle réduit obtenu. Pour pallier ce risque, Petit [55] propose alors une possibilité d'adjonction de p modes correcteurs qui pratiquement revient à un ajustement des matrices $\tilde{F}$ et $\tilde{H}$. L'auteur développe ensuite des cas particuliers de la méthode en prenant divers types de matrices f pour la pondération sur les sorties.

2.3.5 Récapitulatif sur les méthodes de minimisation

Les méthodes de minimisation que nous venons de décrire font apparaître une démarche plus robuste que la troncature modale. Les points que l'on peut retenir des méthodes présentées peuvent se résumer comme suit:

En ce qui concerne l'agrégation [48], l'idée est de relier le modèle réduit au modèle détaillé de façon linéaire. Le modèle agrégé peut alors permettre une analyse riche du comportement thermique du système. Néanmoins, le choix des modes dominants s'opère jusqu'à présent de manière indépendante de l'étape de minimisation proprement dite.

La méthode d'Eitelberg [24] requiert le choix de quelques nœuds importants. Ceci fait appel à l'expérience du modélisateur et on doit refaire la réduction pour chaque nouveau problème. Opérant dans l'espace physique (la méthode n'est pas modale), cette méthode permet la prise en compte de quelques non-linéarités telles que les échanges radiatifs.

L'approche mixte réduction-identification [55], plus récente que les autres, a l'avantage d'être aussi bien adaptée à l'identification à partir de données expérimentales. Le modèle réduit obtenu n'a pas de lien explicite avec le modèle détaillé.

Au sujet de l'automatisation de la réduction, il y a deux solutions proposée. Michailenco [48] propose de fixer la taille "n" du modèle réduit de telle sorte que $\mathcal{C}_{nn} \gg \mathcal{C}_{(n+1)(n+1)}$. La seconde solution due à Petit [55], consiste à incrémenter sur l'ordre "n" du modèle réduit jusqu'à avoir $(\mathcal{M}(\text{ordre } n) \simeq \mathcal{M}(\text{ordre } n+1))$. Ce point est particulièrement difficile et nous y reviendrons au §3.5.

Chapitre 3

Réduction d'un modèle d'état modal par amalgame modal

3.1 Préliminaire

Nous développons dans ce chapitre une méthode complète pour la réduction de modèles d'état modaux décrivant les systèmes thermiques linéaires et réciproques. On garde ici la même notation que précédemment en ce qui concerne les modèles détaillé et réduit.

Modèle détaillé ``dimension  $N$ ``

$$\begin{aligned}\dot{x}(t) &= Fx(t) + B\dot{U}(t) \\ y(t) &= Hx(t) + SU(t)\end{aligned}\tag{3.1}$$

Modèle réduit ``dimension  $n$ ``

$$\begin{aligned}\dot{\tilde{x}}(t) &= \tilde{F}\tilde{x}(t) + \tilde{B}\dot{U}(t) \\ \tilde{y}(t) &= \tilde{H}\tilde{x}(t) + \tilde{S}U(t)\end{aligned}\tag{3.2}$$

Avant de présenter la méthode, nous allons d'abord préciser les choix adoptés et qui ne sont pas sans incidence sur la suite des développements.

a) Le vecteur des sorties $y(t)$ de l'équation (3.1) est maintenant le champ spatio-temporel de température. Au niveau formel, cela revient à prendre comme

matrice d'observation la base modale $\mathbf{P}$. Le modèle d'état modal supposé décrire finement l'évolution thermique de l'intégralité du système est donc:

$$\begin{aligned}\dot{x}(t) &= Fx(t) + B\dot{U}(t) \\ T(M, t) &= \mathbf{P}x(t) + SU(t)\end{aligned}\quad (3.3)$$

Avec $\mathbf{P} = [V_1(M), V_2(M) \cdots V_N(M)]$ est le tableau uniligne composé des N premières fonctions propres de l'opérateur $\mathcal{L}$ de la chaleur (équation (1.33)).

Le choix de s'intéresser à la totalité du champ de température est nécessaire dans beaucoup d'applications. Citons par exemple l'étude de la trempe des métaux, les déformations thermomécaniques telles la dilatation en régime dynamique, etc. La connaissance précise de $T(M, t)$ permet de déduire d'autres quantités telles que les flux de chaleur. De façon générale, toute sortie $y_i(t)$ peut être obtenue de manière approchée par la relation (1.4) en utilisant $\tilde{T}(M, t)$ donné par le modèle réduit

$$\tilde{y}_i(t) = j_i [\tilde{T}(M, t)] + g_i [U(M, t)] \quad (3.4)$$

Le modèle (3.3) repose sur une formulation continue aussi bien dans l'espace que dans le temps. Nous verrons plus loin que l'utilisation de cette formulation nous conduira à des résultats explicites. La démarche préconisée ici consiste en effet à aller aussi loin que possible formellement avant de songer aux méthodes numériques qui, selon nous, doivent relayer les méthodes analytiques.

b) Le modèle réduit recherché se met -conformément au système (3.3)- sous la forme

$$\begin{aligned}\dot{\tilde{x}}(t) &= \tilde{F}\tilde{x}(t) + \tilde{B}\dot{U}(t) \\ \tilde{T}(M, t) &= \tilde{\mathbf{P}}\tilde{x}(t) + \tilde{S}U(t)\end{aligned}\quad (3.5)$$

Avec $\tilde{\mathbf{P}} = [\tilde{V}_1(M), \tilde{V}_2(M) \cdots \tilde{V}_N(M)]$ est le tableau uniligne des pseudo-fonctions propres. Rappelons que les fonctions propres $V_i(M)$ vérifient une propriété fondamentale: elles sont orthonormées relativement à la capacité calorifique, c'est à dire

$$\forall i = 1, 2 \cdots N, \forall j = 1, 2 \cdots N \quad \langle V_i(M), C(M)V_j(M) \rangle = \delta_{ij}$$

De façon similaire, nous imposons à $\tilde{\mathbf{P}}$ de vérifier¹

$$\forall i = 1, 2 \cdots n, \forall j = 1, 2 \cdots n \quad (i \neq j) \quad \langle \tilde{V}_i(M), C(M)\tilde{V}_j(M) \rangle = 0 \quad (3.6)$$

La contrainte (3.6) va être très utile dans le développement formel de la méthode. En effet, beaucoup de propriétés de l'analyse modale reposent sur la propriété d'orthogonalité des fonctions propres. Il en est ainsi par exemple de la forme découplée de l'équation d'état. Les critères de troncature par l'énergie esquissés au §2.2 sont *explicites* et faciles à implémenter pour la même raison. L'introduction de la contrainte (3.6) conduit à préserver au niveau du modèle

¹Sachant que la normalisation des pseudo-modes est chose facile. Nous la ferons ultérieurement.

réduit (3.5) une structure formelle très voisine de celle du système (3.3). Cette contrainte permet également d'envisager ultérieurement l'application de la synthèse modale² [28] à partir de modèles réduits.

c) Dans le modèle réduit (3.5), nous choisissons

$$\tilde{S} = S = [S_1(M), S_2(M) \cdots S_p(M)]$$

où $S_i(M)$ $i = 1 \cdots p$ sont données par (1.12)

Ce choix s'apparente plutôt à une simplification qu'à une contrainte, et permet d'assurer la conservation des régimes statiques.

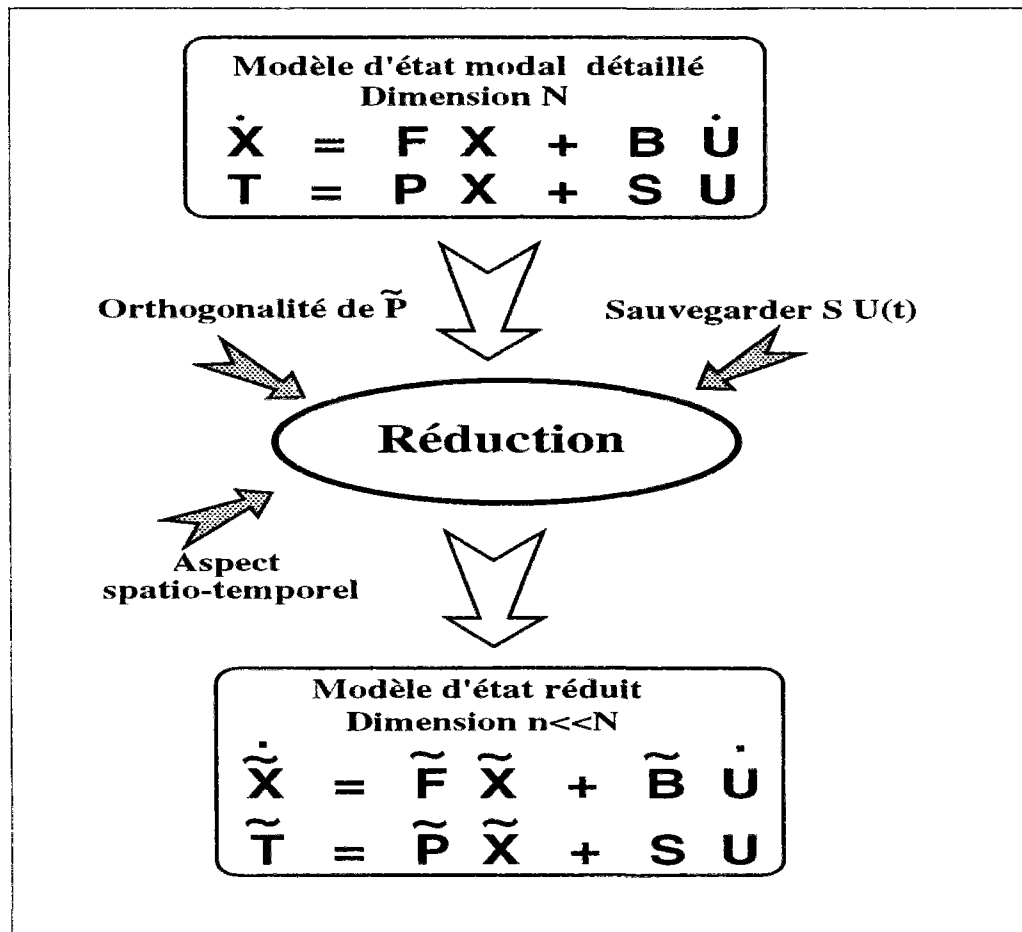


Figure 3.1: Principe de la réduction d'un modèle d'état modal

Le problème de réduction ainsi posé est illustré par la figure 3.1.

²La synthèse modale est une approche permettant la construction de modèles modaux de systèmes thermiques complexes de façon modulaire. L'idée de base est de rechercher les éléments propres d'un système thermique à partir de ceux de ses composants élémentaires. L'orthogonalité (au sens de (3.6)) des modes propres des composants est une condition nécessaire à la synthèse modale.

Notre objectif de réduction étant cerné, nous allons maintenant exposer la méthode d'amalgame modal [53]. On commence par définir au §3.2 une mesure adéquate pour l'erreur de réduction. On se référera à maintes reprises à cette mesure dans la méthode de réduction proposée. La description formelle de la méthode est donnée dans le §3.3. Pour le lecteur désireux d'implémenter la méthode proposée, nous avons résumé au §3.5 les résultats sous forme d'un algorithme pour la réduction automatique des modèles d'état modaux. Enfin, nous allons voir au §3.6 comment reconstituer la "structure modale" du modèle réduit obtenu.

3.2 Mesure de l'erreur de réduction

Nous retenons ici l'idée d'un critère quadratique pour quantifier l'erreur de réduction. Le choix de ce type de critère, largement préconisé dans divers domaines, se justifie principalement par les possibilités de calcul liées aux formes quadratiques.

Sur l'étendue spatiale $\mathcal{D}$ du système étudié et pour une plage temporelle $\mathcal{D}_t$ dans $[0, \infty[$, on définit naturellement l'erreur de réduction par

$$\forall M \in \mathcal{D} \text{ et } \forall t \in \mathcal{D}_t \quad e(M, t) = \tilde{T}(M, t) - T(M, t) \quad (3.7)$$

où $\tilde{T}(M, t)$ et $T(M, t)$ sont les températures au point M et à l'instant t , données respectivement par les modèles réduit et détaillé. D'après les équations des sorties de (3.3) et (3.5), cette erreur peut également s'écrire

$$e(M, t) = \sum_{k=1}^n \tilde{x}_k(t) \tilde{V}_k(M) - \sum_{i=1}^N x_i(t) V_i(M) \quad (3.8)$$

Pour chaque sollicitation $U_j(t)$, les autres étant nulles, nous pouvons calculer les états correspondants $\tilde{x}_{kj}(t)$ ($k = 1 \dots n$) et $x_{ij}(t)$ ($i = 1 \dots N$). Pour chacune des p sollicitations $U_j(t)$, l'erreur correspondante est notée $e^j(M, t)$. Ainsi, on définit le tableau d'écart suivant:

$$\mathcal{E}(M, t) = [e^1(M, t), e^2(M, t), \dots, e^p(M, t)] \quad [1, p] \quad (3.9)$$

telle que

$$e^j(M, t) = \sum_{k=1}^n \tilde{x}_{kj}(t) \tilde{V}_k(M) - \sum_{i=1}^N x_{ij}(t) V_i(M)$$

Nous construisons finalement une mesure de l'erreur de réduction par la double intégrale

$$\mathcal{M} = \int_{\mathcal{D}_t} \int_{\mathcal{D}} \text{tr} [\mathcal{E}^T(M, t) \cdot Q(M, t) \cdot \mathcal{E}(M, t)] dM dt \quad (3.10)$$

où $Q(M, t)$ est une fonction définie positive de l'espace et du temps.

Remarquons que cette mesure est une norme sur l'espace spatio-temporel des tableaux d'écart définis par (3.9). Par ailleurs le fait d'introduire la trace de la matrice $[\mathcal{E}^T Q \mathcal{E}]$ est équivalent à supposer nulles les interactions entre les différentes sollicitations. Les mesures de l'erreur introduites dans les différentes méthodes de réduction exposées au §2.3 peuvent être considérées comme des cas particuliers de (3.10).

3.3 Description de la méthode

3.3.1 Présentation générale

La méthode d'amalgame modal [53] repose sur une partition de l'espace des modes propres en quelques sous-espaces orthogonaux entre eux (au sens de (1.6)). Chaque sous-espace est engendré par un mode *principal* et quelques modes *mineurs*. Le caractère principal ou mineur d'un mode sera déterminé par un critère qu'on donnera au §3.3.4. A partir de chaque sous-espace, nous allons générer un pseudo-mode qui représente un "mélange" -ou encore **amalgame**- des modes propres du même sous-espace car on va réunir les caractères de plusieurs modes en un seul. Par cette opération d'amalgame, il s'agit d'apporter au mode principal de chaque sous-espace une information thermique récupérée sur les modes mineurs. Du point de vue mathématique, l'amalgame des fonctions propres d'un sous-espace se fait en les combinant *linéairement*. Ainsi, la procédure d'amalgame appliquée à tous les sous-espaces de façon indépendante aboutit à l'obtention de n pseudo-modes $\{\tilde{V}_i(M) \ i = 1 \dots n\}$ orthogonaux entre eux. La matrice $\tilde{P}$ du modèle réduit (3.5) est alors composée des modes amalgamés

$$\tilde{P} = [\tilde{V}_1(M), \tilde{V}_2(M) \dots \tilde{V}_n(M)]$$

Le nombre des modes amalgamés est égal à celui des modes principaux qui est aussi celui des sous-espaces orthogonaux.

A partir de la connaissance des modes principaux d'amalgame, on tronque les matrices F et B (les lignes et colonnes conservées sont celles qui correspondent à ces modes) pour obtenir les matrices réduites $\tilde{F}$ et $\tilde{B}$. Le vecteur d'état réduit $\tilde{x}(t)$ contient les états associés aux seuls modes principaux. L'équation d'état réduite de (3.5) devient ainsi entièrement définie.

La matrice statique $\tilde{S}$ est identique à celle du modèle détaillé (S), ce qui permet au modèle réduit de converger vers le même régime asymptotique que celui du modèle détaillé. L'équation des sorties de (3.5) est donc aussi définie.

On donne d'abord au §3.3.2 les notations utiles au suivi de la démarche formelle de la méthode proposée. Les deux étapes qui suivent sont consacrées à la méthode elle-même. On explique au §3.3.3 comment calculer les modes amalgamés en supposant une partition

quelconque de l'espace des modes propres. Cette étape utilise une minimisation classique de la mesure de l'erreur $\mathcal{M}$. L'incidence du choix de la partition de l'espace des modes sur la pertinence du modèle réduit obtenu est discutée au §3.3.4. Le nombre de partitions possibles étant grand (de l'ordre de C_N^n), on développe alors un algorithme original permettant de trouver l'une des meilleures partitions au sens de la mesure de l'erreur.

3.3.2 Notations

Afin de faciliter la lecture de ce qui suit, il est utile de préciser les notations utilisées.

Les sous-espaces orthogonaux d'amalgame -au nombre de n - sont notés H_k pour $k = 1, 2 \dots n$. Chaque sous-espace H_k est généré par un mode principal et quelques modes mineurs. Le nombre des modes associés à H_k est noté n_k ($n_k \geq 1$) de telle sorte que

$$\sum_{k=1}^n n_k = N$$

N étant le nombre total des modes propres du modèle détaillé.

Les n_k modes propres de H_k forment la base $\mathbf{P}_k$ définie par

$$\mathbf{P}_k = \mathbf{P} \pi_k$$

π_k étant une matrice $[N, n_k]$ de projection appropriée

Nous supposons -pour une raison de simplicité- que le mode principal de H_k est placé dans la première colonne de $\mathbf{P}_k$.

Nous définissons aussi une fonction de projection qui donne l'indice dans $\mathbf{P}$ de chaque mode propre de H_k . Selon cette fonction de projection, le mode propre m ($1 \leq m \leq n_k$) de H_k (de classement m dans $\mathbf{P}_k$) possède le classement $u^k(m)$ dans sa base d'origine $\mathbf{P}$. En particulier, le mode principal de H_k est donc $V_{u^k(1)}$. Cette fonction de projection est illustrée par la figure 3.2.

Le mode principal $V_{u^k(1)}$ de H_k est relié à $\mathbf{P}$ par

$$V_{u^k(1)} = \mathbf{P} \xi_k$$

ξ_k étant un vecteur $[N]$ de projection appropriée

Enfin, pour une fonction $f(t)$ quelconque du temps, nous désignerons par $[\Delta f(t)]_{t_2}^{t_1}$ l'écart $[f(t_1) - f(t_2)]$.

Avec ces notations, nous passons maintenant à la méthode proposée en expliquant d'abord les trois étapes qui la composent.

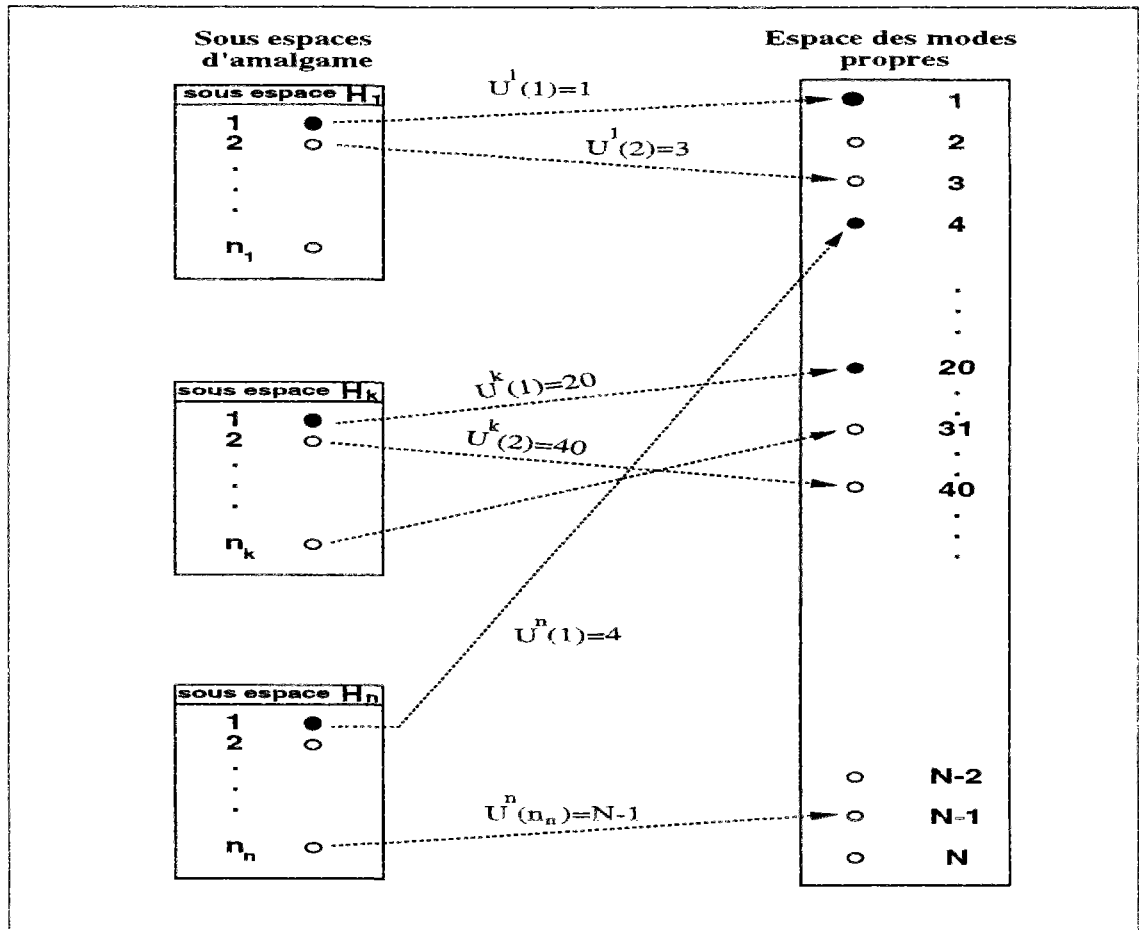


Figure 3.2: Partition de l'espace des modes propres. Les symboles $\bullet$ et $\circ$ désignent respectivement un mode principal et un mode mineur.

3.3.3 Détermination des modes amalgamés

Comme nous l'avons expliqué dans la présentation générale de la méthode, l'amalgame modal consiste à approximer les n_i modes propres de chaque sous-espace H_i par un seul pseudo-mode. L'approche adoptée consiste à combiner linéairement les modes propres de chaque sous-espace. Cette approche est suffisante pour aboutir à des modes amalgamés orthogonaux entre eux au sens de (1.6).

Le calcul des modes amalgamés suppose que nous disposons d'une partition de l'espace des modes propres en n sous-espaces. Cette partition est ici quelconque: les n sous-espaces $H_1, H_2, \dots, H_n$ sont de dimensions $n_1, n_2, \dots, n_n$ respectivement ($\sum_{k=1}^n n_k = N$). C'est au §3.3.4 que nous indiquerons la manière de partitionner l'espace d'état modal.

La combinaison linéaire des modes propres de H_i s'écrit

$$\tilde{V}_i(M) = \mathbf{P}_i \omega_i \quad (3.11)$$

Le problème revient à déterminer les vecteurs des coefficients de décomposition ω_i (dimension n_i) pour tous les sous-espaces H_i $i = 1 \dots n$. $\tilde{V}_i(M)$ est le mode amalgamé censé représenter à lui seul la dynamique des n_i modes de H_i .

Reprenons l'erreur de réduction élémentaire (3.7)

$$\forall M \in \mathcal{D} \text{ et } \forall t \in \mathcal{D}_t \quad e(M, t) = \tilde{T}(M, t) - T(M, t) \quad (3.12)$$

Puisque par hypothèse $\tilde{S} = S$, $e(M, t)$ peut s'écrire également

$$\begin{aligned} e(M, t) &= \sum_{k=1}^n \tilde{x}_k(t) \tilde{V}_k(M) - \sum_{i=1}^N x_i(t) V_i(M) \\ &= \sum_{k=1}^n \mathbf{P}_k \omega_k \tilde{x}(t) - \mathbf{P} x(t) \end{aligned} \quad (3.13)$$

Sachant que $\left\{ \begin{array}{l} \mathbf{P}_k = \mathbf{P} \pi_k \text{ d'après la notation de §3.3.2} \\ \mathbf{P} x(t) = \sum_{k=1}^n \mathbf{P} \left(\pi_k \pi_k^T \right) x(t) \end{array} \right.$

l'expression (3.13) devient

$$e(M, t) = \sum_{k=1}^n \mathbf{P} \pi_k \left(\omega_k \tilde{x}_k(t) - \pi_k^T x(t) \right) \quad (3.14)$$

En considérant successivement les p sollicitations $U_1 \ U_2 \dots U_p$, on forme les deux tableaux

$$\tilde{X}_k(t) = [\tilde{x}_{k1}(t), \tilde{x}_{k2}(t) \dots \tilde{x}_{kp}(t)] \text{ et } X(t) = [x^1(t), x^2(t) \dots x^p(t)]$$

Remarque: $\tilde{x}_k(t)$ est une fonction scalaire qui représente l'état du mode amalgamé $\tilde{V}_k(M)$, alors que $x(t)$ est un vecteur $[N]$ d'état. Aussi, il faut distinguer entre $\tilde{X}_k(t)$ qui est un vecteur $[p]$ uniligne et $X(t)$ qui est une matrice $[N, p]$.

L'expression (3.14) prend alors la forme matricielle

$$\begin{aligned} \mathcal{E}(M, t) &= \sum_{k=1}^n \mathbf{P} \pi_k \alpha_k(t) \quad [1, p] \\ \text{avec} \quad \alpha_k(t) &= \omega_k \tilde{X}_k(t) - \pi_k^T X(t) \end{aligned} \quad (3.15)$$

Aussi, la mesure de l'erreur $\mathcal{M}$ donnée par (3.10) se formule de la façon suivante

$$\mathcal{M} = \sum_{k=1}^n \sum_{r=1}^n \int_{\mathcal{D}_t} \text{tr} \left[\alpha_k^T(t) \pi_k^T \left(\int_{\mathcal{D}} \mathbf{P}^T Q \mathbf{P} dM \right) \pi_r \alpha_r(t) \right] dt \quad (3.16)$$

On voit apparaître dans (3.16) le produit scalaire (voir relation (1.6))

$$\langle \mathbf{P}^T, Q \mathbf{P} \rangle = \int_{\mathcal{D}} \mathbf{P}^T Q \mathbf{P} dM \quad (3.17)$$

On choisit alors la fonction Q égale à la capacité calorifique (fonction de la variable d'espace), ce qui permet d'avoir

$$\langle \mathbf{P}^T, C(M) \mathbf{P} \rangle = I_N \quad \text{matrice identité } [N, N] \quad (3.18)$$

De plus, il est facile de montrer que le produit³ $\pi_k^T \pi_r$ est égal à la matrice I_{n_k} (matrice identité $[n_k, n_k]$) si $r = k$, et nul sinon. D'où finalement

$$\mathcal{M} = \sum_{k=1}^n \int_{\mathcal{D}_t} \text{tr} \left[\alpha_k^T(t) \alpha_k(t) \right] dt \quad (3.19)$$

Il importe de remarquer que cette dernière relation est la somme de n termes positifs: la mesure de l'erreur $\mathcal{M}$ se décompose ainsi en termes de contribution relatifs aux n sous-espaces d'amalgame

$$\mathcal{M} = \sum_{k=1}^n \mathcal{M}_k \quad \text{avec } \mathcal{M}_k \geq 0 \quad (3.20)$$

Donc le minimum de $\mathcal{M}$ par rapport à l'ensemble $\{\omega_1, \omega_2 \dots \omega_n\}$ est atteint lorsque celui de chaque fonctionnelle $\mathcal{M}_k$ l'aura été par rapport à ω_k , ce qui veut dire

$$\min_{\omega_1 \dots \omega_n} \{\mathcal{M}\} = \sum_{k=1}^n \min_{\omega_k} \{\mathcal{M}_k\} \quad (3.21)$$

et le vecteur ω_k qui réalise le minimum de la fonctionnelle $\mathcal{M}_k$ satisfait à la condition

$$\frac{\partial}{\partial \omega_k} (\mathcal{M}_k(\omega_k)) = 0 \quad (3.22)$$

Moyennant quelques opérations de dérivation relativement simples [annexe D], on obtient de (3.22) le résultat général suivant

$$\omega_k = \pi_k^T \frac{\int_{\mathcal{D}_t} \mathbf{X}(t) \tilde{X}_k^T(t) dt}{\int_{\mathcal{D}_t} \tilde{X}_k(t) \tilde{X}_k^T(t) dt} \quad \text{dimension } [n_k] \quad (3.23)$$

³La matrice de projection π_k est définie à partir de ses éléments $(\pi_k)_{ij} = \delta_{i, u^k(j)}$ (symbole de Kronecker).

et les composantes $\omega_k(m)$ $m = 1 \dots n_k$ sont données par

$$\omega_k(m) = \frac{\sum_{j=1}^p \int_{\mathcal{D}_t} x_{u^k(m)_j}(t) \tilde{x}_{kj}(t) dt}{\sum_{j=1}^p \int_{\mathcal{D}_t} \tilde{x}_{kj}^2(t) dt} \quad (3.24)$$

Le vecteur ω_k ainsi calculé donne pour chaque sous-espace d'amalgame H_k les coefficients de décomposition de la combinaison linéaire (3.11) réalisant le minimum de la mesure $\mathcal{M}$ de l'erreur de réduction.

Remarquons que pour établir l'expression des coefficients de décomposition $\omega_k(i)$, il n'a pas été nécessaire d'indiquer la forme précise des sollicitations. De même, il n'a pas été nécessaire de fixer une partition définitive.

Explicitons maintenant (3.23) pour les sollicitations tests

- Echelon de Heavyside (avec $s=0$)
- Impulsion ($s=1$)

pour lesquelles⁴ on a

$$\begin{aligned} X(t) &= F^s e^{Ft} B \\ \tilde{X}(t) &= \xi_k^T F^s e^{Ft} B \end{aligned}$$

où ξ_k est un tableau $[N]$ de projection défini dans §3.3.2. Nous arrivons alors à la solution explicite

$$\begin{aligned} \omega_k &= \frac{\pi_k^T L' \xi_k}{\xi_k^T L' \xi_k} \\ \text{avec } L' &= F^s L F^s \end{aligned} \quad (3.25)$$

L est la matrice $[N, N]$ solution de l'équation de Lyapunov [annexe B] suivante

$$\begin{aligned} L F + F L &= -\mathbb{B} \\ \mathbb{B} &= \left[\Delta e^{Ft} B B^T e^{Ft} \right]_{t_1}^{t_2} \end{aligned} \quad (3.26)$$

La matrice F étant diagonale, il est possible de donner la forme explicite des composantes $\omega_k(i)$ $i = 1 \dots n_k$ de (3.25)

$$\begin{aligned} \omega_k(i) &= -\frac{2\lambda_{z_1}^{1-s}\lambda_{z_i}^s}{\lambda_{z_1} + \lambda_{z_i}} \frac{\sum_{l=1}^p B_{z_1 l} B_{z_i l}}{\sum_{l=1}^p B_{z_1 l}^2} \frac{\left[\Delta e^{(\lambda_{z_1} + \lambda_{z_i})t} \right]_{t_2}^{t_1}}{\left[\Delta e^{(2\lambda_{z_1})t} \right]_{t_2}^{t_1}} \\ &\text{avec les indices } z_i = u^k(1) \text{ et } z_i = u^k(i) \end{aligned} \quad (3.27)$$

⁴Pour simplifier les calculs, nous rajoutons la condition initiale d'un champ de température identiquement nul

Remarque: Nous avons convenu de placer dans la première colonne de $\tilde{\mathbf{P}}$ le mode principal de H_k . Le coefficient de décomposition de $\tilde{V}_k(M)$ sur le mode principal de H_k est donc $\omega_k(1)$. Il est alors facile de vérifier à partir de (3.27) que

$$\omega_k(1) = 1$$

de telle sorte que

$$\tilde{V}_k(M) = V_{u^k(1)} + \sum_{r=2}^{n_k} \omega_k(r) V_{u^k(r)}(M)$$

Dans la pratique, il est possible d'utiliser directement (3.27) ce qui est très avantageux aussi bien en temps calcul qu'en précision. Pour d'autres types de sollicitations il est nécessaire d'intégrer l'expression générale (3.23) lorsque cela est possible⁵, ou bien de recourir à une méthode d'intégration numérique. Pour des sollicitations mesurées, on peut procéder de deux façons: la première consiste à échantillonner (3.23) dans le temps exactement comme le sont les entrées. La seconde façon de procéder [1] repose sur l'approximation de chaque sollicitation par un polynôme (de degré fini) de la variable temporelle puis d'intégrer (3.23). Mais l'hypothèse de la linéarité des échanges thermiques autorise [43] de travailler avec des sollicitations tests (échelon et/ou impulsion) qui peuvent reconstituer la plupart des signaux grâce au théorème de convolution.

3.3.4 Partition de l'espace des modes propres

La démarche que nous venons d'exposer pour le calcul des modes amalgamés est à l'évidence optimale au regard de la mesure $\mathcal{M}$. Les modes amalgamés sont en effet donnés par (3.11) et (3.23) qui minimisent explicitement $\mathcal{M}$ donnée par (3.10). Mais ce caractère optimal de la dernière étape ne doit pas être confondu avec celui du modèle réduit obtenu. En effet, la précision du modèle réduit obtenu dépend aussi de la partition appliquée à l'espace des modes propres. Autrement dit, si l'on utilise une partition quelconque, les modes amalgamés obtenus sont optimaux "relativement" à cette partition mais le modèle réduit reste "sous-optimal".

Nous allons éclaircir ce point en deux étapes. La première est théorique et nous permettra de proposer un algorithme pour fixer la partition. La seconde étape est une expérience numérique pour

- mettre en évidence l'importance du choix de la partition d'une part
- montrer la fiabilité de l'algorithme proposé dans la première étape.

Approche théorique:

Reprenons notre mesure de l'erreur de réduction (3.10)

$$\mathcal{M} = \int_{\mathcal{D}_t} \int_{\mathcal{D}} \text{tr} \left[\mathcal{E}^T(M, t) \cdot Q(M, t) \cdot \mathcal{E}(M, t) \right] dM dt \quad (3.28)$$

⁵Ici, des environnements de calcul formel tels MAPPLE, MATHEMATICA et MACCYMA peuvent être très utiles.

Nous avons vu en §3.3.3 que cette mesure se décompose en n termes positifs, chacun se référant à l'un des n sous-espaces d'amalgame. Cette décomposition est traduite par la relation

$$\mathcal{M} = \sum_{k=1}^n \mathcal{M}_k \quad \text{avec } \mathcal{M}_k \geq 0 \quad (3.29)$$

avec:

$$\left\{ \begin{array}{l} \mathcal{M}_k = \int_{\mathcal{D}_t} \text{tr} \left[\alpha_k^T(t) \alpha_k(t) \right] dt \\ \alpha_k(t) = \omega_k \tilde{X}_k(t) - \pi_k^T X(t) \quad [n_k, p] \end{array} \right.$$

De façon analogue, chaque fonctionnelle $\mathcal{M}_k$ se décompose en n_k termes *positifs*: chaque terme est relatif à un des modes propres de H_k [annexe F]

$$\mathcal{M}_k = \sum_{m=1}^{n_k} \mathcal{M}_{km} \quad (\mathcal{M}_{km} \geq 0) \quad (3.30)$$

Aussi, la mesure de l'erreur de réduction se décompose donc entièrement sur l'ensemble des modes propres car

$$\mathcal{M} = \sum_{k=1}^n \sum_{m=1}^{n_k} \mathcal{M}_{km}$$

Lorsque, pour une partition donnée, le **minimum** $\min_{\omega_k} \{\mathcal{M}_k\}$ **est réalisé** (et seulement dans ce cas, mais c'est celui qui nous intéresse), les termes $\mathcal{M}_{km}$ deviennent [annexe F]

$$\mathcal{M}_{km} = \left[\sum_{j=1}^p \int_{\mathcal{D}_t} x_{u^k(m)j}^2 dt \right] - \left[\omega_k^2(m) \sum_{j=1}^p \int_{\mathcal{D}_t} (\tilde{x}_{kj}(t))^2 dt \right] \quad (3.31)$$

Il est important de saisir la signification de l'égalité (3.31) pour comprendre la manière de choisir les modes principaux ainsi que la procédure de répartition des modes mineurs qui mènent au modèle réduit amalgamé. Commençons par faire trois observations:

Observation 1: $\mathcal{M}_{km} \geq 0$ se présente sous forme de deux termes. Le premier

$$\mathcal{M}_{km}^{\{1\}} = \left[\sum_{j=1}^p \int_{\mathcal{D}_t} x_{u^k(m)j}^2 dt \right]$$

est *positif* et ne dépend que du mode $u^k(m)$. Autrement dit, ce terme est intrinsèque au sens où il est indépendant de la partition. Le second terme

$$\mathcal{M}_{km}^{\{2\}} = - \left[\omega_k^2(m) \sum_{j=1}^p \int_{\mathcal{D}_t} (\tilde{x}_{kj}(t))^2 dt \right]$$

est *négatif* ($|\mathcal{M}_{km}^{\{2\}}| \leq |\mathcal{M}_{km}^{\{1\}}|$) et fait intervenir aussi bien le mode $u^k(m)$ que le mode principal de H_k (via $\omega_k(m)$ et $\tilde{x}_{kj}(t)$). Ce terme négatif diminue la valeur

de $\mathcal{M}_{km}$ et il faut le porter à son maximum (en valeur absolue) en plaçant le mode $u^k(m)$ dans le sous-espace adéquat. Le maximum de $|\mathcal{M}_{km}^{\{2\}}|$ correspond au minimum de $\mathcal{M}_{km}$ car $\mathcal{M}_{km}^{\{1\}}$ est intrinsèque.

Observation 2: Si le champ initial ($t=0$) de température est permanent, alors on a $\tilde{x}_{kj}(t) = x_{u^k(1)j}$. Sachant que $\omega_k(1) = 1$, il est facile de vérifier à partir de (3.31) que

$$\mathcal{M}_{k1} = 0 \quad \forall k = 1 \dots n$$

ce qui veut dire que la contribution à $\mathcal{M}_k$ du mode principal de chaque sous-espace d'amalgame est nulle. La mesure $\mathcal{M}$ est donc une somme des contributions des modes mineurs seulement. En particulier, si $n = N$ ce qui revient à ne pas faire de réduction, on a bien $\mathcal{M} = 0$.

Observation 3: Si l'on évalue la mesure $\mathcal{M}$ pour un modèle tronqué (conformément à §2.2) en considérant les modes principaux comme dominants, on obtient

$$\mathcal{M}' = \sum_{i \notin \{I_D\}} \sum_{j=1}^p \int_{\mathcal{D}_t} x_{ij}^2 dt \quad (3.32)$$

où $\{I_D\}$ désigne l'ensemble des indices des modes dominants. Les modes $i \notin \{I_D\}$ sont les modes mineurs et l'on peut écrire avec les notations de la méthode d'amalgame

$$\mathcal{M}' = \sum_{k=1}^n \mathcal{M}'_k$$

avec $\mathcal{M}'_k = \sum_{m=2}^{n_k} \mathcal{M}_{km}^{\{1\}}$

Ce résultat est facile à vérifier à partir de (3.31). En effet, négliger complètement les modes mineurs revient à considérer au niveau de (3.31)

$$\omega_k(m) = 0 \quad \forall k = 1 \dots n \quad \forall m \geq 2$$

autrement dit

$$\tilde{V}_k(M) = V_{u^k(1)}(M)$$

De ces trois observations, nous déduisons les règles suivantes

Pour le modèle tronqué reposant sur les modes principaux (modes principaux=dominants), le mode mineur $u^k(m)$ $m \geq 2$ qui est complètement négligé, contribue à $\mathcal{M}$ de la quantité (positive) $\mathcal{M}_{km}^{\{1\}}$.

Pour le modèle réduit par la méthode d'amalgame, le même mode mineur qui est ici associé avec un mode principal (en l'affectant dans un sous-espace H_k) contribue à $\mathcal{M}$ de la quantité (positive) $\mathcal{M}_{km}^{\{1\}} + \mathcal{M}_{km}^{\{2\}}$ (avec $\mathcal{M}_{km}^{\{2\}} < 0$).

Il devient maintenant clair que le fait d'affecter un mode mineur dans un sous-espace d'amalgame (au lieu de le négliger complètement comme on le fait dans les méthodes de troncature) a pour conséquence de diminuer sa contribution à la mesure $\mathcal{M}$. C'est en ce sens qu'un mode mineur peut apporter une information thermique à un mode principal.

En vertu de l'observation 2, il faut *annuler* les plus grandes contributions $\mathcal{M}_{km}$ en considérant les modes correspondants comme principaux. Pour cela, on procède comme suit:

On commence par considérer tous les modes propres comme équivalents du point de vue de la dominance. Le premier mode principal (qu'on va associer à H_1) est le mode correspondant au maximum de la quantité

$$\left[\sum_{j=1}^p \int_{\mathcal{D}_t} x_{ij}^2 dt \right] \quad i = 1 \cdots N$$

Parmi les $N - 1$ modes mineurs restants, on prendra comme deuxième mode principal celui qui engendrera la plus grande contribution à $\mathcal{M}_1$. Ce mode correspond donc au maximum de $\mathcal{M}_{1m}$ évaluée pour tous les modes m pris parmi les $N - 1$ modes encore non-principaux à ce stade de la procédure. Pour sélectionner le troisième mode principal, il faut opérer deux étapes:

- (1) Pour chaque mode m encore non-principal, il faut calculer $\mathcal{M}_{1m}$ et $\mathcal{M}_{2m}$. on repérera alors la quantité $\min_m = \min\{\mathcal{M}_{1m}, \mathcal{M}_{2m}\}$
- (2) Le troisième mode principal est alors celui qui correspond au $\max\{\min_m \mid m \text{ balayant les indices des modes non-principaux}\}$

Les autres modes principaux (4,5...) sont déterminés de façon similaire au troisième: possédant le i premiers modes principaux associés à $H_1 \ H_2 \ \cdots \ H_i$, on effectue les deux étapes suivantes:

- (1) Pour chaque mode m encore non-principal, il faut calculer $\mathcal{M}_{1m} \cdots \mathcal{M}_{im}$. on repérera alors la quantité $\min_m = \min\{\mathcal{M}_{1m} \cdots \mathcal{M}_{im}\}$.
- (2) Le "(i+1)^{ème}" mode principal est alors celui qui correspond au $\max\{\min_m \mid m \text{ balayant les indices des modes non-principaux}\}$

Un fois les modes principaux sélectionnés, il reste à répartir les modes mineurs dans les différents sous-espaces. Cette répartition est facile à faire et **de façon optimale**:

Afin de minimiser la mesure $\mathcal{M}$ de l'erreur de réduction, il est nécessaire d'affecter chaque mode mineur dans le sous-espace d'amalgame où sa contribution à l'erreur de réduction est la plus petite. Il faut donc calculer pour chaque mode mineur i les quantités $\mathcal{M}_{ki} \ k = 1 \cdots n$. Le minimum de ces quantités permet de désigner le meilleur sous-espace pour l'affectation du mode mineur i .

Exprimons maintenant $\mathcal{M}_{km}$ dans le cas particulier des sollicitations échelon et impulsion. Sachant que

$$x_{ij}(t) = B_{ij} \lambda_i^s e^{\lambda_i t}$$

avec $s = 0$ pour l'échelon et $s = 1$ pour l'impulsion

On peut alors aisément vérifier d'après (3.27) et (3.31) le résultat

$$\mathcal{M}_{km} = \mathcal{M}_{km}^{\{1\}} + \mathcal{M}_{km}^{\{2\}}$$

avec

$$\left\{ \begin{array}{l} \mathcal{M}_{km}^{\{1\}} = -\frac{\lambda_m^{2s-1}}{2} (\sum_{l=1}^p B_{ml}^2) \left[\Delta e^{2\lambda_m t} \right]_{t_2}^{t_1} \\ \mathcal{M}_{km}^{\{2\}} = \frac{2\lambda_z \lambda_m^{2s}}{(\lambda_z + \lambda_m)^2} \frac{(\sum_{l=1}^p B_{zl} B_{ml})^2}{\sum_{l=1}^p B_{zl}^2} \frac{([\Delta e^{(\lambda_z + \lambda_m)t}]_{t_1}^{t_2})^2}{[\Delta e^{2\lambda_z t}]_{t_1}^{t_2}} \quad k = 1 \dots n \\ \text{où } z = u^k(1) \text{ est l'indice du mode principal de } H_k \end{array} \right. \quad (3.33)$$

On donne dans la figure 3.3 les principales étapes de la méthode de réduction proposée. On peut y noter les trois niveaux où la mesure $\mathcal{M}$ intervient: sélection des modes principaux, répartition des modes mineurs et amalgame modal dans les sous-espaces.

Résumons ce qui vient d'être dit:

- La répartition des modes mineurs dans les différents sous-espaces est optimale au sens de $\mathcal{M}$, pour une partition donnée.
- Dans chaque sous-espace H_k , la combinaison du mode principal avec les modes mineurs qui l'entourent (pour obtenir un mode amalgamé) est également optimale au sens de $\mathcal{M}$.
- Pour un ordre fixé à l'avance du modèle réduit à construire, on peut imaginer que le choix séquentiel des modes dominants (qui minimise pas à pas l'erreur) aboutit à une solution proche de l'optimum. Cette conjecture semble raisonnable et nous allons le vérifier sur un exemple numérique pris parmi les cas que nous avons testés.

Illustration par l'expérience:

On dispose d'un modèle de dimension $N = 30$. Le système thermique modélisé est un bâtiment bizonne soumis à plusieurs sollicitations (température d'air extérieur, flux solaire, chauffage interne ...). On va chercher à réduire ce modèle à l'ordre $n = 5$. Il faut alors fixer d'abord la partition de l'espace des modes, c'est à dire former cinq sous-espaces orthogonaux.

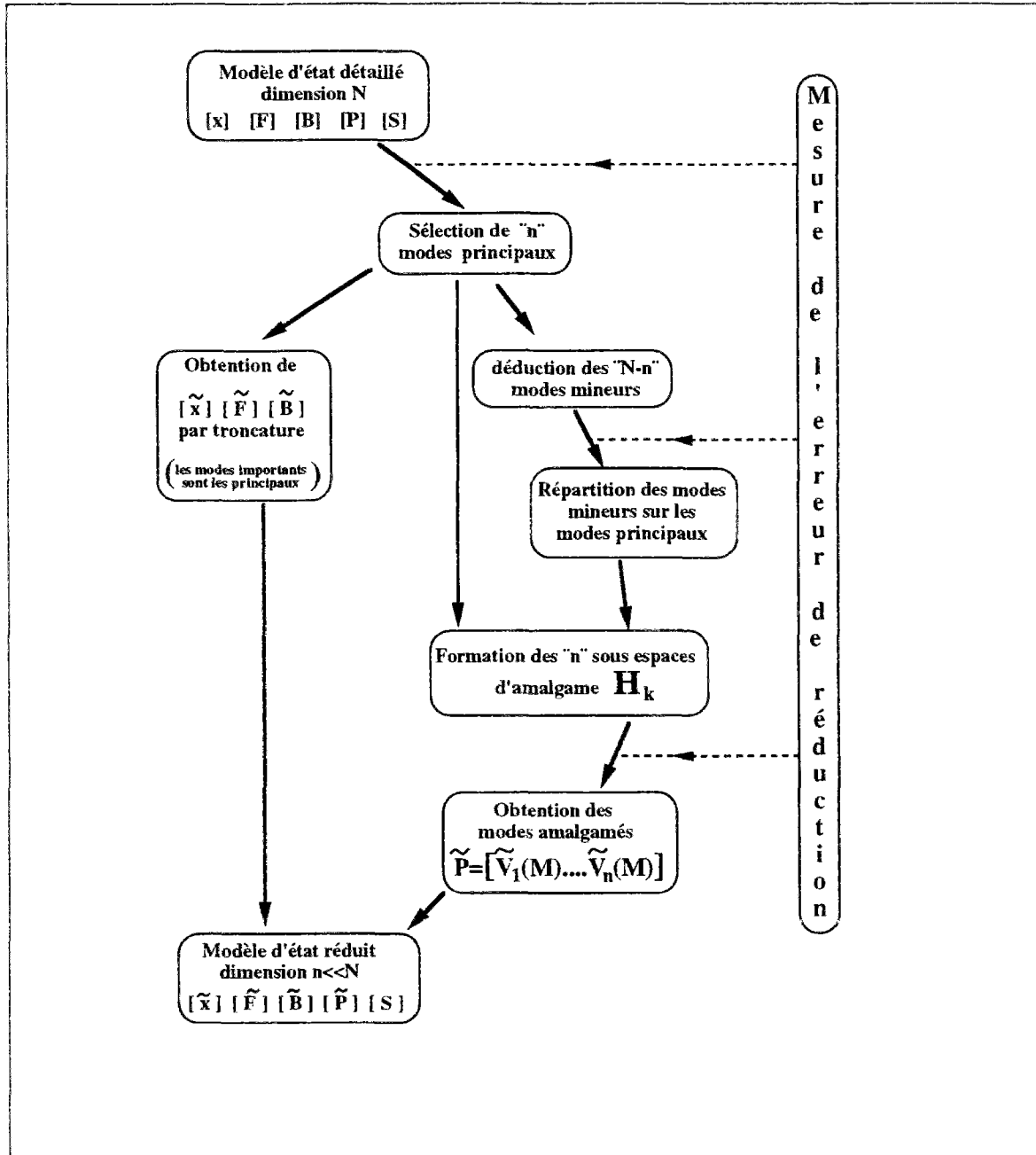


Figure 3.3: Méthode d'amalgame modal: principales étapes

En réalité, le problème du choix de la partition revient à celui des modes principaux seulement. En effet, la répartition des modes mineurs est toujours optimale si l'on suit la règle donnée ci-dessus. Dans cet exemple, on peut choisir les cinq modes principaux de C_{30}^5 manières ($C_{30}^5 = 142506$).

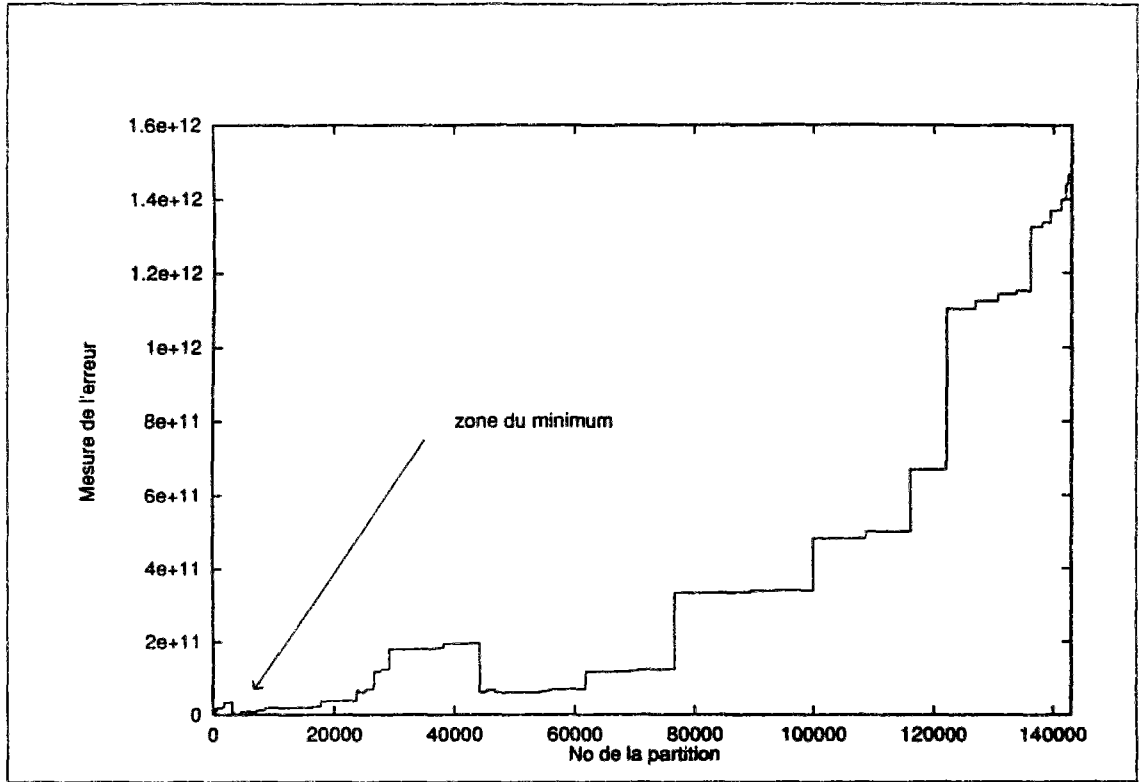


Figure 3.4: Influence de la partition des modes utilisée sur la réduction.

Nous avons formé l'ensemble des 142506 combinaisons possibles pour les modes principaux. Pour chaque combinaison, nous avons réparti de façon optimale - selon la procédure indiquée ci-dessus - les modes mineurs autour des modes principaux. On obtient ainsi les 142506 partitions possibles pour l'espace des modes propres. Enfin, nous avons évalué la valeur de $\mathcal{M}$ pour chacune de ces partitions. Le calcul de $\mathcal{M}$ est rapide puisqu'il se fait à partir de relations explicites:

$$\mathcal{M} = \sum_{k=1}^5 \sum_{m=2}^{n_k} \mathcal{M}_{km}$$

où $\mathcal{M}_{km}$ est donnée directement par (3.33) dans le cas des sollicitations échelons que nous considérons ici.

L'évolution de $\mathcal{M}$ lorsqu'on décrit les 142506 partitions dans un certain ordre est donnée dans la figure 3.4. La figure 3.5 est un grossissement de la partie où on peut trouver

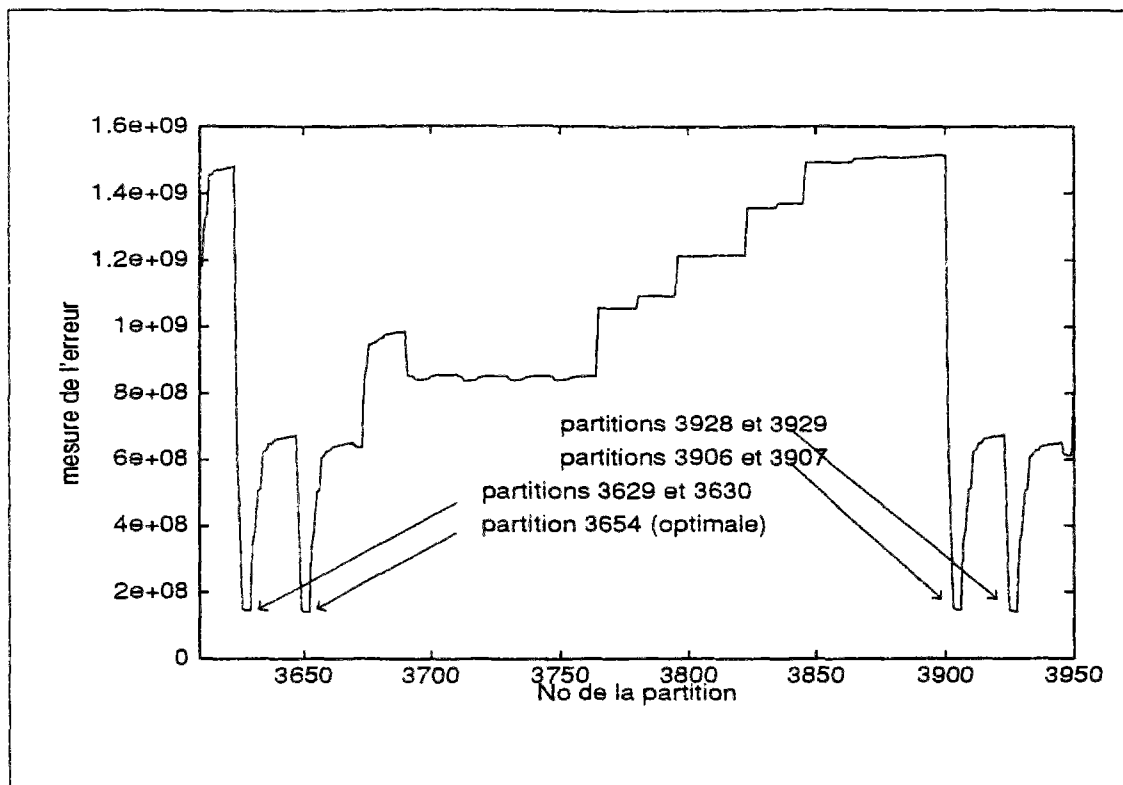


Figure 3.5: Grossissement autour de la **zone de minimum**.

le minimum de $\mathcal{M}$ (autour de la valeur 3700 de l'axe de abscisses). D'autre part, nous avons appliqué l'algorithme proposé (voir approche théorique ci-dessus) pour la sélection des modes principaux: les modes principaux sont (1,3,5,7 et 11) et correspondent à la partition N°3629.

L'examen des figures 3.4 et 3.5 nécessite quelques commentaires:

- Il existe des "paliers" horizontaux sur la figure 3.4. Ceci veut dire qu'il existe des combinaisons de modes principaux équivalentes du point de vue numérique. De façon générale, ces paliers correspondent au remplacement d'un mode faible (resp. dominant) par un autre mode aussi faible (resp. dominant) dans la combinaison des modes principaux. Du point de vue théorique, il est difficile d'affirmer l'existence de partitions strictement identiques au sens de $\mathcal{M}$.
- Il existe des "sauts verticaux" de $\mathcal{M}$. Leur interprétation est simple: lorsqu'un mode mineur est remplacé par un mode dominant, il se produit une chute verticale de la mesure $\mathcal{M}$. Inversement, lorsqu'un mode dominant est remplacé par un mode mineur, on obtient une montée verticale de $\mathcal{M}$. Il est d'ailleurs possible d'identifier, par une analyse fine de la figure 3.4, les numéros des modes dominants: ce sont ceux qui causent les "sauts verticaux importants" dont nous venons de parler.
- La figure 3.5 montre qu'il existe - numériquement - un minimum absolu: C'est

la combinaison 3654 désignée sur la figure 3.5. Mais il y a aussi quelques autres combinaisons (3629, 3654, 3907, 3929, etc) suffisamment voisines, ce qui rend difficile d'affirmer l'unicité de la solution minimale. Là aussi, du point de vue théorique, on n'est pas sûre de l'unicité de la solution optimale. On peut seulement dire qu'il y a un nombre fini (moins d'une dizaine sur cet exemple) de "meilleures combinaisons".

L'application de l'algorithme proposé dans l'approche théorique ci-dessus a retenu une des meilleures partitions. De toutes façon, la différence (au sens de $\mathcal{M}$) entre ces "meilleures combinaisons" est tellement faible que l'obtention de l'une parmi elles est suffisante.

Au niveau du temps calcul, avec une machine SUN SPARC Station 10:

- l'obtention de la combinaison 3629 par notre algorithme a mis environ 5 secondes (temps CPU).
- l'obtention de l'ensemble des 142506 partitions possibles a demandé plus de 2300 secondes (temps CPU).

Les tests effectués sur d'autres modèles (structures planes composites, pont thermique) ont donné des résultats similaires à ceux que nous venons de voir. L'algorithme proposé pour sélectionner une partition aboutit à l'une des meilleures partitions possibles et ne demande pas de temps de calcul prolongé.

Dans chaque sous-espace d'amalgame, le choix d'affecter la constante de temps principale au mode amalgamé est en réalité une contrainte. On peut libérer cette contrainte en ajustant la constante de temps de chaque sous-espace d'amalgame. Cet ajustement a pour but de diminuer encore la mesure de l'erreur $\mathcal{M}$. C'est l'objet de la section suivante.

3.4 Ajustement temporel du modèle réduit

Une fois les modes amalgamés calculés, il est possible d'adjoindre une étape supplémentaire qui consiste à ajuster les constantes de temps du modèle amalgamé. Mais cette étape n'est pas vraiment une nécessité car il n'y a pas une grande amélioration possible sur le modèle amalgamé. En revanche, l'ajustement temporel peut être très utile au cas où les modes principaux sont "mal choisis".

Le principe de cette étape est simple: une fois les modes amalgamés calculés, on suppose qu'on possède tous les paramètres du modèle réduit (3.2) **sauf la matrice $\tilde{F}$** . Il faut alors opérer une minimisation de la mesure de l'erreur (3.10). On montre [annexe §F] que cette minimisation se ramène à la résolution d'une équation transcendante pour chaque sous-espace d'amalgame. Pour des sollicitations en échelon⁶ et pour le sous-espace H_k ,

⁶Pour les sollicitations en impulsion, l'expression analytique possède une forme complexe. Nous l'avons volontairement omise sachant que le lecteur peut retrouver le résultat en suivant la démarche donnée pour l'échelon [annexe §F].

Cette équation transcendante est donnée par

$$\sum_{j=1}^{n_k} \dot{\mathcal{M}}_{kj} = 0 \quad (3.34)$$

avec

$$\begin{aligned} \dot{\mathcal{M}}_{kj} = -a_j \frac{\left[\Delta \left((\tilde{\lambda} + \lambda_{u^k(j)})t - 1 \right) e^{(\tilde{\lambda} + \lambda_{u^k(j)})t} \right]_{t_2}^{t_1}}{(\tilde{\lambda} + \lambda_{u^k(j)})^2} \\ + b_j \frac{\left[\Delta \left(2\tilde{\lambda}t - 1 \right) t e^{2\tilde{\lambda}t} \right]_{t_2}^{t_1}}{\tilde{\lambda}^2} \end{aligned} \quad (3.35)$$

$\tilde{\lambda}$ est, au signe près, l'inverse de la constante de temps recherchée

$$\begin{aligned} a_j &= 2\omega_k(j) \sum_{l=1}^p \tilde{B}_{kl} B_{u^k(j)l} \\ b_j &= \frac{\omega_k^2(j)}{2} \sum_{l=1}^p \tilde{B}_{kl}^2 \end{aligned}$$

Reprenons l'exemple qui a servi dans la figure 3.4: le modèle modale de dimension $N = 30$ du bâtiment bizoné. Nous reprenons strictement la même démarche qui a permis d'avoir la figure 3.4 pour laquelle nous rajoutons maintenant un ajustement temporel des constantes de temps.

La figure 3.6 montre la diminution de $\mathcal{M}$ liée à l'ajustement temporel. On constate que cette diminution est importante pour les partitions inadéquates. En revanche, dans la zone du minimum, c'est à dire lorsque la partition est optimale (ou proche), la diminution de $\mathcal{M}$ reste faible. La figure 3.7 montre un élargissement autour de la zone du minimum et on constate bien le rôle négligeable de l'ajustement temporel dans cette zone.

3.5 Automatisation de la réduction

3.5.1 Normalisation de l'erreur

La possibilité d'obtenir des modèles réduits par un algorithme automatique est très avantageuse. Elle doit permettre de réduire un modèle détaillé en imposant deux contraintes:

- la première concerne la dimension " n " du modèle réduit. On peut imposer au modèle réduit une dimension **maximale** n_{max} telle que: $n \leq n_{max}$.

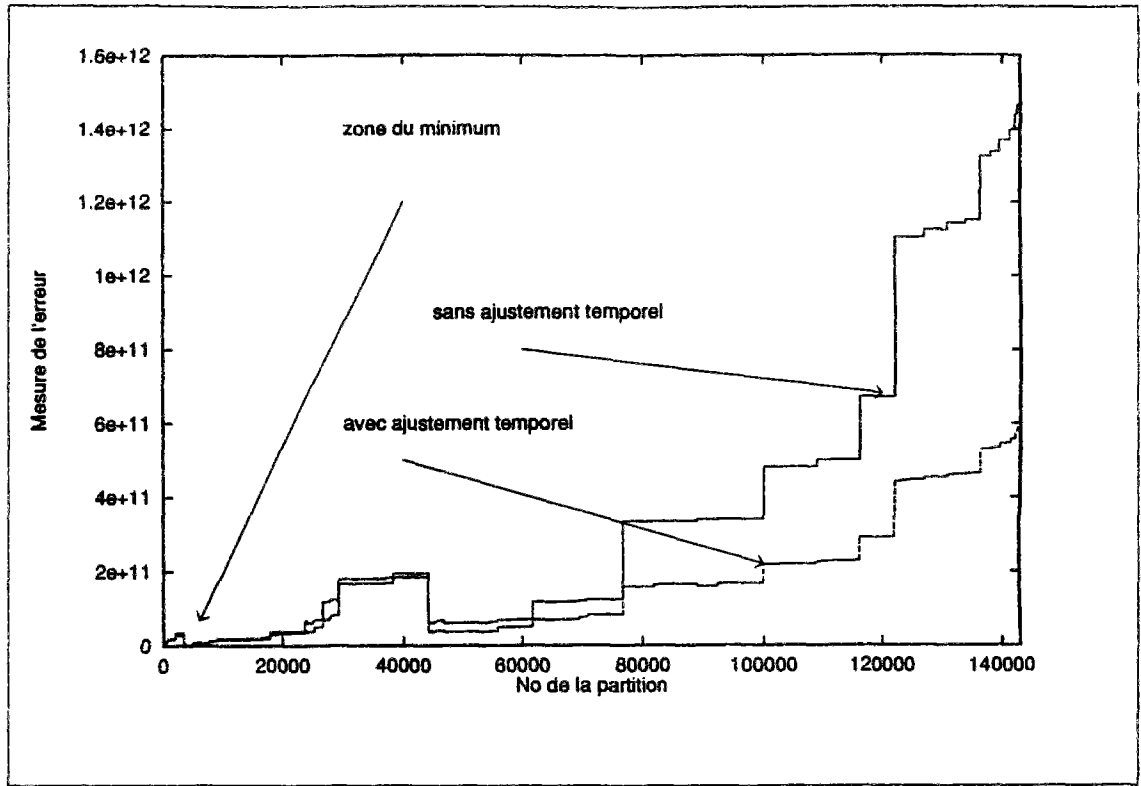


Figure 3.6: Illustration du rôle de l'ajustement temporel.

► La seconde contrainte, plus intéressante mais plus compliquée à réaliser, concerne la mesure de l'erreur de réduction. Il est souhaitable de pouvoir rechercher un modèle réduit vérifiant $\mathcal{M} \leq \mathcal{M}_{max}$, où $\mathcal{M}_{max}$ est l'erreur de réduction *maximale* tolérée.

Remarque: Il peut arriver que les deux contraintes soient impossibles à réaliser simultanément. Autrement dit, on peut obtenir $\mathcal{M} > \mathcal{M}_{max}$ pour $n = n_{max}$. Dans ce cas, il est nécessaire de faire un compromis, soit en augmentant n_{max} ou $\mathcal{M}_{max}$ (ou les deux à la fois). Si l'on impose seulement l'une des deux contraintes, alors il n'y a pas de compromis à faire et la contrainte doit être respectée par le modèle réduit.

Reprenons la mesure de l'erreur de réduction retenue dans la méthode d'amalgame en prenant $Q(M, t) = C(M)$.

$$\mathcal{M} = \int_{\mathcal{D}_t} \int_{\mathcal{D}} \text{tr} [\mathcal{E}^T(M, t) \cdot C(M) \cdot \mathcal{E}(M, t)] dM dt \quad (3.36)$$

qui s'écrit aussi

$$\mathcal{M} = \sum_{j=1}^p \int_{\mathcal{D}_t} \int_{\mathcal{D}} (e^j(M, t))^2 C(M) dM dt$$

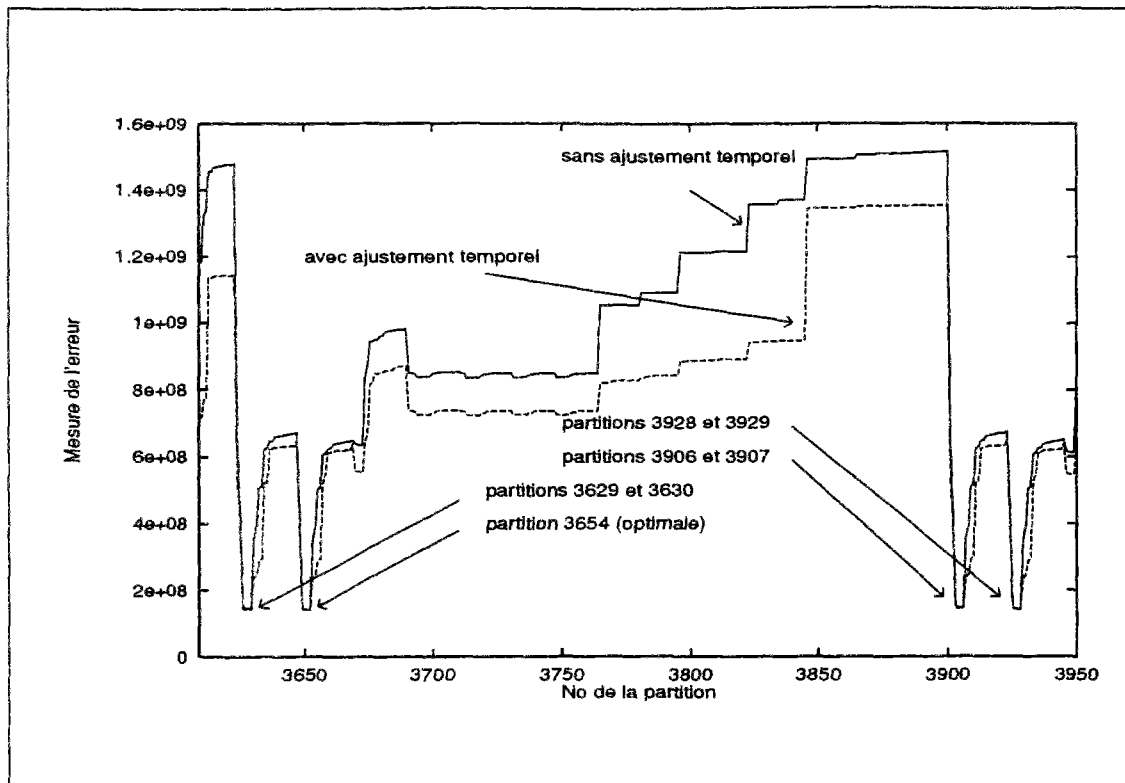


Figure 3.7: Ajustement temporel: grossissement autour de la **zone de minimum**.

En réalité, si la mise en œuvre numérique de cette mesure est possible, elle reste inexploitable. En effet, au niveau pratique, le domaine d'appartenance de $\mathcal{M}$ peut aller d'une valeur proche de zéro à une autre très grande ($9 \cdot 10^{10}$, 10^{11} , ...) selon le système étudié et le nombre de sollicitations prises en compte dans le modèle. Une simple analyse dimensionnelle de (3.10) montre qu'elle a un caractère **extrinsèque**: si -pour simplifier- nous choisissons $Q(\mathcal{M}, t) = 1$ (Fonction identité) alors la dimension de $\mathcal{M}$ est [$K^2 m^2 s$]. La valeur numérique de $\mathcal{M}$ dépend donc de plusieurs paramètres qui sont:

- la constante de temps principale (par $\mathcal{D}_t$).
- les dimensions géométriques (par $\mathcal{D}$).
- le nombre de sollicitations appliquées au système (par la $\sum_{j=1}^p$).
- le choix des unités pour exprimer les sollicitations.
- les pondérations utilisées.

Autrement dit, deux modèles réduits (de deux systèmes thermiques différents) ayant la même valeur numérique de $\mathcal{M}$ ne sont pas nécessairement aussi satisfaisants en ce qui concerne la représentation thermique de leurs systèmes respectifs. Cette mesure telle qu'elle est donnée par (3.36) est donc inadaptée à l'automatisation et ne peut guider l'utilisateur s'il souhaite obtenir le plus petit modèle réduit respectant une contrainte donnée sur l'erreur.

Pour remédier à cette difficulté, nous allons améliorer (3.36) de façon à la rendre plus intrinsèque. Pour cela, on se base sur les cinq points cités ci-dessus

1 ► Pour éliminer l'effet du nombre de sollicitations, on divise $\mathcal{M}$ par le nombre p .

2 ► Pour enlever la dépendance vis-à-vis de la constante de temps principale, on divise $\mathcal{M}$ sur un domaine temporel Δt_0 . L'étendue de Δt_0 est liée à celle de $\mathcal{D}_t = [t_1, t_2]$. Si $t_2 \geq t_{perm}$ (t_{perm} étant le temps d'établissement du régime permanent qui est voisin de $4\tau_1$) alors $\Delta t_0 = (t_{perm} - t_1)$ sinon $\Delta t_0 = (t_2 - t_1)$. En effet en raison du choix $\tilde{S} = S$, le régime permanent est garanti au niveau du modèle réduit (3.5) quel que soit le point considéré de la structure thermique. Le modèle réduit (3.5) ne peut donc être biaisé une fois que le régime permanent est établi et on a

$$\int_{t_1}^{t_2 \geq t_{perm}} \int_{\mathcal{D}} tr [\mathcal{E}^T(M, t).Q(M, t).\mathcal{E}(M, t)] dM dt$$

$$\# \int_{t_1}^{t_{perm}} \int_{\mathcal{D}} tr [\mathcal{E}^T(M, t).Q(M, t).\mathcal{E}(M, t)] dM dt$$

3 ► Il faut aussi remédier à l'effet de l'étendu spatiale et la pondération par $C(M)$. On divise alors $\mathcal{M}$ par le terme intégral $\int_{\mathcal{D}} C(M) dM$. Il correspond à l'énergie accumulée dans le système lorsque sa température s'élève de 1K.

4 ► Examinons l'effet du choix des unités pour les sollicitations. La mesure $\mathcal{M}$ est une somme de contributions relatives aux divers sollicitations appliquées au système (à cause de la trace $tr[.]$). Pour beaucoup de systèmes, les sollicitations intervenant dans son évolution sont diverses: températures [K], puissances [W], densités de flux [Wm^{-2}], etc. De ce fait, l'utilisation d'échelons unitaires lors des différentes manipulations sur $\mathcal{M}$ (minimisation, évaluation) peut conduire à des résultats insuffisants. En effet, dans ce cas les contributions à $\mathcal{M}$ des diverses sollicitations ne sont plus équilibrées. Il faut donc pondérer les sollicitations. On propose une méthode de pondération simple à utiliser, mais elle n'est pas unique.

Considérons le régime glissant donné par (1.12)

$$T_g(M, t) = \sum_{j=1}^p S_j(M) U_j(t)$$

Pour que les p sollicitations aient approximativement la même contribution au niveau de $T_g(M, t)$, il faut avoir

$$S_j(M) U_j(t) \# cte \quad \forall j = 1 \dots p$$

où cte est une constante arbitraire que l'on choisit égale à 1. Dans le cas de sollicitations en forme d'échelons unitaires, on définit les quantités

$$U_j^{\text{équilibrée}} = \frac{1}{\overline{S}_j}$$

avec
$$\overline{S}_j = \frac{\int_{\mathcal{D}} S_j(M) dM}{\int_{\mathcal{D}} dM}$$

Ces quantités sont les valeurs des sollicitations échelons (qui ne sont plus unitaires maintenant) qui permettent d'avoir des contributions du même ordre de grandeur au niveau de $T_g(M, t)$. Mais dans la pratique, il est plus simple de garder des valeurs unitaires pour les échelons et de remplacer la matrice $[B]$ par $[B\mathcal{P}]$, $\mathcal{P}$ désignant la matrice $[p, p]$ des pondérations donnée par $\mathcal{P}_{ij} = (\bar{S}_j)^{-1} \delta_{ij}$ (matrice diagonale). Cette manipulation laisse inchangées les expressions formelles de $\mathcal{M}$ (voir par exemple (3.33)).

5 ► Hormis la pondération sur les sollicitations dont nous venons de parler et qui est indispensable, on peut aussi privilégier une plage temporelle donnée (sous-domaine de $[0, \infty[$). Cette possibilité se réalise simplement en remplaçant $[B]$ par $[B\mathcal{F}(t)]$, $\mathcal{F}(t)$ étant une fonction de pondération sur le temps. Notons qu'on peut introduire une pondération à la fois sur les sollicitations et le temps: on remplace simplement $[B]$ par $[B\mathcal{P}\mathcal{F}(t)]$.

Les deux derniers points ont pour objectif l'amélioration de la précision du modèle réduit. Ce sont des choix liés à l'utilisateur qui se pose la question: quelles sont les facteurs de pondération à introduire **au niveau du modèle de référence** pour obtenir un bon modèle réduit vis-à-vis de toutes les sollicitations et/ou sur une plage temporelle donnée. Au contraire, les trois premiers points sont spécifiques aux systèmes thermiques et la mesure $\mathcal{M}$ en dépend. Aussi, les trois premiers points aboutissent à une mesure dont la dimension est $[K^2]$ et il est préférable de considérer sa racine carrée. La nouvelle mesure $\bar{\mathcal{M}}$ s'écrit

$$\bar{\mathcal{M}} = \sqrt{\frac{1}{p} \frac{\mathcal{M}}{\Delta t_0 \int_{\mathcal{D}} C(M) dM}} \quad \text{en Kelvin} \quad (3.37)$$

La mesure $\bar{\mathcal{M}}$ est reprise dans l'annexe E où elle est interprétée. Le domaine d'appartenance de $\bar{\mathcal{M}}$ est globalement entre zéro et la moyenne de $T(M, t)$ dans $\mathcal{D}$ lorsque le régime asymptotique est atteint. Par exemple, dans le cas d'une paroi homogène soumise à deux échelons unitaires de ses deux côtés, $\bar{\mathcal{M}}$ décroît de 1 (cas du modèle le plus biaisé, ie: $n = 0$) à zéro (cas du modèle le plus précis, ie: $n = N$).

Lorsque la taille "n" du modèle réduit augmente, $\mathcal{M}$ varie qualitativement en $\frac{1}{n^2}$ et $\bar{\mathcal{M}}$ en $\frac{1}{n}$. Il s'ensuit que $\mathcal{M}$ a une décroissance plus rapide que $\bar{\mathcal{M}}$. La figure 3.8 donne l'évolution de $\mathcal{M}$ et $\bar{\mathcal{M}}$ d'un bâtiment bizon (2 chambres) pour "n" variant de 1 à 15. Les valeurs de $\mathcal{M}$ et de $\bar{\mathcal{M}}$ ont été **rapportées à la même échelle** pour pouvoir faire une comparaison qualitative. On constate que les deux fonctionnelles n'évoluent presque plus au delà de $n = 12$.

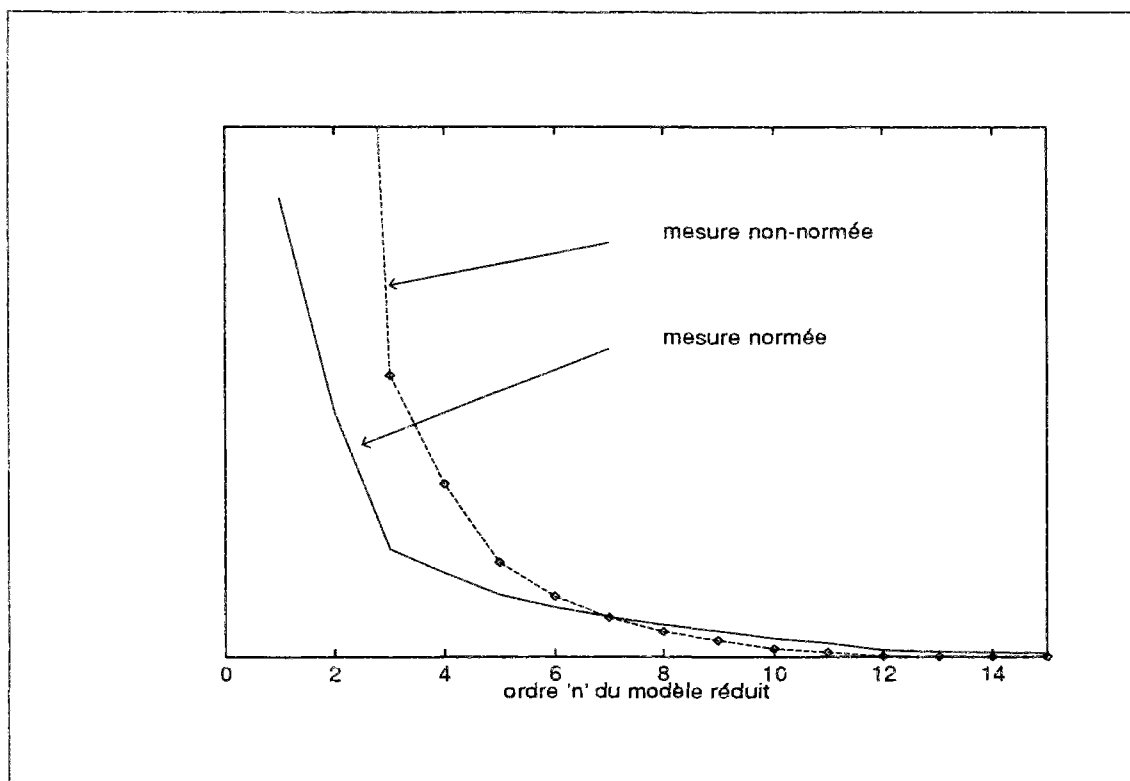


Figure 3.8: Décroissances qualitatives de $\mathcal{M}$ et $\overline{\mathcal{M}}$ en fonction de "n".

3.5.2 Algorithme d'amalgame

Il est maintenant possible de résumer notre travail en vue d'une implémentation informatique aisée. La démarche automatique de réduction doit suivre les étapes **A** à **G** résumées ci-dessous. Les étapes **B** et **C** sont données ⁷ pour des sollicitations échelons pour lesquelles il faut prendre $s = 0$ et impulsions pour lesquelles $s = 1$.

A Fixation des contraintes

On commence par imposer les deux contraintes: taille n_{max} du modèle à ne pas dépasser et erreur maximale de réduction $\overline{\mathcal{M}}_{max}$ tolérée.

B Choix des modes principaux d'amalgame

B.1

⁷A partir des résultats généraux de §3.3.4, il est possible de réécrire un algorithme de réduction valable quel que soit le type des sollicitations. Les relations intervenant dans l'algorithme correspondant font intervenir des intégrales temporelles.

Le mode principal $V_1(M)$ associé à H_1 (le premier sous-espace d'amalgame) est le mode correspondant au maximum des quantités

$$-\frac{\lambda_m^{2s-1}}{2} \left(\sum_{l=1}^p B_{ml}^2 \right) \left[\Delta e^{2\lambda_m t} \right]_{t_2}^{t_1} \quad m = 1 \cdots N$$

lorsqu'on fait varier l'indice m de 1 à N .

B.2

Les modes principaux des sous-espaces d'amalgame H_k ($k = 2 \cdots n$) sont sélectionnés récursivement. Pour cela on se base sur la relation donnant l'erreur produite par le mode "m" lorsqu'il est affecté comme mode mineur au mode principal "j"

$$\begin{aligned} \mathcal{M}_{mj} = & -\frac{\lambda_m^{2s-1}}{2} \left(\sum_{l=1}^p B_{ml}^2 \right) \left[\Delta e^{2\lambda_m t} \right]_{t_2}^{t_1} + \\ & \frac{2\lambda_z \lambda_m^{2s}}{(\lambda_z + \lambda_m)^2} \frac{(\sum_{l=1}^p B_{zl} B_{ml})^2}{\sum_{l=1}^p B_{zl}^2} \left(\frac{\left[\Delta e^{(\lambda_z + \lambda_m)t} \right]_{t_1}^{t_2}}{\left[\Delta e^{2\lambda_z t} \right]_{t_1}^{t_2}} \right)^2 \\ & \text{avec } z = u^j(1) \end{aligned}$$

Pour chaque mode "m" encore non-principal, on repère la quantité

$$\min(m) = \min\{\mathcal{M}_{mj}, j \text{ balayant les modes principaux } 1 \cdots i \text{ déjà fixés}\}$$

Le mode principal "i+1" est celui qui réalise

$$\max\{\min(m) \text{ } m \text{ balayant les indices des modes non-principaux}\}$$

B.3

Les modes mineurs sont ceux qui ne sont pas sélectionnés dans **B.1** et **B.2**.

C Distribution des modes mineurs sur les sous-espaces

Pour chaque mode mineur V_m , on calcule les n quantités suivantes ($q = 1 \cdots n$) relatives aux différents sous-espaces d'amalgame.

$$\begin{aligned} \mathcal{M}_{mq} = & -\frac{\lambda_m^{2s-1}}{2} \left(\sum_{l=1}^p B_{ml}^2 \right) \left[\Delta e^{2\lambda_m t} \right]_{t_2}^{t_1} + \\ & \frac{2\lambda_z \lambda_m^{2s}}{(\lambda_z + \lambda_m)^2} \frac{(\sum_{l=1}^p B_{zl} B_{ml})^2}{\sum_{l=1}^p B_{zl}^2} \left(\frac{\left[\Delta e^{(\lambda_z + \lambda_m)t} \right]_{t_1}^{t_2}}{\left[\Delta e^{2\lambda_z t} \right]_{t_1}^{t_2}} \right)^2 \\ & \text{avec } z = u^q(1) \end{aligned}$$

L'indice q correspondant à la plus petite de ces quantités détermine le meilleur sous-espace H_q d'affectation du mode mineur m .

[D] Evaluation quantitative de l'erreur de réduction

On calcule par (3.37) l'erreur occasionnée par la réduction. Si le modèle est jugé satisfaisant au regard des contraintes de taille et de précision imposées, alors on passe à l'étape suivante **[E]**. Notons que la contrainte de précision peut être satisfaite pour une dimension n' inférieure à la taille maximale n_{max} autorisée pour le modèle amalgamé. Dans cette situation, on peut reprendre en **[C]** en considérant la nouvelle taille n' proposée par l'algorithme. Au contraire, si la contrainte de précision est dépassée pour l'ordre souhaité du modèle, alors on peut décider d'augmenter n_{max} ou M_{max} (ou les deux à la fois). Dans ce dernier cas, le programme est repris en **B.2** ou poursuivi en **[E]** selon le cas.

[E] Calcul des modes amalgamés

On calcule la matrice $\tilde{P}$ en générant un mode amalgamé à partir de chaque sous-espace d'amalgame. Il suffit alors d'utiliser les relations (3.11) et (3.23) pour chaque sous-espace. Dans le cas des sollicitations tests (échelon et impulsion), la relation (3.23) est explicitée par (3.27).

[F] Calcul de $[\tilde{F}]$ et $[\tilde{B}]$

$[\tilde{F}]$ et $[\tilde{B}]$ s'obtiennent par troncature de F et B respectivement: on ne garde que les colonnes et lignes qui correspondent aux modes principaux. Si l'on souhaite opérer un ajustement temporel, il faut alors résoudre l'équation transcendante (3.34) pour chaque sous-espace d'amalgame. Les valeurs propres localisées (grâce à une dichotomie par exemple) forment la matrice $[\tilde{F}]$.

[G] Obtention du modèle réduit

Le modèle réduit est finalement donné par le système (3.5).

L'organigramme de la figure 3.9 développe l'étape **[D]** de l'algorithme ci-dessus pour faciliter la compréhension de la réduction automatique. Il faut noter que dans cette étape, le passage de l'ordre " n " à " $n+1$ " est simple: il suffit de chercher le mode principal $n+1$ et de refaire la distribution des modes mineurs. Le nombre d'opérations peut être élevé mais le temps de calcul reste faible car les relations utilisées sont explicites.

L'ensemble des étapes **[A]** à **[D]** a été implémenté dans un logiciel nommé AnaRed.

3.6 Structure modale du modèle réduit

L'importance de la similitude de la structure du modèle réduit avec celle du modèle de référence a été identifiée depuis déjà quelques décennies pour les réseaux d'énergie dans le domaine de l'électronique [61]. En thermique, il y a les travaux de synthèse modale [28] dans laquelle le modèle réduit élaboré pour un composant d'un système thermique complexe ne peut être raccordé à ceux d'autres composants si ses modes ne sont pas orthogonaux entre eux au sens de (1.6).

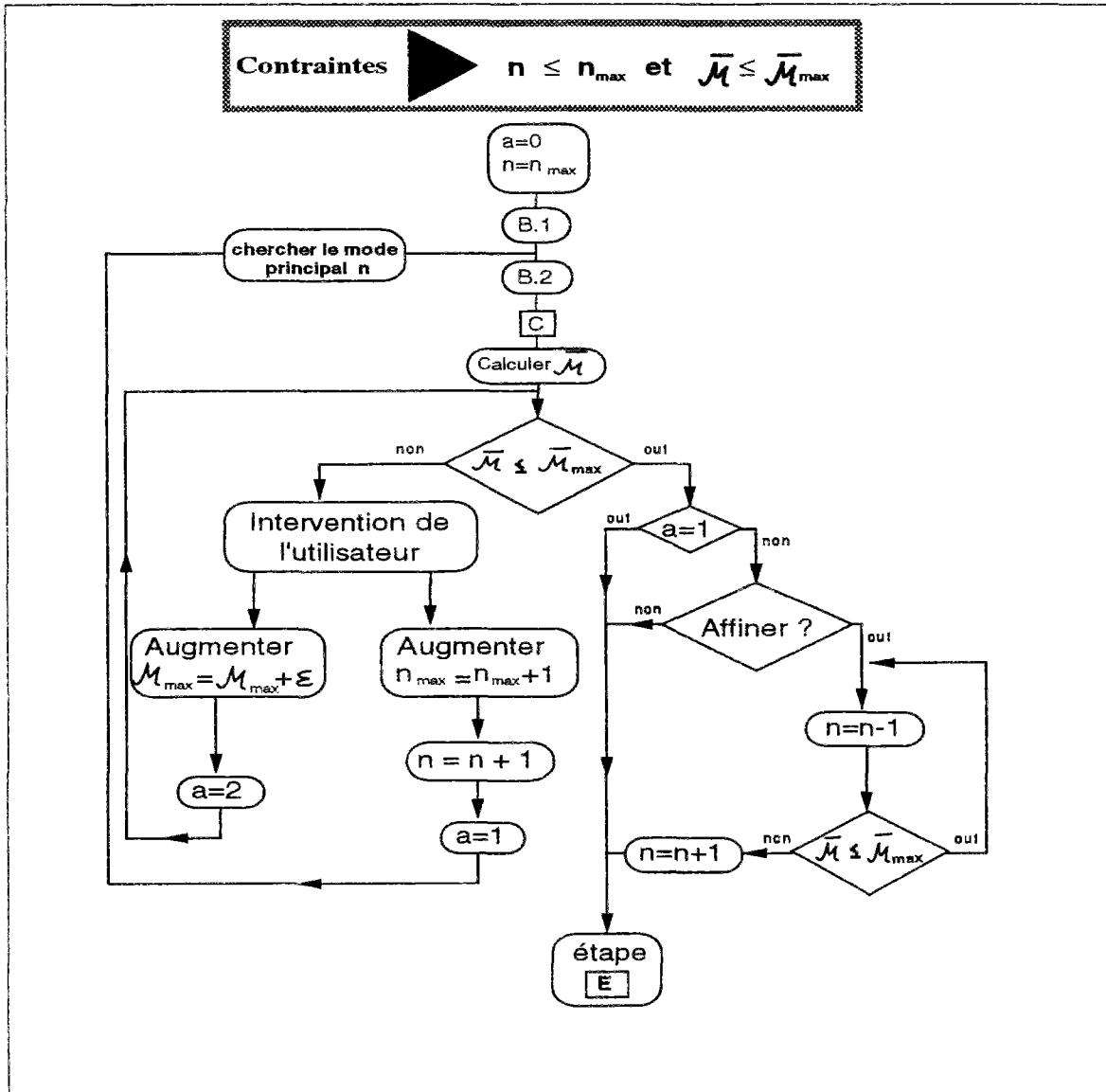


Figure 3.9: Algorithme de principe pour la réduction automatique.

Dans la forme donnée jusqu'à présent, le modèle réduit par amalgame n'est pas normé vis-à-vis de la norme associée au produit scalaire introduit par (1.6). On va donc effectuer une normalisation afin de rendre au modèle réduit une forme "standard" identique au modèle de référence.

Modèle détaillé "dimension N"

$$\begin{aligned}\dot{x}(t) &= Fx(t) + B\dot{U}(t) \\ T(M,t) &= Px(t) + SU(t)\end{aligned}\tag{3.38}$$

Modèle réduit "dimension n"

$$\begin{aligned}\dot{\tilde{x}}(t) &= \tilde{F}\tilde{x}(t) + \tilde{B}\dot{U}(t) \\ \tilde{T}(M,t) &= \tilde{P}\tilde{x}(t) + SU(t)\end{aligned}\tag{3.39}$$

Les modes amalgamés de ce dernier modèle sont orthogonaux mais non-orthonormés. On peut transformer (3.39) en

$$\begin{aligned}\dot{\tilde{x}}^*(t) &= \tilde{F}^*\tilde{x}^*(t) + \tilde{B}^*\dot{U}(t) \\ \tilde{T}(M,t) &= \tilde{P}^*\tilde{x}^*(t) + SU(t)\end{aligned}\tag{3.40}$$

où la relation d'orthonormalité des modes amalgamés est vérifiée. Pour cela cherchons $\tilde{F}^*$, $\tilde{B}^*$, $\tilde{P}^*$ et $\tilde{x}^*(t)$.

Les modes amalgamés $\tilde{V}_i(M)$ $i = 1 \dots n$ sont donnés par (voir (3.11))

$$\tilde{V}_i(M) = P_i \omega_i\tag{3.41}$$

et vérifient

$$\begin{aligned}\forall i = 1, 2 \dots n \quad \forall j = 1, 2 \dots n \\ \langle \tilde{V}_i(M), C(M)\tilde{V}_j(M) \rangle = \delta_{ij} \sum_{m=1}^{n_i} \omega_i^2(m)\end{aligned}\tag{3.42}$$

Posons

$$\tilde{V}_i^*(M) = \frac{\tilde{V}_i(M)}{\mathcal{N}_i} \quad \forall i = 1 \dots n\tag{3.43}$$

$$\text{avec } \mathcal{N}_i = \sqrt{\omega_i^T \omega_i}$$

Les $\tilde{V}_i^*(M)$ sont les modes amalgamés orthonormés. Ils sont associés aux constantes de temps⁸ $\tilde{\tau}_i$ $i = 1 \dots n$ qui sont celles des modes principaux des sous-espaces d'amalgame. Autrement dit

$$\tilde{\tau}_i = \tau_{u^1(i)} \quad \forall i = 1 \dots n$$

On pose

$$\boxed{\tilde{F}^* = \tilde{F} \quad \text{et} \quad \tilde{P}^* = \tilde{P} \tilde{\mathcal{N}}^{-1}} \quad (3.44)$$

où $\tilde{\mathcal{N}}$ est la matrice *diagonale* $[n, n]$ composée des normes ($\mathcal{N}_i$ $i = 1 \dots n$) des modes amalgamés.

Remarque: Les éléments "propres" $(\tilde{V}_i^*(M), \tilde{\lambda}_i)$ ne sont pas solution du problème aux valeurs propres (1.5). Ils vérifient⁹

$$\begin{aligned} \mathcal{L}(\tilde{V}_i^*(M)) &= \mathcal{L} \left(\sum_{m=1}^{n_k} \left(\frac{\omega_i(m)}{\mathcal{N}_i} V_{u^i(m)}(M) \right) \right) \\ &= \frac{\omega_i(m)}{\mathcal{N}_i} \sum_{m=1}^{n_k} \mathcal{L} \left(V_{u^i(m)}(M) \right) \quad (\text{linéarité de } \mathcal{L}) \\ &= \frac{\omega_i(m)}{\mathcal{N}_i} C(M) \sum_{m=1}^{n_k} \lambda_{u^i(m)} V_{u^i(m)} \end{aligned}$$

Ecrivons l'expression de la température avec le modèle (3.39)

$$\begin{aligned} \tilde{T}(M, t) &= \tilde{P} \tilde{x}(t) + S U(t) \\ &= \tilde{P}^* \tilde{\mathcal{N}} \tilde{x}(t) + S U(t) \end{aligned}$$

On pose alors

$$\boxed{\tilde{x}^*(t) = \tilde{\mathcal{N}} \tilde{x}(t)} \quad (3.45)$$

Multiplions (à gauche) par $\tilde{\mathcal{N}}$ l'équation d'état de (3.39)

$$\dot{\tilde{x}}^*(t) = \tilde{\mathcal{N}} \left(\tilde{F} \tilde{x}(t) + \tilde{B} \dot{U}(t) \right)$$

⁸Les constantes de temps τ_i sont les inverses (au signe près) des valeurs propres λ_i

$$\tau_i = -\frac{1}{\lambda_i}$$

⁹Chaque couple $(\lambda_i, V_i(M))$ vérifie $\mathcal{L}(V_i(M)) = \lambda_i C(M) V_i(M)$

Le produit de $\tilde{\mathcal{N}}$ et $\tilde{F}$ étant commutatif, on pose finalement

$$\boxed{\tilde{B}^* = \tilde{\mathcal{N}}\tilde{B}} \quad (3.46)$$

et les composantes de $\tilde{B}^*$ sont formellement données par (d'après (1.16))

$$B_{i,j}^* = -\mathcal{N}_i \int_{\mathcal{D}} C(M) S_j(M) V_{u^i(1)}(M) dM \quad (3.47)$$

Ainsi s'achève la définition du modèle amalgamé normé.

3.7 Estimation du temps de calcul

Il est possible d'avoir une idée du temps des calculs nécessaires au fonctionnement de la méthode présentée. Pour cela, nous analysons les étapes **B**, **C**, **D** et **E**. Nous considérons pour simplifier les sollicitations de type échelon.

B La recherche de "n" modes principaux nécessite l'utilisation de la relation explicite (3.33) (sauf pour le 1^{er} mode principal). Cette relation doit être évaluée $N + (N - 1) + 2(N - 2) + 3(N - 3) \cdots n(N - n)$ fois, c'est à dire

$$N + \sum_{i=1}^n i(N - i) \quad \text{fois}$$

C La répartition des $(N - n)$ modes mineurs autour des modes principaux nécessite aussi l'utilisation de la relation explicite (3.33). Cette relation doit être évaluée

$$(N - n)n \quad \text{fois}$$

D La mesure $\mathcal{M}$ (ou $\overline{\mathcal{M}}$) est la somme des contributions de tous les modes mineurs. Or ces contributions sont évaluées dans l'étape **C**. Il suffit alors de faire la somme

$$\mathcal{M} = \sum_{k=1}^n \sum_{m=1}^{n_k} \mathcal{M}_{km}$$

Il n'y a pas de termes nouveaux à calculer dans cette étape.

E Les modes amalgamés $\tilde{V}_k(M)$ $k = 1 \cdots n$ sont complètement déterminés par la connaissance des coefficients $\{\omega_k(m) \mid k = 1 \cdots n, m = 2 \cdots n_k\}$. Ces coefficients, au nombre de $(N - n)$ (car $\omega_k(1) = 1 \quad \forall k = 1 \cdots n$) sont donnés explicitement par (3.27). Il faut donc évaluer (3.27)

$$(N - n) \quad \text{fois}$$

Par ailleurs les relations (3.33) et (3.27) sont, du point de vue numérique, du même ordre de complexité. Si l'on désigne par

δt_1 le temps moyen (en CPU) nécessaire pour calculer l'expression (3.33)

δt_2 le temps moyen (en CPU) nécessaire pour calculer l'expression (3.27)

alors le temps t nécessaire à la machine pour exécuter les étapes **B** à **E** de l'algorithme est d'environ

$$t \simeq \left(N + n(N - n) + \sum_{i=1}^n i(N - i) \right) \delta t_1 + (N - n) \delta t_2$$

Sur la machine SUN SPARC Station 10 que nous utilisons, on trouve

$\delta t_1 \# 20 \cdot 10^{-4} s$ (temps en CPU qui est une moyenne calculée sur 100000 termes)

$\delta t_2 \# 14 \cdot 10^{-4} s$ (temps en CPU qui est une moyenne calculée sur 100000 termes)

Pour $N = 200$ et $n = 5$, on a: $t = 9s$ environ.

Evidemment ces chiffres sont indicatifs et le temps réel pour la réduction d'un modèle thermique est bien plus grand. Indépendamment du type de la machine, la structure du programme et le langage utilisé conditionnent grandement le temps global de la méthode. A cela peuvent s'ajouter tous les tests informatiques nécessaires, les entrées-sorties du programme, etc. Mais dans tous les cas, reposant sur des relations explicites, la méthode d'amalgame modal est certainement performante en temps de calcul. Avec $N = 200$ et $n = 5$, le programme **AnaRed** (langage ADA, environ 3000 lignes) qui utilise la méthode d'amalgame modal, le temps global (du chargement en mémoire du modèle de référence jusqu'au stockage du modèle amalgamé **inclus**) nécessaire à la réduction est d'environ 130s (temps CPU et toujours sur la même machine).



Exemples de réduction de modèles thermiques

Chapitre 4

Paroi tricouche

4.1 Description de la paroi

L'exemple traité ici est une structure tricouche en conduction unidimensionnelle. Cet exemple est volontairement simple (dans un but pédagogique). Il servira en effet à illustrer les étapes importantes de la méthode de réduction par amalgame modal.

La paroi traitée est une structure symétrique et composée de trois couches: béton-isolant-béton. Deux lames d'air de $1mm$ d'épaisseur chacune, sont introduites de part et d'autre de l'isolant. Ceci permet de tenir compte - de façon approximative - de l'air enfermé dans les aspérités de l'isolant (Polystyrène) en contact avec le béton. La figure 4.1 illustre la paroi tandis que le tableau 4.1 donne les paramètres thermophysiques de chaque couche homogène.

Les coefficients d'échange **globaux**¹ sont fixés à: $h_g = 8Wm^{-1}K^{-1}$ et $h_d = 18Wm^{-1}K^{-1}$. Ces valeurs sont à peu près celles recommandées par le CSTB² pour des murs extérieurs d'une enveloppe de bâtiment.

Le modèle d'état modal de référence est obtenu à partir du logiciel **MurAna**³ [54]. Il repose

¹On regroupe dans les coefficients h_g et h_d ceux de la convection et du rayonnement linéarisés.

²Centre Scientifique et Techniques du Bâtiment

³Nous avons développé, à partir d'une formulation analytique, un logiciel nommé **MurAna**, capable de générer un modèle d'état modal d'une paroi plane multicouche. Le contact thermique entre les couches peut être imparfait. La démarche utilisée consiste à résoudre une équation transcendante aux valeurs propres issue du formalisme des matrices de transfert. Le calcul des modes propres associés ne pose pas de problème particulier. Cependant des algorithmes appropriés sont introduits pour fiabiliser le programme informatique. Les algorithmes importants introduits dans **MurAna** permettent

- de ne pas rater de racines de l'équation aux valeurs propres, qui peuvent être numériquement très proches ou simplement groupées.

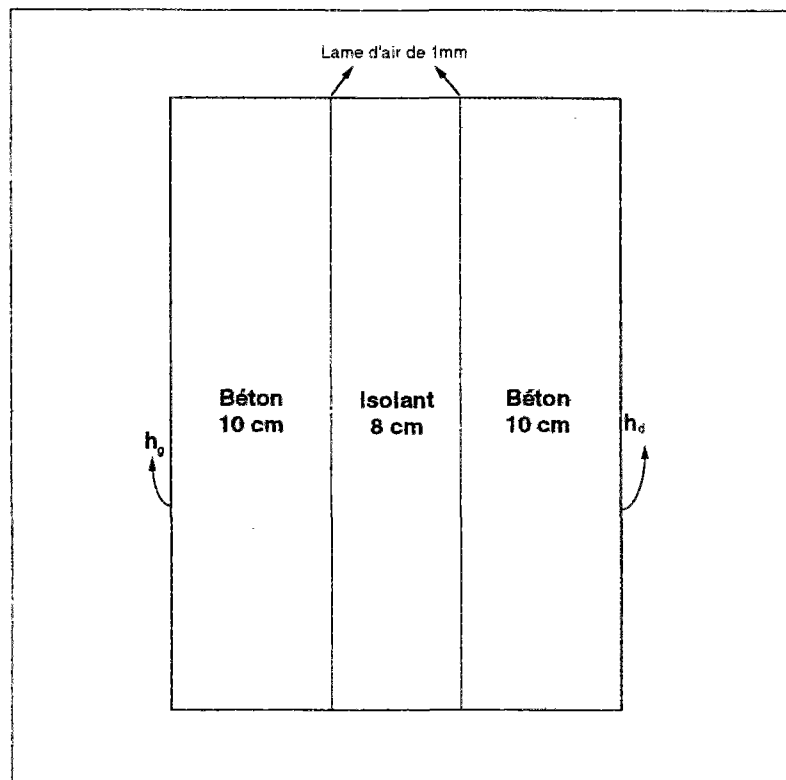


Figure 4.1: Structure multicouche

ici sur 250 éléments propres. La constante de temps principale est $\tau_1 = 7\text{h } 42\text{mn } 40.98\text{s}$. La dernière constante de temps calculée est $\tau_{250} = 0.13\text{s}$, ce qui autorise des simulations précises à partir de $t = 1\text{s}$ (4 à 5 fois τ_{250}). Une fois les fonctions propres connues analytiquement, il est possible de calculer leurs valeurs en des points quelconques de la structure, ce qui permet aussi de travailler sur un maillage non uniforme. Nous avons alors opéré un échantillonnage de trois nœuds par centimètre, ce qui revient à prendre 57 nœuds pour l'ensemble de la paroi (la paroi a une épaisseur totale de 28cm).

Les sollicitations agissant sur la structure sont au nombre de deux: températures des ambiances côté gauche et droit. Les sorties de notre modèle sont toutes les températures aux nœuds considérés (au total 57 nœuds).

Les résultats qui suivent ont été obtenus pour les conditions suivantes:

- d'asservir le pas d'exploration par une euristique originale, ce qui permet un gain appréciable en temps de calcul.

MurAna est doté d'une interface graphique interactive qui permet de visualiser l'évolution temporelle d'une sortie mais également d'analyser ou de réduire le modèle d'état modal.

matériaux	béton	Air	isolant	Air	béton
conductivité($Wm^{-1}K^{-1}$)	1.75	0.025	0.04	0.0625	1.75
masse vol.(kgm^{-3})	2500	1.24	30	1.24	2500
C_p ($Jkg^{-1}K^{-1}$)	820	1000.0	850	1000.0	820
épaisseur(m)	0.1	0.001	0.08	0.001	0.1

Tableau 4.1: Composition de la structure multicouche

- Transferts conductifs unidimensionnels dans la paroi.
- Plage temporelle $\mathcal{D}_t = [0, \infty]$.
- Sollicitations de type échelon.
- Champ de température initial nul.

4.2 Réduction modale

4.2.1 Les sous-espaces d'amalgame

La réduction du modèle d'état modal de la paroi tricouche est ici développée. Nous appliquons l'algorithme de réduction présenté au §3.5 **en imposant simplement la dimension du modèle réduit $n = 5$** . Pour cette dimension, on obtient les valeurs de $\mathcal{M}$ et $\overline{\mathcal{M}}$ suivantes

mesure de l'erreur $\mathcal{M}$ ($J^\circ C s$)	erreur $\overline{\mathcal{M}}$ ($^\circ C$)
8639.63	$3.07 \cdot 10^{-4}$

Tableau 4.2: Valeurs de $\mathcal{M}$ et $\overline{\mathcal{M}}$ pour $n = 5$.

Les sous-espaces d'amalgames issus de la méthode d'amalgame sont donnés dans le tableau 4.3. Les modes principaux des différents sous-espaces sont encadrés $\boxed{\bullet}$.

Les modes principaux d'amalgame ne sont pas ceux de la troncature de Marshall. En effet, le cinquième mode principal est le mode sept. Notons que pour cet exemple, on obtient trois sous-espaces d'amalgame monodimensionnels: les sous-espaces 1,2 et 4. La plupart des modes mineurs sont affectés dans le dernier sous-espace d'amalgame.

sous-espace	1	2	3	4	5
composition en modes	1	2	3 6	4	7 5, 8, 9, 10 11, 12 ... 250

Tableau 4.3: répartition de l'espace des modes propres en 5 sous-espaces d'amalgame

Remarque: Dans le cadre de l'étude thermique des enveloppes de bâtiments, Lefebvre [41] tronque le modèle modal puis corrige le dernier élément propre grâce à une information récupérée sur les modes négligés. Le dernier sous-espace d'amalgame obtenu pour notre paroi tricouche regroupe la plupart des modes mineurs. Ce résultat peut être rapproché de la technique de Lefebvre. Mais ceci est loin de constituer une règle générale même pour les parois de bâtiment. En effet, pour la même paroi tricouche décrite ci-dessus, en changeant simplement les valeurs numériques des coefficients d'échanges convectifs ($h_g = 11 \frac{W}{mK}$ et $h_d = 19 \frac{W}{mK}$), nous obtenons les sous-espaces suivants

Sous-espace 1: **1** (dimension 1)

Sous-espace 2: **2** (dimension 1)

Sous-espace 3: **3** 7, 9, 12, 17, 25, 30, 33, 35, 47, 49, 52 ... (dimension 62)

Sous-espace 4: **4** (dimension 1)

Sous-espace 5: **6** 5, 8, 10, 11, 13, 14, 15, 16, 18, 19, 20 ... (dimension 185)

Nous aurons l'occasion de voir d'autres sous-espaces d'amalgames lorsque nous aborderons les prochains exemples.

4.2.2 Les modes amalgamés

Pour avoir une illustration simple, intéressons-nous seulement au troisième sous-espace d'amalgame qui se compose du mode principal **3** et du mode mineur 6. A partir de ce sous-espace, nous allons générer le troisième mode amalgamé $\tilde{V}_3(M)$ défini par (voir §3.3.3)

$$\tilde{V}_3(M) = \omega_3(1)V_3(M) + \omega_3(2)V_6(M)$$

Les coefficients $\omega_3(1)$ et $\omega_3(2)$ qui réalisent le minimum de la fonctionnelle $\mathcal{M}_3$ sont

$$\omega_3(1) = 1$$

$$\omega_3(2) = 0.087$$

Le coefficient $\omega_3(2)$ étant faible, le mode mineur $V_6(M)$ apporte peu d'information thermique au mode principal $V_3(M)$. L'interprétation géométrique du mode amalgamé $\tilde{V}_3(M)$ est donnée dans figure 4.2. On remarquera que le mode principal est peu modifié.

En revanche, le dernier sous-espace d'amalgame est de dimension 245. Le mode principal $V_7(M)$ récupère une information thermique de 244 modes mineurs. Voyons si ce mode principal est modifié de façon notable ou non. Pour cela, nous traçons dans la figure 4.3 le mode principal $V_7(M)$ et le mode amalgamé $\tilde{V}_5(M)$ issu du cinquième sous-espace d'amalgame. La différence entre les deux modes est particulièrement importante dans la couche d'isolant. La correction apportée par l'amalgame modal dans le dernier sous-espace concerne donc l'isolant. En examinant le mode mineur 5 qui est donné dans la figure 4.4, on s'aperçoit qu'il a une amplitude importante dans l'isolant et il est difficile de le négliger complètement. Le coefficient $\omega_5(2)$ associé au mode mineur $V_5(M)$ est égal à 0.35 et est justement le plus grand parmi $\omega_5(m)$ $m = 2 \cdots 245$.

La forme de chaque mode amalgamé permet de déduire globalement son rôle dans le modèle réduit. Examinons les modes amalgamés dans l'ordre:

- le premier mode amalgamé est principalement localisé dans la couche de béton du côté gauche de la paroi. Il est donc important pour la dynamique de cette couche. Son tracé est donné dans la figure 4.5.
- Le mode amalgamé $\tilde{V}_2(M)$ est approximativement le symétrique (au signe près) de $\tilde{V}_1(M)$ et il concerne la dynamique de la couche de béton du côté droit de paroi. Son tracé est donné dans la figure 4.5.
- Le troisième mode amalgamé $\tilde{V}_3(M)$ est réparti entre la couche de béton côté gauche et l'isolant. Ce mode intervient vraisemblablement pour assurer le "lien thermique" (ou continuité thermique) entre ces deux couches. Son tracé est donné dans la figure 4.2.
- Le mode amalgamé $\tilde{V}_4(M)$ peut aussi être considéré (au signe près) comme le symétrique de $\tilde{V}_3(M)$. De façon similaire, ce mode assure le lien thermique entre la couche du béton côté droit et l'isolant. Son tracé est donné dans la figure 4.5.
- Le dernier mode amalgamé est principalement localisé dans la centre de la paroi. L'opération d'amalgame modal dans le dernier sous-espace a donc eu pour effet de représenter l'isolant. Son tracé est donné dans la figure 4.3.

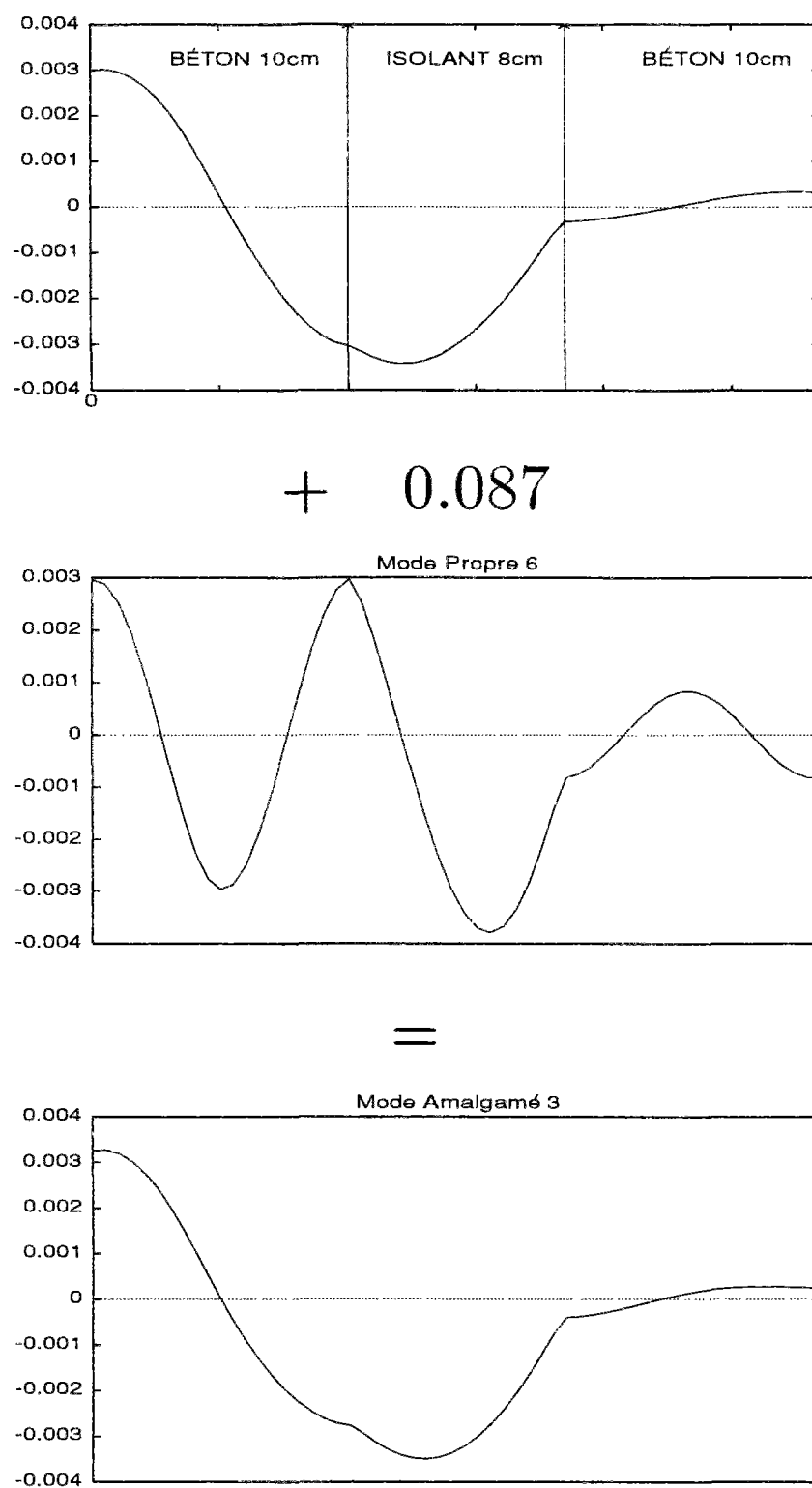


Figure 4.2: Interprétation géométrique du troisième mode amalgamé

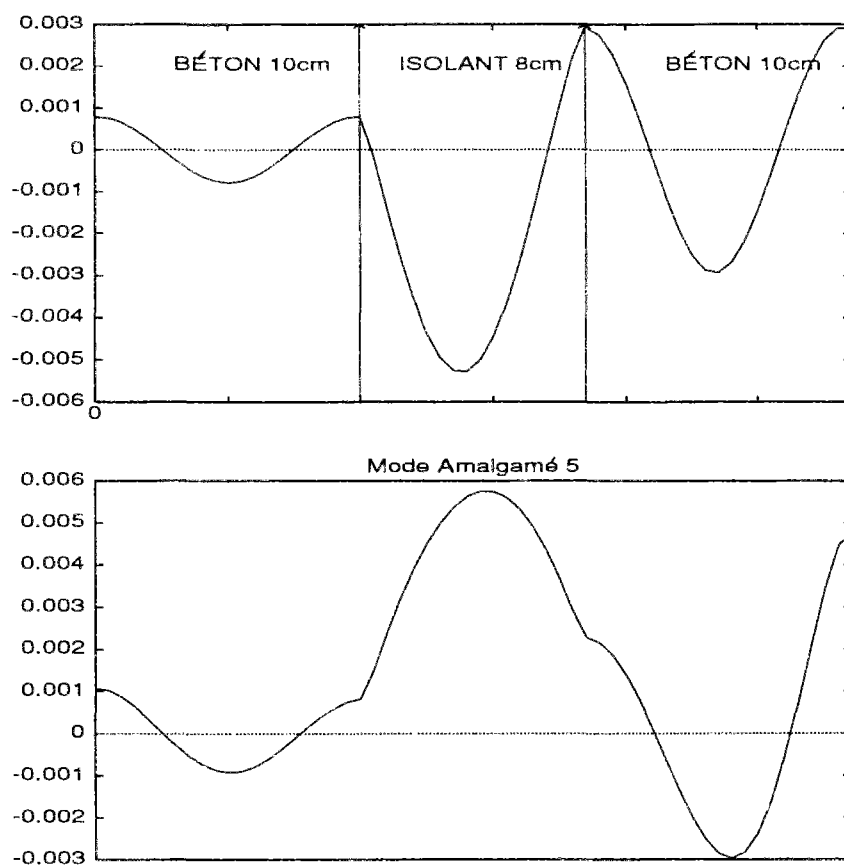


Figure 4.3: Sous-espace 5: comparaison du mode principal $V_7(M)$ et du mode amalgamé $\tilde{V}_5(M)$.

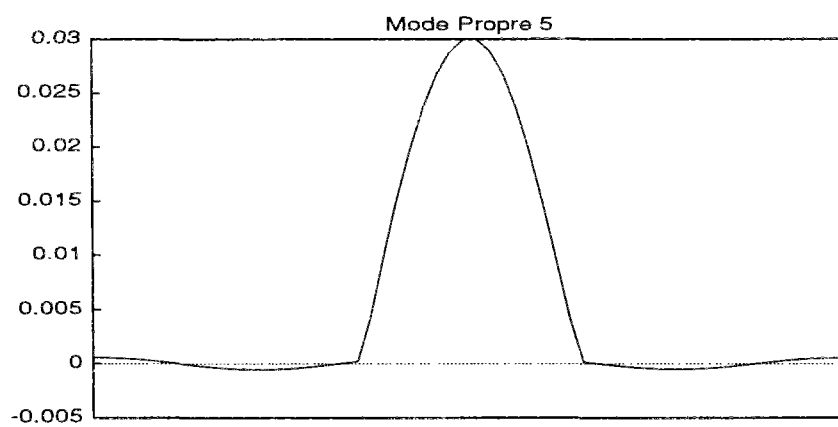


Figure 4.4: Mode Propre 5

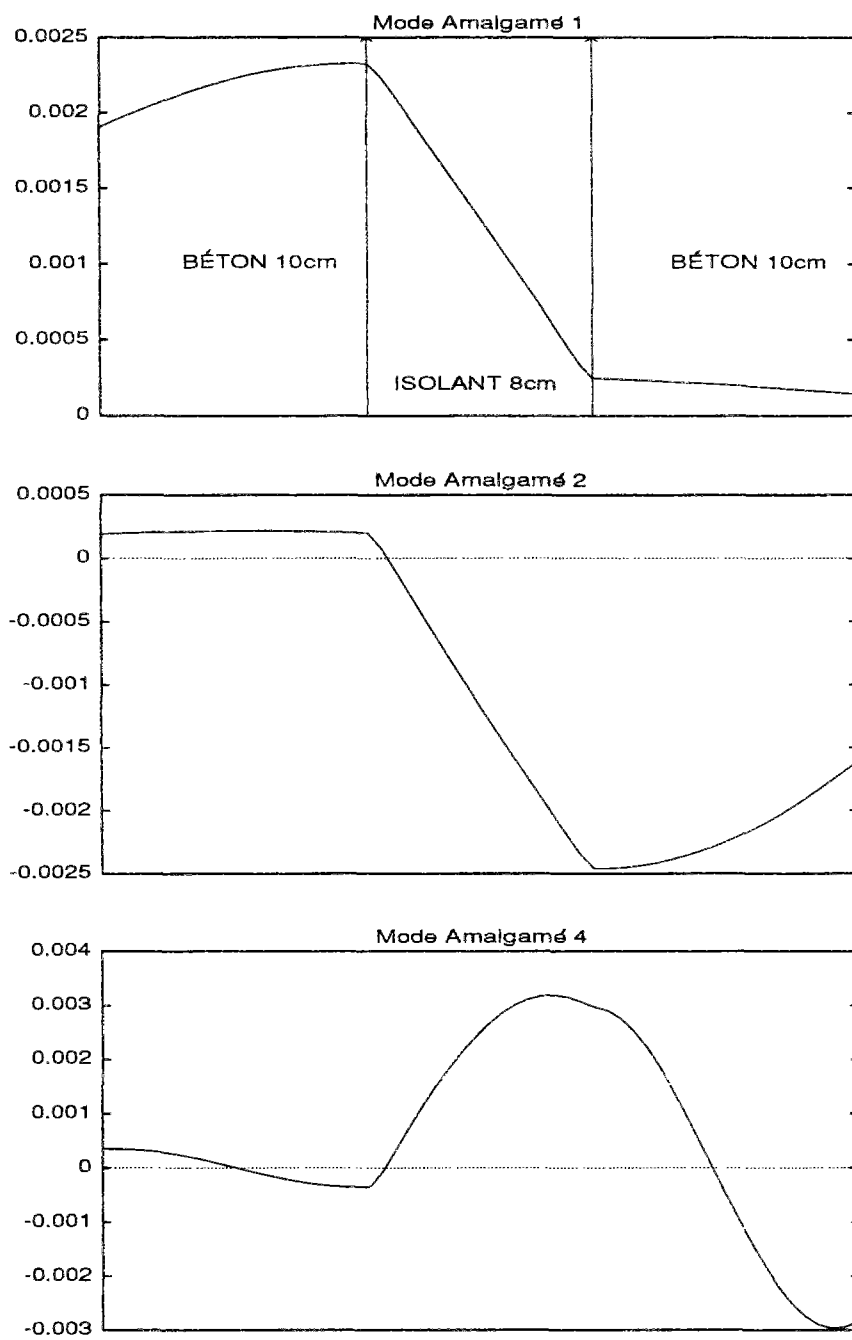


Figure 4.5: Modes amalgamés 1,2 et 4

4.2.3 Simulation de l'évolution thermique

Afin de montrer l'efficacité de la méthode de réduction par amalgame modal, il est nécessaire d'effectuer la simulation de l'évolution thermique de la paroi. En fait nous limiterons les calculs aux points d'échantillonnage au nombre de 57 pour cet exemple. Pour cela, nous considérons la paroi tricouche lorsqu'elle est sollicitée par un échelon unitaire à travers sa face droite (la sollicitation sur la face gauche est nulle). Pour simplifier, le champ initial de température est considéré comme étant identiquement nul et le domaine temporel $\mathcal{D}_t$ s'étale de zéro à l'infini.

La représentation graphique utilisée est tridimensionnelle: L'axe "x" représente l'épaisseur matérielle de la paroi, c'est à dire la paroi elle-même. L'axe "y" est le temps. Cet axe utilise une échelle de type logarithmique⁴, ce qui permet de mieux distinguer les temps d'évolution rapide où se trouvent généralement les erreurs des différents modèles. Les axes "x" et "y" forment un plan horizontal. La direction verticale est la température (réponse des différents nœuds à la sollicitation échelon).

La figure 4.6 donne l'évolution thermique de l'ensemble des nœuds de la paroi tricouche. Ce champ spatio-temporel de température est calculé à partir du modèle de référence (**M.Ref**), qui repose sur 250 éléments propres. Le temps d'établissement du régime permanent est d'environ 28h. Remarquons au passage le rôle de l'isolant dans l'atténuation de la perturbation thermique à travers la paroi.

Considérons maintenant le modèle principal (**M.Pri**), c'est à dire le **modèle tronqué** reposant sur les modes principaux seulement (modes 1,2,3,4 et 7). La figure 4.7 montre l'écart (en valeur absolue) entre les champs spatio-temporels de température donnés par (**M.Ref**) et (**M.Pri**) respectivement. La localisation des erreurs dans la couche de l'isolant est nette.

On suit la même démarche en ce qui concerne le modèle amalgamé (**M.Ama**). Examinons la figure 4.8 qui donne l'écart (en valeur absolue) entre les champs spatio-temporels de température donnés par (**M.Ref**) et (**M.Ama**) respectivement. Les erreurs du modèle principal sont maintenant bien atténuées grâce à la prise en compte de l'effet des modes mineurs. Remarquons que la correction apportée par l'amalgame modal concerne surtout le domaine de l'espace occupé par l'isolant. Il est possible d'expliquer ce résultat: la mesure de l'erreur $\mathcal{M}$ fait intervenir une pondération par la capacité calorifique (voir §3.3.3). Logiquement, les modes propres d'amplitude importante dans la couche d'isolant

⁴ Rappel: La variable temporelle t^* utilisée pour l'axe "y" est définie par

$$t^* = \frac{\text{Ln} \left(1 + 100 \frac{t}{t_{ref}} \right)}{\text{Ln}(101)}$$

Pour $t = t_{ref}$, on obtient $t^* = 1$. La variation du paramètre t_{ref} permet d'avoir une échelle dilatable de telle sorte que les erreurs puissent être distinguées aux faibles instants.

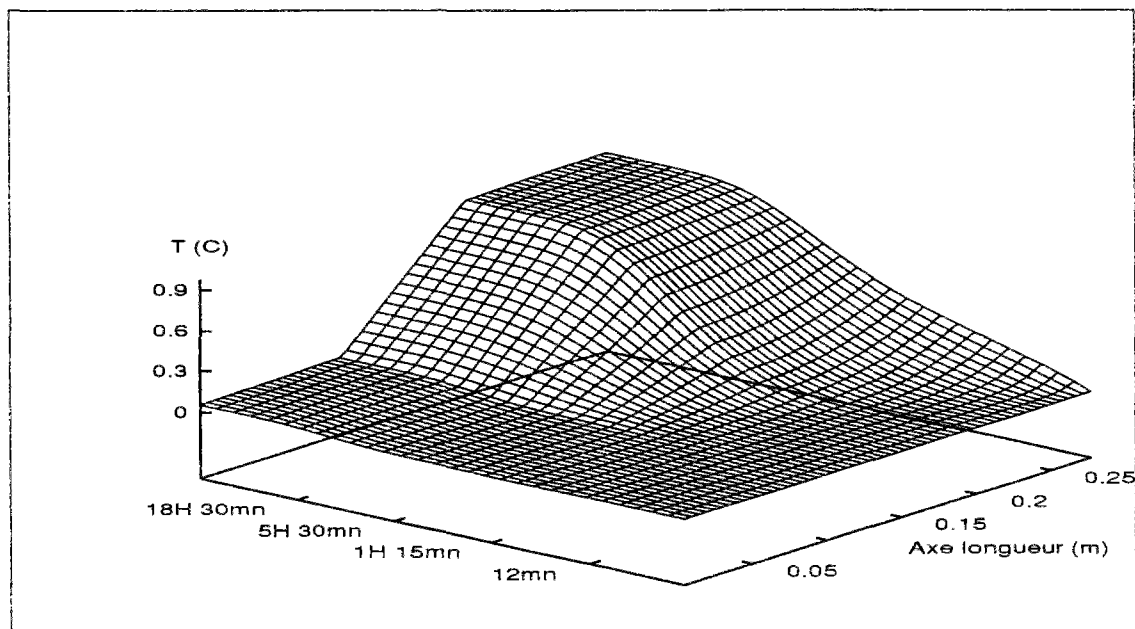


Figure 4.6: Réponse en température de la paroi à une sollicitation échelon sur son côté droit (modèle détaillé).

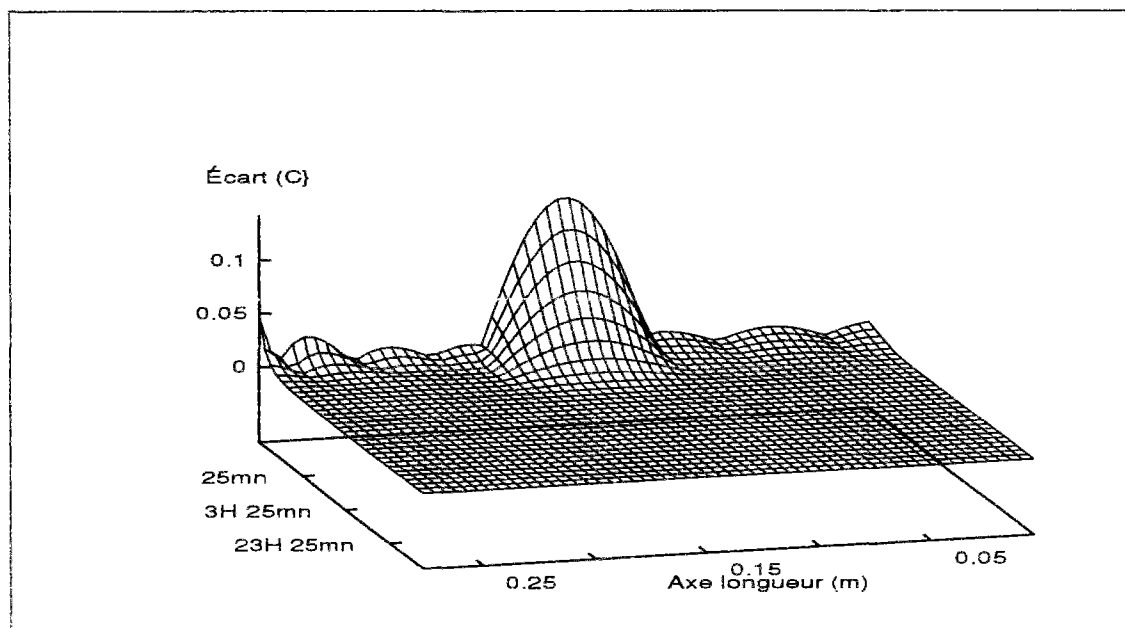


Figure 4.7: Ecart entre la réponse spatio-temporelle de (M.Ref) et celle de (M.Pri) obtenu pour $\mathcal{D}_t = [0, \infty[$.

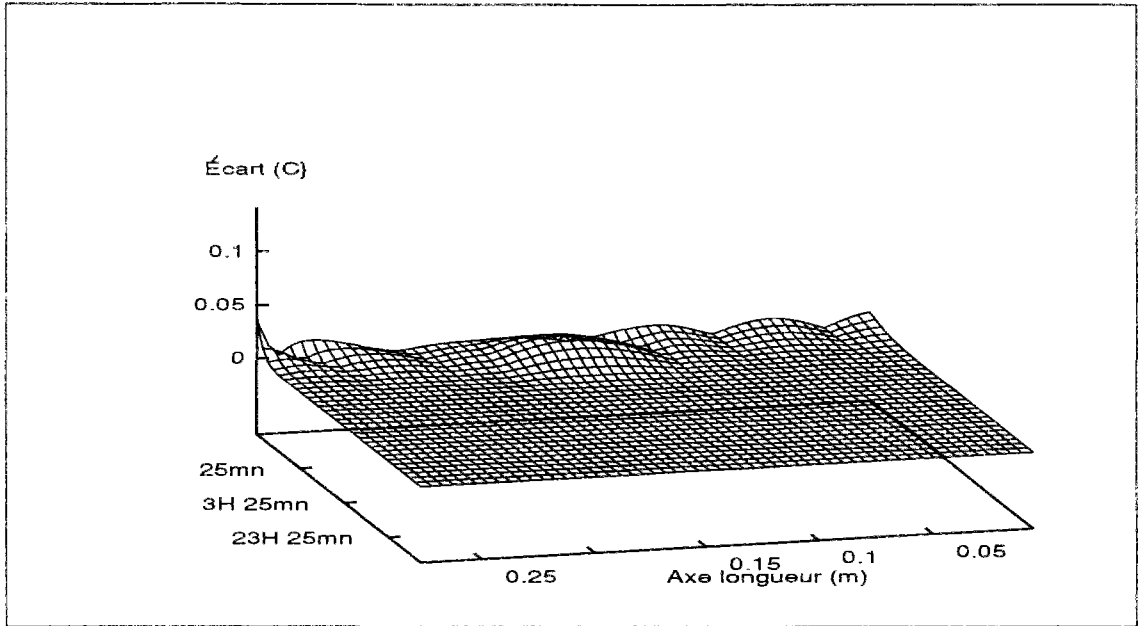


Figure 4.8: Écart entre la réponse spatio-temporelle de (M.Ref) et celle de (M.Ama) obtenu pour $\mathcal{D}_t = [0, \infty[$.

ont peu de chance d'être parmi les modes principaux. Ceci explique les erreurs importantes du (M.Pri) au niveau de l'isolant. La méthode d'amalgame modal a donc "concentré" la correction au niveau de la couche de l'isolant.

Remarque: Afin de pouvoir comparer les précisions des modèles (M.Pri) et (M.Ama), nous avons utilisé les mêmes échelles dans les figures 4.7 et 4.8.

4.3 Erreurs de réduction

Nous avons défini en §3.5 une mesure de l'erreur de réduction $\overline{\mathcal{M}}$ que nous pouvons évaluer aussi bien pour (M.Ama) que pour (M.Pri) ou le modèle de Marshall⁵ (M.Mar).

Nous avons effectué des réductions successives du (M.Ref) à l'ordre 1,2,3... Pour chacune de ces dimensions, nous avons évalué numériquement la mesure $\overline{\mathcal{M}}$ pour les trois modèles: (M.Mar), (M.Pri) et (M.Ama). Les résultats obtenus sont résumés dans la figure 4.9 et permettent de faire deux types de comparaisons

- La première consiste à comparer les courbes associées au (M.Pri) et (M.Ama). Cette comparaison permet de confirmer le travail accompli: le (M.Ama) est

⁵Rappel: Le modèle de Marshall est obtenu par simple troncature du (M.Ref) en ne gardant que les n premiers modes propres dans l'ordre décroissant des constantes de temps.

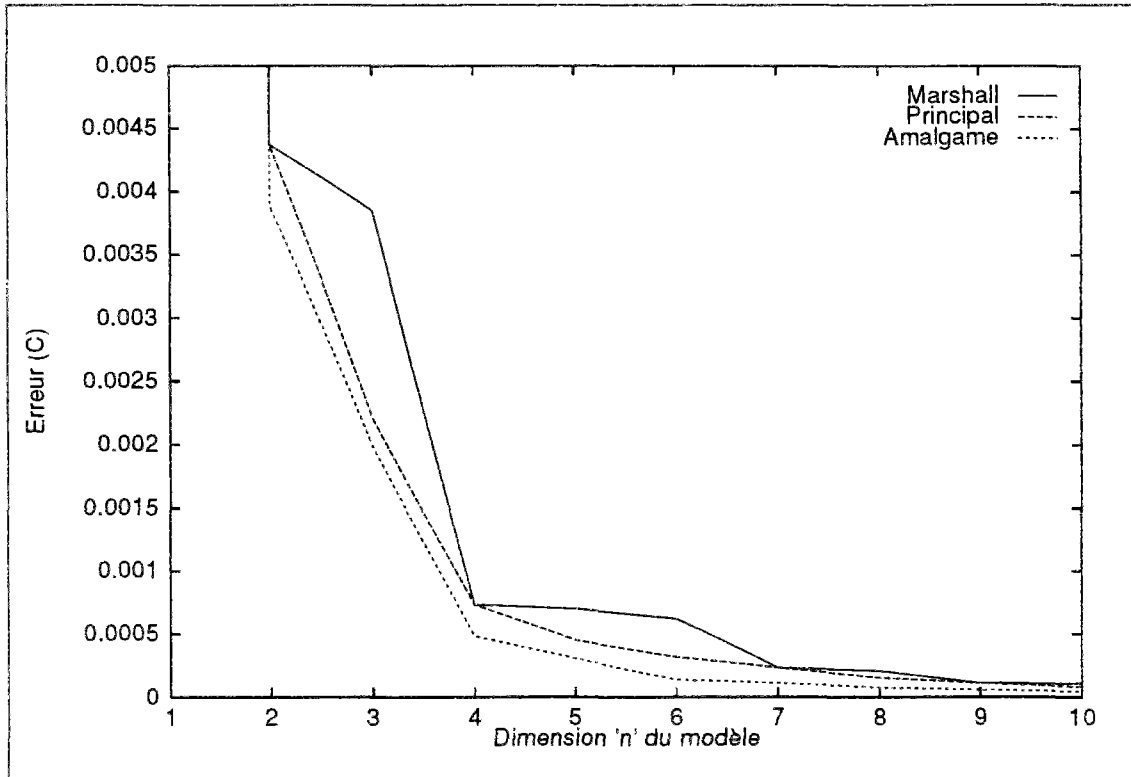


Figure 4.9: Evolution de la mesure $\overline{\mathcal{M}}$ en fonction de n .

plus précis que le (M.Pri) quel que soit le degré de la réduction.

► Du point de vue de la troncature modale, il est aussi possible de conclure que le (M.Pri) est plus précis que le (M.Mar) quelle que soit la dimension du modèle réduit. Le critère de troncature permettant le choix des modes principaux (le critère est de type énergétique) est donc plus apte à sélectionner les modes dominants. Les "paliers horizontaux" de la courbe d'erreur associée au (M.Mar) s'expliquent par le rajout de modes faibles au sens de $\mathcal{M}$. Par contre les courbes d'erreurs associées aux (M.Pri) et (M.Ama) sont plus régulières et présentent un taux de décroissance monotone lorsque n augmente.

4.4 Influence du domaine temporel $\mathcal{D}_t$

Tous les résultats établis dans la méthode de réduction par amalgame modal ont été donnés dans le cas général $\mathcal{D}_t = [t_1, t_2]$. Cette généralité nous a conduit à obtenir des expressions contenant des termes $[\Delta f(t)]_{t_2}^{t_1}$ qui désignent l'écart $[f(t_1) - f(t_2)]$. Ces termes sont réduits à 1 dans le cas où $\mathcal{D}_t = [0, \infty]$.

D'autre part, tous les modèles réduits présentent un défaut commun: les erreurs sont

globalement situées dans le début de l'évolution thermique du système. Dans le cas de la paroi tricouche étudiée ci-dessus, on peut vérifier sur la figure 4.8 que les erreurs (même faibles) sont approximativement réparties entre $t = 0$ et $t = 25mn$ ($25mn$ est de l'ordre de $10^{-2}\tau_1$). Aussi nous appliquons la méthode d'amalgame pour obtenir un modèle réduit d'ordre $n = 5$ et pour les mêmes conditions imposées ci-dessus sauf pour le domaine temporel, c'est à dire

- Transferts conductifs unidimensionnels dans la paroi.
- Plage temporelle $\mathcal{D}_t = [10^{-4}\tau_1, 10^{-2}\tau_1]$.
- Sollicitations de type échelon.
- Champ de température initial nul.

Le fait d'appliquer la méthode d'amalgame sur $\mathcal{D}_t = [10^{-4}\tau_1, 10^{-2}\tau_1]$ n'aboutit pas aux mêmes sous-espaces d'amalgame que ceux de §4.2.1 ci-dessus. On donne dans le tableau 4.4 ci-dessous la nouvelle partition de l'espace des modes propres.

sous-espace	1	2	3	4	5
composition en modes	1 , 3, 11	2	4	6 , 8, 16, 24, 32 ... (dim 33)	7 , 5, 9, 10, 12, 13 ... (dim 212)

Tableau 4.4: répartition de l'espace des modes propres en 5 sous-espaces d'amalgame lorsque $\mathcal{D}_t = [10^{-4}\tau_1, 10^{-2}\tau_1]$

Notons que cette partition se distingue de celle du tableau 4.3 sur deux points essentiels:

Le mode $V_3(M)$ devient ici mineur et se place dans le sous-espace 1, mais le coefficient $\omega_1(2)$ qui lui est associé est de 0.055, ce qui signifie qu'il modifie peu le mode principal $V_1(M)$.

Le mode $V_6(M)$ devient ici principal et récupère une information thermique de la part de 32 modes mineurs.

Ce réarrangement a une conséquence sur la précision du modèle amalgamé obtenu. Pour comparer ce dernier à celui obtenu avec $\mathcal{D}_t = [0, \infty[$ dans §4.2.3, il faut tracer l'écart (en valeur absolue) entre les champs spatio-temporels de température donnés par **(M.Ref)** et **(M.Ama)** respectivement. Cet écart est illustré par la figure 4.10 dans laquelle nous avons gardé les mêmes échelles que les figures 4.7 et 4.8. L'effet de modification du domaine temporel $\mathcal{D}_t$ peut ici être résumé comme suit:

- Les erreurs du modèle amalgamé sont encore plus atténuées aux faibles instants, c'est à dire jusqu'à $10^{-2}\tau_1$ environ.
- Il y a apparition de faibles erreurs dans la zone du transitoire (jusqu'à 6h environ sur la figure 4.10). Mais cette détérioration est insignifiante puisque l'écart est de l'ordre de $10^{-3}^\circ C$.

► L'erreur aux frontières de la paroi est nettement atténuée. Ces zones sont instantanément atteintes par la sollicitation. Il est donc logique que l'ajustement temporel proposé ici améliore la qualité du modèle réduit sur les limites de la paroi. Ainsi, le modèle réduit par amalgame permet de simuler un comportement qui s'apparente dans les premiers instants à celui d'un massif semi-infini. C'est une des difficultés traditionnelles des décompositions sous forme de série infinie qui trouve ici une solution. Notons que la connaissance précise de l'évolution de température aux frontières de la paroi permet de déduire les flux traversant ces dernières.

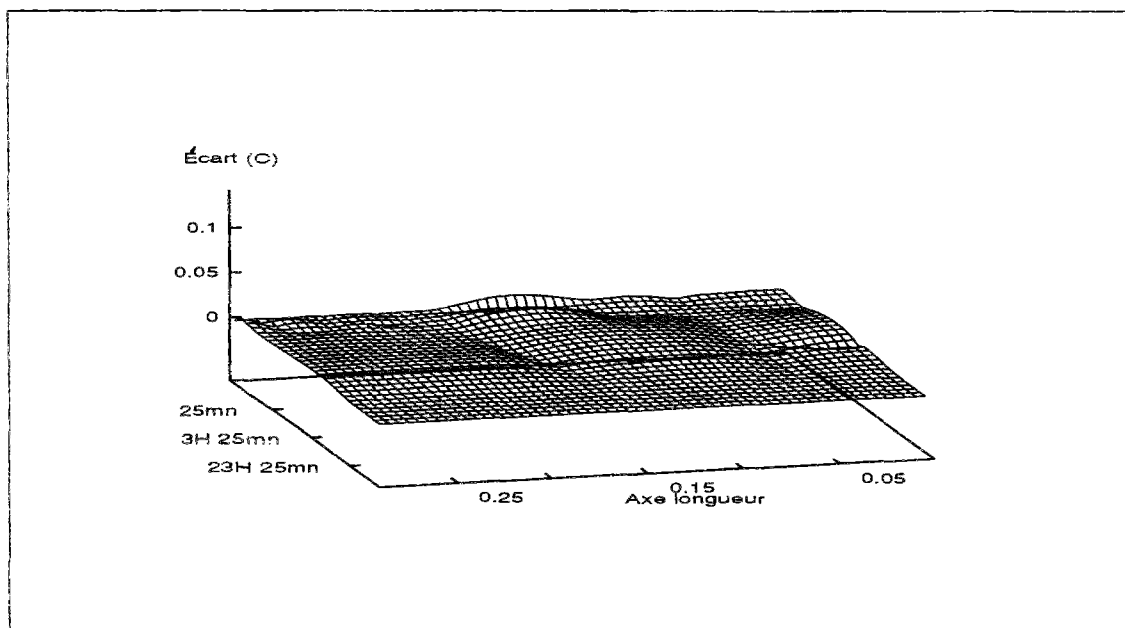


Figure 4.10: Ecart entre la réponse spatio-temporelle de (M.Ref) et celle de (M.Ama) obtenu avec $\mathcal{D}_t = [10^{-4}\tau_1, 10^{-2}\tau_1]$.

Que se passe-t-il si l'on choisit $\mathcal{D}_t$ encore plus petit? Nous avons vérifié que le résultat est similaire à celui que nous avons obtenu ci-dessus: le modèle amalgamé a une grande précision dans $\mathcal{D}_t$ et en dehors de ce domaine, il est possible de voir apparaître des erreurs.

Il ressort de cette première étude que la méthode de réduction par amalgame modal peut générer un petit modèle suffisamment précis avec possibilité de privilégier un horizon temporel.

Chapitre 5

Bâtiment bizona

5.1 Description du bâtiment

Nous abordons maintenant une structure thermique plus complexe que la paroi tricouche traitée précédemment. Il s'agit d'un bâtiment bizona qui, comme son nom l'indique, se compose de deux pièces (on dit aussi deux zones). La zone *nord* est assez inerte (isolation et petite surface vitrée) et la zone *sud* largement vitrée (15% de sa surface sud). Le mur mitoyen est constitué d'une couche de béton de 5cm seulement afin de favoriser les échanges dynamiques entre les deux zones. Une porte supposée ouverte sépare les deux zones et fait intervenir un échange d'air de $0.001m^3s^{-1}$ réciproque entre les deux zones. La figure 5.1 donne la description du bâtiment étudié.

L'intérêt de cet exemple est d'étudier le comportement de la méthode d'amalgame lorsque plusieurs composants sont thermiquement couplés et sollicités par plusieurs excitations (flux solaire et température d'air extérieur sur l'ensemble des composants de l'enveloppe du bâtiment). Mais avant de présenter les résultats de notre démarche, nous allons faire une étude préliminaire de la réduction au moyen de la troncature modale. On retrouvera alors une partie de nos conclusions sur la troncature modale (§2.2).

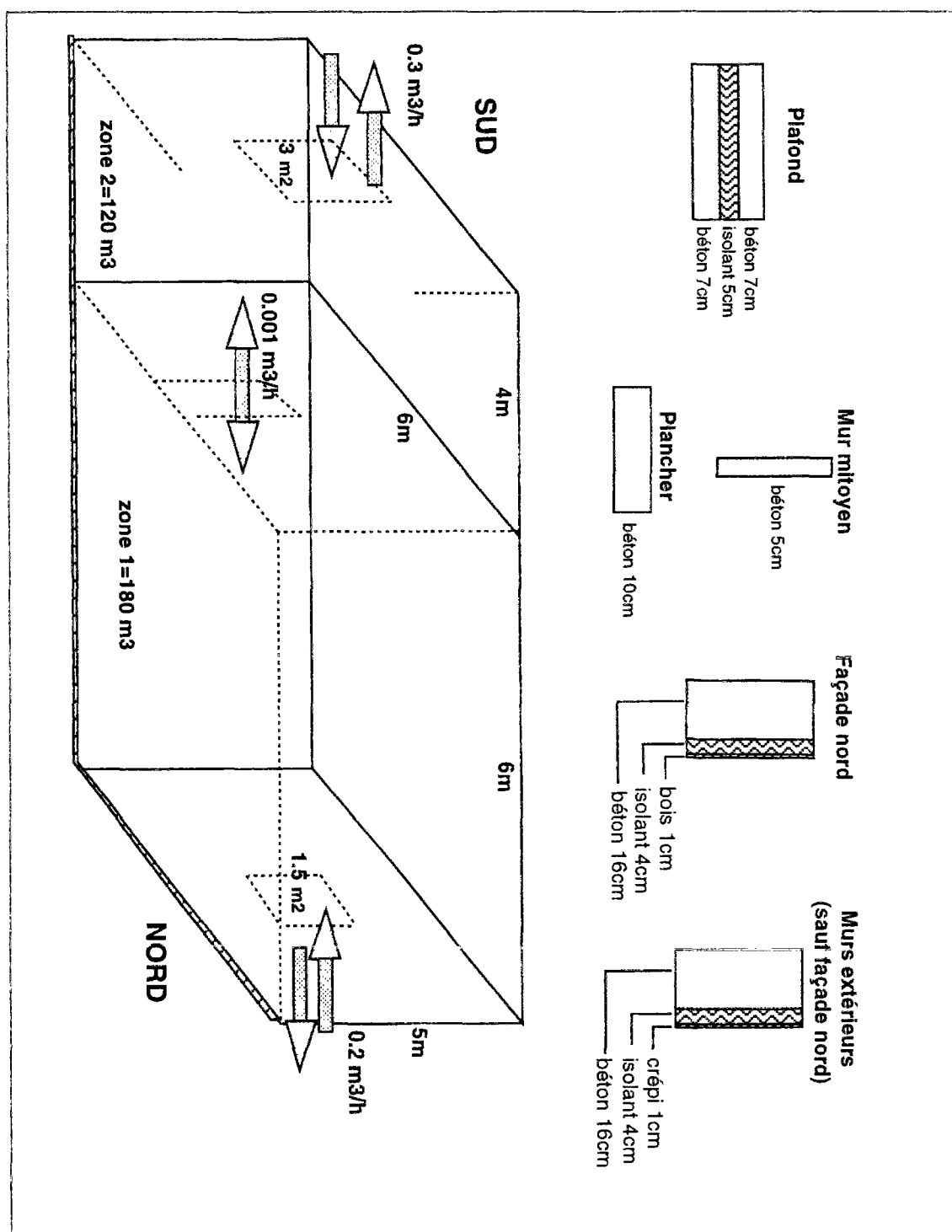


Figure 5.1: Description du bâtiment bizoné.

numéro	Heures	Minutes	Secondes	numéro	Heures	Minutes	Secondes
1	11	59	5.29	11	2	51	45.74
2	8	6	2.94	12	2	45	52.92
3	7	44	0.57	13	2	0	44.60
4	7	42	20.94	14	0	44	22.66
5	7	42	8.87	15	0	41	32.97
6	7	42	4.95	16	0	40	4.35
7	7	27	18.58	17	0	40	4.34
8	6	24	30.18	18	0	40	4.31
9	3	48	20.24	...			
10	3	45	54.01	174	0	0	10.23

Tableau 5.1: Constantes de temps du bâtiment bizone.

5.2 Etude préliminaire

Le modèle d'état modal du bâtiment bizone a été obtenu grâce au logiciel *m2m*¹ qui utilise la méthode des différences finies. Le modèle est de dimension 174, c'est à dire reposant sur 174 modes propres. Le nombre de nœuds est égal à celui des modes propres. On donne dans le tableau 5.1 quelques constantes de temps caractéristiques du bâtiment étudié. Les modes propres $V_1(M) \cdots V_{10}(M)$ sont donnés dans les figures 5.2 et 5.3. Pour faciliter la lecture des modes propres et des figures qui vont suivre, les différents composants du bâtiment ont été numérotés comme indiqué sur le tableau 5.2 ci-après. On retrouve cette numérotation sur le mode 1.

L'examen des constantes de temps et des modes propres du bâtiment permet de faire deux remarques:

- Il existe des constantes de temps numériquement voisines. Tel est le cas pour les cinq constantes de temps $\tau_3 \cdots \tau_7$. D'autre part, on trouve dans le bâtiment cinq parois identiques dans leur structure: il s'agit des composants 4,6,7,8, et 9 du tableau 5.2. Les modes propres associés à ces cinq constantes de temps, c'est à dire $V_3(M) \cdots V_7(M)$, sont spatialement localisés dans les cinq composants cités. Si les cinq constantes de temps sont numériquement voisines, il n'en est pas de même pour l'aspect spatial des modes propres associés. En fait, on peut montrer [49] que s'il y a m parois identiques alors il y a des valeurs propres multiples. L'ordre de multiplicité est $m-1$. Les modes propres associés aux valeurs propres multiples sont localisés strictement dans les parois, en particulier les volumes d'air ne sont pas excités dans ces modes (on dit que ces modes sont inobservables sur l'air). Il existe un mode dont la valeur propre est proche de la valeur multiple mais distincte. Le mode associé est cette fois observable sur l'air

¹ *m2m*, Logiciel générateur de modèles modaux de bâtiments multizones, Groupe Informatique et Systèmes Energétiques (GISE), Ecole des Mines de Paris - Ecole Nationale des Ponts et chaussées.
Auteur: G. Lefebvre.

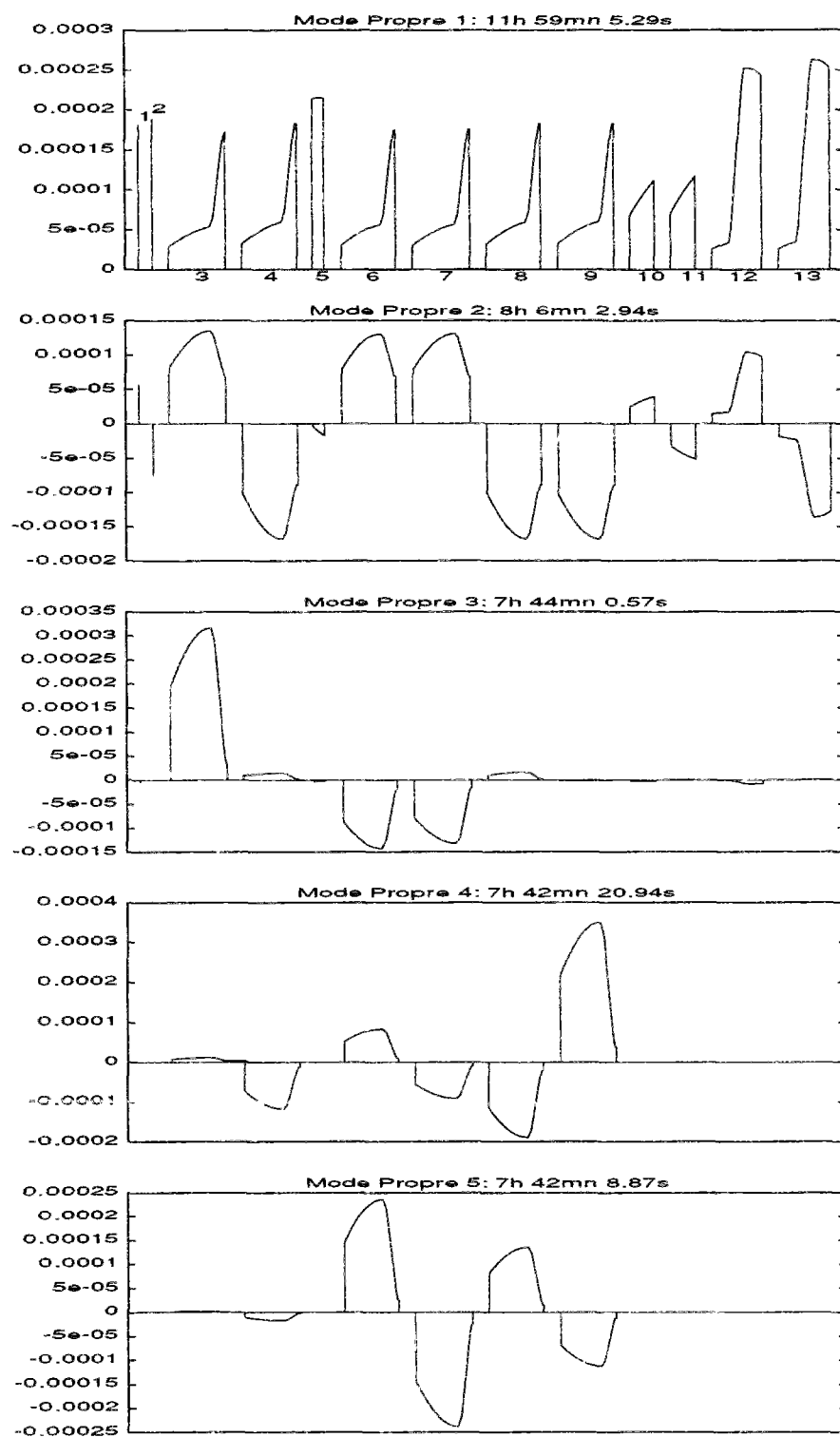


Figure 5.2: Modes propres 1 à 5 du bâtiment bizone.

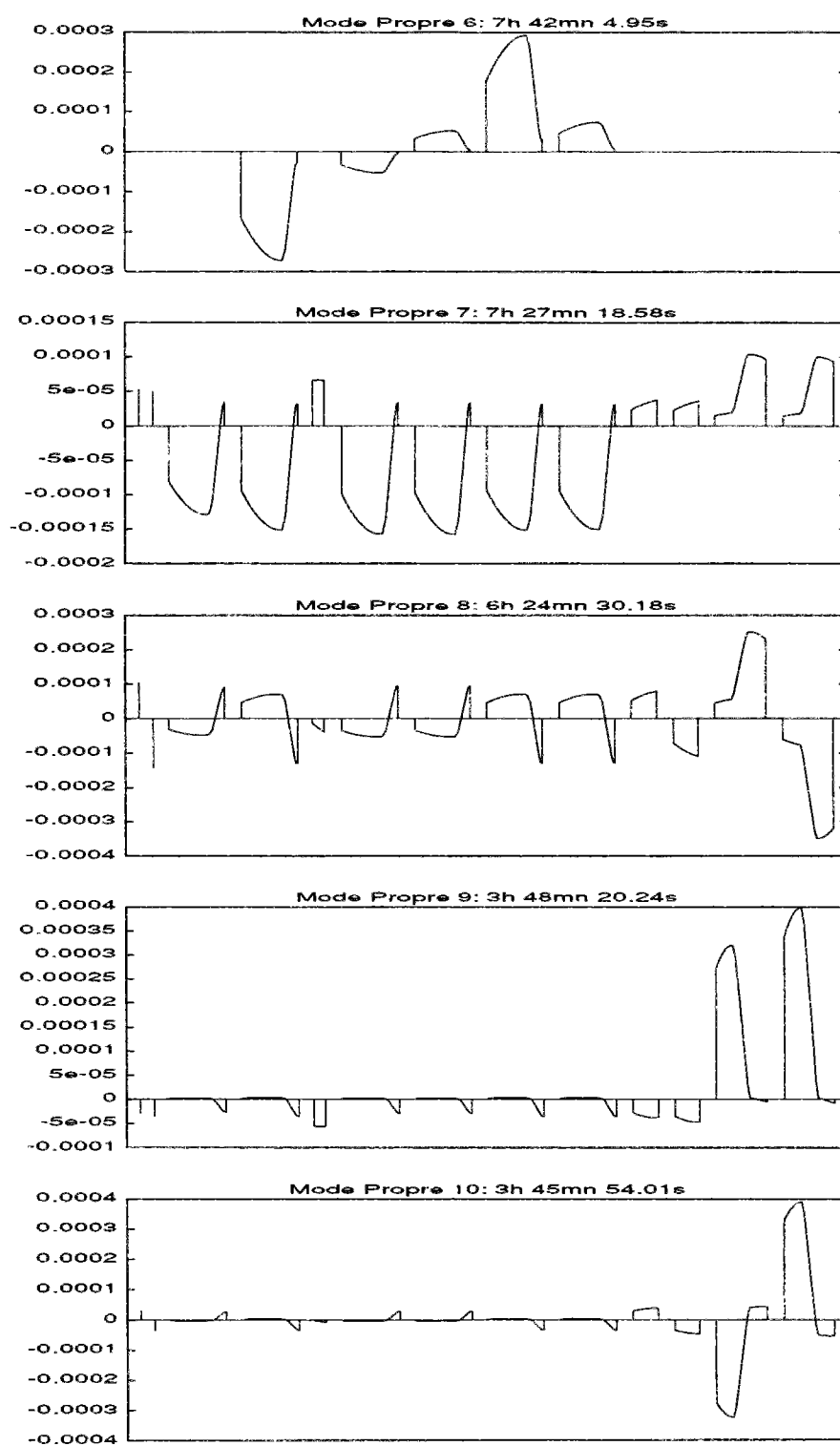


Figure 5.3: Modes propres 6 à 10 du bâtiment bizone.

Composant	Numéro
zone 1	1
zone 2	2
mur nord	3
mur sud	4
mur mitoyen	5
mur ouest 1 (chambre nord)	6
mur est 1 (chambre nord)	7
mur ouest 2 (chambre sud)	8
mur est 2 (chambre sud)	9
Plancher 1 (chambre nord)	10
Plancher 2 (chambre sud)	11
Plafond 1 (chambre nord)	12
Plafond 2 (chambre sud)	13

Tableau 5.2: Numérotation des composants du bâtiment bizonne.

et sa forme spatiale dans les parois identiques est semblable (cas du mode 7). La faible densité de notre maillage explique le fait que numériquement, $\tau_3 \cdots \tau_6$ sont différentes. La démonstration n'est rigoureusement établie que dans le cas où il n'existe pas de couplage par rayonnement entre les parois, mais seulement des échanges de type convection. Nous reprendrons ce point ci-dessous.

► La décroissance des constantes de temps n'est pas rapide. En effet, la dixième constante de temps a encore une valeur de $0.3\tau_1$.

Les deux remarques précédentes permettent de prévoir dès maintenant l'effet de la troncature modale sur le modèle. En effet, le caractère localisé des modes propres autorise généralement une troncature pour une (ou quelques) sortie donnée. Il est alors préférable de sélectionner les modes propres importants relativement au couple entrée-sortie en question (méthode de type Litz décrite au §2.2.3). Pour illustrer ceci, considérons le couple

sortie: température d'air de la zone 1 (T_{zone1}).

entrée: échelon unitaire de température extérieure (T_{ext}).

pour lequel nous mettons en œuvre une réduction à l'ordre 6 par troncature modale. Le spectre de la réponse indicielle² associé au couple entrée-sortie est donné dans la figure

²Rappel: lorsque toutes les sollicitations sont nulles sauf la $j^{\text{ème}}$ composante qui est égale à un échelon de Heaviside, la $i^{\text{ème}}$ sortie est donnée par

$$y_{ij} = S_{ij} \left(1 - \sum_{m=1}^N \frac{H_{im} B_{mj}}{S_{ij}} e^{\lambda_m t} \right)$$

et les quantités réelles $\left(\alpha_{ij}^m = \frac{H_{im} B_{mj}}{S_{ij}} \quad m = 1 \cdots N \right)$ sont les valeurs des raies associées aux différents

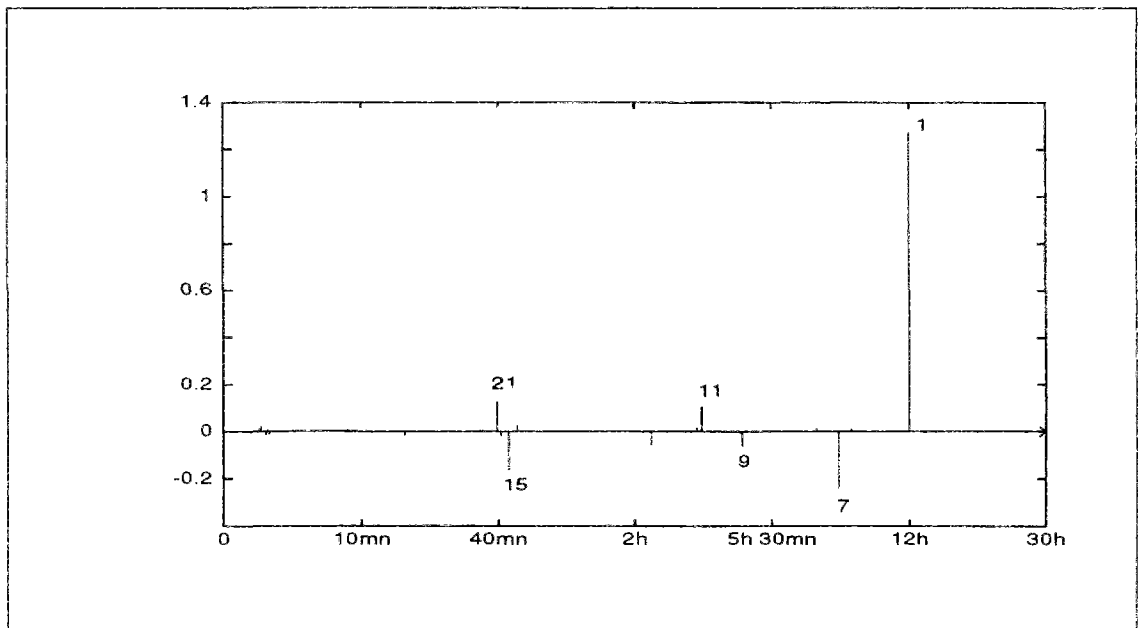


Figure 5.4: Spectre de réponse indicelle de la zone 1 à un échelon unitaire de T_{ext} .

5.4: dans l'ordre décroissant des hauteurs de raies, les modes propres importants sont **1,7,15,21,11** et **9**.

Remarquons que les modes 3,4,5 et 6 ne sont pas sélectionnés par le critère de Litz. En fait, en regardant leurs formes spatiales (figures 5.2 et 5.3), on voit qu'ils ont des amplitudes insignifiantes au niveau des deux zones. Il sont donc inobservables pour ces deux sorties. Mais l'inobservabilité n'est pas intrinsèque car elle est relative à un petit nombre de sorties. Le critère de Litz reposant sur l'observabilité (et la commandabilité), les modes 3,4,5 et 6 ne peuvent être sélectionnés. Par contre, pour d'autres critères (les critères énergétiques par exemple), on ne peut pas garantir que les modes de ce type aient systématiquement des dominances faibles.

Les différents critères de troncature esquissés au §2.2 sont testés ici. Les modes sélectionnés par les différents critères sont résumés dans le tableau 5.3. Les critères de type énergétique (dominance et dominance pondérée) donnent ici les mêmes résultats.

La figure 5.5 montre les évolutions temporelles de la sortie T_{zone1} données par les quatre méthodes de troncature que nous venons de citer. On y remarque que la troncature selon le critère de Marshall est la plus dégradée de $t = 0$ à $t = 28h$ environ: les constantes de temps 1 à 6 concernent les temps d'évolution lente. Le meilleur modèle tronqué est donné par le critère de Litz qui sélectionne des modes propres mieux répartis dans le temps et

modes propres. La méthode de Litz convient alors de considérer comme importants les n modes ayant les raies les plus grandes.

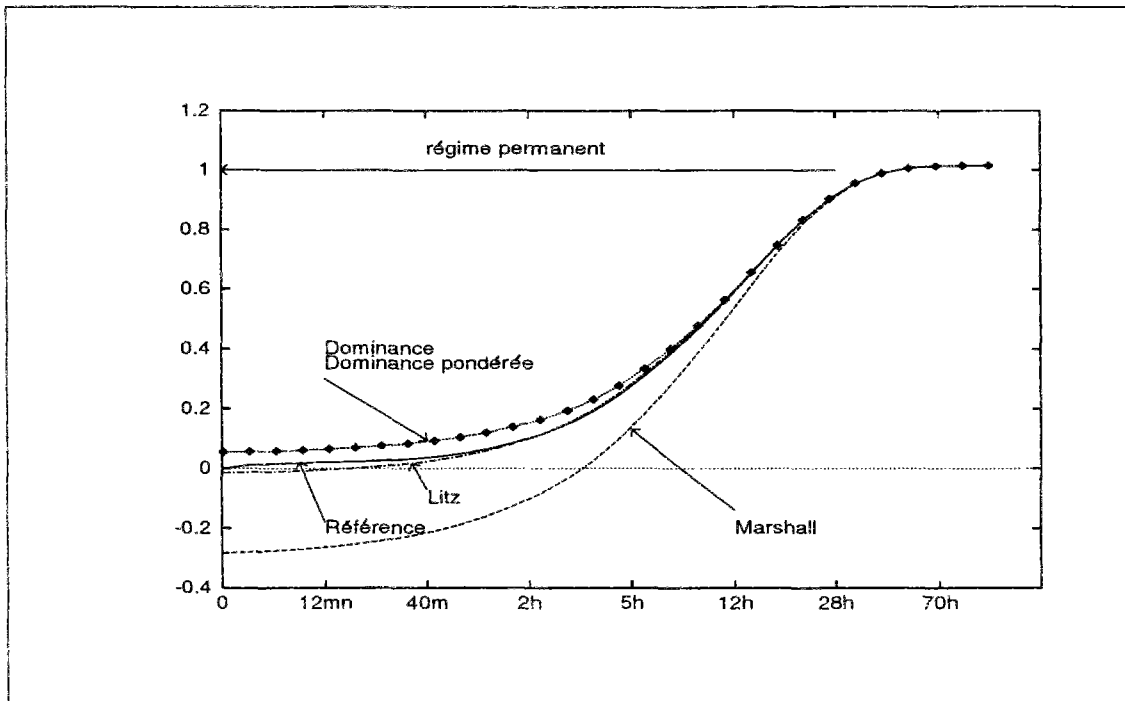


Figure 5.5: Réponse indicielle de la T_{zone1} donnée par les différents modèles tronqués.

critère	modes importants dans l'ordre de sélection
Marshall	1, 2, 3, 4, 5, 6
Litz	1, 7, 15, 21, 11, 9
Dominance	1, 7, 5, 9, 3, 6
Dominance pondérée	1, 7, 9, 5, 3, 6

Tableau 5.3: Sélection des modes par les différents critères de troncature modale

d'amplitude importante au niveau de la zone 1. Les critères énergétiques (dominance et dominance pondérée) donne une courbe biaisée mais acceptable car l'écart ponctuel avec la courbe de référence ne dépasse pas 0.05°C . L'avantage des critères énergétiques est de présenter des résultats similaires à ceux de la figure 5.5 quel que soit le point de sortie. En effet, les critères énergétiques ordonnent les modes propres d'après leurs rôles pour l'ensemble des nœuds et des sollicitations. Ceci n'est pas le cas pour le critère de Litz qui demande une sélection de modes pour chaque sortie observée.

5.3 Amalgame modal

Le modèle modal élaboré (dimension 174) va être réduit en utilisant ici la version automatique de la méthode d'amalgame. Autrement dit, nous allons rechercher le plus petit modèle réduit vérifiant la contrainte³

$$\overline{\mathcal{M}} < \overline{\mathcal{M}}_{max}$$

Le tableau 5.4 ci-dessous les dimensions n respectant quelques contraintes d'erreur.

contrainte $\overline{\mathcal{M}}_{max}$ ($^{\circ}C$)	ordre n du modèle réduit
0.1	1
0.05	2
0.01	6
0.005	9
0.001	14

Tableau 5.4: Exemples de contraintes d'erreur et de valeurs "n" correspondant.

5.3.1 Les sous-espaces d'amalgame

A titre d'exemple, considérons le cas où $\overline{\mathcal{M}}_{max} = 0.01$ $^{\circ}C$. La plus petite dimension du modèle amalgamé respectant cette contrainte d'erreur est $n = 6$ (ce choix permet ici de comparer au modèle de §5.2). Les sous-espaces d'amalgame correspondants sont donnés dans le tableau 5.5 ci-dessous où les modes principaux sont encadrés. Les sous-espaces sont donnés selon l'ordre dans lequel les modes principaux ont été trouvés. Seul le premier mode propre n'est pas modifié. Remarquons que les modes **1,5,7 et 9** ont été aussi sélectionnés par les deux critères énergétiques.

5.3.2 Les modes amalgamés

Nous nous intéressons au troisième sous-espace d'amalgame composé du mode principal $V_{11}(M)$ et des modes mineurs $V_{12}(M)$ et $V_{13}(M)$. Nous avons choisi ce sous-espace parce qu'il contient un nombre faible de modes propres ce qui facilite notre illustration. Le mode amalgamé résultant de ce sous-espace est donné par

$$\tilde{V}_3(M) = \omega_3(1) V_{11}(M) + \omega_3(2) V_{12}(M) + \omega_3(3) V_{13}(M)$$

³Pour la signification de $\overline{\mathcal{M}}$, on peut se référer au §3.5.

sous-espace	1	2	3	4	5	6
composition	1	7 2, 3	11 12, 13	9 10	15 14, 19	5 4, 6
en modes		8			20, 21, 25, 26	16, 17, 18, 22
					...	23, 24, 30, 31
						40, 70, 71, 72
						76, 77, 78, 86
	dim 1	dim 4	dim 3	dim 2	dim 145	dim 19

Tableau 5.5: répartition de l'espace des modes propres en 6 sous-espaces d'amalgame

$$\text{avec: } \omega_3(1) = 1.000 \quad \omega_3(2) = -0.161 \quad \omega_3(3) = -0.303$$

Cette combinaison linéaire est représentée dans la figure 5.6. La modification apportée au mode principal est intéressante à interpréter:

Les deux zones évoluent quasiment de la même manière. Les thermogrammes des modes 11, 12 et 13 montrent que les composants 5 (mur mitoyen), 10 (plancher 1), 11 (plancher 2), 12 (plafond 1) et 13 (plafond 2) sont fortement excités. La combinaison linéaire est réalisée de telle sorte que l'effet du mur mitoyen soit "effacé". En effet, le mur mitoyen, ne joue plus le rôle de barrière thermique si les deux zones évoluent de la même façon. Le mode amalgamé simplifie le bâtiment en "une zone équivalente couplée aux composants 10, 11, 12 et 13. Les composants 10 et 11 ont la même amplitude dans le mode amalgamé. Il en est de même pour les composants 12 et 13.

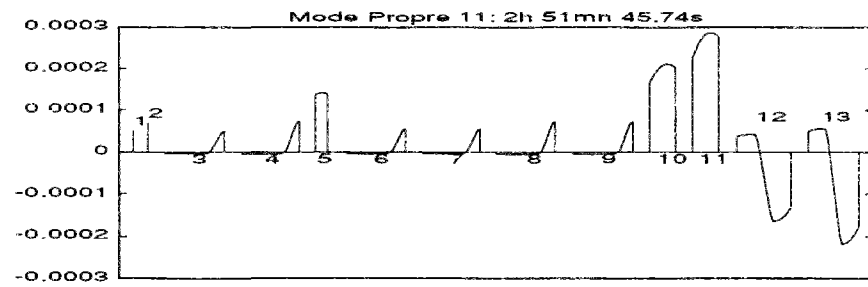
5.3.3 Simulation de l'évolution thermique

Nous nous intéressons ici à l'évolution thermique de l'ensemble des composants du bâtiment bizone. Les résultats qui suivent sont obtenus pour la situation suivante

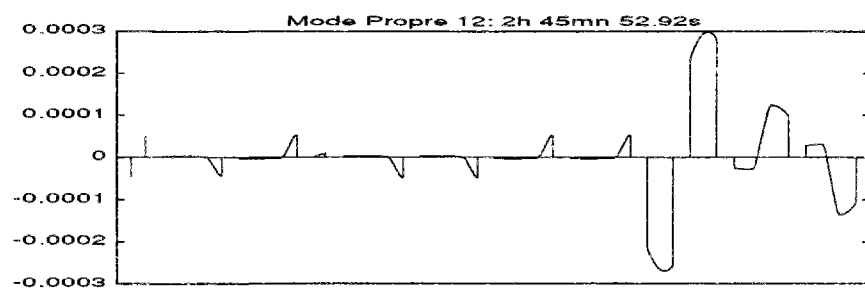
- Plage temporelle $\mathcal{D}_t = [0, \infty]$.
- La seule entrée non-nulle est T_{ext} qui est un échelon unitaire.
- Champ de température initial nul.

Dans une première étape, les modèles que nous testons sont au nombre de trois:

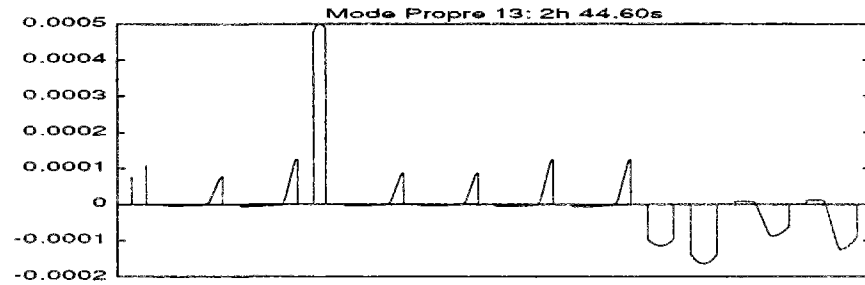
- Le modèle tronqué obtenu par le critère de Marshall, et qui repose sur les modes 1,2,3,4,5 et 6. Ce modèle est noté **(M.Mar)**



- 0.161



- 0.303



=

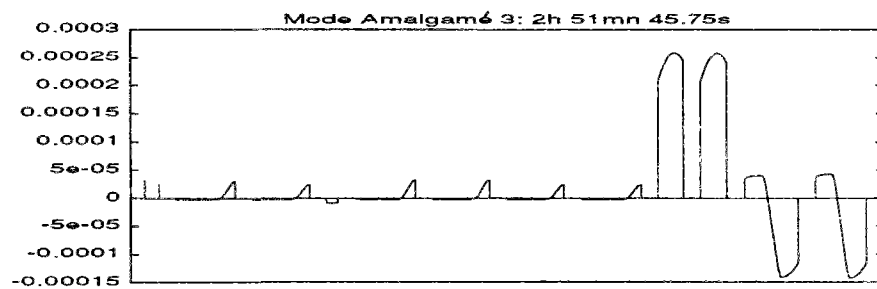


Figure 5.6: Obtention du troisième mode amalgamé

- Le modèle principal reposant sur les 6 modes principaux qui sont 1,5,7,9,11 et 15. On le note (M.Pri)
- Le modèle amalgamé d'ordre 6 qui vérifie $\overline{M} < 0.01^\circ C$. On le désignera par (M.Ama)

Comme auparavant, pour repérer les différents composants du bâtiment, nous gardons la numérotation du tableau 5.2. On place sur l'axe "x" les composants du bâtiment⁴. Sur cet axe, on laissera entre deux composants successifs un espace symbolique mais nécessaire à la lisibilité de la figure. L'axe "y" représente le temps. L'échelle utilisée est de type logarithmique. L'axe "z" (vertical) est la température.

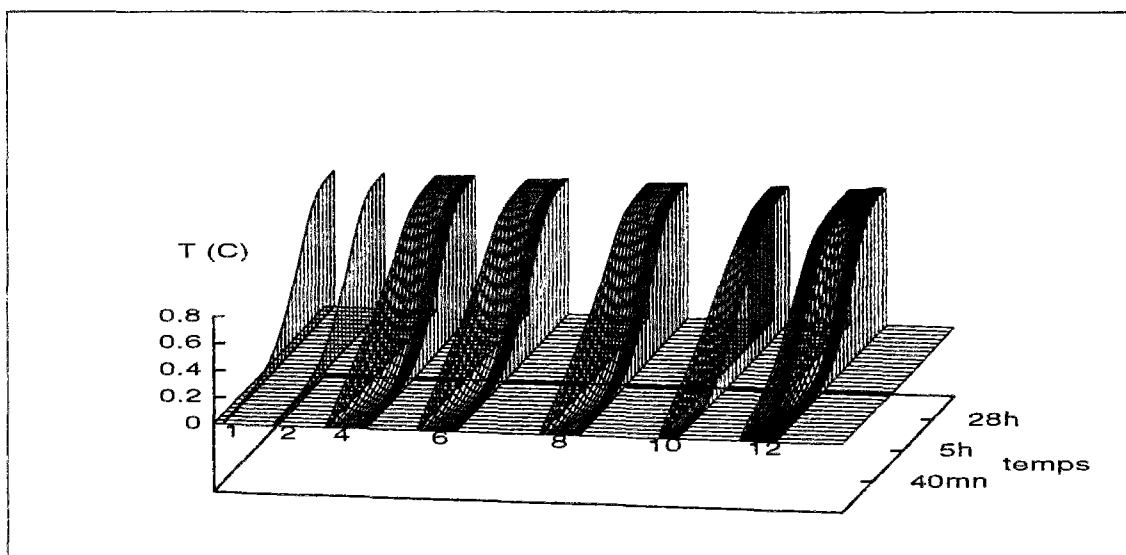


Figure 5.7: Réponse en température de quelques composants du bâtiment (M.Ref).

On peut voir sur la figure 5.7 l'évolution dans le temps des composants 1,2,4,6,8,10 et 12 du bâtiment. Le nombre de composants est volontairement limité à sept par souci de clarté de la figure. Le régime permanent est atteint par l'ensemble des composants à partir de $t = 50h$ environ.

⁴En fait sur l'axe des "x", chaque composant est figuré proportionnellement à sa dimension caractéristique (son épaisseur). Cela suffit car la diffusion thermique est ici supposée être monodimensionnelle dans les parois. Pour les zones d'air, elle sont représentées chacune par un point puisqu'on admet l'homogénéité de la température d'air.

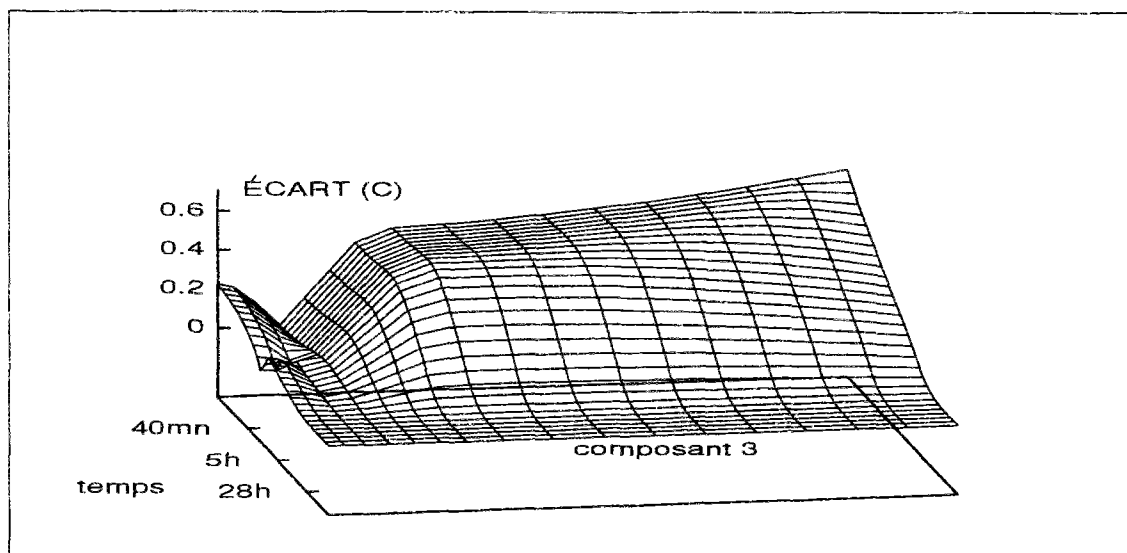


Figure 5.8: Ecart spatio-temporel de température entre (M.Mar) et (M.Ref).

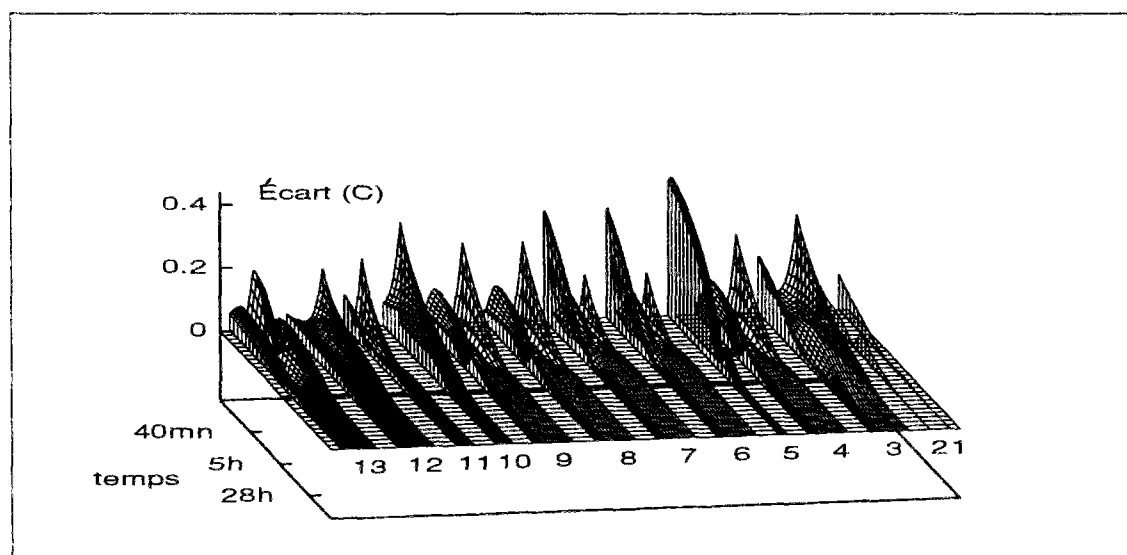


Figure 5.9: Ecart spatio-temporel de température entre (M.Pri) et (M.Ref).

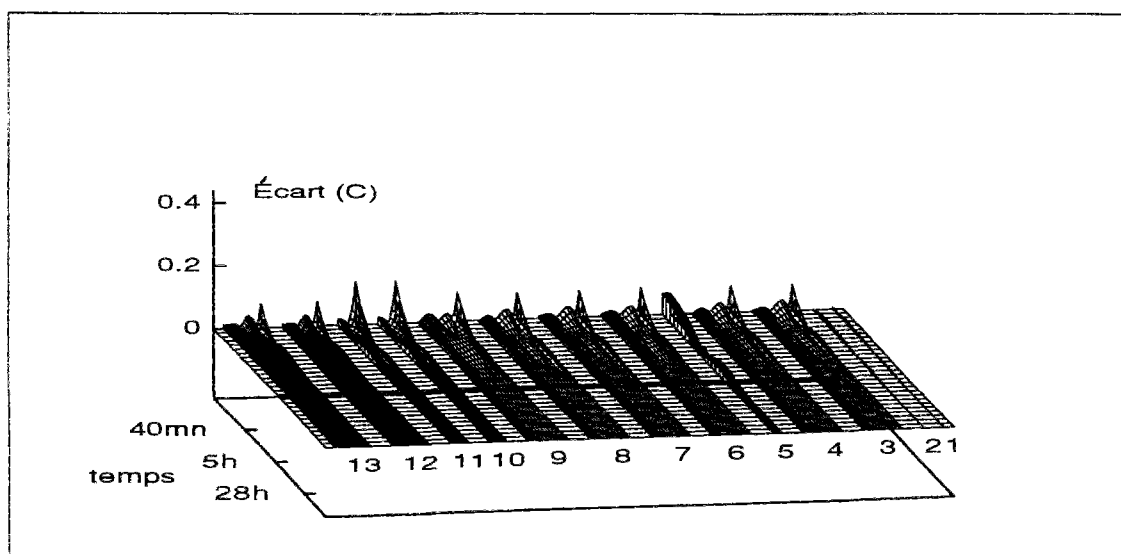


Figure 5.10: Ecart spatio-temporel de température entre (M.Ama) et (M.Ref).

Pour comparer les quatre modèles cités ci-dessus, nous allons procéder comme pour le cas de la paroi tricoche étudiée: tracer le champ écart spatio-temporel entre chaque modèle et la référence (M.Ref).

La figure 5.8 montre l'écart spatio-temporel de température entre (M.Mar) et (M.Ref) pour le composant 3 du bâtiment (mur nord). Par souci de clarté, nous avons limité cette figure à un seul composant. Les résultats sont similaires avec tous les autres composants. Les erreurs du (M.Mar) dépassent 0.6°C au début de l'évolution thermique.

L'écart entre le (M.Pri) et (M.Ref) est donné sur la figure 5.9. Le maximum d'erreur se situe sur le composant 5 (mur mitoyen) et est d'environ 0.4°C . Les erreurs restent grandes mais décroissent beaucoup plus rapidement que dans le cas précédent (critère de Marshall).

A l'exception du premier mode principal, les autres modes principaux sont corrigés par amalgame modal. Le tableau 5.5 permet de voir la répartition des modes mineurs autour des modes principaux avant la correction. La figure 5.10 montre l'écart spatio-temporel entre (M.Ama) et (M.Ref). On remarque immédiatement que le niveau général de l'erreur est beaucoup plus faible que dans les deux cas précédents: le maximum d'erreur qui se situe au début de la simulation reste inférieur⁵ à 0.1°C , et ceci dans chaque composant. Dès que le temps de simulation dépasse 10mn, l'erreur reste inférieure à 0.005°C .

⁵l'écart maximal constaté numériquement est de 0.08°C .

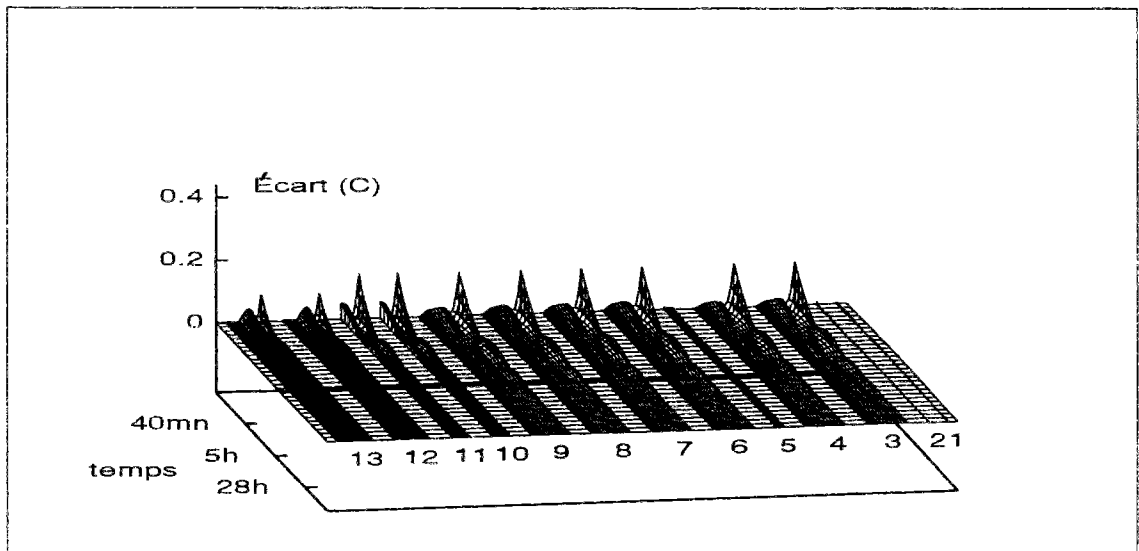


Figure 5.11: Écart spatio-temporel de température entre (M.Agr) et (M.Ref).

5.3.4 Réduction par l'agrégation

Nous avons mis en œuvre la méthode d'agrégation de Michăilescu résumée au §2.3.2. Afin de pouvoir effectuer une simple comparaison avec notre démarche, le modèle réduit recherché est d'ordre 6. La méthode d'agrégation a été mise en œuvre en considérant les sorties en température aux 174 nœuds de la discrétisation. Les modes sélectionnés par ordre décroissant de leurs contributions énergétiques sont: **1,7,9,11,13 et 15**. Sur cet exemple, les modes dominants sélectionnés par la méthode d'agrégation ne diffèrent des modes principaux que peu: c'est le mode 13 qui est retenu au lieu du mode 5 (voir tableau 5.5). L'écart spatio-temporel de température entre (M.Agr) et (M.Ref) est représenté sur la figure 5.11. On notera que les résultats sont de même qualité que ceux de la figure 5.10. L'erreur très légèrement supérieure obtenue avec l'agrégation est liée au caractère quasi-optimal de la cette dernière.

Dans une autre étape, au lieu de sélectionner les modes dominants d'après leurs contributions énergétiques, nous avons **forcé** l'algorithme d'agrégation à choisir ses modes dominants identiques à ceux des modes principaux d'amalgame (modes **1,5,7,9,11 et 15**). Les résultats obtenus sont donnés dans la figure 5.12. Contrairement à la figure 5.11, le modèle agrégé est ici pratiquement équivalent au modèle amalgamé sauf au niveau du mur mitoyen (composant 5). Il est alors possible de donner deux conclusions:

- La première concerne le choix des modes dominants dans la méthode d'agrégation. En effet lorsqu'on remplace les modes de plus grande contribution énergétique par les modes principaux, l'écart spatio-temporel de température entre (M.Agr) et (M.Ref) diminue légèrement (voir les figures 5.11 et 5.12). Le critère qui permet de choisir les modes principaux est donc pertinent.
- La méthode d'agrégation comporte une étape (étape 2 du §2.3.2) qui per-

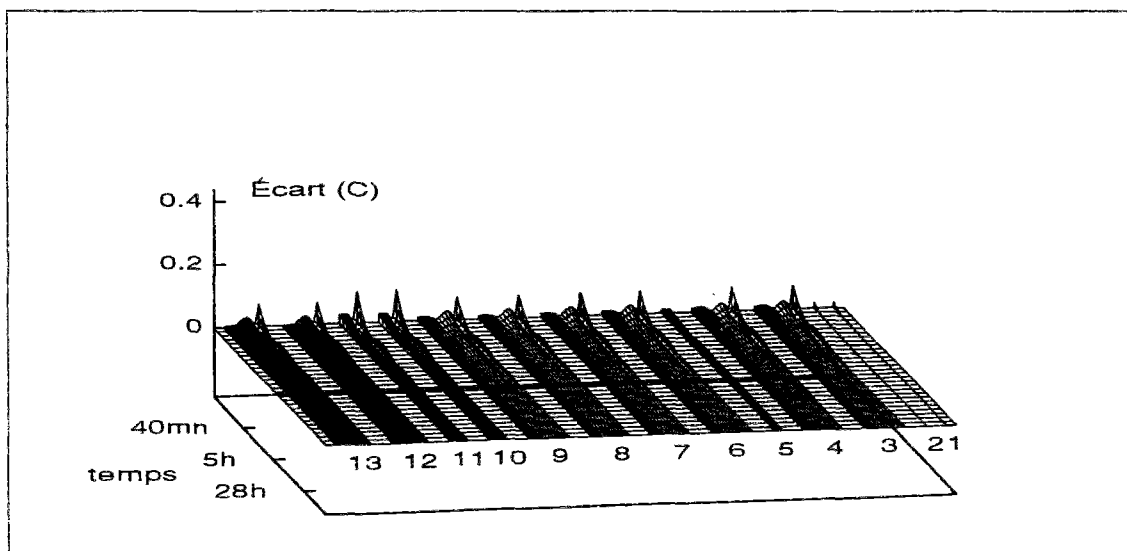


Figure 5.12: Écart spatio-temporel de température entre (M.Agr) et (M.Ref). Le (M.Agr) repose sur les modes dominants 1,5,7,9,11 et 15

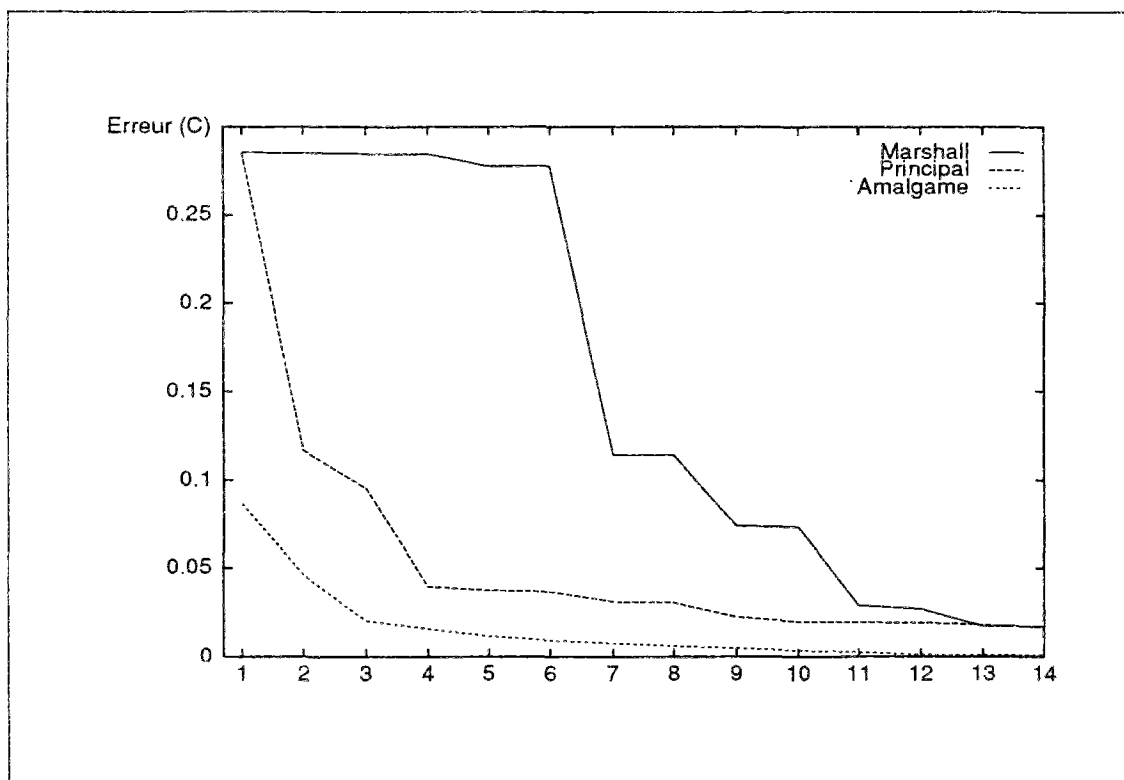
met de calculer des "pseudo-modes"⁶ au nombre de six dans cet exemple. Ces pseudo-modes s'apparentent alors à une synthèse de l'espace des modes du (M.Ref). Cette synthèse est faite de manière optimale puisqu'elle est obtenue par minimisation directe et explicite du critère d'erreur. Mais elle est faite de façon globale: c'est l'ensemble des modes propres du (M.Ref) qui sont "condensés" dans quelques pseudo-modes, ceux-ci n'étant plus orthogonaux. Au contraire, dans la méthode d'amalgame, les modes mineurs sont d'abord répartis autour des modes principaux pour former des sous-espaces, et la correction se fait indépendamment sur ces derniers. L'équivalence des résultats présentés sur les figures 5.10 et 5.12 prouvent que les démarches d'amélioration dans les deux méthodes sont aussi pertinentes. On peut conclure que la contrainte d'orthogonalité dans la méthode d'amalgame ne dégrade pas la qualité des résultats.

5.4 Erreurs de réduction

Comme pour la paroi tricouche étudiée précédemment, nous étudions maintenant l'évolution de la mesure $\overline{M}$ lorsque la taille du modèle réduit augmente $n = 1, 2, 3 \dots$. la figure 5.13 montre la fonctionnelle $\overline{M}(n)$ pour les modèles (M.Mar), (M.Pri) et (M.Ama).

La courbe $\overline{M}(n)$ associée au (M.Mar) est en forme d'escalier. Les "paliers horizontaux"

⁶Pour l'agrégation, nous utilisons la dénomination "pseudo-modes" tout en sachant qu'ils ne vérifient pas les propriétés des modes propres. En particulier, il ne sont pas orthogonaux entre eux.

Figure 5.13: Evolution de la mesure $\overline{\mathcal{M}}$ en fonction de n .

sont dûs aux modes négligeables au sens de $\overline{\mathcal{M}}$. Au contraire, chaque "rampe décroissante" correspond au rajout d'un mode principal. En réalité, ce type de courbe est aussi un moyen efficace pour repérer les modes importants d'un modèle thermique, et tout repose sur le choix de la mesure de l'erreur exploitée. Sur la courbe de $\overline{\mathcal{M}}(n)$ associée au (M.Pri), les "paliers horizontaux" disparaissent car seuls les modes qui engendrent une rampe décroissante de l'erreur sont sélectionnés. Ceci fait que, à ordre égal, le (M.Pri) est plus précis que le (M.Mar).

On peut noter sur la figure la différence entre les modèles (M.Mar) et (M.Pri). Par ailleurs, l'évolution de l'erreur du modèle amalgamé est plus régulière que celle des deux autres modèles, ce qui est un avantage. En effet, on est certain qu'une augmentation de l'ordre conduit bien à une amélioration du modèle (ceci est particulièrement vrai aux ordres faibles). Nous avons vérifié sur cet exemple que le (M.Ama) d'ordre 6 est équivalent au (M.Mar) d'ordre 50 environ.

5.5 Rôle de la pondération sur les sollicitations

La mise en œuvre pratique de la méthode d'amalgame - comme toutes les méthodes de minimisation - nécessite le choix du type des sollicitations. Pour notre part, la méthode a été implémentée pour des sollicitations échelons unitaires. Par conséquent, si la sollicitation utilisée dans la simulation n'a pas la forme d'un échelon, il n'est pas sûr que le modèle réduit soit optimal relativement à celle-ci. On va donc soumettre les différents modèles réduits à des sollicitations variables dans le temps.

Par ailleurs, le modèle de référence comporte au total neuf sollicitations:

- la température extérieure $[K]$
- deux puissances de chauffe (deux chambres) $[W]$
- six densités de flux solaire (sur les différentes façades) $[Wm^{-2}]$.

Nous avons ici un modèle thermique comprenant des sollicitations variées. Aussi, comme nous l'avons expliqué au §3.5.1, il est nécessaire d'effectuer des pondérations sur les sollicitations si l'on veut que le modèle réduit soit aussi précis quelle que soit la sollicitation. D'après la technique donnée au §3.5.1, le facteur de pondération de la sollicitation "j" est

$$P_{jj} = \bar{S}_j^{-1} \quad \text{avec} \quad \bar{S}_j = \frac{1}{q} \sum_{i=1}^q S_{ij}$$

où q est le nombre de points de la discrétisation spatiale (ici $q=174$). A titre d'exemple, les facteurs de pondération de la sollicitation T_{ext} et des deux puissances de chauffe sont donnés dans le tableau 5.6 ci-dessous.

sollicitations	T_{ext}	puiss. de chauffe 1	puiss. de chauffe 2
facteurs de pondération	0.9912	255.045	214.823

Tableau 5.6: Facteurs de pondérations des sollicitations.

Exceptionnellement pour les bâtiments, on peut aussi faire un autre calcul "simplifié" pour obtenir **un facteur de pondération commun** pour les deux sollicitations en puissance de chauffe. On se base sur la loi des déperditions thermiques du bâtiment en régime permanent. On trouve⁷ alors un facteur de pondération compris entre 150 et 300.

⁷En régime permanent, les échanges de chaleur entre l'intérieur et l'extérieur du bâtiment peuvent être approchés par [41]

$$\text{Déperditions [Watts]} = G V \Delta T$$

$G [Wm^{-3}K^{-1}]$ est le coefficient de déperdition en régime permanent dont la valeur est généralement située entre 0.5 et $1Wm^{-3}K^{-1}$. Le volume $V [m^3]$ du bâtiment est ici de $300m^3$. $\Delta T [^{\circ}C]$ est l'écart de température entre l'intérieur et l'extérieur du bâtiment.

Pour avoir un écart ΔT voisin de $1^{\circ}C$, il faut fournir une puissance de chauffe comprise entre 150 (pour

Avec les pondérations du tableau ci-dessus, nous appliquons l'amalgame modal. Nous désignons maintenant le modèle amalgamé obtenu par **(M.Ama-Pondéré)** pour le distinguer du **(M.Ama)** obtenu sans pondération. La partition de l'espace d'état modal réalisée n'est plus celle de §5.3. Elle est maintenant donnée dans le tableau 5.7.

sous-espace	1	2	3	4	5	6
composition	1	7 3,4	8 , 2, 6, 71	9 10	11 12, 13	15 14, 17
en modes		5	76, 80, 117, 129 131, 132, 133, 145 148, 149, 150, 151 152, 155, 156, 159 161, 162, ...		16, 23	18, 19, 20, 21 22, 24, 25, 26 27, 28, 29, 30 31, 32, 33, 34 35, 36 ...
	dim 1	dim 4	dim 31	dim 2	dim 5	dim 131

Tableau 5.7: répartition de l'espace des modes propres en 6 sous-espaces d'amalgame

La comparaison des thermogrammes spatiaux des modes amalgamés des **(M.Ama)** et **(M.Ama-Pondéré)** montre que seuls deux modes sont différents (les autres sont identiques ou très proches):

- $\tilde{V}_2(M)$ du **(M.Ama)**. Ce mode résulte de la combinaison des modes du sous-espace H_2 (voir tableau 5.5).
- $\tilde{V}_3(M)$ du **(M.Ama-Pondéré)** qui est la combinaison des modes du sous-espace H_3 (voir tableau 5.7).

Les thermogrammes spatiaux de ces deux modes sont donnés dans la figure 5.14. Contrairement à $\tilde{V}_2(M)$, les deux zones du bâtiment sont excités dans $\tilde{V}_3(M)$. Alors que le mode $\tilde{V}_2(M)$ du **(M.Ama)** est inobservable au niveau de la sortie considérée, $\tilde{V}_3(M)$ est très excité au niveau des deux zones. Par ailleurs sur ce dernier mode on constate que

- les deux raies associées aux zones 1 et 2 sont de signe contraire
- le mur mitoyen (composant 5 excité) présente un gradient thermique

Ces constatations expriment un effet de retard entre les deux zones (conduction dans le mur mitoyen) d'une part et un déséquilibre thermique entre elles d'autre part (une zone qui se réchauffe au dépend de l'autre).

On peut aussi illustrer cet effet de retard par l'utilisation des spectres de réponse indicielle pour le couple (sortie: T_{zone2} , entrée: puiss. de chauffe dans la zone 1). Le spectre de la figure 5.15 est obtenu à partir du **(M.Ama)** alors que celui de la figure 5.16 à partir du **(M.Ama-Pondéré)**. Sur le premier spectre, la raie 2 - celle associée à $\tilde{V}_2(M)$ du

$G=0.5Wm^{-3}K^{-1}$) et 300W (pour $G=1Wm^{-3}K^{-1}$). Si l'on veut utiliser des échelons unitaires pour les sollicitations en puissance de chauffe, il faut alors pondérer ces dernières par un facteur compris entre 150 et 300.

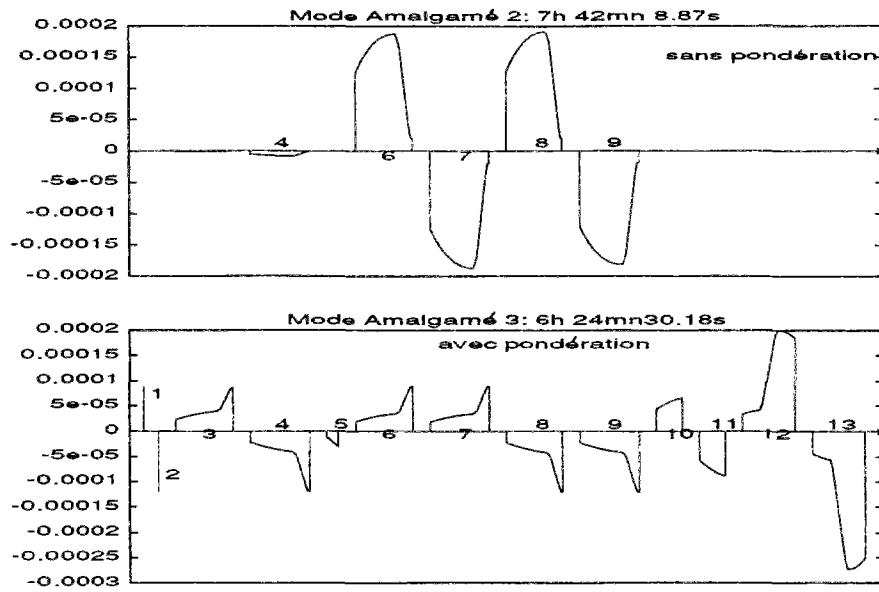


Figure 5.14: Modes amalgamés: $\tilde{V}_2(M)$ de (M.Ama) et $\tilde{V}_3(M)$ de (M.Ama-Pondéré)

(M.Ama) - est quasiment nulle. Ce mode est donc inobservable au niveau de la sortie considérée. Au contraire sur le second spectre, la raie 3 - associée à $\tilde{V}_3(M)$ du (M.Ama-Pondéré) - est très importante. Son signe est négatif: elle introduit donc un retard au niveau de la sortie observée.

A) Sollicitation en température

La première sollicitation, la température extérieure en période d'été, varie entre 20 à 30°C selon une sinusoïde décrite⁸ par la figure 5.17. La période du signal (T_{ext}) est de 24h.

Intéressons-nous à l'évolution de la température de la zone 1. Dans la figure 5.18, on compare les modèles tronqués au (M.Ref). Les modèles tronqués - tous d'ordre 6 - testés sont au nombre de cinq:

- le modèle (M.Mar) issu du critère de Marshall (§2.2.2).
- le modèle (M.Lit) issu du critère de Litz (§2.2.3) (soll: T_{ext} , sortie: T_{zone1}).
- le modèle (M.Dom1) issu du critère de dominance (§2.2.4).
- le modèle (M.Dom2) issu du critère de dominance pondéré dans le temps (§2.2.5).
- le modèle (M.Pri) reposant sur les modes principaux d'amalgame (voir tableau 5.7)

La simulation est faite pour les 3ème et 4ème jour pour éliminer l'effet de la condition initiale. L'examen de la figure 5.18 montre que le (M.Mar) est le moins précis. Le

⁸La loi de variation utilisée ici est donnée par $T_{ext} (^{\circ}C) = 25 + 5 \sin\left(\frac{2\pi}{T}t - \frac{\pi}{2}\right)$ où $T = 1$ Jour est la périodes du signal

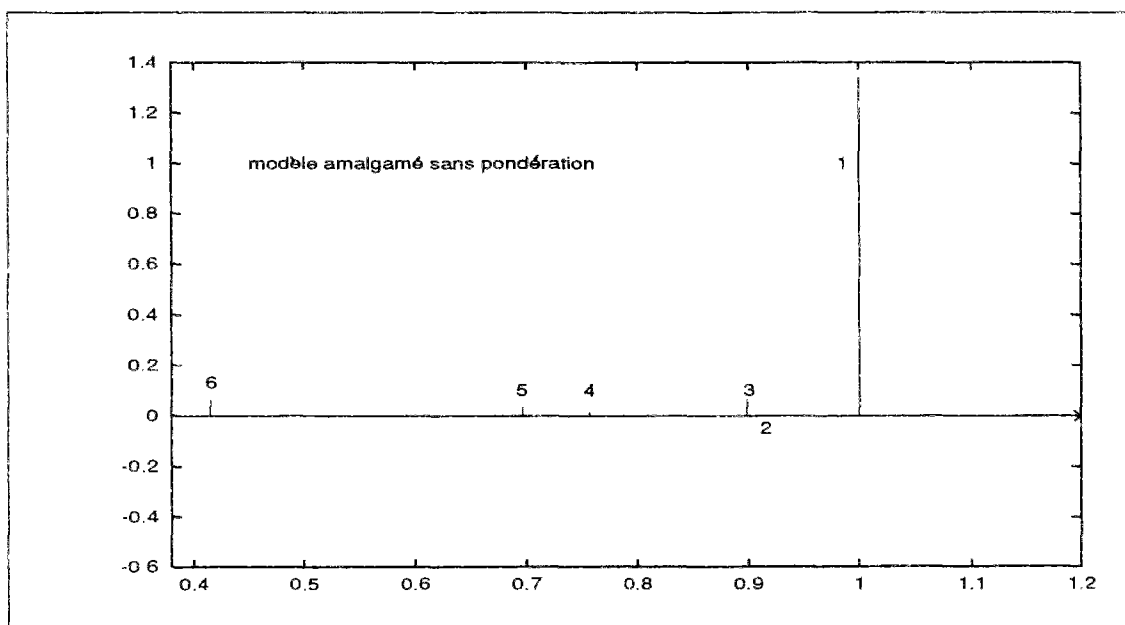


Figure 5.15: Spectre de réponse indicielle obtenu du (M.Ama) (sollicitation: puissance de chauffe dans la zone 1, réponse: T_{zone2}).

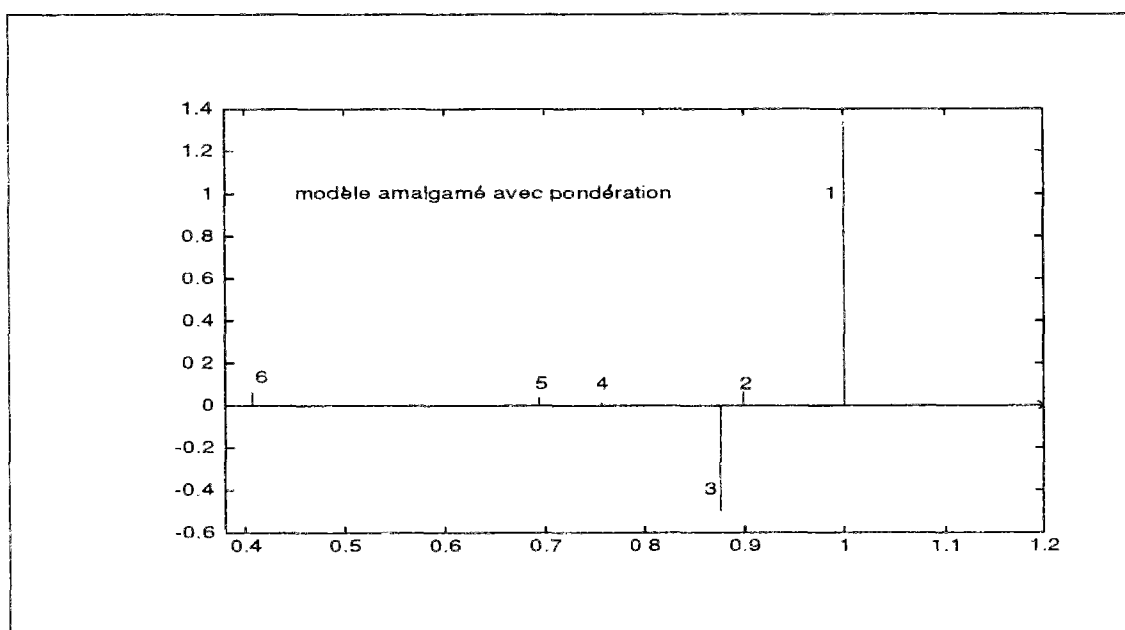


Figure 5.16: Spectre de réponse indicielle obtenu du (M.Ama-Pondéré) (sollicitation: puissance de chauffe dans la zone 1, réponse: T_{zone2}).

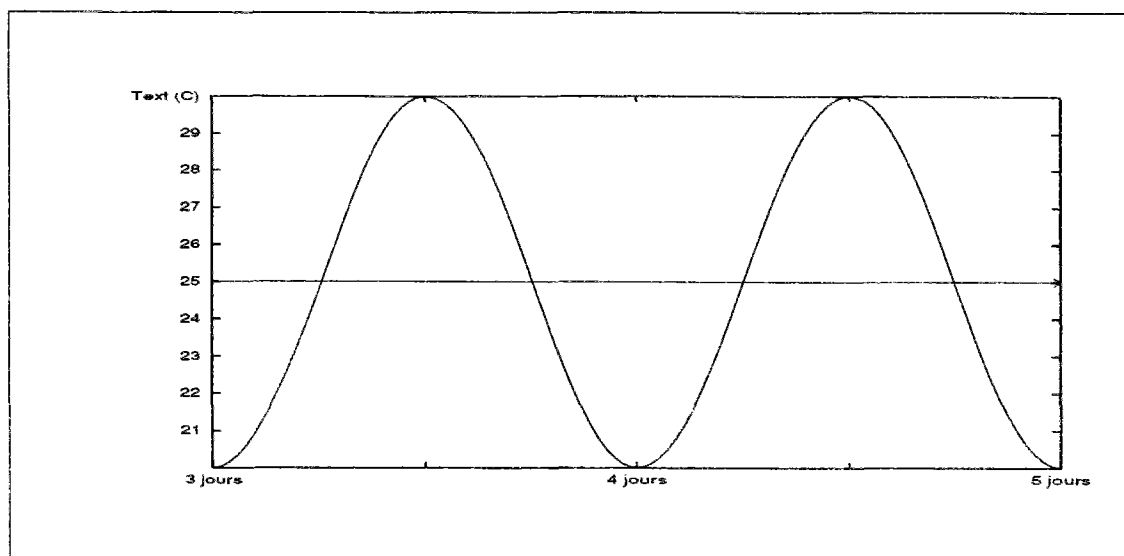


Figure 5.17: Sollicitation en température extérieure.

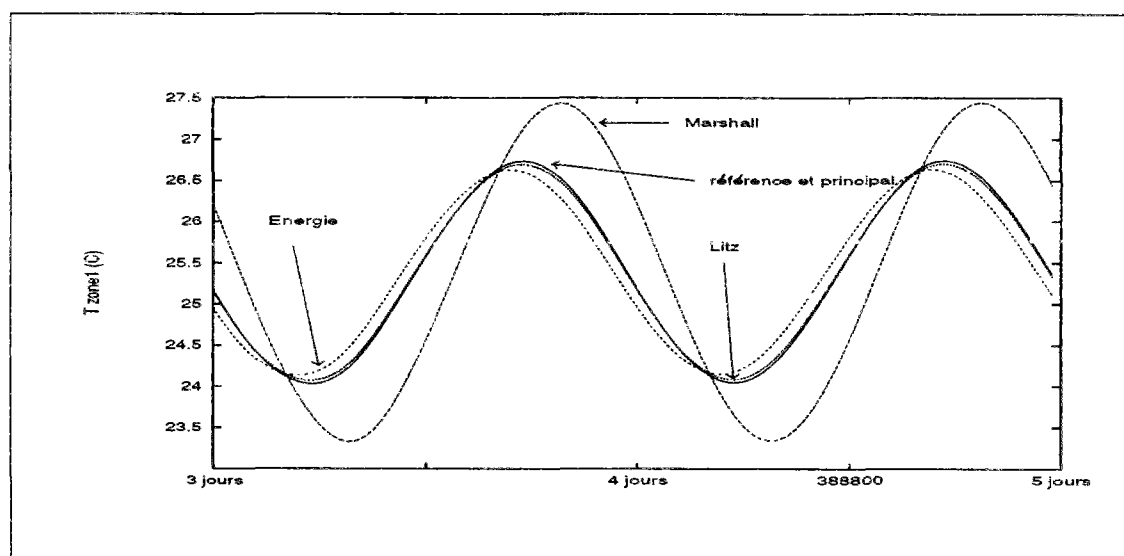


Figure 5.18: Simulation de T_{zone1} par des modèles tronqués.

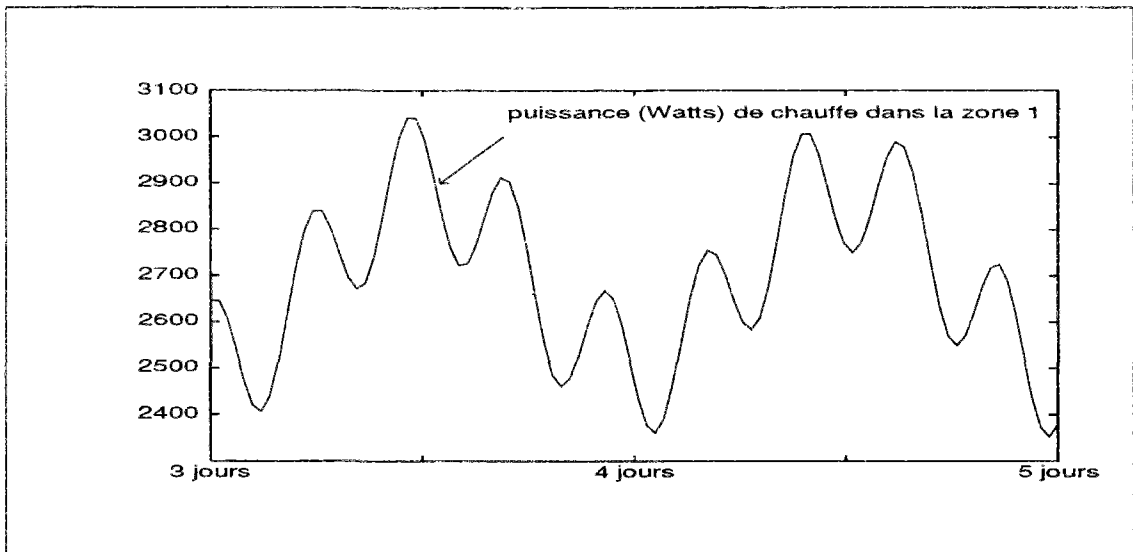


Figure 5.19: Sollicitation en puissance de chauffe de la zone 1.

(M.Lit) est très proche de la référence. Le critère de Litz a sélectionné pour le couple entrée/sortie considéré les modes: 1, 7, 15, 21, 11 et 9. Les (M.Dom1) et (M.Dom2) sont ici identiques et reposent sur les modes: 1, 7, 5, 9, 3, et 6. Sur le figure 5.18, la courbe de réponse associée à ces deux critères est peu biaisée. En ce qui concerne le (M.Pri), il est quasiment équivalent au (M.Ref).

Le (M.Pri) étant équivalent au (M.Ref), nous ne montrons pas ici la simulation effectuée à partir du (M.Ama-Pondéré). En effet ce dernier ne peut être que plus précis que le (M.Pri). Notons que le (M.Agr) donne aussi d'excellents résultats.

B) Sollicitation en puissance de chauffe

La deuxième sollicitation est la puissance de chauffe dans la zone 1 en période d'hiver. La loi de variation - volontairement choisie⁹ - de la sollicitation est donnée dans la figure 5.19.

Avec le (M.Ama-Pondéré), on simule l'évolution T_{zone2} sous l'effet de la puissance de chauffe dans la zone 1. Les résultats obtenus à partir des (M.Ama-Pondéré) et (M.Agr-Pondéré)¹⁰ sont donnés dans la figure 5.20. Ces résultats sont très proches de la référence. Notons que (M.Ama-Pondéré) est un peu plus précis que le (M.Agr-Pondéré).

⁹Cette loi est donnée par une superpositions de deux signaux (fonctions sinus et cosinus)

$$Puissance (Watts) = 2700 \sin \left(\frac{2\pi}{T_1} t - \frac{\pi}{2} \right) + 150 \cos \left(\frac{2\pi}{T_2} t \right)$$

où $T_1 = 1\text{Jour}$ et $T_2 \simeq 5h\ 30mn$ sont les périodes des deux signaux.

¹⁰le modèle agrégé obtenu après pondération sur les sollicitations en puissance de chauffe.

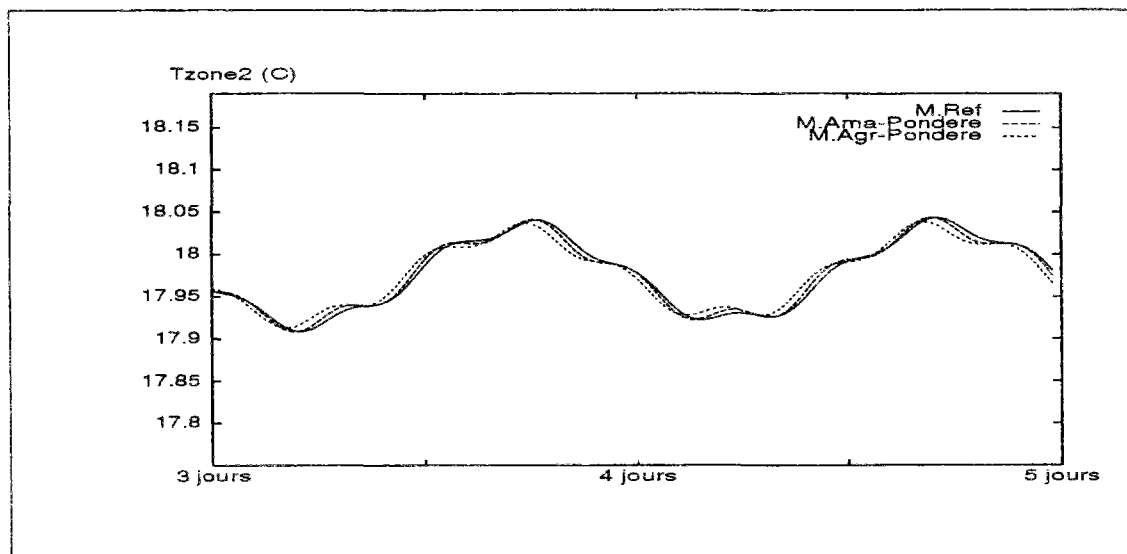


Figure 5.20: Simulation de T_{zone2} par les (M.Ama-Pondéré) et (M.Agr-Pondéré)

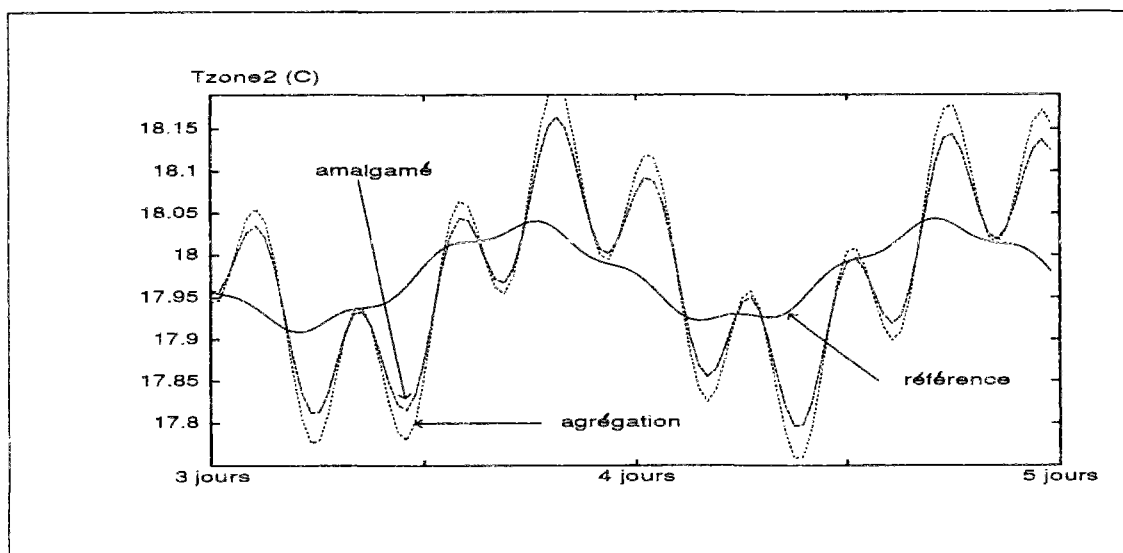


Figure 5.21: Simulation T_{zone2} par les (M.Ama) et (M.Agr).

Pour montrer l'intérêt de la pondération sur les sollicitations, on donne dans la figure 5.21 les résultats obtenus à partir du **(M.Ama)** (sans pondérations). Les résultats ne sont guère surprenants compte tenu de l'analyse précédente sur les modes amalgamés:

- le **(M.Ama)** considère que les deux zones évoluent de la même manière en éliminant le rôle du mur mitoyen.
- le **(M.Ama-Pondéré)** retrouve un mode amalgamé qui crée un retard entre les deux zones.

Chapitre 6

Pont thermique

Parmi les singularités thermiques que l'on rencontre dans un bâtiment, il y a celles qui sont liées à une discontinuité de l'isolation. C'est notamment le cas pour les liaisons plancher/mur, refends/mur ... où, pour des raisons pratiques (surtout la solidité de structure), la discontinuité de l'isolation est difficile à éviter. Dans ces parties, les échanges ne sont plus monodimensionnels mais bi, voire tridimensionnels.

Ce type de singularités sont appelées "ponts thermiques" car il favorisent les échanges (fuites) thermiques entre l'intérieur du bâtiment et le milieu extérieur. La prise en compte des ponts thermiques dans la modélisation du comportement thermique d'un bâtiment ne peut qu'augmenter la fiabilité des modèles. Le système étudié ici a déjà été traité par Salgon [58] dans le but d'appliquer les outils de l'analyse modale (spectres, analyse des modes, réduction ...). Il s'agit d'une liaison plancher-terrasse-mur et est représenté sur la figure 6.1. Le domaine étudié est interrompu à 60cm de l'angle intérieur, distance à laquelle on peut considérer les transferts comme monodimensionnels et par conséquent imposer une condition de flux nul sur cette frontière.

6.1 Modèle de référence

Les échanges thermiques sont ici bidimensionnels (2D). Pour ce système thermique, nous avons élaboré¹ un modèle modal de dimension 311 (311 nœuds et 311 modes propres). Les sollicitations sont au nombre de deux:

- ▶ Température extérieure (T_{ext})
- ▶ Température intérieure (T_{int}).

¹Le modèle de départ est obtenu avec la technique des éléments finis: 266 éléments et 311 nœuds. Le modèle modal est ensuite obtenu par transposition des matrices éléments finis dans l'espace d'état modal (technique exposée au chapitre 1).

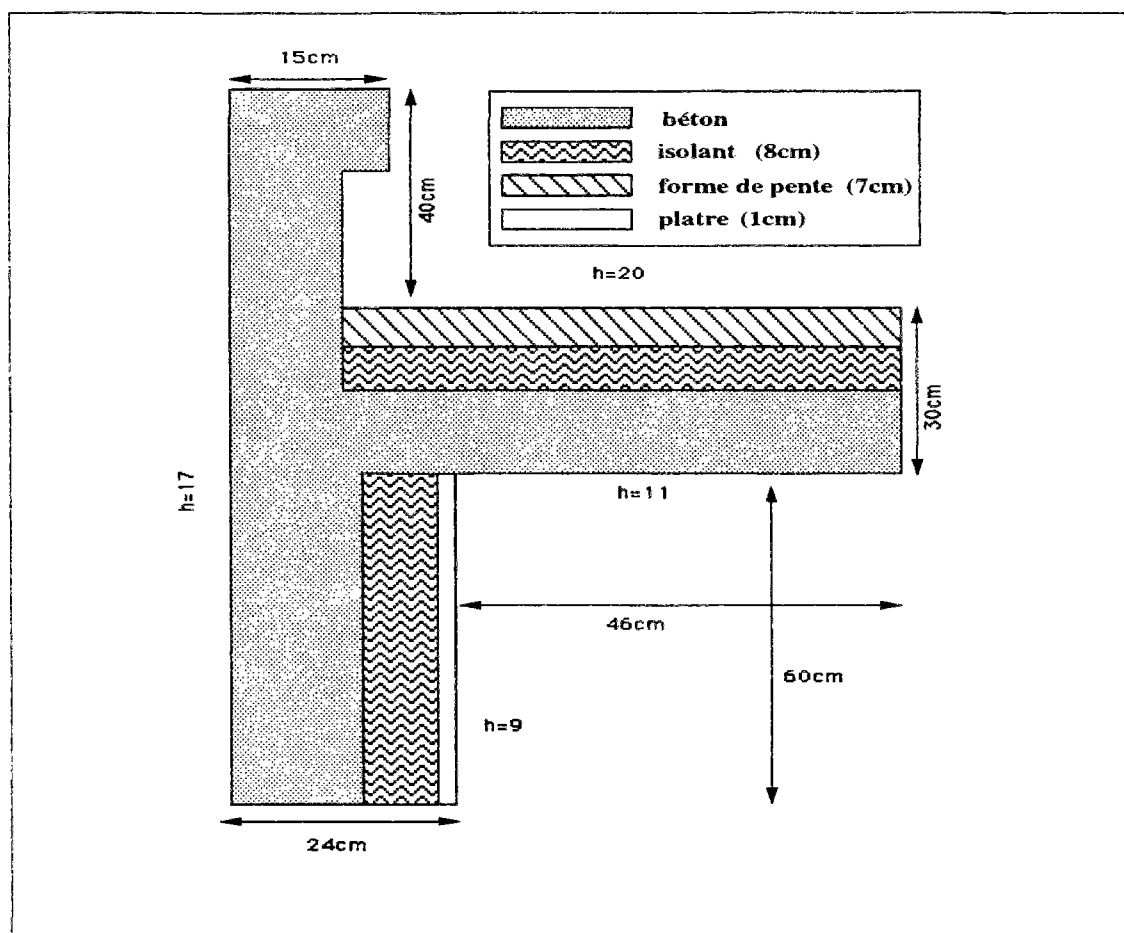


Figure 6.1: Description du pont thermique.

Les valeurs des paramètres thermophysiques des différents matériaux contenus dans le pont thermique sont donnés dans le tableau 6.1. La constante de temps principale est voisine de $\tau_1 = 10h$. La dernière constante de temps calculée est $\tau_{311} = 19.04s$. Le tableau 6.2 donne quelques constantes de temps du pont thermique.

matériaux	béton	isolant	Forme de pente	Platre
conductivité ($Wm^{-1}K^{-1}$)	1.75	0.041	1.15	0.35
masse vol. (kgm^{-3})	2450	30	1800	800
C_p ($Jkg^{-1}K^{-1}$)	920	920	845	1045

Tableau 6.1: Paramètres thermophysiques.

numéro	Heure	Minute	Seconde	numéro	Heure	Minute	Seconde
1	9	59	14.67	11	1	30	9.79
2	7	29	23.70	12	1	13	48.89
3	6	9	24.64	13	1	12	42.11
4	4	0	31.85	14	1	4	47.19
5	3	3	22.69	15	0	55	49.11
6	2	12	56.90	16	0	51	49.21
7	2	1	9.61	17	0	47	12.48
8	1	57	17.84	18	0	40	32.00
9	1	50	19.19	...			
10	1	37	48.10	311	0	0	19.04

Tableau 6.2: Constantes de temps du pont thermique.

6.2 Simulation avec le modèle de référence

L'évolution thermique de la structure ne peut pas être donnée de façon continue dans le temps car la diffusion thermique est ici bidimensionnelle. On se contente de donner les thermogrammes spatiaux de température à trois instants différents

- ▶ $t_1 = 6mn$
- ▶ $t_2 = 30mn$
- ▶ $t_3 = 2h$

Les résultats qui suivent ont été obtenus pour les conditions suivantes:

- ▶ Plage temporelle $\mathcal{D}_t = [0, \infty[$.
- ▶ Champ de température initial nul.
- ▶ Sollicitations échelons unitaires.

On donne pour chacun des trois instants t_1, t_2 et t_3 deux thermogrammes de température obtenus à partir du **(M.Ref)** et correspondant à

- ▶ T_{ext} est une échelon unitaire et $T_{int} = 0$
- ▶ T_{int} est une échelon unitaire et $T_{ext} = 0$

Ces six thermogrammes sont illustrés par les figures 6.2 à 6.7. On remarque que jusqu'à l'instant t_3 , la zone intérieure (respectivement extérieure) n'est pas concernée par la perturbation extérieure (respectivement intérieure). On peut aussi constater sur les figures 6.6 et 6.7 le gradient thermique engendré par la rupture de l'isolation.

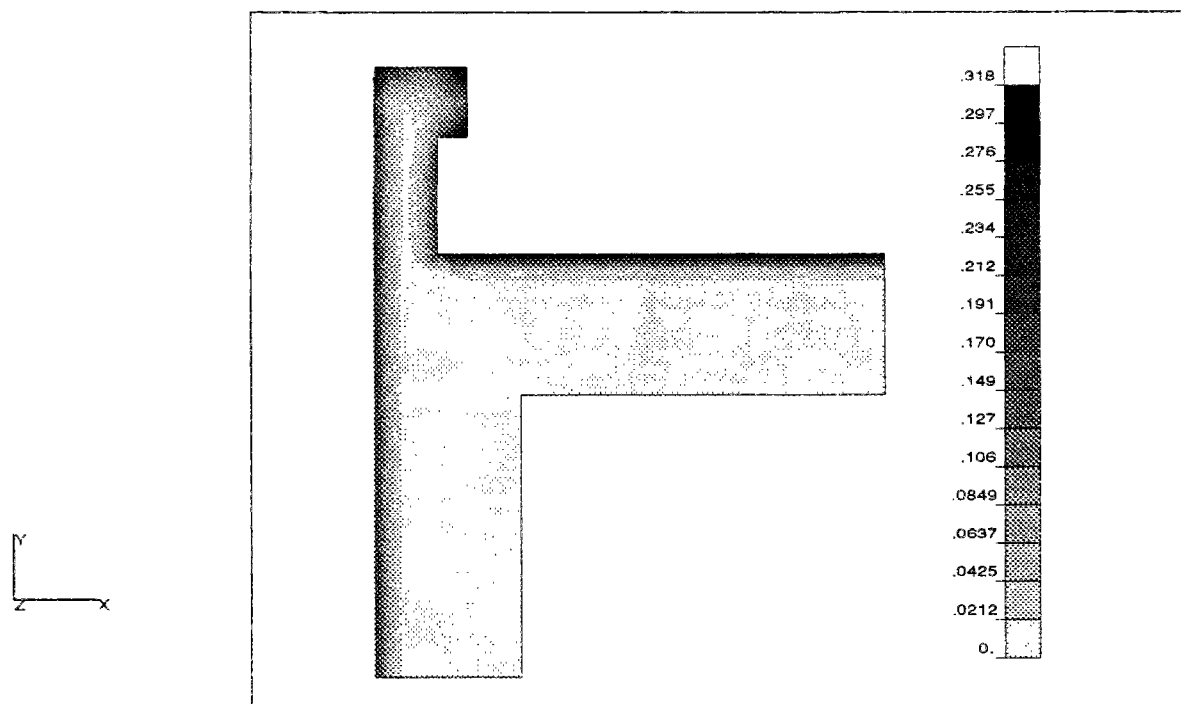


Figure 6.2: Sollicitation T_{ext} : champ de température (en $^{\circ}C$) à t_1 .

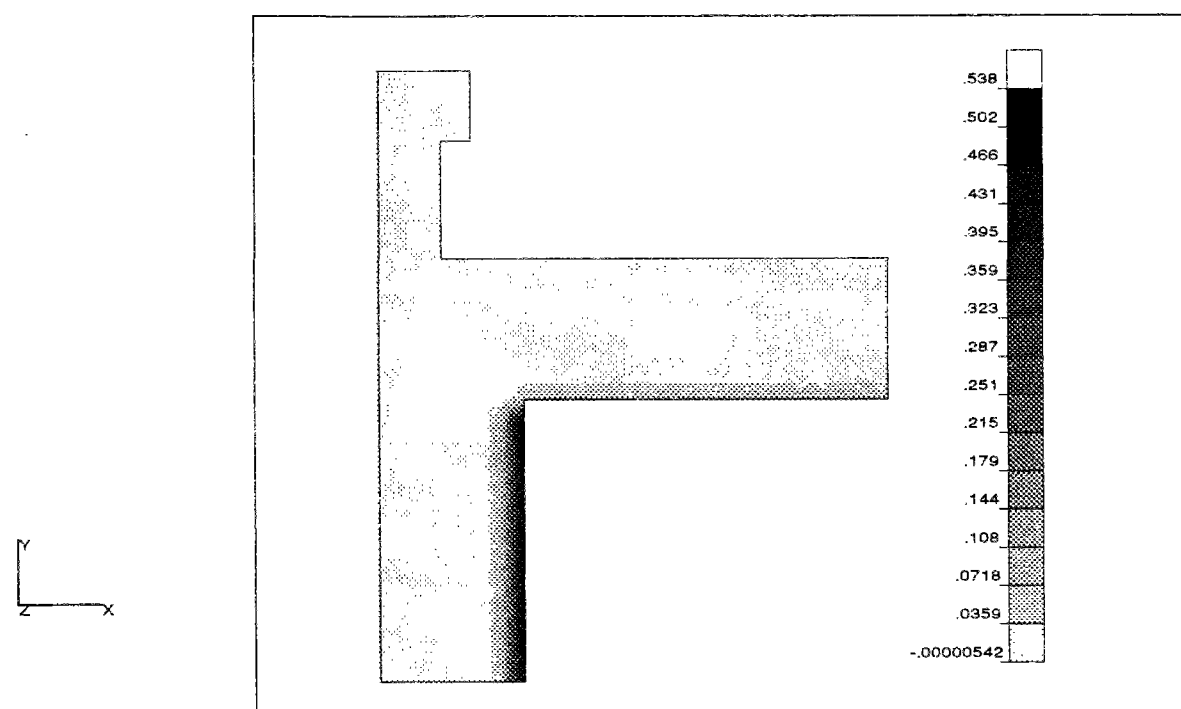


Figure 6.3: Sollicitation T_{int} : champ de température (en $^{\circ}C$) à t_1 .

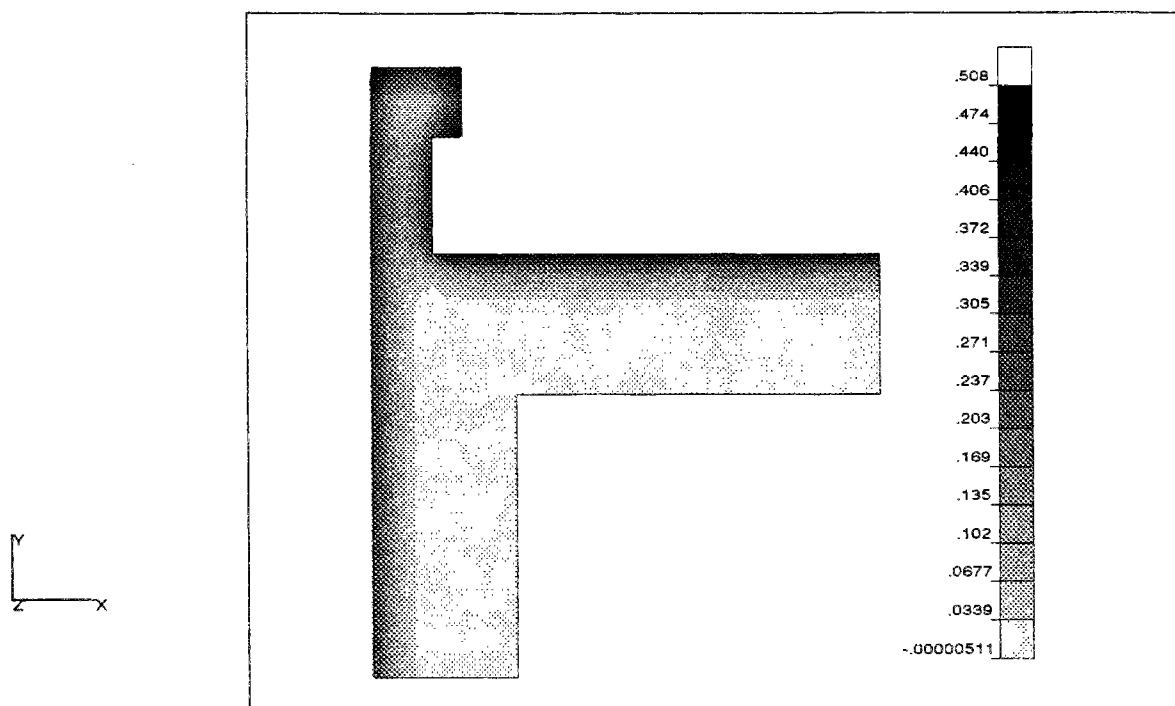


Figure 6.4: Sollicitation T_{ext} : champ de température (en $^{\circ}C$) à $t_2=30mn$.

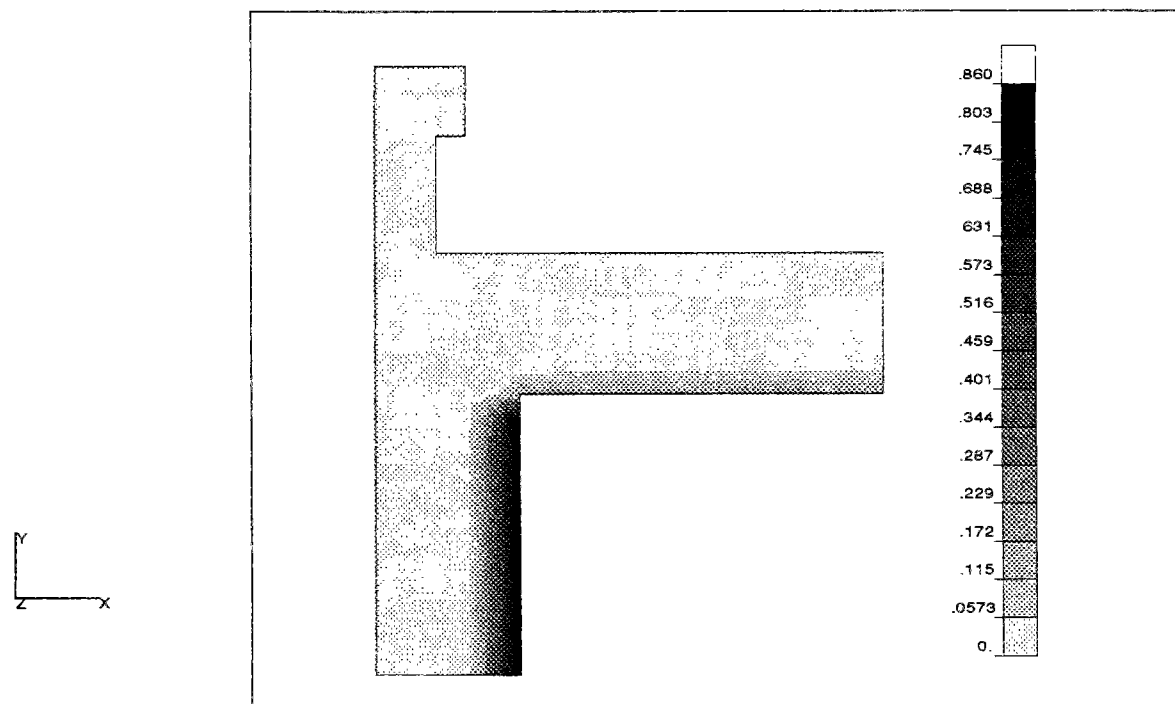


Figure 6.5: Sollicitation T_{int} : champ de température (en $^{\circ}C$) à $t_2=30mn$.

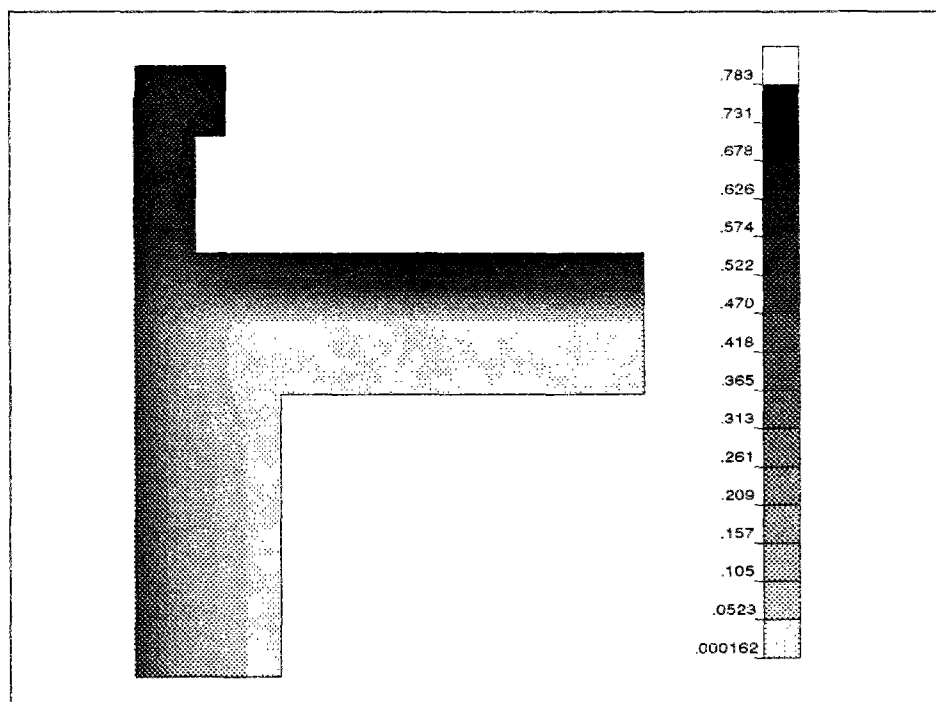


Figure 6.6: Sollicitation T_{ext} : champ de température (en $^{\circ}C$) à t_3 .

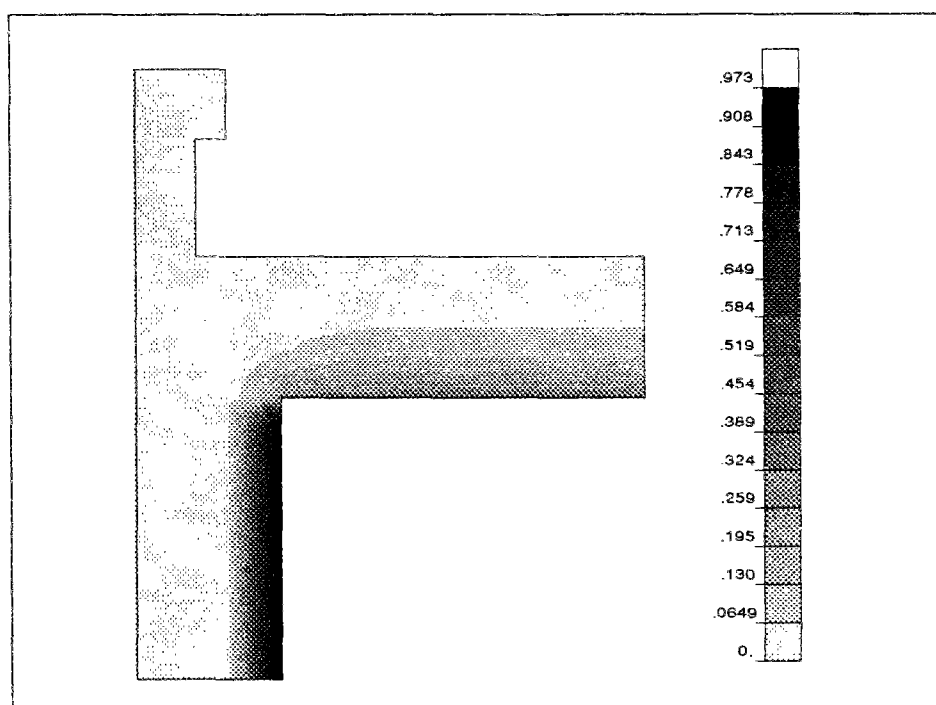


Figure 6.7: Sollicitation T_{int} : champ de température (en $^{\circ}C$) à t_3 .

6.3 Amalgame modal

Le modèle modal du pont thermique va être réduit avec la méthode d'amalgame modal. On donne dans le tableaux 6.3 les valeurs de $\mathcal{M}$ et $\overline{\mathcal{M}}$ lorsque l'ordre du modèle réduit augmente de 1 à 15.

ordre n	$\mathcal{M} (J^{\circ}C s)$	$\overline{\mathcal{M}} (^{\circ}C)$
1	$2.693 \cdot 10^{+9}$	$1.271 \cdot 10^{-1}$
2	$2.485 \cdot 10^{+8}$	$3.863 \cdot 10^{-2}$
3	$2.168 \cdot 10^{+7}$	$1.841 \cdot 10^{-2}$
4	$1.204 \cdot 10^{+7}$	$1.103 \cdot 10^{-2}$
5	$1.015 \cdot 10^{+7}$	$7.808 \cdot 10^{-3}$
6	$6.587 \cdot 10^{+6}$	$6.289 \cdot 10^{-3}$
7	$4.403 \cdot 10^{+6}$	$5.142 \cdot 10^{-3}$
8	$2.994 \cdot 10^{+6}$	$4.240 \cdot 10^{-3}$
9	$1.864 \cdot 10^{+6}$	$3.346 \cdot 10^{-3}$
10	$1.469 \cdot 10^{+6}$	$2.970 \cdot 10^{-3}$
11	$1.092 \cdot 10^{+6}$	$2.560 \cdot 10^{-3}$
12	$7.287 \cdot 10^{+5}$	$2.091 \cdot 10^{-3}$
13	$5.632 \cdot 10^{+5}$	$1.839 \cdot 10^{-3}$
14	$2.838 \cdot 10^{+5}$	$1.305 \cdot 10^{-3}$
15	$2.342 \cdot 10^{+5}$	$1.185 \cdot 10^{-3}$

Tableau 6.3: Mesure $\mathcal{M}$ et erreur $\overline{\mathcal{M}}$ pour $n = 1 \dots 15$.

On constate que pour avoir un modèle réduit d'une précision équivalente à celle de l'exemple précédent (bâtiment bizone), c'est à dire $\overline{\mathcal{M}}_{max} \leq 0.01^{\circ}C$, il est nécessaire de prendre au moins $n=5$. Considérons donc le modèle amalgamé (M.Ama) d'ordre 5.

Les sous-espaces d'amalgame obtenus sont donnés dans le tableau 6.4. Pour chaque sous-espace d'amalgame, on donne une partie des modes propres qui le composent avec les coefficients de décomposition associés (valeurs (.)). Les modes principaux sont encadrés $\boxed{\cdot}$.

Le premier sous-espace contient le mode $V_1(M)$ de plus grande constante de temps ($\tau_1 \# 10h$) qui est entouré de 26 modes mineurs. Le thermogramme du premier mode amalgamé est proche du mode principal $V_1(M)$ car les coefficients de décomposition associés aux modes mineurs qui l'entourent sont faibles. Les cinq modes amalgamés issus des cinq sous-espaces d'amalgame sont donnés dans les figures 6.8 à 6.12.

Comme les deux exemples précédents, on trace directement les champs écarts de température (en valeur absolue) entre les (M.Ama) et (M.Ref). Ceci permet d'avoir une vue rapide et complète des erreurs du modèle réduit. Au trois instants t_1, t_2 et t_3 , les champs

sous-espaces d'amalgame				
1	2	3	4	5
mode	mode	mode	mode	mode
1 (1.0000)	2 (1.0000)	3 (1.0000)	4 (1.0000)	6 (1.0000)
21 (0.0113)		53 (0.0074)	5 (-0.6895)	7 (-0.5400)
40 (0.0068)		59 (0.0074)	11 (0.0135)	8 (-0.4198)
42 (-0.0021)		69 (-0.0038)	29 (-0.0027)	9 (-0.3594)
44 (-0.0012)		129 (-0.0004)	54 (0.0030)	10 (0.1358)
48 (-0.0007)		136 (-0.0025)	86 (0.0003)	12 (-0.0139)
75 (0.0006)		141 (0.0007)	139 (-0.0002)	13 (-0.0318)
92 (0.0001)		158 (0.0004)	149 (-0.0001)	14 (0.0179)
...		...	...	...
dim 27	dim 1	dim 11	dim 9	dim 263

Tableau 6.4: Répartition de l'espace des modes propres en 5 sous-espaces d'amalgame

écarts de température sont donnés dans les figure 6.13 à 6.18.

On constate que à l'instant t_1 , lorsque la structure est excitée par un échelon de T_{int} , le modèle amalgamé est fortement biaisé au niveau de l'isolant situé sur le mur (voir figure 6.14). Les autres thermogrammes présentent nettement moins d'erreurs. Cependant, l'aspect général des champs écarts montre que les erreurs (même faibles) tendent à se situer au niveau des zones peu capacitives (isolant dans cet exemple). On retrouve ici le même constat établi dans le premier exemple (paroi tricouche). Cette propriété de l'amalgame modal s'explique par la mesure $\mathcal{M}$ de l'erreur retenue dans la méthode qui pondère le champ de température par la capacité calorifique. Cette pondération fait que les modes propres concernant les milieux isolants n'apparaissent pas toujours dans l'espace des modes principaux et la correction apportée par l'amalgame modal peut être insuffisante si la dimension "n" est trop faible.

Analysons le contenu des sous-espaces d'amalgame pour déduire la raison de l'insuffisance du (**M.Ama**) constatée sur la figure 6.14. Pour des structures composées de plusieurs milieux homogènes, il n'est pas rare de rencontrer² un sous-ensemble de modes propres localisés chacun dans un milieu homogène. C'est le cas du mode propre V_{40} donné dans la figure 6.19 qui est complètement localisé dans l'isolant du mur. La forme spatiale de ce mode propre est très proche du premier mode propre d'une seule couche d'isolant en conduction monodimensionnelle et soumise à des conditions aux limites de Dirichlet et Fourier sur ses deux faces respectivement. Par ailleurs, en regardant les sous-espaces d'amalgame du tableau 6.4, on s'aperçoit que $V_{40}(M)$ est un mode mineur affecté pour contribuer à la correction du premier mode principal. Le coefficient de décomposition associé à $V_{40}(M)$ est important lorsqu'on le compare à ceux d'autres modes mineurs mais cela reste insuffisant car la dynamique de H_1 ($\tau_1 \# 10h$) est trop lente comparée à

²Il n'y a pas à l'heure actuelle de justification théorique de cette constatation

$\tau_{40} \# 17mn$. Cette affectation est pourtant optimale: le sous-espace H_1 est celui où la contribution numérique de $V_{40}(M)$ à $\mathcal{M}$ est minimale.

Au niveau de l'isolant situé sur le plafond, le modèle amalgamé est précis. En fait on constate que le mode propre $V_8(M)$ donné dans la figure 6.20 est aussi localisé sur l'isolant du plafond. Ce mode propre est mineur mais il est affecté dans le dernier sous-espace dont la dynamique est proche de la sienne (voir tableau 6.2). D'autre part, le coefficient de décomposition associé à $V_8(M)$ est important (voir tableau 6.4). Ceci explique l'aspect spatial du dernier mode amalgamé $V_3(M)$ qui possède une forte (respectivement faible) amplitude sur l'isolant du plafond (respectivement du mur). Cette correction n'a pas été aussi forte pour le premier mode amalgamé sur l'isolant du mur.

Augmentons l'ordre du modèle réduit à l'ordre 6. Les sous-espaces d'amalgame sont maintenant donnés dans le tableau 6.5. On constate la "migration" du mode $V_{40}(M)$ vers le dernier sous-espace d'amalgame: le mode principal $V_{21}(M)$ est plus proche temporellement de $V_{40}(M)$ que ne l'est $V_1(M)$. De plus, parmi tous les modes mineurs du dernier sous-espace, $V_{40}(M)$ possède le plus grand coefficient de décomposition (tableau 6.5). Le champ écart entre les **(M.Ref)** et **(M.Ama)** à l'instant t_1 est maintenant donné dans la figure 6.21 (pour la sollicitation en T_{int}). L'erreur obtenue dans la zone de l'isolant du mur est cette fois-ci nettement atténuée.

Si on poursuit l'augmentation l'ordre n , le mode $V_{40}(M)$ devient principal à $n=10$. On donne dans la figure 6.22 le champ écart entre les **(M.Ref)** et **(M.Ama)** à l'instant t_1 (pour la sollicitation en T_{int}). Le résultat est encore amélioré.

Ainsi, partant d'une contrainte donnée sur la dimension n et/ou l'erreur $\mathcal{M}$, l'amalgame modal optimise le contenu des sous-espaces d'amalgame de manière à minimiser la mesure de l'erreur. Pour des ordres n faibles, le modèle amalgamé peut présenter une insuffisance au niveau des zones peu capacitives. Cette insuffisance disparaît en augmentant la dimension du **(M.Ama)**.

sous-espaces d'amalgame					
1	2	3	4	5	6
mode	mode	mode	mode	mode	mode
1 (1.0000)	2 (1.0000)	3 (1.0000)	4 (1.0000) 5 (-0.6895)	6 (1.0000) 7 (-0.5400) 8 (-0.4198) 9 (-0.3594) 10 (0.1358) 12 (-0.0139) 13 (-0.0318) 14 (0.0179) ...	21 (1.0000) 11 (0.1224) 23 (0.1532) 29 (-0.0446) 32 (-0.0216) 40 (0.8508) 41 (-0.0313) 42 (-0.2650) ...
dim 1	dim 1	dim 1	dim 2	dim 233	dim 73

Tableau 6.5: Répartition de l'espace des modes propres en 6 sous-espaces d'amalgame

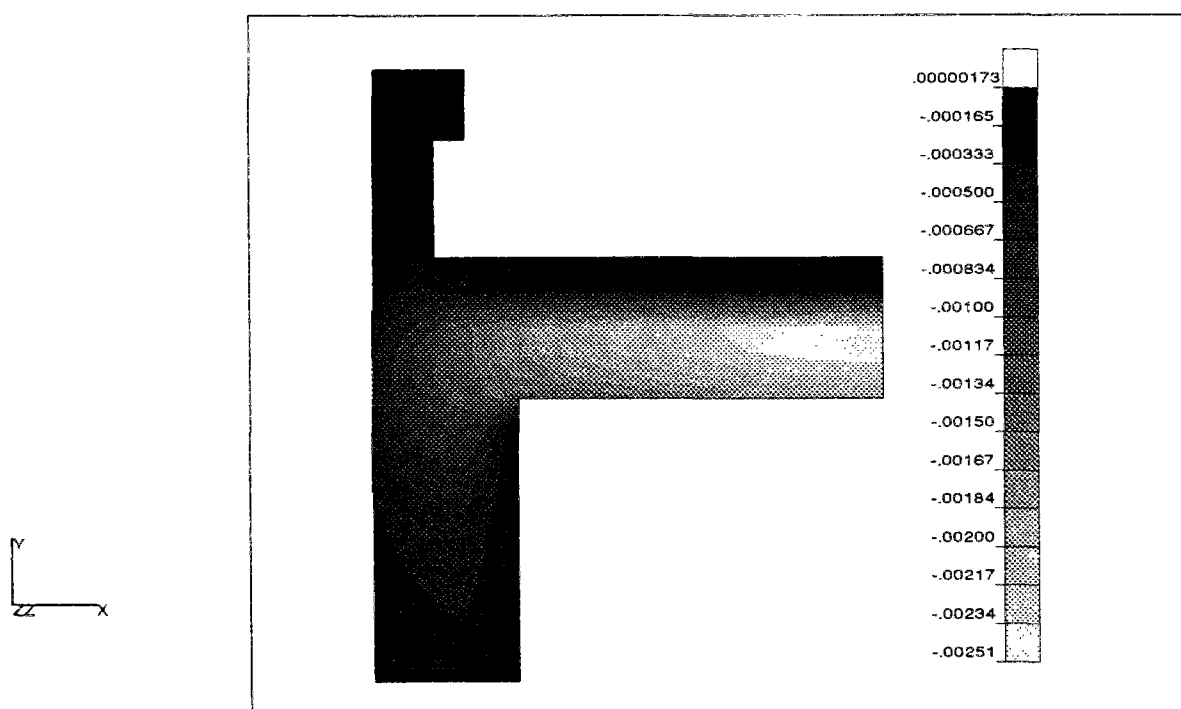


Figure 6.8: Mode amalgamé 1 (9h 59mn 14s).

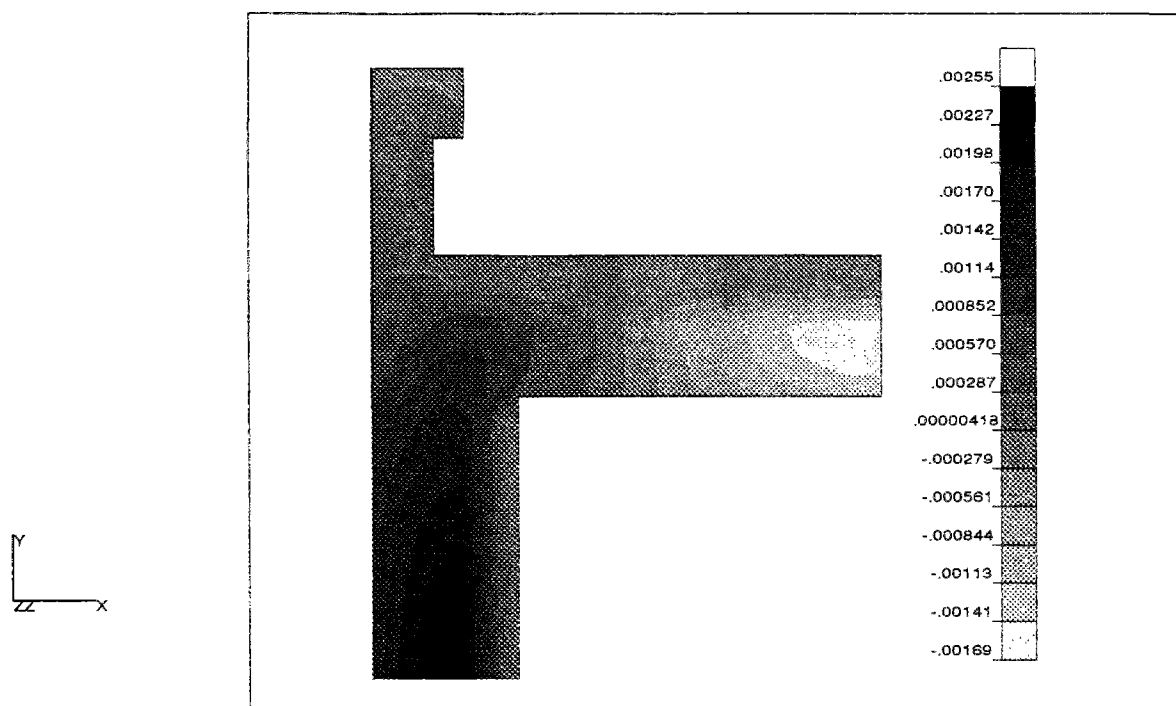


Figure 6.9: Mode amalgamé 2 (7h 29mn 23s).

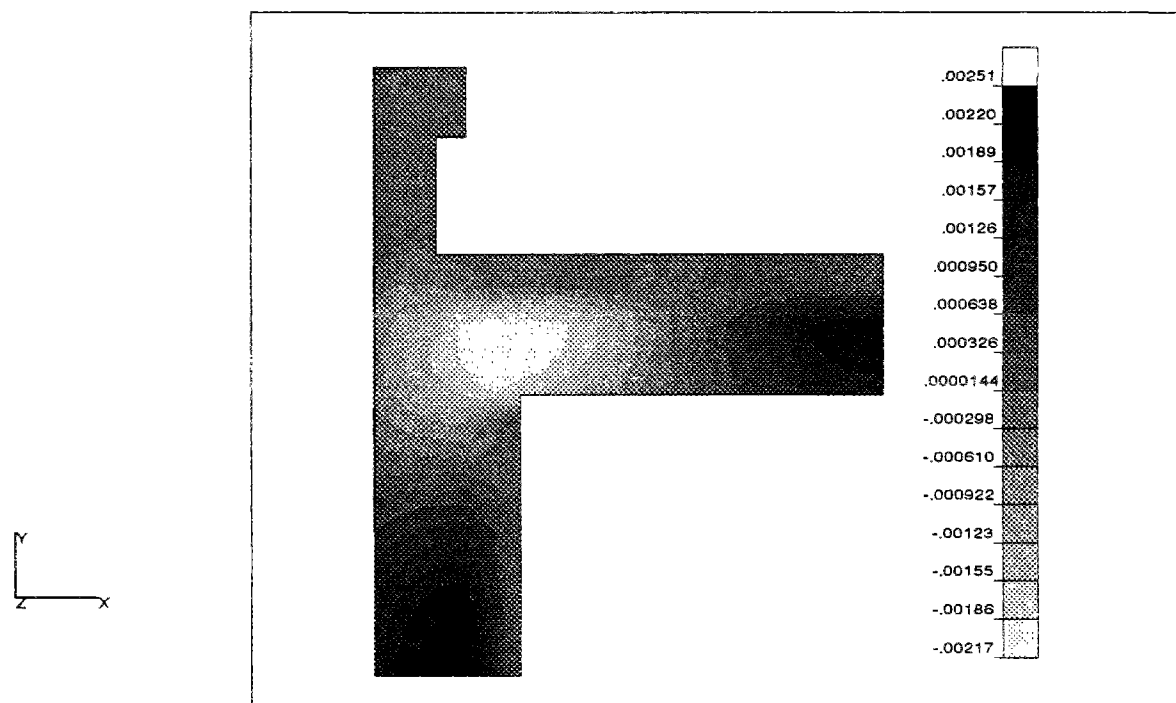


Figure 6.10: Mode amalgamé 3 (6h 9mn 24s).

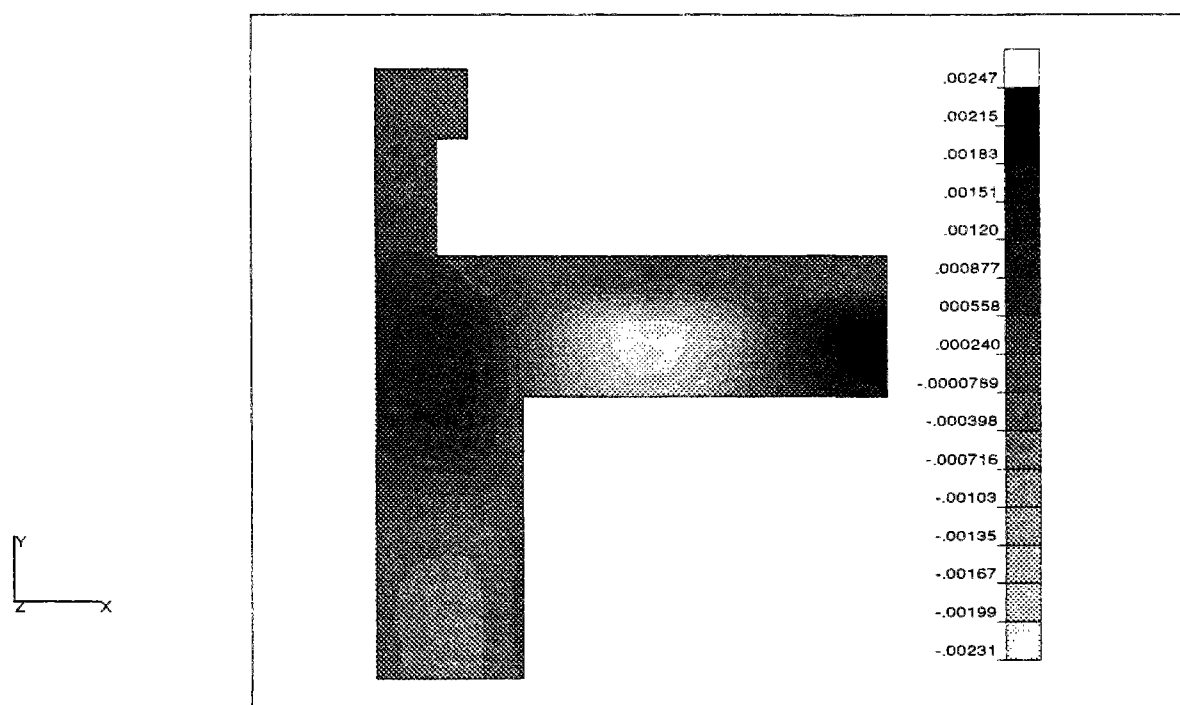


Figure 6.11: Mode amalgamé 4 (4h 31s).

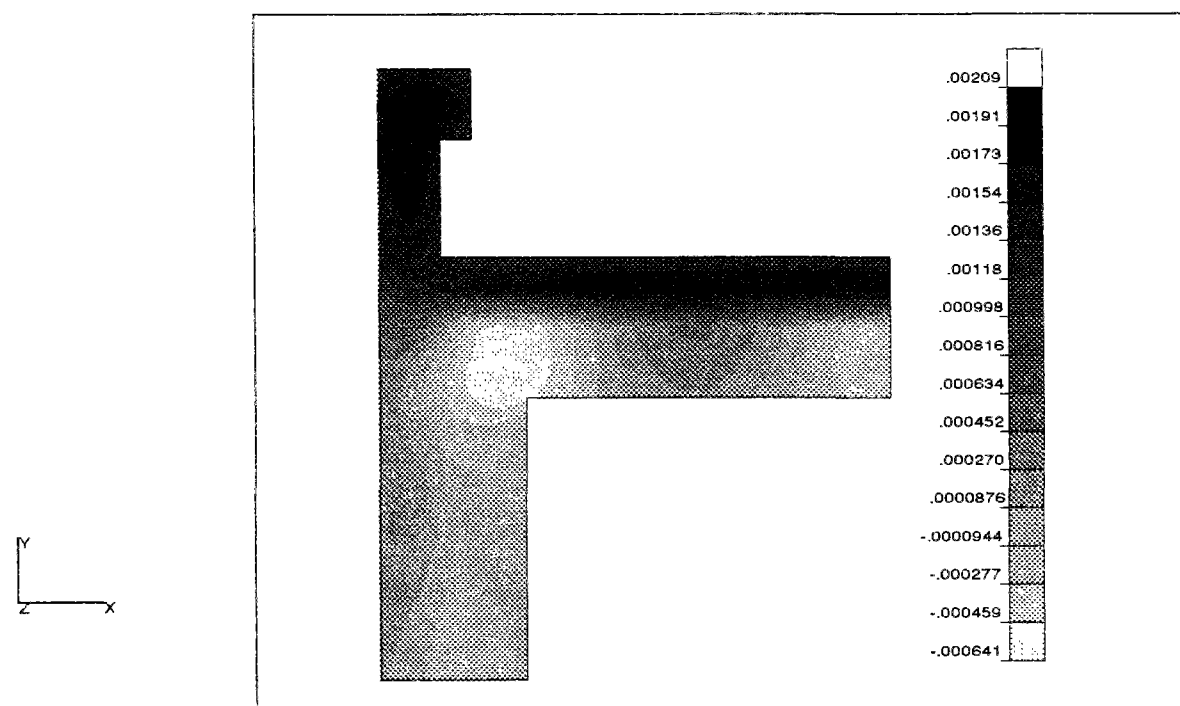


Figure 6.12: Mode amalgamé 5 (2h 12mn 56s).

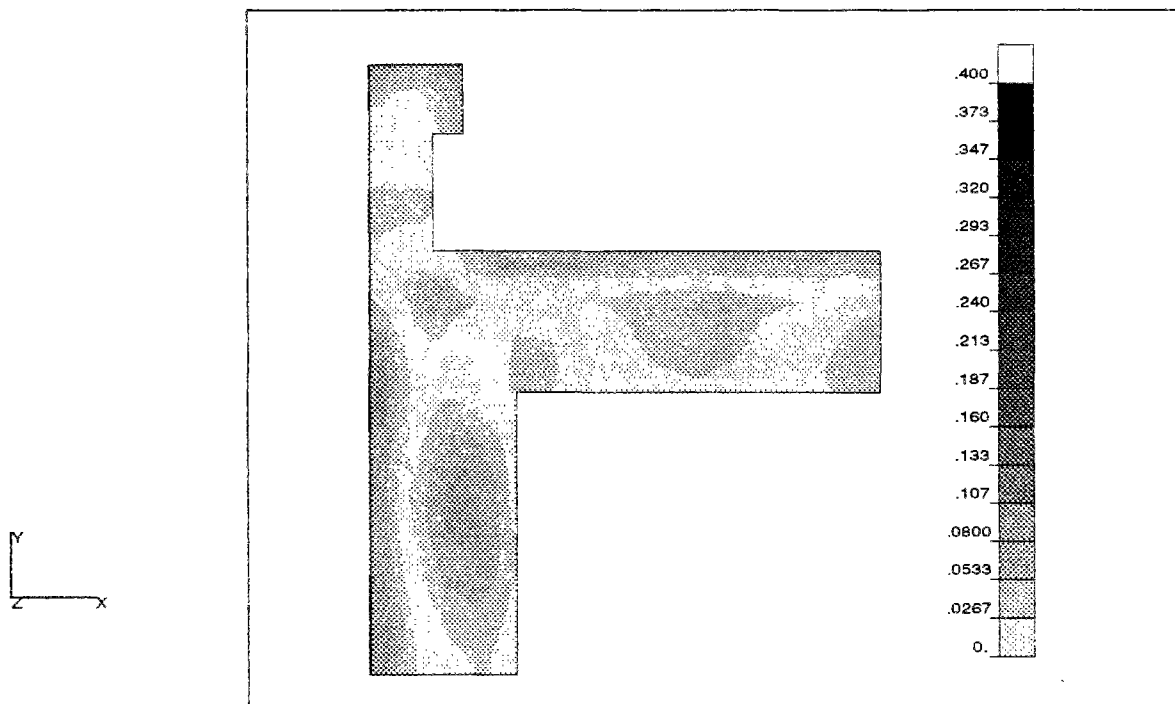


Figure 6.13: Sollicitation T_{ext} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Ama) à t_1 .

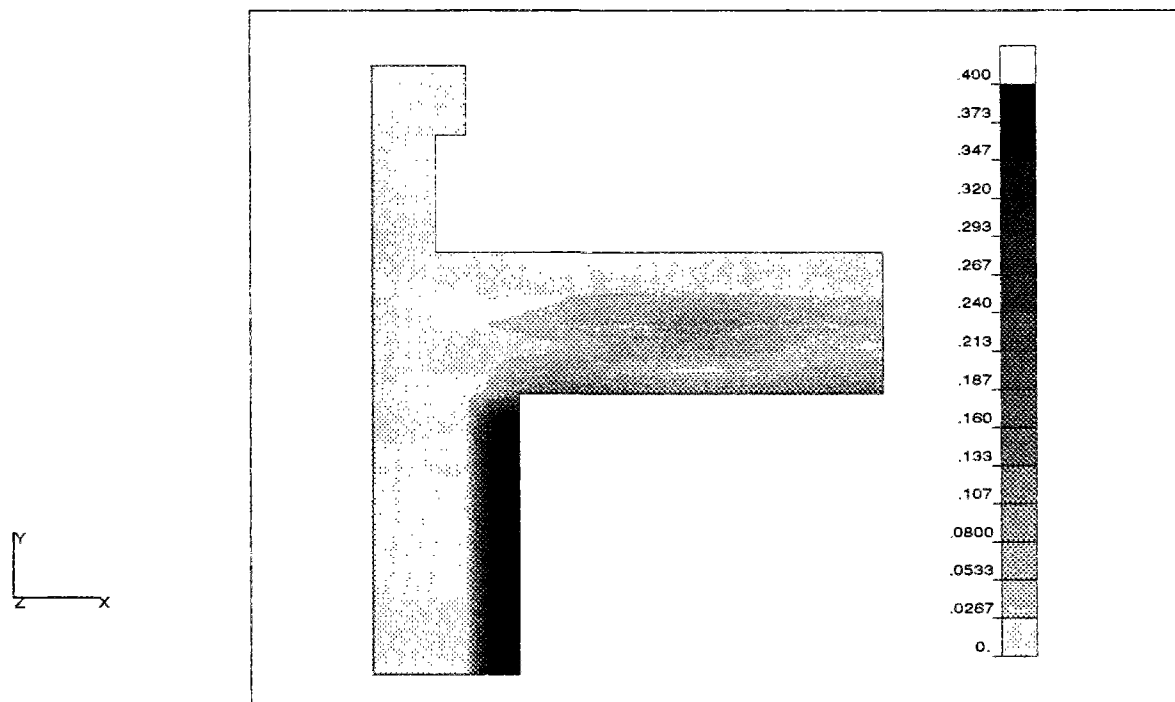


Figure 6.14: Sollicitation T_{int} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Ama) à t_1 .

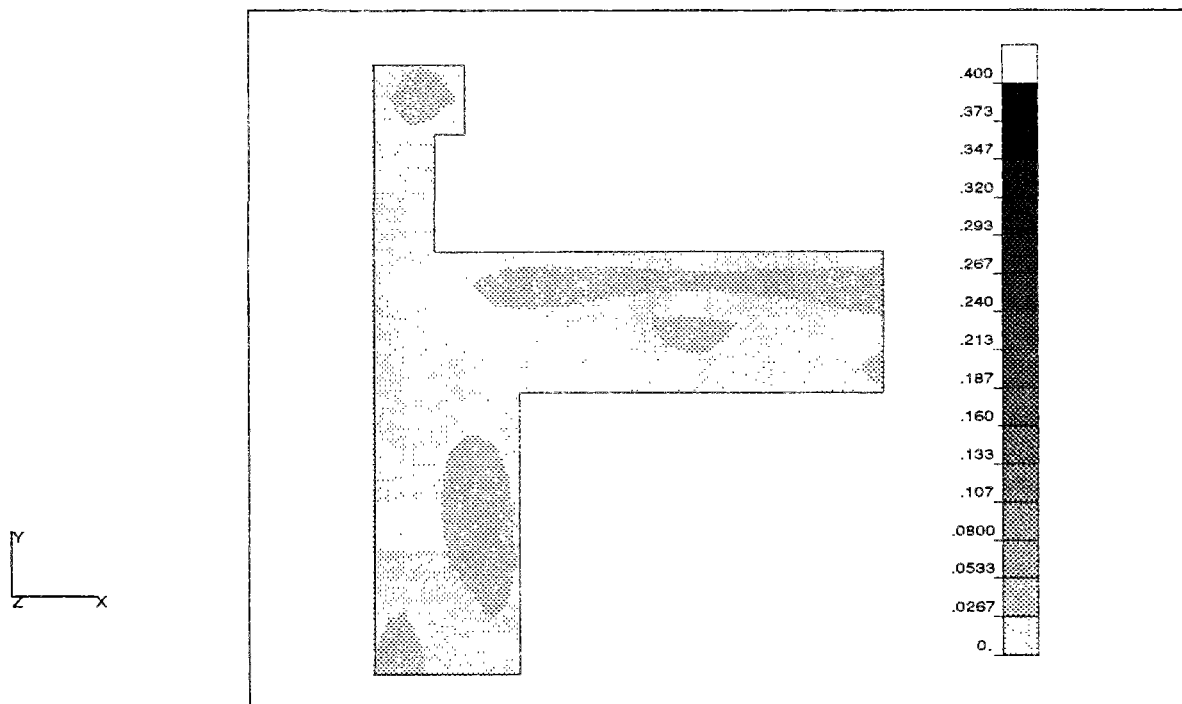


Figure 6.15: Sollicitation T_{ext} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Ama) à t_2 .

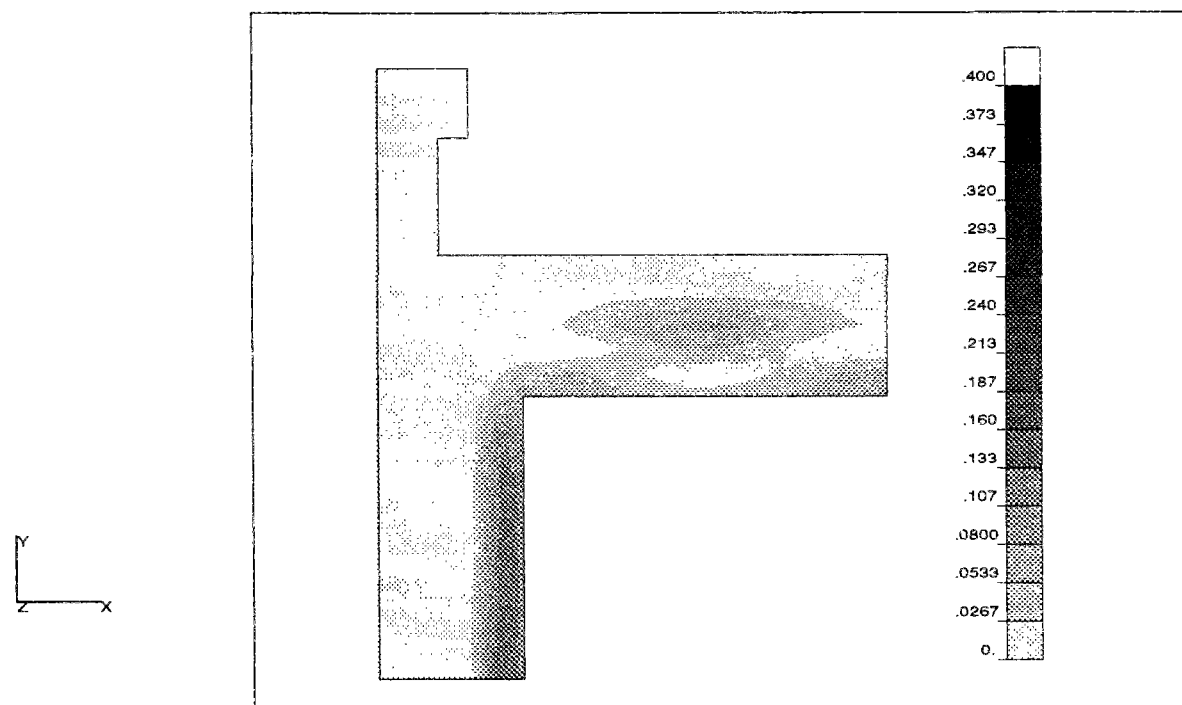


Figure 6.16: Sollicitation T_{int} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Ama) à t_2 .

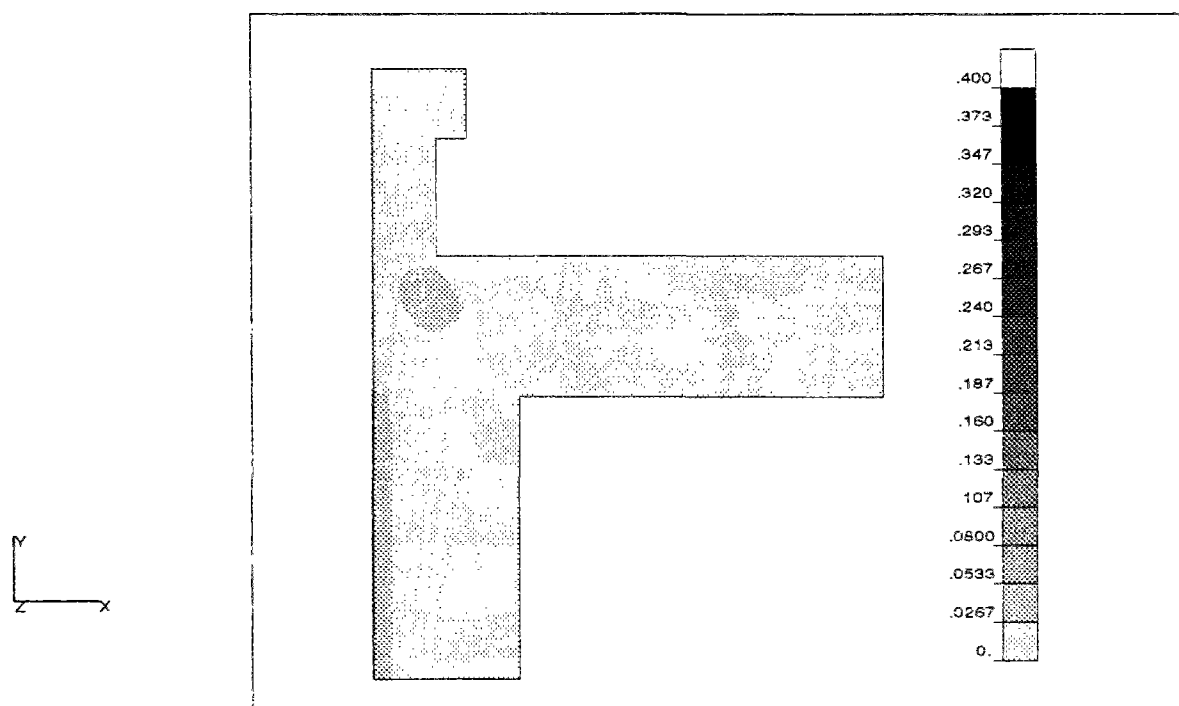


Figure 6.17: Sollicitation T_{ext} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Ama) à t_3 .

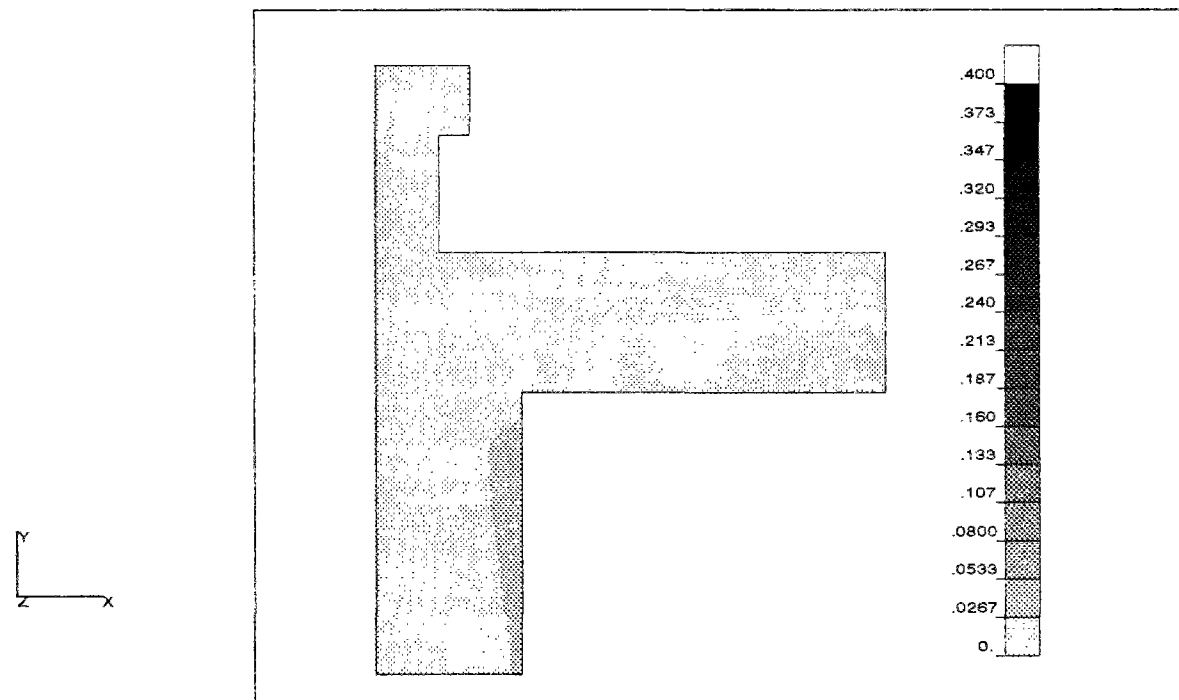


Figure 6.18: Sollicitation T_{int} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Ama) à t_3 .

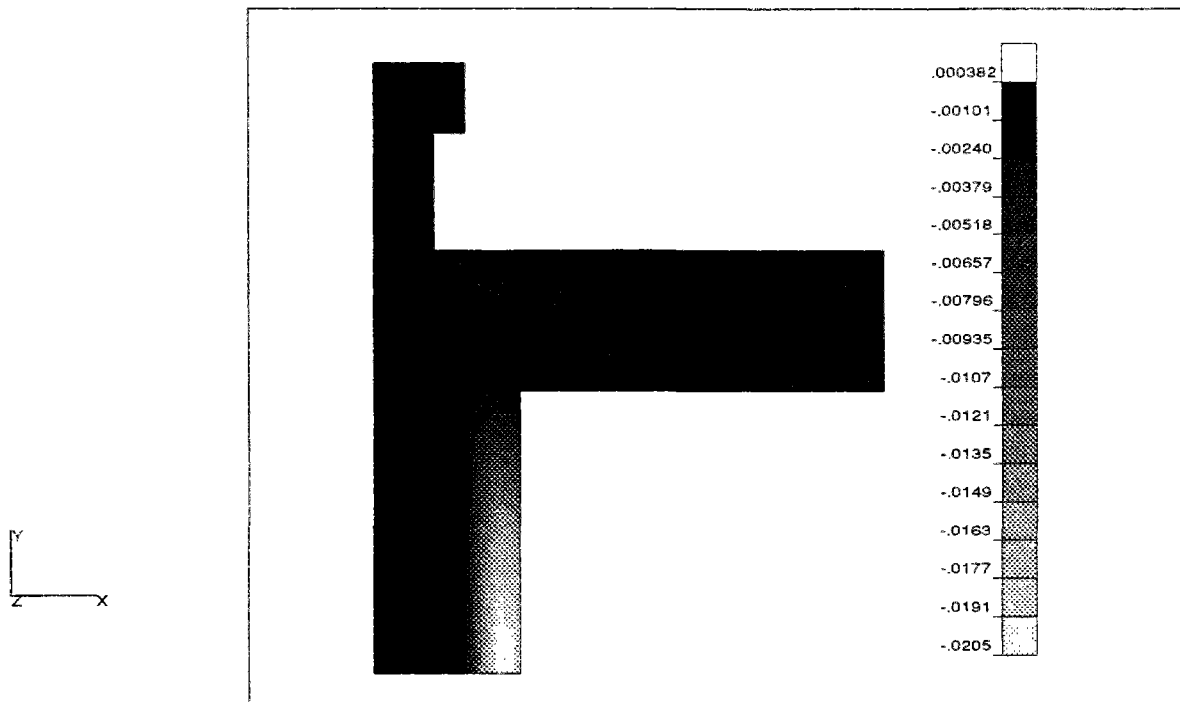


Figure 6.19: Mode propre 40 (17Mn 28s).

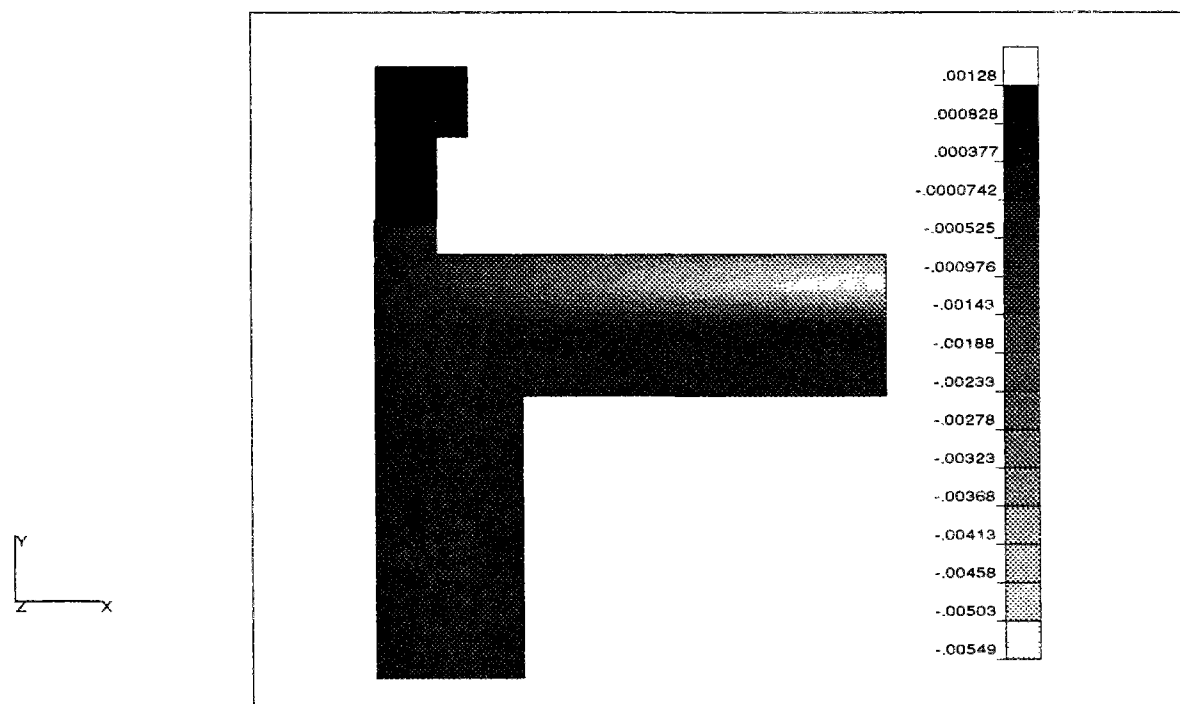


Figure 6.20: Mode propre 8 (1h 57mn 17s).

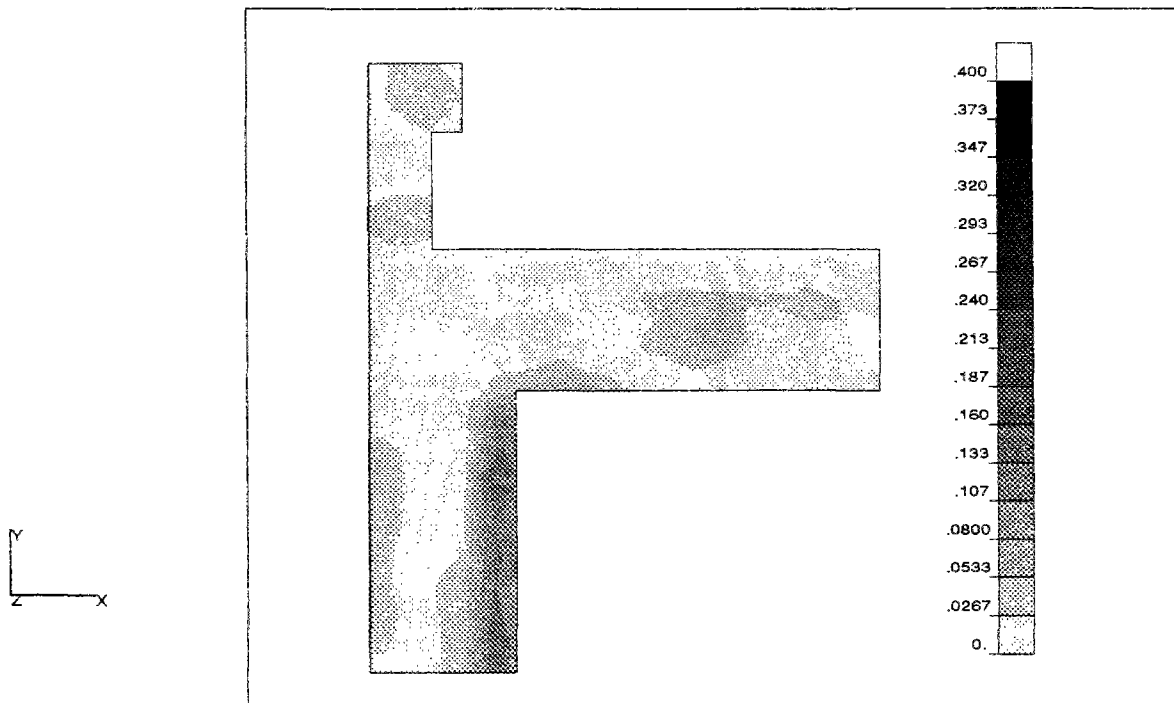


Figure 6.21: Sollicitation T_{int} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Ama) d'ordre 6" à t_1 .

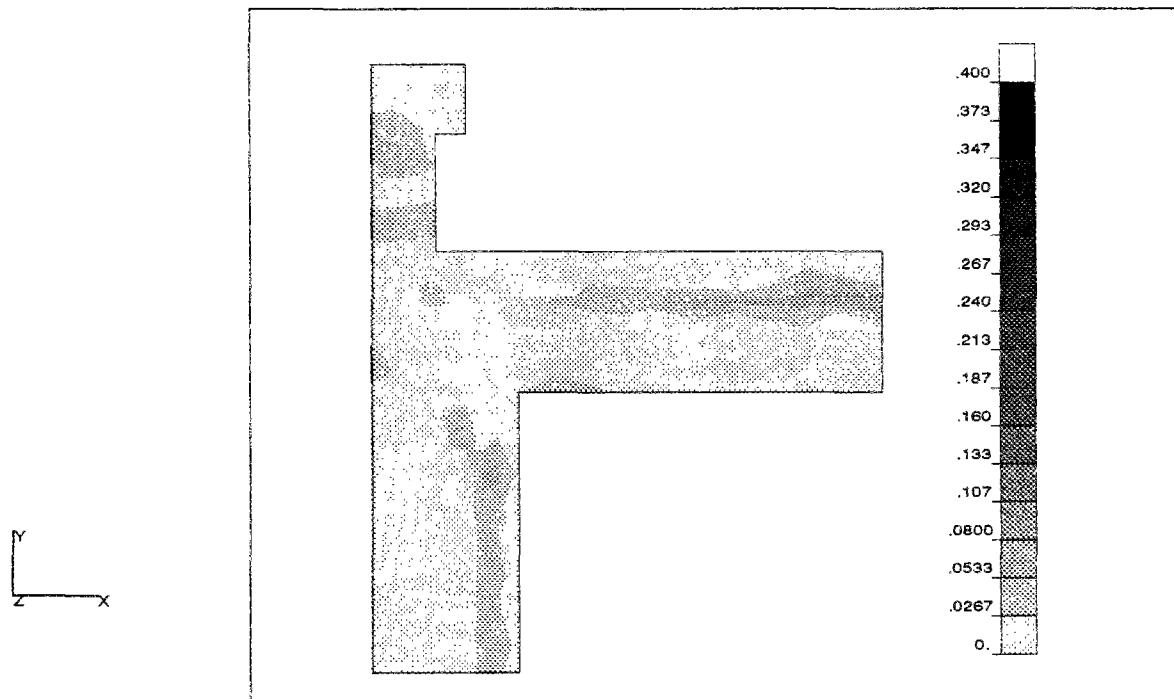


Figure 6.22: Sollicitation T_{int} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Ama) d'ordre 10" à t_1 .

6.4 Méthode mixte réduction-identification

Parmi les méthodes que nous avons décrites dans la partie théorique de cette thèse, il y a la méthode mixte réduction-identification [56]. Nous allons utiliser cette méthode pour réduire le modèle modal du pont thermique.

La méthode mixte (voir §2.3.4) est réalisée en deux étapes indépendantes:

- 1 ► "identification" des matrices $\hat{P}$ et $\tilde{F}$
- 2 ► calcul de la matrice de commande $\tilde{B}$

L'identification des pseudo-éléments propres a été faite³ à partir de thermogrammes des températures aux 311 nœuds du système et obtenus à partir du (M.Ref). Ces thermogrammes sont calculés à des instants allant de $t = 0$ à $t = 10\tau_1$ avec un pas de temps de 500s. Au total cela donne 289 thermogrammes de température. Chaque thermogramme est un vecteur de température à un instant donné et en tous les nœuds **lorsque les deux sollicitations (T_{ext} et T_{int}) sont des échelons unitaires.**

L'identification des pseudo-modes et des constantes de temps associées se fait par la minimisation du critère

$$\mathcal{J} = \sum_{i=1}^q \sum_{m=1}^r [\Delta T_i(t_m)]^2$$

où $\Delta T_i(t_m) = T_i(t_m) - \tilde{T}_i(t_m)$ représente l'écart de température au nœud i , à l'instant t_m , entre le (M.Ref) et le modèle réduit (M.Mixte). "q" est ici le nombre de nœuds de la discrétisation spatiale (311) et "r" le nombre de points de la discrétisation temporelle (289).

Le tableau 6.6 donne les constantes de temps identifiées en incrémentant⁴ l'ordre n du modèle réduit de 1 à 5. Le paramètre σ représente l'écart type qui sert ici pour quantifier la précision du (M.Mixte). Ce paramètre est donné par

$$\sigma = \sqrt{\frac{\mathcal{J}}{q r}} \quad (\text{en } ^\circ C)$$

On peut faire plusieurs remarques:

- à partir de $n=4$, la première constante de temps est très proche de celle de la référence qui est de 9h 59mn 14s. Pour $n=5$, la quatrième constante de temps est proche de la deuxième de la référence qui est de 7h 29mn 24s.
- Les constantes de temps identifiées contiennent pratiquement toutes les dynamiques.

³Nous ne disposons pas des programmes informatiques pour cette méthode. Cette identification n'aurait pas été possible sans la participation personnelle de Monsieur D. PETIT de l'IUSTI-université de Provence (Aix- Marseille I).Qu'il trouve ici l'expression de notre profonde reconnaissance.

⁴En fait la méthode mixte opère par étapes successives: on identifie d'abord $\tilde{P}$ et $\tilde{F}$ pour $n=1$, ensuite ces résultats servent d'entrées pour identifier $\tilde{P}$ et $\tilde{F}$ pour $n=2$, etc.

- on obtient une constantes de temps inférieure au pas de temps utilisé qui est de 8mn 20s. Il semble [56] que cette valeur (qui n'a pas de sens physique) sert à ajuster le premier point de chaque thermogramme ($t=0$).

Pour effectuer une comparaison avec la méthode d'amalgame, nous utilisons ici le (M.Mixte) d'ordre 5. Les cinq pseudo-modes identifiés sont donnés dans les figures 6.23 à 6.27. On peut constater que:

- Le premier mode identifié a globalement la même allure que le premier mode amalgamé de la figure 6.8 (ce dernier étant très proche du premier mode $V_1(M)$). Les constantes de temps des deux modes sont peu différentes.
- Le mode identifié 2 a aussi une allure voisine de celle du cinquième mode amalgamé. Comme le mode amalgamé 5, ce mode identifié sert à garantir une certaine précision du (M.Mixte) au niveau de l'isolant du plafond. La constante de temps du deuxième mode identifié est de 2h 5mn 59s et celle du cinquième mode amalgamé est de 2h 12mn 56s. Ces deux valeurs sont voisines.
- Le mode identifié 4 ressemble au mode amalgamé $\tilde{V}_2(M)$. Les constantes de temps de ces deux modes sont 6h 50mn 51s et 7h 29mn 23s respectivement.
- Les modes identifiés 3 et 5 n'ont pas de ressemblance avec les modes amalgamés. Notons toutefois que le mode identifié 3 est particulièrement localisé dans l'isolant du mur. D'ailleurs, son allure ressemble à celle du champ écart de la figure 6.14 et au mode propre 40 (figure 6.19) dont il est très proche temporellement. Ce mode a vraisemblablement un rôle important pour la précision du (M.Mixte) au niveau de l'isolant du mur. Ce mode identifié n'a pas été retrouvé dans la méthode d'amalgame, ce qui explique les erreurs de la figure 6.14.
- Le dernier mode identifié est aussi localisé dans l'isolant du mur. Sa constante de temps est faible (6mn 24s). Il est vraisemblable que ce mode "améliore" le (M.Mixte) au niveau de l'isolant du mur aux instants très courts.

Traçons les champs écarts de température (en valeur absolue) pour les deux sollicitations prises individuellement et aux trois instants t_1 , t_2 et t_3 . Les thermogrammes correspondants sont donnés dans les figures 6.28 à 6.33. On peut les comparer à ceux obtenus par la méthode d'amalgame modal (figures 6.13 à 6.33). On remarque qu'il n'y pas de localisation d'erreurs dans l'isolant. En particulier, la figure 6.29 ne comporte pas la localisation des erreurs constatées sur la figure 6.14 de l'amalgame modal. En revanche, hormis la figure 6.14, toutes les autres présentent des erreurs un peu plus grandes que celles de l'amalgame modal.

Les matrices $\tilde{P}$ et $\tilde{F}$ ont été obtenues à partir de sollicitations échelons unitaires. Aussi, lorsqu'on sollicite la structure à la fois par les deux sollicitations (un cas peu physique)

- (T_{ext}) en forme d'échelon unitaire
- (T_{int}) en forme d'échelon unitaire

alors les méthodes mixte et d'amalgame deviennent globalement équivalentes. Ce cas de figure correspond exactement aux conditions dans lesquelles les matrices $\tilde{P}$ et $\tilde{F}$ ont été identifiées. On pense que la dominance des modes identifiés dans la méthode mixte dépend de la façon dont les thermogrammes de température ont été obtenus à partir du (M.Ref). On peut (peut être !) aboutir à (M.Mixte) plus performant en utilisant des thermogrammes de température correspondant aux sollicitations

- (T_{ext}) en forme d'échelon unitaire (de 1 à 0)
- (T_{int}) en forme d'échelon unitaire (de -1 à 0)

Seuls les tests numériques pourront affirmer ou non cette remarque. Néanmoins, il est logique de pouvoir identifier les modes "les plus dominants" en utilisant des thermogrammes "riches", c'est à dire où il y aurait un maximum de modes excités. En particulier, il faut être sûr qu'il n'y a pas d'effet de "compensation de modes" dans les thermogrammes utilisés. Dans le cas de structures avec un nombre de sollicitations important (comme le bâtiment bizonne que nous avons étudié avec 9 sollicitations), le choix des sollicitations pour obtenir les thermogrammes est délicat.

Les deux méthodes présentées ne sont pas vraiment comparables en termes de précision car chacune a ses avantages. Elles aboutissent à des modèles réduits performants. La méthode mixte s'adapte aussi aux problèmes d'identification à partir de thermogrammes expérimentaux. En revanche, le temps de calcul de la méthode d'amalgame est nettement inférieur à celui de la méthode mixte. Sur l'exemple du pont thermique que venons d'étudier, la méthode d'amalgame modal a nécessité environ⁵ 150s et la méthode mixte environ⁶ 3h.

	ordre 1	ordre 2	ordre 3	ordre 4	ordre 5
constantes de temps	$\tau_1=8h\ 48s$	$\tau_1=9h\ 12mn\ 52s$ $\tau_2=2h\ 47s$	$\tau_1=9h\ 15mn\ 36s$ $\tau_2=2h\ 18mn\ 36s$ $\tau_3=18mn\ 26s$	$\tau_1=9h\ 56mn\ 29s$ $\tau_2=2h\ 6mn\ 7s$ $\tau_3=17mn\ 17s$ $\tau_4=6h\ 49mn\ 2s$	$\tau_1=9h\ 58mn\ 1s$ $\tau_2=2h\ 5mn\ 59s$ $\tau_3=18mn\ 16s$ $\tau_4=6h\ 50mn\ 51s$ $\tau_5=6mn\ 24s$
$\sigma\ (^{\circ}C)$	0.05735	0.01802	0.00852	0.00337	0.00146

Tableau 6.6: Constantes de temps identifiées et écarts types.

⁵ temps CPU sur la machine SUN SPARC Station 10

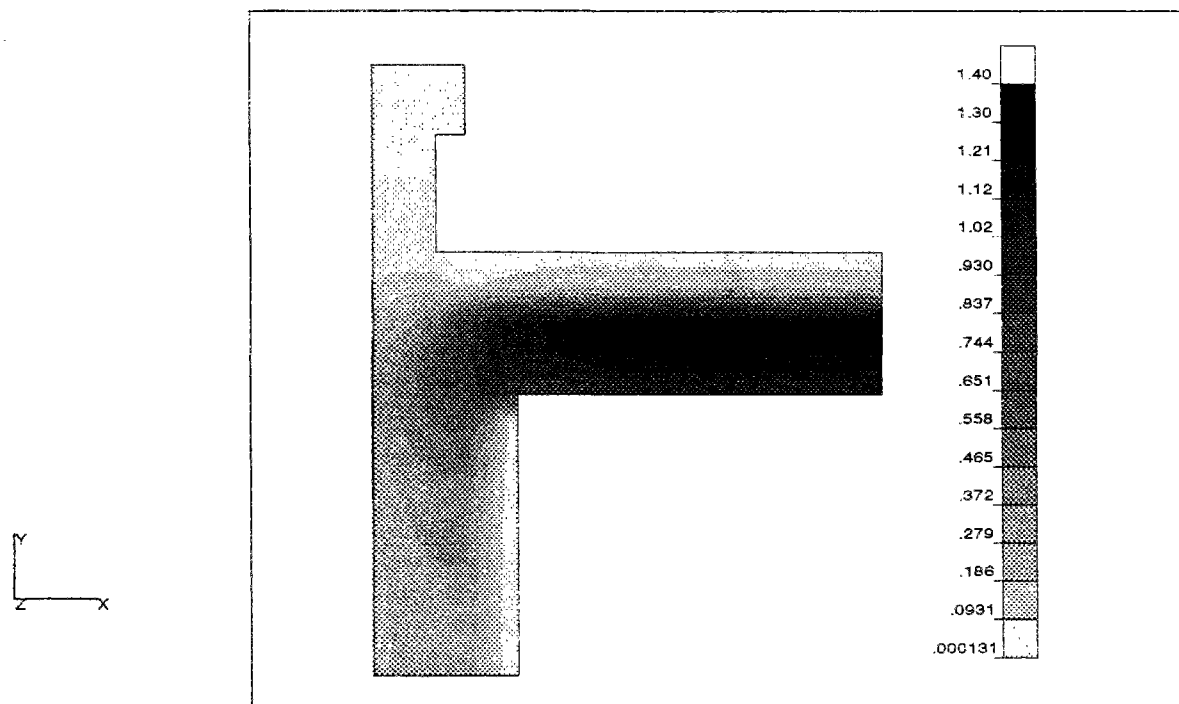


Figure 6.23: Mode identifié 1 (9h 58mn 1s).

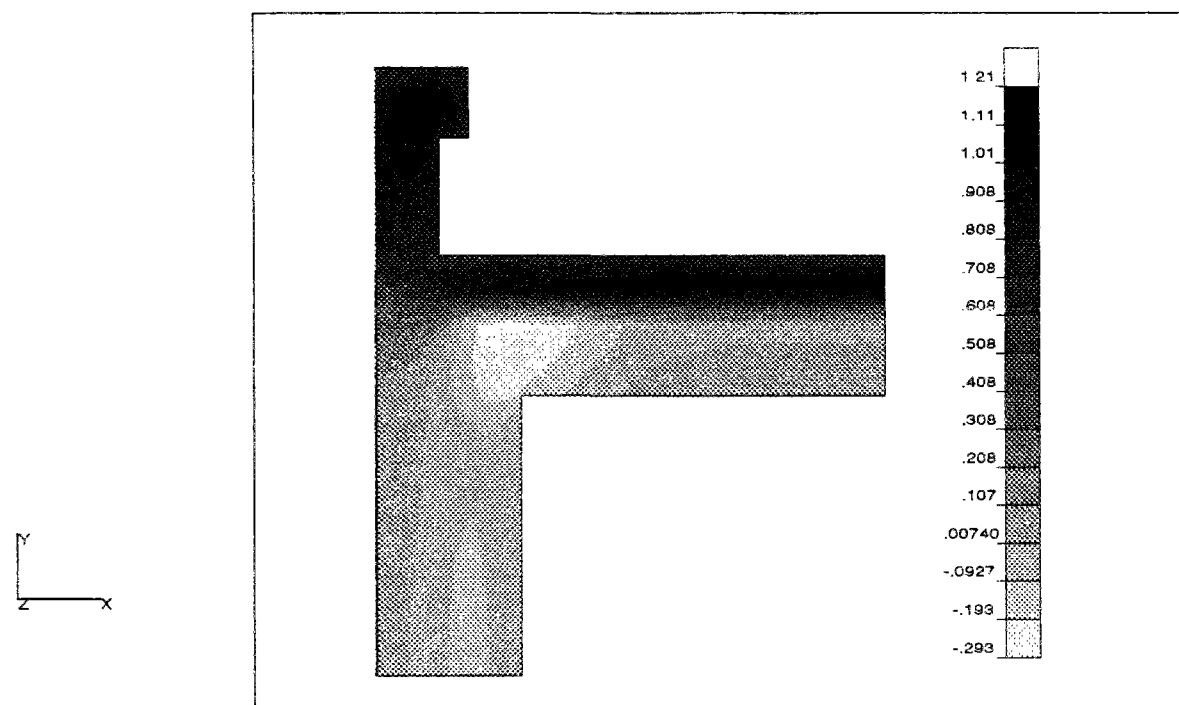


Figure 6.24: Mode identifié 2 (2h 5mn 59s).

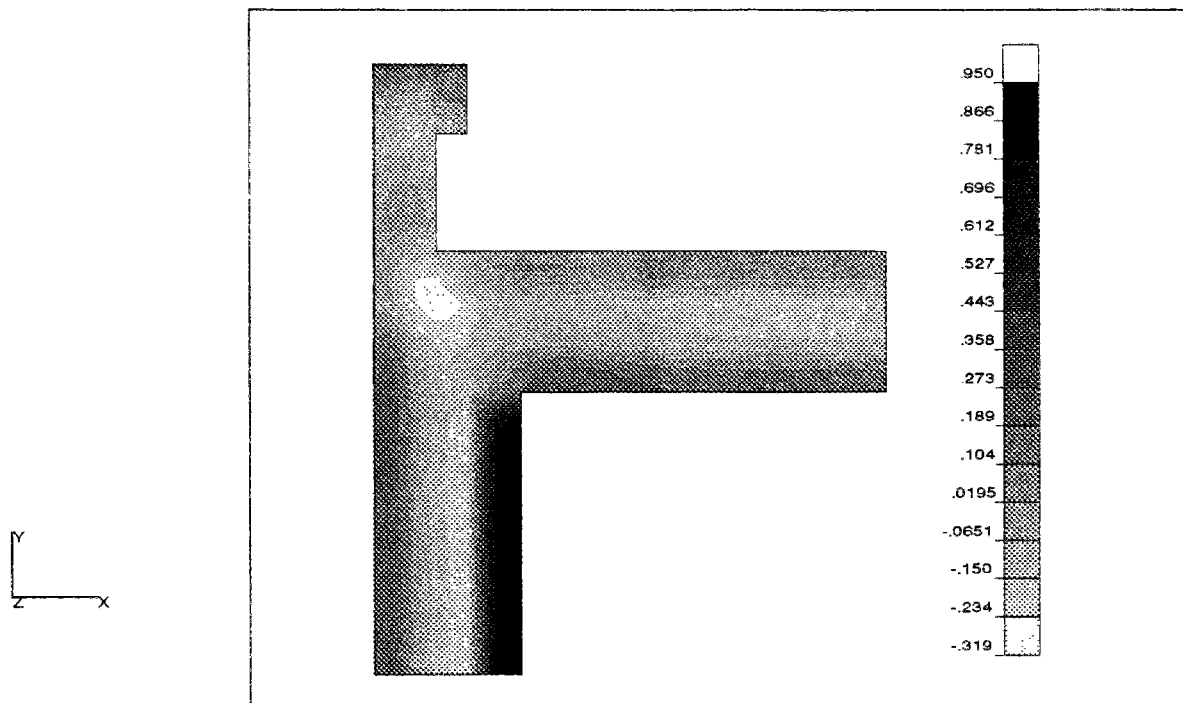


Figure 6.25: Mode identifié 3 (18mn 16s).

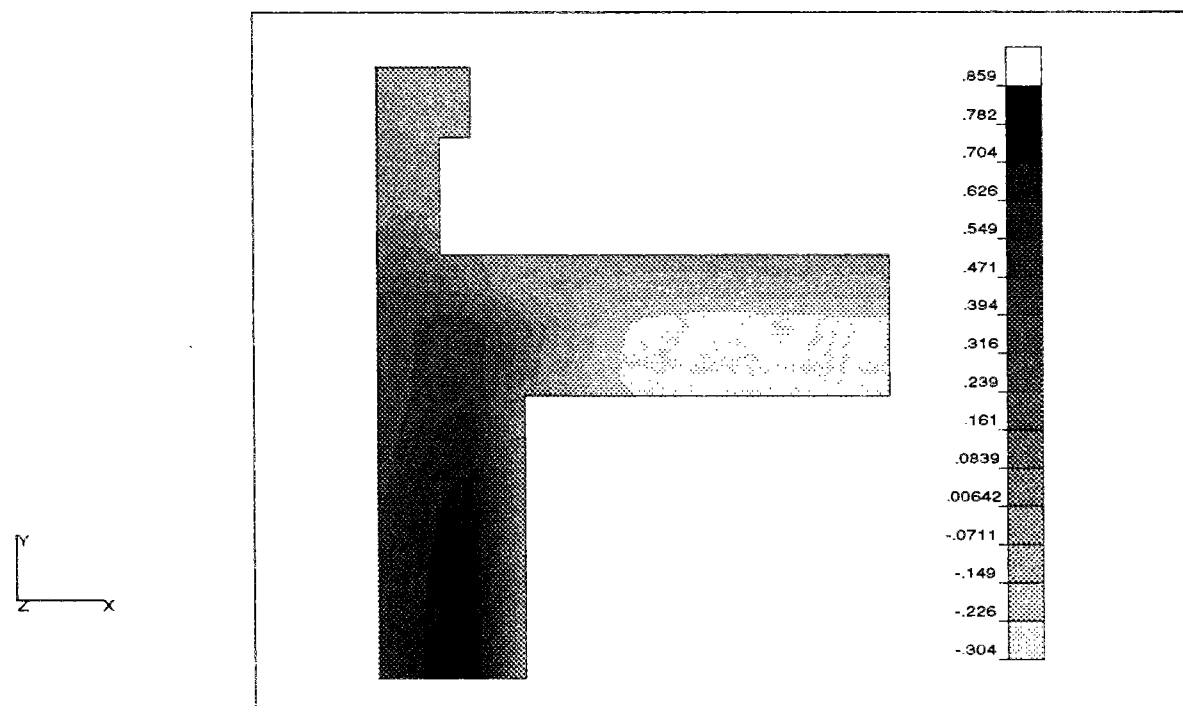


Figure 6.26: Mode identifié 4 (6h 50mn 51s).

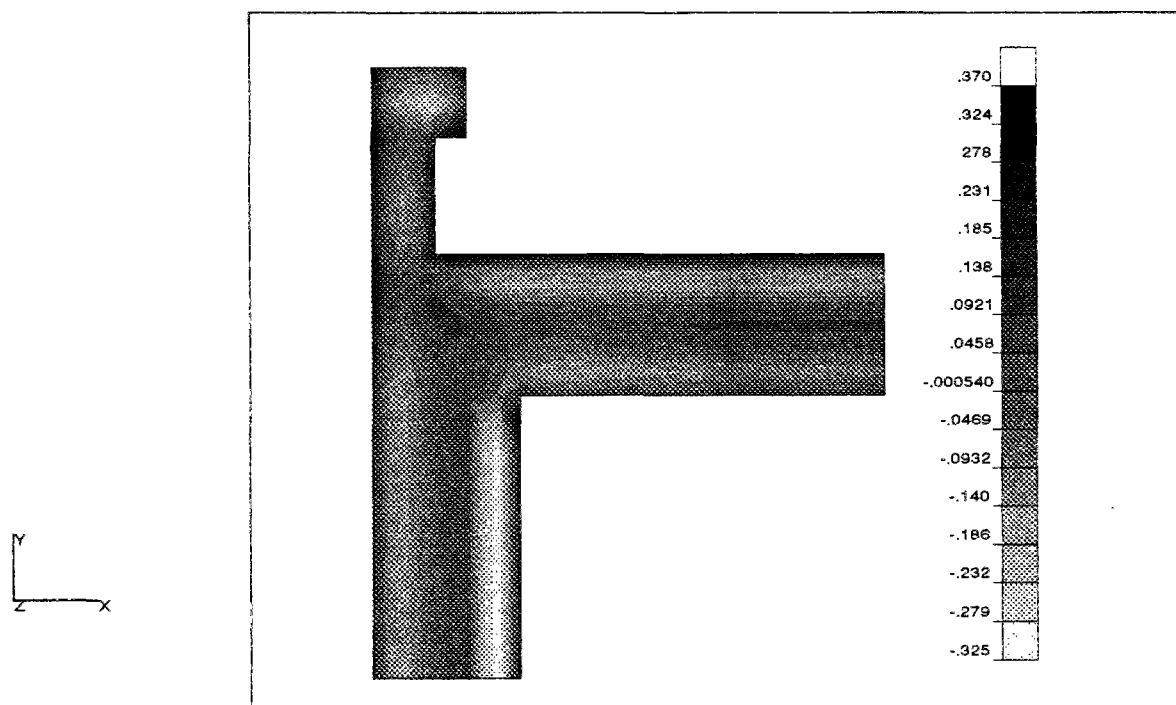
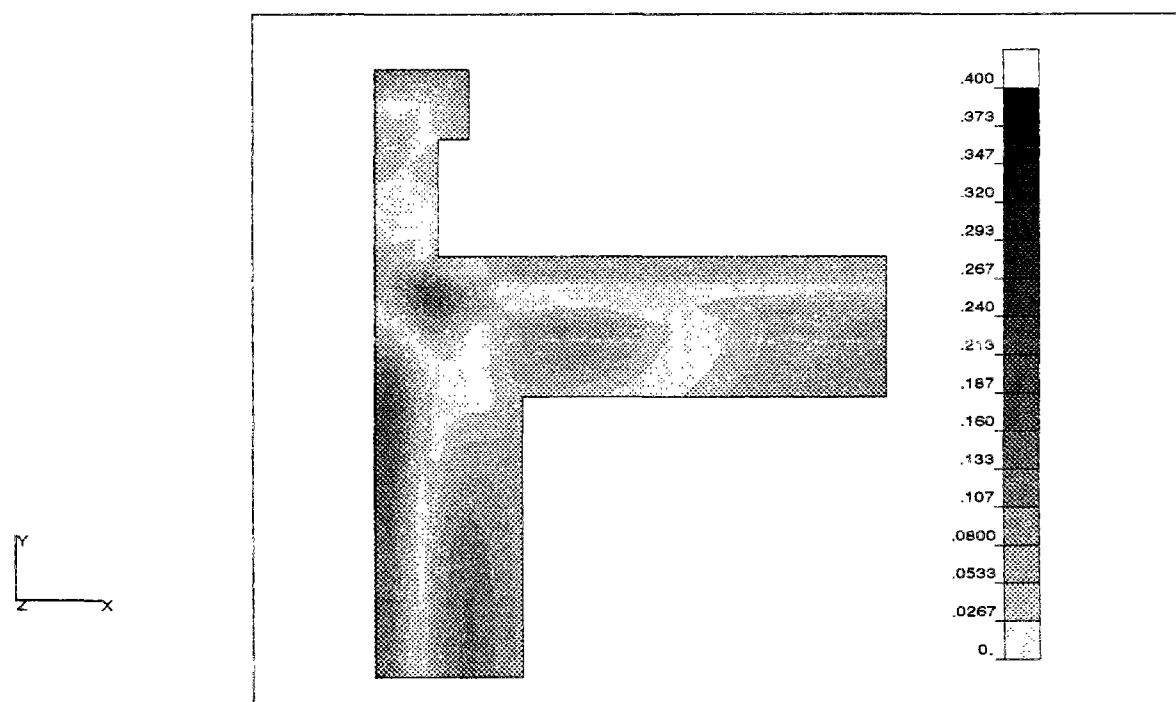


Figure 6.27: Mode identifié 5 (6mn 24s).

Figure 6.28: Sollicitation T_{ext} : écart de température (en °C) entre (M.Ref) et (M.Mixte) à t_1 .

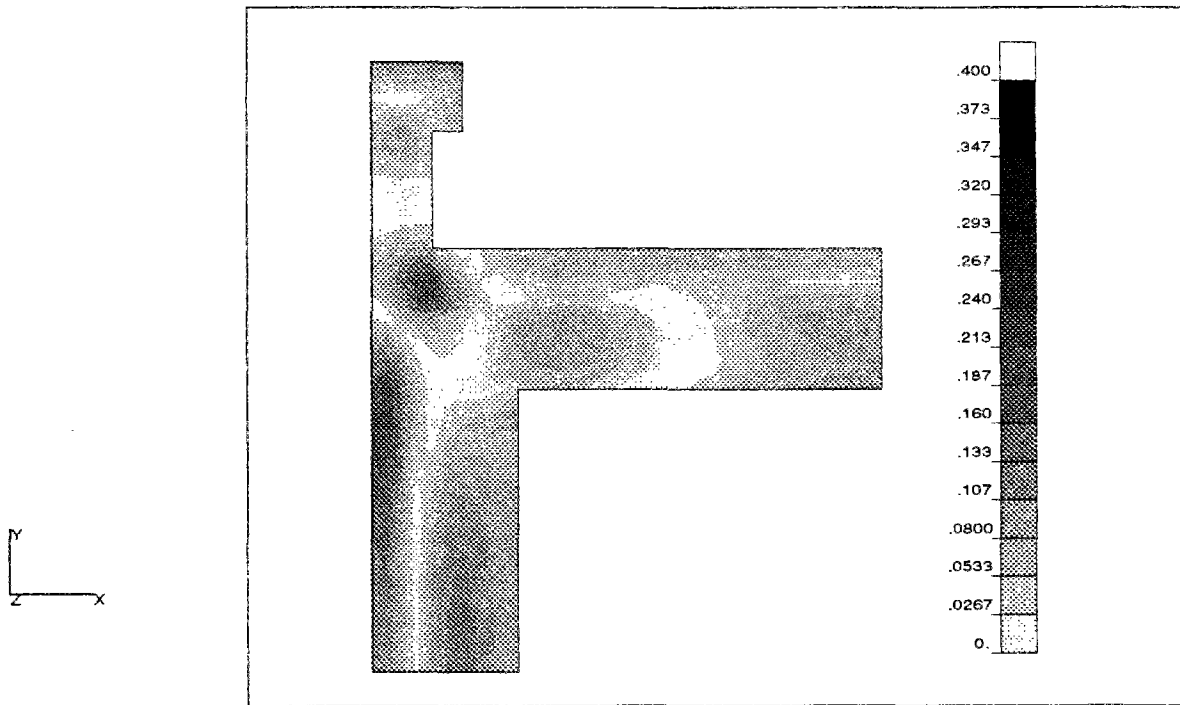


Figure 6.29: Sollicitation T_{int} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Mixte) à t_1 .

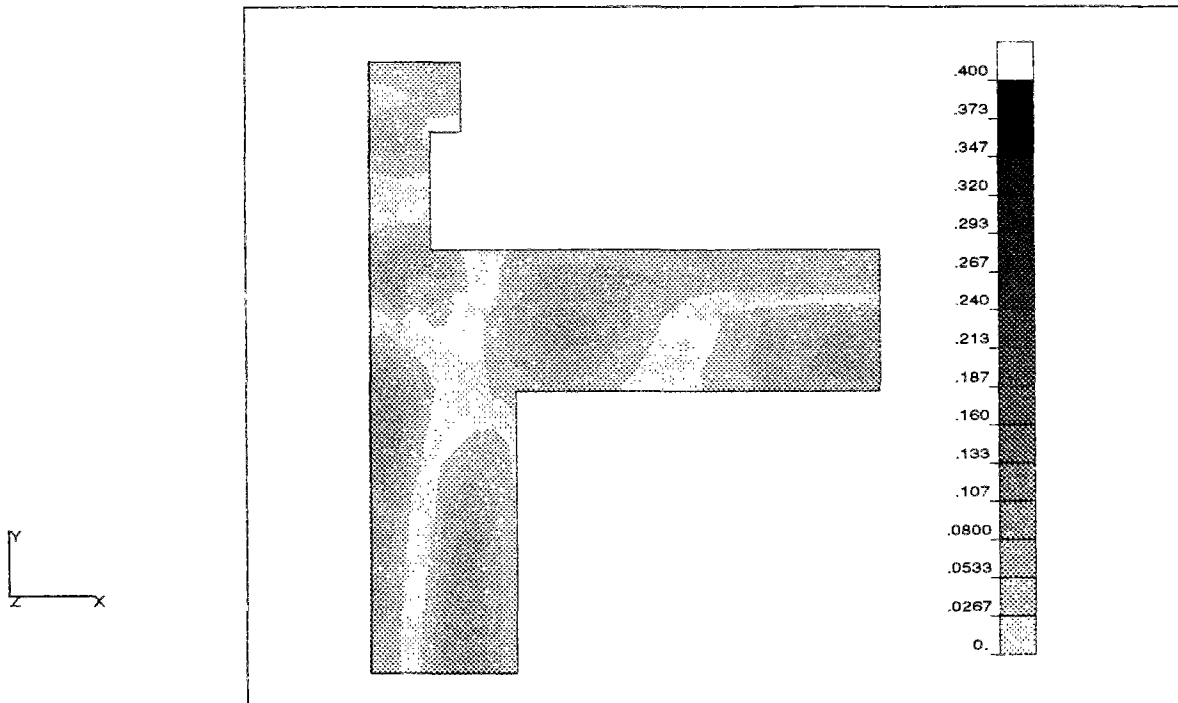


Figure 6.30: Sollicitation T_{ext} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Mixte) à t_2 .

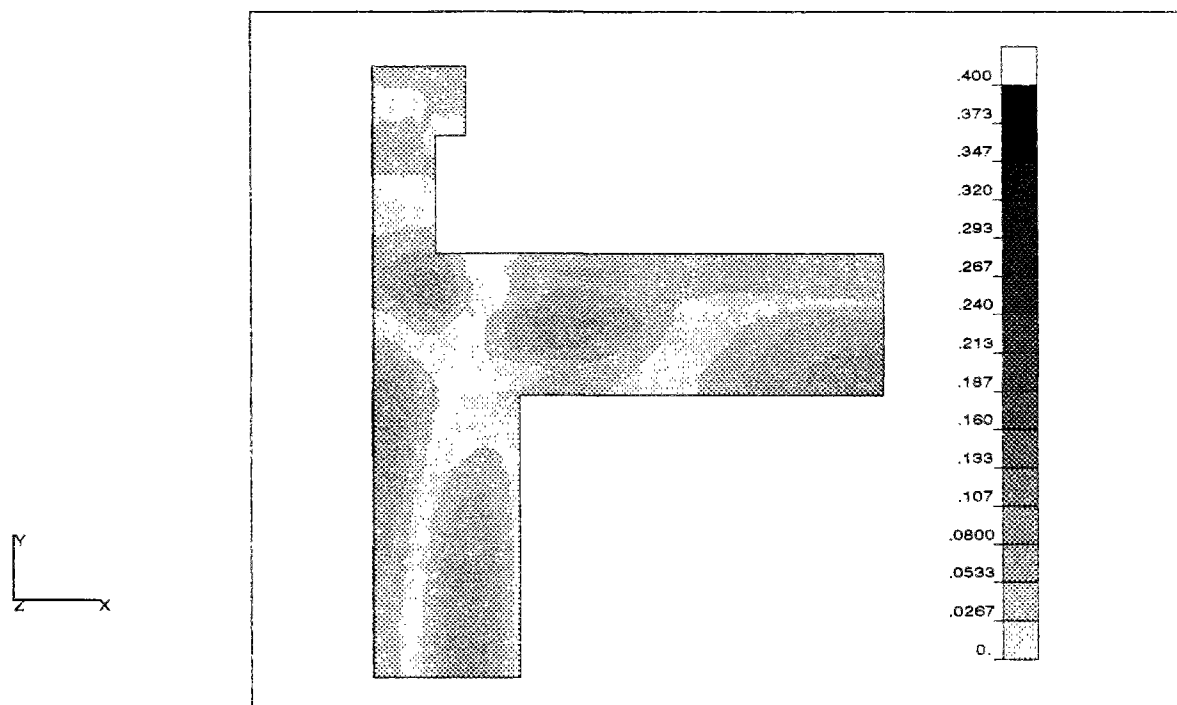


Figure 6.31: Sollicitation T_{int} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Mixte) à t_2 .

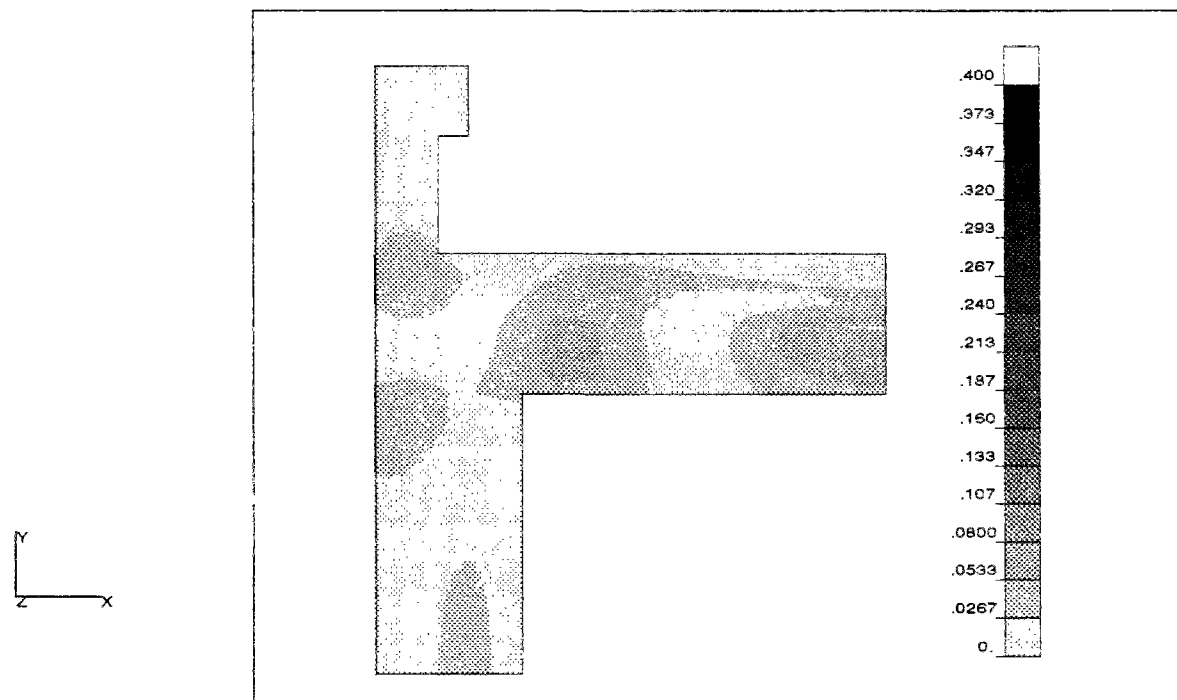


Figure 6.32: Sollicitation T_{ext} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Mixte) à t_3 .

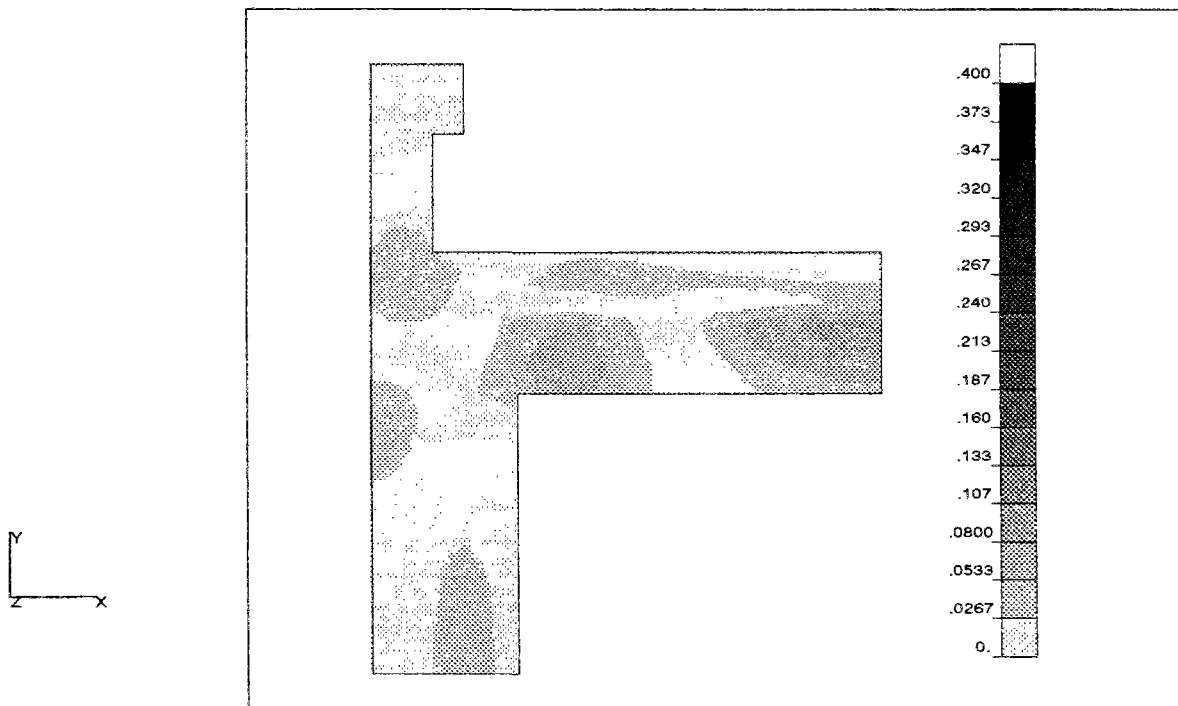


Figure 6.33: Sollicitation T_{int} : écart de température (en $^{\circ}C$) entre (M.Ref) et (M.Mixte) à t_3 .

Chapitre 7

Vilebrequin

7.1 Modèle de référence

Le système thermique que nous abordons ici est un vilebrequin d'un moteur automobile. Ce système a été déjà étudié par Flament [28] pour appliquer la technique de la synthèse modale. A travers cet exemple, nous voulons montrer la possibilité d'appliquer la méthode de réduction proposée en géométrie tridimensionnelle.

La forme et les dimensions du vilebrequin sont données dans la figure 7.1. Le métal constitutif de la structure est l'acier trempé dont les caractéristiques thermophysiques sont données dans le tableau 7.1.

Acier trempé		
Conductivité	masse volumique	C_p
45 [$Wm^{-1}K^{-1}$]	7800 [kgm^{-3}]	460 [$Jkg^{-1}K^{-1}$]

Tableau 7.1: Propriétés thermophysiques du vilebrequin.

Pour modéliser ce système, nous avons retenu des conditions aux limites de type Dirichlet sur les faces (externes) du vilebrequin en contact avec l'environnement extérieur.

Avant d'appliquer la réduction par amalgame modal, il faut d'abord disposer d'un modèle

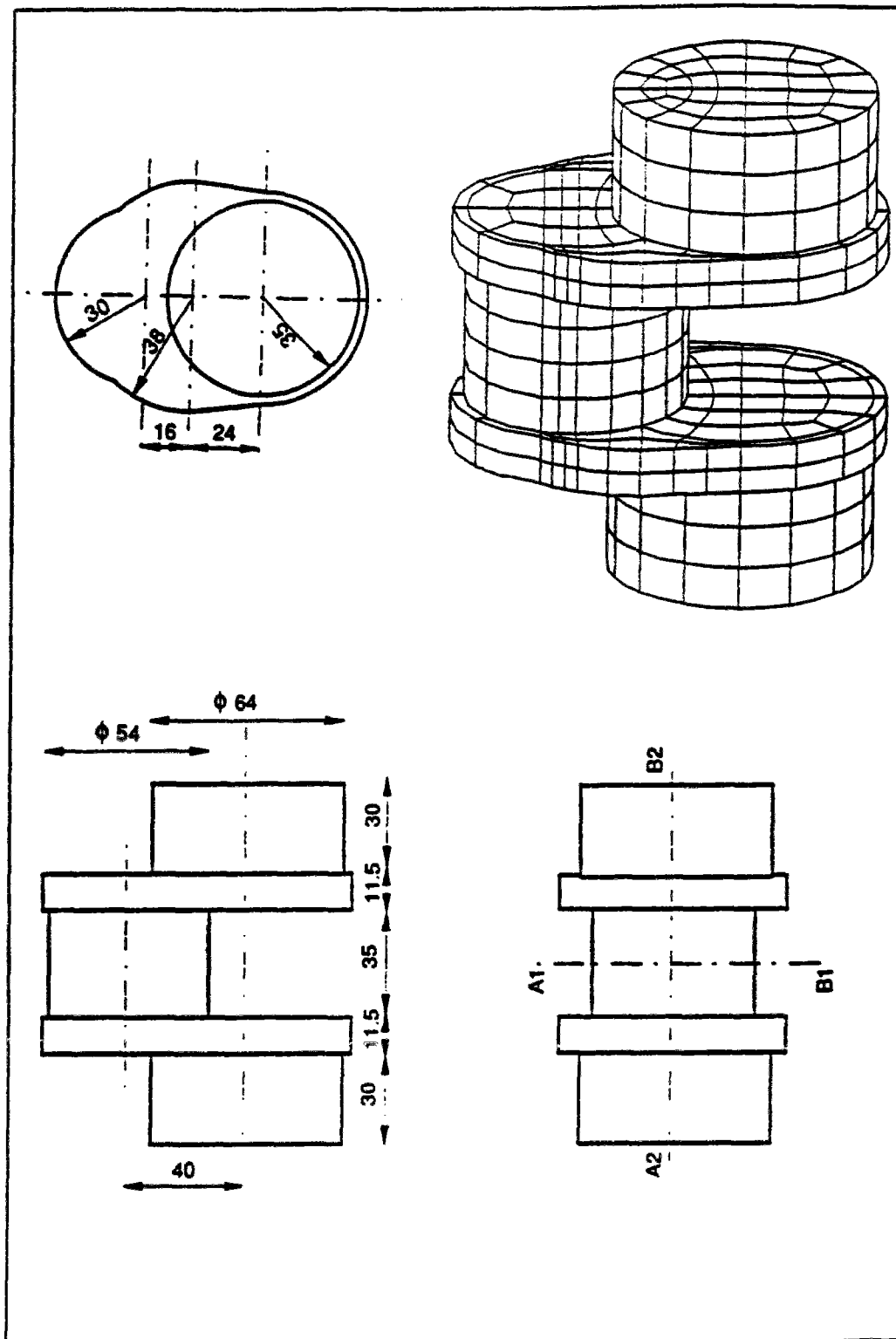


Figure 7.1: Description du vilebrequin.

modal pour le vilebrequin. On a alors construit un modèle modal¹ de dimension 50, c'est à dire reposant sur 50 modes propres. Le nombre de nœuds de la discrétisation spatiale est de 4015 (2165 nœuds libres et 1850 nœuds de surface). Le nombre de modes est relativement faible² mais n'influe pas sur l'application de la méthode de réduction. En revanche, la simulation du système ne peut être précise qu'à partir de trois à quatre fois la dernière constante de temps calculée.

On donne dans le tableau 7.2 les 50 constantes de temps du vilebrequin (dans l'ordre décroissant de leurs valeurs). Elles sont de l'ordre de la seconde. On obtient ici des groupements de constantes de temps multiples comme (6,7), (20,21), etc. A l'heure actuelle, nous ne connaissons pas de justification théorique de ce résultat. La multiplicité des valeurs propres obtenues dans l'exemple 1D du bâtiment bizone (voir §5) a été démontrée [48] et vérifiée dans d'autres travaux [13] On peut penser que la multiplicité des valeurs propres obtenues sur cet exemple est liée aux plans de symétrie du vilebrequin.

constantes de temps									
N°	τ_i	N°	τ_i	N°	τ_i	N°	τ_i	N°	τ_i
1	7.54s	2	7.10s	3	7.02s	4	4.36s	5	4.10s
6	4.00s	7	4.00s	8	3.92s	9	3.51s	10	3.47s
11	2.90s	12	2.84s	13	2.67s	14	2.62s	15	2.60s
16	2.56s	17	2.51s	18	2.48s	19	2.41s	20	2.24s
21	2.24s	22	2.18s	23	2.18s	24	2.16s	25	2.15s
26	2.05s	27	2.02s	28	1.82s	29	1.82s	30	1.82s
31	1.78s	32	1.76s	33	1.74s	34	1.73s	35	1.72s
36	1.72s	37	1.66s	38	1.63s	39	1.62s	40	1.62s
41	1.57s	42	1.52s	43	1.50s	44	1.48s	45	1.47s
46	1.47s	47	1.47s	48	1.42s	49	1.39s	50	1.38s

Tableau 7.2: Constantes de temps du Vilebrequin.

7.2 Amalgame modal

On veut réduire ici le modèle détaillé construit ci-dessus par amalgame modal. On impose simplement la dimension du modèle réduit égale à 8. A cet ordre ($n=8$), les valeurs de $\mathcal{M}$ et $\overline{\mathcal{M}}$ sont respectivement $0.01979 J^\circ C s$ et $0.00029^\circ C$.

¹Le modèle de départ discrétisé dans l'espace a été obtenu par la méthode d'approximation par éléments finis (éléments de type Lagrange à 20 nœuds). Le modèle modal est ensuite obtenu par transposition des matrices éléments finis dans l'espace d'état modal. La principale étape de cette transposition consiste à diagonaliser le couple des matrices d'échange et de capacité (technique exposée au chapitre 1).

²En fait, la limitation sur les dimensions des modèles est liée à l'espace disque des machines et leurs mémoires vives. Avec 50 modes et 4015 nœuds, le modèle occupe un espace-disque voisin 3 mégaoctets. On a alors préféré de garder un maillage fin pour avoir une meilleure précision dans des résultats dans l'espace, quitte à simuler la structure à partir d'un instant relativement grand.

sous-espaces d'amalgame			
1	2	3	4
modes	modes	modes	modes
1 (1.0000)	3 (1.0000) 2 (0.0006)	5 (1.0000) 4 (-0.0012) 6 (0.0002) 7 (0.0004) 8 (-0.0008) 9 (-0.2787) 10 (-0.0014)	12 (1.0000) 11 (-0.0013) 13 (-0.0007) 14 (-0.0060)

5	6	7	8
modes	modes	modes	modes
19 (1.0000) 15 (-0.3765) 16 (-0.0006) 17 (0.0011) 18 (0.0004)	21 (1.0000) 20 (-0.0039) 22 (-0.3813) 23 (0.0003) 24 (0.0008) 25 (-0.0001) 26 (0.0002) 27 (-0.0175)	32 (1.0000) 28 (-0.0024) 29 (0.2323) 30 (-0.0011) 31 (-0.0050) 33 (-0.2724) 34 (0.0059) 35 (0.0036) 36 (0.0015) 37 (0.0011) 38 (0.0013) 39 (0.0007) 40 (-0.2223) 41 (0.0004)	50 (1.0000) 42 (-0.1115) 43 (0.0001) 44 (0.0001) 45 (-0.0005) 46 (-0.0008) 47 (-0.1634) 48 (-0.0001) 49 (-0.0035)

Tableau 7.3: Répartition de l'espace des modes propres en 8 sous-espaces d'amalgame

Les sous-espaces d'amalgame sont donnés dans le tableau 7.3. Les modes principaux sont encadrés $\boxed{}$ et les valeurs (.) représentent les coefficients de décomposition associés aux modes. On constate que les modes principaux d'amalgame sont répartis selon toutes les dynamiques. En particulier les premier et dernier modes ($V_1(M)$ et $V_{50}(M)$) appartiennent à l'espace des modes principaux. De plus, la répartition des modes mineurs est ici facile à interpréter: globalement chaque mode mineur est affecté au mode principal qui lui est le plus proche temporellement. On peut penser logiquement que cette partition aboutit à un bon modèle réduit et nous allons le vérifier ci-dessous avec des simulations. Les partitions obtenues lors des exemples traités précédemment (milieux homogènes par morceaux) sont moins évidentes à interpréter.

Note: Pour les tracés des différents thermogrammes (modes propres, champ de

température ou champ écart) qui suivent, nous utilisons le plan repéré A_2B_2 de la figure 7.1.

On donne dans la figure 7.2 le premier mode propre du vilebrequin. Ce mode concerne l'ensemble de la structure et présente une symétrie évidente. Ce mode propre est aussi le premier mode amalgamé puisqu'il est principal et n'est entouré d'aucun mode mineur.

Il n'est pas facile de faire une analyse fine de la manière dont les modes principaux sont modifiés par l'amalgame car il faut visualiser les modes en trois dimensions. Nous allons néanmoins faire quelques constatations intéressantes sur les sous-espaces d'amalgame H_5 et H_6 (se référer au tableau 7.3).

► sous-espace H_5 : ce sous-espace repose sur le mode principal $V_{19}(M)$ (figure 7.3) et les modes mineurs 15, 16, 17 et 18. On obtient de ce sous-espace le mode amalgamé $\tilde{V}_5(M)$ donné dans la figure 7.4. La comparaison du mode principal et du mode amalgamé montre qu'il ne sont pas très différents dans leurs formes (sur le plan A_2B_2). Remarquons toutefois que l'amalgame modal (voir figure 7.4) conserve parfaitement la symétrie des modes propres.

► sous-espace H_6 : ce sous-espace est formé du mode principal $V_{21}(M)$ et de sept modes mineurs. Parmi ces modes, il y a deux couples de constantes de temps numériquement très proches: les couples (20,21) et (22,23). Les modes 21 et 22 ont les plus grands coefficients de décomposition (1.0 et -0.3813 respectivement) du sous-espace, alors que ceux des modes 20 et 23 sont faibles. On retrouve ici une similitude avec les résultats de l'exemple du bâtiment: chaque groupement de modes multiples fait apparaître un seul mode dominant.

Malgré la présence de quelques modes non-symétriques (mode 20 par exemple), le mode amalgamé résultant est symétrique. Ceci est tout à fait logique vu que le vilebrequin est soumis à une condition aux limites de Dirichlet sur toute sa surface externe. Dans cet exemple, les modes non-symétriques n'interviennent que très faiblement dans l'amalgame modal.

Les modes 21 et 22 sont donnés dans les figures 7.5 et 7.6. D'après les coefficients de décomposition du tableau 7.3, on peut considérer que le mode amalgamé $\tilde{V}_6(M)$ est une combinaison de V_{21} et V_{22} seulement. Il est alors possible de montrer que la correction du mode principal par amalgame modal est forte. On considère que H_6 est généré par seulement V_{21} et V_{22} . H_6 est donc un sous-espace de dimension 2. On peut effectuer un changement de base qui revient ici à une rotation du repère formé par V_{21} et V_{21} . On impose que l'un des deux vecteurs de la nouvelle base soit $\tilde{V}_6(M)$. Ce dernier est approché par

$$\tilde{V}_6(M) = V_{21}(M) - 0.3813V_{22}(M)$$

Le deuxième vecteur (on le note $\tilde{V}_{6\perp}(M)$) de la nouvelle base peut être facilement calculé car il est orthogonal à $\tilde{V}_6(M)$ tel que

$$\langle \tilde{V}_6(M), C(M)\tilde{V}_{6\perp}(M) \rangle = 0$$

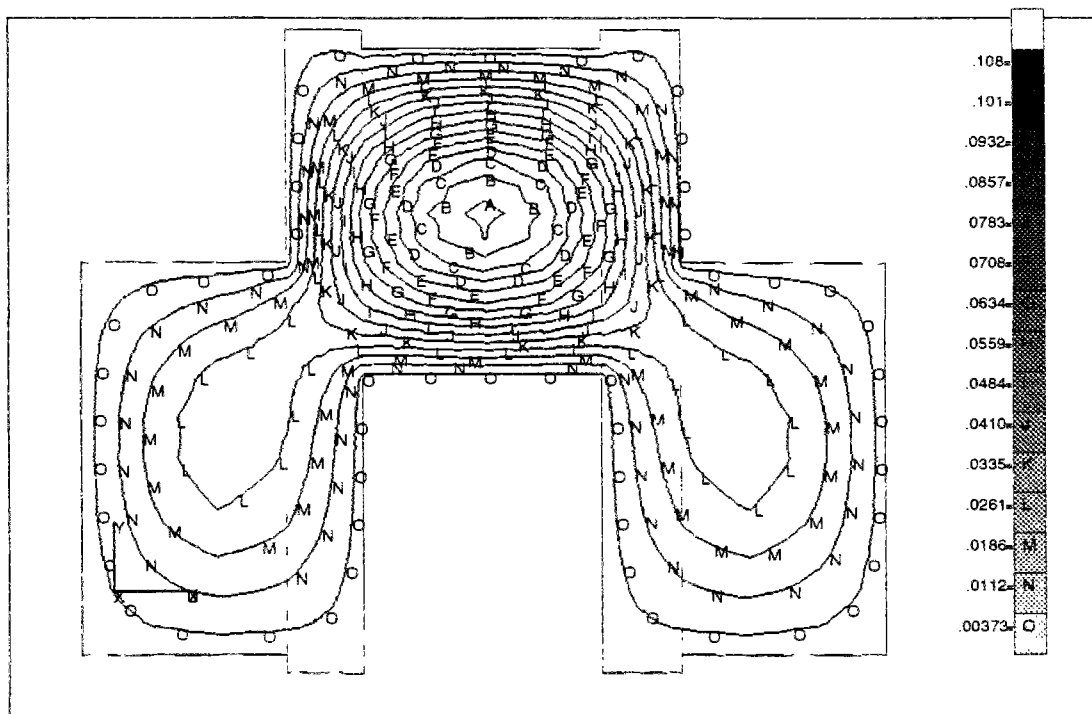


Figure 7.2: Mode propre 1 (identique au mode amalgamé 1). Isovaleur maximale $A=0.108$, écart entre deux isovaleurs 0.0072 .

L'amalgame modal est équivalent à approximer la direction du sous-espace H_6 par celle de $\tilde{V}_6(M)$. Cette approximation est d'autant plus réaliste que l'amplitude spatiale de $\tilde{V}_{6\perp}(M)$ est faible. On donne dans la figure 7.8 l'allure du "mode" $\tilde{V}_{6\perp}(M)$ selon le plan A_2B_2 et avec la même échelle que celle de $\tilde{V}_6(M)$: on constate que son amplitude est quasi-nulle dans tout le plan A_2B_2 . En fait ce résultat montre que la direction de $\tilde{V}_6(M)$ est bien celle du sous-espace H_6 . Les "modes" $\tilde{V}_6(M)$ et $\tilde{V}_{6\perp}(M)$ sont tracés dans la figure 7.9 dans l'espace 3D avec des surfaces d'isovaleurs repérées par des niveaux de gris différents. L'interprétation à partir de cette figure n'est pas simple.

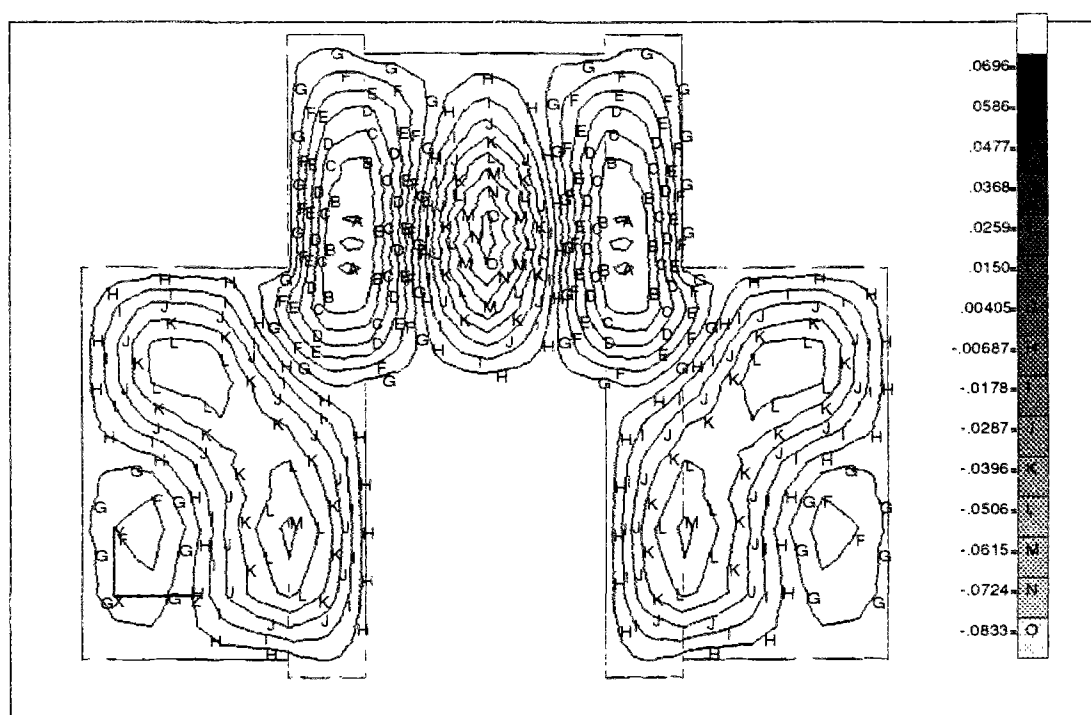


Figure 7.3: Mode propre 19. Isovaleur maximale $A=0.0696$, écart entre deux iso-valeurs 0.00464.

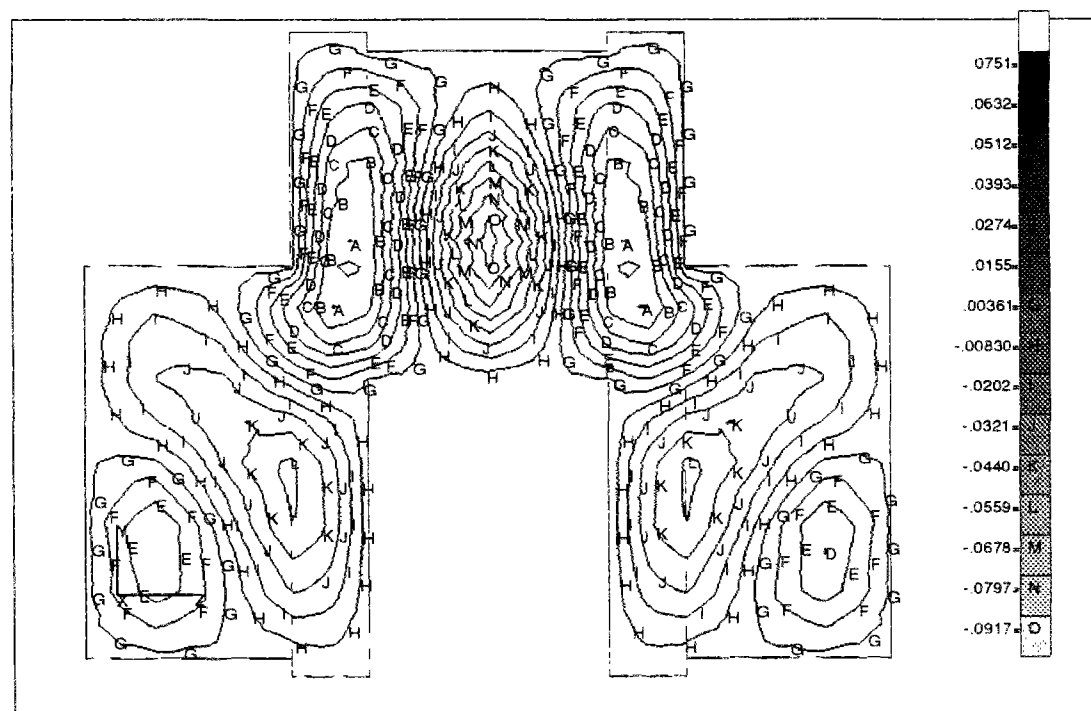


Figure 7.4: Mode amalgamé 5. Isovaleur maximale $A=0.075$, écart entre deux iso-valeurs 0.005.

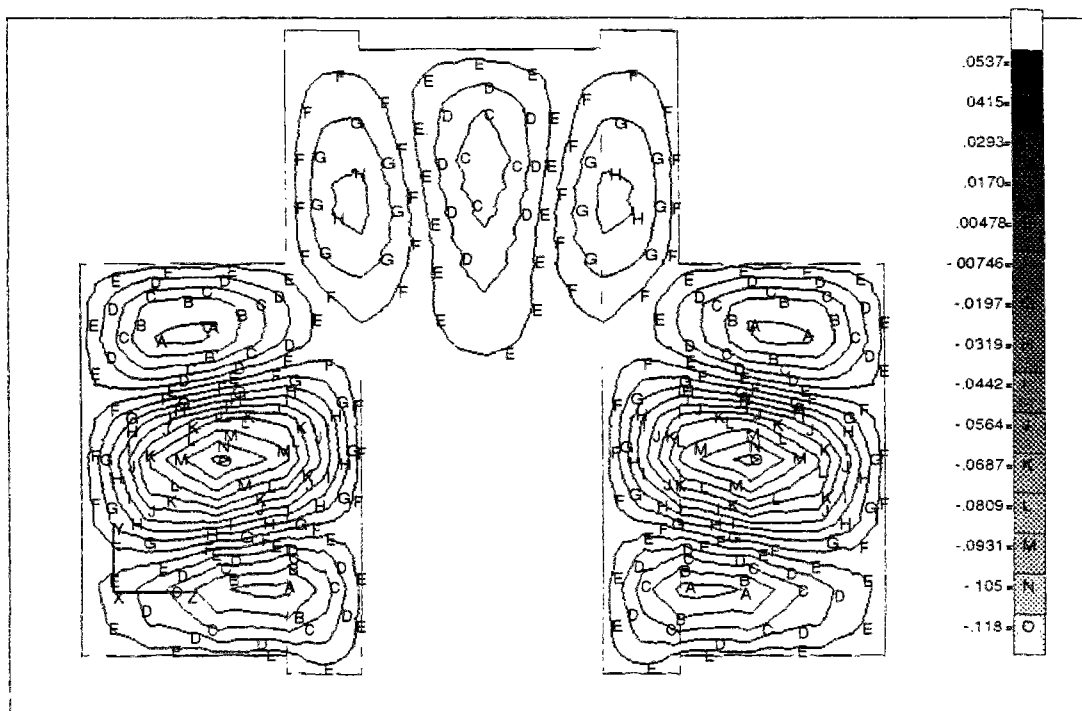


Figure 7.5: Mode propre 21. Isovaleurs max: A=0.0537 min: O=-0.118.

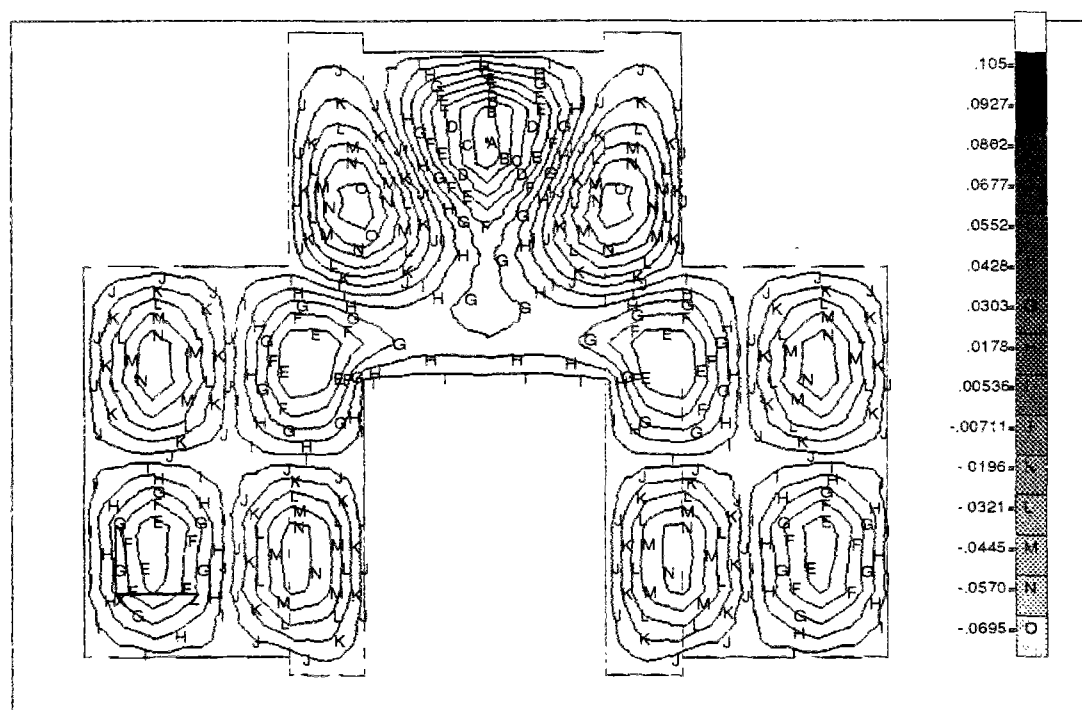


Figure 7.6: Mode propre 22. Isovaleur max: A=0.105 min: O=-0.0695.

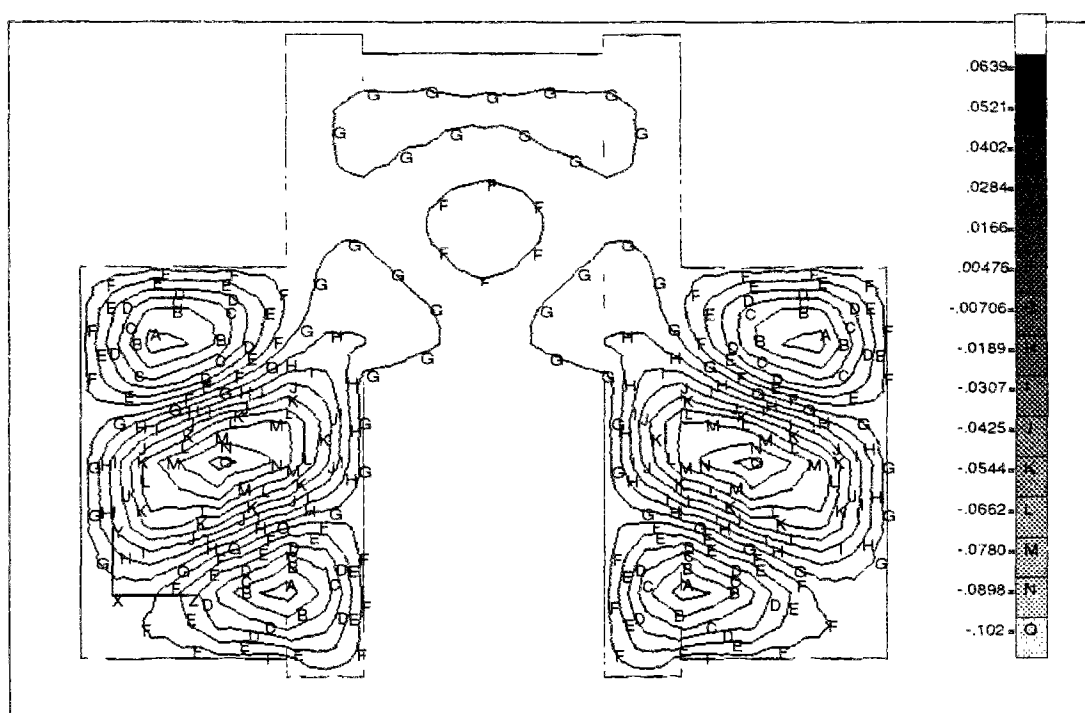


Figure 7.7: Mode amalgamé 6. Isovaleur max: A=0.0639 min: O=-0.102.

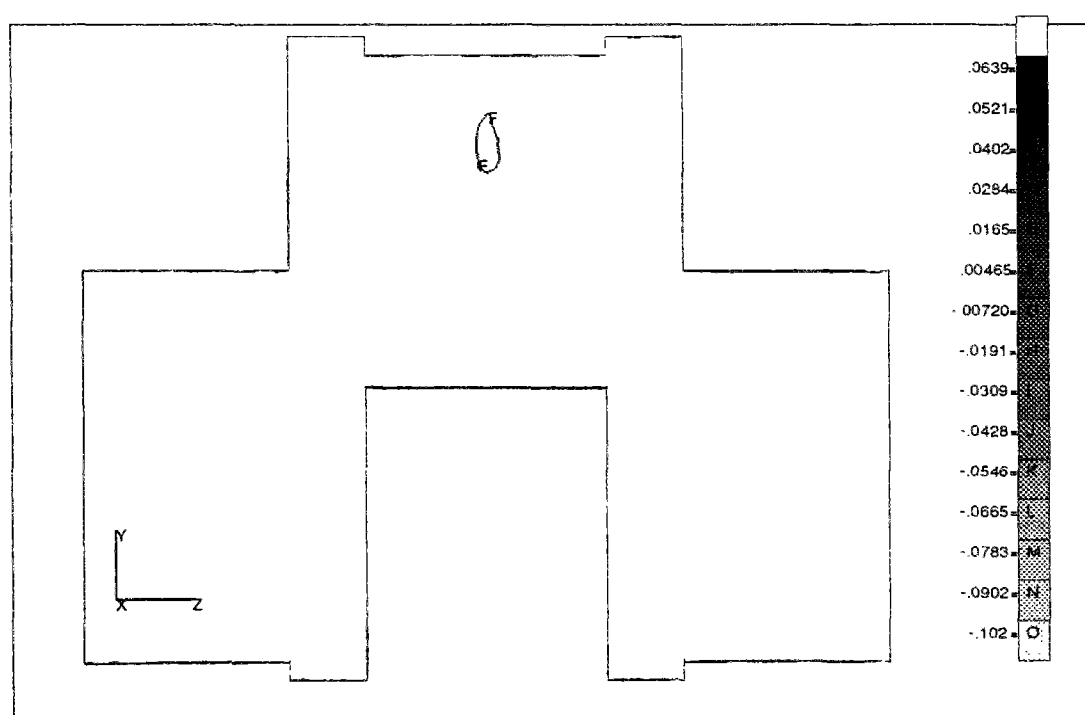


Figure 7.8: Mode amalgamé 6 $\perp$. Isovaleur max: A=0.0639 min: O=-0.102.

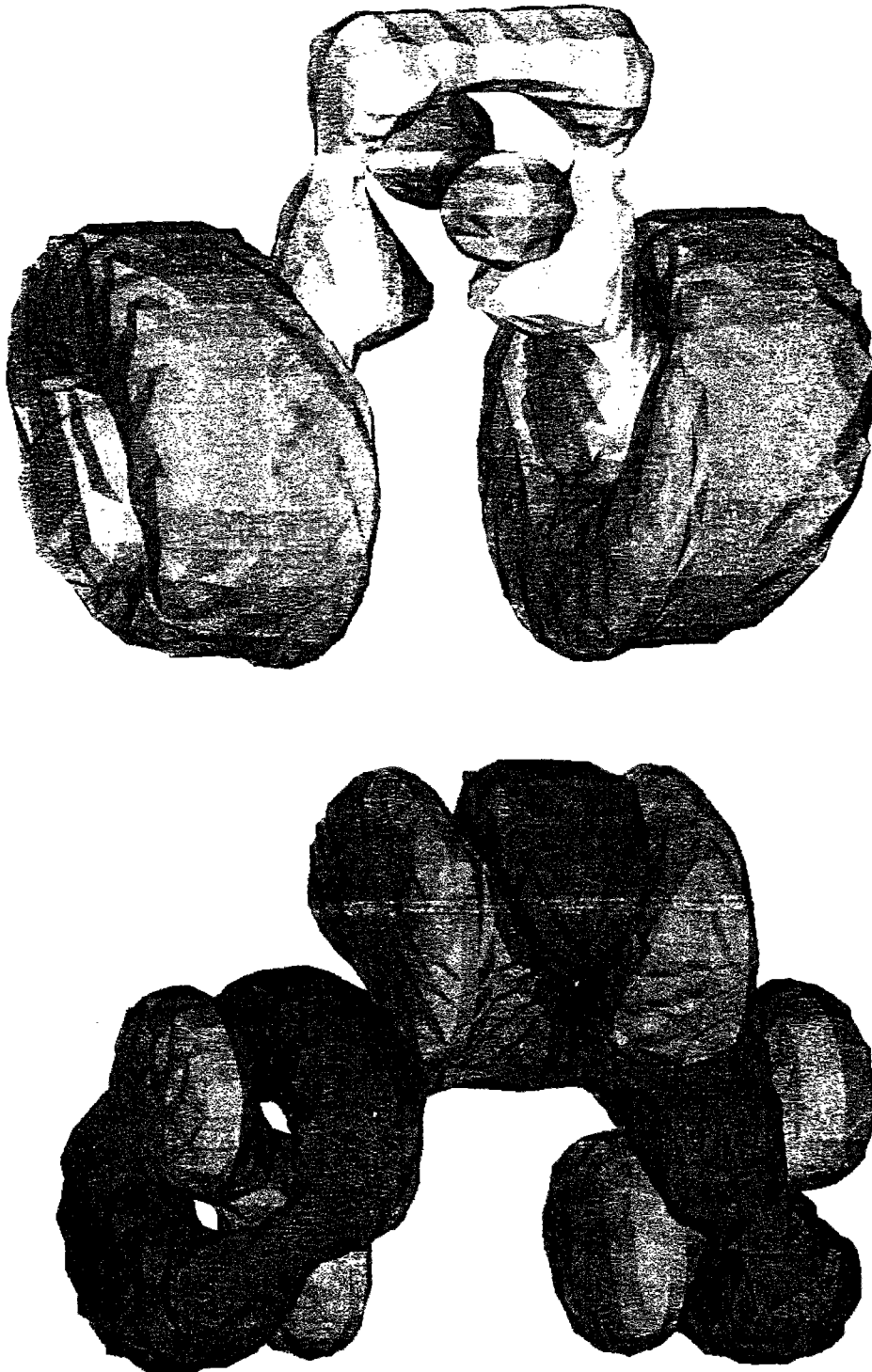


Figure 7.9: Les isovaleurs sont représentés par des niveaux de gris dans l'ordre décroissant. En haut le mode amalgamé $V_6(M)$ (isovaleurs max:0.07 ,min:-0.11). En bas le "mode" $V_{6\perp}(M)$ (isovaleurs max:0.005 ,min:-0.003).

7.3 Simulation de l'évolution thermique

Nous allons effectuer une simulation d'un cas de trempe en milieu froid du vilebrequin. Il ne s'agit pas d'un problème réel de trempe en raison des fortes non-linéarités (variations de λ , ρ , C_p et phénomènes complexes liés aux changements de phases), et dans ce cas le formalisme d'état modal donné au §1 ne s'applique pas. Toutefois, à partir d'un certain instant, les gradients thermiques deviennent relativement faibles et l'hypothèse de linéarité peut être admise. C'est dans cette dernière situation que nous nous plaçons dans cette section.

Les résultats qui suivent sont obtenus pour les conditions de trempe suivantes:

- ▶ la température initiale du vilebrequin est $800^\circ C$ (sortie du four)
- ▶ la température du bain de trempe est $140^\circ C$.

La dernière constante de temps τ_{50} étant de 1.38s, nous montrons les résultats aux deux instants suivants:

- ▶ $t_1 = 5s$
- ▶ $t_2 = 8s$

On donne dans les figures 7.10 et 7.11 les champs de température du vilebrequin (suivant le plan de symétrie A_2B_2 de la figure 7.1). Les figures 7.12 et 7.13 représentent les champs écarts (en valeur absolue) entre les modèles de référence et amalgamé. L'écart maximal observé est de $3^\circ C$. Cet écart n'est pas grand par rapport aux valeurs des températures à l'instant t_1 qui atteignent $700^\circ C$. A l'instant t_2 , ces erreurs sont bien atténuées.

Ces résultats montrent la faisabilité d'appliquer la méthode d'amalgame modal pour réduire un modèle thermique de grande dimension en géométrie tridimensionnelle. Ceci était prévisible car la formulation modale est standard: un modèle modal a strictement la même forme mathématique et les mêmes propriétés quelle que soit la géométrie (1D, 2D, 3D). Un outil informatique pour la visualisation de champs tridimensionnels (modes, température, champs écart, etc) serait sans doute d'une grande utilité.

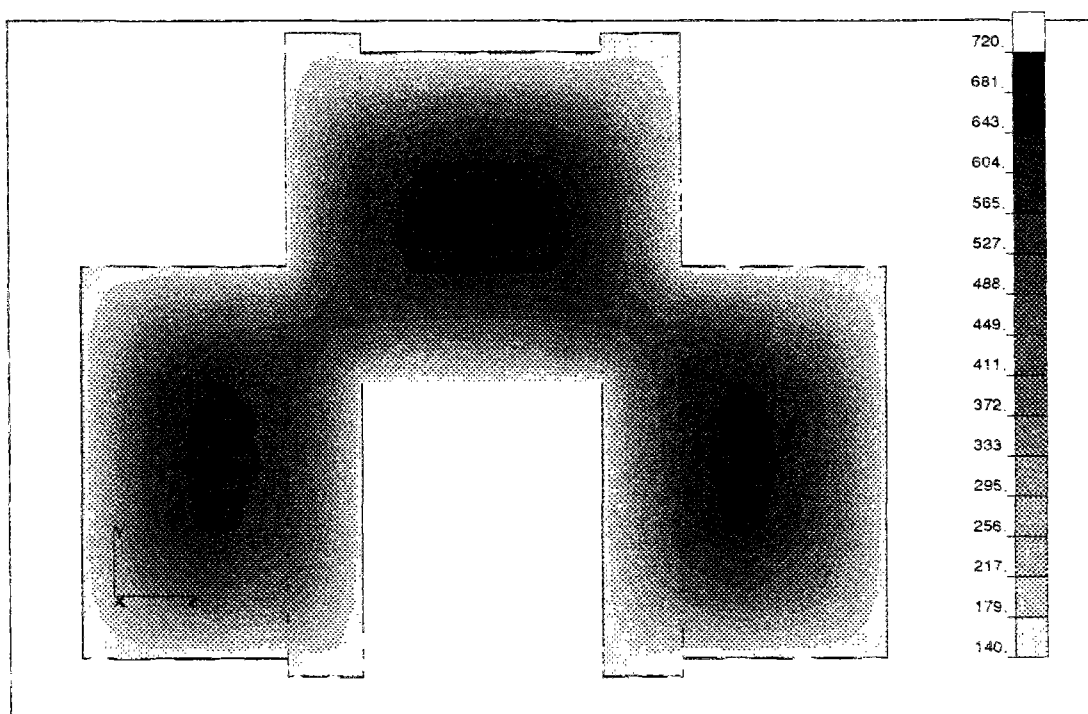


Figure 7.10: Champ de température à l'instant t_1 .

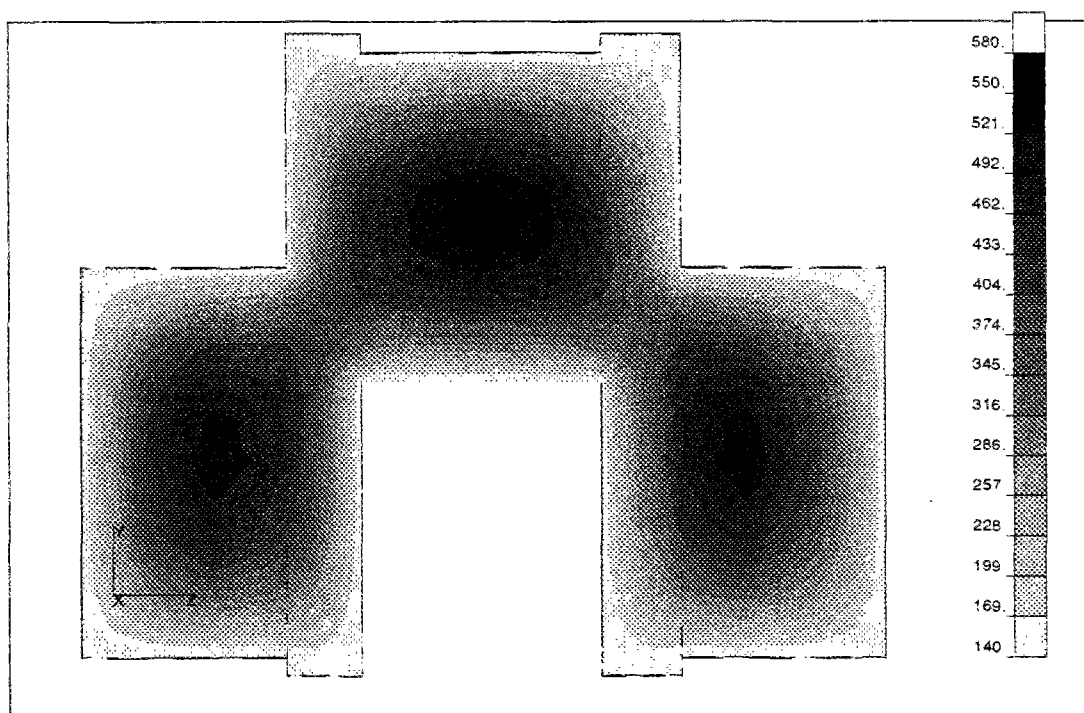
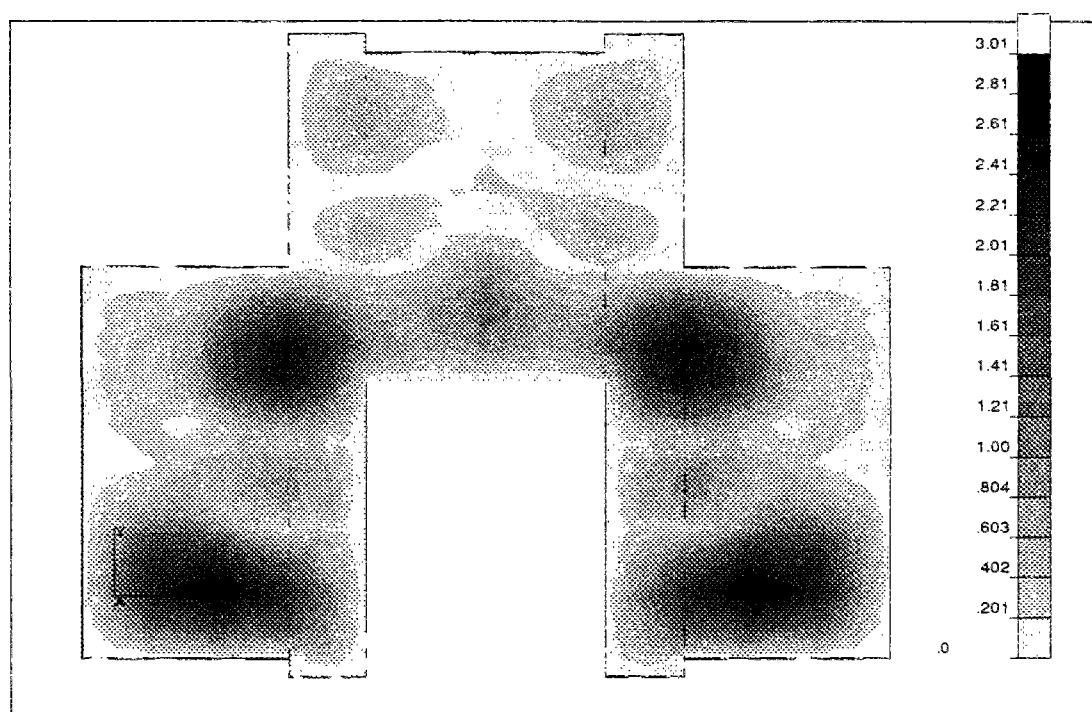
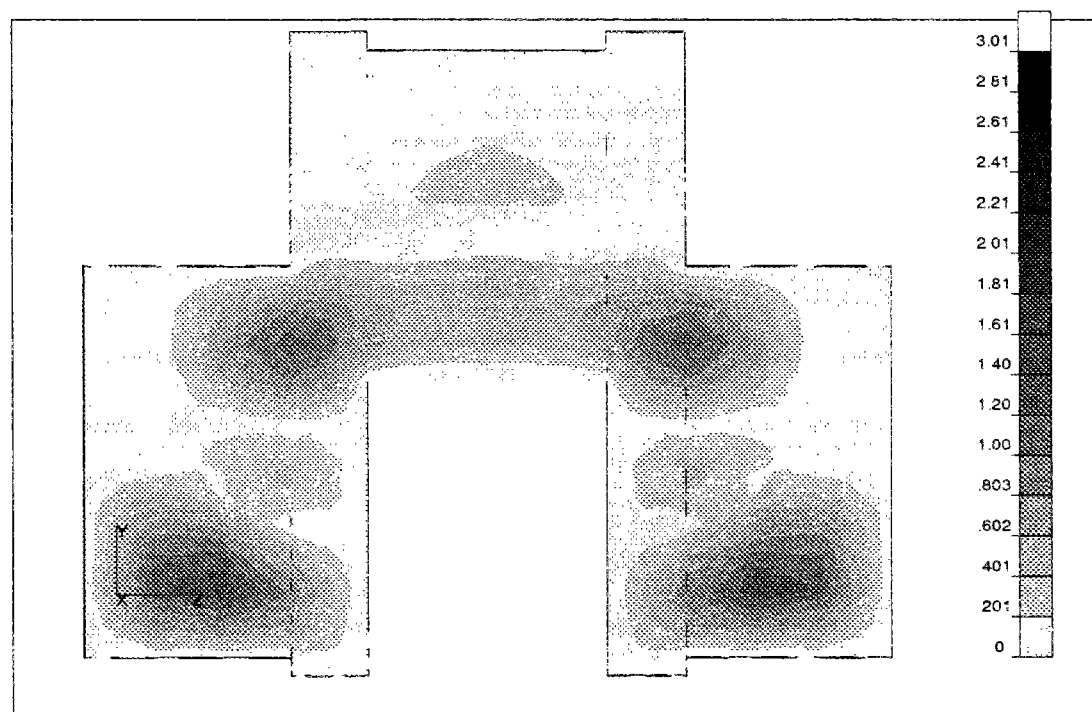


Figure 7.11: Champ de température à l'instant t_2 .

Figure 7.12: Champ écart de température à t_1 .Figure 7.13: Champ écart de température à t_2 .

Chapitre 8

Conclusion & perspectives

La réduction de modèles thermiques est un domaine relativement récent. Les méthodes développées à ce sujet sont cependant déjà nombreuses. Elles ont toutes contribué - chacune à sa manière - à enrichir le sujet. Mais leur diversité témoigne aussi de la difficulté de condenser l'information thermique dans un nombre très restreint de paramètres et de variables.

Dans le travail présenté ici, nous nous sommes intéressés principalement à deux classes de méthodes:

- ▶ les méthodes de troncature qui ordonnent les modes propres selon un critère. Seuls les premiers modes du classement obtenu sont gardés.
- ▶ les méthodes de minimisation qui sont plus sophistiquées que les précédentes. Elles font appel à des techniques de minimisation d'un critère souvent quadratique.

Les critères de troncature décrits et testés sont au nombre de quatre: Marshall qui ordonne les modes propres dans l'ordre décroissant de leurs constantes de temps, le critère de Litz qui repose sur l'examen du spectre de réponse indicielle d'un couple entrée-sortie et enfin les critères énergétiques (dominance et dominance pondérée). Le choix d'un critère se justifie surtout par le but de la réduction. Si l'on s'intéresse à un couple entrée-sortie, alors le critère de Litz est généralement efficace. Notons que le critère de Litz peut être amélioré en tenant compte des compensations éventuelles des raies proches temporellement d'une part et de l'amortissement du terme exponentiel d'autre part. Les autres critères sont globaux, c'est à dire ne sont pas spécifiques à un couple entrée sortie. De façon général, le critère de Marshall aboutit à un modèle réduit insuffisant pour les instants courts. Le critère de dominance, reposant sur l'aspect à la fois spatial et temporel des modes propres, aboutit à des modèles réduits certes biaisés, mais plus précis que ceux du premier critère. Nous avons ensuite introduit un nouveau critère (dominance pondérée) qui permet globalement

d'éviter la localisation des biais dans des plages temporelles étroites. Indépendamment de la qualité de ces critères, on ne peut écarter définitivement la troncature modale: par sa simplicité, elle reste d'une grande utilité dans les travaux relevant de l'estimation simplifiée (évolution qualitative d'une sortie, besoin de chauffage, déperdition thermiques...).

En ce qui concerne la seconde classe de méthodes (minimisation), nous avons résumé formellement trois d'entre elles: les méthodes d'agrégation, d'Eitelberg et mixte "réduction-identification". Ces méthodes sont robustes et aboutissent à des modèles réduits précis.

L'approche d'Eitelberg qui opère dans l'espace physique, peut permettre la prise en compte de quelques non-linéarités (échanges radiatifs par exemple), en revanche elle nécessite le choix de quelques nœuds "importants". Ceci fait appel à l'expérience du modélisateur. La méthode n'étant pas modale, nous ne l'avons pas expérimentée.

L'approche d'agrégation, plus proche que les autres de la méthode d'amalgame, nécessite la sélection de quelques modes dans une première étape. Cette sélection est effectuée au moyen d'un critère dit "de contributions énergétiques". Nous avons montré sur l'exemple du bâtiment bizona que cette sélection n'est pas rigoureusement optimale. Par ailleurs, les "pseudo-modes" issus de l'agrégation ne sont pas orthogonaux entre eux relativement à la mesure $C(M)$. Ces pseudo-modes s'apparentent à une synthèse - optimale mais globale - de l'espace des modes propres. Les modèles agrégés que nous avons testés sont performants dans la reconstitution du comportement thermique.

L'approche mixte réduction-identification est plus récente. Elle s'adapte aussi bien à la réduction d'un modèle de connaissance qu'à l'identification de modèles réduits à partir de données expérimentales. Comparée aux autres méthodes d'identification, celle-ci permet d'identifier un grand nombre de paramètres, notamment d'utiliser un vecteur des sorties comprenant beaucoup de composantes (le champ de température complet par exemple). Nous nous sommes intéressé à la version adaptée à la réduction de modèles de connaissance. On identifie d'abord les modes propres et les constantes de temps associées. Cette phase est réalisable à partir de thermogrammes obtenus par simulation des sorties grâce au modèle de connaissance. En fait par cette étape on évite aussi la sélection des modes propres à partir du modèle détaillé, les modes identifiés apparaissant "naturellement" lors d'expérimentations numériques. En revanche, comme nous l'avons dit, les modes identifiables sont liés au thermogrammes intervenant dans les expérimentations numériques. Le nombre et la dominance des modes identifiés dépendent des modes excités par les sollicitations lors de la génération des thermogrammes. Le modèle réduit obtenu ici n'a pas de lien explicite avec le modèle d'origine et les temps de calculs ne sont pas négligeables.

Pour notre objectif où l'on s'intéresse à la totalité du champ de température, la troncature modale ne permet pas d'avoir à la fois des modèles fortement réduits et suffisamment pertinents. Les méthodes de minimisation citées, bien que performantes, ne répondent pas à nos contraintes (notamment sur la conservation de l'orthogonalité et le temps de

calcul). Ceci nous a donc amené à développer une nouvelle méthode de réduction: la méthode d'amalgame modal.

La méthode d'amalgame modal appartient à la deuxième classe de méthodes (minimisation). Nous avons d'abord précisé la mesure de l'erreur sur laquelle la méthode est bâtie. Ensuite, nous avons présenté la méthode d'amalgame elle-même. Rappelons que la méthode contient deux étapes fondamentales:

- la réalisation d'une partition de l'espace d'état modal. Dans la pratique cette étape est menée en deux phases:
 - la sélection des modes principaux.
 - la répartition des modes mineurs autour des modes principaux pour former les sous-espaces orthogonaux.

La seconde partie (répartition des modes mineurs) est optimale au sens de la mesure retenue. La première (sélection des modes principaux) peut être réalisée de façon optimale, mais avec un temps de calcul très pénalisant. Nous avons alors trouvé un algorithme pour la sélection des modes principaux qui est très rapide et qui permet d'aboutir

- soit à la meilleure combinaison
 - ou bien à une combinaison suffisamment proche de l'optimum
- au sens de la mesure de l'erreur retenue, et qui correspond à un minimum local.

- l'obtention d'un mode amalgamé à partir de chaque sous-espace d'amalgame. On récupère ici une information sur tous les modes mineurs qui va ensuite "corriger" le mode principal. Cette correction est optimale au sens de la mesure de l'erreur.

Ces différentes étapes de la méthode sont -dans la pratique- indépendantes les unes des autres. De plus, la réalisation de ces étapes se fait à partir de relations totalement explicites. Rappelons qu'il ne s'agit pas ici d'une simplification du problème: c'est une conséquence qui découle du choix (contrainte) sur l'orthogonalité des modes amalgamés. En effet cette contrainte a pour effet de faire apparaître formellement la relation d'orthonormalité des modes propres, ce qui découple les contributions modales à la mesure de l'erreur.

Une fois l'amalgame modal présenté, nous avons donné une méthode pour "ajuster" les constantes de temps. Cependant, les tests effectués ont montré que la diminution de la mesure de l'erreur résultant de cet ajustement reste insignifiante pour la méthode proposée. Etant donné que cette étape n'a pratiquement pas d'incidence sur le temps de calcul, nous avons jugé intéressant de l'implémenter dans la méthode.

Dans le cas où le modèle de référence comporte des sollicitations de différentes natures, la pondération sur les sollicitations est nécessaire. L'exemple du bâtiment bizone a permis de montrer l'importance de cette pondération.

Nous avons enfin effectué une synthèse de notre travail en proposant un algorithme général pour la méthode d'amalgame. Une estimation du temps de calcul minimal nécessaire au fonctionnement de la méthode a été donnée.

Les exemples traités ont permis d'appliquer la méthode proposée en géométries mono et multidimensionnelle. Notre contribution au sujet de la réduction de modèles modaux est d'avoir formulé une nouvelle méthode qui

est performante en temps de calcul car reposant sur une formulation explicite

permet d'augmenter facilement l'ordre du modèle réduit jusqu'à ce que ce dernier soit jugé satisfaisant

aboutit à un modèle réduit qui sauvegarde la "forme modale". En particulier, les paramètres du modèle réduit sont reliés explicitement à ceux du modèle de référence.

L'un des points qu'on souhaite améliorer est l'algorithme de sélection des modes principaux. Pour garantir de façon "certaine" la sélection de la "meilleure" (au sens de la mesure de l'erreur) combinaison de modes principaux, nous ne connaissons à l'heure actuelle que la technique triviale: comparer toutes les combinaisons possibles. Or il semble que des algorithmes spécifiques [34] existent pour ce type de problèmes combinatoires.

Le choix des modes principaux ainsi que la répartition des modes mineurs dans les différents sous-espaces sont réalisés en considérant des sollicitations en formes d'échelons unitaires. Or le modèle réduit est souvent destiné à servir pour d'autres formes de sollicitations ou des relevés expérimentaux. Il est probable que l'utilisation des sollicitations stochastiques conduit à des modèles réduits plus pertinents. Ce point va être étudié prochainement en collaboration avec l'Instituto de Energia Renovable¹.

La méthode modale a été déjà généralisée aux problèmes thermiques faisant intervenir un transfert de masse. L'opérateur de la chaleur est alors non-autoadjoint et ses solutions sont complexes (modes propres et valeurs propres). La manipulation de tels opérateurs n'est pas facile. Les éléments propres de l'opérateur de la chaleur vérifient une relation dite de bi-orthogonalité. Si, comme nous l'avons fait au §3.3.3, on arrive à exploiter cette nouvelle relation pour éliminer la variable spatiale de la mesure de l'erreur, alors il est probable qu'on pourra généraliser notre travail aux systèmes thermiques comprenant un transport d'enthalpie.

La synthèse modale [28] permet la construction de modèles modaux de systèmes thermiques à partir de ceux (dits modèles locaux) de composants plus petit. La réduction des modèles locaux associée à la synthèse modale est une voie séduisante. Au stade actuel, la méthode de synthèse modale opère à partir des modèles locaux réduits par troncature, car les éléments propres doivent être solution de l'opérateur $\mathcal{L}$ de la chaleur. On peut tenter de reprendre formellement les équations de la synthèse modale en considérant les

¹Instituto de Energia Renovable - CIEMAT
Avd. Complutense, 22. 28040-Madrid (Espagne)

modèles réduits par amalgame modal. En effet un mode amalgamé s'exprime linéairement en fonction des solutions de $\mathcal{L}$. Mais la synthèse modale soulève d'autres problèmes. En particulier, elle fait apparaître de nouveaux modes, dit "de couplage", et qui sont associés à un opérateur de "raccordement" (opérateur de type Poincaré-Steklov). Ici encore il existe des relations d'orthogonalité. L'amalgame modal est-il transposable? Les travaux à venir apporteront la réponse.

Références bibliographiques

Bibliographie

- [1] H. Li. *Intégration numérique de l'équation différentielle modale du modèle de l'évolution thermique d'un bâtiment*. Rapport de DEA, Université Paris VI, Sept 1992.
- [2] N.I Akhiezer. *Theory of linear operators in Hilbert Space*. Vol I, Frederick Ungar, New-York, Sept 1966.
- [3] J.D Alpevich. *Gradient methods for optimal linear system reduction*. Int. J. of Contr. Vol 18, P 767-772, 1973.
- [4] M. Aoki. *Control of large-Scale dynamic systems by aggregation*. IEEE Vol AC-13, P246-253, June 1968.
- [5] M. Athans. *The matrix minimum principal*. Information and control, Vol.11 P 592-606, 1968.
- [6] P. Bacot. *Identification des modèles de comportement des systèmes thermiques*. Revue générale de thermique N 277 P 15-21, 1985.
- [7] P. Bacot. *Analyse modale des systèmes thermiques*. Thèse de Docteur-Ingénieur. Université Paris VI, Fev. 1984.
- [8] P. Bacot, A. Neveu, and T. Sicard. *Analyse modale des phénomènes thermiques en régime variable dans le bâtiment*. Revue Générale de Thermique, N 267, P 189-201, Mars 1984.
- [9] R.H. Bartels and G.W. Stewart. *Solution of the matrix equation $AX+XB=C$* . Communications of the ACM Vol 15, N 9, Sept 1972.
- [10] M.T. Ben Jaafar. *Réduction de modèle de diffusion thermique: application des méthodes de Marshall, d'Eitelberg et d'Agrégation à un modèle de collecteur caloporteur*. Thèse. Université de Provence, Nov 1987.
- [11] M.T. Ben Jaafar, R. Pasquetti, and D. Petit. *Model reduction for thermal diffusion: Application to a heat transmission tube model*. Int. J. for Numerical methods of engineering, 1990.
- [12] G. J. Bierman. *Weighted least square stationary approximations to linear systems*. IEEE trans. automatic control, April 1972.

- [13] I. Blanc Sommereux. *Etude du couplage dynamique de composants du bâtiment par synthèse modale*. Thèse de Doctorat . Ecole des Mines de Paris, Oct 1989.
- [14] CENREG-EMP. *Document des résumés des communications et conférences du Colloque de thermique* . Ecole des Mines de Paris, Sophia-Antipolis, 20-21 Mai 1992.
- [15] M. Chen. *Nouvelle solution modale pour les parois de bâtiment*. Thèse de Doctorat . Ecole des Mines de Paris, Oct 1989.
- [16] M.R. Chidambara. *Two simple techniques for the simplification of large dynamic systems* . Proc. JACC, Université de Colorado, P 669-674, 1969.
- [17] P.G. Ciarlet. *Introduction à l'analyse numérique matriciel et à l'optimisation* . Editions Masson, 1990.
- [18] R. Dautray. *Analyse mathématique et calcul numérique pour les sciences et les techniques* . Collection CEA Masson, 1984.
- [19] E.J. Davison. *A method for simplifying linear dynamic systems* . IEEE TAC, Vol AC 11, P 93-101, 1966.
- [20] E.J. Davison. *A new method for simplifying large dynamic systems* . IEEE TAC, Vol AC 13, P 214-215, 1968.
- [21] M. Decoster, E. Noldus, and A. Van Cauwenberghe. *Réduction des systèmes linéaires*. RAIRO Automatique, P 47-66, 1976.
- [22] M. Decoster, E. Noldus, and A. R. Van Cauwenberghe. *Réduction des systèmes linéaires stationnaires et continus*. Journal A, Vol 17, N 3 P 58-74, 1976.
- [23] G. Duc. *Etude de grands système par modèles agrégés et méthode de perturbation* . Thèse de Docteur Ingénieur. Orsy, Décembre 1981.
- [24] E. Eitelberg. *Interactive Model reduction by minimizing the weighted equation error*. IFAC Symposium, Toulouse, 1980.
- [25] E. Eitelberg. *Comments on model reduction by minimizing the equation error*. IEEE Trans. Autom. Contr. Vol AC-27, N 4, p 1000-1002, Août 1982.
- [26] E. Eitelberg, J.C. Blada, and R.G. Harley. *A model reduction technique suitable for optimal controller design in high-order power systems*. Electric Power System Research 12, P 51-62, 1987.
- [27] F. Fer. *Thermodynamique macroscopique*. . Tome I, Gordon & Breach, Paris, 1970.
- [28] B. Flament. *Synthèse modale de systèmes thermiques*. Thèse de Doctorat, Ecole des Mines de Paris, Janvier 93.
- [29] B. Flament, F. Bourquin, and A. Neveu. *Synthèse modale: une méthode de sous-structuration dynamique pour la modélisation des systèmes thermiques linéaires*. Int. J. H. Mass. Transfer, A paraître.

- [30] A. J. Fossard and J. F. Magani. *Modélisation, commande et applications des systèmes à échelles de temps multiples*. RAIRO Automatic / Syst. Anal. And Cont., Vol 16 N1, P 5-23, 1982.
- [31] J. Fourier. *Théorie analytique de la chaleur*. Paris, 1822.
- [32] F.D. Galiana. *On the approximation of the multi-input, multi-output constant linear models*. Int. J. of Contr. Vol 17, P 1313-1324, 1973.
- [33] R. Giquel. *Approche systémique et thermique instationnaire*. Résumé des communications et conférences de synthèse, colloque de thermique SFT-EMP, Sophia-Antipolis, 20-21 Mai 1992.
- [34] GISE/CENERG. *Le projet SYMBOL: Synthèse Modale et Boîte à Outils Logiciels*. Rapport AFME/ARMINES-EMP, 1990.
- [35] M. Gondran and M. Minoux. *Graphes et algorithmes*. Eyrolles, Paris, 1979.
- [36] D. Guinin, F. Aubonnet, and B. Joppin. *Précis de mathématiques, cours-exercices résolus, Analyse 2*. Bréal, Nov. 1988.
- [37] GUT-CET. *COURS de Modélisation numérique en thermique*. Institut des Etudes Scientifiques de Cargèse. SFT-CNRS/PIRSEM-AFME, 29 Juin-4 Juillet 1992.
- [38] Imbert. *Analyse des structures par éléments finis*. Edition 2 Cépadués, 1984.
- [39] A. Jameson. *Solution of the equation $AX+XB=C$ by inversion of an (m,m) or (n,n) matrix*. SIAM J. Appl. Math, Vol 16 P198-201, 1968.
- [40] K. El Khoury. *Formulation modale de problèmes de diffusion thermique avec transport*. Thèse de Doctorat. Ecole des Mines de Paris, 1989.
- [41] J.P. Lebacque. *Systèmes dynamiques en dimension infinie. Première partie: Généralités*. Cahier du CERMA N12, ENPC P109-141, Juin 1991.
- [42] G. Lefebvre. *Analyse et réduction modales d'un modèle de comportement thermique de bâtiment*. Thèse de Doctorat. Université Paris VI, Nov. 1987.
- [43] G. Lefebvre, J. Bransier, and A. Neveu. *Simulation du comportement thermique d'un local par des méthodes numériques d'ordre réduit*. Revue générale de thermique N 302 P 106-114, Fev 1987.
- [44] J. Liferman. *Systèmes linéaires. Variables d'état*. Masson & Cie, 1972.
- [45] L. Litz. *Order reduction of linear state space models via optimal approximation of the nondominant modes*. North-Holland Publishing Company Large Scale system 2, P 171-184, 1981.
- [46] R. Marlinvaud. *L'agrégation dans les modèles économiques*. Cahiers des séminaires d'économétrie, N 4, 1956.

- [47] S. A. Marshall. *An approximate method for reducing the order of a linear system*. Control, P 642-643, 1966.
- [48] P. Merour. *La réduction des modèles en thermique. Application à l'étude d'un circuit intégré électronique*. Thèse de Doctorat, Poitiers, 1986.
- [49] G. P. Michaïlesco. *Approximation des systèmes complexes par des modèles de dimension réduite*. Thèse d'état. Orsay, Avril 1979.
- [50] F.P. Neirac. *Approche théorique et expérimentale des modèles réduits du comportement thermique des bâtiments*. Thèse de Doctorat, EMP, 1989.
- [51] A. Neveu. *Sur la perturbation des modèles modaux par une variation des paramètres physiques initiaux*. Note interne, CENERG-EMP, 1988.
- [52] G. Obinata and H. Inooka. *A method for modeling linear time invariant systems by linear systems of low order*. IEEE Transactions on Automatic Control, Août 1976.
- [53] A. Oulefki and A. Neveu. *Réduction et correction par amalgame modal d'un modèle thermique*. Communication, Colloque de Thermique, EMP Sophia-Antipolis, 20-21 Mai 1992.
- [54] A. Oulefki and A. Neveu. *Réduction par amalgame modal d'un modèle thermique*. Journal de Physique Volume 3, N 2, P 303-320, Février 1993.
- [55] A. Oulefki and A. Neveu. *Sur la détermination analytique des éléments propres d'une paroi plane composite*. Rapport interne E.N.P.C - G.I.S.E, Paris, sept. 1990.
- [56] D. Petit. *Réduction de modèles de connaissance et identification de modèles d'ordre réduit. Application au processus de diffusion thermique*. Thèse de Doctorat d'Etat ès Sciences. Université de Provence, IUSTI, 1991.
- [57] D. Petit and R. Pasquetti. *Réduction de modèles par identification de modes dominants: application à un modèle bidimensionnel de diffusion thermique*. Revue de Physique Appliquée N 25 P 831-842, Août 1990.
- [58] H. Reinard. *Equations différentielles. Fondements et applications. Chap I V*. gauthier-Villard. Paris, 1982.
- [59] J. J. Salgon. *Analyse modale par discrétisation spatiale. Application à la modélisation des ponts thermiques dans le bâtiment*. Thèse. Université Paris IV, 1987.
- [60] J.B. Saulnier. *La réduction de modèles en thermique*. AI 83 IASTED Symposium, Lille, Mars 1983.
- [61] J. Sicard, P. Bacot, and A. Neveu. *Analyse modale des échanges thermiques dans le bâtiment*. Int. J. Heat Transfer Vol.28 P 111-123, 1985.
- [62] J.M. Undrill and A.E Turner. *Construction of power systems electromechanical equivalents by modal analysis*. IEEE TPAS, Vol PAS 90, P 2049-2059, 1971.
- [63] D.A Wilson. *Optimum solution of model reduction problem*. Proc. IEEE, Vol 117, P 1161-1165, 1970.

Annexes

Annexe A

Etude énergétique de la troncature

A.1 Dominances des modes propres

La réduction de modèles d'état modaux par la troncature repose sur le classement des modes propres d'après leur importance selon un critère quantitatif donné. Parmi les critères de classement, ceux qui se réfèrent à l'énergie du système thermique sont les plus efficaces.

L'énergie associée à un système thermique sur un domaine $\mathcal{D}_t = [t_1, t_2]$ et lorsque seule la $j^{\text{ème}}$ composante du vecteur des sollicitations $U(t)$ n'est pas nulle est par définition

$$E_j = \int_{\mathcal{D}_t} \int_{\mathcal{D}} \left(T_d^j(M, t) \right)^2 C(M) \, dM \, dt \quad (\text{A.1})$$

Où $T_d^j(M, t)$ représente la composante dynamique de la température au point M et à l'instant t , qui s'écrit

$$T_d^j(M, t) = \sum_{i=1}^N x_{ij}(t) V_i(M) \quad (\text{A.2})$$

Sachant que les fonctions propres $V_i(M)$ vérifient une relation d'orthonormalité telle que

$$\forall i = 1, 2 \dots N \quad \forall j = 1, 2 \dots N$$

$$\langle V_i(M), C(M)V_j(M) \rangle = \int_{\mathcal{D}} V_i(M)C(M)V_j(M)dM = \delta_{ij} \quad (\text{A.3})$$

alors (A.1) devient

$$E_j = \sum_{i=1}^N \int_{\mathcal{D}_i} (x_{ij}(t))^2 dt \quad (\text{A.4})$$

Où $x_{ij}(t)$ est l'état $x_i(t)$ du mode "i" lorsque seule la composante $U_j(t)$ n'est pas nulle. Cette dernière relation se présente comme une somme de N quantités positives. Aussi, la contribution de chaque mode à l'énergie E_j est donnée par

$$\forall i \in 1..N \quad \epsilon_i^j = \int_{\mathcal{D}_i} (x_{ij}(t))^2 dt \quad (\text{A.5})$$

Si l'on tient compte maintenant de l'ensemble des p sollicitations de façon indépendante, alors on définit la quantité e_i par

$$\forall i \in 1..N \quad e_i = \sum_{j=1}^p \epsilon_i^j \quad (\text{A.6})$$

On définit finalement la *dominance* (ou l'énergie) d_i de chaque mode propre en normalisant la relation (A.5) de la façon suivante

$$\forall i = 1 \dots N \quad d_i = \frac{1}{p} \sum_{j=1}^p E_{ij} \quad \text{avec} \quad E_{ij} = \frac{\epsilon_i^j}{E_j} \quad (\text{A.7})$$

et l'on a naturellement

$$0 \leq d_i \leq 1 \quad \text{et} \quad \sum_{i=1}^N d_i = 1$$

ce qui signifie que *l'énergie du système thermique est répartie sur l'ensemble des N modes propres*.

Le calcul effectif des quantités d_i ne nécessite en réalité que celui des ϵ_i^j qui dépendent des états x_{ij} des modes propres. Le calcul de ces états peut se faire de deux manières, ce qui donne les deux critères ci-après.

A.1.1 Critère de dominance

Considérons les deux types de sollicitations tests

- Echelon unitaire de Heaviside
- Impulsion unitaire de Dirac

pour lesquelles les quantités x_{ij} s'écrivent

$$x_{ij} = B_{ij} \lambda_i^s e^{\lambda_i t}$$

Avec $s = 0$ pour l'échelon et $s = 1$ pour le Dirac. On obtient d'après (A.5)

$$\forall i \in 1 \cdots N \quad \forall j = 1 \cdots p \quad e_i^j = -\frac{B_{ij}^2 \lambda_i^{2s-1}}{2} \left[e^{2\lambda_i t_1} - e^{2\lambda_i t_2} \right] \quad (\text{A.8})$$

A.1.2 Critère de dominance pondérée

Dans le tracé des spectres de réponse indicielle (voir figure 2.2 du §2.2), on utilise une échelle t^* de type logarithmique telle que

$$t^* = \frac{\text{Ln} \left(1 + 100 \frac{t}{t_{ref}} \right)}{\text{Ln}(101)} \quad (\text{A.9})$$

qui a pour but de dilater l'axe temporel pour les temps d'évolution rapide. En effet, les modèles réduits par troncature présentent généralement une insuffisance au début de l'évolution thermique du système. Le paramètre t_{ref} est arbitraire et permet d'avoir une échelle aussi dilatable que nécessaire pour mieux analyser le début de l'évolution. A l'instant $t = t_{ref}$, on a $t^* = 1$. Dans la pratique $t_{ref} = \tau_1$ est généralement une bonne valeur.

Si l'on remplace dans (A.1) la variable t par la nouvelle variable $t^* = g(t)$ et $\mathcal{D}_t$ par $\mathcal{D}_{t^*} = [g(t_1), g(t_2)]$, alors on a

$$E_j = \int_{\mathcal{D}_{t^*}} \int_{\mathcal{D}} \left(T_d^j(M, t) \right)^2 C(M) \frac{\partial g(t)}{\partial t} dM dt \quad (\text{A.10})$$

Cette opération simple s'apparente donc à une introduction d'une fonction *poids* $\frac{\partial g(t)}{\partial t}$ qui tient mieux compte de la dynamique rapide que (A.1). La fonction poids $\partial g(t)/\partial t$ est immédiate à partir de (A.9)

$$\frac{\partial g(t)}{\partial t} = \frac{1}{\text{Ln}(101)} \frac{100}{t_{ref}} \frac{1}{\text{Ln}\left(1 + 100 \frac{t}{t_{ref}}\right)} \quad (\text{A.11})$$

Comme en §A.1.1, on considère les états

$$x_{ij} = B_{ij} \lambda_i^s e^{\lambda_i t} \quad s = 0 : \text{échelon} \quad \text{et} \quad s = 1 : \text{Dirac}$$

et l'on obtient d'après (A.5)

$$e_i^j = \frac{1}{\text{Ln}(101)} B_{ij}^2 \lambda_i^{2s} e^{\frac{2}{\beta \tau_i}} \int_{y_1}^{y_2} \frac{e^{c_i y}}{y} dy \quad (\text{A.12})$$

avec $\beta = \frac{100}{t_{ref}}$, $c_i = -\frac{2}{\beta \tau_i}$ et $(y_k = 1 + \beta t_k \quad k = 1, 2)$

Remarque: L'intégrale intervenant dans (A.12) a comme primitive

$$\int \frac{e^{c_i y}}{y} dy = \text{Ln}(y) + \sum_{k=1}^{\infty} \frac{(c_i y)^k}{k \cdot k!} + Cte$$

A.1.3 Erreur de troncature

La troncature d'un modèle d'état modal au moyen d'un critère énergétique revient à garder les "n" éléments propres les plus dominants selon (A.7). Les modes dominants engendrent le *sous-espace dominant* S_n . Notons par $\{I_D\}$ l'ensemble des indices des modes dominants engendrant S_n . On définit alors la dominance D_{S_n} du sous-espace dominant S_n par

$$D_{S_n} = \sum_{m \in \{I_D\}} d_m \quad \text{avec} \quad 0 \leq d_m \leq 1 \quad (\text{A.13})$$

On peut considérer que le sous-espace S_n est d'autant plus dominant que sa dominance D_{S_n} est proche de 1. Pour avoir un sous-espace S_n suffisamment dominant, on peut être amené à sélectionner peu ou beaucoup de modes, selon le système thermique étudié et ses conditions aux limites. On dit alors que la dominance est *concentrée* ou *diffuse* selon que la valeur "n" nécessaire pour que D_{S_n} soit proche de 1 est respectivement faible ou grande.

Remarque: $D_{S_N} = 1$ et $D_{S_0} = 0$

L'erreur de troncature $\|e\|^2$ peut être alors définie comme étant l'écart entre la dominance de S_N (qui est égale à 1 puisque le modèle détaillé est supposé contenir *toute* l'énergie du système thermique) et celle de S_n . On écrit donc

$$\|e\|^2 = 1 - D_{S_n} \quad (\text{A.14})$$

qu'on peut écrire également

$$\|e\|^2 = \sum_{m \notin \{I_D\}} d_m \quad (\text{A.15})$$

A.1.4 Troncature automatique

Il est possible, au moins sur le plan théorique, d'automatiser la troncature des modèles d'état modaux par des critères énergétiques. Pour cela, on peut imposer au sous-espace dominant S_n de réaliser la contrainte

$$D_{S_n} \gg D_{min} \quad (\text{A.16})$$

Où D_{min} est la dominance minimale que doit contenir le sous-espace S_n . La procédure d'automatisation revient donc à chercher le modèle réduit ayant la plus petite dimension "n" et vérifiant la contrainte d'erreur ci-dessus. Dans la pratique, cette manière d'automatiser la troncature se heurte à la difficulté de fixer la valeur de D_{min} .

Annexe B

Equation de Lyapunov

Soit un domaine temporel $\mathcal{D}_t = [t_1, t_2]$ et trois matrices $A[n, n]$, $B[n, m]$ et $C[m, m]$. Calculons l'intégrale de Lyapunov définie par

$$L = \int_{\mathcal{D}_t} \begin{bmatrix} e^{At} & B & e^{Ct} \end{bmatrix} dt \quad [n, m] \quad (\text{B.1})$$

Si l'on pose

$$l(t) = \begin{bmatrix} e^{At} & B & e^{Ct} \end{bmatrix} \quad (\text{B.2})$$

on vérifie facilement que

$$l'(t) = \begin{bmatrix} A & e^{At} & B & e^{Ct} \end{bmatrix} + \begin{bmatrix} e^{At} & B & C & e^{Ct} \end{bmatrix} \quad (\text{B.3})$$

On peut écrire

- $\int_{\mathcal{D}_t} l'(t) dt = l(t_2) - l(t_1)$
Sachant ¹ que $e^{Ct}C = Ce^{Ct}$ on a aussi
- $\int_{\mathcal{D}_t} l'(t) dt = A L + L C$ d'après (B.2) et (B.3)

¹Car le développement en série de e^{Ct} est $e^{Ct} = \sum_{i=1}^{\infty} \frac{(Ct)^i}{i!}$. Comme t est une variable scalaire, la multiplication de e^{Ct} par C peut se faire à droite ou à gauche.

En notant $\mathbb{B} = l(t_2) - l(t_1)$ on déduit facilement que

$$AL + LC = -\mathbb{B} \quad (\text{B.4})$$

L'équation (B.4) est dite alors: équation de Lyapunov associée à l'intégrale de Lyapunov (B.1) Le calcul de L se fait souvent par la résolution du système matriciel (B.4) qui dépend de la forme des matrices A et C . Nous donnons ici la solution explicite de (B.4) dans le cas où A et C sont deux matrices diagonales.

$$L_{ij} = -\frac{\mathbb{B}_{ij}}{A_{ii} + C_{ii}} \quad (\text{B.5})$$

Remarque Si A et C ne sont pas diagonales mais diagonalisables, il est facile de résoudre l'équation de Lyapunov (B.4) de façon explicite. On peut se rapporter à la référence [55] à ce sujet. Si maintenant ces deux matrices ne sont pas diagonalisables, il existe des algorithmes spécifiques [9] [38] pour l'obtention de la solution de B.4.

La théorie de Lyapunov [57] repose principalement sur l'étude et l'approximation des attracteurs de systèmes dynamiques. On montre [40] que l'intégrale de Lyapunov donne des indications potentielles sur l'attracteur d'un système dynamique. Les mathématiciens considèrent que l'apparition des intégrales de Lyapunov dans un problème d'optimisation est plutôt un bon signe qui découle d'une bonne formulation du problème!

Annexe C

Formulation modale pour les régimes forcés

La formulation modale sur laquelle repose les travaux de cette thèse fait intervenir la dérivée première $\dot{U}(t)$ du vecteur des sollicitations par rapport au temps. On rappelle son écriture

$$\begin{cases} \dot{x}(t) = Fx(t) + B\dot{U}(t) \\ y(t) = Hx(t) + SU(t) \end{cases} \quad (C.1)$$

A partir de la formulation modale (C.1), il est possible [55] d'obtenir des représentations faisant intervenir les dérivées $U^{(i)}$ $i \geq 2$. L'équation d'état de (C.1) donne le vecteur d'états asymptotique pour $U^{(1)} = \text{Cte}$ qui s'écrit¹

$$x_{(1)}^{asym} = -F^{-1}BU^{(1)}$$

Pour obtenir la formulation² en $U^{(2)}$, posons

$$x_{(1)}(t) = x_{(2)}(t) - F^{-1}BU^{(1)} \quad (C.2)$$

On obtient alors à partir de (C.1)

$$\begin{cases} \dot{x}_{(2)}(t) = Fx_{(2)}(t) + B_{(2)}U^{(2)}(t) \\ y(t) = Hx_{(2)}(t) + g(U, U^{(1)}) \end{cases} \quad (C.3)$$

$$\text{avec } B_{(2)} = F^{-1}B \text{ et } g(U(t), U^{(1)}(t)) = -HF^{-1}BU^{(1)}(t) + SU(t)$$

¹On désigne par $U^{(i)}$ la $i^{\text{ème}}$ dérivée de $U(t)$ par rapport au temps t .

²On suppose que les composantes du vecteur d'entrée $U(t)$ vérifient les conditions de dérivabilité nécessaires à la suite des calculs.

En répétant cette démarche plusieurs fois on aboutit à la relation générale

$$x_{(j)}(t) = x_{(j+1)}(t) - F^{-j} B U^{(j)}$$

et le modèle d'état modal faisant intervenir $U^{(j+1)}$ s'écrit

$$\begin{cases} \dot{x}_{(j+1)}(t) = F x_{(j+1)}(t) + B_{(j+1)} U^{(j+1)}(t) \\ y(t) = H x_{(j+1)}(t) + g(U, U^{(1)}, \dots, U^{(j)}) \end{cases} \quad (C.4)$$

avec

$$B_{(j+1)} = F^{-j} B \quad (C.5)$$

$$\text{et } g(U, U^{(1)}, \dots, U^{(j)}) = -\sum_{k=1}^j H F^{-k} B U^{(k)}(t) + S U(t)$$

La résolution de (C.4) fait intervenir les états $x_{(k)}(t)$. Aussi, nous allons voir comment ils s'expriment en fonction des $U^{(k)}(t)$.

Le système (C.1) a pour solution

$$x_{(1)}(t) = e^{Ft} x_{(1)}(0) + \int_0^t e^{F(t-\xi)} B_{(1)} U^{(1)}(\xi) d\xi \quad (C.6)$$

où l'intégrale temporelle peut se faire par partie, ce qui donne

$$x_{(1)}(t) = e^{Ft} x_{(1)}(0) + \left[-e^{F(t-\xi)} F^{-1} B U^{(1)}(\xi) \right]_0^t + \int_0^t e^{F(t-\xi)} B_{(2)} U^{(2)}(\xi) d\xi \quad (C.7)$$

où $B_{(2)}$ est donné par (C.5). En introduisant le changement de variable (C.2), cette dernière relation s'écrit

$$x_{(2)}(t) = e^{Ft} x_{(2)}(0) + \int_0^t e^{F(t-\xi)} B_{(2)} U^{(2)}(\xi) d\xi \quad (C.8)$$

avec

$$x_{(2)}(0) = x_{(1)}(0) + F^{-1} B U^{(1)}(0) \quad (C.9)$$

En reprenant la démarche ci-dessus plusieurs fois, on obtient le résultat général

$$x_{(i)}(t) = e^{Ft} x_{(i)}(0) + \int_0^t e^{F(t-\xi)} B_{(i)} U^{(i)}(\xi) d\xi \quad (C.10)$$

Le problème de réduction peut alors être généralisé de la façon suivante

modèle détaillé (dimension N)

$$\begin{cases} \dot{x}_{(i)}(t) = F x_{(i)}(t) + B_{(i)} U^{(i)}(t) \\ y(t) = H x_{(i)}(t) + g(U, U^{(1)} \dots U^{(i-1)}) \end{cases} \quad (C.11)$$

modèle réduit (dimension $n \ll N$)

$$\left\{ \begin{array}{l} \dot{\tilde{x}}_{(i)}(t) = \tilde{F}\tilde{x}_{(i)}(t) + \tilde{B}_{(i)}U^{(i)}(t) \\ \tilde{y}(t) = \tilde{H}\tilde{x}_{(i)}(t) + g(U, U^{(1)} \dots U^{(i-1)}) \end{array} \right. \quad (\text{C.12})$$

Le problème de réduction ainsi posé s'applique seulement à la partie dynamique de la réponse, ce qui permet de sauvegarder les régimes permanent ($U = Cte$), rampe ($U^{(1)} = Cte$) et régimes asymptotiques complexes correspondant à ($U^{(i)} = Cte \quad i > 1$). On trouvera dans [55] la reformulation de plusieurs méthodes de réduction utilisant les dérivées successives du vecteur d'entrée.

Annexe D

Rappels sur la dérivation tensorielle

Nous donnons ici quelques règles pratiques de dérivation, qui sont utiles pour le lecteur désireux de rétablir les résultats des différentes méthodes de réduction insérées dans cette thèse.

Notations

$\{\}$	Vecteur
$()$	Matrice
$\cdot$	Produit scalaire
T	Transposition
tr	Trace

Dérivation tensorielle

Soient A et B deux tenseurs d'ordres n et m respectivement. On a

$$C = \frac{\partial A}{\partial B} = \frac{\partial A_{i_1, i_2 \dots i_n}}{\partial B_{j_1, j_2 \dots j_m}}$$

C est un tenseur d'ordre $(n + m)$.

Règles utiles

Soient: α scalaire; V, W et U trois vecteurs; A et B deux matrices. Les dimensions de V, W, U, A et B sont appropriées pour chacune des situations suivantes.

$$\blacktriangleright \left\{ \frac{\partial \alpha}{\partial V} \right\} = \frac{\partial \alpha}{\partial V_i}$$

$$\blacktriangleright \left(\frac{\partial W}{\partial V} \right) = \frac{\partial W_i}{\partial V_j}$$

$$\begin{aligned} \blacktriangleright \left(\frac{\partial [\alpha W]}{\partial V} \right) &= \frac{\partial \alpha}{\partial V_j} W_i + \alpha \frac{\partial W_i}{\partial V_j} \\ &= W \left\{ \frac{\partial \alpha}{\partial V} \right\}^T + \alpha \left(\frac{\partial W}{\partial V} \right) \end{aligned}$$

$$\begin{aligned} \blacktriangleright \left\{ \frac{\partial [V W]}{\partial U} \right\} &= \frac{\partial V_i W_i}{\partial U_j} = \frac{\partial V_i}{\partial U_j} W_i + V_i \frac{\partial W_i}{\partial U_j} \\ &= \left(\frac{\partial V}{\partial U} \right) \cdot W + \left(\frac{\partial W}{\partial U} \right)^T \cdot V \end{aligned}$$

$$\blacktriangleright \frac{\partial [tr[V W^T]]}{\partial U} = \frac{\partial [V W]}{\partial U}$$

Soient maintenant X et A deux matrices de dimensions appropriées pour *chacune* des situations suivantes

$$\blacktriangleright tr[AX] = tr[XA]$$

$$\blacktriangleright \left(\frac{\partial}{\partial X} tr[AX] \right) = \frac{\partial}{\partial X_{ij}} A_{iq} X_{ql} = A_{ji}$$

$$\text{donc} \quad \left(\frac{\partial}{\partial X} tr[AX] \right) = A^T$$

$$\blacktriangleright \left(\frac{\partial}{\partial X} tr[X^T A] \right) = \frac{\partial}{\partial X_{ij}} X_{ql} A_{ql} = A_{ij}$$

$$\text{donc} \quad \left(\frac{\partial}{\partial X} tr[X^T A] \right) = A$$

$$\blacktriangleright \left(\frac{\partial}{\partial X} tr[X^T AX] \right) = \frac{\partial}{\partial X_{ij}} [X_{ql} X_{ml} A_{qm}] = X_{mj} A_{im} + X_{qj} A_{qi}$$

$$\text{donc} \quad \left(\frac{\partial}{\partial X} tr[X^T AX] \right) = (A + A^T)X$$

$$\blacktriangleright \left(\frac{\partial}{\partial X} tr[XAX^T] \right) = \frac{\partial}{\partial X_{ij}} X_{ql} A_{lm} X_{qm} = A_{jm} X_{im} + X_{il} A_{lj}$$

$$\text{donc} \quad \left(\frac{\partial}{\partial X} tr[XAX^T] \right) = X(A^T + A)$$

Soient quatre matrices: $A[m, p]$ $B[n, n]$ $C[p, n]$ et $X[n, p]$

$$\begin{aligned}
\blacktriangleright \left(\frac{\partial}{\partial X} \text{tr}[AX^T BXC] \right) &= \frac{\partial}{\partial X_{ij}} A_{ql} X_{ml} B_{mk} X_{kn} C_{nq} \\
&= A_{qj} B_{ik} X_{kn} C_{nq} + A_{ql} X_{ml} B_{mi} C_{jq}
\end{aligned}$$

donec

$$\left(\frac{\partial}{\partial X} \text{tr}[AX^T BXC] \right) = BXC A + B^T X (CA)^T$$

Annexe E

Mesure de l'erreur

Nous avons montré au §3.5 pourquoi la mesure $\mathcal{M}$ ne peut pas être retenue pour l'automatisation de la réduction. Ensuite, nous avons proposé une mesure $\overline{\mathcal{M}}$ permettant d'imposer une contrainte d'erreur à partir de laquelle on cherche un modèle réduit de dimension appropriée (respectant cette contrainte). L'objectif est ici de donner la signification de $\overline{\mathcal{M}}$.

Pour alléger les calculs, considérons d'abord le cas où le système thermique est soumis à l'action d'une seule sollicitation. On a alors

$$\overline{\mathcal{M}}^2 = \frac{\int_{\mathcal{D}_t} \int_{\mathcal{D}} \mathcal{E}^2(M, t) C(M) dM dt}{\Delta t_0 \int_{\mathcal{D}} C(M) dM} \quad (\text{E.1})$$

où $\mathcal{E}^2(M, t)$ est ici une fonction scalaire que nous désignerons par “**carré de l'écart**”. Comme nous l'avons expliqué au §3.5, Δt_0 n'est pas nécessairement égal à $(t_2 - t_1)$ (rappel: $\mathcal{D}_t = [t_1, t_2]$). Si $t_2 \geq t_{perm}$ (t_{perm} étant le temps d'établissement du régime permanent qui est voisin de $4\tau_1$) alors $\Delta t_0 = t_{perm} - t_1$ sinon $\Delta t_0 = t_2 - t_1$. En effet, on a

$$\int_{t_1 \leq t_{perm}}^{t_2 \geq t_{perm}} \int_{\mathcal{D}} \mathcal{E}^2(M, t) C(M) dM dt \neq \int_{t_1}^{t_{perm}} \int_{\mathcal{D}} \mathcal{E}^2(M, t) C(M) dM dt$$

L'expression (E.1) peut aussi s'écrire

$$\overline{\mathcal{M}}^2 = \frac{1}{\Delta t_0} \int_{\mathcal{D}_t} \left(\frac{\int_{\mathcal{D}} \mathcal{E}^2(M, t) C(M) dM}{\int_{\mathcal{D}} C(M) dM} \right) dt \quad (\text{E.2})$$

Supposons que le système thermique se compose de plusieurs milieux homogènes. Dans ce cas, la dernière expression s'écrit aussi

$$\overline{\mathcal{M}}^2 = \frac{1}{\Delta t_0} \int_{\mathcal{D}_t} \left(\frac{\sum_{i=1}^{nch} C_i \int_{\mathcal{D}_i} \mathcal{E}^2(M, t) dM}{\sum_{i=1}^{nch} C_i \int_{\mathcal{D}_i} dM} \right) dt \quad (\text{E.3})$$

où

- * nch est le nombre de composants homogènes du système thermique.
- * C_i est la capacité calorifique ($C_i = \rho_i C_{p_i}$) du $i^{\text{ème}}$ composant homogène.
- * $\mathcal{D}_i$ est le sous-domaine spatial occupé par le $i^{\text{ème}}$ composant homogène.

Si $nch = 1$ alors le terme entre parenthèse (.) de la dernière expression est par définition "la moyenne sur $\mathcal{D}$ du carré de l'écart à l'instant t " car dans ce cas la mesure C_i se simplifie. Si maintenant $nch > 1$ alors le terme (.) est par définition "la moyenne sur $\mathcal{D}$, à l'instant t , du carré de l'écart vis-à-vis de la mesure $C(M)$ ". L'intégration de (.) sur $\mathcal{D}_t$ associée à la division par Δt_0 donne la moyenne de (.) sur Δt_0 .

Lorsqu'il y a plusieurs sollicitations (au nombre de p), $\mathcal{E}(M, t)$ devient un vecteur tel que

$$\mathcal{E}(M, t) = [e^1(M, t), e^2(M, t) \cdots, e^p(M, t)] \quad [1, p]$$

En négligeant les interactions entre les différentes sollicitations, on pose

$$\overline{\mathcal{M}}^2 = \frac{1}{p} \sum_{j=1}^p \left(\frac{1}{\Delta t_0} \int_{\mathcal{D}_t} \left(\frac{\sum_{i=1}^{nch} C_i \int_{\mathcal{D}_i} (e^j(M, t))^2 dM}{\sum_{i=1}^{nch} C_i \int_{\mathcal{D}_i} dM} \right) dt \right) \quad (\text{E.4})$$

Aussi, $\overline{\mathcal{M}}^2$ est maintenant "la moyenne sur Δt_0 et par sollicitation de (.)". $\overline{\mathcal{M}}^2$ a la dimension de (Kelvin)². En ce qui concerne la méthode d'amalgame, dans les différentes étapes qui mènent à l'obtention du modèle réduit, les sollicitations utilisées sont des échelons **unitaires**. Numériquement, nous avons constaté que les valeurs de $\overline{\mathcal{M}}$ (en K) sont comprises entre 0 et 1. Des exemples de valeurs de $\overline{\mathcal{M}}$ sont données dans la partie réservée aux exemples (voir l'exemple du bâtiment bizonne ou du pont thermique). La relation (3.37) donnée au §3.5.1 est identique à (E.4).

En ce qui concerne la mesure de l'erreur $\mathcal{M}$, elle est donnée par

$$\mathcal{M} = \sum_{j=1}^p \int_{\mathcal{D}_t} \int_{\mathcal{D}} (e^j(M, t))^2 C(M) dM dt$$

et sa signification est claire: elle totalise le carré de l'écart pondéré par $C(M)$ sur $\mathcal{D}$, sur $\mathcal{D}_t$ et pour l'ensemble des sollicitations.

Remarque:

Dans la méthode d'amalgame, les diverses étapes menant au modèle réduit ont pour objectif d'approcher au mieux le minimum de $\mathcal{M}$. Une fois que l'espace d'état modal est partitionné, on évalue $\overline{\mathcal{M}}$ pour avoir une idée sur la pertinence

du modèle réduit. En fait, ceci est possible car $\overline{\mathcal{M}}$ possède exactement les mêmes extrémums que $\mathcal{M}$. En effet, $\overline{\mathcal{M}}$ est directement proportionnelle à la racine carrée de $\mathcal{M}$ (ie: $\overline{\mathcal{M}} = cte \cdot \sqrt{\mathcal{M}}$). Comme la racine carrée est une fonction strictement croissante dans $\mathbb{R}_+$, alors la recherche du minimum de $\mathcal{M}$ est rigoureusement équivalente à celle du minimum de $\overline{\mathcal{M}}$. Cependant la manipulation de $\mathcal{M}$ est plus aisée que celle de $\overline{\mathcal{M}}$.

Annexe F

généralités sur la méthode d'amalgame

Dans la méthode d'amalgame présentée au §3, nous avons essayé d'alléger la présentation en évitant de donner les preuves théoriques de quelques résultats. L'objet de cette annexe est de donner l'essentiel de la démarche qui permet d'obtenir ces résultats.

F.1 Détermination des modes amalgamés

L'espace des modes propres est partitionné en n sous-espaces H_k $k = 1 \cdots n$ (la dimension de H_k est n_k) qui sont orthogonaux entre eux. Chaque sous-espace est engendré par un mode principal et quelques modes mineurs. A partir de chaque sous-espace, on calcule un seul mode amalgamé par combinaison linéaire des modes qui le composent. On traduit ceci par la relation

$$\tilde{V}_k(M) = P_k \omega_k \quad \forall k = 1 \cdots n \quad (\text{F.1})$$

Dans chaque sous-espace H_k , le mode recherché $\tilde{V}_k(M)$ est connu si l'on détermine le vecteur ω_k (ω_k a n_k composantes) des coefficients de la combinaison (F.1).

Reprenons la mesure de l'erreur sous sa forme

$$\mathcal{M} = \sum_{k=1}^n \mathcal{M}_k \quad (\text{F.2})$$

Avec

$$\left| \begin{array}{l} \mathcal{M}_k = \int_{\mathcal{D}_t} \text{tr} \left[\alpha_k^T(t) \alpha_k(t) \right] dt \\ \alpha_k(t) = \omega_k \tilde{X}_k(t) - \pi_k^T X(t) \quad [n_k, p] \end{array} \right. \quad (\text{F.3})$$

Comme $\text{tr} \left[\alpha_k^T(t) \alpha_k(t) \right] \geq 0$, on a nécessairement

$$\min_{\omega_1 \cdots \omega_n} \{ \mathcal{M} \} = \sum_{k=1}^n \min_{\omega_k} \{ \mathcal{M}_k \} \quad (\text{F.4})$$

Calcul de ω_k : Dans le sous-espace H_k , ω_k est tel que

$$\frac{\partial}{\partial \omega_k} \mathcal{M}_k(\omega_k) = 0$$

Or

$$\text{tr} [\alpha_k^T(t) \alpha_k(t)] = \text{tr} [\tilde{X}_k^T \omega_k^T \omega_k \tilde{X}_k] - \text{tr} [\tilde{X}_k^T \omega_k^T \pi_k^T X] - \text{tr} [X^T \pi_k \omega_k \tilde{X}_k] + \text{tr} [X^T \pi_k \pi_k^T X]$$

Pour alléger les calculs, on désigne respectivement par t_1 t_2 t_3 et t_4 les quatre termes du membre de droite de la dernière relation et l'on a

$$\mathcal{M}_k = \sum_{i=1}^4 \int_{\mathcal{D}_t} t_i(t) dt \quad (\text{F.5})$$

$$\text{et } \frac{\partial}{\partial \omega_k} \left[\sum_{i=1}^4 \int_{\mathcal{D}_t} t_i(t) dt \right] = 0 \quad (\text{F.6})$$

On peut vérifier au moyen de la technique de dérivation tensorielle exposée en annexe D que

$$\begin{aligned} \bullet \quad \frac{\partial}{\partial \omega_k} t_1(t) &= 2\omega_k \left(\tilde{X}_k \tilde{X}_k^T \right) \\ \bullet \quad \frac{\partial}{\partial \omega_k} t_2(t) &= -\pi_k^T X \tilde{X}_k^T \\ \bullet \quad \frac{\partial}{\partial \omega_k} t_3(t) &= \frac{\partial}{\partial \omega_k} t_2(t) \\ \bullet \quad \frac{\partial}{\partial \omega_k} t_4(t) &= 0 \end{aligned} \quad (\text{F.7})$$

d'où finalement d'après (F.6)

$$2\omega_k \int_{\mathcal{D}_t} \left(\tilde{X}_k \tilde{X}_k^T \right) dt - 2\pi_k^T \int_{\mathcal{D}_t} \left(X \tilde{X}_k^T \right) dt = 0 \quad (\text{F.8})$$

ou encore

$$\omega_k = \pi_k^T \frac{\int_{\mathcal{D}_t} \mathbf{X}(t) \tilde{X}_k^T(t) dt}{\int_{\mathcal{D}_t} \tilde{X}_k(t) \tilde{X}_k^T(t) dt} \quad \text{dimension } [n_k] \quad (\text{F.9})$$

Les modes amalgamés sont entièrement déterminés par l'évaluation de (F.9) pour $k = 1 \dots n$.

Il est par ailleurs possible de prouver que ω_k donné par (F.9) correspond au *minimum absolu* de $\mathcal{M}_k$ et non à un point col. Pour cela, il est suffisant¹ de vérifier que

$$\forall i, j = 1 \dots n_k \quad \frac{\partial^2}{\partial \omega_k^2(i)} \mathcal{M}_k > 0 \quad \text{et} \quad \frac{\partial^2}{\partial \omega_k(j) \partial \omega_k(i)} \mathcal{M}_k = 0$$

¹**Théorème:** ([35] P277 théorème 19)

Soit $f: E \mapsto \mathbb{R}$ de classe C^2 et $a \in E$ un point critique de f tel que

A partir de (F.7), on vérifie bien que

$$\frac{\partial}{\partial \omega_k(i)} \mathcal{M}_k = 2 \left(\tilde{X}_k \tilde{X}_k^T \right) \omega_k(i) \quad (\text{scalaire}) \quad (\text{F.10})$$

d'où

$$\left| \begin{array}{l} \frac{\partial^2}{\partial \omega_k^2(i)} \mathcal{M}_k = 2 \left(\tilde{X}_k \tilde{X}_k^T \right) > 0 \\ \text{et } \frac{\partial^2}{\partial \omega_k(j) \partial \omega_k(i)} \mathcal{M}_k = 0 \end{array} \right. \quad (\text{F.11})$$

F.2 Contribution à $\mathcal{M}$ des modes propres

L'expression (F.2) montre que la mesure de l'erreur est une somme de n termes positifs, chacun étant relatif à l'un des sous-espaces H_k $k = 1 \dots n$. Montrons que $\mathcal{M}_k$ se décompose à son tour en termes de contribution des n_k modes propres de H_k . Reprenons (F.5)

$$\mathcal{M}_k = \sum_{i=1}^4 \int_{\mathcal{D}_t} t_i(t) dt$$

$$\begin{aligned} \bullet \int_{\mathcal{D}_t} t_1(t) dt &= \int_{\mathcal{D}_t} \text{tr} \left[\tilde{X}_k^T \left(\omega_k^T \omega_k \right) \tilde{X}_k \right] dt \\ &= \left(\omega_k^T \omega_k \right) \int_{\mathcal{D}_t} \left(\tilde{X}_k \tilde{X}_k^T \right) dt \\ \bullet \int_{\mathcal{D}_t} t_2(t) dt &= - \int_{\mathcal{D}_t} \text{tr} \left[\tilde{X}_k^T \omega_k^T \pi_k^T X \right] dt \\ &= - \text{tr} \left[\pi_k^T \int_{\mathcal{D}_t} \left(X \tilde{X}_k^T \right) dt \omega_k^T \right] \end{aligned}$$

en tenant compte de (F.9) on obtient

la forme quadratique q_a soit non dégénérée

$$q_a : u \longmapsto \sum_{i=1}^n u_i^2 \frac{\partial^2 f}{\partial x_i^2}(a) + 2 \sum_{i,j=1}^n u_i u_j \frac{\partial^2 f}{\partial x_i \partial x_j}(a)$$

Si $\text{sgn } q_a = (n, 0)$ (ie: q_a définie positive), f présente un minimum en a
 Si $\text{sgn } q_a = (0, n)$ (ie: q_a définie négative), f présente un maximum en a
 Sinon $\text{sgn } q_a = (s, n-s)$, ($1 \leq s \leq n-1$), f ne présente en a ni maximum, ni minimum: a est un point col.

$$\int_{\mathcal{D}_t} t_2(t) dt = - \left(\omega_k^T \omega_k \right) \int_{\mathcal{D}_t} \left(\tilde{X}_k \tilde{X}_k^T \right) dt$$

donc

$$\int_{\mathcal{D}_t} t_2(t) dt = - \int_{\mathcal{D}_t} t_1(t) dt$$

$$\begin{aligned} \bullet \int_{\mathcal{D}_t} t_3(t) dt &= - \int_{\mathcal{D}_t} \text{tr} \left[X^T \pi_k \omega_k \tilde{X}_k \right] dt \\ &= - \text{tr} \left[\pi_k^T \int_{\mathcal{D}_t} \left(X \tilde{X}_k^T \right) dt \omega_k^T \right] \end{aligned}$$

en tenant compte de (F.9) on obtient aussi

$$\begin{aligned} \int_{\mathcal{D}_t} t_3(t) dt &= - \int_{\mathcal{D}_t} t_1(t) dt \\ \bullet \int_{\mathcal{D}_t} t_4(t) dt &= \int_{\mathcal{D}_t} \text{tr} \left[X^T \pi_k \pi_k^T X \right] dt \end{aligned}$$

On a finalement

$$\begin{aligned} \mathcal{M}_k &= \int_{\mathcal{D}_t} t_4(t) dt - \int_{\mathcal{D}_t} t_1(t) dt \\ &= \int_{\mathcal{D}_t} \text{tr} \left[X^T \pi_k \pi_k^T X \right] dt - \left(\omega_k^T \omega_k \right) \int_{\mathcal{D}_t} \left(\tilde{X}_k \tilde{X}_k^T \right) dt \end{aligned} \quad (\text{F.12})$$

qui s'écrit aussi

$$\begin{aligned} \mathcal{M}_k &= \sum_{m=1}^{n_k} \mathcal{M}_{km} \\ \text{avec } \mathcal{M}_{km} &= \left[\sum_{j=1}^p \int_{\mathcal{D}_t} x_{u^k(m)_j}^2 dt \right] - \left[\omega_k^2(m) \sum_{j=1}^p \int_{\mathcal{D}_t} (\tilde{x}_{kj}(t))^2 dt \right] \end{aligned} \quad (\text{F.13})$$

Cette dernière relation est importante dans la méthode de réduction par amalgame modal car elle prouve qu'un mode mineur peut apporter une information thermique à un mode principal. Le partitionnement de l'espace des modes doit alors se faire de façon à ce que cet apport d'information soit maximal.

F.3 Expression de $\mathcal{M}$ pour les différents modèles

Nous profitons des résultats que nous venons d'établir pour résumer l'expression de la mesure de l'erreur pour les modèles tronqués et amalgamés.

Modèles tronqués

Désignons par $\{I_D\}$ l'ensemble des indices des modes jugés dominants d'après un critère quelconque de troncature. La mesure de l'erreur s'écrit alors

$$\mathcal{M} = \sum_{i \notin \{I_D\}} \sum_{j=1}^p \int_{\mathcal{D}_t} x_{ij}^2 dt \quad (\text{F.14})$$

qui se réduit dans le cas des sollicitations tests (échelon: $s=0$ et Dirac: $s=1$)

$$\mathcal{M} = \sum_{i \notin \{I_D\}} \left[-\frac{\lambda_i^{2s-1}}{2} \left(\sum_{l=1}^p B_{il}^2 \right) [\Delta e^{2\lambda_i t}]_{t_2}^{t_1} \right] \quad (\text{F.15})$$

Remarque: Le terme $[\cdot]$ de la dernière relation est l'énergie du mode "i" (voire §2.2.4 ou l'annexe A). Si l'on veut obtenir un bon modèle tronqué au sens de la mesure $\mathcal{M}$, alors il faut négliger les modes de faible énergie.

Modèle amalgamé

La mesure $\mathcal{M}$ s'écrit pour un modèle amalgamé

$$\mathcal{M} = \sum_{k=1}^n \mathcal{M}_k$$

avec:

$$\mathcal{M}_k = \sum_{m=2}^{n_k} \left(\left[\sum_{j=1}^p \int_{\mathcal{D}_t} x_{u^k(m)j}^2 dt \right] - \left[\omega_k^2(m) \sum_{j=1}^p \int_{\mathcal{D}_t} (\tilde{x}_{kj}(t))^2 dt \right] \right) \quad (\text{F.16})$$

et se réduit dans le cas des sollicitation tests (échelon: $s=0$ et Dirac: $s=1$)

$$\mathcal{M}_k = \sum_{m=2}^{n_k} \left[-\frac{\lambda_m^{2s-1}}{2} \left(\sum_{l=1}^p B_{ml}^2 \right) [\Delta e^{2\lambda_m t}]_{t_2}^{t_1} + \left[\frac{2\lambda_z \lambda_m^{2s}}{(\lambda_z + \lambda_m)^2} \frac{(\sum_{l=1}^p B_{zl} B_{ml})^2}{\sum_{l=1}^p B_{zl}^2} \frac{([\Delta e^{(\lambda_z + \lambda_m)t}]_{t_1}^{t_2})^2}{[\Delta e^{2\lambda_z t}]_{t_1}^{t_2}} \right] \right] \quad (\text{F.17})$$

F.4 Ajustement temporel du modèle amalgamé

On se propose ici d'ajuster les constantes de temps de chaque mode amalgamé afin de diminuer la valeur numérique de la mesure de l'erreur. Il faut bien noter que cette étape n'intervient qu'une fois la méthode d'amalgame appliquée dans son intégralité. En particulier, tous les sous-espaces d'amalgame H_k $k = 1 \dots n$ sont définis et les coefficients de décomposition $\omega_k(i)$ $k = 1 \dots n$ $i = 1 \dots n_k$ connus.

Considérons le sous-espace H_k et sa contribution $\mathcal{M}_k$ à l'erreur de réduction. Nous avons vu ci-dessus que

$$\mathcal{M}_k = \sum_{i=1}^4 \int_{\mathcal{D}_t} t_i(t) dt \quad (\text{F.18})$$

Dans (F.18), nous considérons être en possession de tous les paramètres sauf la valeur propre $\tilde{\lambda}_k$ associée au mode amalgamé $\tilde{V}_k(M)$. La fonctionnelle $\mathcal{M}_k$ a alors un minimum absolu $\tilde{\lambda}_k^{min}$ que nous recherchons ici. Ce minimum vérifie

$$\left(\frac{\partial}{\partial \tilde{\lambda}_k} \mathcal{M}_k \right)_{\tilde{\lambda}_k = \tilde{\lambda}_k^{min}} = 0 \quad (\text{F.19})$$

Pour calculer la dérivée de $\mathcal{M}_k$ par rapport à $\tilde{\lambda}_k$, il est nécessaire d'explicitier X et $\tilde{X}_k$. En prenant² des sollicitations de type échelon et avec les notations utilisées dans la méthode d'amalgame (voire §3.3.2) nous obtenons

$$\frac{\partial}{\partial \tilde{\lambda}_k} \mathcal{M}_k = \sum_{j=1}^{n_k} \dot{\mathcal{M}}_{kj} \quad (\text{F.20})$$

avec

$$\begin{aligned} \dot{\mathcal{M}}_{kj} = -a_j \frac{\left[\Delta \left((\tilde{\lambda} + \lambda_{u^k(j)})t - 1 \right) e^{(\tilde{\lambda} + \lambda_{u^k(j)})t} \right]_{t_2}^{t_1}}{(\tilde{\lambda} + \lambda_{u^k(j)})^2} \\ + b_j \frac{\left[\Delta \left(2\tilde{\lambda}t - 1 \right) t e^{2\tilde{\lambda}t} \right]_{t_2}^{t_1}}{\tilde{\lambda}^2} \quad (\text{F.21}) \end{aligned}$$

$$a_j = 2\omega_k(j) \sum_{l=1}^p \tilde{B}_{kl} B_{u^k(j)l}$$

$$b_j = \frac{\omega_k^2(j)}{2} \sum_{l=1}^p \tilde{B}_{kl}^2$$

Remarque: On rappelle que $\tilde{B}_{kl} = B_{u^k(1)l}$

On ne peut donc donner explicitement $\tilde{\lambda}_k^{min}$ et une résolution numérique de (F.21) s'impose. Cependant, il n'est pas difficile de localiser le zéro de (F.21) et en ce qui nous concerne, nous avons appliqué une simple dichotomie.

Soulignons que cet ajustement n'a d'intérêt que si les modes principaux sont mal choisis. En effet, dans la pratique les constantes de temps ajustées sont très peu différentes de celles des modes principaux. L'amélioration apportée au modèle amalgamé est insignifiante. Ceci prouve encore une fois que le modèle amalgamé ne peut être amélioré au sens de la mesure de l'erreur retenue. Autrement dit, le modèle amalgamé est optimal.

²Nous avons volontairement omis de considérer ici les sollicitations de type impulsion car le résultat de la dérivation est long. Mais il n'y a pas de difficulté pour le faire car il s'agit d'une dérivation classique.