



Modélisation numérique et approche expérimentale du contact en dynamique : Application au contact aubes/carter de turboréacteur

Etienne Arnoult

► To cite this version:

Etienne Arnoult. Modélisation numérique et approche expérimentale du contact en dynamique : Application au contact aubes/carter de turboréacteur. Mécanique des structures [physics.class-ph]. Université de Nantes, 2000. Français. NNT : . tel-01792898

HAL Id: tel-01792898

<https://theses.hal.science/tel-01792898>

Submitted on 16 May 2018

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.

UNIVERSITE DE NANTES

**ÉCOLE DOCTORALE
SCIENCES POUR L'INGÉNIEUR DE NANTES**

Année 2000

THÈSE DE DOCTORAT

Discipline : Sciences Pour l'Ingénieur
Spécialité : Génie Mécanique

présentée et soutenue publiquement par

ARNOULT Etienne

le 13 janvier 2000

à l'École Centrale de Nantes

**MODELISATION NUMERIQUE ET APPROCHE EXPERIMENTALE
DU CONTACT EN DYNAMIQUE :
APPLICATION AU CONTACT AUBES/CARTER DE TURBOREACTEUR**

JURY :

Président	G. Coffignal	Professeur, ENSAM, Paris
Rapporteurs	L. Jézéquel	Professeur, Ecole Centrale de Lyon, Ecully
	J. Piranda	Professeur, Université de Franche Comté, Besançon
Examineurs	M. Berthillier	Ingénieur de Recherche, SNECMA, Villaroche
	J. Bonini	Ingénieur de Recherche, SNECMA, Villaroche
	H. Bung	Ingénieur de Recherche, CEA Saclay
	B. Cochelin	Professeur, LMA, Marseille
	B. Peseux	Professeur, Ecole Centrale Nantes
	C. Wielgosz	Professeur, LGC Nantes – St Nazaire

Directeur de thèse : Bernard PESEUX
Laboratoire Mécanique et Matériaux
Division Structure
Ecole Centrale de Nantes
BP 92101
44321 Nantes Cedex 03

N° ED : 82-434

Table des Matières

Introduction générale	5
I Détection du contact et modélisation spline	9
Chapitre I Position du problème - Fonctions splines	11
I.1 Cadre de l'étude	12
I.2 Environnement mathématique	14
I.2.1 Définitions	15
I.2.2 B-splines	15
Chapitre II Modélisation spline des courbes et des surfaces	19
II.1 Eléments de construction	20
II.1.1 Choix de la représentation	20
II.1.2 Méthodes d'élaboration	20
II.1.3 Construction des vecteurs de noeuds	21
II.1.3.1 Influence de la multiplicité d'un paramètre	21
II.1.3.2 Paramétrages intrinsèques	22
II.1.3.3 Paramétrages référents	23
II.2 Modélisation des courbes gauches	25
II.2.1 Méthode directe	25
II.2.2 Interpolation	27
II.2.3 Lissage	28
II.2.3.1 Rappels théoriques	29
II.2.3.2 Mise en place	29
II.2.4 Remarques générales sur les splines	32
II.2.4.1 Une quatrième dimension: les NURBS	32
II.2.4.2 Les courbes fermées	34
II.2.4.3 Evaluation d'un point sur une courbe	37
II.3 Modélisation des surfaces	41
II.3.1 Méthode directe	41
II.3.2 Interpolation	42

II.3.2.1	Calcul des points de contrôle	43
II.3.2.2	Programmation	43
II.3.3	Lissage	44
II.3.3.1	Calcul des points de contrôle	45
II.3.3.2	Programmation	46
II.3.4	Résolution des systèmes	47
II.3.5	Prise en compte de l'épaisseur	48
Chapitre III	Recherche d'intersection courbe/surface	51
III.1	Recherche d'intersection courbe/surface	52
III.1.1	Méthodes couramment utilisées en éléments finis	52
III.1.1.1	Approches hiérarchiques	52
III.1.1.2	Approches par voisinage	53
III.1.2	Principe des méthodes analytiques	54
III.1.3	Méthode proposée	54
III.1.4	Difficultés rencontrées	57
III.1.4.1	Exploration d'une zone	57
III.1.4.2	Taille des incréments	58
III.1.4.3	Gestion des incertitudes	61
III.1.4.4	Choix des paramètres définissant les incréments	64
III.1.4.5	Remarques	65
III.1.5	Exemples	66
III.1.5.1	Intersections multiples et franches	66
III.1.5.2	Points doubles	66
III.1.5.3	Zone de tangence	67
III.1.6	Remarque: l'intersection surface/surface	68
III.2	Pré et post-traitement des résultats	69
III.2.1	Données requises	69
III.2.1.1	Les données indispensables	70
III.2.1.2	Les critères cachés	70
III.2.2	Evaluation de la pénétration	71
III.2.2.1	Définition	71
III.2.2.2	Organisation du calcul	71
III.2.2.3	Exemple	72
III.2.3	Reprojection sur une facette	72
III.2.3.1	Calcul des efforts de contact: présentation théorique	73
III.2.3.2	Coefficients barycentriques	76
III.2.3.3	Vecteur normal	79
Chapitre IV	Résultats numériques	81
IV.1	Cas du code SAMCEF	82
IV.1.1	Exemple traité	82

IV.1.2 Résultats	83
IV.1.2.1 Examen de la force de contact	84
IV.1.2.2 Comparaison des temps de calcul	87
IV.2 Cas du code PLEXUS	89
IV.2.1 L'implémentation	89
IV.2.2 Exemples traités	89
IV.2.2.1 Premier exemple	90
IV.2.2.2 Second exemple	91
IV.2.3 Conditionnement des matrices	93
IV.2.3.1 Rappels théoriques	93
IV.2.3.2 Cas de l'interpolation	95
IV.2.3.3 Cas du lissage	97
IV.2.3.4 Comparaison interpolation / lissage	99
Conclusion	101
 II Approche expérimentale	 103
Chapitre V Banc d'essai et analyse modale	105
V.1 Description du banc d'essai	106
V.1.1 Le carter	106
V.1.2 Le montage du laser	107
V.1.2.1 Acquisition modale et vélocimétrie	107
V.1.2.2 Montage pratique et réglages	109
V.1.3 Chaîne d'acquisition	111
V.2 Identification modale du carter	113
V.2.1 Modes des structures axisymétriques	113
V.2.2 Démarche générale	114
V.2.3 Résultats	116
V.2.4 Modèle éléments finis	116
V.2.4.1 Premier maillage	118
V.2.4.2 Prise en compte du défaut	120
V.2.4.3 Etude d'influence	123
V.2.4.4 Remarque: effet de la rotation	129
 Chapitre VI Essais de contact	 131
VI.1 Notion d'interaction	132
VI.1.1 Théorie	132
VI.1.2 Application au banc d'essai	134
VI.2 Création du contact et premiers résultats	134
VI.2.1 Montage	134
VI.2.2 Détermination négative	136

VI.2.3	Détermination positive	140
VI.2.3.1	Essais dans la zone médiane	140
VI.2.3.2	Essais dans la zone haute	142
VI.2.4	Modification du montage et nouveaux résultats	143
Conclusion		145
Conclusion générale		147
Références bibliographiques		149
Annexe A	Expression analytique complète des B-splines élémentaires de bas degré	155
Annexe B	Projection d'un point sur une droite ou un plan	159
Annexe C	Cas des courbes planes cartésiennes	163
Annexe D	Algorithme de recherche du contact	167
Annexe E	Jeu de données PLEXUS	171
Annexe F	Recherche d'intersection surface/surface	175
Annexe G	Montage expérimental	179
Annexe H	Vélocimétrie laser	183
Annexe I	Rappels sur la théorie des plans d'expérience	191

Introduction générale

L'objet de ce mémoire est de présenter les travaux réalisés dans le cadre d'une thèse financée par le Ministère de l'Enseignement Supérieur et de la Recherche, en partenariat avec SNECMA (Société Nationale d'Etude et de Construction des Moteurs d'Avions), et préparée au sein de la Division Structures du Laboratoire Mécanique et Matériaux de l'Ecole Centrale de Nantes.

LE CONTEXTE

Le sujet, proposé par SNECMA, consiste en une étude des problèmes de contact qui peuvent survenir dans les turboréacteurs, notamment lorsqu'une roue aubagée, depuis l'entrée du moteur (roue fan) jusqu'à la sortie (compresseurs intermédiaires) est amenée à entrer en contact avec le carter qui la contient.

Ce type de contact, dont l'origine peut être la conséquence d'une manœuvre (contact par effet gyroscopique), ou bien l'apparition d'un balourd au niveau du rotor (ingestion de glace ou d'oiseaux), peut avoir des conséquences importantes sur l'ensemble du moteur et remettre en cause l'intégrité de l'appareil, et par-là même la sécurité des passagers. C'est pourquoi une attention toute particulière est portée à l'étude de ces phénomènes dans les axes actuels de recherche du bureau des méthodes de Snecma.

En fait, les problèmes de contact peuvent se scinder en deux catégories:

- ◇ d'une part, les phénomènes transitoires (quelques millisecondes), qui peuvent générer des contraintes très importantes dans l'ensemble de la structure; c'est le cas, par exemple, lors d'une perte d'aube, à la suite de laquelle le moteur s'arrête après avoir subi des taux de décélération de l'ordre de 35000 tr/min par seconde. Ces problèmes ne seront pas abordés dans ce mémoire;
- ◇ d'autre part, les phénomènes qui s'étendent sur des échelles de temps plus longues (légers contacts provoqués par exemple par un faible balourd dû aux petites différences géométriques entre les aubes), mais qui peuvent tout autant induire des contraintes importantes dans le moteur, par sollicitation des modes propres de vibration des composants.

Une bonne compréhension de ces derniers phénomènes constitue la motivation première

du sujet proposé. Il a en effet été observé que, suite à des problèmes de contact, certains phénomènes dynamiques pouvaient apparaître, comme:

- ◇ une excitation forcée des modes propres de vibration d'une aube, avec pour conséquence possible des fragmentations en sommet d'aube;
- ◇ un couplage, suite au frottement, entre les modes vibratoires du carter et ceux de la roue aubagée;
- ◇ une instabilité dynamique, également induite par le frottement, conduisant à des phénomènes de précession inverse.

Il est aisément concevable que la solution consistant à augmenter le jeu existant entre aube et carter, pour réduire ces contacts, n'est pas acceptable: les performances du moteur s'en trouveraient affectées. L'objectif d'optimisation des rendements des turboréacteurs passe même par la diminution de ces jeux. L'idée implicitement contenue dans cet objectif est donc la conception de systèmes aubes/carter tolérants au contact: la nécessité de bien comprendre la genèse des problèmes dynamiques décrits ci-dessus n'en prend que plus d'importance.

La démarche adoptée par accéder à cette meilleure compréhension passe tout d'abord par l'observation, sur un cas concret aisément manipulable, de phénomènes vibratoires induits par le contact, puis par leur modélisation à l'aide d'outils numériques adaptés; une phase de confrontation des résultats fournis par les deux méthodes est alors possible, dont le double objectif est d'une part la validation des outils numériques, et d'autre part l'apport d'un éclairage nouveau sur les observations expérimentales, permettant d'en approfondir ainsi l'analyse. Tels étaient les objectifs que nous nous étions fixés.

Au cours de la thèse, le programme présenté n'a pu être réalisé dans son intégralité. Il a donc été choisi de ne présenter ici que les deux phases abordées, la confrontation numérique/expérimentale constituant la prolongation naturelle de ces travaux. Le plan du mémoire va donc s'articuler suivant deux parties qui seront présentées de la façon suivante:

- ◇ la première, numérique, aura pour objet de développer un outil utilisable dans un code éléments finis en dynamique afin d'améliorer les performances des simulations;
- ◇ la seconde, expérimentale, aura pour double objectif, tout d'abord de fournir un ensemble de données pouvant donner lieu à une simulation numérique permettant de valider les développements précédents, et ensuite d'essayer de mieux comprendre un phénomène de couplage vibratoire entre le carter et la roue aubagée (interaction).

L'APPROCHE NUMÉRIQUE

La première partie du rapport concerne les développements numériques réalisés autour de la détection du contact. Cette phase est très importante pour la simulation d'un

mouvement entre deux corps puisque d'elle dépend les efforts appliqués à chaque pas de temps à chacun de ces corps, efforts qui conditionnent leurs mouvements ultérieurs. Une mauvaise prise en compte d'une pénétration, une évaluation erronée d'un vecteur normal peuvent ainsi orienter le calcul, au mieux vers des configurations qui ne correspondent à aucune réalité physique observée, et au pire à une divergence numérique.

Malheureusement, le soin requis pour la bonne gestion du contact coûte souvent cher en temps de calcul. La principale difficulté à résoudre provient du problème de facétisation, qui rend discontinu le vecteur normal aux surfaces, et qui ne prend pas en compte de façon correcte la courbure des structures impliquées. L'augmentation de la finesse du maillage allant de pair avec une augmentation du temps de calcul, une nouvelle méthode a été recherchée (chapitre I): il s'agit, à chaque fois qu'une détection de contact doit avoir lieu, de construire des entités définies analytiquement, d'en chercher les intersections et de reprojeter les résultats sur le maillage éléments finis. Les fonctions splines ont été retenues pour procéder à la modélisation de ces entités.

Le chapitre II rappelle les notions importantes nécessaires à la mise en place de la construction de courbes et de surfaces à l'aide de fonctions splines. Trois types d'approches sont notamment présentés en détail: méthode directe, interpolation et lissage au sens des moindres carrés.

Le chapitre III aborde ensuite la recherche de contact proprement dite. L'algorithme itératif que nous avons développé consiste à faire glisser le long de la courbe une boule dont le rayon diminue à chaque itération, et à regarder l'évolution de la trace éventuelle laissée par cette boule sur la surface. Le processus permet ainsi d'obtenir les points d'intersection entre courbe et surface, et par la suite, la liste de tous les nœuds de la courbe ayant pénétré la surface. Un ensemble de grandeurs sont alors calculées, qui sont ensuite transmises au code pour lui permettre de poursuivre le calcul.

Le chapitre IV présente deux exemples d'intégration des routines de recherche de contact dans des codes éléments finis de structures très différentes:

- ◊ Samcef (paragraphe IV.1), code implicite, dont l'accès est possible via un élément utilisateur déjà existant;
- ◊ Plexus (paragraphe IV.2), code explicite de dynamique rapide, pour lequel les développements ont été menés directement dans les fichiers source.

Les résultats obtenus sont encourageants: d'une part notre méthode de recherche de contact se révèle peu sensible aux variations de densité en éléments des maillages, et d'autre part, le surcoût en termes de temps de calcul semble ne pas être trop important, et en particulier peut être compensé par la possibilité de réduire la taille du modèle éléments finis.

L'APPROCHE EXPÉRIMENTALE

La seconde partie du rapport aborde la phase expérimentale du problème. Un banc d'essai a été mis en place afin de recréer le phénomène de contact entre une partie de carter de turboréacteur et un modèle simplifié d'aube.

Le chapitre V décrit le dispositif expérimental: le carter est fixé sur le plateau d'un tour vertical dont la vitesse de rotation est pilotée, et l'aube est rendue solidaire du porte-outil. La première phase consiste à analyser le banc d'essai, et notamment à procéder à l'analyse modale du carter, à l'aide d'une méthode utilisant un vélocimètre laser. Un modèle éléments finis est ensuite construit et comparé en termes de comportement modal avec le banc d'essai. Les différences qui apparaissent sont attribuables à un défaut de forme du carter, dont la géométrie exacte est ensuite relevée. Un nouveau maillage est réalisé, et le modèle obtenu fait alors l'objet d'une étude d'influence, par l'intermédiaire de la mise en place d'un plan d'expérience dont les facteurs sont ses caractéristiques géométriques et matériaux.

Les essais de contact sont ensuite abordés dans le chapitre VI. L'objectif, au-delà de l'obtention de données expérimentales qui serviront ultérieurement à la création d'un cas test permettant de valider la méthode de recherche numérique du contact, consiste à essayer de créer un phénomène de couplage vibratoire entre l'aube et le carter. Ce phénomène, dénommé *interaction* par les motoristes, correspond à l'installation dans l'aube et dans le carter de vibrations de même nature, qui s'entretiennent et peuvent éventuellement causer de sévères dégâts dans les moteurs. Le phénomène est rare, et sa genèse mal connue. Dans la configuration simplifiée qui est la nôtre, une relation nécessaire d'existence relie les caractéristiques modales des éléments impliqués et la vitesse de rotation du système.

Après des essais sur quelques valeurs possibles de paramètres d'interaction, il apparaît que le défaut de forme constitue un obstacle à l'observation du phénomène. Le dépôt d'un matériau abrasable sur la paroi intérieure du carter semble constituer un bon palliatif, mais les nouveaux résultats obtenus ne permettent cependant pas de conclure sur les conditions d'existence du phénomène.

Partie I

Détection du contact et modélisation spline

Cette partie regroupe l'ensemble des développements réalisés autour de la recherche analytique de contact dans un code éléments finis en dynamique.

La mise en place de cette démarche particulière est justifiée dans le paragraphe I.1. Il apparaît que les configurations traitées peuvent se ramener à des recherches d'intersection courbe/surface modélisées par splines. La théorie des splines est rappelée dans le paragraphe I.2, et la construction des courbes et des surfaces est détaillée respectivement dans les paragraphes II.2 et II.3.

Les éléments géométriques nécessaires à la réalisation de la détection du contact étant précisés, la recherche d'intersection que nous proposons est alors abordée dans le paragraphe III.1, et le traitement des résultats obtenus dans le paragraphe III.2. Les développements que nous avons réalisés dans deux codes éléments finis en dynamique sont finalement présentés dans le chapitre IV.

Chapitre I

Position du problème - Fonctions splines

Ce chapitre commence par situer le cadre de l'étude (paragraphe I.1): le problème consiste à essayer d'améliorer la simulation numérique de mouvements pour lesquels la prise en compte du contact conditionne fortement le résultat.

Pour cela, nous nous proposons de mettre en place un algorithme de traitement du contact qui s'appuie, dans la phase de détection, sur des entités analytiques construites à l'aide de fonctions splines. Quelques rappels très généraux sur l'environnement mathématique de ces fonctions particulières, et notamment sur les bases formées des B-splines, sont rassemblés dans le paragraphe I.2.

I.1 Cadre de l'étude

Les contacts dynamiques dans les turboréacteurs font l'objet de nombreuses études, tant sur le plan expérimental que sur le plan numérique. Pour ce dernier point, différents algorithmes de recherche et traitement du contact ont été développés. Ces algorithmes viennent tout naturellement s'intégrer dans des ensembles plus généraux tels que les codes éléments finis capables de gérer l'évolution globale d'une structure complexe.

Cependant, à l'intérieur même d'un code, il arrive que la performance des outils numériques se trouve diminuée du fait de leur utilisation dans un environnement non-optimal. Les limitations peuvent être d'origine extérieure au code, comme c'est le cas lorsque le processeur en charge des calculs n'est pas assez puissant; mais ces limitations peuvent aussi être dues à des problèmes internes, lorsqu'il y a interaction entre l'utilisation de plusieurs outils.

Dans le cas des outils dévolus au traitement du contact, la facétisation constitue un bon exemple de ce dernier type de problème. La facétisation résulte en fait de la discrétisation géométrique de la structure étudiée. Lors des études dynamiques, en particulier lorsque le contact intervient, deux situations sont révélatrices des problèmes occasionnés:

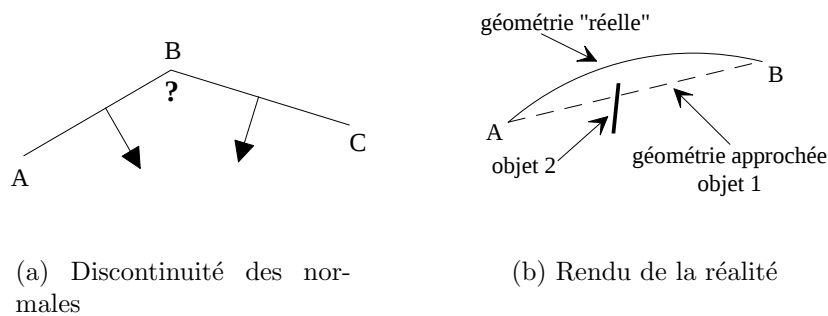


Figure I.1 : Problèmes liés à la discrétisation

- ◇ sur le schéma I.I.1, les normales aux segments AB et BC sont parfaitement définies sur chaque segment pris séparément; pourtant, la normale au point B n'est pas connue;
- ◇ sur le schéma I.I.1, un algorithme de détection du contact trouverait sur le segment AB une intersection qui n'existe pas réellement.

Ces difficultés peuvent être contournées: le problème de la détermination de la normale est classiquement résolu en attribuant au nœud B une normale portée par la somme des deux normales adjacentes (voir figure I.I.1); quant aux effets de la seconde difficulté, ils peuvent être atténués en ajoutant des nœuds supplémentaires entre les nœuds A et B, ce qui revient à augmenter la taille du maillage (voir figure I.I.1).

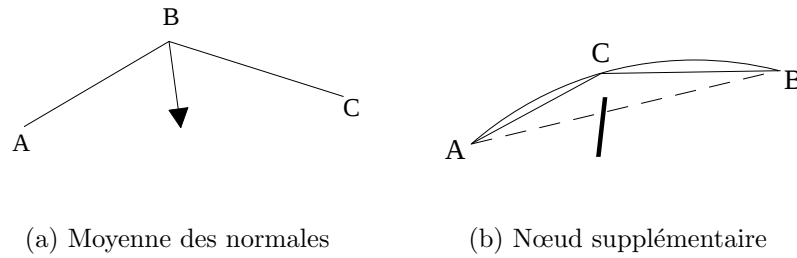


Figure I.2 : Solutions couramment envisagées

Cependant, si la première des deux solutions exposées constitue un palliatif acceptable, la seconde apparaît, dans la pratique, peu raisonnable, car la taille du maillage, ainsi que la taille propre des éléments, a une influence importante sur les temps de calcul.

L'objectif de l'étude menée ici est de proposer une méthode de recherche des points d'intersection de deux entités géométriques dans un cas bien particulier: celui où une aube va rencontrer un carter. Les problèmes de facétisation dans cette configuration entraînent en effet une modélisation très mauvaise de la rotation de la roue aubagée dans le carter, et peuvent conduire à un blocage *numérique* des pièces l'une par rapport à l'autre du fait du trop petit nombre de nœuds généralement retenus pour modéliser le carter. C'est pourquoi la démarche suivante est proposée; elle vient s'insérer dans la phase de calcul des codes éléments finis dédiée à la recherche du contact et consiste en:

Etape 1: recherche grossière des zones potentielles de contact:

- ◇ aubes impliquées,
- ◇ partie du carter en vis-à-vis;

Etape 2: pour chaque zone retenue, construction d'éléments géométriques s'appuyant sur les nœuds du maillage:

- ◇ courbe gauche pour le sommet de l'aube,
- ◇ surface pour le carter;

Etape 3: recherche de l'intersection des deux géométries;

Etape 4: projection des résultats sur le maillage.

Les étapes importantes de cette démarche sont les étapes 2 à 4. Le choix de la méthode de construction des entités géométriques est en particulier fondamental: il faut en effet pouvoir construire des éléments qui assurent une continuité suffisante pour supprimer les problèmes de facétisation, tout en restant aisément manipulables. Les fonctions *SPLINES*, déjà éprouvées et souvent utilisées dans bien des domaines, offrent une souplesse et une stabilité numérique qui justifient grandement leur choix.

I.2 Environnement mathématique

Le terme de *spline* (littéralement *baguette souple*) est apparu pour la première fois dans les travaux de Schoenberg en 1946 [SCH46a] [SCH46b]. Pourtant, la notion mathématique associée à ce terme est beaucoup plus ancienne puisqu'elle est déjà présente dans quelques travaux du début du siècle [HAM91], les premières approches des fonctions appelées aujourd'hui *B-splines* ayant même été réalisées par N. Lobatchevsky, au début du XIX^{ème} siècle [SHE97]. Malgré tout, la théorie, telle qu'elle est connue actuellement, ne s'est vraiment développée qu'à la fin des années 50 et au début des années 60, avec les travaux successifs de P. de Casteljau (chez Citroën) [DC93] et de P.E. Bézier (chez Renault) [BEZ93], qui ont d'ailleurs tous les deux laissés leurs noms à des outils splines encore fréquemment utilisés de nos jours (algorithme de calcul d'un point sur une courbe pour le premier, courbe spline particulière pour le second). La mise en forme moderne de la théorie doit beaucoup aux contributions de C. de Boor [DB78] et M. Cox [COX93], vers la fin des années 70.

Actuellement, les splines sont fréquemment utilisées dans les outils informatiques manipulant des éléments géométriques (logiciels de DAO et CAO, pré-processeurs et mailleurs de codes éléments finis par exemple), puisqu'elles permettent de créer simplement des entités alliant souplesse de manipulation et esthétique. Mais elles sont également utilisées, parfois indirectement, dans des domaines plus théoriques. En mécanique, le développement de leur utilisation dans la phase même de calcul est assez récent, et cherche encore des justifications [PAT98]. Des fonctions splines ont déjà été mises à contribution, notamment pour le caractère de continuité qu'elles fournissent à moindre coût, dans des domaines aussi variés que la méthode des éléments de frontière (éléments de Overhauser par exemple [HAL90] [HIL90]), la résolution des équations de propagation d'ondes ([DAG99], [KAS98]), et la construction d'éléments finis, dans lesquels elles apparaissent pour l'établissement des fonctions de forme (poutre [PAT99], plaque [YUE99], coque [BEN94a]).

Une présentation rigoureuse de l'environnement mathématique des splines ferait appel à des notions complexes telles que les espaces faibles de Tchebycheff. Cependant, l'objectif de cette section n'est pas de reproduire d'excellents ouvrages tels que ceux de de Boor [DB78] ou de Nürnberger [NUR89] qui traitent en détail des splines dans un cadre mathématique très général et très formel, mais d'en extraire les grands résultats qui permettent d'approcher plus aisément le problème à résoudre. En cela, le cours du professeur C.K. Shene, de l'Université Technologique du Michigan, a été très utile [SHE97].

Une façon simple de présenter les splines est de dire qu'il s'agit de fonctions définies par morceaux, telles que tous les morceaux se recollent "bien" (c'est-à-dire jusqu'à un certain ordre de dérivation). Seules seront abordées ici les splines polynomiales: sur chaque intervalle, la fonction est définie par un polynôme.

I.2.1 Définitions

Soit $[a, b]$ un intervalle de $\mathbb{R}$, et soit $\Omega_{npt} = (t_i)_{i=0, npt}$ une partition, appelée *vecteur de nœuds*¹, de cet intervalle telle que $t_0 = a$, $t_{npt} = b$ et $t_0 < \dots < t_{npt}$.

La fonction $s : [a, b] \rightarrow \mathbb{R}$ est une fonction spline polynomiale de degré n ($n > 0$) si elle vérifie les deux conditions suivantes:

$$(i) \quad s \in \mathcal{C}^{n-1}([a, b])^2$$

$$(ii) \quad \forall i \in [0..npt - 1] \quad \forall t \in [t_i, t_{i+1}] \quad s \in \mathcal{P}_n^3$$

L'ensemble des splines polynomiales de degré n (il est fréquent de parler de *l'ordre* de la spline, qui est par définition $k \triangleq n + 1$) définies sur Ω_{npt} sera noté $\mathcal{S}_n(\Omega_{npt})$. Cet ensemble est un espace vectoriel de dimension $npt + n$.

I.2.2 B-splines

Parmi toutes les bases possibles de $\mathcal{S}_n(\Omega_{npt})$, il en est une particulière constituée des B-splines (Basic Splines) de degré n [DB93].

Une B-spline de degré n est un élément de $\mathcal{S}_n(\Omega_\infty)$ qui vérifie⁴:

$$\exists i \in \mathbb{Z} \quad \forall t \in \mathbb{R}$$

$$\diamond \quad s(t) = 0 \quad \text{sauf sur} \quad [t_i, t_{i+n+1}]$$

$$\diamond \quad \forall t \in [t_i, t_{i+n+1}] \quad s(t) \geq 0$$

$$\diamond \quad s \in \mathcal{C}^{n-1}([t_i, t_{i+n+1}])$$

La définition des B-splines utilise la notion de différence divisée. Ainsi, la B-spline de degré n “accrochée” au point t_i , notée $B_{ni}(t)$, est définie par (voir figure I.3 pour les notations):

$$B_{ni}(t) = (t_{i+n+1} - t_i) \cdot [t_i..t_{i+n+1}] \cdot q_n(\bullet, t)$$

avec:

$$\diamond \quad q_n(t_i, t) = \begin{cases} (t_i - t)^n & \text{si } t_i \geq t \\ 0 & \text{si } t_i < t \end{cases}$$

¹l'ensemble Ω_{npt} est également parfois appelé *séquence nodale*

² $\mathcal{C}^{n-1}([a, b])$ est l'ensemble des fonctions $n - 1$ fois dérivables sur $[a, b]$, la dérivée $n - 1$ étant continue

³ $\mathcal{P}_n$ est l'ensemble des polynômes de degré inférieur ou égal à n

⁴ $\Omega_\infty = (t_i)_{i \in \mathbb{Z}}$ est une partition particulière telle que:

$$\lim_{i \rightarrow -\infty} t_i = -\infty \quad \text{et} \quad \lim_{i \rightarrow +\infty} t_i = +\infty$$

$$\diamond [t_0 t_1].f(\bullet) = \frac{f(t_1)-f(t_0)}{t_1-t_0}$$

$$\diamond [t_0..t_k].f(\bullet) = \frac{[t_1..t_k].f(\bullet)-[t_0..t_{k-1}].f(\bullet)}{t_k-t_0}$$

Cette définition, bien que suffisante en elle-même, n'est pas souvent utilisée sous cette forme. Il lui est préféré la relation de récurrence suivante, dite relation de Cox-de Boor, qui en est une conséquence:

$$\diamond B_{00}(t) = \begin{cases} 1 & \text{si } t \in [t_0, t_1] \\ 0 & \text{si } t \notin [t_0, t_1] \end{cases}$$

$$\diamond B_{ni}(t) = \begin{cases} \frac{t-t_i}{t_{i+n}-t_i}B_{n-1,i}(t) + \frac{t_{i+n+1}-t}{t_{i+n+1}-t_{i+1}}B_{n-1,i+1}(t) & \text{si } t \in [t_i, t_{i+n+1}] \\ 0 & \text{si } t \notin [t_i, t_{i+n+1}] \end{cases}$$

Cette relation de récurrence permet de construire une B-spline de n'importe quel degré, sur un vecteur de nœuds quelconque.

Remarque: il est rare, au cours d'un problème, de recourir à des B-splines de degrés variés. Bien souvent, ce sont des B-splines de degré fixé une fois pour toutes qui sont utilisées, et ce degré n'excède pratiquement jamais 3. A titre indicatif, figurent en annexe A les expressions des B-splines pour les degrés 0 à 3 inclus, pour des vecteurs de nœuds quelconques.

Un exemple de spline cubique (degré 3) est présenté sur la figure I.3.

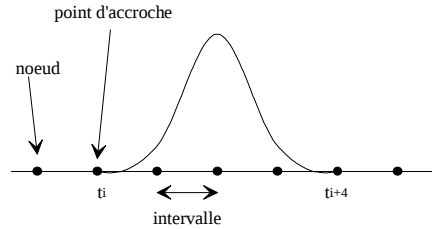


Figure I.3 : Exemple de B-splines pour $n = 3$

Il convient de noter que la relation de récurrence précédente suppose a priori distinctes les différentes valeurs de t_i . En fait, l'utilisation de la convention $0/0=0$ permet de considérer des vecteurs de nœuds ayant des points confondus⁵. Par ailleurs, la définition des B-splines entraîne la relation suivante, connue sous le nom de "partition de l'unité":

$$\forall t \in \mathbb{R} \quad \sum_{i \in \mathbb{Z}} B_{ni}(t) = 1$$

⁵dans le cas de nœuds confondus, il y a perte du caractère $\mathcal{C}^{n-1}$ de la B-spline en ces nœuds

L'ensemble $(B_{ni})_{i=-n, \dots, npt-1}$ ⁶ constitue une base de $\mathcal{S}_n(\Omega_{npt})$ (la démonstration de ce fait n'est pas présentée ici). Il en résulte le théorème suivant:

Théorème de projection : *Pour toute spline s de $\mathcal{S}_n(\Omega_{npt})$, il existe une unique famille de réels $(\alpha_i)_{i=-n, \dots, npt-1}$ telle que:*

$$s(t) = \sum_{i=-n}^{npt-1} \alpha_i B_{ni}(t)$$

Il est possible de donner une interprétation géométrique à ce théorème: la grandeur $s(t)$ peut être considérée comme le barycentre des grandeurs α_i affectées des coefficients $B_{ni}(t)$. La partition de l'unité et la positivité des B-splines assurent de plus que $s(t)$ se trouve à l'intérieur de l'enveloppe convexe des α_i .

D'un point de vue informatique, la programmation des B-splines constitue le premier pas indispensable à la réalisation de la démarche proposée plus haut. Les données nécessaires et les sorties sont regroupées dans le tableau I.1.

Entrées	Sorties
degré n vecteur de nœuds de taille $n + 2$	pour chaque intervalle: $n + 1$ coefficients

Tableau I.1: Entrées / sorties du programme

L'algorithme de calcul est présenté ci-dessous:

```

Lire les données
Initialiser  $B_{0,i}$  pour  $i = 0, \dots, n$            /boucle sur les nœuds/
Pour  $i = 1, \dots, n$                              /boucle sur le degré/
    Pour  $j = 0, \dots, n - i$                        /boucle sur les points d'accroche/
        Pour  $k = 1, \dots, n + 1$                    /boucle sur les intervalles/
            Calculer les  $i + 1$  coefficients de  $B_{ij}$  sur l'intervalle n°  $k$ 
            à l'aide de la formule de récurrence
        Fin pour
    Fin pour
Fin pour

```

⁶pour la définition des $npt + n$ fonctions B_{ni} , il faut adjoindre n nœuds support supplémentaires à $\mathcal{S}_n(\Omega_{npt})$

Chapitre II

Modélisation spline des courbes et des surfaces

Ce chapitre présente les résultats importants nécessaires à la modélisation des courbes et des surfaces construites à l'aide de fonctions splines.

Le paragraphe II.1 regroupe un certain nombre de notions utiles à la construction des courbes et des surfaces: il justifie le choix de leur représentation paramétrée, présente notamment les méthodes mises en œuvre pour approcher un nuage de points, ainsi que les différents types de paramétrages qui peuvent être utilisés.

Les paragraphes II.2 et II.3 détaillent la construction pratique respectivement d'une courbe et d'une surface, soit par une méthode directe, soit par interpolation, soit par lissage. Le paragraphe II.2.4 contient quelques rappels sur des notions plus générales comme les NURBS ou la manipulation de courbes fermées, mais présente surtout une justification importante d'un choix d'un algorithme relatif à l'évaluation des points (paragraphe II.2.4.3).

II.1 Éléments de construction

La construction d'entités géométriques à l'aide des fonctions splines utilise un certain nombre de méthodes communes aux courbes et aux surfaces. Ces points communs sont présentés ici.

II.1.1 Choix de la représentation

Il existe au moins deux façons de représenter une courbe ou une surface:

- ◇ soit à l'aide d'une formulation cartésienne (qu'elle soit ou non explicite): les coordonnées de chaque point $P(x, y, z)$ sont définies par un ensemble de relations de la forme $F(x, y, z) = 0$;
- ◇ soit à l'aide d'une formulation paramétrée: les coordonnées de chaque point s'expriment en fonction d'un ou de plusieurs paramètres $t_1, t_2, \dots, t_k$:

$$P(x(t_1 t_2 \dots t_k), y(t_1 t_2 \dots t_k), z(t_1 t_2 \dots t_k))$$

Le gros avantage de la représentation paramétrique sur la représentation cartésienne est d'autoriser la manipulation de courbes et de surfaces fermées, avec des tangentes quelconques. C'est ce qui motive son choix dans la suite de l'étude (l'annexe C présente quelques résultats quant à la modélisation de courbes à partir d'une formulation cartésienne). Il faut noter cependant que ce mode de représentation implique une difficulté: connaissant la valeur des paramètres, il est aisé de trouver les coordonnées du point associé, alors que la connaissance des coordonnées d'un point ne permet a priori pas de retrouver la valeur des paramètres correspondants.

Remarque: cette non-linéarité n'est pas gênante au stade de la construction, mais constitue une difficulté handicapante lorsque les problèmes d'intersection sont abordés.

II.1.2 Méthodes d'élaboration

Le jeu de données d'un problème de modélisation utilisant des splines comporte un ensemble de points qui sert de support à la définition d'une entité géométrique particulière (une courbe ou une surface). Il existe de nombreuses possibilités pour construire une entité spline à partir de cet ensemble de points, mais trois seulement seront mises en œuvre ici.

Afin de fixer les idées, le cas de la modélisation des courbes est présenté: le jeu de données comporte npt points ordonnés P_i , c'est-à-dire qu'il est possible de doter la courbe gauche d'une abscisse curviligne s telle que l'application de $\mathbb{N}$ dans $\mathbb{R}$ qui a tout i fait correspondre l'abscisse s_i du point P_i soit croissante. Sur un intervalle à définir du paramètre t , et si le degré de la spline choisi pour approcher la courbe est n , l'expression de la courbe gauche sera (écriture vectorielle):

$$c(t) = \sum_{i=0}^{N-1} Q_i B_{ni}(t)$$

Cette écriture fait intervenir N points Q_i qui sont appelés points de contrôle de la courbe, la polygline reliant les points de contrôle dans l'ordre étant appelée polygone de contrôle. Le choix des Q_i peut s'effectuer selon l'une des trois méthodes suivantes:

1. prendre pour Q_i les points du jeu de données P_i : c'est la **méthode directe**, pour laquelle $N = npt$; cette méthode est la moins coûteuse en terme de temps de calcul et bien sûr la plus simple à mettre en place;
2. calculer les Q_i afin que la courbe résultat passe exactement par les points P_i : il existe ainsi npt valeurs particulières $\xi_1, \dots, \xi_{npt}$ du paramètre telles que:

$$P_i = \sum_{j=0}^{npt-1} Q_j B_{nj}(\xi_i)$$

il s'agit donc d'une **interpolation** et $N = npt$; cette méthode nécessite la résolution d'un système linéaire de taille $npt * npt$;

3. calculer les Q_i pour qu'il existe npt valeurs $\xi_1, \dots, \xi_{npt}$ du paramètre qui minimisent la quantité $\sum_{i=0}^{npt-1} (c(\xi_i) - P_i)^2$: c'est la **méthode de lissage** au sens des moindres carrés et la valeur de N est alors indépendante de npt ; cette méthode nécessite la résolution d'un système linéaire de taille $N * N$.

II.1.3 Construction des vecteurs de noeuds

Le choix d'une méthode de modélisation ne suffit pas pour définir une entité spline: il reste en effet à déterminer les vecteurs paramètres qui servent à construire les B-splines. Dans le cas de la méthode directe, un seul vecteur est nécessaire, mais dans le cas de l'interpolation comme dans celui du lissage, deux vecteurs doivent être définis, le premier associé aux points donnés et le second aux points de contrôle qu'il faut calculer.

II.1.3.1 Influence de la multiplicité d'un paramètre

Avant de décrire les procédés de calcul des vecteurs de nœuds, il convient de préciser une propriété intéressante portant sur la multiplicité. En effet, si un nœud apparaît plusieurs fois dans un vecteur paramètre, c'est-à-dire si deux ou plusieurs valeurs successives du vecteur paramètre sont égales, alors le nombre de B-splines non nulles pour ce nœud diminue. En particulier, si la multiplicité du nœud atteint n , degré choisi pour la modélisation, une seule fonction B-spline sera non nulle en ce nœud, et sa valeur sera 1 du fait de la partition de l'unité. Soit t_{i_0} ce paramètre: du fait de sa multiplicité, il vient:

$$c(t_{i_0}) = Q_{i_0} B_{ni_0}(t_{i_0}) = Q_{i_0}$$

De cette remarque somme toute géométrique, il faut retenir que si un paramètre t est inséré n fois dans un vecteur paramètre, alors l'un des points de contrôle sera situé **sur** la courbe.

Une exploitation classique de cette propriété réside dans la constitution du vecteur paramètre. En effet, en prenant $t_0 = t_1 = \dots = t_n$, n étant le degré des B-splines qui seront utilisées, le passage de la courbe par le premier point de contrôle est assuré. Il en sera de même pour le dernier point de contrôle en prenant $t_N = t_{N+1} = \dots = t_{N+n}$.

Au-delà de l'utilisation de cette propriété, le choix du paramétrage est encore très libre. M. Daniel a montré dans sa thèse [DAN89] que la stabilité de telle ou telle méthode dépend en partie du choix du type de construction du vecteur paramètre. Il existe deux grandes familles de vecteurs paramètres:

- ◇ les paramétrages intrinsèques, qui ne font appel à rien d'autre qu'au nombre de points utilisés, et éventuellement à leur répartition dans l'espace;
- ◇ les paramétrages référents, qui se déduisent d'autres vecteurs paramètres déjà existants.

Dans tout ce qui suit, il est supposé que N valeurs doivent être construites (indépendamment de l'éventuelle multiplicité de la première et de la dernière valeur).

II.1.3.2 Paramétrages intrinsèques

PARAMÉTRAGE UNIFORME

Le vecteur paramètre (t_i) est constitué simplement à l'aide d'une succession de valeurs équiréparties:

$$\forall i \quad t_{i+1} - t_i = \text{constante}$$

Dans la suite, cette méthode sera utilisée avec:

$$t_i = i \text{ pour } i \in [0, \dots, N-1]$$

Cette méthode sera référencée Pi1.

PARAMÉTRAGE SELON LES LONGUEURS DE CORDE

Ce vecteur paramètre est construit en respectant l'espacement des N points P_i du jeu de données:

- ◇ $t_0 = 0$
- ◇ $t_i = \text{dilat} * \frac{\sum_{j=0}^{i-1} \|P_j P_{j+1}\|}{L}$ pour $i = 1, \dots, N-1$, avec $L = \sum_{j=0}^{N-2} \|P_j P_{j+1}\|$ est la longueur de la polygline reliant les P_i , et *dilat* est un coefficient qui permet de dilater l'intervalle $[0, 1]$ sur lequel sont calculées les valeurs t_i (dans le cas présent, la valeur de *dilat* sera systématiquement prise égale à $N-1$)

Par la suite, cette méthode sera référencée Pi2.

II.1.3.3 Paramétrages référents

Dans certains cas, deux paramétrages sont nécessaires: le premier pour caractériser des points particuliers pour le calcul (points où l'interpolation est réalisée par exemple), et le second pour construire les B-splines (ce paramétrage est donc lié aux points de contrôle). Si un certain paramétrage ξ_i comprenant N_1 valeurs a déjà été construit (peu importe le schéma de construction), et si $N_2 + n + 1$ valeurs de t_i sont requises pour procéder à la suite des calculs, alors, les t_i sont choisies suivant l'extension de la formulation de de Boor proposée par Daniel dans sa thèse [DAN89]:

$$\begin{aligned} \diamond t_0 &= \dots = t_n \triangleq \xi_0 \\ \diamond t_{i+n+1} &= \frac{\xi_{i+i_1} + \dots + \xi_{i+i_2}}{i_2 - i_1 + 1} \text{ pour } i = 0, \dots, N_2 - n - 2, \text{ avec } i_1 = E(\alpha * i + 0.5) + 1 \text{ et} \\ &\quad i_2 = E(\alpha * i + 0.5) + n, E(.) \text{ désignant la partie entière, et } \alpha = \frac{N_1 - N_2}{N_2 - n - 2} \\ \diamond t_{N_2} &= \dots = t_{N_2+n} \triangleq \xi_{N_1-1} \end{aligned}$$

Cette méthode, qui sera par la suite référencée Pr1, revient globalement à placer les valeurs t_i entre les valeurs ξ_i . Elle introduit une condition sur la valeur de N_2 puisque, du fait de la définition de α , il faut respecter:

$$N_2 \neq n + 2$$

Par la suite, chaque fois que les cas $N_2 = n + 2$ ont été rencontrés, la valeur de N_2 a été remplacée par $N'_2 = N_2 + 1$.

Il convient par ailleurs de noter que dans le cas particulier où $N_1 = N_2 \triangleq N$ (cas de l'interpolation par exemple), ce schéma se simplifie, car alors $\alpha = 0$, et s'écrit:

$$\begin{aligned} \diamond t_0 &= \dots = t_n \triangleq \xi_0 \\ \diamond t_{i+n+1} &= \frac{\xi_{i+1} + \dots + \xi_{i+n}}{n} \text{ pour } i = 0, \dots, N - n - 2 \\ \diamond t_N &= \dots = t_{N+n} \triangleq \xi_{N-1} \end{aligned}$$

Par la suite, cette méthode sera référencée Pr2.

Le tableau II.1 rassemble les identifications des différents paramétrages qui seront utilisés par la suite.

Référence	Description
Pi1	uniforme
Pi2	longueur de corde
Pr1	extension de de Boor ($N \neq n_{pt}$)
Pr2	extension de de Boor ($N = n_{pt}$)

Tableau II.1: Identification des paramétrages

II.2 Modélisation des courbes gauches

Le jeu de données est constitué de npt points P_i , et le degré de modélisation pour les splines est n . Si les N points de contrôle, obtenus par l'une des trois méthodes du paragraphe II.1.2, sont les Q_i , l'expression générique de la courbe spline est:

$$c(t) = \sum_{i=0}^{N-1} Q_i B_{ni}(t)$$

II.2.1 Méthode directe

A l'avantage indéniable de la simplicité de mise en œuvre de cette méthode (il n'y a, en particulier, aucune inversion de matrice à effectuer puisque $Q_i = P_i$ pour tout i) s'oppose le problème du positionnement de la courbe calculée par rapport aux points donnés. En effet, contrairement à la méthode de lissage, qui fait passer la courbe *entre* les points de contrôle, la méthode directe peut générer une courbe assez éloignée du polygone de contrôle, particulièrement si la courbe ne présente aucun point d'inflexion, comme c'est le cas sur la figure II.1.

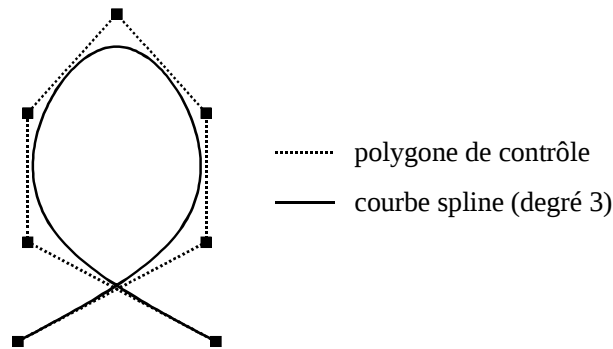


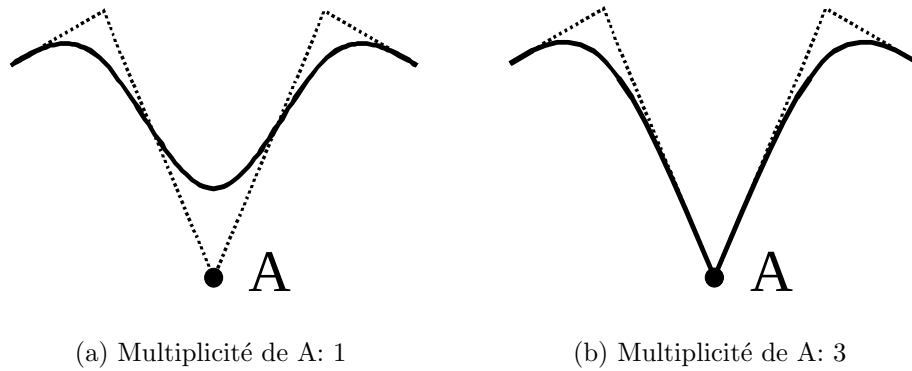
Figure II.1 : Situation d'une spline directe par rapport aux points de contrôle

Deux solutions sont alors possibles:

1. augmenter le nombre de points de contrôle;
2. augmenter la multiplicité de points jugés importants pour la modélisation.

La seconde solution a pour conséquence une perte, au niveau des points multiples, du caractère C^{n-1} de la courbe. Dans le cas extrême où la multiplicité d'un point atteint n , degré choisi pour la modélisation, la courbe passe par ce point, reste continue mais n'est plus dérivable au point considéré (voir figure II.2). Ce procédé est d'ailleurs à rapprocher de celui qui consiste à insérer un paramètre plusieurs fois dans un vecteur de nœuds (voir le paragraphe II.1.3.1).

Le paramétrage utilisé dans ce cas est le paramétrage uniforme (Pi1).

Figure II.2 : Influence de la multiplicité d'un point ($n = 3$)

Ainsi, d'un point de vue informatique, les données nécessaires à la réalisation du calcul et les sorties attendues sont regroupées dans le tableau II.2.

Entrées	Sorties
npt points de contrôle degré n	paramétrage pour chaque intervalle: $n + 1$ coefficients

Tableau II.2: Entrées / sorties du programme

L'algorithme de calcul est présenté ci-dessous¹:

```

Lire les données
Fabriquer le paramétrage
Construire la base des B-splines
Pour chaque coordonnée coord (soit  $x$ ,  $y$  ou  $z$ )
    Pour chaque intervalle du vecteur paramètre
        Calculer les coefficients de la spline
        à l'aide de la relation  $c(t) = \sum_i coord(P_i) \cdot B_{ni}(t)$ 
    Fin pour
Fin pour

```

Le calcul des coordonnées s'effectue alors de la façon suivante (une justification précise de cette approche est fournie dans le paragraphe II.2.4.3):

¹noter que le calcul des coefficients sur l'intervalle $[t_i, t_{i+1}]$ se limite aux indices compris entre $i - n$ et i , qui correspondent aux seules splines non-nulles pour l'indice i

Rechercher j tel que $t \in [t_j, t_{j+1}]$
 A l'aide des coefficients de la spline sur l'intervalle $[t_j, t_{j+1}]$,
 calculer les coordonnées

Un exemple est présenté sur la figure II.II.2.3.2.

II.2.2 Interpolation

Il s'agit ici de déterminer les npt points de contrôle Q_i de façon à ce que la spline passe exactement par les npt points donnés P_i , ce qui revient à écrire les npt relations suivantes:

$$\forall i \in [0..npt - 1] \quad P_i = \sum_{j=0}^{npt-1} Q_j B_{nj}(\xi_i)$$

où les ξ_i sont des valeurs particulières du paramétrage qu'il faudra choisir *avec pertinence*.

L'expression finale de la spline d'interpolation est donc de la forme:

$$c(t) = \sum_{i=0}^{npt-1} Q_i B_{ni}(t)$$

Cette expression utilise npt B-splines, ce qui nécessite donc la construction d'un vecteur paramètre (t_i) comportant $npt + n + 1$ nœuds. Le choix de ces paramètres n'est pas indépendant de la qualité du résultat ([DAN89]). En fait, pour construire un "bon" paramétrage, il est plus simple de raisonner *à l'envers*, c'est-à-dire de se donner de "bonnes" valeurs ξ_i (qui respectent, par exemple, la densité de répartition des points le long de la courbe), puis de construire le paramétrage (t_i) à partir de ces valeurs-là.

C'est pourquoi les valeurs ξ_i sont calculées suivant le schéma Pi2, et les valeurs t_i suivant le schéma Pr2.

Le système d'équations permettant d'accéder aux points de contrôle s'obtient en écrivant que pour chacune des npt valeurs de ξ_i , la relation suivante est vérifiée:

$$P_i = \sum_{j=0}^{npt-1} Q_j B_{nj}(\xi_i)$$

En notant $\{\mathcal{Q}\}$ le vecteur des inconnues (les points Q_i), et $\{\mathcal{P}\}$ le vecteur des points de passage (les points P_i), cette dernière relation peut s'écrire matriciellement:

$$\{\mathcal{P}\} = \mathbb{A} \cdot \{\mathcal{Q}\}$$

où: $\mathbb{A} = [a_{ij}] = [B_{nj}(\xi_i)]$. La résolution de ce système linéaire fournit les points de contrôle recherchés.

Remarque: Toutes les inversions matricielles de cette étude se feront par décomposition LU.

Entrées	Sorties
npt points donnés degré n	paramétrages t et ξ pour chaque intervalle: $n + 1$ coefficients

Tableau II.3: Entrées / sorties du programme

Les données nécessaires à la programmation ainsi que les sorties attendues sont rapportées dans le tableau II.3.

L'algorithme de programmation se présente donc de la façon suivante:

```

Lire les données
Fabriquer les paramétrages  $\xi$  et  $t$ 
Construire la base des B-splines (sur la base  $t$ )
Pour  $i, j = 0, \dots, npt - 1$ 
    Calculer  $a_{ij} = B_{nj}(\xi_i)$ 
Fin pour
Inverser  $\mathbb{A}$ 
Pour chaque coordonnée coord (soit  $x, y$  ou  $z$ )
    Calculer  $\mathbb{A}^{-1} * \{\mathcal{P}_{coord}\}$ 
    Pour chaque intervalle
        Calculer les coefficients de la spline
    Fin pour
Fin pour

```

Un exemple est présenté sur la figure II.II.2.3.2.

II.2.3 Lissage

Le lissage développé ici repose sur l'application d'une méthode des moindres carrés. Comme dans le cas de l'interpolation, les N points de contrôle utilisés ne sont plus les points donnés, mais doivent être déterminés et il n'y a plus de lien a priori entre N et npt . C'est pourquoi là aussi deux paramétrages sont nécessaires: le vecteur de nœuds (ξ_i) sera celui lié aux points donnés (les P_i), et le vecteur (t_i) celui lié aux points de contrôle recherchés (les Q_i). Les paramétrages proposés ici sont ceux préconisés par M. Daniel dans sa thèse [DAN89].

Les valeurs des ξ_i sont choisies uniformément (Pi1). Les valeurs des t_i sont choisies suivant l'extension de la formulation de de Boor (Pr1).

II.2.3.1 Rappels théoriques

Soit $\mathcal{L}^2(\xi_0, \xi_{npt-1})$ l'espace vectoriel des fonctions de carré sommable définies de $[\xi_0, \xi_{npt-1}]$ sur $\mathbb{R}$. Cet espace, muni du produit scalaire:

$$\langle f, g \rangle = \int_{\xi_0}^{\xi_{npt-1}} f(t)g(t)dt$$

est un espace préhilbertien réel. L'ensemble $\mathcal{S}_n(\Omega_m)$ ² des fonctions splines de degré n définies sur $[\xi_0, \xi_{npt-1}]$ est un sous-espace vectoriel de dimension finie de $\mathcal{L}^2(\xi_0, \xi_{npt-1})$. Soit f_x un élément de $\mathcal{L}^2(\xi_0, \xi_{npt-1})$ vérifiant:

$$\forall i \in [0..npt-1] \quad f_x(\xi_i) = x_{P_i}$$

où x_{P_i} désigne l'abscisse de P_i . Un théorème assure alors l'existence d'une fonction s_x de $\mathcal{S}_n(\Omega_m)$ qui soit la meilleure approximation de f_x dans $\mathcal{S}_n(\Omega_m)$ pour la norme associée au produit scalaire $\langle, \rangle$. La famille $(B_{ni})_{i=0,\dots,N-1}$ est une base de $\mathcal{S}_n(\Omega_m)$ sur laquelle s_x se décompose de façon unique:

$$s_x(t) = \sum_{i=0}^{N-1} \lambda_{xi} B_{ni}(t)$$

Le vecteur $\{\lambda_x\}$ est solution du système:

$$\mathbb{A} \cdot \{\lambda_x\} = \{V_x\}$$

où $\mathbb{A}$ est la matrice de terme général $a_{ij} = \langle B_{ni}, B_{nj} \rangle$ et où $\{V_x\}$ est le vecteur de terme général $v_{xi} = \langle B_{ni}, f_x \rangle$. Il en résulte un système $N \times N$ à résoudre.

Remarques:

- ◇ il n'y a mathématiquement aucune raison pour que les points de contrôle extrêmes obtenus par le calcul coïncident avec les points extrêmes initiaux. Néanmoins, ces points sont, dans la pratique, pour les méthodes de lissage et d'interpolation et pour une extrémité donnée, souvent très voisins l'un de l'autre. C'est pourquoi il est fréquent, une fois les calculs faits, de remplacer les points extrêmes obtenus par les points extrêmes initiaux. Il en résulte bien sûr une légère modification de la courbe (en particulier, l'interpolation n'en est plus une) ce qui nous a conduit à ne pas utiliser cet artifice très répandu;
- ◇ il n'y a aucune relation entre la courbe obtenue par lissage en prenant npt pour valeur de N et la courbe obtenue par interpolation.

II.2.3.2 Mise en place

D'un point de vue programmation, le recours au lissage génère davantage de calculs que la méthode directe, et même que la méthode d'interpolation. Au-delà de l'inversion

²la valeur de m est telle que $n + m = N$

matricielle, le simple calcul des termes du second membre pose à lui seul quelques difficultés. En effet, l'évaluation de l'intégrale par une méthode de Gauss (des méthodes de type Simpson se révélant numériquement insatisfaisantes) nécessite la connaissance des valeurs prises par la fonction à intégrer aux points de Gauss. Or si les B-splines peuvent être aisément calculées en tout point de l'intervalle $[\xi_0, \xi_{npt-1}]$, les valeurs des f_x ne sont connues qu'aux points ξ_i : il faut donc procéder à une interpolation entre ces points (l'interpolation choisie ici est linéaire).

Les données requises pour le calcul et les sorties attendues sont présentées dans le tableau II.4.

Entrées	Sorties
npt points donnés	paramétrage t et ξ
degré n	
N points de contrôle	pour chaque intervalle:
recherchés	$n + 1$ coefficients

Tableau II.4: Entrées / sorties du programme

L'algorithme de calcul est présenté ci-dessous:

```

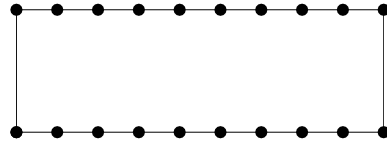
Lire les données
Fabriquer les paramétrages
Construire la base des B-splines
Pour  $i, j = 0, \dots, N - 1$ 
    Calculer  $a_{ij} = \langle B_{ni}, B_{nj} \rangle$ 
Fin pour
Inverser  $\mathbb{A}$ 
Pour chaque coordonnée coor (soit  $x, y$  ou  $z$ )
    Pour  $i = 0, \dots, N - 1$ 
        Calculer  $v_{i,coor} = \langle B_{ni}, f_{coor} \rangle$ 
    Fin pour
    Calculer  $\mathbb{A}^{-1} * \{V_{coor}\}$ 
    Pour chaque intervalle
        Calculer les coefficients de la spline
    Fin pour
Fin pour

```

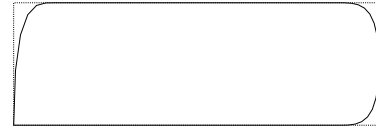
Une fois connus les coefficients de la spline par intervalle, le calcul des coordonnées d'un point à partir de la donnée du paramètre t est alors identique à celui effectué dans le cadre de la méthode directe.

Sur l'exemple ci-dessous, les résultats fournis par les trois méthodes (directe, interpolation et lissage) sont comparés. Cet exemple est présenté en détails par M. Daniel dans

sa thèse [DAN89], et est, selon ses propres termes, “difficile à traiter”. La donnée du problème (voir figure II.II.2.3.2) est un ensemble de 21 points définissant un rectangle: il y a donc 4 points avec changement brutal de direction (discontinuité du vecteur tangent). Des splines de degré 3 sont utilisées.



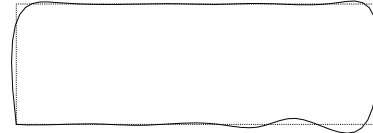
(a) Données



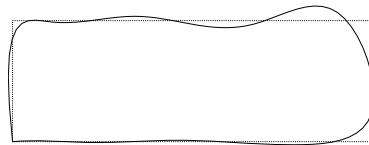
(b) Méthode directe



(c) Interpolation



(d) Lissage - 17 points de contrôle



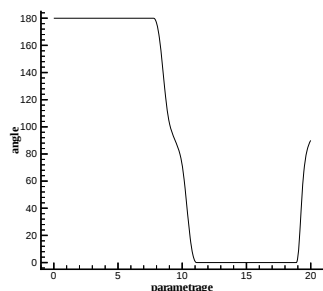
(e) Lissage - 12 points de contrôle

Figure II.3 : Comparaison méthode directe - interpolation - lissage

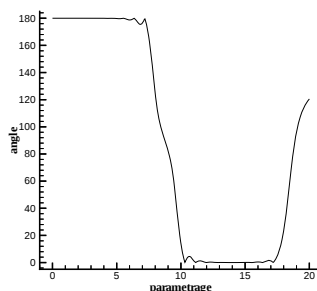
L’interpolation (voir figure II.II.2.3.2) présente des écarts importants avec la courbe initiale, notamment sur les côtés du rectangle.

Les résultats d’un lissage sont présentés sur les figures II.II.2.3.2 et II.II.2.3.2. Sur la figure II.II.2.3.2, 17 points de contrôle ont été demandés, et seulement 12 sur la figure II.II.2.3.2. Il apparaît que des oscillations (visibles sur les figures II.4 où sont représentées les évolutions des valeurs des angles que la normale aux courbes fait avec l’axe vertical) sont générées à chaque passage de singularité. Cet effet n’existe pas lorsque la méthode directe est utilisée (voir figure II.II.2.3.2).

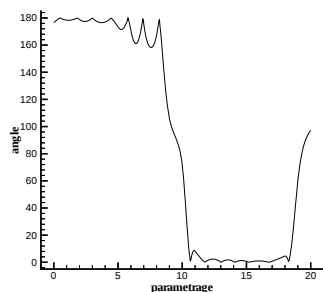
Les géométries qui seront étudiées par la suite (sommets d’aubes et carter) présentent



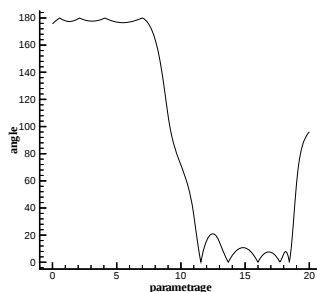
(a) Méthode directe



(b) Interpolation



(c) Lissage - 17 points de contrôle



(d) Lissage - 12 points de contrôle

Figure II.4 : Evolution de la normale aux courbes

moins de singularités que l'exemple du rectangle. Les problèmes d'oscillations ne seront donc pas aussi importants que dans ce cas et nous espérons que le choix de la méthode de modélisation ne souffrira pas de ce problème.

II.2.4 Remarques générales sur les splines

II.2.4.1 Une quatrième dimension: les NURBS

La formulation spline telle qu'elle est présentée ici n'est en fait qu'une forme simplifiée issue d'un contexte beaucoup plus général. Sans redéfinir complètement ce contexte, nous allons ici prendre un peu de recul par rapport aux développements effectués.

PRÉSENTATION

Nous avons vu qu'il était possible de donner de l'importance à un point d'une courbe simplement en augmentant la multiplicité de ce point dans le jeu de données. Cet artifice

a des conséquences qui peuvent être contraignantes par la suite (perte du caractère C^{n-1} de la courbe en ce point notamment). Par ailleurs, pour un degré n donné, il n'existe qu'un nombre limité de possibilités de valoriser un point: lui attribuer une multiplicité 2, puis 3, ... jusqu'à n . Une multiplicité supérieure à n ne sert à rien, puisqu'à partir de cette valeur, les deux parties de la courbe située de chaque côté du point concerné sont complètement indépendantes: le déplacement des points situés sur une des branches de la courbe n'entraîne aucune modification sur l'autre branche, comme cela est visible sur la figure II.5.

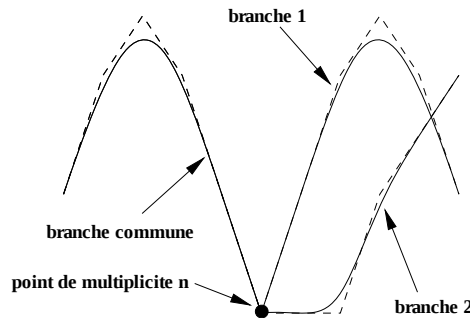


Figure II.5 : Effet d'un point de multiplicité n

Cette constatation laisse insatisfait et appelle des questions quant à la manière d'attribuer de l'importance à un point dans un jeu de données, sans pour autant affecter la continuité de la courbe ou de ses dérivées. Ceci peut être fait en passant dans une dimension supérieure, c'est-à-dire en attribuant une caractéristique supplémentaire aux points traités. Cette grandeur supplémentaire est appelée *poids*: chaque point de contrôle Q_i va donc se voir attribuer un poids ω_i dont l'importance sera fonction du rôle du point dans le problème. En particulier, un point affecté d'un poids nul sera complètement ignoré par la courbe, alors qu'un point affecté d'un grand poids sera attractif pour cette courbe.

La notion, qui vient d'être sommairement évoquée est la notion de NURBS (Non Uniform Rational B-Spline [RIE93]), dont l'expression mathématique générique prend la forme dans le cas d'une courbe de degré n définie par N points de contrôle (Q_i, ω_i) :

$$c(t) = \frac{\sum_{i=0}^{N-1} \omega_i Q_i B_{ni}(t)}{\sum_{i=0}^{N-1} \omega_i B_{ni}(t)}$$

LIEN ENTRE LES FORMULATIONS SPLINES

Il est possible, à partir de cette définition, d'établir le lien entre les différents éléments évoqués jusqu'ici. Une NURBS relie des éléments de $\mathbb{R}^4$. En faisant subir à cette entité une projection dans l'espace $\mathbb{R}^3$ suivant la direction $\omega = 1$ (ce qui revient à prendre le même poids pour tous les points), alors l'expression ci-dessus prend la forme (compte-tenu

de la partition de l'unité):

$$c(t) = \sum_{i=0}^{N-1} Q_i B_{ni}(t)$$

Il s'agit bien de l'expression des splines telle qu'elle a été présentée au début de la section II.2.

Une autre restriction concerne la relation entre le nombre de points npt du jeu de données et le degré n retenu pour modéliser les courbes: dans le cas particulier où $npt = n + 1$, la paramétrisation obtenue est la suivante:

$$t_i = 0 \text{ pour } i \in [0, \dots, npt - 1]$$

$$t_i = 1 \text{ pour } i \in [npt, \dots, 2npt - 1]$$

C'est le principe de la méthode de Bézier, pour laquelle les fonctions B-splines se confondent avec les polynômes de Bernstein, et s'écrivent sous la forme (cas d'une B-spline accrochée en t_i):

$$B_{n,i}(t) = C_n^i t^i (1 - t)^{n-i}$$

A noter qu'il existe des NURBS de Bézier, ainsi que des splines de Bézier.

L'ensemble de cette classification est rassemblée sur le schéma II.6.

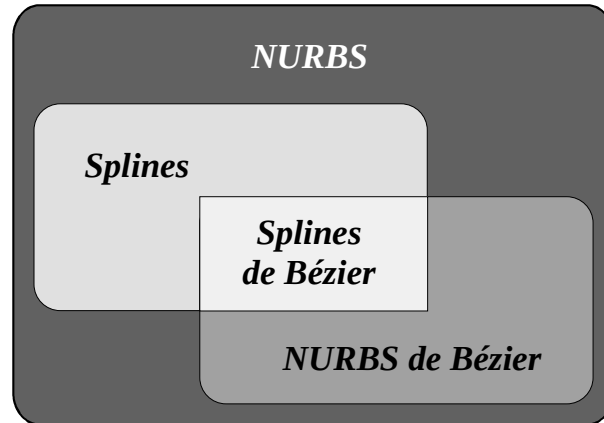


Figure II.6 : Classification des splines

II.2.4.2 Les courbes fermées

Un des arguments forts qui justifie l'emploi des splines est la possibilité qu'elles offrent de pouvoir manipuler n'importe quel type de courbe, en particulier des courbes qui présentent des tangentes quelconques comme les courbes qui se referment (cas du cercle). Les choix des paramétrages effectués, et en particulier l'égalité des paramètres initiaux et terminaux, assurent le passage de la spline obtenue par le premier et le dernier point du jeu de données. La modélisation d'une courbe fermée nécessite donc la double saisie

d'un point du jeu de données, une première fois au début du jeu, et une seconde fois à la fin. Cependant, cette précaution, qui garantit bien le caractère *fermé* de la courbe, ne garantit nullement une quelconque forme de continuité au niveau de la jonction. Plusieurs solutions sont envisageables pour résoudre ce problème.

UTILISATION DE POINTS MULTIPLES

Il est tout d'abord possible de s'arranger pour que dans le jeu de données, il y ait un recouvrement partiel de la courbe, c'est-à-dire que plusieurs points successifs ($n + 1$ de chaque côté du point de jonction pour assurer la continuité) soient rentrés deux fois, formant une liste du type:

$$\{P_{npt-n-1}, \dots, P_{npt-2}, P_{npt-1}, P_0, P_1, P_2, \dots, P_{npt-1}, P_0, P_1, P_2, \dots, P_n, P_{n+1}\}$$

Cette solution, pratique et simple de mise en œuvre, se révèle coûteuse dans le cas des méthodes de lissage et d'interpolation, car la taille des systèmes à inverser est augmentée de $2n + 2$.

UTILISATION DE CONDITIONS AUX LIMITES

Une autre solution consiste à introduire deux points de contrôle supplémentaires de chaque côté du point de jonction, et à imposer des conditions de continuité (portant sur les dérivées première et seconde de la courbe) sur ce point de jonction. Cette solution ne fait donc intervenir que deux équations supplémentaires. Les implications de cette méthode sont précisées ici dans le cas de l'interpolation.

Le jeu de données comprend npt points P_i de telle façon que $P_0 = P_{npt-1}$. Aux npt points de contrôle classiques, les $(Q_i)_{i=0, npt-1}$, viennent s'ajouter deux nouveaux points de contrôle notés Q_{-1} et Q_{npt} : il y a ainsi désormais $npt + 2$ inconnues au problème. L'expression de la spline prend alors la forme:

$$c(t) = \sum_{i=-1}^{npt} Q_i B_{ni}(t)$$

Les paramétrages sont donc légèrement modifiés. Pour les points donnés, le paramétrage (ξ_i) sera identique à celui utilisé dans une configuration "non-fermée" (voir paragraphe II.2.2) pour les indices i compris entre 0 et $npt - 1$. Deux paramètres ξ_{-1} et ξ_{npt} sont ajoutés:

- ◇ $\xi_{-1} = -\xi_1$
- ◇ $\xi_{npt} = \xi_{npt-1} + (\xi_{npt-1} - \xi_{npt-2})$

Pour le paramétrage (t_i) des points de contrôle recherchés, la définition suivante est adoptée (schéma de de Boor):

- ◇ $t_0 = \dots = t_n \triangleq \xi_{-1}$

$$\diamond t_{i+n+1} = \frac{\xi_i + \dots + \xi_{i+n-1}}{n} \text{ pour } i = 0, \dots, npt - n$$

$$\diamond t_{npt+2} = \dots = t_{npt+n+2} \triangleq \xi_{npt}$$

Le système matriciel pour cette interpolation s'écrira donc:

$$\forall i \in [0..npt - 1] \quad P_i = \sum_{j=-1}^{npt} Q_j B_{nj}(\xi_i)$$

ce qui conduit à npt équations, auxquelles vont s'ajouter les deux équations de raccordement:

$$c'(\xi_0) = c'(\xi_{npt-1})$$

$$c''(\xi_0) = c''(\xi_{npt-1})$$

Il convient de noter que cette méthode présuppose que le degré des splines mises en jeu est au moins égal à 2. Sur la figure II.7, un exemple est présenté: le jeu de données est constitué de 5 points qui définissent un carré (le premier et le dernier point sont confondus). La figure II.II.2.4.2 donne le résultat obtenu en utilisant la méthode d'interpolation sans modification (c'est-à-dire que seule la condition $P_0 = P_{npt-1}$ est réalisée): on y observe une discontinuité de la tangente à la courbe au niveau du point de jonction. La figure II.II.2.4.2 donne le résultat obtenu en ajoutant deux points supplémentaires dans la recherche: la courbe est alors continue en tout point, de même que sa tangente et sa courbure.

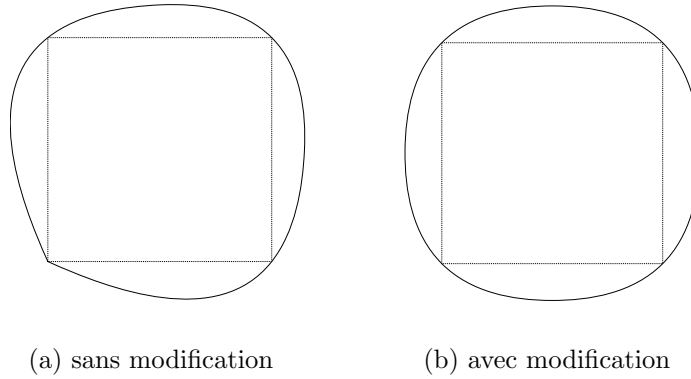


Figure II.7 : Interpolation d'un carré

REMARQUE

Ce paragraphe présente une méthode permettant de travailler avec des courbes fermées, qui sont considérées ici comme des courbes périodiques; elle est généralisable pour la création de surfaces cycliques dans une direction (cas de cylindres, de troncs de cônes...). Ces développements ne seront pas utilisés par la suite dans la phase d'implémentation qui a été réalisée; les contacts qui seront modélisés concernent une aube et la partie de carter

(c'est-à-dire une surface de révolution) en vis-à-vis: il serait donc coûteux, en termes de temps de calcul, de travailler avec l'ensemble du carter. Cependant, cette méthode pourrait être employée ultérieurement pour des modélisations simplifiées de rotation d'arbres dans des paliers.

II.2.4.3 Evaluation d'un point sur une courbe

La question du calcul d'un point à partir de la donnée d'une valeur du paramètre a déjà été abordée dans la section II.2.1. La méthode utilisée est très calculatoire: l'ensemble des coefficients des fonctions splines est tout d'abord calculé, puis stocké dans un tableau. Lors de l'évaluation d'un point, les coefficients nécessaires au calcul de ce point sont recherchés dans le tableau, puis recombinaison pour obtenir le résultat voulu. Il existe en fait une autre méthode, beaucoup plus élégante, qui ne nécessite pas le calcul direct des fonctions splines élémentaires: selon que la spline est de Bézier ou quelconque, l'algorithme porte le nom de de Casteljau ou de de Boor-Cox.

ALGORITHME DE DE BOOR-COX

Seul l'algorithme de de Boor-Cox, qui est une généralisation de l'algorithme de de Casteljau, sera présenté ici. Cet algorithme est celui proposé par de Boor et Cox en 1972 [COX72]; une autre formulation a été avancée en 1978 par de Boor [DB78], qui permet de réduire sensiblement le nombre d'opérations nécessaires, mais le principe de base en est le même. Il repose sur l'exploitation de l'influence de la multiplicité d'un nœud dans le vecteur paramètre (voir paragraphe II.1.3.1): si un paramètre t est inséré n fois dans un vecteur paramètre, alors l'un des points de contrôle sera situé **sur** la courbe.

Il s'agit donc à présent de procéder à l'insertion d'un nœud dans le vecteur paramètre, sans pour autant modifier l'allure générale de la courbe. Pour cela, il faut rappeler que le nombre de nœuds du vecteur paramètre est égal à $npt + n + 1$ (ici, npt est le nombre de points de contrôle: $npt = N$). Si ce nombre augmente, n étant fixé, cela signifie que npt a augmenté: l'insertion d'un paramètre se traduit donc par l'ajout d'un point dans le polygone de contrôle.

Soit $t \in [t_i, t_{i+1}[$: le calcul de $c(t)$ n'implique que les points de contrôle pour lesquels les fonctions B-splines sont non nulles, c'est-à-dire les points $P_{i-n}, \dots, P_i$. Pour insérer un nouveau point de contrôle, ces $n + 1$ points sont remplacés par $n + 2$ nouveaux points, notés $P_j^{(1)}$ définis comme suit:

$$\diamond P_{i-n}^{(1)} = P_{i-n}$$

$$\diamond P_{i+1}^{(1)} = P_i$$

$\diamond$ pour j compris entre $i - n + 1$ et i , le point $P_j^{(1)}$ est situé sur le segment $P_{j-1}P_j$ de telle sorte que:

$$P_j^{(1)} = (1 - \alpha_j)P_{j-1} + \alpha_j P_j$$

avec

$$\alpha_j = \frac{t - t_j}{t_{j+n} - t_j}$$

Cette dernière relation est une conséquence de la formulation par récurrence des B-splines, et assure la non-déformation de la courbe initiale.

Le vecteur de nœuds a alors été modifié, puisque t est devenu le nœud numéro $i+1$, les numéros de tous les nœuds au-delà de $i+1$ ayant été incrémentés de 1. Le processus est alors itéré: t va être à nouveau inséré dans le vecteur de nœuds. Comme t est confondu avec t_{i+1} , seuls les n points $P_{i-n+1}^{(1)}$ à $P_i^{(1)}$ interviennent dans le calcul de $c(t)$ (en effet, $B_{n,i+1}(t_{i+1}) = 0$). Les points $P_j^{(2)}$ sont ensuite construits ($n+1$ en tout). La multiplicité de t est alors 2. L'insertion suivante ne fera donc intervenir que les $n-1$ points $P_{i-n+2}^{(2)}$ à $P_i^{(2)}$.

Il apparaît donc que le nombre de points intervenant dans le calcul se réduit d'une unité à chaque nouvelle insertion. La $n^{\text{ème}}$ insertion ne conduit à la création que d'un seul point: c'est le point recherché. Les différentes étapes, dans le cas $n=3$ sont présentées sur la figure II.8.

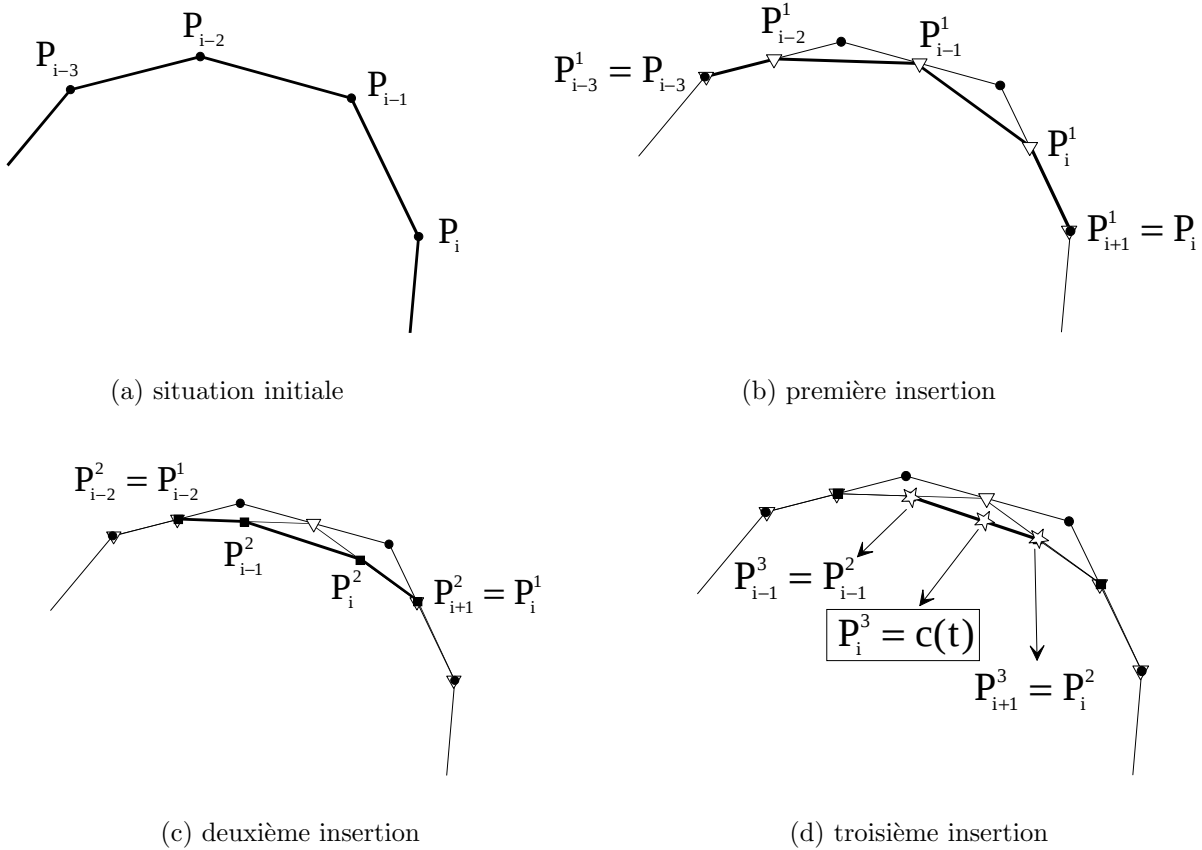


Figure II.8 : Calcul d'un point sur une courbe

CHOIX DE LA MÉTHODE D'ÉVALUATION

Le nombre N_1 d'opérations mathématiques élémentaires (additions, multiplications et divisions) requis pour évaluer un point sur une courbe gauche à l'aide de l'algorithme de de Boor-Cox est égal à [DAN89]:

$$N_1 = 13 \frac{n(n+1)}{2}$$

La méthode que nous utilisons consiste tout d'abord à calculer les coefficients de toutes les B-splines, puis à ne recombinaison que les coefficients nécessaires pour un point. Le nombre N_2 d'opérations requis pour arriver au résultat s'obtient en considérant chaque étape:

- ◇ pour la création des coefficients des B-splines:
 - la formule de récurrence nécessite 9 opérations;
 - la boucle sur les intervalles comporte $n + 1$ itérations;
 - les deux boucles sur le degré et les points d'accroches comportent ensemble $\frac{n(n+1)}{2}$ itérations;
- ◇ il y a autant de B-splines à créer qu'il y a de points de contrôle³, soit n_{pt} ;
- ◇ enfin, la recombinaison des B-splines pour obtenir la valeur cherchée, ce qui revient à évaluer un polynôme de degré n , demande $2n$ opérations (utilisation de l'algorithme de Hörner), à effectuer une fois pour chacune des trois coordonnées.

Ainsi, au total, cette méthode induit un nombre N_2 d'opérations égal à:

$$N_2 = 9 n_{pt} \frac{n(n+1)^2}{2} + 6n$$

Dans le cas où un seul point doit être évalué, il est clair que $N_2 \gg N_1$, ce qui ne rend pas la méthode "globale" très attrayante. Cependant, la méthode de recherche qui va être développée va faire appel un très grand nombre de fois à l'évaluation d'un point: ce nombre est très variable, mais son ordre de grandeur est de 100 à 200 pour la seule recherche d'intersection (sans même parler du post-traitement). Il faut donc comparer non pas N_1 et N_2 directement, mais le nombre total de calculs à réaliser pour $\mathcal{N}$ évaluations:

- ◇ pour la méthode de de Boor-Cox, il faut réaliser un nombre $N_{1\mathcal{N}}$ d'opérations égal à:

$$N_{1\mathcal{N}} = \mathcal{N}.N_1 = \mathcal{N}.13 \frac{n(n+1)}{2}$$

³il est supposé ici que les valeurs n_{pt} (nombre de points dans le jeu de données) et N (nombre de points de contrôle de la spline) sont égaux: ceci est faux en général dans le cas du lissage, mais vrai dans la programmation que nous avons réalisée, et c'est pourquoi le calcul est mené sous cette hypothèse

- ◇ pour la méthode globale, la création des coefficients n'est réalisée qu'une fois pour tous les points à évaluer, et seule la recombinaison doit être effectuée $\mathcal{N}$ fois, ce qui conduit à un nombre $N_{2\mathcal{N}}$ d'opérations égal à:

$$N_{2\mathcal{N}} = 9npt \frac{n(n+1)^2}{2} + 6n\mathcal{N}$$

L'utilisation de la méthode globale devient plus économique que l'utilisation de la méthode de de Boor-Cox dès que $N_{2\mathcal{N}}$ devient inférieur à $N_{1\mathcal{N}}$, ce qui conduit à:

$$\mathcal{N} > 9npt \frac{(n+1)^2}{13n+1}$$

Dans la pratique, les splines de degré 3 sont fréquemment utilisées, ce qui donne:

$$\mathcal{N} > 3,6npt$$

Le nombre de points définissant une courbe dépassera rarement 10: la méthode globale se révèle donc, d'un point de vue nombre d'opérations à effectuer, quasi systématiquement plus intéressante pour notre problème que la méthode pourtant plus élégante de de Boor-Cox, et c'est pourquoi c'est elle qui a été mise en place.

II.3 Modélisation des surfaces

La donnée du problème est un ensemble de $nptu \times nptv$ points P_{ij} appelés points de contrôle de la surface. La mise en place du paramétrage constitue une étape importante: ce paramétrage sera effectué suivant deux directions privilégiées, u et v . Ainsi, si les Q_{ij} désignent les $N_u * N_v$ points de contrôle obtenus par une des trois méthodes, le nuage de points sera approché par une surface dont l'équation (vectorielle) sera :

$$s(u, v) = \sum_{i=0}^{N_u-1} \sum_{j=0}^{N_v-1} Q_{ij} B_{nu,i}(u) B_{nv,j}(v)$$

Comme le montre cette équation vectorielle, la modélisation d'une surface se ramène à la modélisation "croisée" de plusieurs courbes gauches:

- ◇ celles définies par les points Q_{ij} , j étant constant;
- ◇ celles définies par les points Q_{ij} , i étant constant.

Il est ainsi possible d'utiliser des B-splines de degrés différents dans les deux directions (degrés notés nu et nv).

II.3.1 Méthode directe

Deux paramétrages sont nécessaires pour les surfaces, pour chacune des directions: dans les deux cas, un paramétrage uniforme est retenu (Pi1). Le choix de prendre des valeurs de paramètres égales en début et en fin de paramétrage assure le passage de la surface par les quatre points de contrôle qui définissent les angles du jeu de données (soit les points P_{00} , $P_{nptu-1,0}$, $P_{0,nptv-1}$ et $P_{nptu-1,nptv-1}$: voir figure II.9).

D'un point de vue informatique, la modélisation des surfaces ne nécessite que la génération puis la combinaison de modélisations de différentes courbes. Un soin tout particulier doit être apporté à la gestion des différents paramétrages et ensembles de points. Les grandeurs requises et les résultats attendus sont regroupés dans le tableau II.5.

Entrées	Sorties
$nptu$ points de contrôle degré nu	paramétrage suivant u et v
$nptv$ points de contrôle degré nv	pour chaque pavé: $(nu + 1) + (nv + 1)$ coefficients

Tableau II.5: Entrées / sorties du programme

Contrairement au cas des courbes gauches, où le calcul effectif des coefficients du polynôme par intervalle réduit les temps de calcul des coordonnées d'un point (voir le paragraphe II.2.4.3), il est préférable ici de conserver séparément les coefficients des polynômes

suivant chacune des directions u et v , sans chercher à les regrouper dans ce qui deviendrait un polynôme à deux variables: le calcul des coordonnées d'un point à partir de la donnée d'un couple (u, v) ne s'en trouverait pas simplifié. Ce calcul s'effectue alors d'après la définition:

$$s(u, v) = \sum_{i=0}^{nptu-1} \sum_{j=0}^{nptv-1} P_{ij} B_{nu,i}(u) B_{nv,j}(v)$$

Cependant, il est possible de limiter les indices des sommes mises en jeu, car les B-splines ne sont non-nulles que sur certains intervalles. Ainsi, si $u \in [u_{i_0}, u_{i_0+1}]$ et si $v \in [v_{j_0}, v_{j_0+1}]$ alors⁴:

$$s(u, v) = \sum_{i=i_0-nu}^{i_0} \sum_{j=j_0-nv}^{j_0} P_{ij} B_{nu,i}(u) B_{nv,j}(v)$$

La traduction algorithmique de cette équation est:

```

Rechercher  $i_0$  et  $j_0$ 
Pour chaque coordonnée coord (soit  $x$ ,  $y$  ou  $z$ )
  Pour  $i = i_0 - nu, \dots, i_0$ 
    Pour  $j = j_0 - nv, \dots, j_0$ 
      Calculer la quantité  $s_{ij} = coord_{ij} \cdot B_{nu,i}(u) \cdot B_{nv,j}(v)$ 
    Fin pour
  Fin pour
  Additionner tous les  $s_{ij}$ 
Fin pour

```

Un exemple est présenté sur la figure II.II.3.3.2.

II.3.2 Interpolation

Il s'agit ici de déterminer $nptu * nptv$ points de contrôle Q_{ij} de façon à ce que la surface spline passe exactement par les $nptu * nptv$ points donnés P_{ij} , ce qui revient à écrire les $nptu * nptv$ relations suivantes:

$$\forall i_0 \in [0..nptu-1] \quad \forall j_0 \in [0..nptv-1] \quad P_{i_0 j_0} = \sum_{i=0}^{nptu-1} \sum_{j=0}^{nptv-1} Q_{ij} B_{nu,i}(\xi_{i_0}^u) B_{nv,j}(\xi_{j_0}^v)$$

où les ξ_{i_0} et ξ_{j_0} sont des valeurs particulières du paramétrage qu'il faudra choisir *avec pertinence*.

⁴noter que le calcul d'un point de paramètre $u < u_{nu}$ n'est pas utile, car tous les paramètres compris entre u_0 et u_{nu} sont égaux; la quantité $i_0 - nu$ n'a donc aucune raison d'être négative. Il en est de même pour le paramètre v .

Ainsi, deux paramétrages sont nécessaires, à chaque fois dans les deux directions: le premier est attaché aux points P_{ij} donnés (paramétrage noté (ξ_i^u) dans la direction u et (ξ_i^v) dans la direction v) et le second lié aux points de contrôle Q_{ij} cherchés (paramétrage noté (t_i^u) dans la direction u et (t_i^v) dans la direction v). C'est à partir de ce dernier que sont construites les B-splines.

Le choix de ces paramètres n'est pas indépendant de la qualité du résultat ([DAN89]). En fait, pour construire un "bon" paramétrage, il est plus simple de raisonner *à l'envers*, c'est-à-dire de se donner de "bonnes" valeurs ξ_i^u et ξ_i^v , qui respectent, par exemple, la densité de répartition des points dans la direction correspondante (paramétrages selon le schéma Pi2), puis de construire les paramétrages (t_i^u) et (t_i^v) à partir de ces valeurs-là (paramétrages selon le schéma Pr2).

II.3.2.1 Calcul des points de contrôle

L'expression générale de la surface est:

$$s(u, v) = \sum_{i=0}^{nptu-1} \sum_{j=0}^{nptv-1} Q_{ij} B_{nu,i}(u) B_{nv,j}(v)$$

Il s'agit donc de trouver les valeurs des Q_{ij} telles que:

$$\forall i_0 \in [0..nptu - 1] \text{ et } \forall j_0 \in [0..nptv - 1] \quad \sum_{i=0}^{nptu-1} \sum_{j=0}^{nptv-1} Q_{ij} B_{nu,i}(\xi_{i_0}^u) B_{nv,j}(\xi_{j_0}^v) = P_{i_0 j_0}$$

Cette équation induit la résolution d'un système linéaire de taille $(nptu * nptv) * (nptu * nptv)$. Afin de réduire les temps de calcul, ce système peut être abordé suivant la démarche décrite au paragraphe II.3.4.

II.3.2.2 Programmation

Les entrées et sorties du calcul sont regroupées dans le tableau II.6.

Entrées	Sorties
$nptu * nptv$ points de contrôle degré nu et nv	paramétrages ξ^u, ξ^v, t^u et t^v pour chaque pavé: $(nu + 1) + (nv + 1)$ coefficients

Tableau II.6: Entrées / sorties du programme

Le programme s'organise de la façon suivante:

```

Pour  $i_0 = 0, nptu - 1$ 
  Pour  $j = 0, nptv - 1$ 
     $b_j = P_{i_0j}$ 
    Pour  $k = 0, nptv - 1$ 
      Calculer  $a_{jk} = B_{nv,k}(\xi_j^v)$ 
    Fin pour
  Fin pour
Inverser  $\mathbb{A}$ 
Calculer  $m_{i_0}^L = \mathbb{A}^{-1} * b$ 
  remarque: à ce stade, les vecteurs  $m_{i_0}^L$ 
  sont les vecteurs ligne de la matrice  $\mathbb{M}$ 
Fin pour
Pour  $j_0 = 0, nptv - 1$ 
  démarche analogue au cas ci-dessus
  en résolvant les systèmes  $\mathbb{A}\mathbb{Q} = \mathbb{M}$ 
  avec cette fois:  $a_{ik} = B_{nu,k}(\xi_j^u)$ 
Fin pour

```

Un exemple est présenté sur la figure II.II.3.3.2.

II.3.3 Lissage

La mise en place de l'algorithme de calcul des surfaces dans le cas du lissage est délicate, non d'un point de vue théorique, mais d'un point de vue pratique, l'ensemble des paramètres à gérer étant important. Les choix présentés ici sont, encore une fois, issus des résultats obtenus par M. Daniel [DAN89].

Comme dans le cas de l'interpolation, les $N_u * N_v$ points de contrôle utilisés ne sont plus les points donnés, mais doivent être déterminés et il n'y a plus de lien a priori entre N_u et N_v d'une part, et $nptu$ et $nptv$ de l'autre. C'est pourquoi, comme précédemment, deux paramétrages sont nécessaires.

Le paramétrage pour les $nptu * nptv$ points P_{ij} donnés est constitué de deux vecteurs paramètre ξ_i^u et ξ_i^v pour les deux directions de modélisation, construits suivant le schéma Pi2. Le paramétrage pour les $N_u * N_v$ points Q_{ij} recherchés est également constitué de deux vecteurs paramètre t_i^u et t_i^v pour les deux directions, construits selon le schéma Pr1. C'est sur le paramétrages t_i^u et t_i^v que sont construites les B-splines. Rappelons que la définition des paramètres a conduit notamment aux conditions suivantes (dues à la définition des α):

$$\diamond N_u \neq nu + 2$$

$$\diamond N_v \neq nv + 2$$

Si les cas d'égalité devaient être rencontrés, il a été choisi de travailler automatiquement avec $N'_u = N_u + 1$ et $N'_v = N_v + 1$.

II.3.3.1 Calcul des points de contrôle

La démarche théorique est détaillée pour les courbes au paragraphe II.2.3.1, et ne sera donc pas présentée ici.

L'expression générale de la surface est:

$$s(u, v) = \sum_{i=0}^{N_u-1} \sum_{j=0}^{N_v-1} Q_{ij} B_{nu,i}(u) B_{nv,j}(v)$$

Après projection sur les fonctions B-splines B_{nu,i_0} et B_{nv,j_0} , il vient:

$$\begin{aligned} \int_{\xi_0^u}^{\xi_{nptu-1}^u} \int_{\xi_0^v}^{\xi_{nptv-1}^v} s(u, v) B_{nu,i_0}(u) B_{nv,j_0}(v) du dv = \\ \sum_{i=0}^{N_u-1} \sum_{j=0}^{N_v-1} Q_{ij} \langle B_{nu,i}, B_{nu,i_0} \rangle_u \langle B_{nv,j}, B_{nv,j_0} \rangle_v \end{aligned}$$

avec:

$$\langle f, g \rangle_u = \int_{t_0^u}^{t_{N_u+nu}^u} f(u).g(u) du$$

et

$$\langle f, g \rangle_v = \int_{t_0^v}^{t_{N_v+nv}^v} f(v).g(v) dv$$

En introduisant $I_{i_0j_0}$ tel que⁵:

$$I_{i_0j_0} = \int_{\xi_0^u}^{\xi_{nptu-1}^u} \int_{\xi_0^v}^{\xi_{nptv-1}^v} s(u, v) B_{nu,i_0}(u) B_{nv,j_0}(v) du dv$$

le système s'écrit alors:

$$\forall i_0 \in [0..N_u - 1] \text{ et } \forall j_0 \in [0..N_v - 1] \quad \sum_{i=0}^{N_u-1} \sum_{j=0}^{N_v-1} Q_{ij} \langle B_{nu,i}, B_{nu,i_0} \rangle_u \langle B_{nv,j}, B_{nv,j_0} \rangle_v$$

Les points Q_{ij} sont obtenus en inversant un système de taille $N_u N_v * N_u N_v$. Comme dans le paragraphe précédent, ce système est décomposé en sous-systèmes de tailles inférieures, suivant la démarche exposée dans le paragraphe II.3.4.

Remarque: il n'y a mathématiquement aucune raison pour que les points de contrôle extrêmes obtenus par le calcul (à savoir les Q_{0j} et les Q_{i0}) coïncident avec les points extrêmes initiaux (c'est-à-dire les P_{0j} et les P_{i0}). Néanmoins, ces points sont, dans la pratique, et pour une extrémité donnée, souvent très voisins les uns des l'autre. C'est pourquoi il est fréquent, une fois les calculs faits (interpolation ou lissage), de remplacer les points extrêmes obtenus par les points extrêmes initiaux. Il en résulte bien sûr une légère modification de la surface (la surface interpolée notamment n'en est plus une !) ce qui nous a conduit à ne pas utiliser cette méthode.

⁵les termes $I_{i_0j_0}$ se calculent à l'aide d'une intégration numérique utilisant les points de Gauss, la surface étant alors approchée localement par un plan

II.3.3.2 Programmation

Les entrées et sorties du calcul sont regroupées dans le tableau II.7.

Entrées	Sorties
$nptu * nptv$ points de contrôle degré nu et nv	paramétrages ξ^u, ξ^v, t^u et t^v
$N_u * N_v$ points de contrôle recherchés	pour chaque pavé: $(N_u + 1) + (N_v + 1)$ coefficients

Tableau II.7: Entrées / sorties du programme

Le programme s'organise de la façon suivante:

```

Pour  $i_0 = 0, N_u - 1$  et  $j_0 = 0, N_v - 1$ 
  Calculer  $I_{i_0 j_0}$ 
Fin pour
Pour  $i_0 = 0, N_u - 1$ 
  Pour  $j = 0, N_v - 1$ 
     $b_j = I_{i_0 j}$ 
    Pour  $k = 0, N_v - 1$ 
      Calculer  $a_{jk} = \langle B_{nv,j}, B_{nv,k} \rangle_v$ 
    Fin pour
  Fin pour
  Inverser  $\mathbb{A}$ 
  Calculer  $m_{i_0}^L = \mathbb{A}^{-1} * b$ 
  remarque: à ce stade, les vecteurs  $m_{i_0}^L$ 
  sont les vecteurs ligne de la matrice  $\mathbb{M}$ 
Fin pour
Pour  $j_0 = 0, N_v - 1$ 
  démarche analogue au cas ci-dessus
  en résolvant les systèmes  $\mathbb{A}\mathbb{Q} = \mathbb{M}$ 
  avec cette fois:  $a_{ik} = \langle B_{nu,i}, B_{nu,k} \rangle_u$ 
Fin pour

```

Un exemple de calcul par les trois méthodes proposées est présenté sur les figures II.9. Le jeu de données est représenté sur la figure II.II.3.3.2 (les points de ce jeu de données apparaissent explicitement sur les figures II.II.3.3.2, II.II.3.3.2 et II.II.3.3.2).

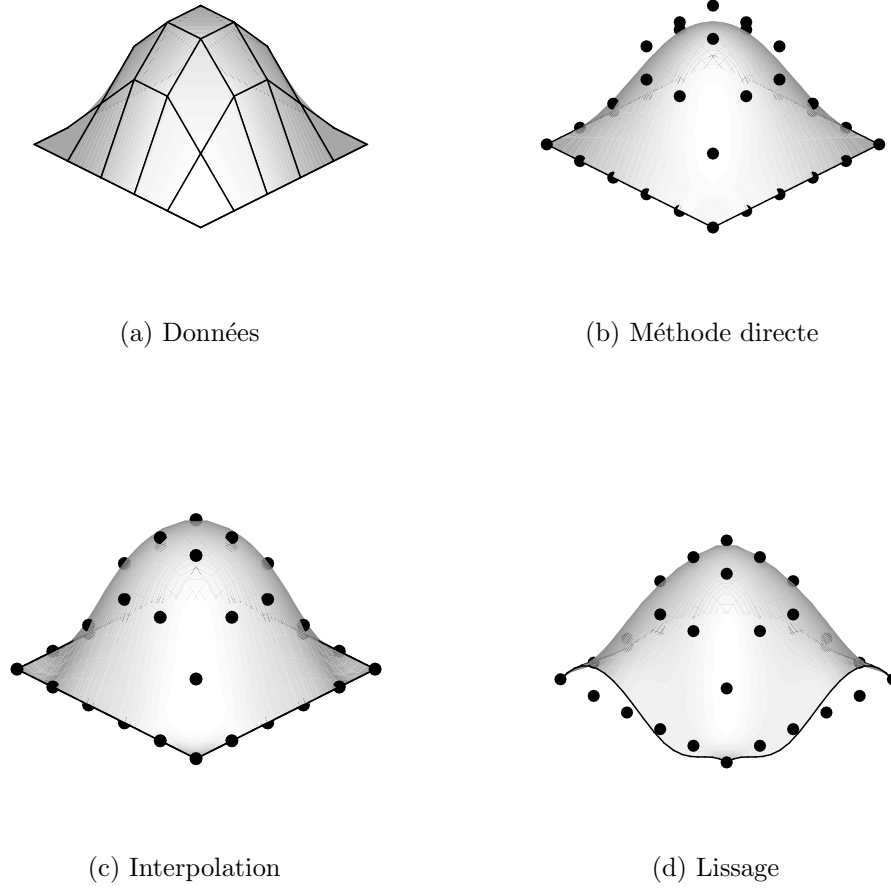


Figure II.9 : Comparaison méthode directe - interpolation - lissage

II.3.4 Résolution des systèmes

Dans le cas de l'interpolation comme dans le cas du lissage, il s'agit de trouver des valeurs Q_{ij} telles que:

$$\forall i_0 \in [0..nptu - 1] \text{ et } \forall j_0 \in [0..nptv - 1] \quad \sum_{i=0}^{nptu-1} \sum_{j=0}^{nptv-1} Q_{ij} B_{nu,i}(\xi_{i_0}^u) B_{nv,j}(\xi_{j_0}^v) = I_{i_0 j_0}$$

Par une résolution directe, cette équation induit la résolution d'un système linéaire de taille $(nptu * nptv) * (nptu * nptv)$. Afin de réduire les temps de calcul, ce système est décomposé en introduisant les quantités $M_{i_0 j_0}$ définies de la manière suivante: .

$$M_{i_0 j_0} = \sum_{i=0}^{nptu-1} Q_{i_0 j_0} B_{nu,i}(\xi_{i_0}^u)$$

Le système à résoudre se décompose ainsi en deux systèmes de tailles inférieures qui s'écrivent:

$$\forall i_0 \in [0..nptu - 1] \quad I_{i_0j_0} = \sum_{j=0}^{nptv-1} M_{i_0j_0} B_{nv,j}(\xi_{j_0}^v) \quad \text{pour } j_0 \in [0..nptv - 1]$$

$$\forall j_0 \in [0..nptv - 1] \quad M_{i_0j_0} = \sum_{i=0}^{nptu-1} Q_{i_0j_0} B_{nu,i}(\xi_{i_0}^u) \quad \text{pour } i_0 \in [0..nptu - 1]$$

La première équation induit la résolution de $nptu$ systèmes de taille $nptv * nptv$ qui permettent d'accéder aux points $M_{i_0j_0}$, et la seconde, la résolution de $nptv$ systèmes de taille $nptu * nptu$ qui permettent d'accéder aux points $Q_{i_0j_0}$.

La résolution par une méthode directe d'un système matriciel de dimension $N * N$ requiert un nombre d'opérations de l'ordre de αN^3 , où α est un réel qui dépend de la méthode utilisée (par exemple, $\alpha = 2/3$ pour la méthode de Gauss, $\alpha = 1/3$ pour la méthode de Cholesky). Dans le cas présent, en supposant $nptu = nptv = 10$, une résolution directe du problème conduirait à $\alpha.10^6$ opérations, alors que la décomposition en série de sous-systèmes conduit à $2\alpha.10^4$ opérations seulement.

II.3.5 Prise en compte de l'épaisseur

La démarche proposée ici étant destinée *in fine* à être intégrée dans un code éléments finis, il convient de traiter le problème lié à l'épaisseur des surfaces. En effet, les points qui serviront de support à la définition des géométries seront les nœuds du maillage, c'est-à-dire les points situés sur la surface moyenne des éléments. Si aucune précaution n'est prise, des situations telles que celle présentée sur la figure II.10 peuvent arriver: l'algorithme de recherche de contact ne détectera aucune pénétration, alors qu'il y aura effectivement des nœuds de la courbe présents dans la demi-épaisseur de la surface.

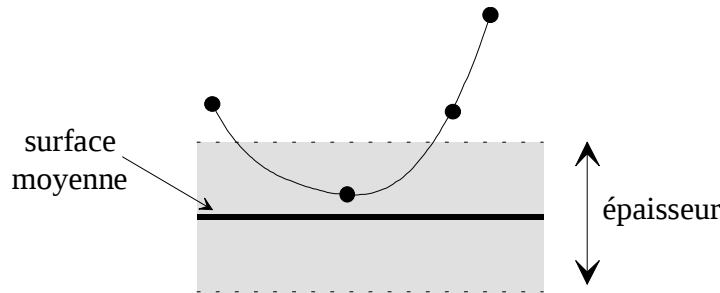


Figure II.10 : Mauvaise détection du contact

C'est pourquoi la recherche d'intersection sera menée sur une surface déduite de la surface moyenne par translation suivant la normale en chaque nœud, d'une quantité égale

à la valeur moyenne des épaisseurs des éléments qui entourent chacun des nœuds. Pour ce faire, deux méthodes peuvent être utilisées.

MÉTHODE CLASSIQUE

La normale de chaque facette est tout d'abord calculée classiquement en effectuant le produit vectoriel des diagonales (le vecteur normal ainsi construit correspond au vecteur normal au plan médian de la facette); chaque nœud se voit ensuite attribué la moyenne des normales des facettes qui l'entourent. Pour que les vecteurs normaux soient tous orientés dans le même sens, il est nécessaire de fixer un point de référence qui détermine les côtés intérieur et extérieur de la surface. Ce sont finalement les points translatés qui servent à la définition de la surface spline (voir figure II.11). Il convient de noter que cette démarche fournit finalement une surface de travail qui se trouve à une distance *inférieure* (ou au plus égale) à la demi-épaisseur des éléments (voir figure II.12).

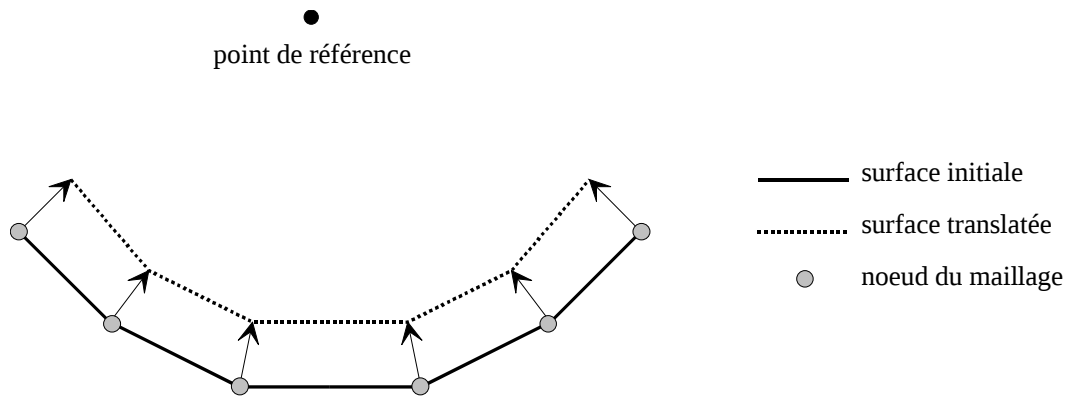


Figure II.11 : Prise en compte de l'épaisseur (méthode classique)

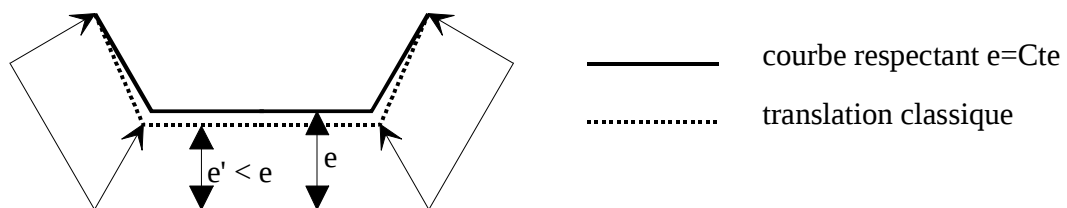


Figure II.12 : Epaisseur réellement prise en compte

UTILISATION DES SPLINES

Pour cette méthode, une première étape consiste à tout d'abord construire la surface spline qui s'appuie directement sur les $n_{ptu} * n_{ptv}$ nœuds du maillage. A partir de cette surface spline, de nouveaux points sont extraits: pour cela, n_{ptu} valeurs du paramétrage dans la direction u et n_{ptv} valeurs dans la direction v sont utilisées (ces

valeurs sont choisies, par exemple, en respectant les longueurs de corde des points originaux). Pour chacun de ces nouveaux points, il est possible, en utilisant la formulation spline, de déterminer le vecteur normal à la surface, orienté grâce à la donnée d'un point extérieur (voir figure II.11). Chaque point peut donc être translaté d'une quantité égale à la demi-épaisseur dans la direction normale: l'ensemble des points translatés sert de jeu de données pour la définition de la surface de travail pour la recherche d'intersection (voir figure II.13).

Cette méthode est plus lourde à mettre en œuvre qu'une méthode classique, et est difficile à appliquer si l'épaisseur des éléments est variable, car les points extraits de la surface initiale ne sont a priori pas confondus avec les nœuds de la surface moyenne⁶.

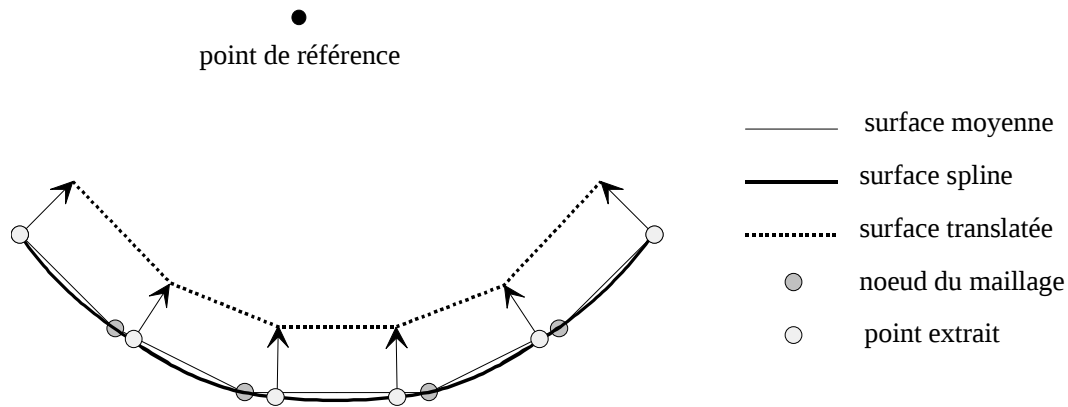


Figure II.13 : Prise en compte de l'épaisseur (méthode spline)

⁶en première approximation, il paraît néanmoins raisonnable de prendre comme épaisseur moyenne en un point la moyenne des épaisseurs des éléments attachés au nœud d'origine de la surface moyenne

Chapitre III

Recherche d'intersection courbe/surface

Dans ce chapitre, les développements réalisés en termes de détection et de traitement du contact sont présentés. Après avoir rappelé quelques méthodes couramment utilisées (paragraphe III.1.1 et III.1.2), la méthode proposée est développée dans les paragraphes III.1.3 (principe) et III.1.4 (détail et points délicats).

Cette méthode est destinée à être implémentée dans des codes éléments finis. Il est donc indispensable de bien préciser non seulement les grandeurs à fournir par le code pour la bonne mise en œuvre de l'algorithme, mais également celles qui doivent être ensuite retournées par l'algorithme pour que le code puisse poursuivre les calculs dans de bonnes conditions. Ces grandeurs (pénétration, vecteur normal, coordonnées barycentriques...) sont précisées dans le paragraphe III.2.

III.1 Recherche d'intersection courbe/surface

Disposant des modélisations d'une courbe gauche et d'une surface, il reste à en définir les possibles points communs.

III.1.1 Méthodes couramment utilisées en éléments finis

Les procédés actuels de recherche d'intersection implémentés dans les codes éléments finis utilisent, pour beaucoup, les deux grandes méthodes suivantes: les approches hiérarchiques et les approches par voisinage ([CES96]). Pour ces deux approches, il est fréquent de trouver une classification des éléments finis mis en jeu en:

- ◇ éléments maîtres;
- ◇ éléments esclaves.

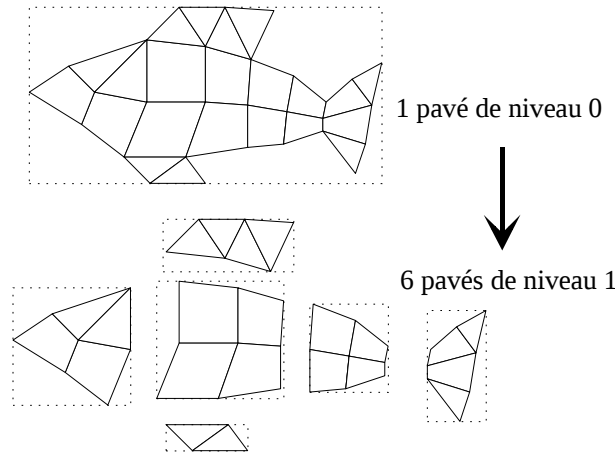
Cette classification, que l'utilisateur fait apparaître dans le jeu de données de son problème, est purement arbitraire et n'a pour but que de désigner les éléments finis qui seront utilisés pour *mener la recherche* dans l'algorithme de détection du contact. Ce dernier est en effet constitué d'une boucle sur l'ensemble ou partie des éléments esclaves: pour chacun d'entre eux, un test est réalisé sur l'ensemble ou partie des éléments maîtres afin de déterminer un potentiel appariement. A la suite de cette première phase, si des couples maître/esclave ont satisfait le premier critère, un second test est réalisé pour confirmer ou invalider précisément l'existence du contact.

III.1.1.1 Approches hiérarchiques

La démarche est la suivante: l'entité maître est entièrement contenue dans un parallélépipède $\mathcal{P}^0$ dit de niveau zéro et repéré par les coordonnées extrêmes de ses sommets, soit $(x_{min}^0, y_{min}^0, z_{min}^0)$ et $(x_{max}^0, y_{max}^0, z_{max}^0)$ ¹. Si le nœud esclave examiné n'appartient pas à $\mathcal{P}^0$, il ne peut y avoir contact; sinon, le contact est possible mais non certain, et il faut procéder à un examen plus fin.

Pour cela, l'ensemble des éléments de $\mathcal{P}^0$ est partitionné en un certain nombre de sous-ensembles disjoints de tailles équivalentes, la sélection pouvant s'opérer par exemple sur la position des centres de gravité des éléments maîtres les uns par rapport aux autres. Ainsi, un nombre N de parallélépipèdes $(\mathcal{P}_i^1)_{i=1,N}$ dits de niveau un sont définis (voir figure III.1). Si le nœud esclave n'appartient à aucun de ces pavés, l'algorithme s'arrête: le nœud testé n'est pas en contact avec l'entité maître. Si au contraire, le nœud esclave appartient à l'un des N pavés, le processus est itéré sur celui-ci (découpage en N nouveaux sous-ensembles de niveau deux...). Le processus s'arrête selon un critère qui peut porter, par exemple, sur la taille des pavés créés, ou bien sur le nombre d'éléments contenus dans chaque pavé.

¹s'il y a lieu, il est prudent d'augmenter la taille des pavés dans chacune des directions d'une quantité fonction de la caractéristique du type d'éléments utilisés, comme par exemple la demi-épaisseur dans le cas de coques

Figure III.1 : Approche hiérarchique (exemple pour $N = 6$)

A ce stade, si le nœud esclave appartient à l'un des pavés restants, un test définitif est réalisé (calcul de distances, calcul d'angles solides...).

Cette méthode présente l'avantage d'être rapide (pour un pas de temps donné): il est fréquent de trouver $N = 2^{dim}$, où dim est la dimension de l'espace de travail. Ainsi, dans l'espace, le pavé de niveau k se voit partitionné en 8 pavés de niveau $k + 1$, ce qui implique qu'un maillage de 250000 éléments est entièrement balayé en 6 étapes au plus. Cependant, si l'entité maître n'est pas rigide, il est nécessaire de réactualiser la définition des pavés à chaque pas de temps, ce qui est finalement coûteux d'un point de vue temps de calcul.

III.1.1.2 Approches par voisinage

Ces méthodes reposent sur le fait que lors d'un calcul dynamique dans lequel apparaît du contact, les pas de temps sont en général très petits. Cela signifie que l'évolution des deux structures en contact entre deux pas de temps est faible. Il en résulte que si un nœud esclave se trouvait en contact avec un élément maître à l'instant t (voir figure III.III.1.1.2), il est fort probable que ce nœud esclave soit encore dans le voisinage de cet élément maître à l'instant $t + \Delta t$ (éléments en grisé sur la figure III.III.1.1.2). Cette hypothèse permet de ne tester qu'un nombre limité d'éléments à chaque pas de temps. Néanmoins, ces méthodes sont inefficaces seules, car elles ne sont utilisables qu'à partir de l'instant où un premier contact a été détecté.

D'un point de vue temps de calcul, les approches par voisinage restent malgré tout intéressantes. C'est pourquoi elles sont souvent utilisées, couplées à une méthode de type approche hiérarchique pour traiter les premiers pas de temps, tant que le contact n'a pas eu lieu. C'est une association de ce type qui figure dans des codes tels que DYNA3D [HAL95], ou PLEXUS [CEA97] (codes explicites dédiés à la dynamique rapide).

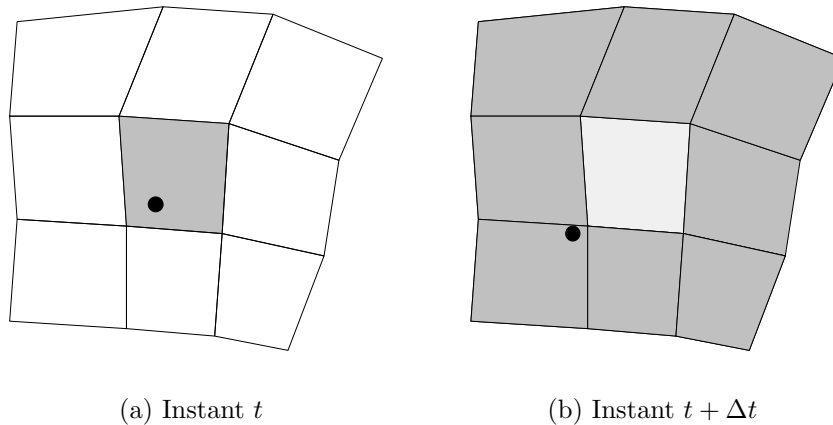


Figure III.2 : Approche par voisinage

III.1.2 Principe des méthodes analytiques

De façon analogue aux approches hiérarchiques présentées ci-dessus, les méthodes de recherche d'intersection entre une courbe et une surface définies analytiquement font appel au principe général suivant, connu sous le nom anglais de “divide and conquer” (c'est-à-dire diviser pour résoudre) [DAN89] [MAL98]: les deux entités à étudier sont englobées dans des boîtes (pavés de l'espace) pour lesquelles il est facile de voir s'il y a intersection ou non. Ces boîtes sont ensuite subdivisées de façon récursive, certaines étant conservées et d'autres éliminées, jusqu'à ce que, un critère de planéité ayant été atteint, le problème puisse être assimilé à un problème (ou groupe de problèmes) dit de type “élémentaire” (intersection droite/plan).

III.1.3 Méthode proposée

Avant même de développer la méthode de balayage mise en place, il convient de rappeler que l'objectif de l'étude n'est pas de transformer les codes éléments finis qui seront utilisés ultérieurement en puissants outils de design et de conception. Il s'agit bien de modifier localement le comportement d'un algorithme de façon à améliorer la fiabilité des simulations qu'il réalise.

L'algorithme a été créé en prenant pour hypothèse que la courbe et la surface ne sont pas supposées être quelconques, mais présenter quelques particularités notamment du point de vue de leur forme. La surface est une modélisation d'un secteur de carter: elle aura donc globalement (même après déformation) la forme d'une partie de cylindre. La courbe est une modélisation du sommet d'une aube qui peut être approchée par une parabole. Ainsi, ces entités ne présentent-elles *a priori* pas de singularités (trous, points multiples, etc), qui nécessiteraient des traitements particuliers.

Une recherche d'intersection n'a de sens que pour une précision donnée: le résultat

de la recherche se présentera donc sous la forme “à l’intérieur d’une boule de rayon ε , il existe au moins deux points appartenant l’un à la surface et l’autre à la courbe”.

La démarche suivante est adoptée: chacun des paramètres mis en jeu (soit t pour la courbe, et u et v pour la surface) est discrétisé uniformément (l’obtention des incréments dt , du et dv sera détaillée plus loin). Par la suite, P_t désignera le point de la courbe gauche associé au paramètre t , et P_{uv} le point de la surface associé au couple (u, v) .

La courbe gauche sert de support à la recherche (voir figure III.III.1.3). Pour chaque valeur du paramètre t , la surface est balayée, et pour chaque couple (u, v) , la distance $P_t P_{uv}$ est calculée. Si cette distance devient inférieure à une précision ε_0 , les paramètres sont conservés en mémoire (paramètres dits d’initiation). Les valeurs des paramètres sont ensuite incrémentées. Dès que la distance redevient supérieure à ε_0 , les paramètres sont à nouveau gardés en mémoire (paramètres dits de terminaison). Un certain nombre de zones potentielles d’intersection entre les deux entités géométriques sont ainsi définies.

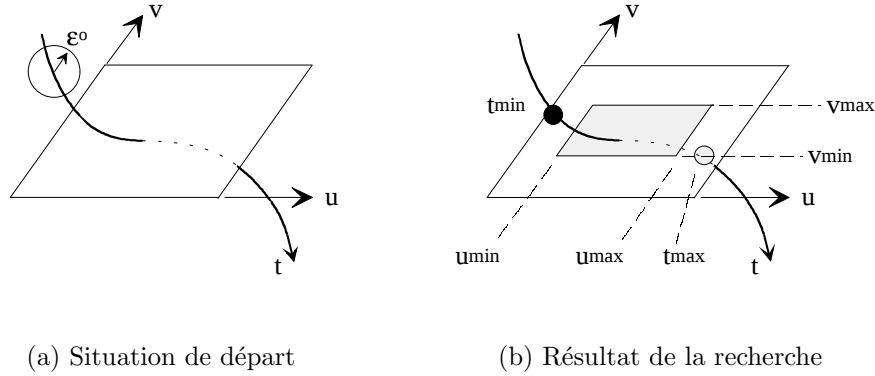


Figure III.3 : Principe de la recherche d'intersection

Sur l’exemple de la figure III.III.1.3, une première recherche conduira à la détermination d’une zone de contact possible, dont les limites sont présentées sur la figure III.III.1.3.

Après avoir terminé un balayage, un autre balayage est effectué avec une précision $\varepsilon_1 < \varepsilon_0$. Cet autre balayage ne porte que sur les éventuelles zones repérées lors du balayage précédent. L’opération est répétée jusqu’à l’obtention de la précision voulue.

Il reste, pour pouvoir initier la recherche, à définir les bornes de la première zone à explorer. Une solution simple consiste à prendre les zones les plus grandes possibles, c’est-à-dire à prendre pour bornes les valeurs extrêmes des paramètres. Cette méthode peut s’avérer coûteuse dans les cas où de grandes surfaces (par rapport à la précision) sont étudiées. Une seconde solution peut être envisagée qui utilise l’interprétation géométrique du théorème de projection des B-splines évoquée page 17: pour un paramètre t , le point P_t d’une courbe peut être considéré comme le barycentre de l’ensemble des points P_i affectés des coefficients $B_{ni}(t)$. En fait, il est possible de préciser davantage quels sont les points P_i qui entrent réellement en jeu dans cette approche géométrique: en effet, si $t \in [t_{i_0}, t_{i_0+1}]$, alors seules les B-splines B_{ni} pour $i \in [i_0 - n, i_0]$ sont non-nulles. Donc seuls les points

P_i pour $i \in [i_0 - n, i_0]$ sont à prendre en compte. Ainsi, le point P_t pour $t \in [t_{i_0}, t_{i_0+1}]$ est en fait le barycentre des points $P_{i_0-n}, P_{i_0-n+1}, \dots, P_{i_0}$ et, compte-tenu de la positivité des B-splines, se trouve à l'intérieur de leur enveloppe convexe. Cette caractéristique est illustrée sur la figure III.4 (exemple pour $n = 3$).

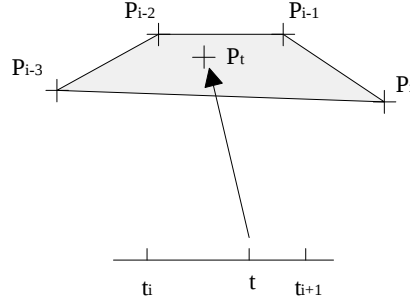


Figure III.4 : Position d'un point par rapport aux points de contrôle

Il est donc possible, pour chaque intervalle paramétré $[t_i, t_{i+1}]$ de construire une boîte de type min-max qui englobe les $n + 1$ points $P_{i-n} \dots P_i$ et à l'intérieur de laquelle se trouvent tous les points P_t pour $t \in [t_i, t_{i+1}]$. Une démarche analogue peut s'appliquer à la surface. Le principe de détermination de la première zone est alors le suivant²:

```

Pour chaque intervalle  $[t_i, t_{i+1}]$ 
  Définir la boîte min-max  $C_i$  de la courbe
  Pour chaque pavé  $[u_j, u_{j+1}] \times [v_k, v_{k+1}]$ 
    Définir la boîte min-max  $S_{jk}$  de la surface
    Chercher l'intersection des boîtes  $C_i$  et  $S_{jk}$ 
    Si l'intersection est non vide,
      conserver en mémoire les paramètres  $i, j, k$ 
    Fin si
  Fin pour
Fin pour

```

Ici, une gestion des éventuels points multiples n'a pas été envisagée. Elle serait beaucoup trop compliquée à ce stade de l'étude. Il en résulte que la zone obtenue à l'issue de ce premier balayage peut très bien ne pas réduire les intervalles de recherche, et donc constituer une étape inutile qui augmente le temps de calcul. En fait, lorsque l'implémentation du module de recherche des points de contact dans un code sera abordée, il sera souhaitable de procéder tout d'abord à une étude min-max avant de générer surface et courbe.

L'algorithme principal de la recherche d'intersection se présente donc sous la forme suivante:

²en fait, le terme de *première zone* est impropre, car le résultat de la recherche peut très bien être constitué de plusieurs zones (cas de l'intersection d'un plan avec une spirale page 66 par exemple)

```

Lire les données (courbe, surface)
Construire les entités géométriques
Lire les précisions (précision initiale  $\varepsilon$  et précision voulue  $\varepsilon_v$ )
Définir la (ou les) zone(s) initiale(s)
Tant que  $\varepsilon > \varepsilon_v$ 
    Effectuer un balayage des zones potentielles d'intersection
    et redéfinir les bornes
    Diminuer  $\varepsilon$ 
Fin tant que
Analyser les résultats

```

L'analyse des résultats consiste principalement en la reprojection, sur le polygone de contrôle, des points d'intersection obtenus. Dans le cas de la courbe, la démarche est la suivante: pour chaque point d'intersection, recherche du projeté orthogonal du point sur les droites définies par deux points consécutifs du polygone de contrôle, et calcul de la distance de projection. Parmi toutes les valeurs de i pour lesquelles la projection se fait à **l'intérieur** du segment P_iP_{i+1} est conservée celle pour laquelle la distance de projection est minimale³. Le détail du calcul de la distance figure en annexe 2.

Il est clair que cette méthode n'a rien de vraiment subtil. Comme toutes les méthodes de type "rouleau compresseur", elle présente l'avantage d'être très efficace (tous les points d'intersection seront trouvés), et l'inconvénient de ne pas toujours être rapide.

III.1.4 Difficultés rencontrées

Les deux principales difficultés qui surgissent lorsqu'il s'agit de mettre en place cet algorithme sont:

1. la gestion de l'exploration d'une zone (problème des bornes et de leur évolution, traitement des points multiples);
2. le calcul de la taille des incréments et les implications sur les critères d'arrêt de l'algorithme de recherche.

III.1.4.1 Exploration d'une zone

L'algorithme utilisé est présenté dans sa version simplifiée ci-après. La version complète, qui figure en annexe D, traite la découverte des zones de contact possibles, la gestion des bornes de ces zones ainsi que les points multiples (cas où à un même t correspond plusieurs zones distinctes (u, v)). Le paramètre ind_{zone} est une variable logique valant 0 en dehors des domaines de contact possibles, et valant 1 l'intérieur de ces domaines.

³la projection orthogonale qui a été choisie ici peut aisément être remplacée par une projection suivant un vecteur donné, comme par exemple un vecteur vitesse

```

Pour chaque zone  $z_i$  déterminée à l'étape précédente
  Lire les bornes en  $t$ ,  $u$  et  $v$  de  $z_i$  et la précision  $\varepsilon$ 
  Calculer les incréments  $dt$ ,  $du$  et  $dv$ 
   $ind_{zone} = 0$ 
  Tant que  $t \leq t_{max}$  faire
    Calculer  $P_t$ 
    Tant que  $u \leq u_{max}$  et  $v \leq v_{max}$  faire
      Calculer  $P_{uv}$ 
      Calculer  $gap = \| P_t P_{uv} \|$ 
      Si  $gap \leq \varepsilon$  et  $ind_{zone} = 0$  alors
         $ind_{zone} = 1$ 
        Mémoriser les valeurs de  $t$ ,  $u$  et  $v$ 
          /paramètres d'initiation/
      Fin si
      Si  $gap \leq \varepsilon$  et  $ind_{zone} = 1$  alors
        Mémoriser les valeurs de  $t$ ,  $u$  et  $v$ 
          /paramètres de terminaison/
      Fin si
    Incréments  $u$  et  $v$ 
  Fin tant que
  Incréments  $t$ 
Fin tant que
Fin pour

```

III.1.4.2 Taille des incréments

La taille des incréments s'est révélée très importante pour la stabilité de la méthode. En effet, des incréments trop grands peuvent conduire à de mauvaises interprétations des situations à traiter. Sur la figure III.III.1.4.2 par exemple, l'algorithme ne va pas trouver d'intersection entre la courbe et la surface car l'incrément dt est trop grand pour le problème. Sur la figure III.III.1.4.2, où la courbe est supposée tangente à la surface (les deux éléments ont été dissociés pour plus de clarté), le résultat de la recherche sera composé de plusieurs zones isolées les unes des autres: en effet, les points associés aux deux paramètres t_{i-1} et u_j sont distants de moins de ε , et l'algorithme va donc conclure qu'ils appartiennent à une zone de contact, par exemple la zone 1. Le point associé au paramètre t_i est quant à lui isolé: aucun point paramétré de la surface ne se trouve à une distance inférieure à ε , et l'algorithme va donc conclure que la zone 1 est close. Enfin, les points associés aux deux paramètres t_{i+1} et u_{j+1} sont distants de moins de ε , et l'algorithme va donc conclure qu'ils initient une nouvelle zone de contact potentiel. Ainsi, du fait du mauvais choix des incréments, l'algorithme va conclure qu'il existe un nombre grand de zones de contact, alors qu'il s'agit en fait d'une seule et même zone.

A l'inverse, des incréments trop petits augmentent inutilement les temps de calcul. La

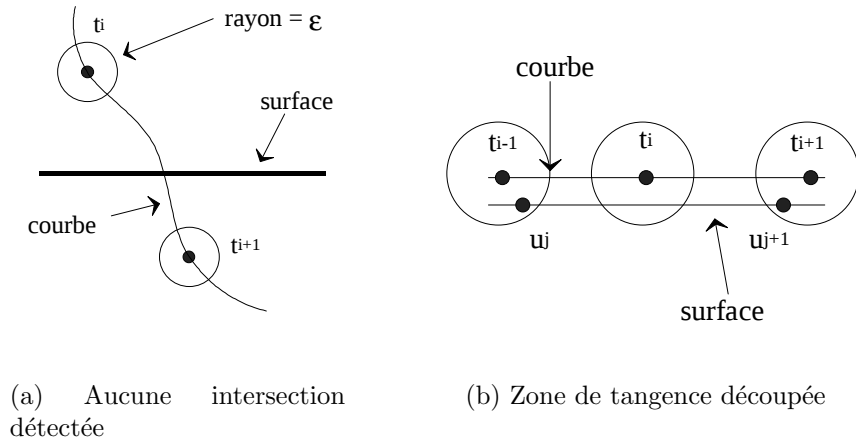


Figure III.5 : Problèmes liés à la taille des incréments

difficulté consiste donc à trouver une relation entre la taille d'un incrément paramétrique et la taille de l'incrément d'espace qui lui est associé, la relation entre les deux n'étant pas linéaire. La démarche adoptée est la suivante.

PRINCIPE DE LA DÉMARCHE

◇ Cas de l'incrément dt

Soit $c(t)$ la définition paramétrique de la courbe gauche. Un développement de Taylor-Young au premier ordre permet d'écrire:

$$c(t + dt) = c(t) + c'(t).dt + o(dt)^2$$

soit, pour dt suffisamment petit, et en supposant que c' ne s'annule pas (ce qui est toujours vrai dans le problème traité, car c' est en fait le vecteur tangent à la courbe):

$$\|dt\| \approx \frac{\|c(t + dt) - c(t)\|}{\|c'(t)\|}$$

La quantité $\|c(t + dt) - c(t)\|$ représente l'accroissement réel de la courbe, c'est-à-dire la distance qui sépare les deux points de paramètre t et $t + dt$. Pour l'instant, cette quantité sera supposée majorée par une grandeur notée $\tilde{\varepsilon}_t$ dont la valeur sera précisée ultérieurement.

Ainsi, le choix de l'incrément dt est conditionné par la relation⁴:

$$\|dt\| \leq \min_{t \in [t_{min}, t_{max}]} \left(\frac{\tilde{\varepsilon}_t}{\|c'(t)\|} \right)$$

⁴le calcul du minimum doit être effectivement réalisé dans tous les cas car il faut rappeler que même si la courbe est de degré 1, la quantité $\|c'(t)\|$ n'est pas constante

◇ **Cas des incréments du et dv**

Soit $s(u, v)$ la définition paramétrique de la surface. Un développement de Taylor-Young au premier ordre donne:

$$s(u + du, v) = s(u, v) + \frac{\partial s(u, v)}{\partial u} \cdot du + o(du)^2$$

Soit, comme précédemment:

$$\|du\| \approx \frac{\|s(u + du, v) - s(u, v)\|}{\left\| \frac{\partial s(u, v)}{\partial u} \right\|}$$

La quantité $s(u + du, v) - s(u, v)$ représente l'accroissement réel de la surface dans la direction u . Cette quantité est majorée par $\tilde{\varepsilon}_u$, dont la valeur sera précisée ultérieurement. Finalement, la majoration obtenue pour du est de la forme:

$$\|du\| \leq \min_{u \in [u_{min}, u_{max}]} \left(\frac{\tilde{\varepsilon}_u}{\left\| \frac{\partial s}{\partial u} \right\|} \right)$$

Le résultat est similaire pour l'expression de dv .

◇ **Remarque**

Ce choix de calcul des paramètres nécessite le calcul de la dérivée de la courbe et de la surface. Pour la courbe, la dérivation est réalisée simplement⁵:

$$c'(t) = \sum_{i=0}^N P_i B'_{ni}(t)$$

Pour la surface, il faut se ramener en fait au calcul de la dérivée dans chacune des deux directions. Ainsi, pour la direction u , il vient:

$$\frac{\partial s(u, v)}{\partial u} = \sum_{i=0}^{nptu-1} \sum_{j=0}^{nptv-1} P_{ij} \frac{\partial B_{nu,i}(u)}{\partial u} B_{nv,j}(v)$$

Le calcul des incréments tel qu'il est présenté jusqu'ici peut en fait être affiné lorsque les degrés de modélisation sont supérieurs ou égaux à 2. En effet, dans ces cas-là, il est possible de prendre en compte la courbure de la spline, puisque la dérivée seconde est non-nulle. C'est l'objet du prochain paragraphe.

⁵on peut montrer que la courbe obtenue par ce calcul de dérivée est encore une courbe spline de degré inférieur appelée *hodographe*

DÉVELOPPEMENTS D'ORDRE SUPÉRIEUR

Le calcul du paragraphe précédent est repris, mais le développement de Taylor est prolongé jusqu'à l'ordre 2.

Pour la courbe, il vient (le nouvel incrément sera noté dt^*):

$$c(t + dt^*) = c(t) + c'(t).dt^* + \frac{c''(t)}{2}.(dt^*)^2 + o(dt^*)^3$$

En négligeant les termes supérieurs à l'ordre 2, cette expression peut être considérée comme une équation du second degré dont l'inconnue est dt^* . La résolution de cette équation conduit à ($\|dt^*\|$ est positif):

$$\|dt^*\| = \frac{-\|c'(t)\| + \sqrt{\|c'(t)\|^2 + 2.\|c(t + dt^*) - c(t)\|.\|c''(t)\|}}{\|c''(t)\|}$$

En utilisant la majoration $\|c(t + dt^*) - c(t)\| \leq \tilde{\varepsilon}_t$, il vient:

$$\|dt^*\| \leq \min_{t \in [t_{min}, t_{max}]} \left(\frac{-\|c'(t)\| + \sqrt{\|c'(t)\|^2 + 2\tilde{\varepsilon}_t.\|c''(t)\|}}{\|c''(t)\|} \right)$$

Des expressions analogues sont obtenues pour $\|du^*\|$ et $\|dv^*\|$. Malgré les calculs supplémentaires qu'ils requièrent, ces évaluations fines se révèlent nécessaires dans le traitement des problèmes.

III.1.4.3 Gestion des incertitudes

Avant de fixer les grandeurs relatives aux incréments, il convient de revenir sur les incertitudes conséquentes à la méthode de recherche. Ces incertitudes vont occasionner des difficultés d'autant plus grandes que la courbe et la surface ne sont pas en contact, mais sont quasiment tangentes, la distance entre les deux devenant voisine de la précision de balayage.

Dans ce cas en effet, et selon la façon dont la recherche a été menée, les deux situations représentées sur les figures III.III.1.4.3 et III.III.1.4.3 peuvent survenir.

Sur la figure III.III.1.4.3 (la courbe est repérée par $\mathcal{C}$ et la surface par $\mathcal{S}$), la disposition des incréments est telle que l'algorithme conclura à l'existence d'une zone de contact potentielle, ce qui est attendu puisque courbe et surface sont effectivement distantes d'une quantité égale à la précision fixée. A l'inverse, sur la figure III.III.1.4.3, l'algorithme conclura à la disjonction de la courbe et de la surface, puisque les points paramétrés sont disposés de telle sorte que la distance minimum obtenue entre les deux entités est supérieure à la précision de balayage. Cette conclusion n'est pas fausse non plus puisque dans ce cas, la courbe et la surface sont effectivement disjointes.

Ce problème pourrait ne pas en être un si les deux situations étaient abordées séparément, les conclusions auxquelles conduit l'algorithme étant acceptables dans les deux cas. Mais

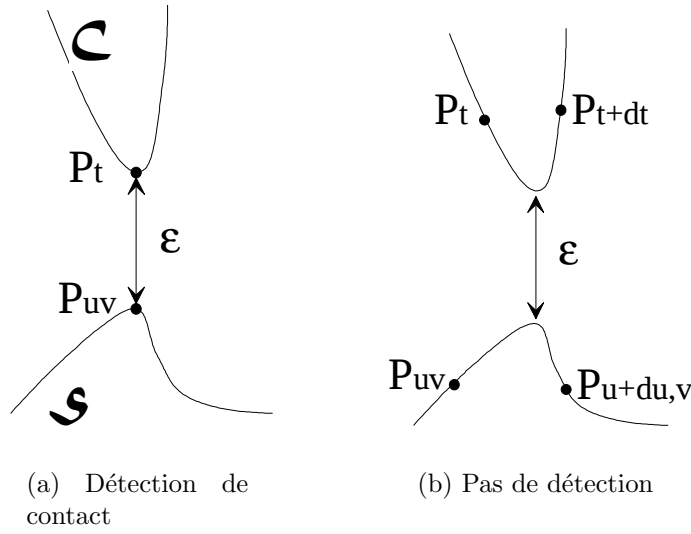


Figure III.6 : Problème lié à la position des points paramétrés

l'objectif de l'ensemble de l'étude est *in fine* d'implémenter la méthode dans un calcul dynamique, c'est-à-dire dans un process qui va gérer des centaines, voire des milliers de situations de ce type les unes à la suite des autres. Il faut pouvoir assurer que pour deux pas de temps successifs, entre lesquels les points de contrôle ne se seront que peu déplacés, le résultat fourni par l'algorithme soit le même, sous peine de voir les efforts de contact, dont l'existence dépend du résultat fourni par l'algorithme, évoluer de manière chaotique.

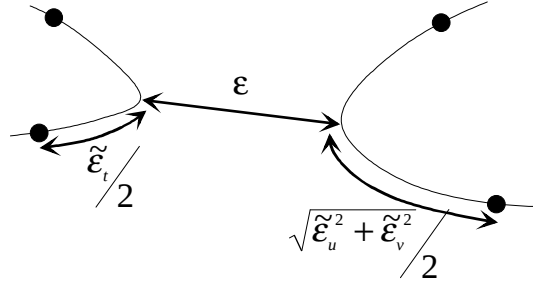
Pour cette raison, la valeur du critère de test sur la distance entre un point de la courbe et un point de la surface, noté ε^* , sera non pas égal à ε , mais calculé de sorte à prendre en compte (voir figure III.7):

- ◇ la distance maximale d'un point quelconque de la courbe au point paramétré le plus proche (soit $\tilde{\varepsilon}_t/2$);
- ◇ la distance maximale de détection sans incertitude (soit ε);
- ◇ la distance maximale d'un point quelconque de la surface au point paramétré le plus proche (soit $\sqrt{\tilde{\varepsilon}_u^2 + \tilde{\varepsilon}_v^2}/2$).

Soit, finalement:

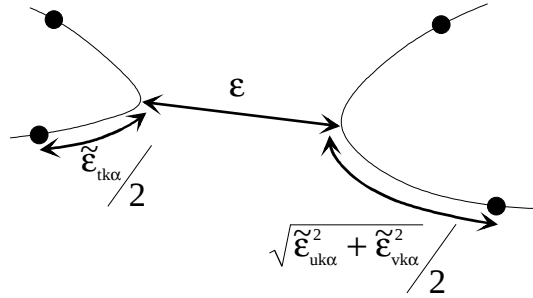
$$\varepsilon^* = \frac{\tilde{\varepsilon}_t}{2} + \varepsilon + \frac{\sqrt{\tilde{\varepsilon}_u^2 + \tilde{\varepsilon}_v^2}}{2}$$

La prise en compte dans l'algorithme de cette précision corrigée doit s'accompagner d'une procédure supplémentaire. En effet, des zones de grandes tailles (par rapport à la précision réelle) vont être mises à jour sur la surface, qui contiendront des points dont la distance à la courbe pourra être très supérieure à la précision recherchée. Il s'agit donc de vérifier dans chacune de ces zones l'existence d'au moins un point distant de la courbe d'au plus ε .

Figure III.7 : Définition de ε^*

Si un point répondant à ce critère existe, l'algorithme peut poursuivre normalement. Mais si aucun point ne répond à ce critère, il est nécessaire de procéder à un affinage de cette zone: il faut effectuer un rebalayage, en conservant la même précision de recherche, mais en utilisant la grandeur $\tilde{\varepsilon}_{t,k,\alpha} = \alpha^k \tilde{\varepsilon}_t$ où α est un réel strictement positif et inférieur à 1, et k représente le nombre de balayages supplémentaires réalisés pour l'affinage en cours (le cas $k = 0$ correspond au balayage "normal"). Cette modification entraîne l'introduction des paramètres $\tilde{\varepsilon}_{u,k,\alpha}$ et $\tilde{\varepsilon}_{v,k,\alpha}$ dont les valeurs seront précisées plus loin. Si cela est nécessaire, le processus peut être répété plusieurs fois. Afin d'éviter de définir des incréments trop petits par rapport à la précision de ces balayages, il est préférable de modifier également ε^* comme le montre la figure III.8, en prenant:

$$\varepsilon_{k,\alpha}^* = \frac{\tilde{\varepsilon}_{t,k,\alpha}}{2} + \varepsilon + \frac{\sqrt{\tilde{\varepsilon}_{u,k,\alpha}^2 + \tilde{\varepsilon}_{v,k,\alpha}^2}}{2}$$

Figure III.8 : Modification de ε^*

Cette définition *dynamique* du critère $\varepsilon_{k,\alpha}^*$ permet de restreindre la taille des zones à explorer au cours d'une phase d'affinage. Afin de prendre en compte de façon identique les situations rencontrées, il a été choisi de considérer la précision ε au sens strict, c'est-à-dire de ne pas conserver les points de la surface qui se trouveraient à une distance **égale** à ε de la courbe.

III.1.4.4 Choix des paramètres définissant les incréments

Le calcul des incréments n'est pas complet tant que les grandeurs $\tilde{\varepsilon}_t$, $\tilde{\varepsilon}_u$ et $\tilde{\varepsilon}_v$ n'ont pas été précisées. Ces grandeurs vont être définies en fonction de la précision ε du balayage en cours.

Le choix de $\tilde{\varepsilon}_t$ doit être d'une part effectué de façon à ne laisser aucune partie de la courbe non incluse dans l'une des différentes boules du balayage, la surface pouvant passer dans l'un de ces "trous" (voir figure III.9). Autrement dit, il importe que deux points paramétrés de la courbe ne soit pas distants d'une quantité supérieure à 2ε . D'autre part, la taille des incréments ayant une influence non-négligeable sur les temps de calcul, il est préférable de ne pas définir d'incréments trop petits par rapport à la précision du balayage en cours.

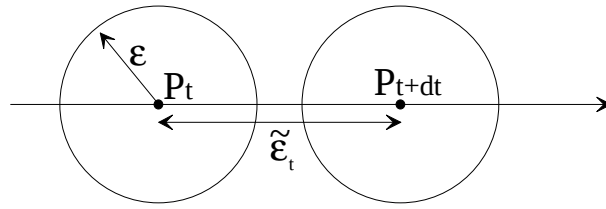


Figure III.9 : Mauvais choix de $\tilde{\varepsilon}_t$

En ce qui concerne la surface, et par souci de simplification, il sera choisi de prendre:

$$\tilde{\varepsilon}_u = \tilde{\varepsilon}_v$$

Par ailleurs, le choix de $\tilde{\varepsilon}_t$ conditionne le choix de $\tilde{\varepsilon}_u$ et $\tilde{\varepsilon}_v$. En effet, comme le montre la figure III.III.1.4.4, la distance maximale séparant deux points paramétrés de la surface, notée d_{max} , ne doit pas excéder $2\sqrt{\varepsilon^2 - \frac{\tilde{\varepsilon}_t^2}{4}}$. L'hypothèse d'égalité des grandeurs $\tilde{\varepsilon}_u$ et $\tilde{\varepsilon}_v$ entraîne finalement la relation (voir figure III.III.1.4.4):

$$\tilde{\varepsilon}_u = \sqrt{2}\sqrt{\varepsilon^2 - \frac{\tilde{\varepsilon}_t^2}{4}}$$

Il apparaît ainsi que le choix de l'ensemble des incréments est ramené au choix de la grandeur $\tilde{\varepsilon}_t$. Compte tenu des remarques formulées au début de ce paragraphe, il a été finalement choisi de prendre $\tilde{\varepsilon}_t = \tilde{\varepsilon}_u$, ce qui entraîne:

$$\tilde{\varepsilon}_t = \tilde{\varepsilon}_u = \tilde{\varepsilon}_v = \frac{2}{\sqrt{3}} \varepsilon \approx 1,155 \varepsilon$$

Il en résulte également:

$$\varepsilon^* = \frac{\varepsilon}{\sqrt{3}} (1 + \sqrt{2} + \sqrt{3}) \approx 2,394 \varepsilon$$

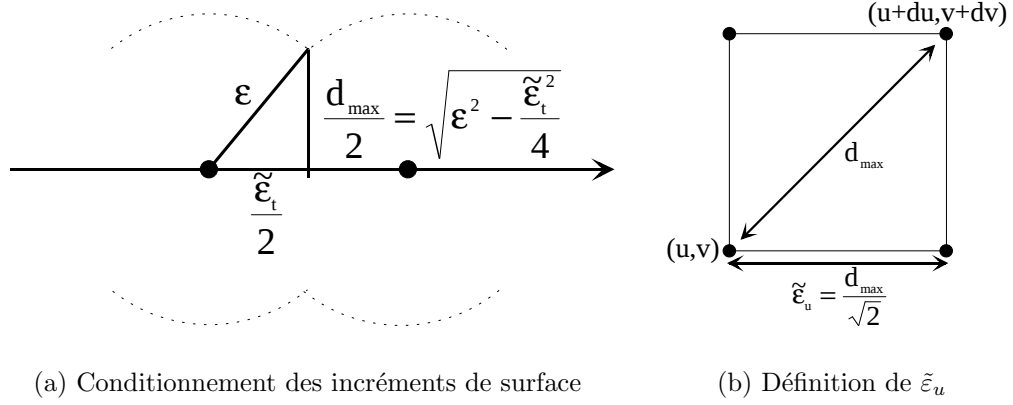


Figure III.10 : Conditions sur les paramètres incrémentaux

Dans les cas où l'affinage est nécessaire, et compte-tenu de la relation à respecter pour définir les incréments de surface par rapport à ceux de la courbe, il vient:

$$\tilde{\epsilon}_{u,k,\alpha} = \tilde{\epsilon}_{v,k,\alpha} = \sqrt{2} \sqrt{\epsilon^2 - \frac{\tilde{\epsilon}_{t,k,\alpha}^2}{4}}$$

soit, tous calculs faits (noter que si $k = 0$, alors $\epsilon_{k,\alpha}^* = \epsilon^*$):

$$\epsilon_{k,\alpha}^* = \frac{\epsilon}{\sqrt{3}} \left(\alpha^k + \sqrt{3 - \alpha^{2k}} + \sqrt{3} \right)$$

Dans les applications réalisées ici, il a été choisi $\alpha = 1/2$, et un test arrête les affinages au-delà d'un nombre fix (dans nos développements, ce nombre vaut 5): si le doute subsiste encore après celui-ci, il sera supposé que le contact n'existe pas. La variable $\tilde{\epsilon}_{t,k,\alpha}$ tendant vers 0, alors que $\epsilon_{k,\alpha}^*$ tend vers 2ϵ , un grand nombre d'affinages engendre un nombre d'incrément important pour chaque zone, et par conséquent des temps de calcul également prohibitifs.

III.1.4.5 Remarques

1. Cas des zones de tangence:

Il est important de voir que le choix de l'incrément tel qu'il a été présenté et la méthode d'examen des zones potentielles de contact utilisée génèrent des calculs de plus en plus longs lors de l'étude de cas tel que l'analyse trop fine d'une zone de tangence (voir cas traité sur la figure III.13). En effet, les bornes de ces zones vont très peu évoluer alors que la taille des incréments va diminuer. L'algorithme finira donc par balayer des zones très grandes avec des incréments très petits, d'où le problème. Un développement est actuellement en cours de réalisation dont le principe est le suivant: une fois les cas de tangence identifiés, les balayages ne

portent que sur les bornes de la zone de tangence (l'intérieur du domaine n'est pas réexplorer lorsque la précision diminue).

2. Evolution de la précision entre deux balayages:

Pour limiter les calculs, il a été observé qu'il vaudrait mieux diminuer "lentement" la précision entre deux balayages successifs (voir tableau III.1). En effet, une diminution trop brutale de la précision entre deux étapes entraîne l'exploration de zones souvent démesurées par rapport à la taille des incréments. A l'inverse, une diminution vraiment trop lente (coefficient voisin de 1) n'est évidemment pas non plus souhaitable. Dans la programmation effectuée, la précision est divisée par 2 à chaque étape de l'étude.

III.1.5 Exemples

Bien qu'il ait été supposé que les cas traités présenteraient un caractère quasi académique (la surface est un cylindre, la courbe et un morceau de parabole), des cas plus difficiles à traiter ont été envisagés: présence non pas d'un seul point d'intersection, mais de plusieurs points distincts, voire d'une zone de tangence, et même de points multiples (boucles).

Quelques exemples sont présentés ci-dessous. Pour chacun d'entre eux, la modélisation de la courbe et de la surface est faite: lorsque des temps de calculs (CPU) sont donnés⁶, ils correspondent à la lecture des paramètres définissant les splines (nombre de points de contrôle, coordonnées, degré souhaité pour la modélisation), à la création des entités géométriques et à la recherche d'intersection proprement dite. Il convient de souligner que le temps de calcul nécessaire à la génération des courbes et des surfaces est négligeable devant le temps nécessaire à la recherche d'intersection.

III.1.5.1 Intersections multiples et franches

Le test consiste à chercher l'intersection d'une spirale (30 points de contrôle) avec un plan (10 x 20 points de contrôle) (voir figure III.III.1.5.1). La dimension du plan est environ 10 x 70. Les splines utilisées sont de degré 1 pour le plan, et de degré 3 pour la spirale. La précision initiale est 1.0, et la précision demandée est 0.001. Les points d'intersection, obtenus pour un temps de calcul de l'ordre de 2 secondes, sont présentés figure III.III.1.5.1.

III.1.5.2 Points doubles

Le test consiste à chercher l'intersection d'une droite (10 points de contrôle, spline de degré 1) avec une surface cylindrique dont un parallèle a la forme d'une boucle (15 x 5 points de contrôle, spline de degré 3) comme le montre la figure III.III.1.5.2. La précision initiale est 1.0, et la précision demandée est 0.001. Les points d'intersection obtenus sont

⁶tous les calculs ont été réalisés sur une même machine, ce qui autorise la comparaison des temps de calcul entre eux

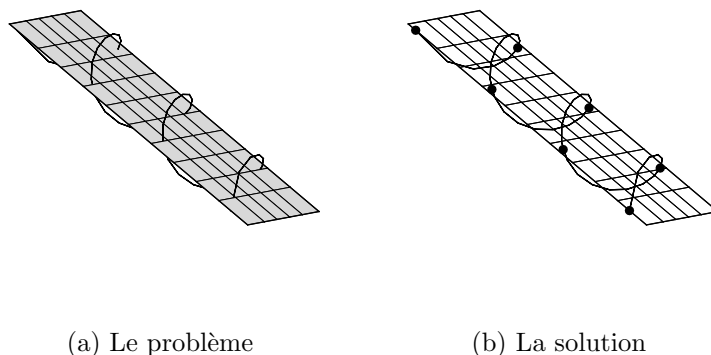


Figure III.11 : Exemple d'intersections multiples

présentés figure III.III.1.5.2. Un autre essai a été réalisé en n'utilisant que 2 points de contrôle pour la droite. Ils conduisent au même résultat.

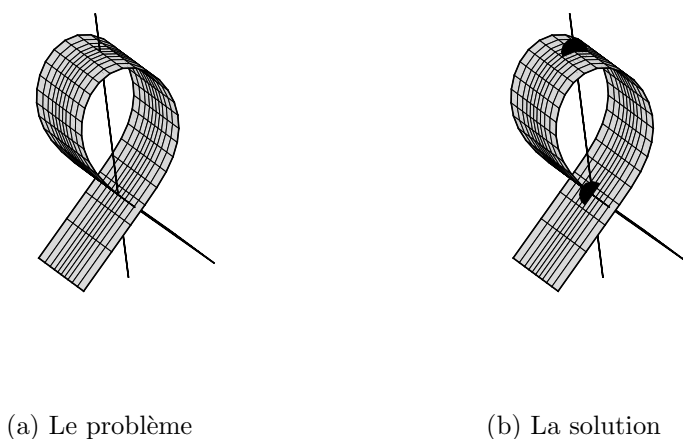


Figure III.12 : Exemple de point double

III.1.5.3 Zone de tangence

Le test consiste à chercher l'intersection d'une droite avec une surface cylindrique dont un parallèle a la forme d'une boucle. La droite est tangente à une partie de la boucle, comme le montre la figure III.III.1.5.3. La précision initiale est 1.0, et la précision demandée est 0.001. La zone d'intersection obtenue est présentée en gras sur la figure III.III.1.5.3. Ce test a été l'occasion de comparer les temps de calcul suivant la façon dont la précision diminuait d'une étape à l'autre. Les résultats sont présentés dans le tableau III.1.

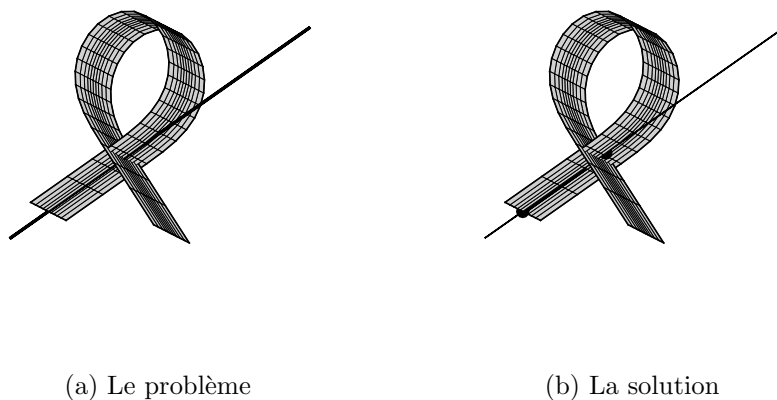


Figure III.13 : Exemple de tangence

A chaque étape la précision est divisée par	Temps de calcul (en minutes-secondes)
1.5	6' 05"
2	1' 48"
3	3' 45"
5	4' 33"
10	22' 56"

Tableau III.1: Comparaison des temps de calcul

III.1.6 Remarque: l'intersection surface/surface

Un prolongement possible des développements réalisés pour la recherche d'intersection courbe/surface est la recherche d'intersection surface/surface. Il suffit en effet pour cela de considérer qu'une surface est constituée de courbes mises les unes à côté des autres. La recherche d'intersection entre deux surfaces peut donc se ramener à la recherche d'intersection d'une surface avec un grand nombre de courbes extraites de l'autre surface. Le résultat se présente donc sous la forme non pas d'une courbe d'intersection, mais d'un ensemble de points communs aux deux surfaces.

Le gros inconvénient de cette méthode dans sa forme actuelle réside dans les temps de calculs qu'elle induit (voir annexe F). Elle est donc inadaptée au problème à résoudre, et ne sera pour l'instant pas exploitée.

III.2 Pré et post-traitement des résultats

La recherche d'intersection qui a été développée ici sera testée dans deux codes éléments finis différents:

- ◊ SAMCEF (développé par Samtech): code implicite, qui n'est pas ouvert a priori, mais qui offre la possibilité à l'utilisateur d'implémenter un élément dont il est libre de fixer la loi de comportement et les caractéristiques (nombre de nœuds, de degrés de liberté...): c'est par cette voie que la méthode de contact par splines sera introduite dans le code, en utilisant comme support l'élément de contact développé par I. Guilleaume [GUI99a] [GUI99b];
- ◊ PLEXUS (développé par le CEA): code explicite, dans lequel les développements ont pu être implantés avec l'accord du CEA et avec l'aide de H. Bung, principal développeur du code [CEA99].

Les données requises par l'algorithme de recherche du contact sont rassemblées dans le paragraphe III.2.1.

Le traitement du contact dans un code éléments finis comprend la phase de recherche des points de contact proprement dite, puis une phase de traitement du résultat obtenu, consistant principalement en l'identification des facettes impactées, et en la localisation des points d'impact dans un repère lié à chacune des facettes (repère barycentrique). Différentes grandeurs sont alors requises, qui dépendent des codes utilisés⁷:

- ◊ facette impactée, position relative du point d'impact sur cette facette, et valeur de la profondeur de pénétration des deux entités (un effort qui dépend de cette dernière valeur est alors attribué à chacun des nœuds de la facette): c'est le cas pour SAMCEF;
- ◊ liste de tous les nœuds ayant pénétré, facette de reprojection pour chacun de ces nœuds et coordonnées barycentriques dans chacune des facettes: c'est le cas pour PLEXUS.

Ces points seront développés dans les paragraphes III.2.2 et III.2.3.

III.2.1 Données requises

Le risque qu'encourt une méthode telle que celle qui a été développée réside dans le nombre élevé de paramètres qu'elle requiert en termes de données de modélisation, de critères d'arrêt, de nombres d'itérations tolérés dans certaines boucles ou de précisions à atteindre: quels sont ceux dont il faut laisser la maîtrise à l'utilisateur ? quels sont ceux qu'il vaut mieux imposer ? Le choix s'impose de lui-même pour certains d'entre eux (paragraphe III.2.1.1), mais reste à la libre disposition de l'utilisateur pour d'autres (il faut comprendre que dans ce cas, des paramètres par défaut ont été prévus, voir paragraphe III.2.1.2).

⁷en fait, les quantités requises dépendent surtout des schémas d'intégration temporelle utilisés par chacun des codes

III.2.1.1 Les données indispensables

Les informations suivantes sont données sous l'hypothèse que le programme principal du code a identifié un certain nombre de domaines pour lesquels la recherche de contact doit être réalisée. Les éléments qui doivent impérativement être fournis à la routine de recherche du contact par le programme principal du code sont alors:

◇ **pour la courbe:**

- les coordonnées des points (éventuellement leur nombre);
- un tableau de connectivité;
- le degré et le type de méthode choisis pour la modélisation spline;

◇ **pour la surface:**

- les coordonnées des points et leur nombre dans les deux directions;
- un tableau de connectivité;
- le degré et le type de méthode choisis pour la modélisation spline (le degré pouvant dépendre de la direction);
- l'épaisseur et un point définissant l'extérieur de la surface.

III.2.1.2 Les critères cachés

Parmi les grandeurs importantes qui régulent le bon déroulement de la recherche de contact, il faut citer les quatre suivantes:

- ◇ **le critère d'arrêt du balayage:** ce paramètre est fondamental puisqu'il détermine la finesse du résultat et conditionne la durée du calcul (plus il est petit, plus le balayage est long); ce paramètre doit être le même d'une zone de contact à une autre, et ne doit surtout pas dépendre de l'instant où le calcul est mené: un critère d'arrêt variable va se traduire par des problèmes dans la rupture du contact; par ailleurs, il est difficile de définir une valeur intrinsèque puisque l'unité de longueur n'est pas connue; ici, il a été choisi de prendre ce critère égal à une fraction de l'épaisseur de la surface (par exemple 1/1000);
- ◇ **la précision initiale du balayage:** là encore, ce paramètre conditionne la rapidité du calcul (trop petit, il constitue un facteur qui ralentit le calcul); ici, il a été pris égal à la moitié de la longueur de la polyligne reliant les points qui définissent la courbe;
- ◇ **le critère d'arrêt de l'affinage:** (voir paragraphe III.1.4.4) il s'agit d'un entier, fixé à 5 dans tous nos développements;

- ◇ **le critère de définition des zones de tangences:** en effet, les résultats sont traités différemment selon que la courbe tangente la surface ou bien qu'elle la coupe franchement; ici, une zone de contact est traitée comme une zone de tangence lorsque plus de 5 incréments d'espace séparent les deux bornes qui définissent la zone de contact sur la courbe (soit $t_{max} - t_{min} > 5.dt$).

III.2.2 Evaluation de la pénétration

Cette grandeur n'est nécessaire que pour le code SAMCEF (élément de contact développé par I. Guilloteau [GUI99b]).

III.2.2.1 Définition

La définition intrinsèque de la pénétration n'est pas aisée, et le choix de calcul qui est présenté ici repose sur des cas de figure susceptibles d'être rencontrés par la suite dans la configuration d'étude aube/carter. C'est pourquoi la définition suivante est retenue: si P_1 et P_2 sont deux points d'intersection consécutifs entre la courbe $\mathcal{C}$ et la surface $\mathcal{S}$, alors la pénétration gap de la courbe dans la surface est (voir figure III.III.2.2.1):

$$gap = \max_{P \in [P_1, P_2]} \left(\min_{Q \in \mathcal{S}} \|\overrightarrow{PQ}\| \right)$$

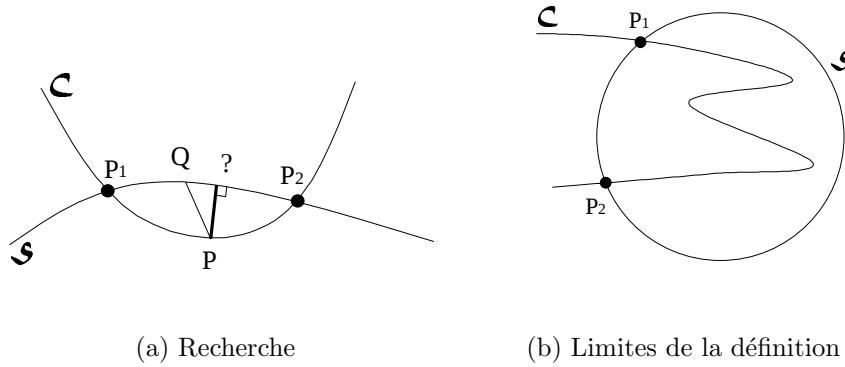


Figure III.14 : Définition de la pénétration

La figure III.III.2.2.1 montre combien cette définition ne saurait être suffisante dans tous les cas. Cependant, nous travaillons en faisant l'hypothèse que les cas rencontrés seront tous du type de celui présenté sur la figure III.III.2.2.1.

III.2.2.2 Organisation du calcul

Chaque point d'intersection est tout d'abord classé dans une des quatre catégories représentées sur les figures III.15, sur lesquelles l'intérieur de la surface est supposé

être la partie inférieure des schémas (les courbes sont orientées dans le sens croissant du paramétrage).

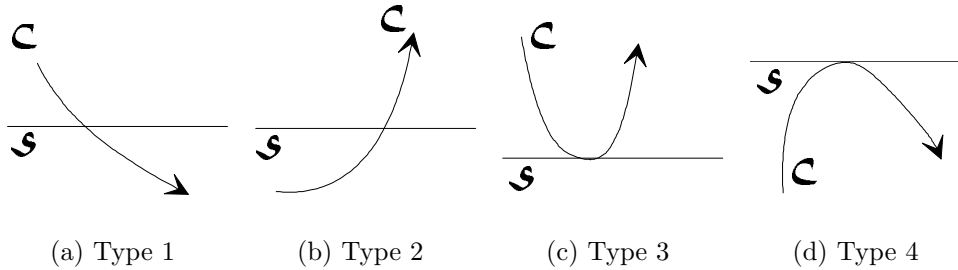


Figure III.15 : Classification des points d'intersection

Si deux points d'intersection consécutifs forment un des couples (1,2), (1,4), (4,2) ou (4,4), alors la partie de courbe située entre ces deux points se trouve à l'intérieur de la surface, et il y a donc possibilité de calculer la pénétration.

Le calcul de cette pénétration se fait selon un algorithme similaire à celui exposé pour la recherche des points d'intersection: pour chaque point P_t compris entre deux points d'intersection, la surface est balayée, et les distances $P_t P_{uv}$ sont calculées. Ces distances sont comparées à une valeur ε , ce qui permet de restreindre progressivement les zones de recherche sur la surface (le procédé est décrit sur la figure III.3). L'objectif est ici de déterminer le point de la surface tangent à une sphère centrée sur P_t et de rayon le plus petit possible. Ainsi, lorsque tous les points de la surface sont situés à une distance supérieure à ε , cette quantité est augmentée (voir figure III.III.2.2.2). Au contraire, si le balayage détecte des points de la surface situés à une distance inférieure à ε , cette quantité est diminuée (voir figure III.III.2.2.2). Augmentation et diminution se font suivant un principe de dichotomie, et le procédé se termine lorsque l'écart relatif entre deux précisions successives devient inférieur à une valeur donnée. La dernière valeur obtenue pour ε est considérée comme la distance du point P à la surface (voir figure III.III.2.2.2).

III.2.2.3 Exemple

L'exemple qui est traité est celui de la figure III.11: il s'agit de l'intersection d'une spirale avec un plan. Une demi-épaisseur a été attribuée au plan, et le résultat du calcul de la pénétration est présenté sur la figure III.17, qui montre une vue de côté du problème. Un segment reliant un point extrême à sa projection est indiqué en trait épais.

III.2.3 Reprojection sur une facette

Les résultats obtenus dans ce paragraphe servent pour les deux codes testés ici, à savoir SAMCEF et PLEXUS, mais c'est la présentation des calculs relatifs à l'emploi de ce dernier qui vont être développés ici.

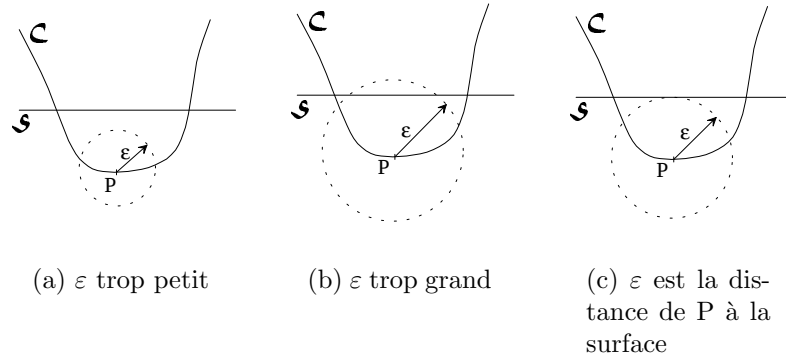


Figure III.16 : Evolution de la recherche

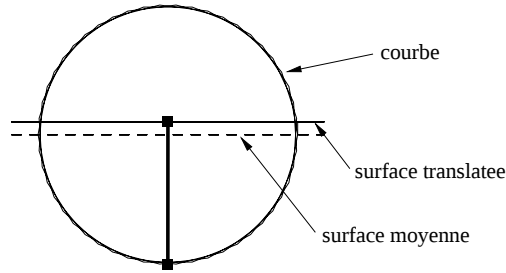


Figure III.17 : Recherche de pénétration

III.2.3.1 Calcul des efforts de contact: présentation théorique

Dans le cas de PLEXUS, les équations du mouvement sont traitées sous la forme:

$$\mathbb{M}.\ddot{X} = F^{ext} - F^{int}$$

Le contact entre deux entités se traduit par l'ajout à cette relation d'une équation de liaison portant sur les vitesses des nœuds des facettes impactées d'une part, et des nœuds impactants de l'autre [BON95] [BON97]. Dans le déroulement global du calcul dynamique, les coefficients de cette relation sont établis après la prise en compte des efforts externes, et avant le calcul des nouvelles accélérations (voir diagramme III.18).

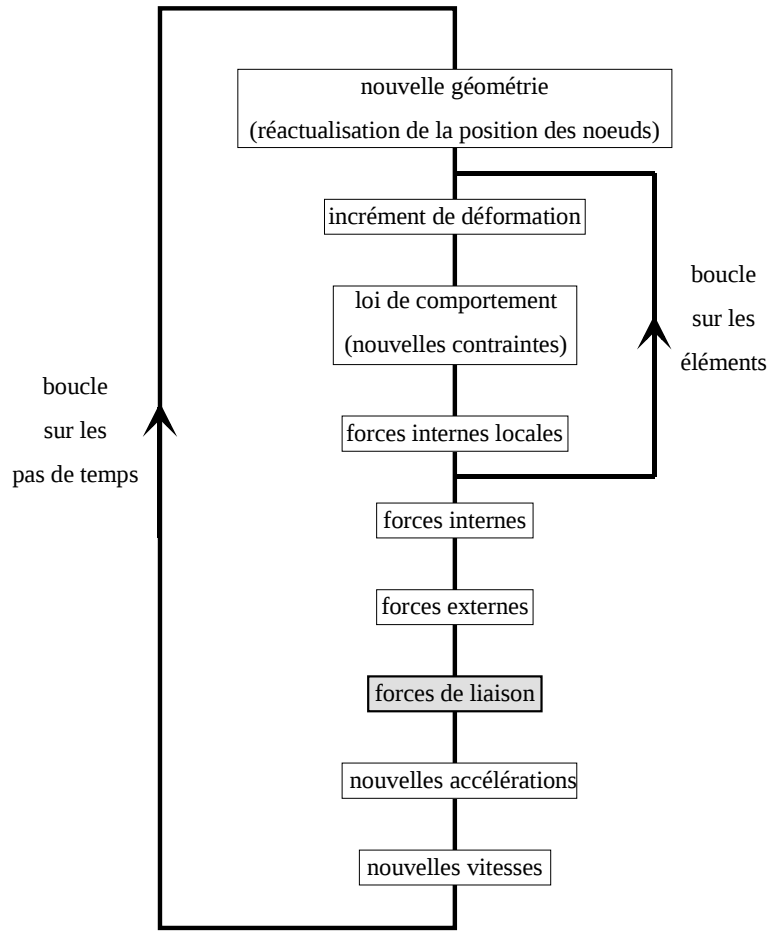


Figure III.18 : Processus général du calcul par éléments finis dans PLEXUS

Le cas d'une liaison est développé ici, avec les notations suivantes⁸:

- P : nœud impactant (nœud de la courbe)
- $N_1N_2N_3N_4$: facette impactée (élément de la surface)
- $\alpha_1\alpha_2\alpha_3\alpha_4$: coefficients barycentriques de la projection de P dans $N_1N_2N_3N_4$
- $\vec{n}$: vecteur normal à la surface au point d'impact P
- $\overrightarrow{V_i(M)}$: vecteur vitesse d'un point M au temps t_i

La relation en vitesse s'écrit alors (l'indice de temps sera précisé ultérieurement):

$$\overrightarrow{V(P)} \cdot \vec{n} = \sum_{i=1}^4 \alpha_i \overrightarrow{V(N_i)} \cdot \vec{n}$$

⁸seuls des éléments à 4 nœuds peuvent être utilisés pour le maillage de la surface

soit encore:

$$\overrightarrow{V(P)} \cdot \vec{n} - \sum_{i=1}^4 \alpha_i \overrightarrow{V(N_i)} \cdot \vec{n} = 0$$

Pour chaque nœud ayant impacté la surface, une relation de liaison similaire est écrite. L'ensemble de ces relations est ensuite regroupé sous la forme matricielle suivante:

$$\mathbb{C} \cdot \dot{X} = 0$$

où X désigne le vecteur des degrés de liberté (ddl) de la structure, et $\mathbb{C}$ est une matrice qui compte autant de lignes que de liaisons à écrire, et autant de colonnes qu'il y a de ddl dans la structure.

L'équation du mouvement se voit donc légèrement modifiée par l'ajout d'une force supplémentaire due à la liaison⁹:

$$\mathbb{M} \cdot \ddot{X}_{n+1} = F_{n+1}^{ext} - F_{n+1}^{int} + F_{n+1}^{liais}$$

où F_{n+1}^{liais} est supposée être de la forme:

$$F_{n+1}^{liais} = \mathbb{C}_{n+1}^t \cdot \lambda_{n+1}$$

et λ_{n+1} est la quantité à déterminer (vecteur des multiplicateurs de Lagrange).

En utilisant le schéma des différences finies centrales entre les instants $t_{n+\frac{1}{2}}$ et $t_{n+\frac{3}{2}}$ (voir figure III.19), il vient:

$$\ddot{X}_{n+1} = \frac{\dot{X}_{n+\frac{3}{2}} - \dot{X}_{n+\frac{1}{2}}}{\Delta t}$$

soit:

$$\mathbb{C}_{n+1} \cdot \ddot{X}_{n+1} = \mathbb{C}_{n+1} \cdot \left(\frac{\dot{X}_{n+\frac{3}{2}} - \dot{X}_{n+\frac{1}{2}}}{\Delta t} \right)$$

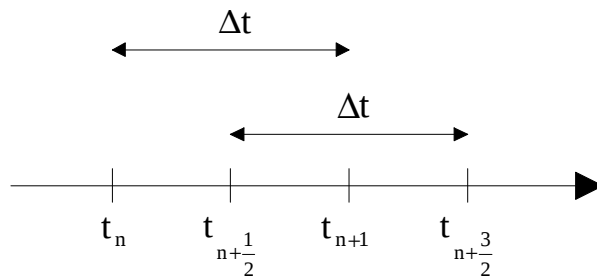


Figure III.19 : Pas de temps

La relation de liaison est écrite au temps $t_{n+\frac{3}{2}}$ sous la forme:

$$\mathbb{C}_{n+1} \cdot \dot{X}_{n+\frac{3}{2}} = 0$$

⁹ X_{n+1} désigne le vecteur des valeurs nodales à l'instant t_{n+1}

ce qui conduit donc à:

$$\mathbb{C}_{n+1} \cdot \ddot{X}_{n+1} \triangleq \mathbb{D}_{n+1} = -\frac{1}{\Delta t} \mathbb{C}_{n+1} \cdot \dot{X}_{n+\frac{1}{2}}$$

De l'équation du mouvement, il vient alors:

$$\ddot{X}_{n+1} = \mathbb{M}^{-1} \cdot (F_{n+1}^{ext} - F_{n+1}^{int} + \mathbb{C}_{n+1}^t \cdot \lambda_{n+1})$$

d'où:

$$\mathbb{C}_{n+1} \cdot \ddot{X}_{n+1} = \mathbb{D}_{n+1} = \mathbb{C}_{n+1} \cdot \mathbb{M}^{-1} \cdot (F_{n+1}^{ext} - F_{n+1}^{int}) + \mathbb{C}_{n+1} \cdot \mathbb{M}^{-1} \cdot \mathbb{C}_{n+1}^t \cdot \lambda_{n+1}$$

soit finalement, en posant $\mathbb{H}_{n+1} = \mathbb{C}_{n+1} \cdot \mathbb{M}^{-1} \cdot \mathbb{C}_{n+1}^t$:

$$\lambda_{n+1} = \mathbb{H}_{n+1}^{-1} (\mathbb{D}_{n+1} - \mathbb{C}_{n+1} \cdot \mathbb{M}^{-1} \cdot (F_{n+1}^{ext} - F_{n+1}^{int}))$$

La quantité λ_{n+1} est alors calculable, et les autres grandeurs du problème à l'instant t_{n+1} le sont également. Ainsi, pour obtenir ce résultat est-il nécessaire de connaître préalablement le vecteur normal $\vec{n}$ et les quatre coefficients $(\alpha_1, \alpha_2, \alpha_3, \alpha_4)$ pour chacun des points de la courbe ayant impacté la surface.

III.2.3.2 Coefficients barycentriques

Il faut ici préciser la notion de facette de reprojection: les éléments finis se déforment au cours du temps, et les 4 nœuds qui définissent un élément quadrangle par exemple peuvent très bien ne pas être tous situés dans un même plan. Or le calcul d'une reprojection d'un point dans une direction donnée sur une surface *quelconque* se révèle trop complexe pour être résolu simplement de façon numérique. On va donc chercher à rapporter ce calcul à un cas plus simple: celui de la projection orthogonale sur un plan¹⁰. Pour cela, le plan médian de l'élément (plan passant par les milieux des 4 côtés) est utilisé. Comme le montre la figure III.20, chacun des nœuds mis en jeu est reprojeté dans le plan médian: le nœud de la courbe P est projeté en H , et les quatre sommets N_1, N_2, N_3 et N_4 sont reprojetés respectivement en N'_1, N'_2, N'_3 et N'_4 .

Le point H est repéré dans le quadrilatère $N'_1 N'_2 N'_3 N'_4$ à l'aide de ses coordonnées barycentriques, c'est-à-dire à l'aide du quadruplet $(\alpha_1, \alpha_2, \alpha_3, \alpha_4)$ tel que:

$$\alpha_1 \overrightarrow{HN'_1} + \alpha_2 \overrightarrow{HN'_2} + \alpha_3 \overrightarrow{HN'_3} + \alpha_4 \overrightarrow{HN'_4} = \vec{0}$$

On impose également au quadruplet de vérifier:

$$\begin{aligned} \alpha_1 + \alpha_2 + \alpha_3 + \alpha_4 &= 1 \\ \alpha_1 &\geq 0 \\ \alpha_2 &\geq 0 \\ \alpha_3 &\geq 0 \\ \alpha_4 &\geq 0 \end{aligned}$$

¹⁰il serait tout aussi simple de reprojetter suivant une direction donnée, par exemple suivant le vecteur vitesse du nœud concerné

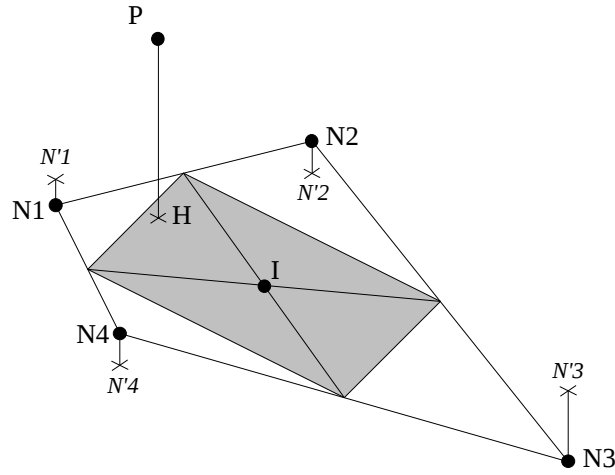


Figure III.20 : Projection des noeuds sur le plan médian

Les valeurs α_i sont obtenues par des calculs d'aires dans le quadrilatère concerné¹¹. Pour cela, il est nécessaire d'introduire le point I , point d'intersection des médianes de $N'_1N'_2N'_3N'_4$. Pour fixer les idées, supposons que H appartienne au triangle $IN'_1N'_4$ (voir figure III.21).

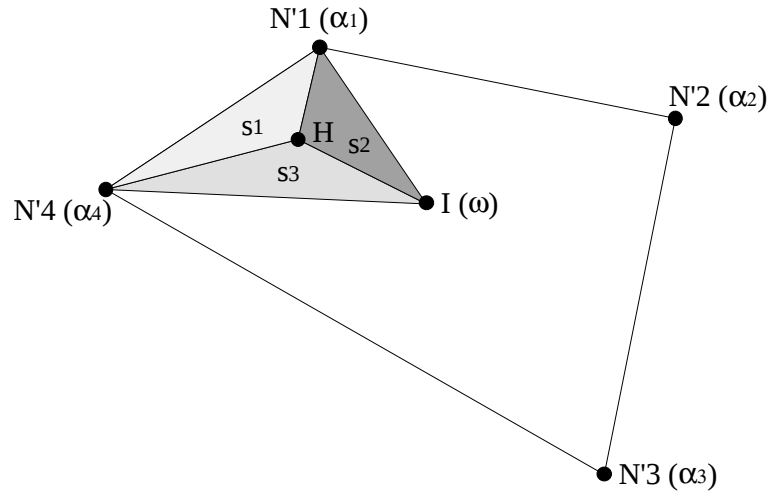


Figure III.21 : Coordonnées barycentriques de H

¹¹il n'est pas toujours possible d'exploiter directement la relation vectorielle $\alpha_1 \overrightarrow{HN_1} + \alpha_2 \overrightarrow{HN_2} + \alpha_3 \overrightarrow{HN_3} + \alpha_4 \overrightarrow{HN_4} = \vec{0}$; en effet, en projection sur trois axes, cette relation peut conduire à des équations inexploitable de type $0 = 0$: il suffit, par exemple, que tous les points soient situés dans un plan de coordonnée constante

Notons:

$$\begin{aligned}
 S &= \text{aire de } IN'_1N'_4 \\
 s_1 &= \text{aire de } HN'_1N'_4 \\
 s_2 &= \text{aire de } HIN'_1 \\
 s_3 &= \text{aire de } HN'_4I
 \end{aligned}$$

Le point I est affecté d'un coefficient ω égal à $\frac{s_1}{S}$. Les coefficients recherchés sont alors:

$$\begin{aligned}
 \alpha_1 &= \frac{s_3}{S} + \frac{\omega}{4} \\
 \alpha_2 &= \frac{\omega}{4} \\
 \alpha_3 &= \frac{\omega}{4} \\
 \alpha_4 &= \frac{s_2}{S} + \frac{\omega}{4}
 \end{aligned}$$

Quatre cas particuliers peuvent alors être envisagés (voir tableau III.2, avec les notations de la figure III.21), qui correspondent aux configurations pour lesquelles l'une des valeurs s_i devient nulle¹².

cas	exemple	conséquence	ω	α_1	α_2	α_3	α_4
$H \in [IN'_i]$	$H \in [IN'_1]$	$s_2 = 0$	$\frac{HN'_1}{IN'_1}$	$\frac{IH}{IN'_1} + \frac{\omega}{4}$	$\frac{\omega}{4}$	$\frac{\omega}{4}$	$\frac{\omega}{4}$
$H \in [N'_iN'_j]$	$H \in [N'_1N'_4]$	$s_1 = 0$	0	$\frac{HN'_4}{N'_1N'_4}$	0	0	$\frac{HN'_1}{N'_1N'_4}$
$H \equiv N'_i$	$H \equiv N'_1$	$s_1 = s_2 = 0$	0	1	0	0	0
$H \equiv I$	$H \equiv I$	$s_2 = s_3 = 0$	1	$\frac{1}{4}$	$\frac{1}{4}$	$\frac{1}{4}$	$\frac{1}{4}$

Tableau III.2: Cas particuliers

Il convient ici de faire une remarque importante: jusqu'à présent, la projection du nœud impactant a été présentée sans s'inquiéter de son existence. Or certaines configurations ne permettent pas de trouver une facette à l'intérieur de laquelle le nœud se reprojetterait, comme le montre la figure III.III.2.3.2: à cause de la discontinuité de la surface formée par la réunion des facettes 1 et 2, il existe une partie d'espace pour laquelle les points ne se reprojetent pas à l'intérieur de ces facettes. Pour résoudre ce problème, on choisit, si une telle situation se présente, de projeter le point sur l'un des segments du maillage.

Dès lors, il faut également envisager l'éventualité où la projection sur un segment n'est pas non plus possible, comme c'est le cas à l'intérieur du cône grisé de la figure III.III.2.3.2. Dans cette configuration, le point est alors "reprojeté" sur le nœud du maillage le plus proche.

¹²noter que ces cas particuliers ne correspondent pas à des *singularités de calcul* pour la démarche qui est présentée ici

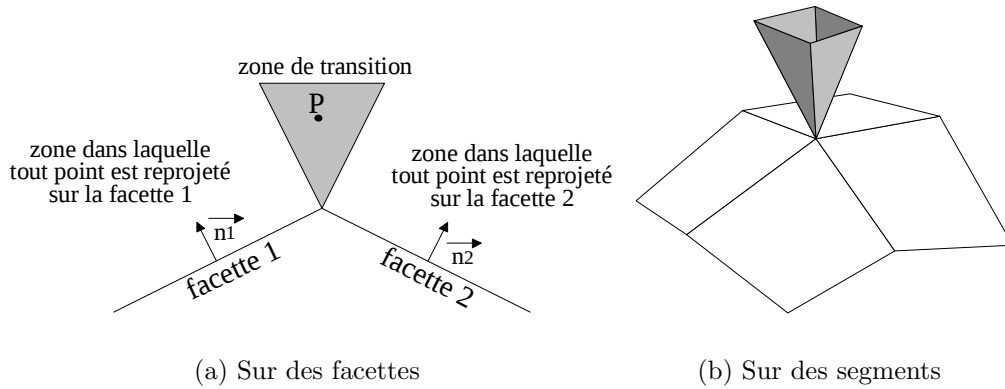


Figure III.22 : Problèmes de reprojection

III.2.3.3 Vecteur normal

Le calcul du vecteur normal constitue un réel problème. En effet, en toute rigueur, il faudrait, pour chaque nœud P ayant impacté la surface, trouver le projeté suivant la direction d'incidence de P sur la surface, et, en utilisant la formulation spline, déterminer le vecteur normal à la surface au point de projection. Malheureusement, le calcul exact nécessiterait la mise en œuvre pour la seconde fois au cours du même pas de temps, de l'algorithme de recherche du contact (cette reprojection revenant en fait à trouver l'intersection d'une droite et d'une surface).

Une solution simple consiste à utiliser le vecteur normal à la facette élément fini de reprojection de P . Il serait cependant désastreux d'utiliser, à ce stade, ce vecteur normal: cela, particulièrement sous PLEXUS, ruinerait l'effort de recherche fine du contact. Un compromis a donc été adopté, qui repose sur les hypothèses suivantes:

1. des cas avec des nœuds confondus, ou pour le moins très proches, ainsi que des courbes avec des points multiples ne seront pas rencontrés en cours de calcul;
2. entre deux pas de temps, la position des nœuds évolue suffisamment peu, ce qui suppose que les pénétrations restent faibles.

Ce compromis consiste:

- ◇ dans le cas où le nœud impactant est seul et quasiment tangent à la surface, la normale à la surface est calculée au centre de la dernière zone de contact définie sur celle-ci (c'est-à-dire au centre du domaine paramétré $[u_{min}, u_{max}] \times [v_{min}, v_{max}]$);
- ◇ si plusieurs nœuds ont impacté la surface, les normales sont calculées aux points paramétrés répartis dans le domaine $[u_{min}, u_{max}] \times [v_{min}, v_{max}]$ en fonction des longueurs de corde de la polygône reliant ces nœuds.

Chapitre IV

Résultats numériques

La méthode de recherche de contact développée au chapitre précédent est implémentée dans deux codes éléments finis:

- ◊ SAMCEF (société Samtech), via un élément utilisateur (paragraphe IV.1);
- ◊ PLEXUS (développé par le CEA), par implémentation directe dans les fichiers source du code (paragraphe IV.2).

Quelques exemples sont présentés, les calculs étant menés d'une part avec les méthodes présentes à l'origine dans ces codes, puis avec la méthode s'appuyant sur les fonctions splines.

IV.1 Cas du code SAMCEF

IV.1.1 Exemple traité

En ce qui concerne le code SAMCEF, un seul exemple a été abordé. Cet exemple a été proposé par SNECMA pour valider l'élément utilisateur développé pour réaliser une gestion simplifiée du contact [GUI99a]. Pour cela, le cas test défini ne fait intervenir que les phénomènes dont la modélisation est recherchée. Dans le cas présent, l'étude des grandes rotations, ou de la plastification des sommets d'aubes ne constitue pas une priorité: l'attention est portée uniquement sur le contact.

Le cas test ainsi proposé concerne une aube, au pied de laquelle est attaché un balourd¹, mise en rotation dans un carter (la courbe de montée en vitesse est représentée sur la figure IV.1). L'utilisation d'une configuration gyroscopique permet de s'affranchir de la rotation *réelle* de l'aube (une matrice gyroscopique est prise en compte dans les calculs, mais l'ensemble du problème est traité sous l'hypothèse des petits déplacements et des petites déformations). Le calcul est réalisé avec le module MECANO, et les routines utilisateurs sont compilées avec l'ensemble du code.

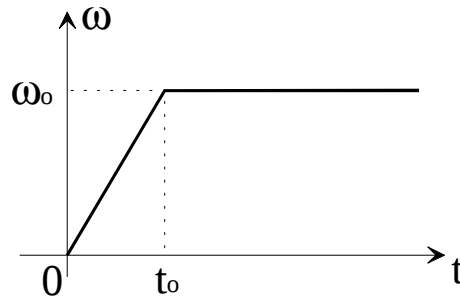


Figure IV.1 : Vitesse de rotation du rotor

Le carter est représenté par un tronc de cylindre encastré sur la partie arrière. Il est modélisé avec des éléments coque de Mindlin (éléments SAMCEF de type 29, possédant quatre nœuds, et six degrés de liberté par nœud), équi-répartis sur toute la circonférence. L'aube est modélisée par un élément ressort de type Bushing (les coefficients de sa matrice raideur sont introduits par l'utilisateur) dont les caractéristiques ont été déterminées au moyen de simulations sur des modèles plus complets. Cette aube est reliée à un rotor modélisé par des éléments poutre (éléments SAMCEF de type 22, possédant deux nœuds, et six degrés de liberté par nœud); ce rotor est encastré au niveau d'un support palier.

¹dans le code SAMCEF [SAM97], un balourd est introduit sous la forme d'une masse ponctuelle m située à une distance e de l'axe de rotation (supposé être l'axe $\vec{z}$), et repéré par un angle α ; si $\dot{\theta}(t)$ représente la vitesse de rotation du mouvement, le balourd induit des efforts f_x et f_y tels que:

$$\begin{aligned} f_x &= m e \dot{\theta}^2 \cos(\theta + \alpha) + m e \ddot{\theta} \sin(\theta + \alpha) \\ f_y &= m e \dot{\theta}^2 \sin(\theta + \alpha) - m e \ddot{\theta} \cos(\theta + \alpha) \end{aligned}$$

L'élément utilisateur développé relie les nœuds de l'aube à certains nœuds du carter (ceux des facettes susceptibles d'être impliquées dans le contact). Tous les matériaux sont supposés rester en phase élastique. Le maillage est représenté sur la figure IV.2.

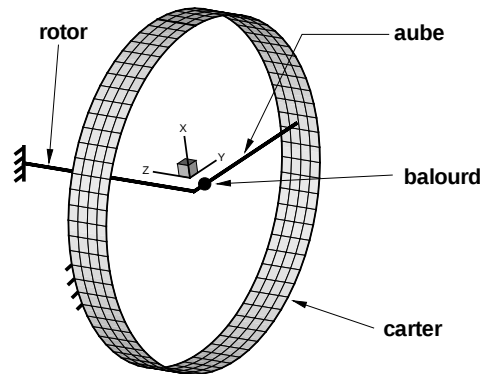


Figure IV.2 : Maillage du cas test SAMCEF

Dans sa formulation première, la recherche de contact utilisée par I. Guilloteau consiste en une recherche d'intersection de type droite/plan (droite formée par les deux nœuds du ressort, et plan formé par trois des quatre nœuds d'un élément du carter). Cette recherche du contact est donc remplacée par la recherche de l'intersection d'une courbe et d'une surface modélisées par splines. Pour l'aube, l'utilisation du degré 1 pour la modélisation des splines s'impose. Le choix de la méthode directe en découle puisqu'elle ne demande aucun calcul matriciel important. Pour le carter, des splines de degré 3 dans la direction orthoradiale et de degré 1 dans la direction de l'axe du cylindre sont utilisées. Ce choix privilégie la bonne prise en compte de la courbure du cylindre. Deux méthodes de mise en œuvre sont retenues pour le carter: interpolation et lissage. En effet, compte-tenu du positionnement des splines par rapport à leur polygone de contrôle, le choix de la méthode directe se révèle inintéressant, puisqu'il conduirait à sous-évaluer de façon importante le rayon réel du carter.

IV.1.2 Résultats

Deux grandeurs sont examinées: la première est l'évolution de la force de contact au cours du temps, et la seconde est le temps CPU nécessaire à la réalisation du calcul. Les paramètres qui varient sont d'une part le nombre d'éléments utilisés sur la circonférence du carter (trois configurations sont étudiées: 72, 36 et 18 éléments), et d'autre part le type de méthode mise en œuvre (classique, lissage et interpolation).

Remarques:

- ◇ toutes les forces présentées sont normalisées en fonction de l'amplitude

maximale de la première force de contact obtenue par interpolation dans la configuration 72 éléments;

- ◇ les temps de calculs sont tous normalisés par rapport au temps de calcul nécessaire dans la méthode classique avec 72 éléments.

IV.1.2.1 Examen de la force de contact

COMPARAISON DES MÉTHODES ENTRE ELLES

Les figures IV.3, IV.4 et IV.5 montrent les résultats fournis par les différentes méthodes selon la densité de maillage du carter.

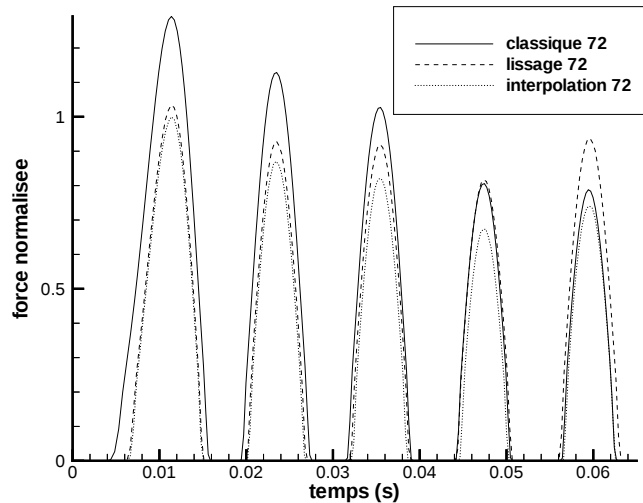


Figure IV.3 : Carter avec 72 éléments dans la circonférence

Les résultats suivants se dégagent de l'examen de ces courbes:

- ◇ pour un maillage de la circonférence du carter avec 72 éléments:
 - les méthodes utilisant des splines donnent des résultats très voisins pour ce qui est de l'amplitude maximale de la force au premier contact (environ 3% d'écart);
 - la méthode classique conduit à évaluer une amplitude maximale supérieure de près de 30% aux méthodes splines. Cette différence provient de l'approximation géométrique des éléments finis qui sous-évalue systématiquement le rayon du carter lors de la recherche d'intersection (voir figure IV.IV.1.2.1). Une autre conséquence de cette sous-évaluation réside dans l'instant du premier contact: avec une méthode classique, un premier contact est détecté à $t = 3,9$ ms, alors que cette valeur est portée à 6,2 ms dans le cas du lissage par exemple;

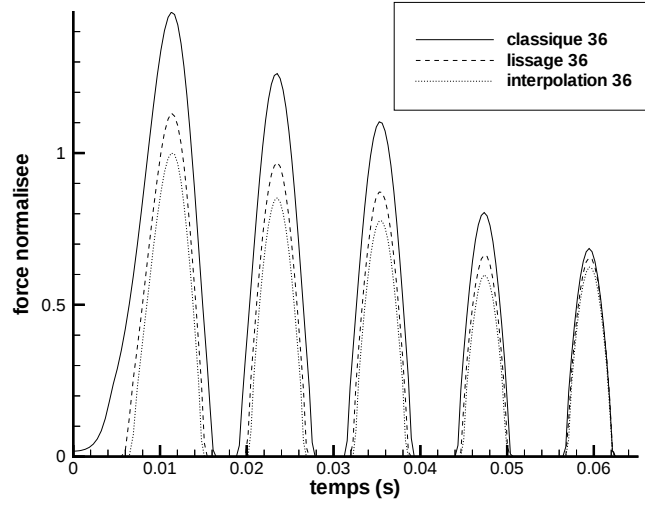


Figure IV.4 : Carter avec 36 éléments dans la circonférence

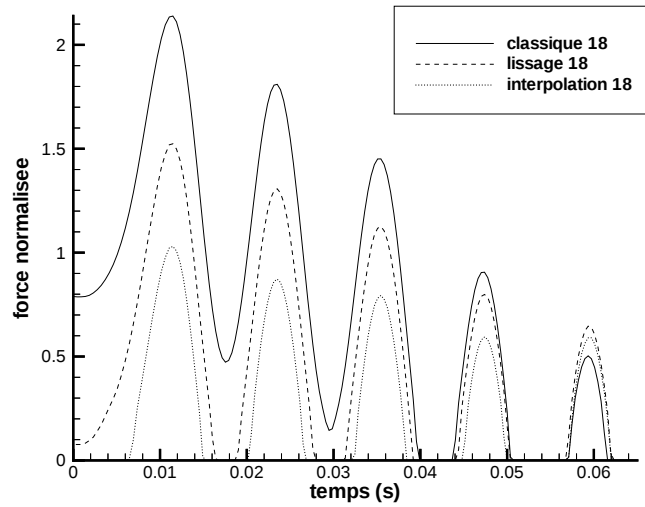


Figure IV.5 : Carter avec 18 éléments dans la circonférence

- ◇ la différence importante entre le rayon réel du carter et son rayon apparent est bien entendu accentuée lorsque le carter n'est plus modélisé que par 36 éléments finis. La force de contact non nulle dans le cas classique à l'instant $t = 0$ traduit même l'existence d'une pénétration initiale, dont l'origine est présentée sur la figure IV.IV.1.2.1. L'écart sur l'amplitude maximale de la force au premier contact entre la méthode classique et la méthode d'interpolation atteint alors presque 50%;

- ◇ la méthode de lissage présente davantage d'écart avec la méthode d'interpolation que précédemment: la raison en est là aussi la différence entre rayon réel et rayon apparent (voir figure IV.7);
- ◇ dans le cas où seuls 18 éléments modélisent la circonférence du carter, les écarts entre les trois méthodes deviennent très importants: plus de 100% entre méthode classique et interpolation, et près de 50% entre lissage et interpolation. Dans cette configuration, la méthode classique présente une très grande pénétration initiale (il faut d'ailleurs trois tours pour que le contact soit rompu entre aube et carter); il en est de même, à une échelle plus faible, pour le lissage.



Figure IV.6 : Problèmes liés à la méthode classique

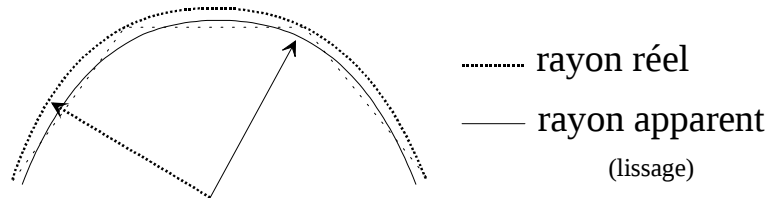


Figure IV.7 : Rayon apparent (lissage)

EVOLUTION EN FONCTION DU NOMBRE D'ÉLÉMENTS

Il est intéressant de comparer les performances de chacune des méthodes lorsque le nombre d'éléments varie. Les résultats sont présentés sur les courbes IV.8, IV.9 et IV.10.

La figure IV.8 confirme combien une recherche classique de contact ne peut se satisfaire d'un maillage quelconque: les résultats sur l'amplitude maximale de la première force de contact semblent converger, mais assez lentement. À l'inverse, la figure IV.10 met en évidence qu'une méthode d'interpolation par spline se révèle stable par rapport à la variation du nombre d'éléments définissant la géométrie des structures impliquées.

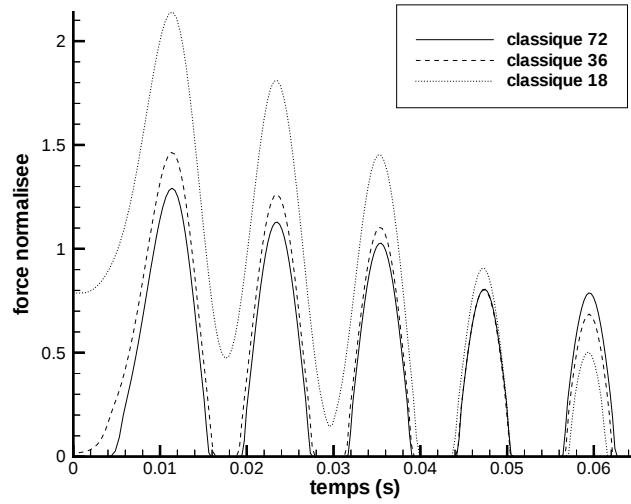


Figure IV.8 : Méthode classique

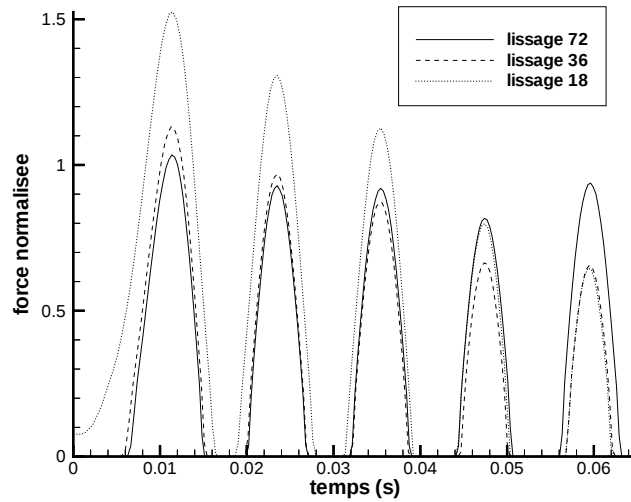


Figure IV.9 : Méthode de lissage

IV.1.2.2 Comparaison des temps de calcul

La figure IV.11 compare les temps nécessaires à la réalisation des simulations selon les méthodes employées.

Les méthodes splines apparaissent plus coûteuses, à configuration égale, en temps de calcul: respectivement 24% et 27% sur la configuration à 72 éléments. Ce résultat était bien sûr attendu, compte-tenu du nombre importants d'opérations supplémentaires

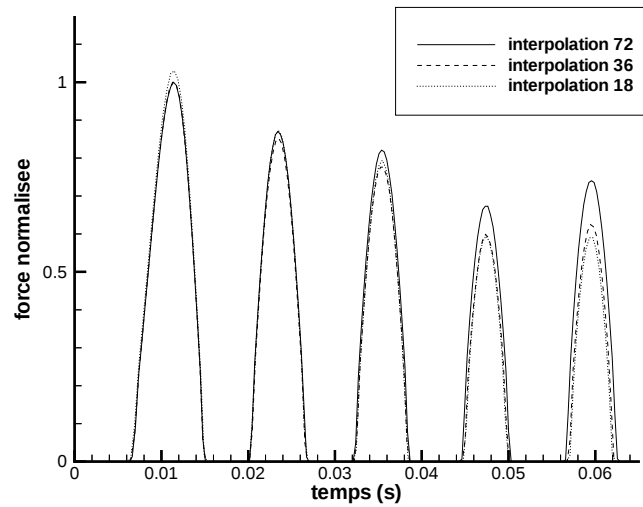


Figure IV.10 : Méthode d'interpolation

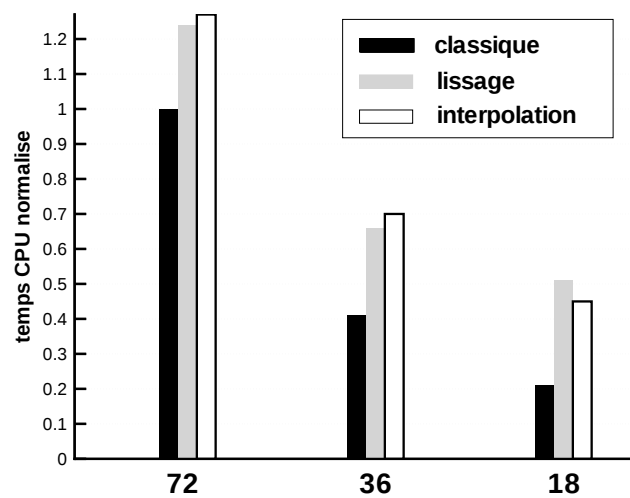


Figure IV.11 : Comparaison des temps de calcul

qu'elles engendrent.

Pour la configuration à 18 éléments, les temps CPU sont relativement faibles, ce qui, en terme de pourcentage, se traduit par des différences assez grandes, mais peu significatives.

Il paraît plus intéressant de noter que le temps de calcul nécessaire pour une interpolation avec 36 éléments est 30% inférieur à celui obtenu dans la configuration classique avec 72 éléments.

IV.2 Cas du code PLEXUS

IV.2.1 L'implémentation

Il est à ce jour possible pour tout utilisateur du code PLEXUS de recourir à la méthode de recherche de contact par splines car tous les développements nécessaires ont été réalisés dans les fichiers sources du code. Cependant, dans sa forme actuelle, seuls certains types de problèmes peuvent être abordés. La principale restriction est la suivante: la structure sur laquelle s'appuiera la surface doit être de révolution autour de l'axe z du repère global, la structure qui contiendra les nœuds définissant la courbe (ou éventuellement les courbes) devant être située à l'intérieur de cette surface. Au-delà de cette restriction, l'option "SPLINE", qui a été développée en collaboration avec H. Bung, présente de nombreux avantages, dont certains liés au langage objet du système CASTEM, qui la rendent souple d'utilisation. Cette option est détaillée en annexe E.

Lors de son activation, la phase du calcul dédiée aux forces de liaison utilise les routines splines pour déterminer ces forces, suivant le processus rappelé sur le diagramme IV.12.

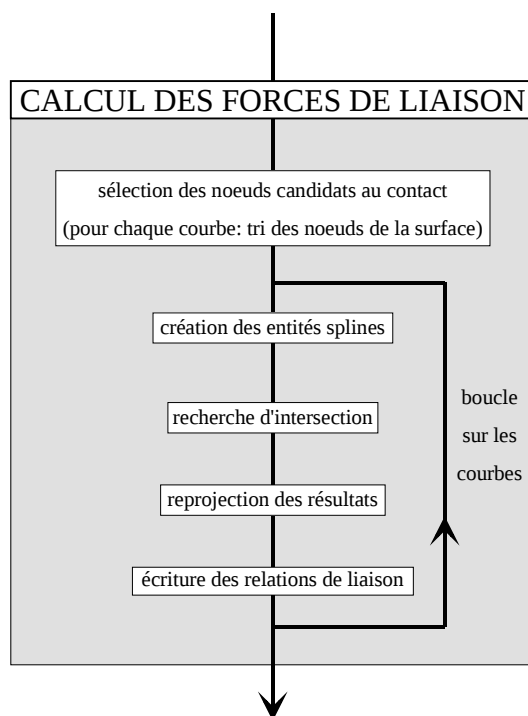


Figure IV.12 : Processus du calcul des forces de liaison

IV.2.2 Exemples traités

L'objectif principal de l'exemple traité est de permettre de valider l'implémentation des développements dans PLEXUS.

IV.2.2.1 Premier exemple

Cet exemple a pour objectif de montrer que l'utilisation des splines permet effectivement de mieux prendre en compte la courbure des structures étudiées.

DESCRIPTION

Le test consiste en une plaque (l'“aube”) animée d'une vitesse de translation initiale $\vec{V}_0 = V_0 \vec{x}$, se déplaçant vers la paroi d'un cylindre de révolution d'axe $\vec{z}$ (le “carter”) sur laquelle elle rebondit (voir figure IV.IV.2.2.1). Le carter est encastré à ses deux extrémités, et l'aube n'est soumise à aucune condition aux limites. L'ensemble est maillé à l'aide d'éléments coque Q4GS (éléments développés par Batoz [BAT90], avec 4 points d'intégration dans le plan). De façon à pouvoir valider les résultats donnés par le calcul par splines en le comparant aux résultats du code, il a été choisi de modéliser la surface et le sommet d'aube par des splines linéaires, à l'aide d'une méthode directe, puis en utilisant l'interpolation et le lissage.

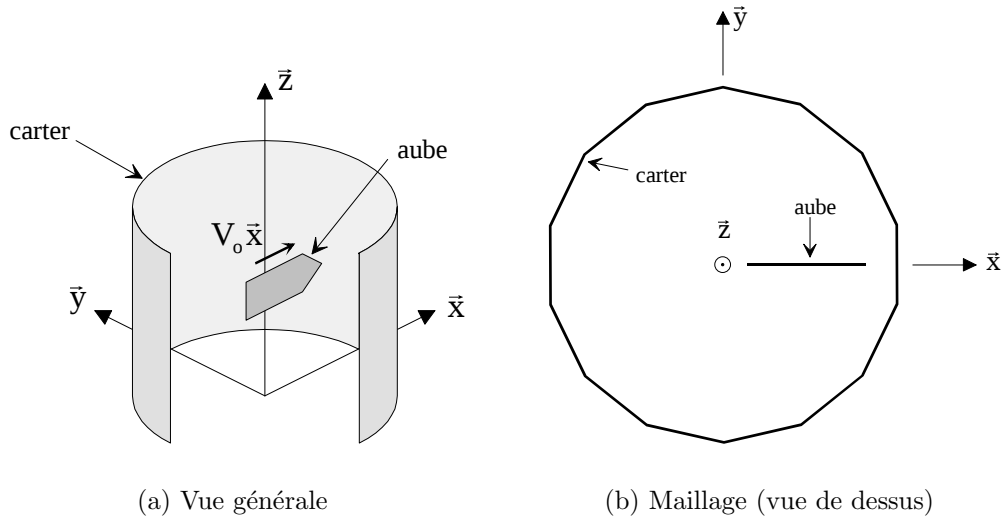


Figure IV.13 : Cas test sous PLEXUS

RÉSULTATS

Les résultats relatifs à la force de contact au sommet de l'aube sont présentés sur la figure IV.IV.2.2.1.

Remarque: le pic très important observé en début de réponse ne fait que traduire la discontinuité en vitesse: l'aube passant de V_0 à 0 au moment du contact, l'accélération devient infinie

De façon générale, les forces obtenues par les deux types de méthodes sont identiques d'un point de vue amplitude, les différents calculs avec les splines étant même indiscernables. Un léger décalage dans le temps apparaît entre les méthodes splines d'une part

et la méthode classique de l'autre (voir le détail sur la figure IV.IV.2.2.1). Ce décalage s'explique par la valeur de l'épaisseur de contact prise en compte. En effet, avec les méthodes splines, la valeur de l'épaisseur introduite correspond au maximum de translation possible de la surface moyenne, alors que pour la méthode classique, elle correspond exactement à la valeur de cette translation (voir à ce sujet le paragraphe II.3.5, et notamment la figure II.12 de la page 49). Il en résulte que la surface spline apparaît moins épaisse que la surface classique: le contact a donc lieu plus tardivement.

D'autres calculs ont ensuite été réalisés en utilisant cette fois des splines de degré 2 et 3 pour prendre en compte la courbure du carter. Les résultats sont présentés sur la figure IV.15.

La méthode d'interpolation donne deux courbes superposées, ce qui valide une bonne prise en compte de la courbure du carter, alors que le lissage fournit deux courbes décalées dans le temps (conséquence du fait que le lissage minimise le rayon apparent, comme cela a été évoqué sur la figure IV.7).

Il apparaît par ailleurs que la force de contact est *bruitée* avec la méthode d'interpolation: un peu avant la séparation, l'algorithme traite un pas de temps pour lequel la conclusion est l'absence de contact, la force de contact associée étant donc nulle pour ce pas de temps. Or deux incréments en temps après, du contact est à nouveau détecté: il en résulte une discontinuité des vitesses, et donc une force de forte amplitude qui apparaît très nettement sur le résultat. L'origine du problème *semble* être la suivante: entre deux pas de temps successifs, les faibles déplacements des nœuds du carter engendrent une grande différence entre les deux surfaces splines qu'ils définissent. Nous avons supposé que ce problème était en partie dû au conditionnement des matrices à inverser. Les quelques éléments donnés dans le paragraphe IV.2.3 permettent d'invalider cette hypothèse.

IV.2.2.2 Second exemple

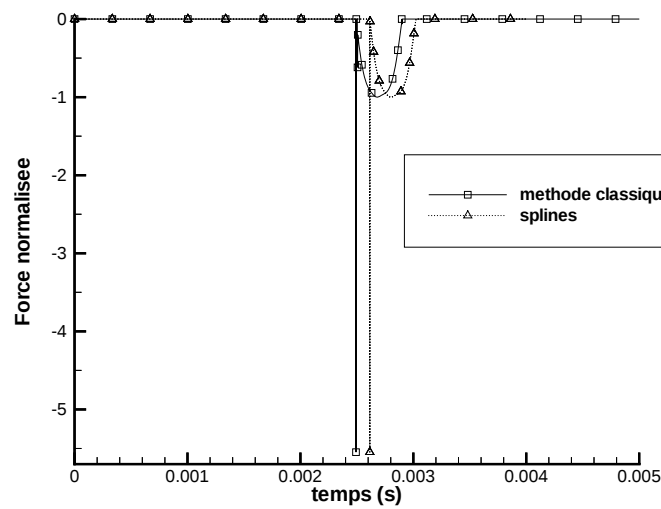
Il s'agit ici de valider la bonne prise en compte de la direction de reprojection (et donc la définition du vecteur normal).

Ce second test reprend le système du paragraphe IV.2.2.1, mais cette fois, le maillage se présente de telle façon que l'aube est dirigée vers une arête du maillage du carter (voir figure IV.16).

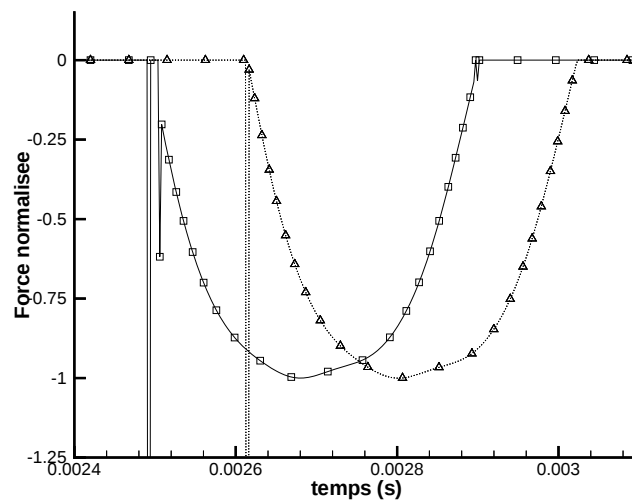
Les figures IV.17 comparent les évolutions de la force de contact calculée par la méthode classique puis par une méthode spline d'interpolation (splines de degré 3) au cours du mouvement.

Dans la direction $\vec{x}$ (figure IV.IV.2.2.2), la force de contact obtenue par une méthode classique paraît bruitée: les pics observés correspondent non pas à des pertes de contact, mais à des changements d'orientation de la vitesse dus au fait que la facette sur laquelle la reprojection se fait n'est pas toujours la même (brusques variations de la normale: voir figure IV.IV.2.2.2). Ce problème est clairement visible sur la figure IV.IV.2.2.2: dans la direction $\vec{y}$, la force change constamment de signe.

La méthode utilisant les splines est insensible à ce type de problème, puisque la normale ne varie pas d'un pas de temps à l'autre (voir figure IV.IV.2.2.2). Il en résulte une valeur



(a) Force de contact normalisée



(b) Détail

Figure IV.14 : Force de contact selon les différentes méthodes

des forces obtenues plus conforme à l'intuition; en particulier, dans la direction $\vec{y}$, la force de contact doit être nulle. Dans la pratique, elle ne l'est pas complètement, car le vecteur normal n'est pas exactement colinéaire au vecteur $\vec{x}$. Cependant, le rapport des forces

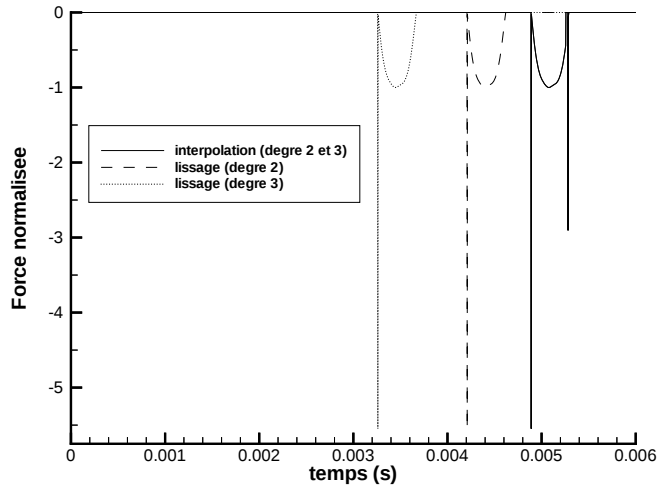


Figure IV.15 : Force de contact pour des splines de degré 2 et 3

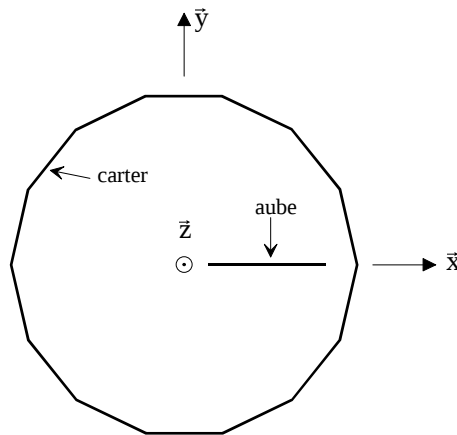


Figure IV.16 : Second cas test sous PLEXUS (vue de dessus)

dans les deux directions $\vec{x}$ et $\vec{y}$ est supérieur à 100 avec les splines: il est inférieur à 4 avec la méthode classique.

IV.2.3 Conditionnement des matrices

IV.2.3.1 Rappels théoriques

Soit x la solution du système linéaire $\mathbb{A}x = b$, où $\mathbb{A}$ est une matrice de $\mathbb{R}^{N \times N}$. Il est possible de définir un facteur de conditionnement, généralement noté $\kappa(\mathbb{A})$, qui permet de quantifier l'influence d'une petite perturbation du système, tant sur les éléments de $\mathbb{A}$ que sur le second membre, sur la réponse [GOL89]. Ainsi, si la matrice varie de $\delta\mathbb{A}$, et le

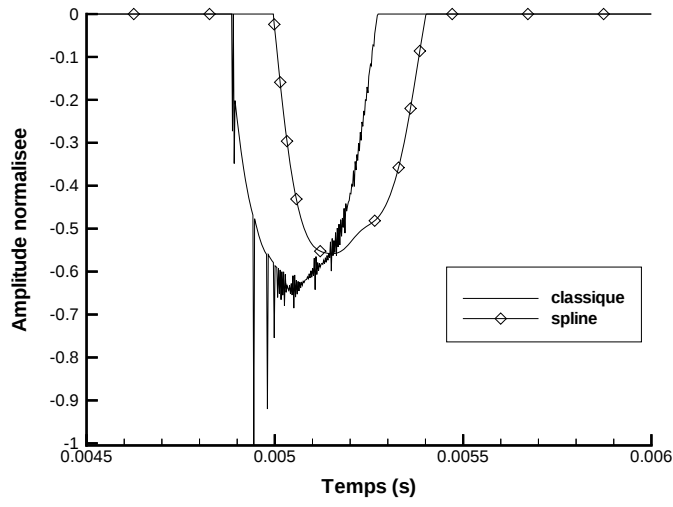
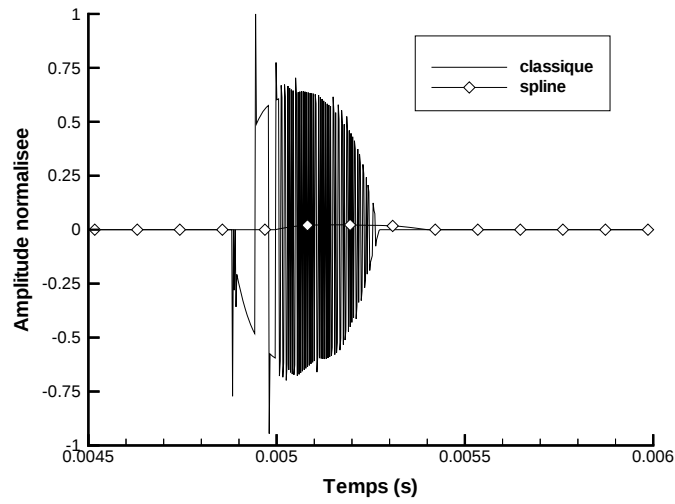
(a) Direction $\vec{x}$ (b) Direction $\vec{y}$

Figure IV.17 : Forces de contact comparées

second membre de δb , alors:

$$\frac{\|\delta x\|}{\|x\|} \leq \kappa(\mathbb{A}) \left(\frac{\|\delta \mathbb{A}\|}{\|\mathbb{A}\|} + \frac{\|\delta b\|}{\|b\|} \right) + \text{termes du second ordre}$$

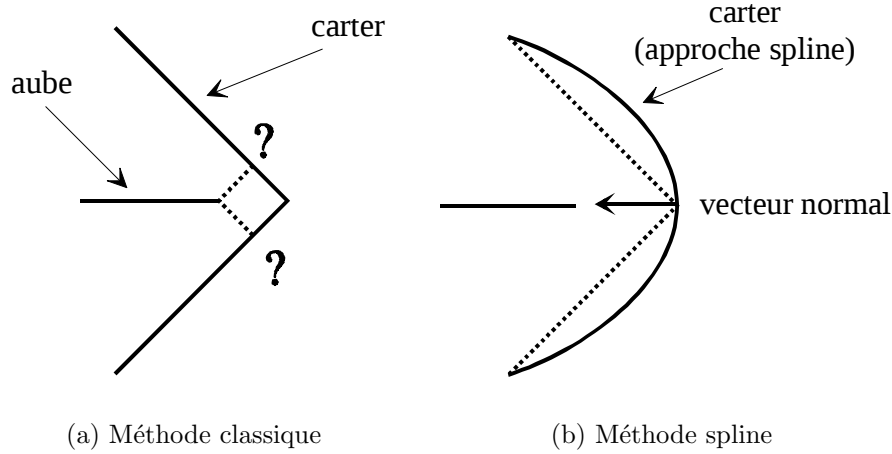


Figure IV.18 : Reprojection

La valeur $\kappa(\mathbb{A})$ est définie par:

$$\kappa(\mathbb{A}) = \|\mathbb{A}\| \cdot \|\mathbb{A}^{-1}\|$$

où $\|\cdot\|$ est une norme matricielle. Le choix de la norme conditionne la valeur de $\kappa(\mathbb{A})$, cependant, il est possible de montrer que les facteurs de conditionnement issus de l'emploi de deux normes différentes sont équivalents. En effet, si $\kappa_{N1}(\cdot)$ et $\kappa_{N2}(\cdot)$ désignent les facteurs de conditionnement pour deux normes matricielles N_1 et N_2 , alors il existe deux réels r_1 et r_2 tels que:

$$r_1 \kappa_{N1}(\mathbb{A}) \leq \kappa_{N2}(\mathbb{A}) \leq r_2 \kappa_{N1}(\mathbb{A})$$

Une valeur de $\kappa(\mathbb{A})$ proche de 1 traduit un bon conditionnement, et donc une faible variation de la réponse pour une faible perturbation. Une grande valeur de $\kappa(\mathbb{A})$ traduit au contraire un mauvais conditionnement. Daubisse a donné dans sa thèse [DAU84] une relation entre $\kappa(\mathbb{A})$ et le nombre C_s de chiffres décimaux significatifs de la solution. Si C_{smax} est le nombre de chiffres décimaux significatifs correspondant aux m bits de la mantisse, alors la plus petite perturbation représentable par la machine est égale à: $2^{-m} = 10^{-C_{smax}}$, et la valeur de C_s peut être approchée par:

$$C_s \approx C_{smax} - \log_{10}(\kappa(\mathbb{A}))$$

Dans le cas présent, les mots sont codés en double précision, ce qui correspond à $m = 64bits$, soit $C_{smax} \approx 19$.

IV.2.3.2 Cas de l'interpolation

Dans le cas de l'interpolation, les matrices à inverser sont des matrices du type:

$$\mathbb{A} = [a_{ij}] = [B_{nj}(\xi_i)] \quad \text{où} \quad \mathbb{A} \in \mathbb{R}^{npt \times npt}$$

La partition de l'unité conduit tout naturellement à utiliser la norme infinie pour calculer le facteur de conditionnement. En effet, par définition, la norme infinie d'une matrice est:

$$\|\mathbb{A}\|_{\infty} = \max_{i=1, npt} \sum_{j=1}^{npt} |a_{ij}|$$

ce qui, dans le cas présent donne:

$$\|\mathbb{A}\|_{\infty} = \max_{i=1, npt} \sum_{j=1}^{npt} |B_{nj}(\xi_i)| = 1$$

Il en résulte que $\kappa(\mathbb{A}) = \|\mathbb{A}^{-1}\|_{\infty}$. Malheureusement, il n'existe pas de lien simple entre $\mathbb{A}$ et la norme de sa matrice inverse.

C'est pourquoi quelques cas particuliers ont été envisagés, de façon à souligner quelques tendances sur le conditionnement. Dans le cas où les npt points du jeu de données induisent un paramétrage ξ régulier (c'est-à-dire que pour tout i , $\xi_{i+1} - \xi_i = \text{constante} = 1$ par exemple), et que le degré de modélisation des splines est 1, il est facile de voir que le paramétrage t définissant les B-splines est de la forme:

$$t_0 = t_1 \triangleq \xi_0$$

$$t_i = \xi_{i-1} \text{ pour } i = 2, \dots, npt - 1$$

$$t_{npt} = t_{npt+1} \triangleq \xi_{npt-1}$$

Il en résulte que la quantité $B_{nj}(\xi_i)$ est égale à δ_{ij} ², et par conséquent que la matrice $\mathbb{A}$ est l'identité de $\mathbb{R}^{npt \times npt}$. Ainsi, dans ce cas particulier, $\kappa(\mathbb{A}) = 1$ (le système est très bien conditionné).

La première ligne du tableau IV.1 regroupe les valeurs de $\kappa(\mathbb{A})$ obtenues sur un exemple simple, formé de 20 points équirépartis sur un segment (voir figure IV.IV.2.3.2), pour différents degrés de modélisation. Il y apparaît que plus le degré augmente, plus le conditionnement des matrices est mauvais.

Deux courtes études de sensibilité ont été réalisées:

1. sensibilité au nombre de points dans le jeu de données: 40 points équirépartis sur un segment ont été pris au lieu des 20 initiaux; il apparaît (voir tableau IV.1) que le facteur de conditionnement se révèle peu sensible au nombre de points mis en jeu;
2. sensibilité à la répartition des points: les 20 points utilisés ont été répartis uniformément sur un profil en dent de scie (voir figure IV.IV.2.3.2), puis sur une parabole (la répartition n'étant alors plus uniforme, voir figure IV.IV.2.3.2), et enfin sur une courbe oscillante (répartition très hétérogène des points, voir figure IV.IV.2.3.2).

De ces quelques exemples, il ressort que:

² δ_{ij} est le symbole de Kronecker: il vaut 1 si $i = j$ et 0 sinon

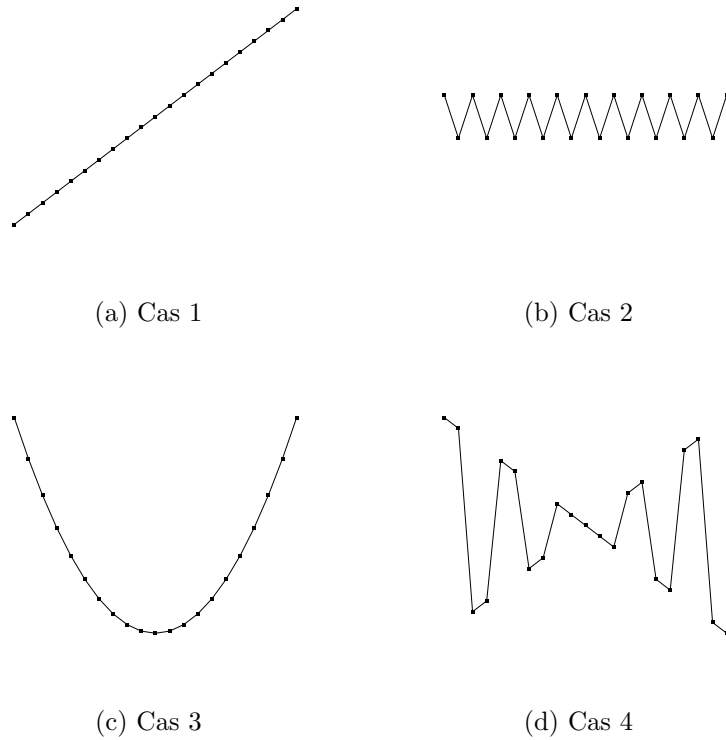


Figure IV.19 : Exemples utilisés pour l'étude de sensibilité

Degré	1	2	3	4	5
cas 1 (20 points)	1	2,6	4,7	10,0	20,9
cas 1 (40 points)	1	2,6	4,7	10,0	20,9
cas 2 (20 points)	1	2,6	4,7	10,0	20,9
cas 3 (20 points)	1	2,8	6,3	30,1	515,4
cas 4 (20 points)	1	2,8	6,3	30,1	515,4

Tableau IV.1: Conditionnement: sensibilité de l'interpolation

- ◇ l'augmentation du degré de modélisation des splines dégrade légèrement le conditionnement des matrices d'interpolation;
- ◇ une répartition non uniforme des points du jeu de données favorise un mauvais conditionnement des matrices d'interpolation.

IV.2.3.3 Cas du lissage

Pour cette méthode, le choix de la norme infinie a été conservé.

Les matrices à inverser sont de terme générique $a_{ij} = \langle B_{ni}, B_{nj} \rangle$, où le produit

scalaire est défini par:

$$\langle B_{ni}, B_{nj} \rangle = \int_{\xi_0}^{\xi_{npt-1}} B_{ni}(t) B_{nj}(t) dt$$

La définition de la norme infinie de $\mathbb{A}$ donne dans le cas présent:

$$\begin{aligned} \|\mathbb{A}\|_{\infty} &= \max_{i=1, npt} \left(\sum_{j=1}^{npt} |\langle B_{ni}, B_{nj} \rangle| \right) \\ &= \max_{i=1, npt} \left(\sum_{j=1}^{npt} \int_{\xi_0}^{\xi_{npt-1}} B_{ni}(t) B_{nj}(t) dt \right) \\ &= \max_{i=1, npt} \left(\int_{\xi_0}^{\xi_{npt-1}} B_{ni}(t) \left(\sum_{j=1}^{npt} B_{nj}(t) \right) dt \right) \\ &= \max_{i=1, npt} \left(\int_{\xi_0}^{\xi_{npt-1}} B_{ni}(t) dt \right) \quad (\text{partition de l'unité}) \end{aligned}$$

Il reste donc à calculer cette dernière intégrale. Par récurrence, il est possible d'obtenir le résultat suivant:

Proposition - Relations d'intégration : *Pour toute B-spline B_{ni} de degré n définie sur l'intervalle $[t_i, t_{i+n+1}]$ par la relation de récurrence de la page 16, les relations suivantes sont vérifiées:*

$$\begin{aligned} \int_{t_i}^{t_{i+n+1}} B_{ni}(t) dt &= \frac{1}{n+1} \cdot (t_{i+n+1} - t_i) \\ \int_{t_i}^{t_{i+n+1}} t \cdot B_{ni}(t) dt &= \frac{1}{(n+1)(n+2)} \cdot (t_{i+n+1} - t_i)(t_i + t_{i+1} + \dots + t_{i+n+1}) \end{aligned}$$

Il en résulte:

$$\|\mathbb{A}\|_{\infty} = \max_{i=1, npt} \left(\frac{1}{n+1} \cdot (t_{i+n+1} - t_i) \right)$$

Cette relation montre qu'il est possible de diminuer la valeur de la norme infinie en travaillant avec des paramétrages adéquats. Mais il apparaît que la norme de la matrice inverse, qu'il n'est, là non plus, pas possible de déterminer simplement, augmente, de sorte que le facteur de conditionnement reste insensible aux changements de paramétrages.

Comme pour le cas de l'interpolation, deux études de sensibilité ont été réalisées: le tableau IV.2 donne les valeurs du facteur de conditionnement pour les configurations représentées sur la figure IV.19 (20 points de contrôle sont demandés).

Il ressort de ces quelques cas:

- ◇ le facteur de conditionnement augmente avec le degré des splines utilisées;
- ◇ les matrices sont d'autant plus mal conditionnées que la répartition des points dans le jeu de données est non uniforme.

Degré	1	2	3	4	5
cas 1 (20 points)	4,7	13,3	37,5	117,9	346,3
cas 1 (40 points)	7,1	15,16	37,5	117,9	1101,0
cas 2 (20 points)	4,7	13,3	37,5	117,9	342,5
cas 3 (20 points)	29,6	45,4	70,7	137,3	445,6
cas 4 (20 points)	29,6	45,4	70,7	137,3	445,6

Tableau IV.2: Conditionnement: sensibilité du lissage

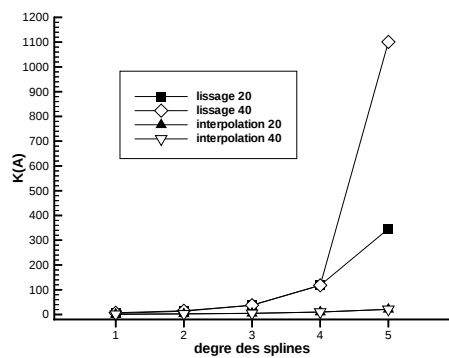
IV.2.3.4 Comparaison interpolation / lissage

Afin de compléter les études de sensibilité réalisées, d'autres calculs ont été menés avec les cas proposés sur la figure IV.19, mais avec des jeux de données de plus grande taille (40 points). L'ensemble des résultats est représenté sur les figures IV.20.

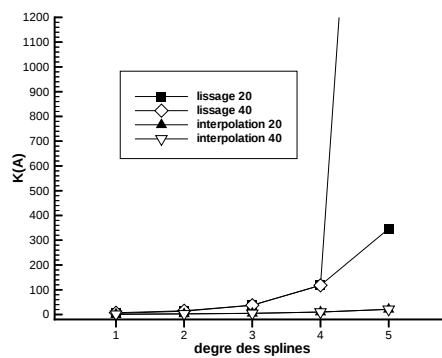
De ces quelques exemples, il ressort que la méthode de lissage semble, à configuration identique, conduire à des matrices moins bien conditionnées que celles obtenues avec l'interpolation. Pour les cas 3 et 4 notamment, le facteur de conditionnement pour des splines de degré 5 atteint respectivement 10^6 et 10^7 . Paradoxalement, sur l'exemple PLEXUS traité, c'est la méthode d'interpolation qui conduit à des instabilités dans le résultat.

Par ailleurs, les facteurs de conditionnement ne dépassent, dans les cas présentés et pour les degrés utilisés dans PLEXUS, jamais 10^3 , ce qui signifie que le nombre C_s de chiffres significatifs de la solution est peu différent de C_{smax} . Ces remarques conduisent à penser que le conditionnement des matrices n'est sans doute pas à l'origine du problème soulevé dans le paragraphe IV.2.2.

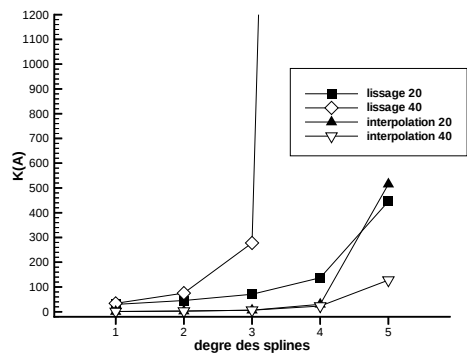
Il faut donc chercher une autre cause possible pour le problème observé. Actuellement, un développement est en cours qui permet de s'attacher plus particulièrement à l'instant de séparation. Ce développement s'appuie sur la prise en compte de l'histoire du phénomène de contact.



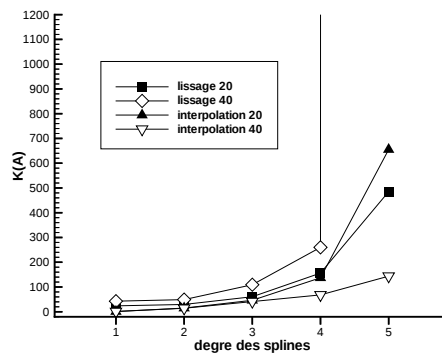
(a) Cas 1



(b) Cas 2



(c) Cas 3



(d) Cas 4

Figure IV.20 : Sensibilités comparées pour des jeux de données de 20 et 40 points

Conclusion

Une méthode de recherche du contact entre une courbe et une surface représentées par des fonctions splines a été développée. Validée sur des exemples statiques, elle a également été insérée dans le contexte d'un calcul dynamique par éléments finis, et ce, dans deux configurations (implicite et explicite). Son rôle est de mieux traiter les problèmes inhérents à la facétisation, notamment les problèmes de discontinuité du vecteur normal à la frontière des éléments, ainsi que les problèmes liés à l'approximation géométrique des structures.

Cette méthode repose sur un principe géométrique qui consiste à faire glisser une boule de rayon de plus en plus petit le long de la courbe, et à regarder l'évolution de la trace laissée par cette boule sur la surface. Le choix d'un bon pas de balayage se révèle important. Certaines des données transmises au code sont avantageusement calculées grâce à la formulation spline, ce qui garantit le caractère de continuité qui fait défaut aux méthodes classiques.

L'implémentation dans Samcef, via la modification d'un élément utilisateur existant, a montré que cette méthode, particulièrement lorsque l'interpolation est utilisée, se révélait peu sensible à la finesse du maillage de la zone de contact de la structure étudiée. Par là même, elle offre la possibilité de réduire les temps de calcul en utilisant des maillages comportant moins d'éléments, ce qui compense le surcoût qu'elle engendre dans ce domaine.

L'implémentation dans Plexus a donné lieu au développement de commandes spécifiques intégrées au code, qui seront prochainement rendues accessibles à tous les utilisateurs. Cette phase de développement requiert encore quelques travaux, notamment pour lever des problèmes sans doute liés à la mauvaise gestion de la séparation des structures en contact.

Partie II

Approche expérimentale

Dans cette partie, une approche expérimentale du contact est présentée. Un banc d'essai a été développé, dont la description fait l'objet du paragraphe V.1: il s'agit de réaliser des essais de contact entre une aube et un carter. L'analyse modale du banc, détaillée dans le paragraphe V.2, a permis de mettre en place un modèle éléments finis, qui est ensuite recalé (paragraphe V.2.4). Les essais de contact, qui cherchent à reproduire un phénomène d'interaction, sont présentés dans le chapitre VI.

Chapitre V

Banc d'essai et analyse modale

Ce chapitre contient dans un premier temps (paragraphe V.1) la description du banc d'essai: éléments constitutifs et chaîne de mesure. Dans un deuxième temps, une analyse modale du banc par vélocimétrie laser est réalisée, en utilisant notamment une propriété des modes des structures à symétrie de révolution (paragraphe V.2). Enfin, les résultats obtenus permettent de recalibrer un modèle éléments finis de l'élément principal (le carter), le recalage comprenant une première phase d'analyse fine de la géométrie (paragraphe V.2.4.2) et une seconde phase d'étude d'influence à l'aide de la méthode des plans d'expérience (paragraphe V.2.4.3).

V.1 Description du banc d'essai

Le banc d'essai comporte quatre parties:

1. un tour vertical dont les caractéristiques sont rassemblées dans l'annexe G;
2. l'ensemble des pièces étudiées: il s'agit d'une partie de carter et de son support;
3. le dispositif associé à l'utilisation du vélocimètre laser;
4. le système d'acquisition.

V.1.1 Le carter

Le carter à étudier est en fait une partie de carter de turboréacteur fournie par SNECMA. Cette pièce se présente sous la forme d'une partie tronconique dont l'angle de conicité est évalué à 9 degrés, et à la base de laquelle est située une bride percée de 72 trous (voir figure V.1). Ce carter est composé d'un matériau au carbure de tungstène à base de cobalt et de chrome, référencé KC20WNx (norme Afnor) (voir annexe G).

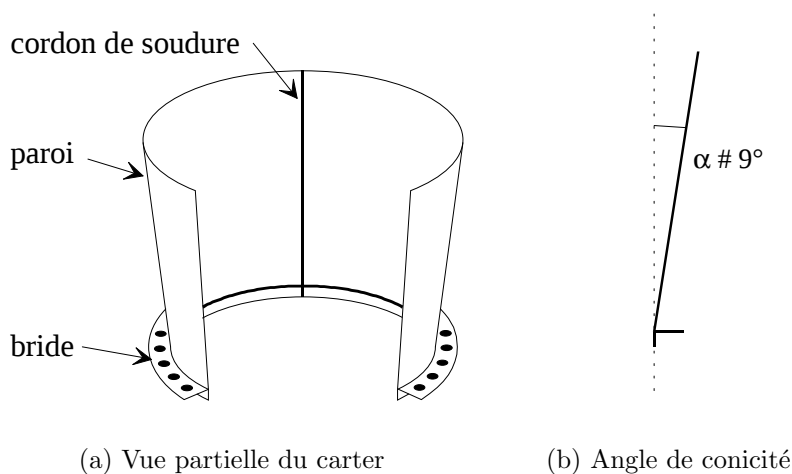


Figure V.1 : Carter

Le diamètre de base du carter étant supérieur au diamètre du plateau du tour vertical qui l'accueille, une pièce intermédiaire a dû être construite pour assurer la liaison. Cette pièce est une couronne aménagée pour accueillir la bride (perçement de 72 trous), et possédant huit points de serrage équirépartis sur le plateau du tour (voir figure V.2). Le montage sur le plateau du tour a été assuré via un serrage à couple constant avec clef dynamométrique des huit points d'ancrage.

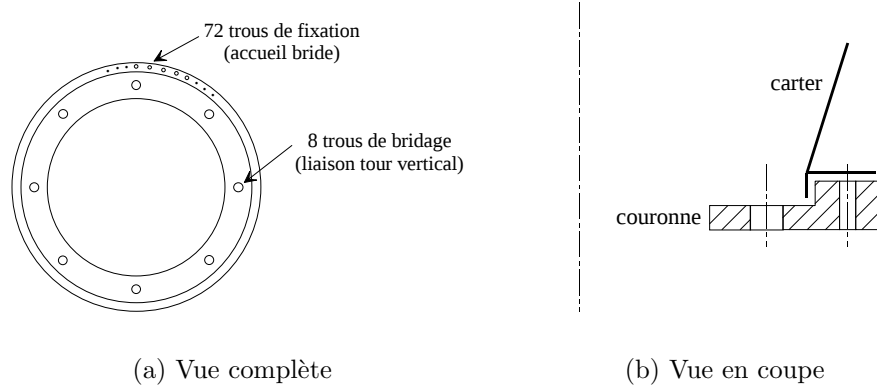


Figure V.2 : Couronne

V.1.2 Le montage du laser

Le principal appareil de mesure utilisé est un vélocimètre laser Polytec OFV 2200. Le principe général de fonctionnement de l'appareil, détaillé dans l'annexe H, repose sur l'utilisation du phénomène Doppler appliqué aux fréquences optiques: le décalage en fréquence qu'introduit une structure en mouvement sur un rayon lumineux est proportionnel à la vitesse de déplacement de cette structure dans la direction de mesure.

D'un point de vue pratique, le vélocimètre va tout d'abord servir à la recherche des modes du carter, avant d'être utilisé pour l'acquisition de la réponse au contact de l'aube.

V.1.2.1 Acquisition modale et vélocimétrie

Cette section présente une méthode de recherche des modes de vibration (fréquence et déformée) d'une structure axisymétrique à l'aide d'un vélocimètre laser. Stanbridge [STA95] l'a mise en œuvre sur des disques, et Billet et Moreno [BIL96] sur un cylindre de révolution. Elle sera ici appliquée à un tronc de cône. Cette méthode utilise le fait que chaque mode de vibration de telles structures est, comme cela sera rappelé dans la section V.2.1, caractérisé par le nombre n de ses diamètres nodaux.

Soit Φ_{1n} et Φ_{2n} les deux déformées modales du mode à n diamètres, et ϕ un paramètre de repérage angulaire tel que:

$$\begin{aligned}\Phi_{1n}(\phi) &= \cos(n(\phi - \phi_{0n})) \\ \Phi_{2n}(\phi) &= \sin(n(\phi - \phi_{0n}))\end{aligned}$$

où ϕ_{0n} caractérise l'orientation du mode. Soit F une force ponctuelle appliquée radialement au point repéré par l'angle ϕ_F (voir figure V.3):

$$F(\phi, t) = F(t) \cdot \delta(\phi - \phi_F)$$

Les forces généralisées $\mathcal{F}_{1n}$ et $\mathcal{F}_{2n}$ associées à F sont données par:

$$\mathcal{F}_{in}(\phi, t) = \int_{\phi} F(\phi, t) \cdot \Phi_{in}(\phi) d\phi \quad \text{pour } i \in \{1, 2\}$$

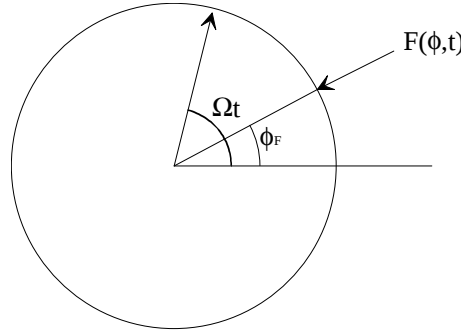


Figure V.3 : Repérage angulaire

soit:

$$\begin{aligned}\mathcal{F}_{1n}(\phi, t) &= F(t) \cdot \cos(n(\phi_F - \phi_{0n})) \\ \mathcal{F}_{2n}(\phi, t) &= F(t) \cdot \sin(n(\phi_F - \phi_{0n}))\end{aligned}$$

Si l'observation ne se fait pas dans le repère fixe, mais dans un repère mobile, en rotation à la vitesse angulaire constante Ω par rapport au repère fixe, alors le point d'application de la force sera repéré par l'angle θ_F défini par:

$$\theta_F = \phi_F - \Omega t$$

Si, de plus, l'excitation $F(t)$ est harmonique, de pulsation ω et d'amplitude F_0 , alors:

$$\begin{aligned}\mathcal{F}_{1n}(\phi, t) &= F_0 \cos(\omega t) \cdot \cos(n(\theta_F + \Omega t - \phi_{0n})) \\ \mathcal{F}_{2n}(\phi, t) &= F_0 \cos(\omega t) \cdot \sin(n(\theta_F + \Omega t - \phi_{0n}))\end{aligned}$$

ce qui, après transformation en somme du produit des fonctions circulaires, donne:

$$\begin{aligned}\mathcal{F}_{1n}(\phi, t) &= \frac{F_0}{2} \{ \cos(\omega t + n\Omega t + n(\theta_F - \phi_{0n})) + \cos(\omega t - n\Omega t - n(\theta_F - \phi_{0n})) \} \\ \mathcal{F}_{2n}(\phi, t) &= \frac{F_0}{2} \{ \sin(\omega t + n\Omega t + n(\theta_F - \phi_{0n})) - \sin(\omega t - n\Omega t - n(\theta_F - \phi_{0n})) \}\end{aligned}$$

Il apparaît ainsi que dans le repère tournant, la réponse de la structure sera la somme de deux oscillations pulsées respectivement à $\omega - n\Omega$ et à $\omega + n\Omega$ (voir figure V.4). Cette propriété est exploitée en faisant tourner le faisceau sortant du vélocimètre à l'intérieur de la structure étudiée, de sorte que le spot laser décrive un grand diamètre de celle-ci (voir figure V.5). L'analyse de la réponse en fréquence du vélocimètre permet, connaissant la vitesse de rotation du faisceau, d'associer une valeur de n à chaque fréquence propre identifiée.

Les fréquences propres peuvent pour leur part être déterminées soit par une acquisition préalable, au marteau, faisceau arrêté (analyse fréquentielle de la structure soumise à une impulsion), ou en utilisant le principe évoqué ci-dessus, avec une excitation de type bruit blanc, pour différentes vitesses de rotation du faisceau. Cette dernière méthode permet d'obtenir des faisceaux de droites dont les intersections avec l'axe $\Omega = 0$ donnent les fréquences propres (voir figure V.6).

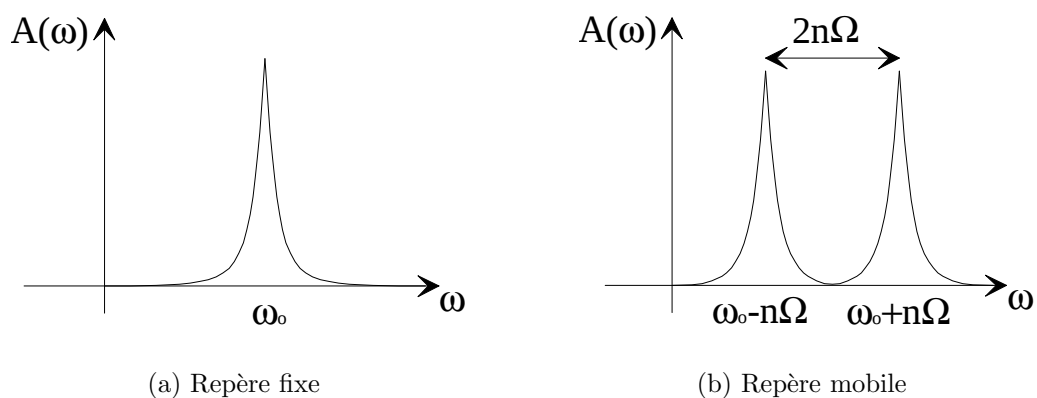


Figure V.4 : Effet de la rotation du point de mesure

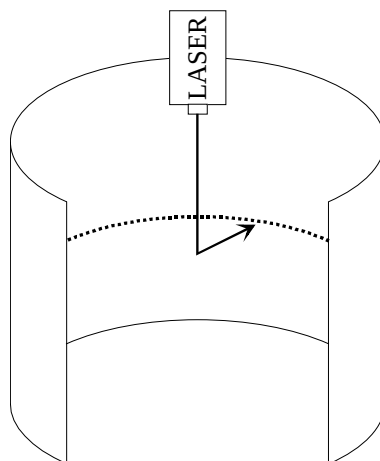


Figure V.5 : Principe de la mesure

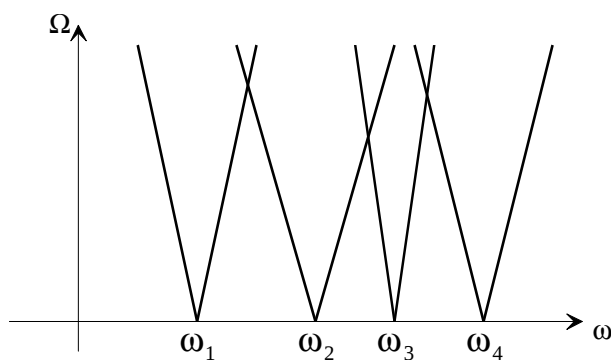


Figure V.6 : Diagramme vitesse de rotation / fréquence

V.1.2.2 Montage pratique et réglages

Afin de mettre en œuvre une méthode d'analyse modale expérimentale reposant sur le décalage en fréquence perçu par le vélocimètre, un certain nombre de dispositions ont

été prises sur le montage.

Tout d'abord, un moteur a été installé sur le plateau du tour vertical, de sorte que son axe de rotation coïncide avec l'axe du tour, et par-là même, avec l'axe présumé de la pièce (voir figure V.7). La vitesse de rotation de l'arbre de sortie est pilotable entre 0 et 2000 tr/min par le manipulateur à l'aide d'un potentiomètre. Aucun affichage ne permet cependant de connaître la vitesse choisie (nécessité d'effectuer des mesures complémentaires). Sur l'arbre de sortie du moteur, un petit dispositif permet de fixer un miroir face avant qui renverra le faisceau du vélocimètre vers la paroi du carter. Ce miroir face avant ne génère pas de dédoublement du faisceau au passage de la couche de verre, comme le ferait un miroir classique, puisque la couche réfléchissante n'est pas protégée. L'orientation du miroir est telle que le faisceau est renvoyé perpendiculairement à la paroi (voir figure V.V.1.3).

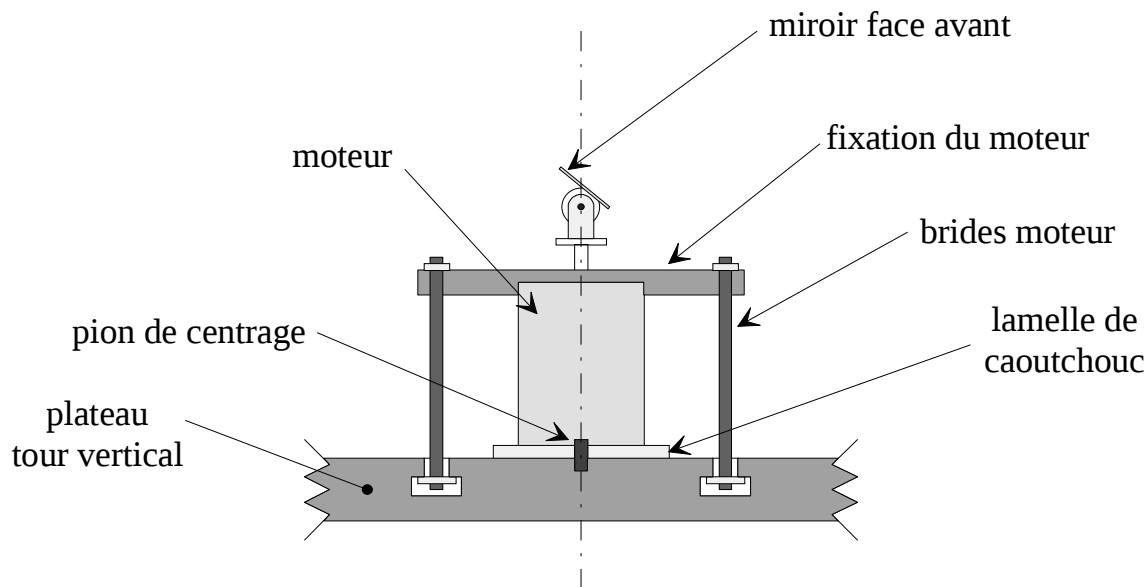


Figure V.7 : Fixation du moteur

Le corps du vélocimètre est pour sa part monté sur un portique qui surplombe le plateau du tour vertical, de sorte que le faisceau soit colinéaire à l'axe du tour (voir figure V.8). Le portique est isolé du sol par l'intermédiaire de lamelles en caoutchouc. Le réglage fin du centrage du laser se fait notamment via l'utilisation d'une pointe biseautée semi-réfléchissante placée au centre du plateau du tour (voir figure V.9).

Toutes ces précautions sont nécessaires pour limiter un certain nombre d'effets parasites qu'engendrerait la non coaxialité du faisceau du laser et de l'axe de rotation du moteur supportant le miroir. Ces effets ont été notamment décrits par Rothberg et Halliwell [ROT94].

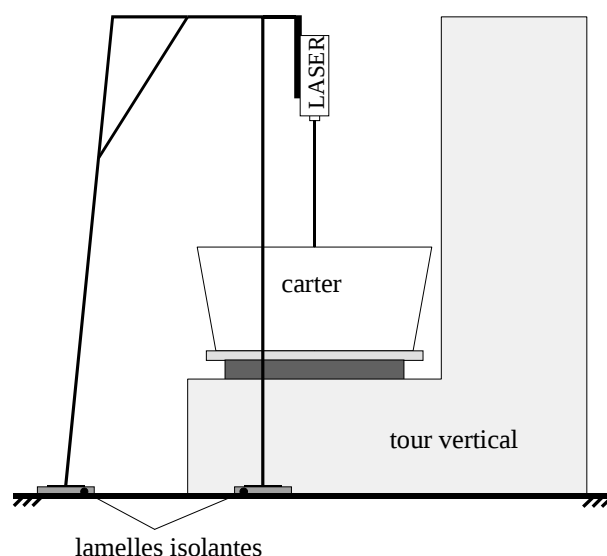


Figure V.8 : Montage du vélocimètre

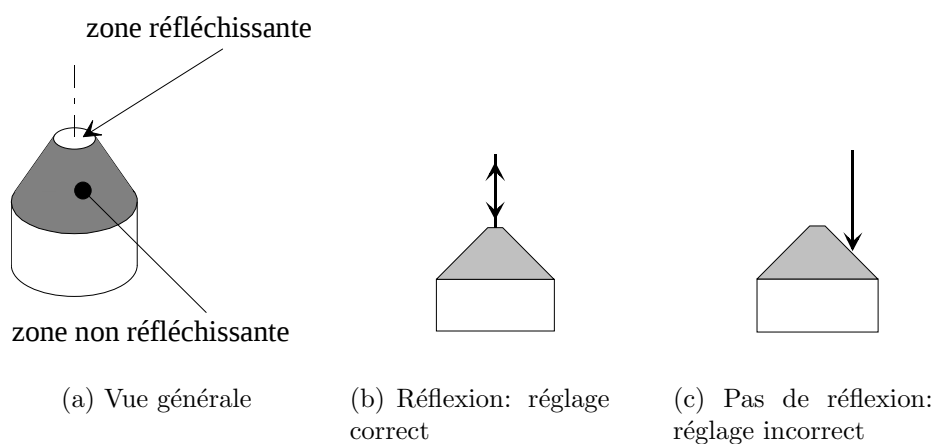


Figure V.9 : Pion de réglage du laser

V.1.3 Chaîne d'acquisition

Le carter est sollicité par l'intermédiaire d'un excitateur B&K (type 4809) fixé en haut de la paroi, de sorte que l'excitation soit perpendiculaire à celle-ci (voir figure V.V.1.3). Un capteur de force B&K (type 8200) lié à l'excitateur permet d'accéder au signal effectivement délivré dans la structure. Le signal issu du vélocimètre est préalablement amplifié par un amplificateur de charge B&K (type 2635) puis transmis, avec le signal du capteur de force, à un analyseur deux voies HP3562A. Enfin, un petit accéléromètre B&K (type 4367) a été également occasionnellement utilisé. L'ensemble de la chaîne d'acquisition est représentée sur la figure V.11.

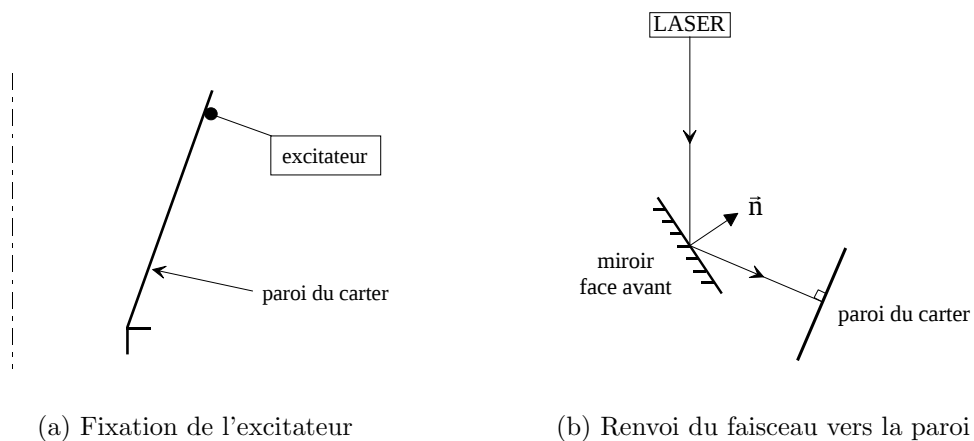


Figure V.10 : Montage de l'excitateur et miroir de renvoi

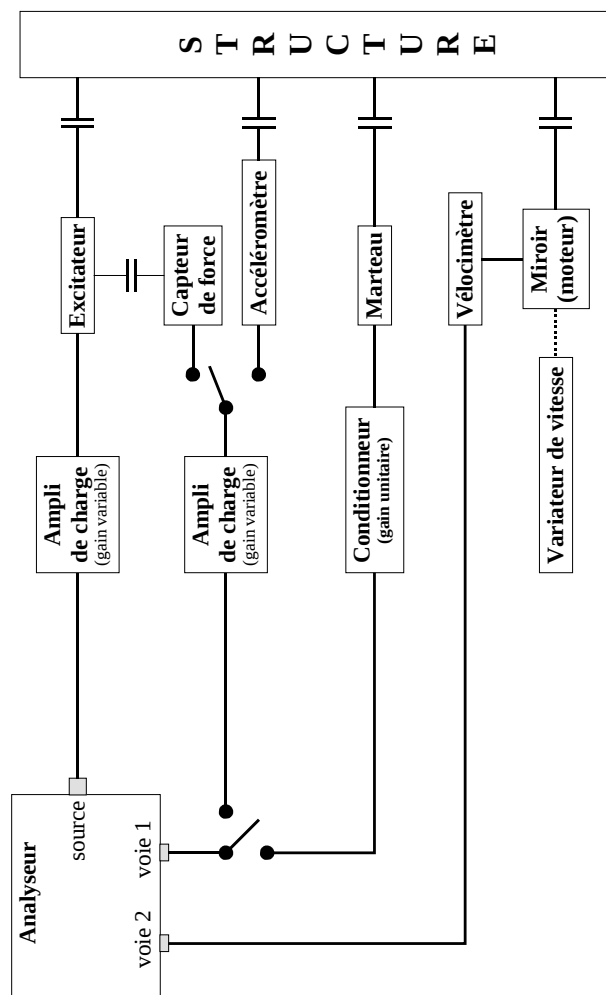


Figure V.11 : Chaîne d'acquisition

V.2 Identification modale du carter

V.2.1 Modes des structures axisymétriques

La mise en équations des coques de révolution, dans le cadre de l'hypothèse des petites perturbations, a donné lieu à un grand nombre de travaux. Leissa [LEI73] a consacré un ouvrage entier à ce problème, dans lequel il recense les résultats de nombreux chercheurs tels que Love, Timoshenko, Flügge, Novozhilov... Il s'attache notamment à comparer les différentes formulations, en particulier les hypothèses utilisées pour établir la mise en équations, ainsi que les résultats auxquels ces méthodes conduisent sur des exemples *simples* tels que la sphère, le cylindre ou le cône.

Les équations locales dynamiques utilisées pour rechercher les modes de ces structures peuvent se mettre sous la forme:

$$\mathcal{M}(\ddot{u}) + \mathcal{K}(u) = 0$$

où $\mathcal{M}$ et $\mathcal{K}$ sont des opérateurs de masse et de raideur, et u représente les déplacements radiaux, orthoradiaux et longitudinaux d'un point du plan moyen.

De façon générale, lorsque la structure est parfaitement axisymétrique, les déformées modales associées à la composante radiale sont doubles, c'est-à-dire qu'à une même fréquence propre correspond deux déformées modales différentes, et sont de la forme (θ désigne un paramètre de position angulaire, voir figure V.12):

$$\Phi_n = \begin{Bmatrix} \Phi_{1n} \\ \Phi_{2n} \end{Bmatrix} = \begin{Bmatrix} \cos(n(\theta - \theta_{0n})) \\ \sin(n(\theta - \theta_{0n})) \end{Bmatrix}$$

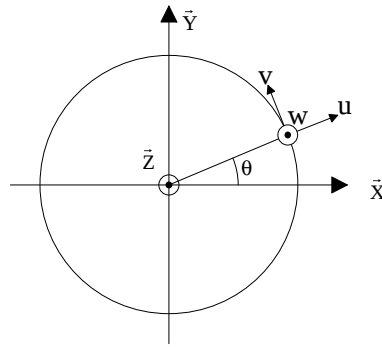


Figure V.12 : Plan méridien d'une structure de révolution

Physiquement, l'entier n représente le nombre de diamètres nodaux du mode, ou encore le nombre de lobes de la déformée (voir figure V.13). Le réel θ_{0n} définit l'orientation du mode. Lorsque la structure est parfaitement de révolution, ce réel n'est pas défini (aucune orientation privilégiée pour les modes), et les deux modes échangent leurs nœuds et leurs ventres de vibration.

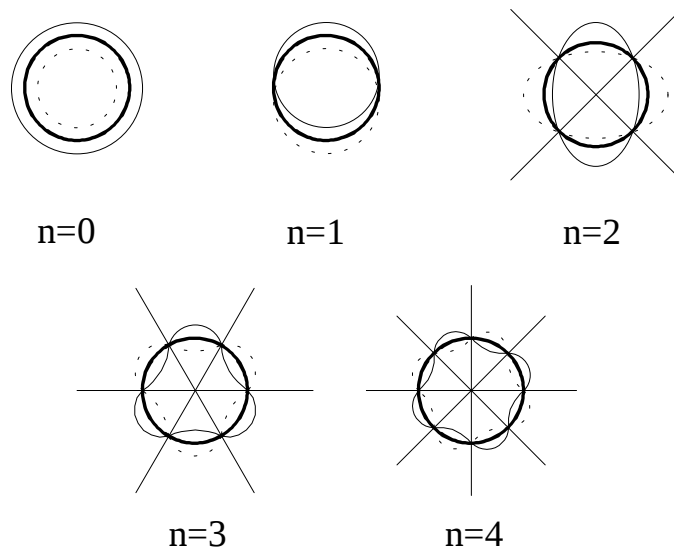


Figure V.13 : Exemples de déformées modales

Dans la pratique, les structures étudiées ne sont jamais parfaitement de révolution: que ce soit pour des raisons géométriques (le rayon dépend de θ par exemple) ou structurales (hétérogénéité du matériau composant la structure, existence de raideurs locales), l'axisymétrie parfaite n'existe pas. Les conséquences de la rupture de symétrie (il est fréquent de parler de désaccordage) est une dissociation des déformées modales, ainsi qu'un dédoublement fréquentiel, chaque déformée étant liée à une fréquence propre différente de l'autre.

Ce phénomène a lui aussi donné lieu à de nombreuses études, depuis la fin du siècle dernier puisque Zenneck a montré en 1899 [ZEN99] que les modes se trouvaient affectés par la présence d'un défaut de symétrie. Ewins a également consacré une partie de ses travaux au désaccordage. Il observe [EWI69] qu'en général, si le défaut est faible, seul le dédoublement de fréquence apparaît (phénomène de *split*, voir figure V.14) sans que la forme globale des modes ne soit affectée, bien que des variations d'amplitude puissent être observables d'un lobe à l'autre.

Si le défaut est plus important, le dédoublement fréquentiel devient très grand, et les déformées modales peuvent présenter des allures très différentes des formes originales.

V.2.2 Démarche générale

Dans un premier temps, la recherche de modes se fait tour à l'arrêt, c'est-à-dire sans que le moteur du tour soit mis en route. La liaison entre le plateau du tour et la couronne est assurée par un bridage en huit points, mais le diamètre de la couronne est légèrement inférieur à celui du plateau. Par ailleurs, le plateau ne se présente pas sous la forme d'une surface plane: il s'agit d'un disque comportant 16 rainures, destinées à accueillir les systèmes de bridage. Il en résulte que la couronne n'est pas en appui sur toute sa surface.

La démarche proposée pour réaliser l'acquisition se décompose en deux temps:

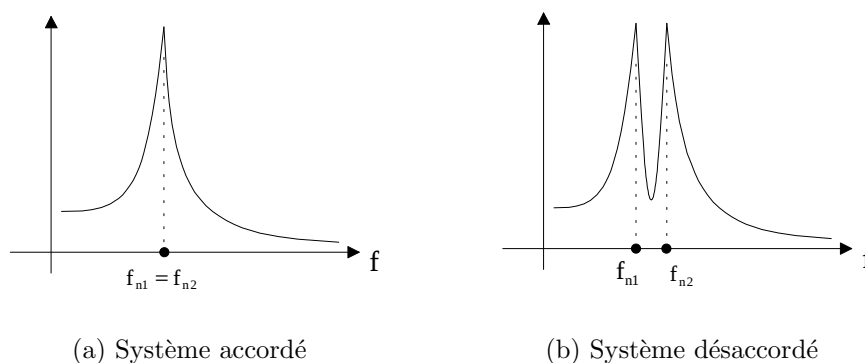


Figure V.14 : Influence de la rupture de symétrie sur la fréquence propre d'un mode

1. recherche des fréquences propres;
2. identification des modes.

La recherche des premières fréquences propres (entre 0 et 300 Hz) a tout d'abord été réalisée en analysant les réponses du carter à un bruit blanc, le vélocimètre visant un point fixe. Il est rapidement apparu le problème suivant: selon le point de visée, la réponse en fréquence varie, ce qui est attendu, mais les spectres ne présentent pas les résonances d'amplitude aux mêmes fréquences. Il en résulte une certaine confusion quant à l'interprétation des résultats obtenus, car il n'est pas possible de déterminer une fréquence propre sans lui attribuer un intervalle d'incertitude qui rend très difficile l'obtention de la déformée. Le faible amortissement de la structure, ainsi que la forte densité modale observée en début de spectre ne font qu'amplifier ce problème.

Un autre procédé de recherche a donc été mis en œuvre: la sollicitation du carter se fait par un balayage sinus lent. Avec un incrément fréquentiel petit (de l'ordre de 0,1 Hz), le carter est excité pendant 20 secondes, puis la mesure est effectuée en procédant à une moyenne sur 200 cycles. L'ensemble est piloté par un analyseur SI 1250, avec un logiciel Solartron. Cette méthode, plus longue, nécessite également davantage d'attention. En effet, la structure est sollicitée sur des fréquences dont certaines, en raison du faible incrément fréquentiel du balayage, sont très voisines de fréquences propres. Compte-tenu du faible amortissement, il y a risque pour l'intégrité de la structure d'une part, risque limité par le choix du niveau d'excitation, et pour l'intégrité de la liaison carter/excitateur de l'autre, liaison qu'il a donc fallu surveiller au cours des mesures.

L'identification des modes associés aux fréquences propres obtenues se fait ensuite en excitant le carter successivement sur chacune de ces fréquences. En utilisant le vélocimètre comme décrit dans le paragraphe V.1.2.1, il est alors possible de déterminer le nombre de diamètres de chacun des modes.

V.2.3 Résultats

La figure V.15 montre un exemple de courbe obtenue à la suite du balayage en fréquence: l'amplitude et la phase sont représentées ensemble. Les deux pics d'amplitude sont espacés de moins de 1 Hz, ce qui traduit bien le faible désaccordage de la structure pour ce mode particulier. Le tableau V.1 rassemble les valeurs des fréquences propres identifiées ainsi que le type de modes auxquels elles correspondent: il apparaît que pour certains modes, le désaccordage se traduit par des dédoublement fréquentiels importants (modes à 4 et 5 diamètres par exemple).

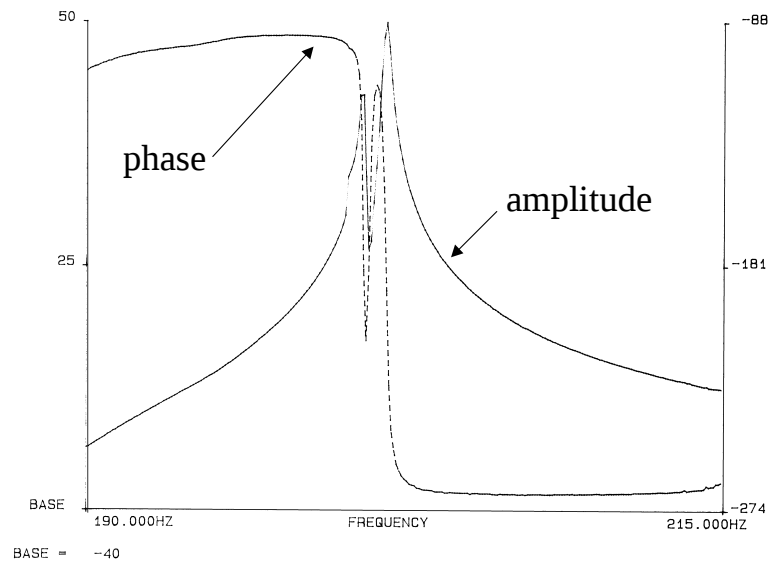


Figure V.15 : Réponse du carter à un balayage sinus (extrait)

D'autres essais ont ensuite été réalisés:

- ◇ tout d'abord en mettant le moteur du tour vertical en route (rotation du plateau pour différentes vitesses): les nouvelles mesures montrent que ce facteur n'a pas d'influence sur le comportement modal du carter (en particulier, la nature des conditions aux limites n'est pas modifiée);
- ◇ ensuite en réalisant des mesures dans la hauteur du carter, à l'aide d'un accéléromèrienctre, afin de s'assurer que les modes, sur la plage de fréquence étudiée, ne présentaient pas de nœuds dans la direction longitudinale: les résultats pour le mode à 10 diamètres figurent sur la figure V.16, les résultats pour les autres modes identifiés étant similaires.

V.2.4 Modèle éléments finis

Afin d'établir un modèle éléments finis du banc, un maillage a été réalisé à l'aide du logiciel SAMCEF. Ce modèle comprend le carter, la couronne sur laquelle il est fixé,

Fréquence propre (Hz)	Nature du mode (ordre de Fourier)
143,10	6 ou 7
144,82	6
147,38	7
153,58	8
154,53	8
155,92	5
162,36	5
172,31	9
174,00	9
200,48	10
201,48	10
218,07	4
235,08	11
235,70	11
261,78	4
273,94	12
275,74	12

Tableau V.1: Résultats modaux expérimentaux

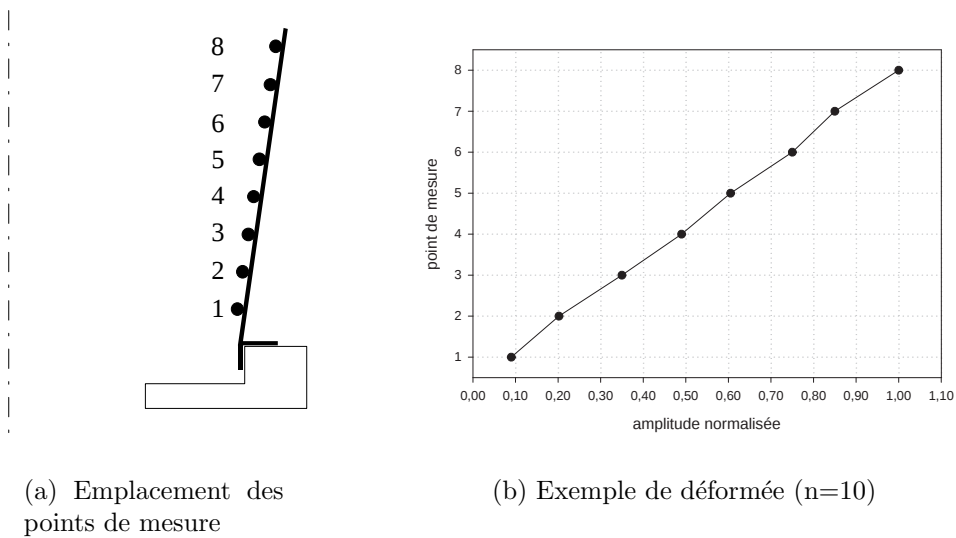


Figure V.16 : Déformée du carter dans le sens longitudinal

et l'excitateur. Les caractéristiques géométriques ont été obtenues par des mesures sur les structures, et les caractéristiques matériau ont été données par les fournisseurs. Les

éléments finis utilisés sont des éléments de Mindlin de trois sortes [SAM97]:

- ◇ éléments poutre (excitateur): type 22, deux nœuds, 6 ddl par nœud (trois rotations, trois translations);
- ◇ éléments coque (carter): type 28 ou 29, trois ou quatre nœuds, 6 ddl par nœud (trois rotations, trois translations);
- ◇ éléments brique (couronne): type 8, huit nœuds, 3 ddl par nœud (trois translations).

V.2.4.1 Premier maillage

Un premier maillage a été réalisé en supposant que la symétrie de révolution de l'ensemble carter + couronne était parfaite (voir figure V.17). Cette hypothèse, justifiée pour la couronne dont le tournage a été fait sur le tour même, et qui n'a pas ensuite été manipulée, est plus hasardeuse en ce qui concerne le carter, puisque celui-ci a été déplacé et manipulé un grand nombre de fois avant d'être installé sur le tour.

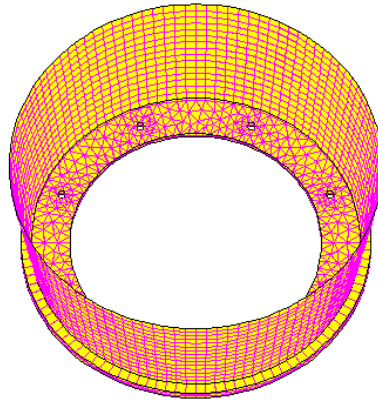


Figure V.17 : Maillage du carter (symétrie de révolution)

Les conditions aux limites sont rassemblées sur la figure V.18:

- ◇ les degrés de liberté de translation entre le carter et la couronne, au niveau de la bride (zone 2 de la figure V.V.2.4.1), sont liés;
- ◇ l'excitateur est encastré à sa base (zone 1 de la figure V.V.2.4.1): tous les degrés de liberté sont bloqués;
- ◇ la liaison avec le tour est assurée en bloquant tous les degrés de liberté de translation dans le plan de la couronne (plan $(\vec{x}, \vec{y})$) dans la zone 3a (voir figure V.V.2.4.1), ainsi que tous les degrés de liberté des nœuds situés au niveau des points d'ancrage (zone 3b de la figure V.V.2.4.1).

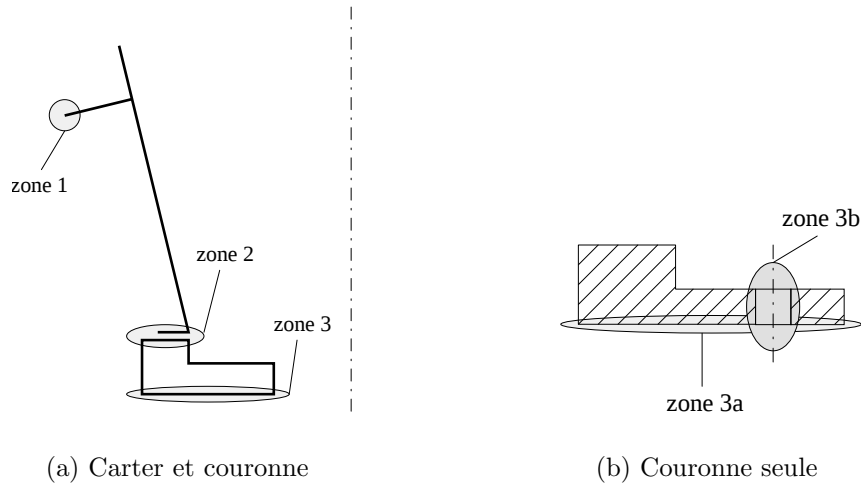


Figure V.18 : Conditions aux limites

Les résultats obtenus (voir tableau V.2) sont conformes aux hypothèses formulées quant à la symétrie de révolution de la structure: les modes sont à diamètres, doubles, le décalage en fréquence étant négligeable (à l'exception du mode à 4 diamètres), et les deux déformées modales échangent leurs nœuds et leurs ventres de vibration.

Le mode à 4 diamètres est très désaccordé: l'origine de ce désaccordage réside dans la condition de fixation du carter sur le tour: les huit points d'ancrage équirépartis définissent quatre diamètres de points *fixes* pour la couronne. Cette fixation définit en fait une orientation particulière pour l'ensemble des modes du carter. Les deux modes à 4 diamètres notamment vont s'orienter de la façon suivante:

- ◇ les diamètres nodaux du premier vont coïncider exactement avec les diamètres de points fixes de la couronne, aucun ventre de vibration n'est contraint;
- ◇ le second mode échange nœuds et ventres de vibration avec le premier mode, ce qui signifie que tous les ventres de vibration de ce second mode vont coïncider avec les diamètres de points fixes de la couronne: ceci est équivalent à l'introduction d'une condition aux limites beaucoup plus contraignante pour ce mode, et entraîne une augmentation importante de sa fréquence propre ainsi qu'une perturbation de la déformée modale (voir figure V.19).

Ce phénomène apparaît pour le mode à 4 diamètres, et dans une proportion beaucoup plus faible, pour le mode à 8 diamètres.

La comparaison entre les résultats éléments finis et les résultats expérimentaux est présentée sur la figure V.20: chaque point (x, y) a pour abscisse une fréquence propre issue de l'expérience, et pour ordonnée la fréquence propre associée dans le modèle éléments finis. Pour chaque point, la corrélation expérience/calcul est d'autant meilleure que la distance à la droite $y = x$ est faible.

Fréquence propre (Hz)	Nature du mode (ordre de Fourier)
153,80	7
153,87	7
159,90	8
161,03	6
161,04	6
161,07	8
179,02	9
179,03	9
183,55	5
183,80	5
206,98	10
206,98	10
210,09	4
242,19	11
242,20	11
283,14	12
283,23	12
294,03	4

Tableau V.2: Résultats modaux éléments finis

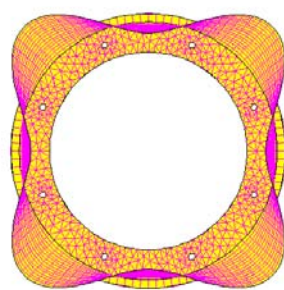
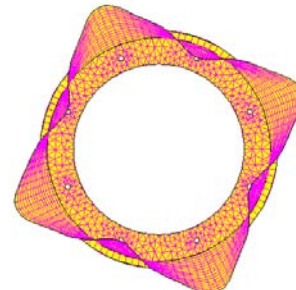
(a) Déformée modale à
210 Hz(b) Déformée modale à
294 Hz

Figure V.19 : Mode à 4 diamètres (vue de dessus)

V.2.4.2 Prise en compte du défaut

Les différences importantes observées entre le modèle éléments finis et les résultats d'expérience conduisent à reconsidérer l'hypothèse de modélisation axisymétrique, en particulier pour le carter. C'est pourquoi un relevé géométrique précis de la paroi a été réalisé.

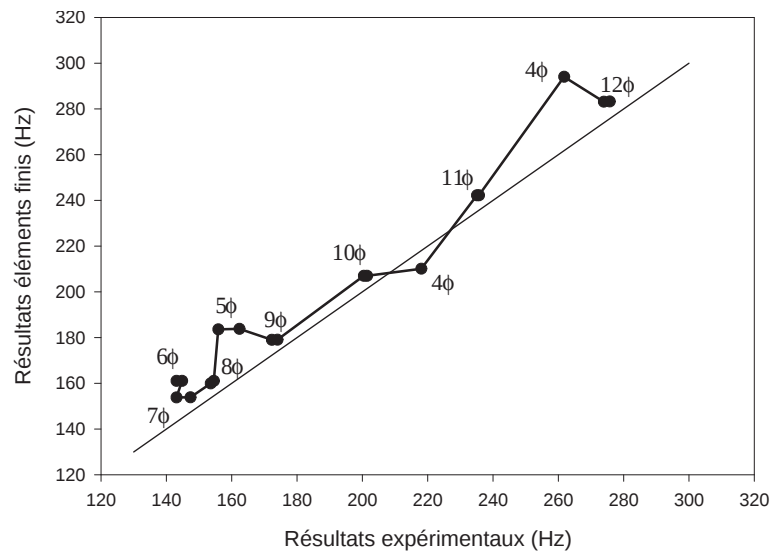
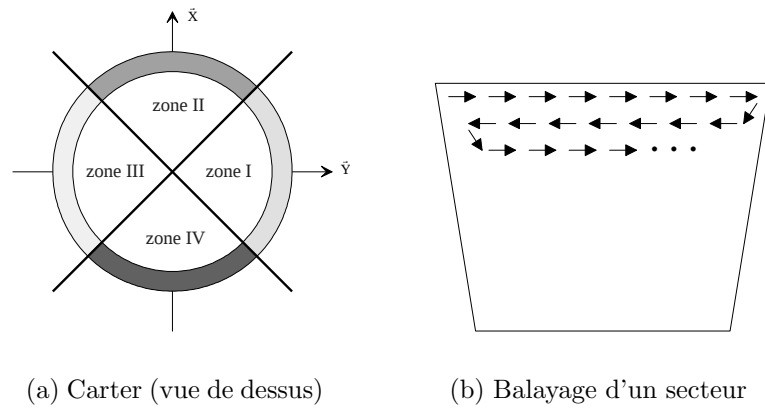


Figure V.20 : Comparaison éléments finis/expérience (maillage sans défaut)

TOPOGRAPHIE DU CARTER

A l'aide d'une machine à mesurer tridimensionnelle (Johansson Cordimatic 1200), une carte topographique fine a été dressée, selon le processus suivant: quatre zones de balayage ont été définies sur la paroi (voir figure V.V.2.4.2), chacune de ces zones ayant été balayée suivant des lignes horizontales en respectant un incrément donné (voir figure V.V.2.4.2).



(a) Carter (vue de dessus)

(b) Balayage d'un secteur

Figure V.21 : Acquisition de la géométrie du carter

Les quatre fichiers de points obtenus ont ensuite été retravaillés afin de supprimer les points superflus (zones de recouvrement) et d'ordonner les points restants. Ceux-ci sont ensuite utilisés comme support pour la définition des nœuds du nouveau maillage éléments finis.

Les courbes V.22 montrent l'évolution du rayon du carter en fonction d'un paramètre

de position angulaire. Chacune des courbes correspond à un relevé pour une altitude donnée. Le défaut de forme, dont l'amplitude maximale atteint 4,5 mm au sommet du carter (soit environ 1%, ce qui est important pour ce type de pièce), apparaît très nettement sur ces courbes. Une analyse de Fourier montre que ce défaut ne se décompose pas simplement sur un ou deux harmoniques (voir figure V.23), ce qui rend difficile l'évaluation de son influence sur les modes de la structures.

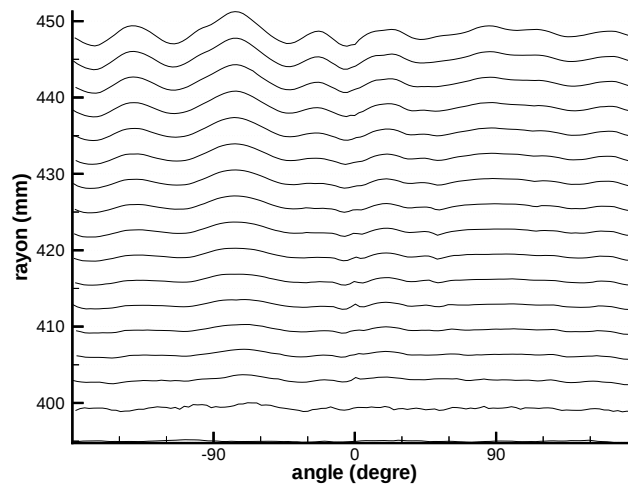


Figure V.22 : Défaut de forme du carter

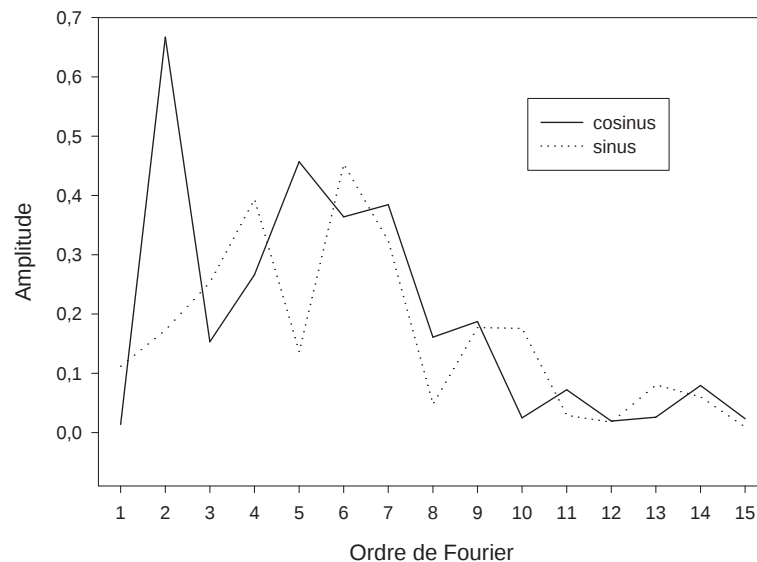


Figure V.23 : Décomposition de Fourier du défaut

NOUVEAU MAILLAGE

Le nouveau maillage réalisé est représenté sur la figure V.24. Il permet d'obtenir les résultats présentés dans le tableau V.3. La figure V.25 présente la comparaison avec les résultats expérimentaux. Il apparaît que la prise en compte fine de la géométrie a permis de mieux représenter le comportement modal de la structure sur l'intervalle d'étude.

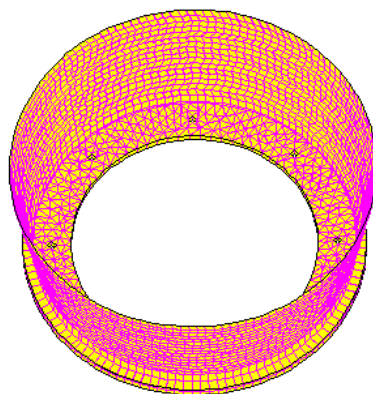


Figure V.24 : Maillage avec prise en compte du défaut

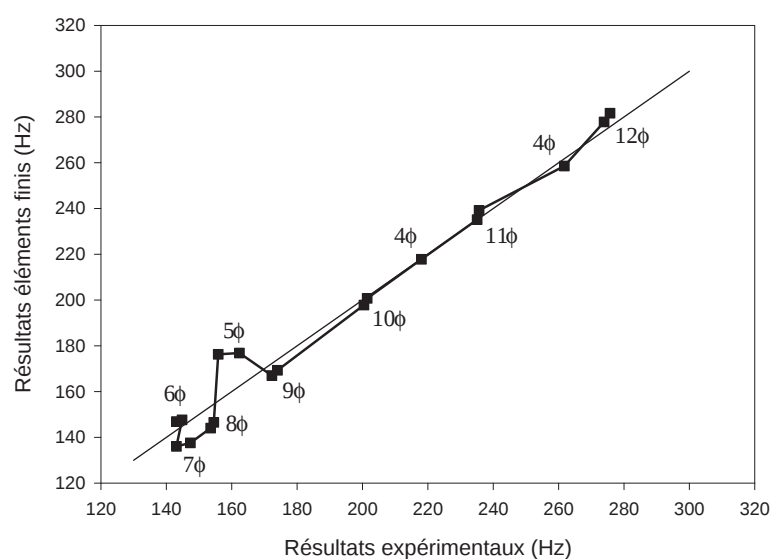


Figure V.25 : Comparaison éléments finis/expérience (maillage avec défaut)

V.2.4.3 Etude d'influence

La figure V.25 montre qu'il existe encore quelques écarts entre le modèle éléments finis et les résultats expérimentaux. C'est pourquoi une étude d'influence, utilisant

Fréquence propre (Hz)	Nature du mode (ordre de Fourier)
136,02	7
137,58	7
144,05	8
146,55	8
146,88	6
147,65	6
167,00	9
169,29	9
176,23	5
176,84	5
197,82	10
200,72	10
217,79	4
235,13	11
239,16	11
258,23	4
277,82	12
281,62	12

Tableau V.3: Résultats modaux éléments finis (avec défaut)

les plans factoriels de la méthode des plans d'expérience, a été menée. Les rappels théoriques nécessaires concernant la mise en œuvre des plans d'expérience sont présentés dans l'annexe I; ils sont extraits des ouvrages [GOU88] et [BEN94b].

Il convient de souligner que cette méthode ne permet que de mettre en évidence l'influence de tel ou tel paramètre sur le comportement du modèle, mais ne permet en aucun cas de mettre en lumière l'absence de prise en compte d'un phénomène par exemple. En ce sens, elle vient compléter la démarche ci-dessus (recherche du défaut de forme notamment).

PRÉSENTATION DE LA MATRICE DES ESSAIS

L'objet d'un plan d'expérience est d'étudier l'influence de la variation d'un certain nombre de *paramètres* sur la *réponse* d'un *modèle*. Il faut donc définir ces trois éléments:

1. **le modèle:** il s'agit du maillage éléments finis établi au paragraphe précédent (maillage avec défaut);
2. **la réponse:** elle est constituée des fréquences propres des modes dont le nombre de diamètres est compris entre 4 et 12, soit 18 valeurs;

3. **les paramètres:** il doivent être choisis en fonction de leur influence supposée sur la réponse; les paramètres suivant ont été retenus:

◇ pour le carter (indices c):

la masse volumique ρ_c ;

le module d'Young E_c ;

l'épaisseur de la paroi tp_c ;

l'épaisseur de la bride tb_c (ce paramètre est *a priori* peu influent, mais il permet de saturer le plan);

◇ pour la couronne (indices r):

la masse volumique ρ_r ;

le module d'Young E_r ;

l'épaisseur de la couronne tb_r ;

Ainsi, sept paramètres ont été retenus. Un plan d'expérience à trois facteurs (numérotés 1 à 3) a été mis en place. Un plan qui ne contiendrait **que** trois facteurs fournirait les influences¹ de ces facteurs sur la réponse du système, ainsi que l'influence des interactions des facteurs entre eux: interaction entre les facteurs 1 et 2, 2 et 3, 1 et 3, ainsi que l'interaction entre tous les facteurs 1, 2 et 3. Ici, il a été choisi de coupler les quatre interactions aux quatre facteurs restant, soit:

le facteur 4 a été couplé à l'interaction 1/2;

le facteur 5 a été couplé à l'interaction 2/3;

le facteur 6 a été couplé à l'interaction 1/3;

le facteur 7 a été couplé à l'interaction 1/2/3.

En termes de plan d'expérience, ce couplage consiste à créer un plan dit fractionnaire 2^{7-4} à sept facteurs, dont quatre sont *aliasés* (c'est-à-dire couplés aux interactions des trois facteurs principaux). Un tel plan est dit *saturé*: il ne serait pas possible en effet d'introduire de facteur supplémentaire, puisque toutes les interactions ont été utilisées. Ainsi, les trois facteurs de base sont²:

1: masse volumique ρ_r ;

2: masse volumique ρ_c ;

3: épaisseur de la couronne tb_r .

¹on parle plus souvent de l'*effet* que de l'influence d'un facteur

²on utilise ici la notation de Box: un facteur est représenté par un numéro, et le couplage d'un facteur avec une interaction s'écrit sous la forme d'une égalité; ainsi $4 = 12$ signifie que le facteur 4 est couplé avec l'interaction entre les facteurs 1 et 2

Les trois facteurs associés aux interactions d'ordre 2 sont:

4=12: épaisseur de la paroi tp_c ;

5=23: module d'Young E_c ;

6=13: module d'Young E_r ;

Enfin, l'épaisseur de la bride tb_c est associée à l'interaction **7=123**. Les domaines de variation des facteurs ont été centrés sur les valeurs fournies (modules d'Young, masses volumiques) ou mesurées (épaisseurs). Les bornes des domaines correspondent aux valeurs centrales plus ou moins 10%. Ces domaines sont précisés dans le tableau V.4.

Paramètre	Borne inférieure -	Centre du domaine 0	Borne supérieure +
ρ_r ($kg.m^{-3}$)	7020,00	7800,00	8580,00
ρ_c ($kg.m^{-3}$)	8217,00	9130,00	10043,00
tb_r (mm)	22,50	25,00	27,50
tp_c (mm)	1,35	1,5	1,65
E_c (MPa)	211500	235000	258500
E_r (MPa)	189000	210000	231000
tb_c (mm)	6,03	6,70	7,37

Tableau V.4: Domaines paramétrés

La matrice des effets, notée $\mathbb{X}$, est rassemblée dans le tableau V.5 ³:

numéro essai	M	1	2	3	4 =12	5 =23	6 =13	7 =123
1	+	+	+	+	+	+	+	+
2	+	+	+	-	+	-	-	-
3	+	+	-	+	-	-	+	-
4	+	+	-	-	-	+	-	+
5	+	-	+	+	-	+	-	-
6	+	-	+	-	-	-	+	+
7	+	-	-	+	+	-	-	+
8	+	-	-	-	+	+	+	-

Tableau V.5: Matrices des effets

³ici, par commodité de lecture, la moyenne sera notée M au lieu de I

CONTRASTES

Les contrastes sont des grandeurs associées à l'effet de la variation d'un facteur (cette notion est détaillée dans l'annexe I).

La figure V.26 représente les contrastes calculés à partir des résultats du plan d'expérience.

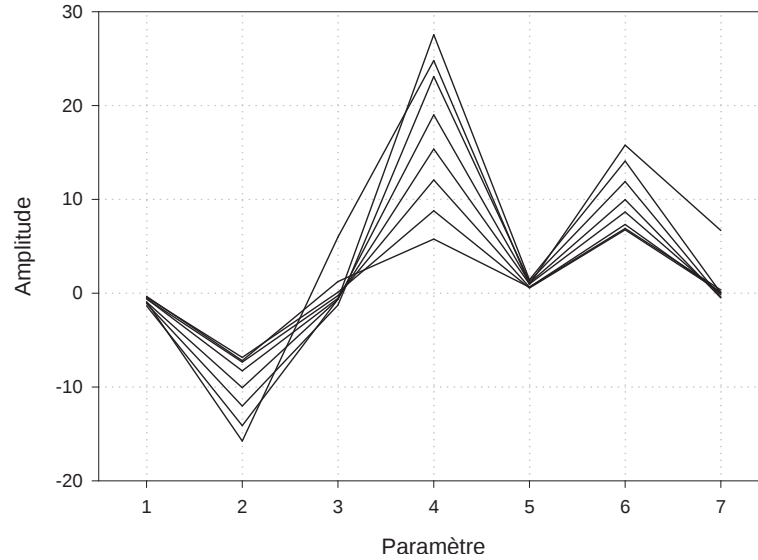


Figure V.26 : Contrastes obtenus (modes 6 à 12 diamètres)

L'ensemble des relations qui définissent les couplages entre facteurs et interactions est appelé *groupe des générateurs d'alias*. Dans le cas présent, le groupe des générateurs d'alias se compose de la façon suivante:

générateurs indépendants: M, 124, 136, 235 et 1237;

générateurs dépendants: 2345, 1345, 347, 1256, 267, 157, 456, 1467, 2457, 3567 et 1234567.

Ainsi, les contrastes à l'ordre 2 s'écrivent:

$$\begin{aligned}
 l_1 &= 1 + 24 + 36 + 57 \\
 l_2 &= 2 + 14 + 35 + 67 \\
 l_3 &= 3 + 16 + 25 + 47 \\
 l_4 &= 4 + 12 + 37 + 56 \\
 l_5 &= 5 + 23 + 17 + 46 \\
 l_6 &= 6 + 13 + 27 + 45 \\
 l_7 &= 7 + 34 + 26 + 15
 \end{aligned}$$

L'analyse de ces contrastes conduit à dégager quelques tendances générales pour l'ensemble des modes à l'exception des modes 4 et 5 diamètres:

- ◇ les paramètres 2 (ρ_c), 4 (tp_c) et 6 (E_r) sont influents; l'influence de ρ_c est négative (la réponse diminue quand le paramètre augmente), alors que les influences de tp_c et E_r sont positives;
- ◇ les paramètres 1 (ρ_r), 3 (tb_r), 5 (E_c) et 7 (tb_c) sont négligeables;
- ◇ les interactions sont négligeables.

Pour les modes à 4 et 5 diamètres, la situation est un peu modifiée (voir figure V.27): les tendances générales observées pour les autres modes se retrouvent à la seule différence du paramètre 3, qui est très influent pour la réponse du premier mode à 4 diamètres, et légèrement influent pour le second mode à 4 diamètres et pour le mode à 5 diamètres.

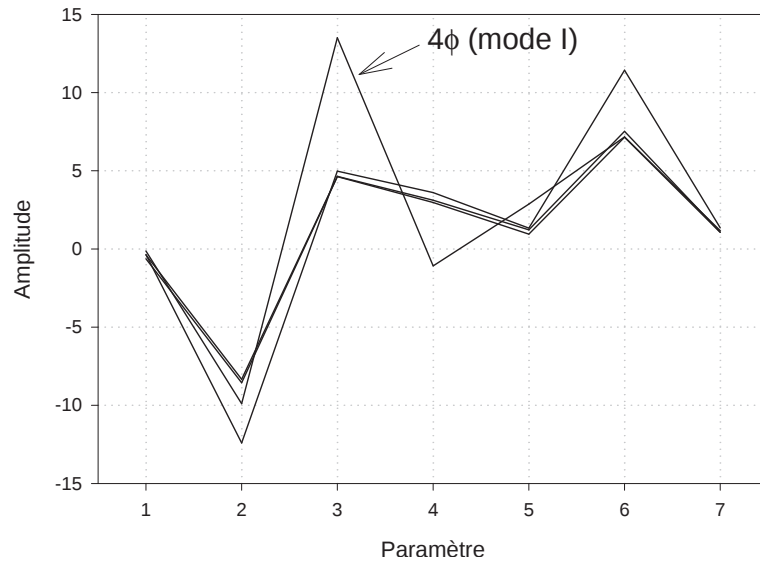


Figure V.27 : Contrastes obtenus (modes 4 et 5 diamètres)

L'écriture des contrastes se trouve donc légèrement modifiée, mais ne fait néanmoins pas intervenir d'interaction. Il n'y a donc pas à réaliser de plan supplémentaire pour désaliaser. En particulier, une représentation mathématique des réponses des modes à 4 et 5 diamètres ne fera intervenir que des monômes de degré 1 en chacun des paramètres, soit (les réponses des modes sont notées r_{iI} et r_{iII} , pour $i \in \{4, 5\}$):

$$\begin{aligned}
 r_{4I} &= M_{4I} + (l_2)_{4I}.x_2 + (l_3)_{4I}.x_3 + (l_4)_{4I}.x_4 + (l_6)_{4I}.x_6 \\
 r_{4II} &= M_{4II} + (l_2)_{4II}.x_2 + (l_3)_{4II}.x_3 + (l_4)_{4II}.x_4 + (l_6)_{4II}.x_6 \\
 r_{5I} &= M_{5I} + (l_2)_{5I}.x_2 + (l_3)_{5I}.x_3 + (l_4)_{5I}.x_4 + (l_6)_{5I}.x_6 \\
 r_{5II} &= M_{5II} + (l_2)_{5II}.x_2 + (l_3)_{5II}.x_3 + (l_4)_{5II}.x_4 + (l_6)_{5II}.x_6
 \end{aligned}$$

Cette formulation peut se mettre sous forme matricielle:

$$\{\mathcal{R}\} = \{\mathcal{M}\} + \mathbb{L}\{\mathcal{X}\}$$

avec:

$$\begin{aligned}
 \{\mathcal{R}\}^t &= \{r_{4I} \ r_{4II} \ r_{5I} \ r_{5II} \} \\
 \{\mathcal{M}\}^t &= \{M_{4I} \ M_{4II} \ M_{5I} \ M_{5II} \} \\
 \{\mathcal{X}\}^t &= \{x_2 \ x_3 \ x_4 \ x_6 \} \\
 \mathbb{L} &= \begin{bmatrix} (l_2)_{4I} & (l_3)_{4I} & (l_4)_{4I} & (l_6)_{4I} \\ (l_2)_{4II} & (l_3)_{4II} & (l_4)_{4II} & (l_6)_{4II} \\ (l_2)_{5I} & (l_3)_{5I} & (l_4)_{5I} & (l_6)_{5I} \\ (l_2)_{5II} & (l_3)_{5II} & (l_4)_{5II} & (l_6)_{5II} \end{bmatrix}
 \end{aligned}$$

Il est *a priori* possible de déterminer le vecteur $\{\mathcal{X}\}$ afin que le vecteur $\{\mathcal{R}\}$ contienne exactement les valeurs $\{\mathcal{R}_{\text{expe}}\}$ obtenues expérimentalement; il suffit pour cela de calculer:

$$\{\mathcal{X}\} = \mathbb{L}^{-1} (\{\mathcal{R}_{\text{expe}}\} - \{\mathcal{M}\})$$

Cette méthode n'est pas la meilleure, car elle conduit, dans ce cas, à des valeurs hors domaine pour les x_i :

$$\{\mathcal{X}\}^t = \{5, 5; -2, 8; -17, 6; 13, 1\}$$

Par ailleurs, les paramètres 2, 4 et 6 étant très influents pour les autres modes, il est préférable de ne pas les modifier (les valeurs moyennes obtenues sont assez proches des valeurs d'essai), et de modifier la valeur du paramètre 3 pour parvenir à rapprocher le modèle de l'expérience. Le contraste du paramètre 3 est le plus important pour le premier mode à 4 diamètres: c'est pourquoi la valeur x_3 est calculée à partir de r_{4I} , en prenant 0 pour tous les autres paramètres. Il en résulte: $x_3 \approx 0,126$, ce qui correspond à $tb_r \approx 25,31 \text{ mm}$. Les résultats fournis par le modèle éléments finis en prenant cette modification en compte sont présentés sur la figure V.28.

La figure V.29 présente les écarts (en pourcentage) entre les valeurs expérimentales et les deux modèles éléments finis (maillage avec défaut non recalé et maillage avec défaut recalé). Cette figure montre que, globalement, les écarts entre modèle et expérience ont été réduits grâce aux résultats de l'étude d'influence, l'amélioration la plus nette se produisant pour le mode à 5 diamètres. Un calcul de l'écart entre les points obtenus et la droite d'équation $y = x$ montre que la valeur diminue de 30% environ entre le modèle non recalé et le modèle recalé.

CONCLUSION

L'analyse d'influence de certains paramètres décrivant le modèle a permis de caler les modes du modèle éléments finis sur une plage fréquentielle d'environ $300Hz$.

V.2.4.4 Remarque: effet de la rotation

Il est légitime de se demander quelle est l'influence de la rotation du tour vertical sur les modes du carter. En fait, dans le cas présent, les vitesses de rotation restent assez faibles (maximum: 5 Hz). Il a donc été possible, sans mettre en place de dispositif lourd,

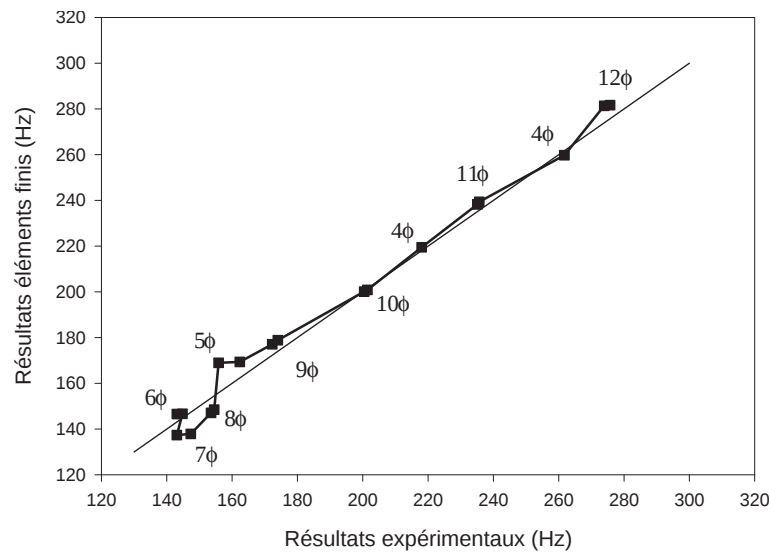


Figure V.28 : Comparaison éléments finis/expérience (maillage recalé)

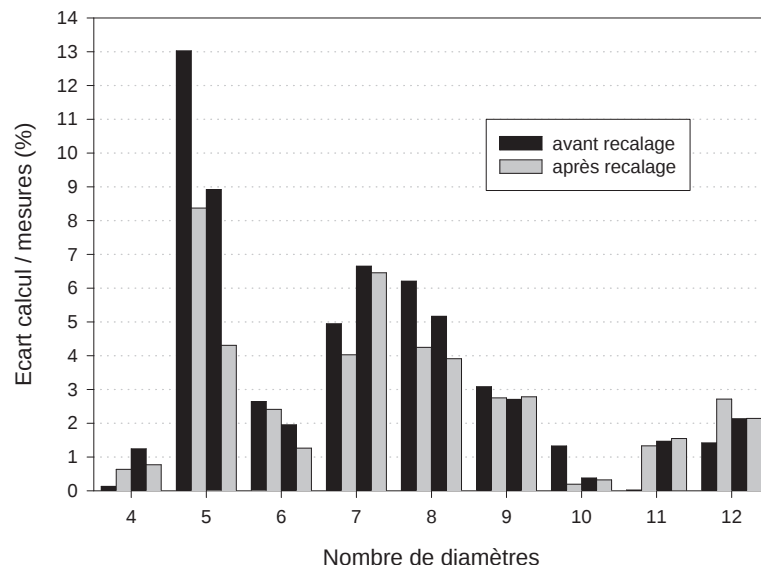


Figure V.29 : Comparaison éléments finis/expérience

d'embarquer du petit matériel sur le plateau du tour: le miroir face avant a été monté (encastrement) sur un cylindre en acier lui-même rendu solidaire du plateau du tour par des brides. Le vélocimètre, installé sur le portique, visait, par l'intermédiaire du miroir, un point de la paroi. Le tour a été ensuite mis en rotation pour différentes vitesses, et des réponses du carter au marteau ont été enregistrées par le vélocimètre. Les réponses obtenues l'ont donc été pour des points fixes du carter (points qu'il était possible de faire varier). Le dépouillement des résultats a montré qu'aux vitesses de rotation atteintes, le comportement modal du carter n'était pas modifié.

Chapitre VI

Essais de contact

Les essais de contact réalisés sont présentés dans ce chapitre. Les premiers essais, qui cherchent à simuler un phénomène d'interaction, se font sur le carter sans modification particulière (paragraphe VI.2). Certaines tendances peuvent se dégager, mais le défaut de forme du carter apparaît comme un obstacle possible à l'apparition nette de l'interaction.

Une deuxième série d'essais est alors réalisée (paragraphe VI.2.4) avec le carter sur lequel une couche de matériau abradable a été déposée.

VI.1 Notion d'interaction

Le tour vertical est utilisé pour effectuer quelques essais de contact entre le carter en rotation et un modèle simplifié d'aube. L'objectif est d'essayer de créer un phénomène d'*interaction* entre les deux structures. La description de ce phénomène fait l'objet du paragraphe suivant.

VI.1.1 Théorie

L'*interaction*, telle qu'elle est traitée ici, correspond à des phénomènes de couplage qui surviennent dans les turboréacteurs notamment lorsque la partie en rotation (la roue aubagée) et la partie fixe (le carter) viennent en contact. L'origine de ce contact peut être:

- ◇ prévisible: au cours de manœuvres par exemple (rotation de l'avion en bout de piste), il peut arriver que la roue aubagée viennent frotter sur le carter. Ce type de contact est bien maîtrisé: l'intérieur du carter est revêtu d'une matière abradable qui amorti ce contact.
- ◇ imprévisible: par exemple lors de l'absorption de corps étrangers (glace, oiseaux...), un balourd se crée qui peut amener la roue aubagée à venir contacter le carter de façon violente; si le contact se prolonge, les deux éléments en contact peuvent entrer simultanément en résonance, et donc subir de fortes contraintes susceptibles de remettre en cause leur intégrité.

Cette résonance simultanée des deux structures correspond à ce qui sera appelé *interaction* par la suite. Il est possible d'écrire une condition mathématique simple sur l'existence du phénomène (par exemple [STA90]). Pour cela, il faut considérer les deux pièces mécaniques en contact comme étant de symétrie de révolution, et, comme telles, possédant des modes propres à diamètres (voir le paragraphe V.2.1). Si le rotor, dont la vitesse de rotation est Ω , est le siège d'une vibration caractérisée par le nombre n de ses diamètres et par une pulsation ω_r , alors, dans un repère lié au rotor, la vibration peut se décomposer en deux ondes de vitesses de propagation opposées égales respectivement à:

$$-\frac{\omega_r}{n} \quad \text{et} \quad +\frac{\omega_r}{n}$$

Dans le repère fixe, en tenant compte de la vitesse de rotation du rotor, ces vitesses s'écrivent:

$$\Omega - \frac{\omega_r}{n} \quad \text{et} \quad \Omega + \frac{\omega_r}{n}$$

Si le stator est lui aussi siège d'une vibration caractérisée par le même nombre n de diamètres, avec une pulsation ω_s , alors là aussi deux ondes se propagent, qui ont, dans le repère fixe, les vitesses:

$$-\frac{\omega_s}{n} \quad \text{et} \quad +\frac{\omega_s}{n}$$

Il y a coïncidence lorsque ces vitesse d'ondes s'égalent, ce qui conduit aux quatre possibilités suivantes:

$$\omega_s - \omega_r = n\Omega$$

$$\omega_s + \omega_r = n\Omega$$

$$\omega_s - \omega_r = -n\Omega$$

$$\omega_s + \omega_r = -n\Omega$$

La dernière relation n'a pas d'interprétation physique. Pour les autres, il est possible de les représenter sur un diagramme type diagramme de Campbell (pulsation en fonction de la vitesse de rotation), comme sur la figure VI.1 (en toute rigueur, la fréquence du rotor évolue avec la vitesse de rotation; ici, pour simplifier, il a été supposé que ce n'était pas le cas).

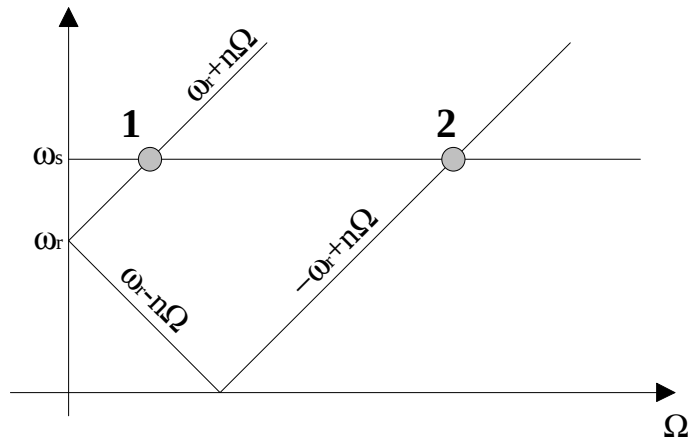


Figure VI.1 : Points d'interaction

Différentes études dont notamment celle de Tobias et Arnold [TOB57] ont montré que seul le cas d'ondes progressives pouvait générer des coïncidences modales d'amplitudes importantes. C'est pourquoi seul le point noté 2 sur la figure VI.1 est considéré comme point critique, et doit donc être évité. Une condition nécessaire de l'interaction est donc:

$$\omega_s + \omega_r = n\Omega$$

L'expérience prouve qu'il ne s'agit pas néanmoins d'une condition suffisante.

Le dimensionnement des moteurs tient compte de ce problème en assurant une marge de sécurité de 10% dans le domaine fréquentiel autour de ces points. La réalisation de cette marge passe par la mise en place de raidisseurs, ce qui a pour conséquence d'augmenter l'encombrement des carters. La performance du moteur s'en trouve affectée. C'est pourquoi l'étude précise des conditions de l'interaction constitue une priorité chez de nombreux motoristes.

VI.1.2 Application au banc d'essai

Dans la suite, il sera plus commode de travailler en termes de fréquence plutôt qu'en termes de vitesse de rotation: ainsi, la fréquence d'un mode du carter sera désignée par f_c , celle d'un mode des aubes par f_a , et la fréquence de la vitesse de rotation sera notée f_{rot} .

La théorie développée est appliquée au banc d'essai dans une configuration très particulière: la partie roue aubagée (rotor) est en effet rapportée à une aube unique. Cette restriction a été choisie en supposant qu'elle permettrait de supprimer la notion de sens de déplacement des ondes dans la roue aubagée, et donc de multiplier les quadruplets (n, f_c, f_a, f_{rot}) satisfaisant la condition d'interaction. Par ailleurs, afin de s'affranchir des effets aérodynamiques, le carter est pris comme partie tournante, et l'aube est fixe. Ainsi, la condition nécessaire d'interaction s'écrit-elle:

$$f_c = f_a \pm n f_{rot}$$

Les modes du carter sont parfaitement identifiés; les modes de l'aube se résument dans cette étude aux modes de flexion d'une poutre droite; la vitesse de rotation du tour, enfin, peut varier de façon discrète dans l'intervalle $[0..300Hz]$ (voir annexe G).

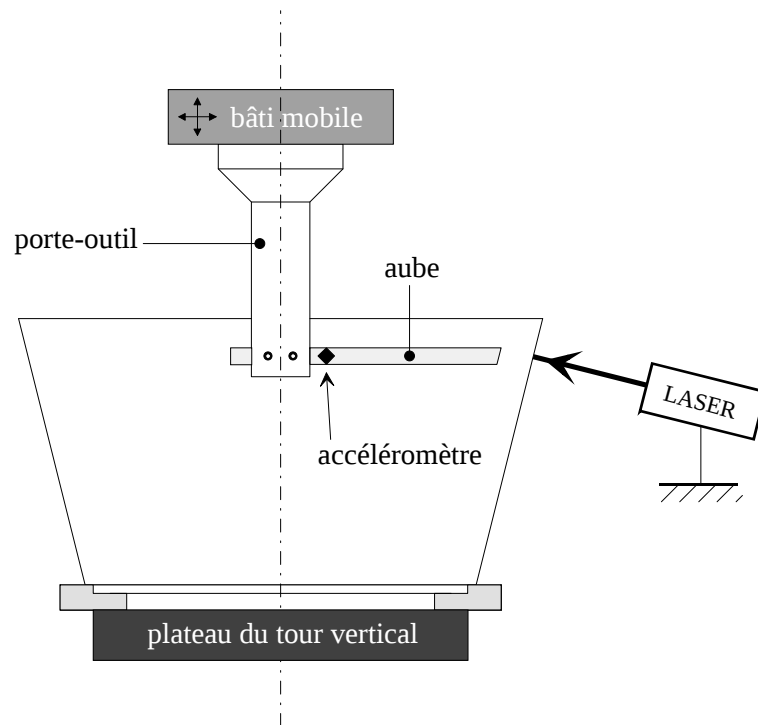
VI.2 Création du contact et premiers résultats

Pour créer l'interaction, il faut choisir un mode du carter: le mode à 10 diamètres à 200 Hz environ a été retenu. Il se trouve en effet assez éloigné des autres modes (écart fréquentiel de 10 % avec les modes adjacents), peu perturbé par le défaut (les déformées modales sont régulières, et le décalage en fréquence entre les deux déformées est inférieur à 0,5 %), et l'amortissement est assez faible.

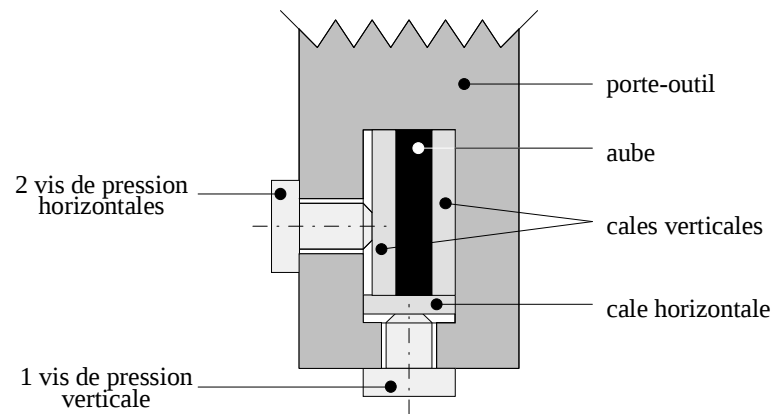
Pour l'aube, c'est le premier mode de flexion qui est retenu (accessible simplement en réglant la longueur libre). Cette aube est constituée d'un matériau qui est un acier classique moins dur que celui du carter afin de ne pas usiner ce dernier lors des essais de contact.

VI.2.1 Montage

La configuration de mise en contact est présentée sur la figure VI.VI.2.1: l'aube est fixée sur le porte-outil du tour vertical: il est ainsi assez simple de régler la valeur de sa première fréquence propre, en augmentant ou en diminuant la longueur entre l'extrémité libre et l'encastrement. Cet encastrement est réalisé par trois points de serrage (voir figure VI.VI.2.1), deux horizontaux, et un vertical. Le premier mode de vibration du porte-outil se trouve à 290 Hz, ce qui est assez éloigné (45 %) de la plage fréquentielle dans laquelle les essais vont être réalisés. Le vélocimètre est installé à l'extérieur du carter, le faisceau, orienté perpendiculairement à la paroi, visant un point situé près du bord libre; il permet donc d'obtenir la réponse du carter dans le repère fixe. Un accéléromètre est collé sur l'aube, près de l'encastrement, dans le sens de flexion.



(a) Vue générale du montage



(b) Fixation aube/porte-outil (vue de face)

Figure VI.2 : Montage de l'aube sur le tour vertical

Il convient de noter qu'ici, le carter étant la partie tournante et l'aube la partie fixe, la condition d'interaction s'écrit:

$$f_a = f_c \pm n f_{rot}$$

VI.2.2 Détermination négative

Le contact est tout d'abord créé de façon à respecter la relation suivante:

$$f_a = f_c - n \cdot f_{rot}$$

La vitesse de rotation choisie est la vitesse maximale du tour vertical: $\Omega = 300$ tr/min (soit 10π rad/s ou 5 Hz). Si la fréquence du mode à $n = 10$ diamètres est approchée par $f_c = 200$ Hz, il en résulte que la fréquence de l'aube est environ $f_a = 150$ Hz.

Une première série de mesures permet de voir évoluer la réponse du carter lorsque la vitesse de rotation du tour augmente. Cette série de mesures a été réalisée en mettant le tour en rotation, puis en amenant l'aube au contact (légère pression: l'aube touche le carter sur une seule zone longue d'environ 30 cm): un enregistrement des réponses de l'aube et du carter est fait sur l'intervalle $[0..300Hz]$ avec une moyenne sur 3 acquisitions. Afin d'essayer de limiter les effets du défaut de forme, le contact a été réalisé dans la partie médiane du carter (voir figure VI.3). Le contact aurait pu être créé dans la partie basse, puisque le rayon y est pratiquement constant, mais trois raisons s'y opposent:

- ◇ la présence de cordons de soudure (le carter est en effet constitué de trois éléments soudés entre eux);
- ◇ la trop grande proximité de l'encastrement, qui pourrait gêner l'apparition du phénomène;
- ◇ l'impossibilité physique pour la traverse du porte outil de descendre en deçà d'une certaine limite.

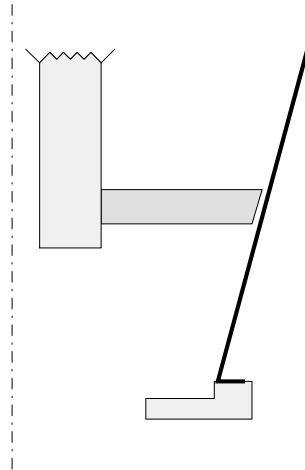


Figure VI.3 : Zone de contact

Les résultats d'une série de mesures sont présentés sur la figure VI.4, pour des vitesses de rotation comprises entre 0 et 2,6 Hz, les réponses associées aux vitesses élevées (5 Hz) seront évoquées plus loin.

Ces résultats sont enregistrés dans le repère fixe: comme cela a été rappelé dans le paragraphe V.1.2.1, toute réponse d'une vibration à n diamètres à la fréquence f dans le repère tournant à la vitesse f_{rot} se traduit par la perception de deux ondes de fréquence $f - n.f_{rot}$ et $f + n.f_{rot}$ dans le repère fixe. Ainsi, le dédoublement observé sur la figure doit être compris comme l'existence d'une vibration dans le carter à la fréquence médiane.

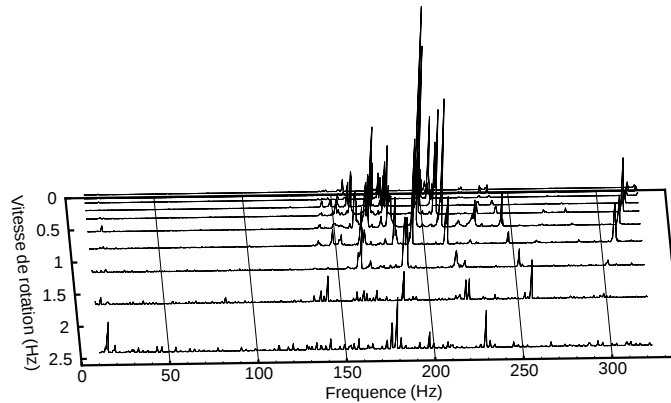


Figure VI.4 : Essai de contact pour $f_a = 150Hz$

Trois remarques peuvent être faites au sujet de cette figure:

- ◇ il semble que le mode à 10 diamètres du carter réponde très nettement à la sollicitation pour de faibles vitesses de rotation (alors que la vitesse de rotation prévue pour l'interaction est 300 tr/min): les réponses des autres modes sont, le plus souvent négligeables (la deuxième réponse nettement observable sur la figure est celle du mode à 9 diamètres);
- ◇ contrairement à ce qui est attendu, la réponse du mode à 10 diamètres semble disparaître avec l'augmentation de la vitesse de rotation;
- ◇ la contribution de l'onde progressive (associée à la fréquence $f_c - n.f_a$ dans le repère fixe) paraît plus importante que la contribution de l'onde rétrograde.

Il convient par ailleurs de noter les points suivants:

- ◇ pour les vitesses de rotation les plus importantes (soit 210 tr/min et 300 tr/min), la contribution des modes autres que le mode à 10 diamètres dans la réponse n'est plus réellement négligeable: la réponse du carter correspond ainsi plus à une réponse d'une structure à un choc; ce problème est attribué au défaut de forme du carter (voir figure VI.5);

◇ par ailleurs, il faut souligner que les courbes ne sont pas vraiment *quantitativement* comparables car le contact n'est pas le même d'une mesure à l'autre, pour deux raisons:

1. le contact créé entre l'aube et le carter peut modifier la nature des géométries des éléments en contact (notamment la géométrie du sommet d'aube), ce qui constitue un premier facteur qui empêche la répétabilité des mesures;
2. entre deux mesures, le tour est arrêté (indispensable pour procéder au changement de vitesse): l'aube est écartée de la paroi du carter avant le freinage du tour, puis est remise au contact une fois la nouvelle vitesse établie; ainsi, d'une mesure à l'autre, la position de l'aube par rapport à la paroi peut varier sensiblement (précision d'avance du tour: 0,1 mm, mais il existe un jeu important dans le volant de commande de déplacement de la traverse qui entraîne le porte outil).

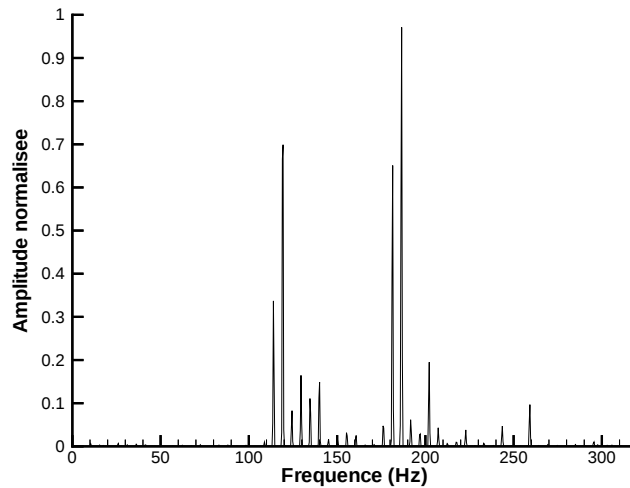
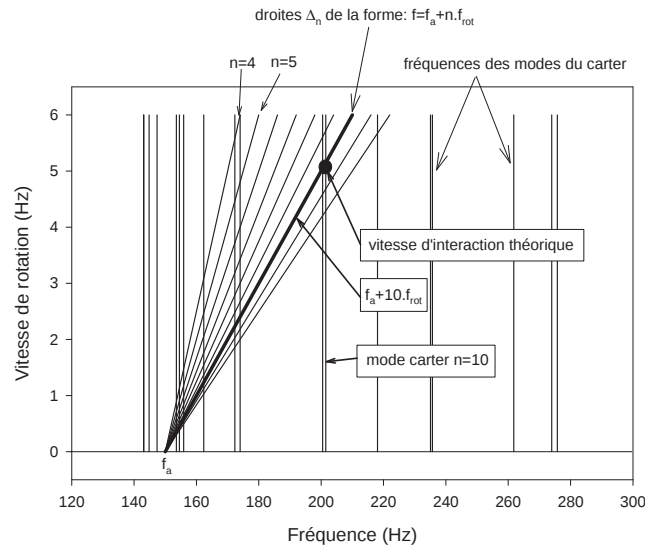


Figure VI.5 : Réponse du carter à 300tr/min

Enfin, il se peut que la fréquence choisie pour l'aube corresponde à une fréquence d'interaction pour un autre ensemble (n, f_c, f_{rot}) . Pour cela, les droites Δ_n définies par $f = f_a + n \cdot f_{rot}$ sont superposées aux droites $f_c = C^{te}$ sur la figure VI.6. Noter que dans ce cas, seuls les ensembles vérifiant $f_a = f_c - n \cdot f_{rot}$ sont recherchés, les fréquences propres du carter étant telles que pour $f_a = 150$ Hz, aucun mode ne peut remplir la condition $f_a = f_c + n \cdot f_{rot}$.

L'examen détaillé de cette figure montre que seuls quelques modes peuvent vérifier une condition d'interaction dans la plage de vitesses de rotation étudiée. Ces modes sont rassemblés dans le tableau VI.1 (une vitesse du tour est dite accessible si elle est éloignée de moins de 10% de la vitesse théorique requise).

Figure VI.6 : Vitesses d'interaction pour $f_a = 150$ Hz

Nature du mode (ordre de Fourier)	Vitesse de rotation d'interaction (tr/min)	Vitesse accessible sur le tour (tr/min)
8	27	30
	34	xxx
5	71	xxx
	148	155
9	148	155
	160	155

Tableau VI.1: Vitesses potentielles d'interaction

Ainsi, les courbes obtenues pour les vitesses de rotation à 30 tr/min et à 155 tr/min devraient révéler les interactions des modes concernés: mais rien de vraiment distinct ne s'en dégage. Pourtant, la réponse du mode à 9 diamètres, et dans une moindre mesure, des modes à 8 et 5 diamètres, est bien visible sur la figure VI.4, alors que les réponses des autres modes sont moins nettes. Nous supposons que la difficulté d'interprétation provient probablement en partie de la présence du défaut de forme qui ne permet pas d'avoir les conditions propices à l'apparition du phénomène d'interaction (en particulier, le contact entre l'aube et le carter non seulement n'est pas continu, mais encore est assimilable à du choc lorsque les vitesses de rotation deviennent importantes). Une série de mesures réalisée avec les mêmes conditions, à la pression de contact près (la pression est plus importante, et le contact a alors lieu sur deux zones) conduit à des résultats comparables (voir figure VI.7), si ce n'est que les réponses de type choc apparaissent dès 155 tr/min .

Ceci est rassurant d'un point de vue répétabilité des mesures.

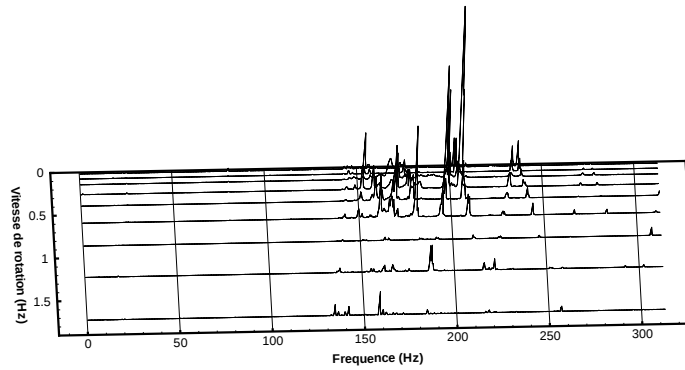


Figure VI.7 : Réponse pour un contact plus important

VI.2.3 Détermination positive

VI.2.3.1 Essais dans la zone médiane

Une deuxième série d'essais a ensuite été réalisée en utilisant cette fois la relation:

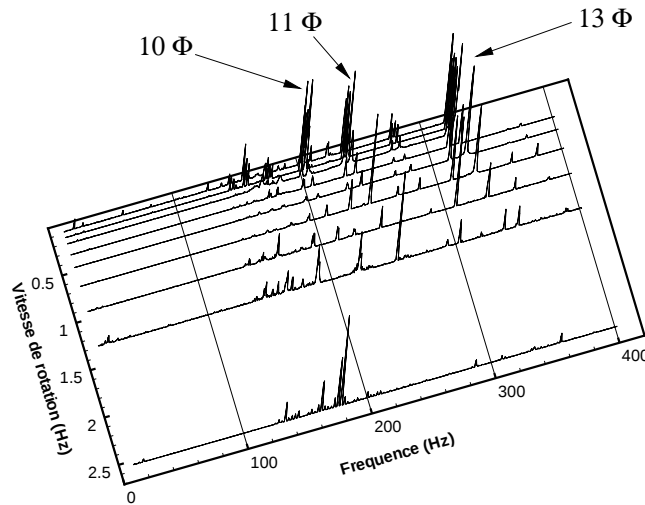
$$f_a = f_c + n \cdot f_{rot}$$

La vitesse de rotation choisie pour l'interaction reste inchangée (soit 5 Hz), le même mode du carter est utilisé (mode à diamètres), et la zone de contact est également conservée. Il en résulte une fréquence pour l'aube d'environ 250 Hz.

Les résultats sont présentés sur la figure VI.8

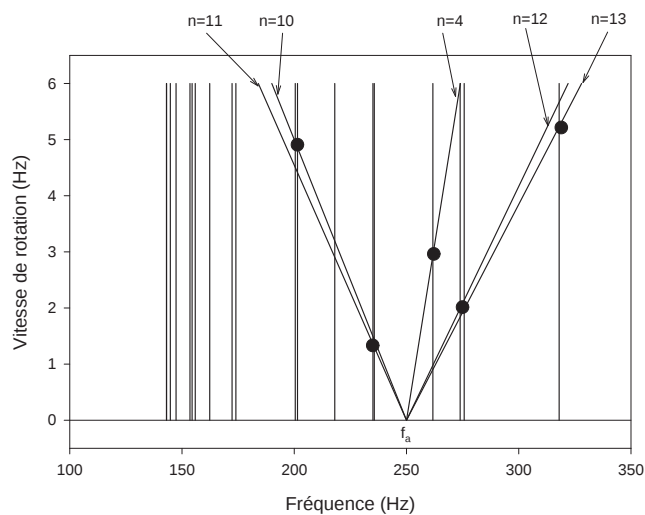
Contrairement au cas précédent, la réponse du mode à 10 diamètres, encore visible pour les petites vitesses, disparaît assez rapidement, et les modes à 5, 9 et 8 diamètres ont des réponses négligeables devant celles de deux autres modes dont les fréquences propres se situent respectivement à environ 235 Hz et 320 Hz. Les résultats de l'analyse modale permettent d'identifier le mode à 235 Hz comme étant le mode à 11 diamètres, et une rapide analyse modale de la plage $[300..400\text{Hz}]$ a montré que la fréquence à 320 Hz correspondait au mode à 13 diamètres du carter (mode double peu perturbé).

Comme dans le cas précédent, la réponse de la structure devient une réponse de type choc pour les hautes vitesses (notamment à 5 Hz), ce qui ne rend pas possible l'observation voulue. Néanmoins, la fréquence de l'aube peut correspondre à une fréquence d'interaction pour d'autres modes. En traçant une figure semblable à la figure VI.6, il est possible de

Figure VI.8 : Essai de contact pour $f_a = 250$ Hz

voir que dans la plage de vitesse étudiée, certains modes peuvent effectivement remplir la condition d'interaction.

Remarque: pour plus de clarté, sur la figure VI.9, seules les droites Δ_n coupant *vraiment* les droites $f_c = C^{te}$ ont été représentées.

Figure VI.9 : Vitesses d'interaction pour $f_a = 250$ Hz

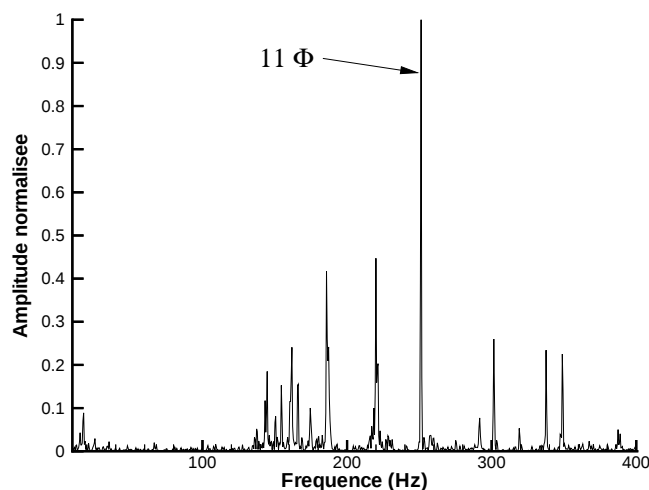
Les points potentiels d'interaction sont rassemblés dans le tableau VI.2.

Ce tableau souligne le fait que les vitesses pour lesquelles se produisent l'interaction

Nature du mode (ordre de Fourier)	Type d'interaction	Vitesse de rotation d'interaction (tr/min)	Vitesse accessible sur le tour (tr/min)
4	-	176	xxx
12	-	120	110
	-	129	xxx
13	-	313	300
11	+	78	80
	+	81	80

Tableau VI.2: Vitesses potentielles d'interaction

ne sont pas toutes accessibles. Sur les courbes obtenues en essais, il n'est possible de remarquer que la réponse importante du mode à 11 diamètres à 80 tr/min (voir figure VI.10), les courbes à 300 tr/min et 110 tr/min n'étant pas exploitables. La mise en place d'un système de variateur de vitesse pourrait être installé sur le moteur du tour, qui permettrait de faire croître continûment la vitesse de rotation.

Figure VI.10 : Réponse à 80 tr/min pour $f_a = 250$ Hz

VI.2.3.2 Essais dans la zone haute

Une série de mesures a été réalisée en conservant les paramètres du paragraphe précédent pour tout ce qui concerne les différentes fréquences utilisées pour l'interaction ($f_c = 200$ Hz, $n = 10$, $f_{rot} = 5$ Hz et $f_a = 250$ Hz), mais l'aube a été cette fois placée au sommet du carter, là où le défaut de forme est le plus important, ce qui se traduit par une bande de contact moins large (environ d'un facteur 2). Les résultats obtenus sont

présentés sur la figure VI.11: il apparaît que la réponse du carter se décompose assez nettement sur ses modes: il s'agit bien d'une réponse de type choc, et aucune tendance ne se dégage du résultat.

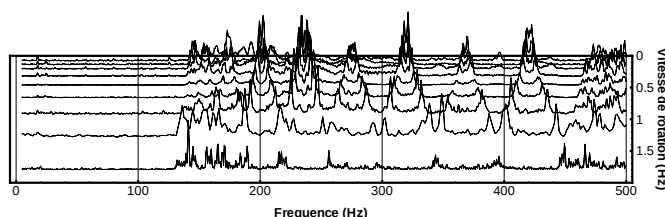


Figure VI.11 : Réponse au contact dans la zone haute du carter

VI.2.4 Modification du montage et nouveaux résultats

L'ensemble des résultats obtenus au cours de la première série d'essais a laissé paraître que le défaut de forme du carter était susceptible d'empêcher la mise en place du phénomène d'interaction. C'est pourquoi l'intérieur du carter a été recouvert d'une couche de matériau abrasable identique à celui utilisé dans les turboréacteurs. L'ensemble a ensuite été usiné à rayon constant, en respectant un angle de conicité moyen sur toute la hauteur, de façon que le contact puisse se produire sur la totalité du carter (voir figure VI.12).

Cette modification géométrique et matérielle a bien sûr entraîné une modification des modes du carter. À ce stade, seule une analyse modale rapide a été réalisée, de façon à avoir un ordre de grandeur des fréquences propres des modes de la nouvelle structure. Le nouveau spectre paraît semblable à celui de la configuration sans abrasable, mais l'ordre des modes semble avoir été modifié (le premier mode serait un mode à quatre diamètres).

Dans ce cas, aucun modèle éléments finis n'a été mis en place: la répartition des masses sur la circonférence n'est plus homogène, et l'utilisation du précédent modèle nécessiterait de prendre en compte l'épaisseur variable de la paroi. Ceci pourra, si nécessaire, faire l'objet d'une étude ultérieure.

Un phénomène surprenant est alors observé: la réponse du carter ne semble plus dépendre de la fréquence de l'aube (voir figure VI.13, obtenues pour deux fréquences d'aubes prises au hasard, à la même vitesse de rotation).

Ces premiers résultats dans la configuration avec abrasable devront être approfondis.

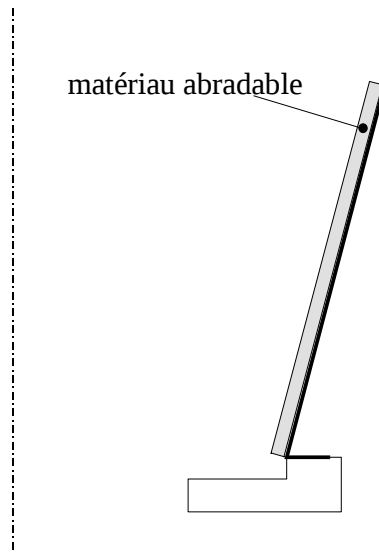
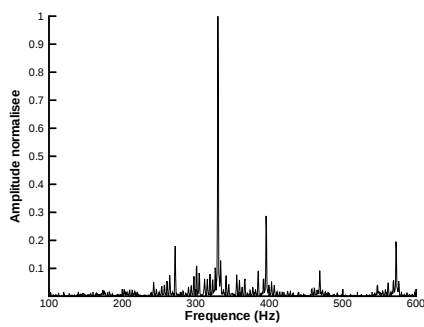
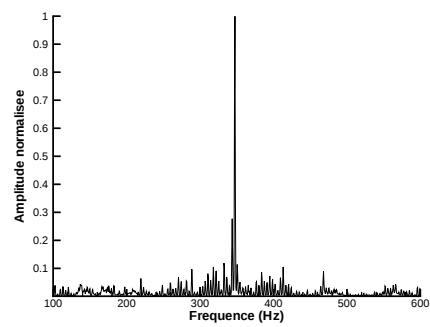


Figure VI.12 : Dpt d'abrasable sur la paroi du carter



(a) Fréquence de l'aube $f_a = 190$ Hz



(b) Fréquence de l'aube $f_a = 295$ Hz

Figure VI.13 : Résultats avec abrasable

Conclusion

Des essais de contact ont été réalisés sur un montage simplifié comprenant une partie mobile constituée d'un tronçon de carter de turboréacteur installé sur le plateau d'un tour vertical, et une partie fixe constituée d'un modèle élémentaire d'aube lié au porte-outil de la machine.

Ces essais avaient tout d'abord pour but de proposer un exemple concret à partir duquel une simulation numérique aurait pu être réalisée pour mettre en œuvre la méthode de recherche de contact s'appuyant sur la modélisation spline développée dans la première partie. Un modèle éléments finis du carter monté sur son support a, dans cette optique, été réalisé. Les résultats de l'analyse modale expérimentale de la pièce ont été, par l'intermédiaire d'une analyse topographique fine et d'un plan d'expérience, utilisés pour recalculer le maillage éléments finis.

Le second objectif des essais était d'essayer de recréer et de mieux comprendre le phénomène d'interaction, qui constitue actuellement une priorité d'étude pour les motoristes, notamment dans l'aviation civile. Cette interaction, contrôlée de façon nécessaire mais non suffisante par une relation entre les fréquences propres des éléments mis en jeu, s'est révélée difficile à observer du fait de la présence d'une irrégularité de forme dans la géométrie du carter. Malgré la mise en œuvre d'un procédé permettant d'assurer une condition de contact permanent au cours du mouvement, l'interaction n'a pu être recréée dans des conditions autorisant son étude.

Conclusion générale

Notre propos était de mettre en place une nouvelle méthode de traitement numérique du contact, puis de procéder, sur un montage simplifié, à des réalisations pratiques, afin d'une part d'établir un ensemble de résultats auxquels seront confrontés par la suite les résultats issus de simulations numériques, et d'autre part, d'essayer de mettre en évidence le phénomène d'interaction.

Dans une première partie, nous avons présenté une méthode de recherche de contact utilisable dans un code éléments finis pour des calculs dynamiques. Cette méthode, qui s'appuie sur la recherche d'intersection d'entités géométriques construites à l'aide de fonctions splines, a été implémentée dans deux codes distincts, et, dans le cas de l'un d'eux, moyennant quelques ajustements, sera bientôt accessible par tout utilisateur. Il apparaît que notre méthode fournit une réponse aux problèmes de facétisation qui nous étaient posés.

Les développements ultérieurs concernent:

- ◇ la gestion de la séparation des corps, en prenant en compte l'histoire du contact, afin de limiter les phénomènes de retour au contact après la séparation;
- ◇ le traitement à donner aux zones de tangence, pour limiter les temps de calcul qu'elles peuvent potentiellement générer;
- ◇ la réalisation d'un test sur le maillage du carter établi dans la partie expérimentale, pour confronter les résultats aux valeurs d'essai;
- ◇ la simulation d'une configuration qui met en échec des méthodes classiques: la rotation d'un arbre dans un palier de diamètre sensiblement égal.

Enfin, l'utilisation de notre méthode a été envisagée pour simuler des problèmes d'emboutissage ainsi que des problèmes de thermoformage.

Dans la seconde partie, nous avons proposé un banc d'essai pour réaliser des configurations simples de contact utilisables pour valider la méthode développée dans la première partie, et essayer d'observer le phénomène d'interaction. Une fois le banc mis en place, une analyse modale fine du carter a été réalisée. Un modèle éléments finis a ensuite été

établi tout d'abord en supposant la structure de révolution, puis en prenant en compte un défaut de forme de la paroi du carter. Grâce à la construction d'un plan d'expérience, le modèle a été recalé dans la plage fréquentielle d'étude. Les essais de contact ont ensuite été réalisés, mais, compte-tenu du défaut de forme, et bien que les conditions nécessaires d'interaction soient remplies, le phénomène n'a pu être observé de façon satisfaisante, même après une modification susceptible d'améliorer les conditions de contact.

Les développements ultérieurs envisageables sont:

- ◇ la mise en forme des données dans l'objectif de les comparer aux résultats de simulation;
- ◇ la modification de la configuration du tour vertical de façon à pouvoir faire varier continûment sa vitesse de rotation;
- ◇ la création d'un système permettant de mieux contrôler les conditions de contact entre aube et carter;
- ◇ l'utilisation de modèles d'aubes plus complexes, ainsi que de systèmes comprenant davantage d'aubes.

Bibliographie générale

- [ARN98] ARNOULT E., PESEUX B., ET BERTHILLIER M., “Aircraft engine blades-casing contact study”, dans *International Conference on Computational Engineering Sciences (ICES’98)*, Vol. 2, pp. 1032–1037, 6-9 octobre 1998. Atlanta (USA).
- [ARN99] ARNOULT E., PESEUX B., ET BONINI J., “Recherche analytique du contact appliquée à un code éléments finis en dynamique”, dans *14ème Congrès Français de Mécanique*, pp. 1–6, 30 août - 3 septembre 1999. Toulouse (France) - papier 218 (CDROM).
- [BAT90] BATOZ J. ET DHATT G., *Modélisation des structures par éléments finis*, Vol. 3 - Coques. Hermès, 1990.
- [BEN94a] BENJEDDOU A. ET HAMDI M. A., “Un nouvel élément fini du type B-spline pour l’analyse dynamique des coques de révolution”, *Revue Européenne des Eléments Finis*, Vol. 3, No. 1, pp. 101–126, 1994.
- [BEN94b] BENOIT D., TOURBIER Y., ET GERMAIN-TOURBIER S., *Plans d’expériences: construction et analyse*. Tec & Doc, Lavoisier, 1994.
- [BEZ93] BEZIER P. E., “The first years of CAD/CAM and the UNISURF CAD system”, dans *Fundamental Developments of Computer-Aided Geometric Modeling*(PIEGL L., ed.), pp. 13–26, Academic Press, 1993.
- [BIL96] BILLET L. ET MORENO J., “Caractérisation des modes de vibration d’une coque de révolution à l’aide d’un vibromètre laser à faisceau tournant”, *Revue Française de Mécanique*, Vol. 1, pp. 25–32, 1996.
- [BON95] BONINI J., *Contribution à la prédiction numérique de l’endommagement de stratifiés composites sous impact basse vitesse*. Thèse, Ecole Nationale Supérieure des Arts et Métiers, Bordeaux, 1995.
- [BON97] BONINI J. ET BUNG H., “Modélisation des problèmes de contact-impact avec frottement en explicite par la méthode des multiplicateurs de Lagrange”, dans *Actes du troisième Colloque National en Calcul des Structures*, Vol. 1, pp. 411 – 416, CSMA, 20-23 mai 1997. Giens (Var).

- [CEA97] CEA, *PLEXUS, Manuel théorique*, 1997.
- [CEA99] CEA, *PLEXUS, notice d'utilisation*, 1999. http://bacchus.lmt.ens-cachan.fr:8080/DOC/PLEXUS/notice_plx.html.
- [CES96] CESCOTTO S., “Contact quasi-statique: Résolution avec régularisation”, dans *Modélisation mécanique et numérique du contact et du frottement*, pp. 251 – 256, IPSI, 15-17 octobre 1996. Paris.
- [CHA90] CHARTIER G., *Les Lasers*. Technologies de pointe, Hermès, 1990.
- [COX72] COX M. G., “The numerical evaluation of B-splines”, *Journal of the Institute of Mathematics and its applications*, Vol. 10, pp. 134–149, 1972.
- [COX93] COX M. G., “Algorithms for spline curves and surfaces”, dans *Fundamental Developments of Computer-Aided Geometric Modeling*(PIEGL L., ed.), pp. 51–76, Academic Press, 1993.
- [DAG99] DAG I., “A quadratic B-spline finite element method for solving nonlinear Schrödinger equation”, *Computer Methods in Applied Mechanics and Engineering*, Vol. 174, pp. 247–258, 1999.
- [DAN89] DANIEL M., *Modélisation de courbes et surfaces par des B-splines - Application à la conception et à la visualisation de formes*. Thèse, Ecole Nationale Supérieure de Mécanique, Nantes, mai 1989.
- [DAU84] DAUBISSE J.-C., *Sur quelques méthodes numériques spécifiques de l'hydrodynamique navale*. Thèse, Ecole Nationale Supérieure de Mécanique, Nantes, juillet 1984.
- [DB78] DE BOOR C., *A practical Guide to Splines*, Vol. 27 de *Applied Mathematical Sciences*. Springer Verlag, 1978.
- [DB93] DE BOOR C., “B(asic)-Spline Basics”, dans *Fundamental Developments of Computer-Aided Geometric Modeling*(PIEGL L., ed.), pp. 27–49, Academic Press, 1993.
- [DC93] DE CASTELJAU P., “Polar forms for curve and surface modeling as used at Citroen”, dans *Fundamental Developments of Computer-Aided Geometric Modeling*(PIEGL L., ed.), pp. 1–12, Academic Press, 1993.
- [Dun93] Dunod, *Pratique des matériaux industriels*, 1993. Tome 2 (Propriétés - Choix - Utilisation), chap. 4.4.1 à 4.4.3.
- [EWI69] EWINS D., “The effect of detuning upon the forced vibrations of bladed disks”, *Journal of Sound and Vibration*, Vol. 9, No. 1, pp. 65–79, 1969.
- [FAR93] FARCY R., *Applications des lasers*. Masson, 1993.

- [GOL89] GOLUB G. ET LOAN C. V., *Matrix computations, second edition*, Vol. 3 de *Johns Hopkins Series in the Mathematical Sciences*. North Oxford Academic, 1989.
- [GOU88] GOUPY J., *La méthode des plans d'expérience*. Bordas, 1988.
- [GUI99a] GUILLOTEAU I., *Modélisation du contact en dynamique: construction d'un élément simplifié de contact et application à l'interaction rotor/stator*. Thèse, Ecole Centrale Nantes, Nantes, novembre 1999.
- [GUI99b] GUILLOTEAU I., PESEUX B., ET BONINI J., "Modélisation du contact rotor-stator en dynamique", dans *Actes du quatrième Colloque National en Calcul des Structures*, Vol. 1, pp. 251 – 256, CSMA, 18 - 21 mai 1999. Giens (Var).
- [HAL90] HALL W. ET HIBBS T., "Continuous, quadrilateral, and triangular surface patches", dans *Applied surface modelling*(CREASY C. ET GILES M., eds.), Chap. 12, pp. 130 – 138, New York: Ellis Horwood, 1990.
- [HAL95] HALLQUIST J., *LS-DYNA3D, Theoretical manual*. Dynalys, 1995.
- [HAM91] HAMMERLIN G. ET HOFFMANN K.-H., *Numerical Mathematics*, Chap. 6, pp. 229 – 271. Undergraduate Texts in Mathematic, Springer Verlag, 1991.
- [HAR91] HARTMANN F., *Les Lasers*, Vol. 1565 de *Que sais-je ?* Presses Universitaires de France, 1991.
- [HIL90] HILL D., FLETCHER E., MOSCARDINI A., ET WILKINSON T., "Overhauser patches for generally curved surface", dans *Applied surface modelling*(CREASY C. ET GILES M., eds.), Chap. 13, pp. 139 – 149, New York: Ellis Horwood, 1990.
- [KAS98] KASHIWAGI M., "A B-spline Galerkin scheme for calculating the hydrostatic response of a very large floating structure in waves", *Journal of Marine Science and Technology*, Vol. 3, pp. 37–49, 1998.
- [LEI73] LEISSA A., "Vibrations of shells", Rapport Techn. SP288, NASA, Washington D.C., 1973. Scientific and Technical Information Office.
- [MAL98] MALGRAS E., *Contribution à la détermination de l'intersection de surfaces paramétriques par des techniques de suivi de contours*. Thèse, Ecole Centrale Nantes, Nantes, janvier 1998.
- [NUR89] NURNBERGER G., *Approximation by Spline Functions*. Springer Verlag, 1989.
- [PAT98] PATSKO N. ET SUBBOTIN Y., "B-splines in the finite element method", *Computational Mathematics and Mathematical Physics*, Vol. 38, No. 1, pp. 11–20, 1998.

- [PAT99] PATEL B., GANAPATHI M., ET SARAVANAN J., “Shear flexible field-consistent curved spline beam element for vibration analysis”, *International Journal for Numerical Methods in Engineering*, Vol. 46, pp. 387–407, 1999.
- [RIE93] RIESENFELD R. F., “Modeling with NURBS curves and surfaces”, dans *Fundamental Developments of Computer-Aided Geometric Modeling*(PIEGL L., ed.), pp. 77–97, Academic Press, 1993.
- [ROT94] ROTHBERG S. ET HALLIWELL N., “Vibration measurements on rotating machinery using laser Doppler velocimetry”, *Journal of Vibration and Acoustics*, Vol. 116, pp. 326–331, 1994.
- [ROT96] ROTHBERG S. ET COUPLAND J., “The new laser Doppler accelerometer for shock and vibration measurement”, *Optics and Lasers in Engineering*, Vol. 25, pp. 217–225, 1996.
- [SAM97] SAMTECH, *SAMCEF, Manuel utilisateur*, 1997. version 7.1.
- [SCH46a] SCHOENBERG I. J., “Contributions to the problem of approximation of equidistant data by analytic functions. Part A: on the problem of smoothing or graduation. A first class of analytic approximation formulae”, *Quarterly of Applied Mathematics*, Vol. 4, pp. 45–99, 1946.
- [SCH46b] SCHOENBERG I. J., “Contributions to the problem of approximation of equidistant data by analytic functions. Part B: on the problem of osculatory interpolation. A second class of analytic approximation formulae”, *Quarterly of Applied Mathematics*, Vol. 4, pp. 112–141, 1946.
- [SHE97] SHENE C., *Introduction to computing with geometry*. course CS390-2, Department of Computer Science: Michigan Technological University, 1997. <http://www.cs.mtu.edu/~shene/COURSES/cs390/NOTES/notes.html>.
- [STA90] STAPLES B. C., “Excitation of travelling wave response in axi-symmetric structures”, dans *Proceedings of the 15th Seminar on Modal Analysis*, pp. 1339 – 1354, 19-21 septembre 1990. Leuven (Belgium).
- [STA95] STANBRIDGE A. ET EWINS D., “Modal testing of rotating discs using a scanning LDV”, *Design Engineering Technical Conferences*, Vol. 3, pp. 1207–1213, 1995. DE-Vol. 84-2.
- [tec] *Techniques de l’Ingénieur*. Manuel M4, pp. M 1165-10.
- [TOB57] TOBIAS S. ET ARNOLD R., “The influence of dynamical imperfection on the vibration of rotating disks”, *Proceedings Institute of Mechanical Engineers*, Vol. 171, pp. 669 – 690, 1957.

-
- [YUE99] YUEN S. ET ERP G. V., “Transient analysis of thin-walled structures using macro spline finite elements”, *Engineering Structures*, Vol. 21, pp. 255–266, 1999.
- [ZEN99] ZENNECK J., “Über die freien Schwingungen nur angenähernd vollkommener kreisförmiger Platten”, *Annalen der Physik*, Vol. 67, pp. 165–184, 1899.

Annexe A

Expression analytique complète des B-splines élémentaires de bas degré

Sont regroupées dans cette annexe les expressions des B-splines élémentaires pour les degrés 0 à 3.

B-spline de degré 0

$$B_{00}(t) = 1 \text{ si } t \in [t_0, t_1]$$

Sur la figure A.1 est représentée une B-spline de degré 0 pour le vecteur de noeuds $(0, 1)$.

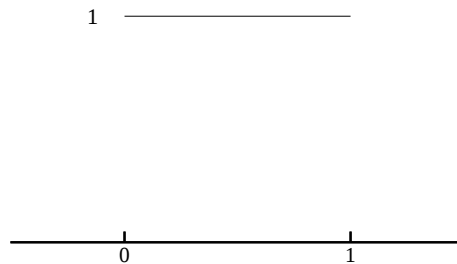


Figure A.1 : B-splines de degré 0

B-spline de degré 1

$$B_{10}(t) = \frac{t - t_0}{t_1 - t_0} \text{ si } t \in [t_0, t_1]$$

$$B_{10}(t) = \frac{t_2 - t}{t_2 - t_1} \text{ si } t \in [t_1, t_2]$$

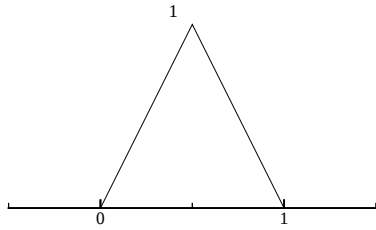
cas particulier d'une subdivision régulière:

si $\forall i \in \{0, 1\}$ on a $t_{i+1} - t_i \triangleq h$ alors:

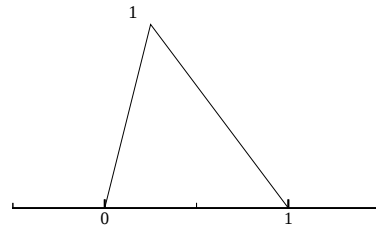
$$B_{10}(t) = \frac{t - t_0}{h} \text{ si } t \in [t_0, t_1]$$

$$B_{10}(t) = \frac{t_2 - t}{h} \text{ si } t \in [t_1, t_2]$$

Sur la figure A.A, le vecteur de nœuds correspond à une subdivision régulière: $(0; \frac{1}{2}; 1)$;
sur la figure A.A, la subdivision n'est pas régulière: $(0; 0.25; 1)$.



(a) Subdivision régulière



(b) Subdivision irrégulière

Figure A.2 : B-spline de degré 1

B-spline de degré 2

$$B_{20}(t) = \frac{(t - t_0)^2}{(t_2 - t_0)(t_1 - t_0)} \text{ si } t \in [t_0, t_1]$$

$$B_{20}(t) = \frac{(t - t_0)(t_2 - t)}{(t_2 - t_0)(t_2 - t_1)} + \frac{(t_3 - t)(t - t_1)}{(t_3 - t_1)(t_2 - t_1)} \text{ si } t \in [t_1, t_2]$$

$$B_{20}(t) = \frac{(t_3 - t)^2}{(t_3 - t_1)(t_3 - t_2)} \text{ si } t \in [t_2, t_3]$$

cas particulier d'une subdivision régulière:

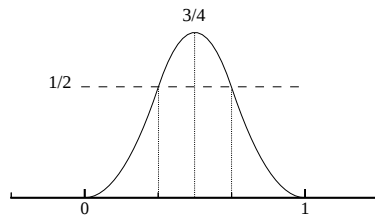
si $\forall i \in \{0, 1, 2\}$ on a $t_{i+1} - t_i \triangleq h$ alors:

$$B_{20}(t) = \frac{(t - t_0)^2}{2h^2} \text{ si } t \in [t_0, t_1]$$

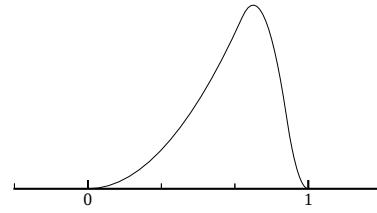
$$B_{20}(t) = \frac{h^2 + 2h(t - t_1) - 2(t - t_1)^2}{2h^2} \quad \text{si } t \in [t_1, t_2]$$

$$B_{20}(t) = \frac{(t_3 - t)^2}{2h^2} \quad \text{si } t \in [t_2, t_3]$$

Sur la figure A.A, le vecteur de nœuds correspond à une subdivision régulière: $(0; \frac{1}{3}; \frac{2}{3}; 1)$; sur la figure A.A, la subdivision n'est pas régulière: $(0; 0.7; 0.9; 1)$.



(a) Subdivision régulière



(b) Subdivision irrégulière

Figure A.3 : B-spline de degré 2

B-spline de degré 3

$$B_{30}(t) = \frac{(t - t_0)^3}{(t_3 - t_0)(t_2 - t_0)(t_1 - t_0)} \quad \text{si } t \in [t_0, t_1]$$

$$B_{30}(t) = \frac{(t - t_0)^2(t_2 - t)}{(t_3 - t_0)(t_2 - t_0)(t_2 - t_1)} + \frac{(t_3 - t)(t - t_1)(t - t_0)}{(t_3 - t_1)(t_2 - t_1)(t_3 - t_0)} \\ + \frac{(t - t_1)^2(t_4 - t)}{(t_3 - t_1)(t_2 - t_1)(t_4 - t_1)} \quad \text{si } t \in [t_1, t_2]$$

$$B_{30}(t) = \frac{(t_3 - t)^2(t - t_0)}{(t_3 - t_0)(t_3 - t_1)(t_3 - t_2)} + \frac{(t_4 - t)(t_3 - t)(t - t_1)}{(t_4 - t_1)(t_3 - t_1)(t_3 - t_2)} \\ + \frac{(t_4 - t)^2(t - t_2)}{(t_4 - t_1)(t_4 - t_2)(t_3 - t_2)} \quad \text{si } t \in [t_2, t_3]$$

$$B_{30}(t) = \frac{(t_4 - t)^3}{(t_4 - t_1)(t_4 - t_2)(t_4 - t_3)} \quad \text{si } t \in [t_3, t_4]$$

cas particulier d'une subdivision régulière:

si $\forall i \in \{0, 1, 2, 3\}$ on a $t_{i+1} - t_i \triangleq h$ alors:

$$B_{30}(t) = \frac{(t - t_0)^3}{6h^3} \text{ si } t \in [t_0, t_1]$$

$$B_{30}(t) = \frac{h^3 + 3h^2(t - t_1) + 3h(t - t_1)^2 - 3(t - t_1)^3}{6h^3} \text{ si } t \in [t_1, t_2]$$

$$B_{30}(t) = \frac{h^3 + 3h^2(t_3 - t) + 3h(t_3 - t)^2 - 3(t_3 - t)^3}{6h^3} \text{ si } t \in [t_2, t_3]$$

$$B_{30}(t) = \frac{(t_4 - t)^3}{6h^3} \text{ si } t \in [t_3, t_4]$$

Sur la figure A.A, le vecteur de nœuds correspond à une subdivision régulière: $(0; \frac{1}{4}; \frac{1}{2}; \frac{3}{4}; 1)$; sur la figure A.A, la subdivision n'est pas régulière: $(0; 0.1; 0.3; 0.5; 1)$. Enfin, sur la figure A.A est représentée une B-spline dont le vecteur de nœuds est : $(0; 1; 1; 1; 1)$ (accumulation en 1).

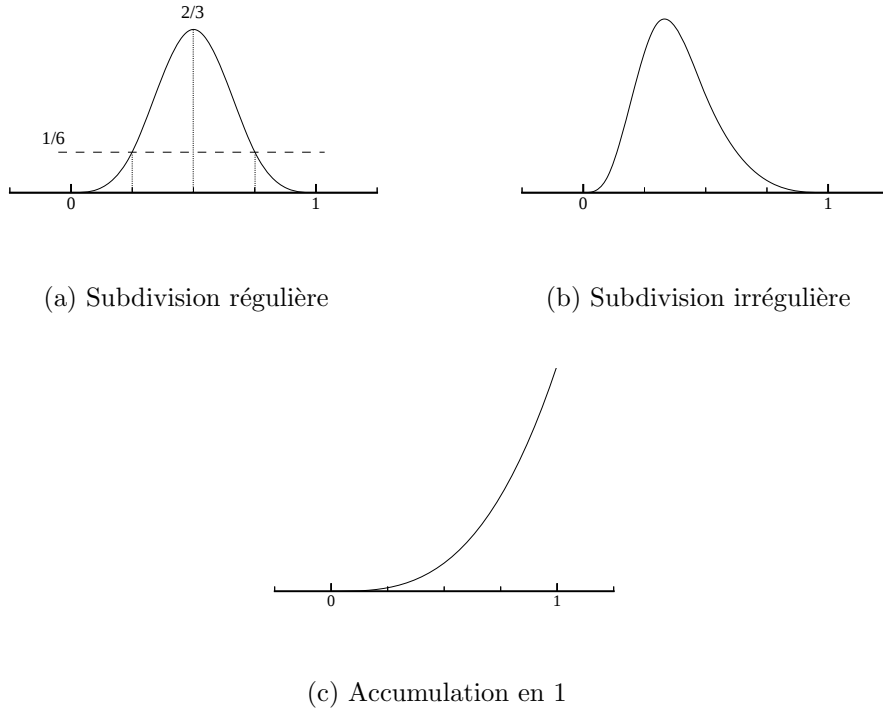


Figure A.4 : B-spline de degré 3

Annexe B

Projection d'un point sur une droite ou un plan

Cette annexe rappelle le détail de la recherche de la projection d'un point sur une droite ou un plan.

Point sur droite (projection orthogonale)

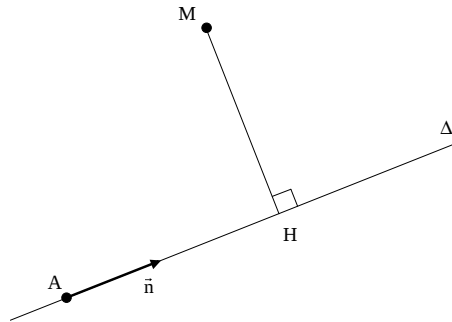


Figure B.1 : Projection point sur droite

Soit M un point de coordonnées (x_0, y_0, z_0) ; soit un point $A(x_A, y_A, z_A)$ et un vecteur (non nécessairement unitaire) $\vec{n}(n_x, n_y, n_z)$ définissant une droite Δ de façon paramétrique, c'est-à-dire que si $P(x, y, z)$ est un point de Δ , alors il existe un réel λ tel que:

$$\begin{cases} x = x_A + \lambda n_x \\ y = y_A + \lambda n_y \\ z = z_A + \lambda n_z \end{cases}$$

Soit H la projection orthogonale de M sur Δ ; H étant un point de Δ , on peut noter λ_H le paramètre qui lui est associé. En écrivant la relation:

$$\overrightarrow{HM} \cdot \vec{n} = 0$$

il vient:

$$\lambda_H = \frac{(x_0 - x_A)n_x + (y_0 - y_A)n_y + (z_0 - z_A)n_z}{\|\vec{n}\|^2}$$

La distance de projection est alors facilement accessible:

$$\|\vec{HM}\|^2 = (x_A - x_0 + \lambda_H n_x)^2 + (y_A - y_0 + \lambda_H n_y)^2 + (z_A - z_0 + \lambda_H n_z)^2$$

Point sur plan (projection orthogonale)

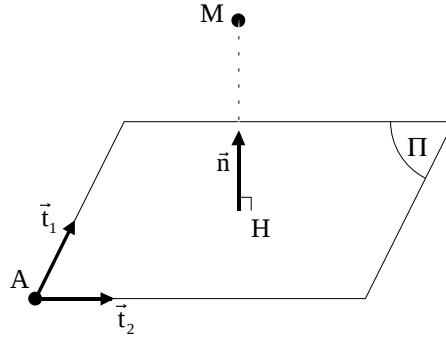


Figure B.2 : Projection orthogonale point sur plan

Soit M un point de coordonnées (x_0, y_0, z_0) ; soit un point $A(x_A, y_A, z_A)$ et deux vecteurs (non nécessairement unitaires) $\vec{t}_1(t_{1x}, t_{1y}, t_{1z})$ et $\vec{t}_2(t_{2x}, t_{2y}, t_{2z})$ définissant un plan Π de façon paramétrique, c'est-à-dire que si $P(x, y, z)$ est un point de Π , alors il existe deux réels λ et μ tels que:

$$\begin{cases} x = x_A + \lambda t_{1x} + \mu t_{2x} \\ y = y_A + \lambda t_{1y} + \mu t_{2y} \\ z = z_A + \lambda t_{1z} + \mu t_{2z} \end{cases}$$

Un vecteur $\vec{n}(n_x, n_y, n_z)$ normal au plan Π est donné par:

$$\vec{n} = \vec{t}_1 \wedge \vec{t}_2$$

Soit Δ la droite définie par M et $\vec{n}$. Un point $P(x, y, z)$ est un point de Δ s'il existe un réel η tel que:

$$\begin{cases} x = x_0 + \eta n_x \\ y = y_0 + \eta n_y \\ z = z_0 + \eta n_z \end{cases}$$

Chercher la projection de M sur Π revient à chercher l'intersection de Δ et de Π c'est-à-dire à résoudre le système suivant dont les inconnues sont les paramètres de repérage du point projeté H :

$$\begin{bmatrix} t_{1x} & t_{2x} & -n_x \\ t_{1y} & t_{2y} & -n_y \\ t_{1z} & t_{2z} & -n_z \end{bmatrix} \cdot \begin{Bmatrix} \lambda_H \\ \mu_H \\ \eta_H \end{Bmatrix} = \begin{Bmatrix} x_0 - x_A \\ y_0 - y_A \\ z_0 - z_A \end{Bmatrix}$$

La résolution donne:

$$\begin{Bmatrix} \lambda_H \\ \mu_H \\ \eta_H \end{Bmatrix} = -\frac{1}{\vec{n}^2} \begin{bmatrix} t_{2z}n_y - t_{2y}n_z & t_{2x}n_z - t_{2z}n_x & t_{2y}n_x - t_{2x}n_y \\ t_{1y}n_z - t_{1z}n_y & t_{1x}n_z - t_{1z}n_x & t_{1x}n_y - t_{1y}n_x \\ -n_x & -n_y & -n_z \end{bmatrix} \cdot \begin{Bmatrix} x_0 - x_A \\ y_0 - y_A \\ z_0 - z_A \end{Bmatrix}$$

La distance de projection est alors donnée par:

$$\|\overrightarrow{HM}\|^2 = (\eta_H)^2 \cdot \|\vec{n}\|^2$$

Point sur plan (direction donnée)

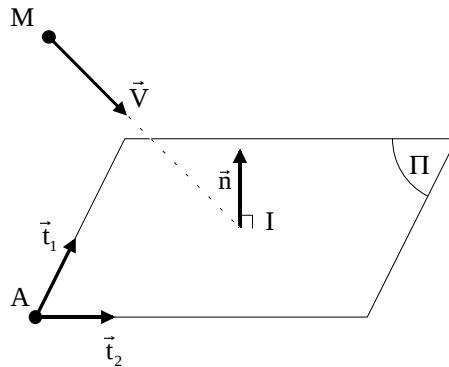


Figure B.3 : Projection point sur plan suivant un vecteur donné

Soit un point $A(x_A, y_A, z_A)$ et deux vecteurs (non nécessairement unitaires) $\vec{t}_1(t_{1x}, t_{1y}, t_{1z})$ et $\vec{t}_2(t_{2x}, t_{2y}, t_{2z})$ définissant un plan Π de façon paramétrique, c'est-à-dire que si $P(x, y, z)$ est un point de Π , alors il existe deux réels λ et μ tels que:

$$\begin{cases} x = x_A + \lambda t_{1x} + \mu t_{2x} \\ y = y_A + \lambda t_{1y} + \mu t_{2y} \\ z = z_A + \lambda t_{1z} + \mu t_{2z} \end{cases}$$

Le vecteur $\vec{n}(n_x, n_y, n_z)$ normal au plan Π est donné par:

$$\vec{n} = \vec{t}_1 \wedge \vec{t}_2$$

Soit Δ la droite définie par $M(x_0, y_0, z_0)$ et $\vec{V}(V_x, V_y, V_z)$ où $\vec{V}$ est un vecteur quelconque. Un point $P(x, y, z)$ est un point de Δ s'il existe un réel η tel que:

$$\begin{cases} x = x_0 + \eta V_x \\ y = y_0 + \eta V_y \\ z = z_0 + \eta V_z \end{cases}$$

Chercher la projection de M sur Π suivant $\vec{V}$ revient à chercher l'intersection de Δ et de Π c'est-à-dire à résoudre le système suivant dont les inconnues sont les paramètres de repérage du point d'intersection I :

$$\begin{bmatrix} t_{1x} & t_{2x} & -V_x \\ t_{1y} & t_{2y} & -V_y \\ t_{1z} & t_{2z} & -V_z \end{bmatrix} \cdot \begin{Bmatrix} \lambda_I \\ \mu_I \\ \eta_I \end{Bmatrix} = \begin{Bmatrix} x_0 - x_A \\ y_0 - y_A \\ z_0 - z_A \end{Bmatrix}$$

La résolution donne:

$$\begin{Bmatrix} \lambda_I \\ \mu_I \\ \eta_I \end{Bmatrix} = -\frac{1}{\vec{n} \cdot \vec{V}} \begin{bmatrix} t_{2z}V_y - t_{2y}V_z & t_{2x}V_z - t_{2z}V_x & t_{2y}V_x - t_{2x}V_y \\ t_{1y}V_z - t_{1z}V_y & t_{1z}V_x - t_{1x}V_z & t_{1x}V_y - t_{1y}V_x \\ -n_x & -n_y & -n_z \end{bmatrix} \cdot \begin{Bmatrix} x_0 - x_A \\ y_0 - y_A \\ z_0 - z_A \end{Bmatrix}$$

Annexe C

Cas des courbes planes cartésiennes

Cette annexe présente le cas des courbes planes cartésiennes (de la forme $y = f(x)$) approchées par des splines.

Généralités

La formulation spline à l'aide de relations paramétriques se traduit par l'établissement d'une relation entre un ensemble de réels (les paramètres) et un ensemble de points de l'espace ($\mathbb{R}^3$ pour les splines, et $\mathbb{R}^4$ pour les NURBS). La représentation cartésienne a été délibérément laissée de côté, principalement parce qu'elle ne permet pas une manipulation souple de toutes les courbes rencontrées (courbes fermées, tangentes quelconques...). Ce paragraphe rappelle les restrictions et les implications induites par une modélisation cartésienne.

Supposons que nous cherchions à manipuler une courbe représentée par une série ordonnée de $npt + 1$ points $(x_i, y_i)_{i=0, npt}$ avec des splines. L'idée de base consiste à conserver la formulation vue dans le cas paramétrique en utilisant la coordonnée x comme paramètre, c'est-à-dire que l'expression analytique de la courbe se mettra sous la forme (n est le degré de modélisation retenu: l'espace de travail sera donc noté $\mathcal{S}_n(\Omega_{npt})$ comme dans le paragraphe I.2.1):

$$y(x) = \sum_{i=-n}^{npt-1} \alpha_i B_{ni}(x)$$

où les (α_i) sont des réels qui jouent le rôle des points de contrôle.

Ainsi, les splines élémentaires B_{ni} seront-elles construites directement à partir du paramétrage (x_i) , suivant la formule de récurrence exposée page 16. Il faut noter cependant que pour avoir la base complète des B-splines sur l'intervalle étudié $[x_0, x_{npt}]$, il faut rajouter n points à chaque extrémité de l'intervalle, donc créer les valeurs $x_{-n}, \dots, x_{-1}$ avant x_0 , et $x_{npt+1}, \dots, x_{npt+n}$ après x_{npt} . Ces valeurs supplémentaires ne sont reliées à aucune valeur de y .

Détermination des points de contrôle

Comme dans le cas des splines paramétriques, les points de contrôle peuvent être calculés de différentes façons: méthode directe, interpolation ou lissage (la liste n'est pas exhaustive). Seul le cas de l'interpolation sera développé ici.

Contrairement à ce que pourrait laisser penser jusqu'à présent la présentation du cas cartésien, sa manipulation par les splines ne se résume pas simplement à prendre l'une des coordonnées pour paramètre. Quelques difficultés vont en effet apparaître.

La dimension de l'espace de travail $\mathcal{S}_n(\Omega_{npt})$ est égale à $npt + n$ (voir le paragraphe I.2.1). Il y a donc $npt + n$ coefficients α_i à déterminer. Les relations d'interpolation fournissent $npt + 1$ équations:

$$\forall i \in [0..npt] \quad y_i = \sum_{j=-n}^{npt-1} \alpha_j B_{nj}(x_i)$$

Il reste donc $n-1$ équations à écrire, et différentes manières peuvent être utilisées. Bien souvent, ce sont les conditions aux limites qui servent à écrire les relations manquantes. Mais deux cas seront à envisager:

- ◇ **n est impair:** dans ce cas $n-1$ est pair, et il est possible d'écrire $(n-1)/2$ conditions à chaque extrémité de la courbe (en x_0 et en x_{npt});
- ◇ **n est pair:** dans ce cas $n-1$ est impair, et il n'est pas possible d'attribuer le même nombre de conditions aux limites à chaque extrémité.

Modélisation avec des degrés impairs

C'est le cas le plus simple. On posera $n = 2p-1$ par la suite. Trois types de conditions aux limites peuvent être utilisées:

1. **Conditions aux limites naturelles:** c'est le cas où les dérivées de la courbe aux points extrêmes sont toutes nulles (aucune contrainte); les conditions aux limites s'écrivent donc:

$$\begin{aligned} y^k(x_0) &= 0 \\ y^k(x_{npt}) &= 0 \text{ pour } k \in [p, ..2p-2] \end{aligned}$$

2. **Conditions aux limites de Hermite:** contrairement au cas précédent, les valeurs des dérivées de la courbe aux points extrêmes sont imposées:

$$\begin{aligned} y^k(x_0) &= f^k(x_0) \\ y^k(x_{npt}) &= f^k(x_{npt}) \text{ pour } k \in [1, .., p-1] \end{aligned}$$

3. **Conditions aux limites périodiques:** elles correspondent au cas où le jeu de données contient 1 période d'une fonction, et il faut donc écrire les conditions de périodicité de la fonction en ses extrémités:

$$y^k(x_0) = y^k(x_{npt}) \text{ pour } k \in [1, .., 2p-2]$$

Modélisation avec des degrés pairs

La dimension de $\mathcal{S}_n(\Omega_{npt})$ est $npt + n$. Compte-tenu des $npt + 1$ relations issues de l'interpolation, il reste $n - 1$ équations à écrire, et n étant pair, $n - 1$ est impair. Il faut donc symétriser le problème. Pour cela, on procède de la façon suivante. Soit:

$$\Omega_{npt-1} = \{\xi_0 = x_0, \xi_1, \dots, \xi_{npt-2}, \xi_{npt-1} = x_{npt}\}$$

Les ξ_i sont choisis de sorte que entre deux valeurs successives de x_i , il se trouve toujours une valeur ξ_i :

$$\xi_0 = x_0 < x_1 < \xi_1 < x_2 < \dots < \xi_{npt-2} < x_{npt-1} < x_{npt} = \xi_{npt-1}$$

Ainsi, $\mathcal{S}_n(\Omega_{npt-1})$ est de dimension $npt + n - 1$. La résolution du problème va donc changer légèrement. Il s'agit à présent de trouver une fonction spline définie sur $\mathcal{S}_n(\Omega_{npt-1})$ (donc définie par $npt + n - 1$ points de contrôle) telle que:

$$\forall i \in [0..npt] \quad y_i = \sum_{j=-n}^{npt-2} \alpha_j B_{nj}(x_i)$$

Il convient de noter que, dans ce cas, les B-splines sont définies sur le paramétrage (ξ_i) et non plus sur (x_i) . Par ailleurs, si $n = 2$, il n'y a pas besoin d'écrire d'équations supplémentaires, car les relations d'interpolation sont au nombre de $npt + 1$, ce qui correspond exactement au nombre de points de contrôle recherchés.

Il est bien sûr possible de procéder de la façon contraire, en construisant un paramétrage (ξ_i) qui englobe les (x_i) , en prenant:

$$\Omega_{npt+1} = \{\xi_0 = x_0, \xi_1, \dots, \xi_{npt}, \xi_{npt+1} = x_{npt}\}$$

de telle sorte que:

$$\xi_0 = x_0 < \xi_1 < x_1 < \xi_2 < \dots < x_{npt-1} < \xi_{npt} < x_{npt} = \xi_{npt+1}$$

Dans ce cas, les relations d'interpolation s'écriront:

$$\forall i \in [0..npt] \quad y_i = \sum_{j=-n}^{npt} \alpha_j B_{nj}(x_i)$$

Il y a donc $npt + 1$ relations dans un espace qui nécessite $npt + n + 1$ points de contrôle. Il reste donc n (pair) équations à écrire, en utilisant les conditions aux limites par exemple.

Lien avec la formulation paramétrée

Comme cela a été précisé, le principe de base de la formulation cartésienne consiste à utiliser une des coordonnées pour établir le paramétrage. Mais la formulation cartésienne

permet de fixer des conditions aux limites au gré de l'utilisateur, ce qui n'est pas le cas de la formulation paramétrée (dans la forme présentée).

Néanmoins, les deux approches se rejoignent dans le cas des conditions aux limites naturelles. En effet, dans le cas paramétré, la constitution des vecteurs de nœuds induit la tangence de la courbe résultat au polygone de contrôle en ses extrémités (voir figure C.1). Cette configuration est donc semblable à celle obtenue avec une spline cartésienne en utilisant les conditions aux limites naturelles.

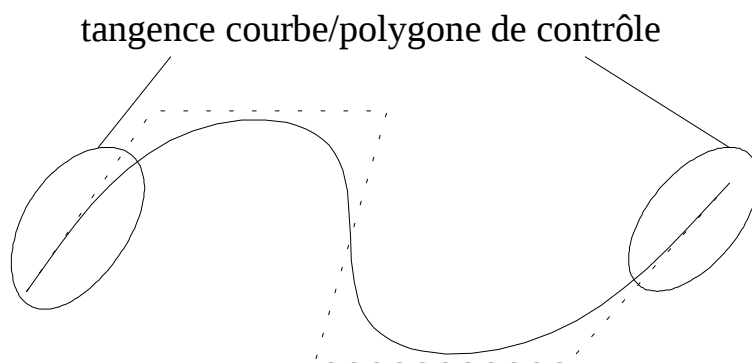


Figure C.1 : Spline paramétrée

La configuration réciproque, c'est-à-dire la construction d'une spline paramétrée respectant des conditions aux limites fixées par l'utilisateur, est plus difficile à établir, car elle suppose la maîtrise de l'influence du vecteur de nœuds sur le résultat. Ce point ne sera pas développé ici.

Annexe D

Algorithme de recherche du contact

L'algorithme complet de recherche du contact est présent dans cette annexe.

Cet algorithme traite la découverte des zones de contact possibles, la gestion des bornes de ces zones ainsi que les points multiples (cas où à un même t correspond plusieurs zones distinctes (u, v)). Pour cela, de nombreux indicateurs sont nécessaires :

- ◇ la variable ind_{zone} est nulle en dehors des zones où l'écart courbe/surface est supérieur à ε , et vaut 1 quand cet écart est inférieur à ε .
- ◇ la variable ind_u vaut 1 si, pour une valeur de t donnée, un balayage complet d'une zone a été effectué sans trouver aucun écart courbe/surface inférieur à ε ; elle vaut 0 sinon.
- ◇ la variable ind_v vaut 1 si, pour une valeur de t donnée, un balayage complet d'un intervalle $[v_{min}, v_{max}]$ a été effectué sans trouver aucun écart courbe/surface inférieur à ε ; elle vaut 0 sinon.
- ◇ la variable ind_{mult} vaut 0 tant que la variable ind_v vaut 0; elle vaut 1 quand une zone a été découverte; elle vaut 2 quand, pour un même paramètre t , une seconde zone a été découverte.
- ◇ la variable $flag$ vaut 1 quand une première zone est explorée: elle autorise alors la modification des bornes des paramètres; elle vaut 0 si, pour un même paramètre t , d'autres zones sont découvertes (c'est-à-dire quand ind_{mult} vaut 2): les bornes de ces nouvelles zones ne sont alors pas enregistrées.

Les numéros qui figurent en italique dans l'algorithme renvoient aux explications ci-après.

Pour chaque zone z_i déterminée à l'étape précédente

 Lire les bornes en t , u et v de z_i et la précision ε

 Calculer les incréments dt , du et dv (1)

$t = t_{min}$

```

 $ind_{zone} = 0, nb_{zone} = 0, ind_{mult} = 0$ 
Tant que  $t \leq t_{max}$  faire
  Calculer  $P_t$ 
   $u = u_{min}$ 
   $mem = u_{min}, flag = 1, ind_u = 1$ 
  Tant que  $u \leq u_{max}$  faire
     $v = v_{min}$ 
     $ind_v = 1$ 
    Tant que  $v \leq v_{max}$  faire
      Calculer  $P_{uv}$ 
      Calculer  $gap = \| P_t P_{uv} \|$ 
      Si  $gap \leq \varepsilon$  et  $ind_{zone} = 0$  et  $ind_{mult} = 0$  alors (2)
         $ind_{zone} = 1, nb_{zone} = nb_{zone} + 1$ 
         $ind_u = 0, ind_v = 0$ 
        Mémoriser les valeurs de  $t, u$  et  $v$ 
        /paramètres d'initiation/
      Fin si
      Si  $gap \leq \varepsilon$  et  $ind_{zone} = 1$  et  $ind_{mult} \neq 1$  et  $ind_v \neq 2$  alors (3)
         $ind_u = 0, ind_v = 0$ 
        Si  $flag = 1$  alors
          Mémoriser les valeurs de  $t, u$  et  $v$ 
          /paramètres de terminaison/
        Fin si
      Fin si
      Si  $gap \leq \varepsilon$  et  $ind_{zone} = 1$  et  $ind_{mult} = 1$  et  $flag = 0$  alors (4)
         $ind_u = 0, ind_v = 0$ 
         $ind_{mult} = 2$ 
      Fin si
      Si  $gap > \varepsilon$  et  $ind_v = 0$  alors
         $ind_v = 2$ 
      Fin si
      Incréments  $v$ 
    Fin tant que
    Si  $ind_v = 1$  et  $ind_{zone} = 1$  et  $ind_{mult} = 0$  alors (5)
       $ind_{mult} = 1$ 
    Fin si
    Si  $ind_v = 1$  alors (6)
      Si  $(u - mem) \geq (du/2)$  alors
         $flag = 0$ 
      Fin si
       $mem = mem + du$ 
    Fin si
    Incréments  $u$ 

```

```

    Fin tant que
    Si  $ind_u = 1$  alors (7)
         $ind_{zone} = 0, ind_{mult} = 0$ 
    Fin si
    Si  $ind_u = 0$  et  $ind_{zone} = 1$  et  $ind_{mult} = 1$  alors (8)
         $ind_{mult} = 2$ 
    Fin si
    Incréments  $t$ 
Fin tant que
Fin pour

```

Annotations pour l'algorithme:

1. le calcul des incréments est fonction de la valeur de la précision (voir le paragraphe III.1.4.2);
2. dans ce cas, aucune zone n'a encore été créée pour le paramètre t en cours: il s'agit donc d'une initiation. Par conséquent, le compteur nb_{zone} est incrémenté, et les variables ind_u et ind_v sont mises égales à 0;
3. dans ce cas, une zone a déjà été découverte pour le paramètre t en cours; s'il s'agit bien de la première zone pour ce paramètre (c'est-à-dire si $ind_{mult} \neq 1$ et $flag = 1$), les bornes de la zone sont modifiées;
4. ce cas est à rapprocher du cas 3, et traduit que le point associé au paramètre t en cours est un point multiple: aucune modification de zone n'est donc apportée;
5. si un intervalle $[t_{min}, t_{max}]$ a été balayé sans rencontrer aucun contact potentiel, bien que pour le paramètre t en cours, une zone ait déjà été découverte, alors ind_{mult} prend la valeur 1: ainsi, la prochaine zone rencontrée ne sera pas prise en compte puisqu'elle traduira la multiplicité du point associé au paramètre t en cours;
6. mem est une variable qui est incrémentée de du chaque fois qu'un intervalle $[t_{min}, t_{max}]$ a été balayé sans rencontrer aucun contact potentiel; s'il existe un décalage entre la valeur de mem et celle du paramètre u , c'est qu'une zone a été détectée: l'indicateur $flag$ qui pilote l'autorisation d'agrandissement des zones, est alors annulé;
7. pour le paramètre t en cours, une zone entière a été balayée sans trouver de zone de contact: les paramètres ind_{zone} et ind_{mult} sont donc réinitialisés;
8. à la fin du balayage selon u , si une zone a été découverte et non-fermée, le fait de mettre $ind_{mult} = 2$ permettra de continuer à traiter normalement cette zone (si elle ne correspond pas à un point multiple) pour l'incrément t suivant.

Annexe E

Jeu de données PLEXUS

Est précisé dans cette annexe le contenu de la commande *SPLINE* du jeu de données PLEXUS.

Le fichier de données

Le fichier de données PLEXUS est composé de la façon suivante (pour plus de détails, se référer à la notice d'utilisation [CEA99]):

- ◇ définition des fichiers de sortie;
- ◇ importation du maillage */le maillage provient par exemple d'un fichier .mesh généré par CASTEM/* ;
- ◇ dimensionnement du problème */nombre d'éléments, nombre de degrés de liberté, nombre de matériaux... ces grandeurs, si elles sont inconnues, sont fournies par le code après une première exécution "à vide" /* ;
- ◇ choix des éléments et de leurs caractéristiques (épaisseur...) et attribution;
- ◇ définition des matériaux et attribution;
- ◇ définition des liaisons */cette commande regroupe aussi bien les conditions de contact que les conditions aux limites, puisque ces phénomènes sont gérés de façon identique dans le code/* ;
- ◇ définition des conditions initiales;
- ◇ définition des paramètres de calcul */c'est-à-dire le temps initial, le temps final, les options relatives au choix du pas de temps, le nombre maximal d'itérations.../* ;
- ◇ sélection des résultats à archiver;
- ◇ définition des sorties graphiques.

La commande SPLINE

L'appel à l'utilisation des splines pour la détection du contact se fait donc dans la commande LIAISON, en utilisant le mot-clef SPLINE.

Si le problème à traiter implique N_s surfaces (qui doit être une surface de révolution d'axe $\vec{z}$), il faut alors saisir:

- ◇ la valeur de N_s (entier); pour chacune de ces N_s surfaces, il faut entrer ensuite:
- ◇ la définition de la surface (liste de nœuds, déléments...);
- ◇ le nombre N_c de courbes pouvant entrer en contact avec cette surface, et leurs définitions (liste de nœuds, déléments...);
- ◇ la méthode de modélisation retenue pour les courbes et pour la surface, respectivement *metc* et *mets* (entiers dont la valeur est 1 pour la méthode directe, 2 pour l'interpolation et 3 pour le lissage);
- ◇ la valeur *nptu* (entier) qui correspond au nombre de nœuds sur une circonférence du carter;
- ◇ les degrés *degc*, *degst* et *degsh* (entiers) des splines utilisées respectivement pour les courbes, pour le carter dans la direction orthoradiale et pour le carter suivant l'axe de révolution (des valeurs comprises entre 1 et 5 sont acceptées);
- ◇ la valeur *ep* (réel) de l'épaisseur à prendre en compte pour la recherche de contact (*note: cette grandeur est indépendante de la valeur de l'épaisseur des coques définie dans la partie "choix des éléments" du jeu de données, et peut donc en différer*);
- ◇ la valeur *f* (entier) de la fréquence de réactualisation de la définition des nœuds de la surface: cette valeur est un nombre de pas de temps.

Finalement, la commande se présente donc sous la forme¹:

```
LIAISON
  SPLINE  $N_s$ 
    SURFACE /LECTURE/
      COURBE  $N_c$ 
        LIGNE /LECTURE/
          LIGNE /LECTURE/
            ... ( $N_c$  fois en tout)
        METC metc METS mets NPTT nptu EPAIS ep
        DEGC degc DEGSH degsh DEGSZ degsh FREQ f
```

¹la procédure /LECTURE/ permet de saisir des informations sous des formats très divers: ainsi, une courbe peut-elle être définie nœud à nœud (saisie d'une liste de numéros), élément par élément (saisie d'une liste de numéros également), ou bien par référence à un nom issu du fichier de maillage (mot)

SURFACE /LECTURE/
... (N_s fois en tout)

Annexe F

Recherche d'intersection surface/surface

Est rapidement présenté dans cette annexe les quelques développements réalisés autour du problème d'intersection surface/surface.

Pour se ramener au problème déjà traité de la recherche d'intersection courbe/surface, il suffit de considérer qu'une surface est constituée de courbes mises les unes à côté des autres. L'évaluation de la pénétration n'est pas abordée dans ce cas précis.

Principe de la méthode

Soit s et t les paramètres utilisés pour décrire la première surface S_{st} , et soit u et v les paramètres de la seconde surface S_{uv} . Soit L_s la courbe issue de S_{st} en fixant le paramètre s et en faisant varier le paramètre t entre ses valeurs minimum et maximum. Le principe de la recherche d'intersection est similaire à celui utilisé pour l'intersection courbe/surface: à chaque étape, une zone de contact potentiel est étudiée. Des courbes L_s sont extraites, et l'intersection de L_s avec S_{uv} est examinée. Les bornes de la zone sont modifiées en fonction du résultat fourni par l'étude de cette intersection. Ainsi, l'algorithme général du programme associé à cette démarche est:

```
Pour chaque zone déterminée à l'étape précédente
  Lire les bornes en  $s$  et la précision  $\varepsilon$ 
  Calculer l'incrément  $ds$ 
   $s = s_{min}$ 
  Tant que  $s \leq s_{max}$  faire
    Chercher les points d'intersection de  $L_s$  et  $S_{uv}$ 
    Si l'intersection est non vide:
      Mémoriser les points d'intersection
      et redéfinir les bornes
    Fin si
    Incrémenter  $s$ 
  Fin tant que
Fin pour
```

Le résultat se présente donc sous la forme non pas d'une courbe d'intersection, mais d'un ensemble de points communs aux deux surfaces.

Le gros inconvénient de cette méthode réside dans les temps de calcul qu'elle induit. Pour les limiter, un vecteur mémoire est créé qui contient les valeurs des paramètres s pour lesquels la recherche d'intersection a déjà été effectuée au cours d'une étape antérieure.

Par ailleurs, deux précisions sont introduites: la première correspond à la précision requise pour les intersections courbe/surface; la seconde correspond à la finesse de balayage de S_{st} . Il est ainsi possible d'obtenir un nombre restreint de points d'intersection définis avec une très bonne précision.

Exemples

Deux exemples sont présentés ici.

1. **Intersection plan / surface:** le test consiste à chercher l'intersection d'un cylindre parabolique avec un plan perpendiculaire à une génératrice du cylindre (voir figure F.1). La dimension du plan est 10×10 . La précision initiale est 1.0; la précision finale pour les courbes et pour le balayage de la surface est de 0.01. Les points obtenus sont présentés sur la figure F.1 (durée du calcul: 4' 25").

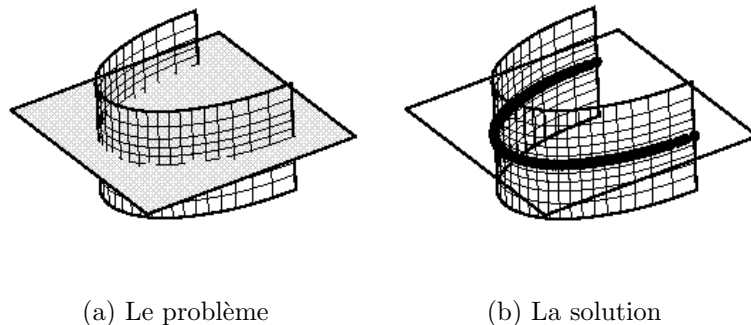


Figure F.1 : Exemple d'intersection plan/surface

2. **Intersection surface / surface:** le test consiste à chercher l'intersection de deux cylindres de diamètres différents, et dont les axes sont concourants (voir figure F.2). Le diamètre du gros cylindre est 8, celui du petit est 5. les résultats sont présentés sur la figure F.2. Cet exemple a été l'occasion de comparer les temps de calcul selon les précisions demandées. La précision initiale est toujours 1.0 alors que la précision finale évolue suivant les cas (voir tableau F.1). Il apparaît qu'il est plus coûteux d'augmenter la précision sur le balayage que sur les points d'intersection.

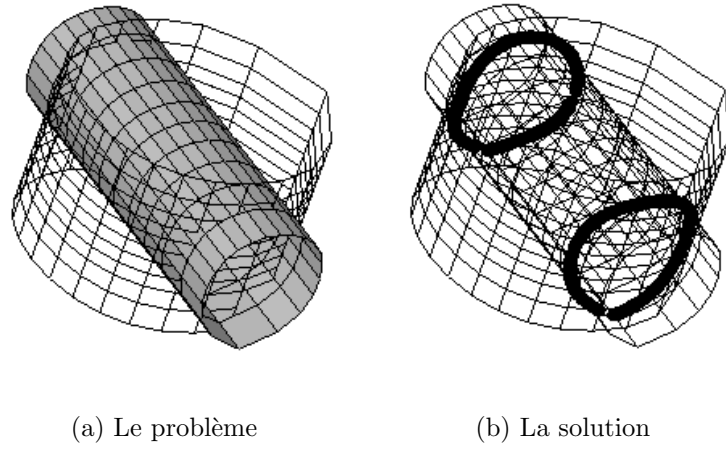


Figure F.2 : Exemple d'intersection surface/surface

Précision finale pour les courbes	Précision finale pour le balayage	Temps de calcul (en minutes-secondes)
0.01	0.01	13' 20"
0.01	0.1	2' 26"
0.001	0.01	15' 13"

Tableau F.1: Comparaison temps de calcul

Annexe G

Montage expérimental

Sont regroupés dans cette annexe les renseignements divers relatifs au montage expérimental.

Matériau du carter

Désignation

[Dun93]

- ◇ *AFNOR*: KC20WN
- ◇ *Haynes*: HS25
- ◇ *Aubert et Duval*: XSH
- ◇ *Astm/Asme (USA)*: F90
- ◇ *Aisi/Uns (USA)*: 5537-5759
- ◇ *Werkstoff (Allemagne)*: 2.4964-2.4967
- ◇ *Aecma (Europe)*: Co-P92HT

Composition

[Dun93][tec]

Il s'agit d'un alliage superréfractaire, utilisé notamment dans les chambres de combustion.

- ◇ *Co*: base
- ◇ *Cr*: 20% (d'où une bonne résistance à la corrosion)
- ◇ *W*: 15%

- ◇ *Ni*: 10%
- ◇ *Fe*: 1,5%
- ◇ *Mn*: 1,5%
- ◇ *Si*: 0,2%
- ◇ *C*: 0,1%

Données mécaniques et physiques

[Dun93]

- ◇ *densité*: 9,13
- ◇ *module d'Young*: 235 000 MPa
- ◇ *limite d'élasticité* $R_{0,002}$: 450 MPa
- ◇ *résistance à la rupture* R_m : 1050 MPa
- ◇ *A* %: 47

Traitement (informations SNECMA)

Le métal est porté à 1205° C pendant 10 à 15 minutes, puis est refroidi à l'air (pas de durcissement structural). Il en résulte une dureté Brinell d'environ 280 (soit environ 300 Vickers), ce qui est une valeur assez élevée (la dureté Brinell d'un acier classique est de l'ordre de 180).

Caractéristiques du tour vertical

Le tour vertical utilisé est un tour Graffenstaden, modèle TV 083. Le diamètre du plateau est de 800 mm, l'encombrement maximum possible d'une pièce ne devant pas dépasser 910 mm de diamètre. Le moteur du tour est un moteur Alsthom (Belfort) asynchrone triphasé de vitesse nominale 1500 tr/min, modèle NP5/77, datant de 1950. La transmission se fait par engrenages.

Les vitesses de rotation du plateau se répartissent en deux gammes:

- ◇ première gamme (vitesse en tr/min): 7,5 - 15 - 30 - 58 - 110 - 210
- ◇ seconde gamme (vitesse en tr/min): 11 - 22 - 42 - 80 - 155 - 300

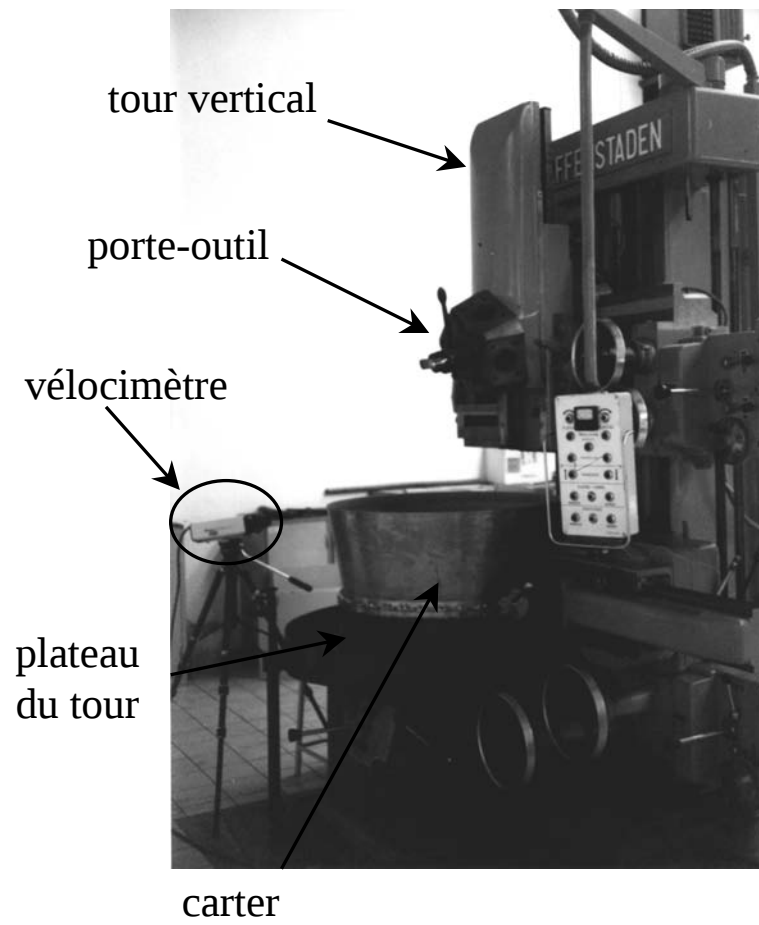


Figure G.1 : Vue générale du banc d'essai

Annexe H

Vélocimétrie laser

L'avantage de pouvoir utiliser un vélocimètre laser réside principalement dans la possibilité d'effectuer des mesures sans contact: il est donc particulièrement intéressant pour des mesures sur des surfaces en rotation, sur des corps difficiles d'accès ou posant des problèmes de température, ainsi que sur de petits éléments (car il n'induit pas d'ajout de masse). A cela, il faut opposer sa relative fragilité (la précision de l'optique est fondamentale), et le coût (conséquence également de la qualité de l'optique).

Le phénomène LASER

Les lasers ont pris dans notre vie quotidienne une place importante, puisqu'on les retrouve dans nombre d'appareils d'usage tant domestiques que professionnels. Les lasers à faible puissance (quelques mW) sont souvent dévolus au transport des informations, que ce soit à travers des appareils de mesure (contrôle de machines outils, mesure de débit sanguin, mesure des longues distances...), de télécommunication ou de design (holographie). Les lasers de plus forte puissance (plusieurs milliers de watts) sont utilisés pour leur aptitude à interagir avec la matière qu'ils contactent (concentration d'énergie très importante sur de petites surfaces, de l'ordre de 1 MW/cm^2): on les retrouve donc dans des procédés qui vont du soudage et de la découpe de matériaux industriels jusqu'aux interventions micro-chirurgicales.

Malgré l'étendue des applications, les lasers sont associés à une technologie très récente, puisque sa finalisation date de la fin des années 50. Le mot LASER lui-même, acronyme pour *Light Amplification by Stimulated Emission of Radiation* (c'est-à-dire littéralement *Amplification de Lumière par Emission Stimulée de Rayonnement*) n'a fait son apparition dans les dictionnaires qu'au début des années 80.

Principe

Les explications apportées ici sont notamment extraites de [CHA90] et [HAR91].

DÉFINITION DU FAISCEAU LASER

Une façon de représenter un faisceau lumineux consiste à dire qu'il s'agit d'un ensemble de particules immatérielles, les photons, se déplaçant à la vitesse de la lumière. Un photon est caractérisé par son énergie $E = h\nu$ et par sa quantité de mouvement $\vec{p} = \frac{h\nu}{c}\vec{i}$ où h est la constante de Planck ($h = 6,625 \cdot 10^{-34} \text{ J.s}$), ν est la fréquence, c est la vitesse de la lumière et $\vec{i}$ désigne la direction de la propagation.

Pour une source lumineuse classique, les photons n'ont pas tous la même fréquence, ni la même quantité de mouvement. Le faisceau laser est défini en opposition à ce type de faisceau, c'est-à-dire que tous les photons ont la même fréquence (le faisceau est dit monochromatique) et même quantité de mouvement (le faisceau est dit directif). Ces propriétés particulières du laser sont une conséquence de son mode de création, qui repose sur le phénomène d'émission stimulée (encore appelée émission induite).

EMISSION STIMULÉE

Comme cela a été précisé, les photons sont notamment caractérisés par leur énergie $E = h\nu$. Albert Einstein a montré en 1917 que l'émission de lumière par un corps correspondait au brusque passage des atomes d'un état excité d'énergie E_2 à un état moins excité d'énergie E_1 , la fréquence du photon émis étant alors égale à:

$$\nu_{12} = \frac{E_2 - E_1}{h}$$

Cette émission peut se produire suivant deux schémas distincts:

1. **émission spontanée** (voir figure H.H): les atomes situés dans le niveau d'énergie E_2 retombent spontanément, aléatoirement et de façon isotrope dans le niveau d'énergie inférieur E_1 , en émettant à chaque fois un photon de fréquence ν_{12} ;
2. **émission stimulée** (voir figure H.H): si un photon de fréquence ν_{12} se trouve déjà présent dans le milieu, il peut induire le passage d'un atome d'énergie E_2 à l'état d'énergie inférieur E_1 , le photon émis ayant non seulement la même fréquence que le photon inducteur, mais également toutes ses autres caractéristiques (direction, polarisation...).

Dans les conditions habituelles, c'est le premier phénomène qui est prépondérant. Afin d'obtenir un faisceau dit LASER, il faut donc s'arranger pour que le processus d'émission induite l'emporte sur l'émission spontanée, ainsi que sur son processus inverse, l'absorption (un atome dans l'état E_1 absorbe un photon de fréquence ν_{12} pour arriver à l'état E_2 , voir figure H.H). La théorie montre que ceci a lieu si la "population" de l'état supérieur E_2 (c'est-à-dire le nombre d'atomes dans cet état présents dans le milieu) est plus élevée que celle de l'état inférieur E_1 . Cette inversion de population est réalisée grâce à un processus de pompage, qui fournit de l'énergie de façon sélective aux atomes du milieu.

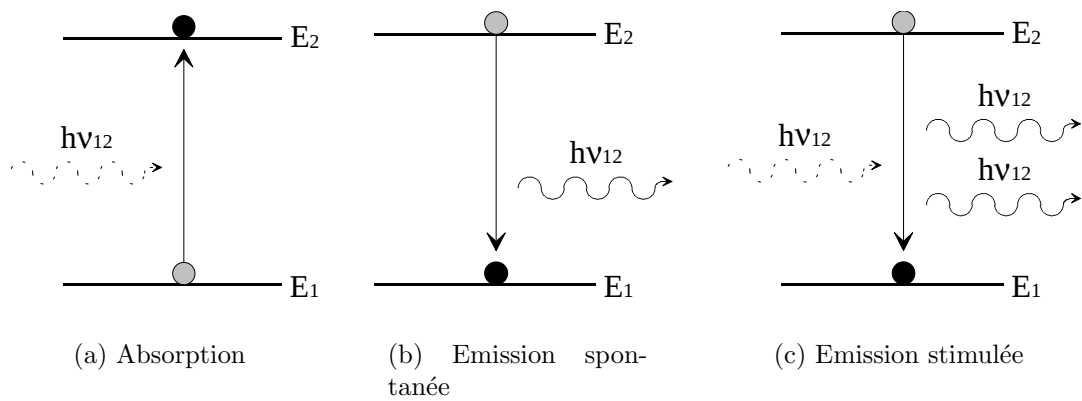
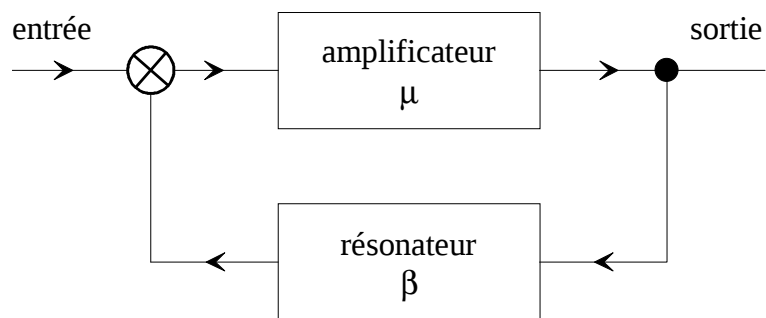


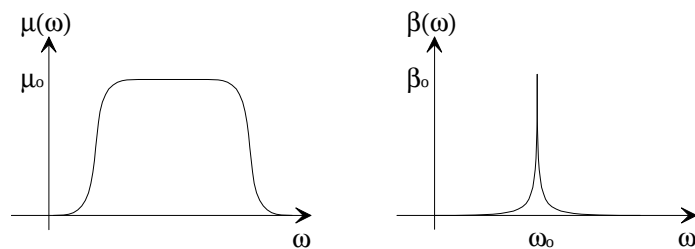
Figure H.1 : Procédés de transitions radiatives entre deux niveaux d'énergie

CRÉATION DU FAISCEAU LASER

Un milieu subissant des pompes, qui facilitent l'émission stimulée, est en fait assimilable à un amplificateur (il est d'ailleurs fréquent d'établir l'analogie entre un laser et un oscillateur électrique, voir figure H.2).



(a) Oscillateur à réaction



(b) Amplification

(c) Résonance

Figure H.2 : Circuit électrique équivalent

Si cet amplificateur est maintenant placé dans une cavité résonnante accordée à la fréquence d'émission, l'effet de réaction conduira à l'apparition d'une oscillation à cette même fréquence. Cette cavité résonnante (voir figure H.3), en général constituée de deux miroirs en vis-à-vis, dont l'un, légèrement transparent, livre passage au faisceau de sortie, imposera notamment à tous les photons d'être émis en phase: le faisceau résultant à la sortie sera donc une onde quasi plane, uni-directionnelle et monochromatique.

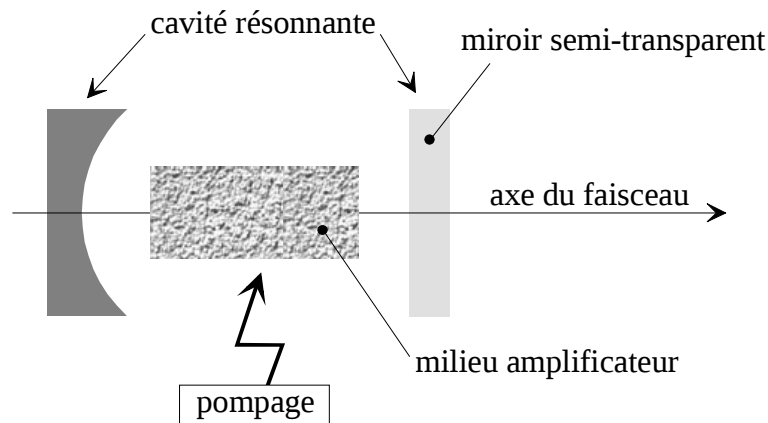


Figure H.3 : Schéma de principe d'un laser

Dans la pratique, de nombreux laser peuvent être rencontrés, dont les longueurs d'onde évoluent depuis l'infra-rouge ($300\ \mu m$) jusqu'à l'ultra-violet ($0,1\ \mu m$). Les milieux actifs qui servent de base à l'amplificateur peuvent être de nature solide (rubis synthétique, néodyme, grenat d'yttrium et d'aluminium (YAG)...) ou gazeuse (hélium-néon, argon ionisé, gaz carbonique...).

Historique

L'histoire de la naissance des faisceaux lasers s'étale sur presque un siècle. En effet, la notion de résonateur existait déjà au siècle passé (résonateur de Pérot-Fabry: 1870), alors que la théorie des amplificateurs ne s'est vraiment développée qu'au XXème siècle. Ce sont les travaux d'Einstein, en 1917, qui exposent le principe d'émission stimulée des atomes sur lequel tout repose. Le pompage optique permettant l'inversion de population a été indiqué par Kastler en 1949. L'idée d'utiliser l'émission stimulée dans un matériau ayant subi une inversion de population pour amplifier des signaux optiques a été émise simultanément et indépendamment en 1953 par les Américains Weber et Townes d'une part (création du MASER, *Microwave Amplifier by Stimulated Emission of Radiation*), et par les Soviétiques Basov et Prokhorov de l'autre. En 1958, Townes et Schawlow concluent qu'il est théoriquement possible de faire fonctionner le dispositif de 1953 à des fréquences très supérieures à celles utilisées jusque là. En juin 1960, l'Américain Maiman utilise le résonateur de Pérot-Fabry, et un rubis synthétique pompé par lampe flash pour donner le jour au premier faisceau LASER.

Effet Doppler et vélocimétrie laser

Pour la mesure, les lasers utilisés sont de faible énergie (inférieure à 1 mW pour respecter les normes).

Principe général

Lorsqu'un corps se déplace suivant une direction $\vec{i}$, avec une vitesse de norme V suivant cette direction, toute onde incidente de fréquence f_0 et de longueur d'onde λ est rétrodiffusée par le corps en une onde de fréquence f_1 différente de f_0 (voir figure H.4). Le décalage en fréquence est connu sous le nom de fréquence Doppler:

$$f_1 - f_0 = f_{Doppler}$$

avec ([FAR93] par exemple):

$$f_{Doppler} = \frac{2V}{\lambda}$$

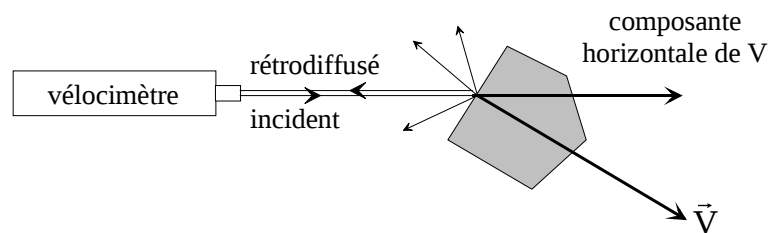


Figure H.4 : Effet Doppler

Les vélocimètres laser utilisent ce décalage pour recueillir des informations en vitesse. Deux difficultés surgissent:

1. la fréquence de l'onde incidente est de l'ordre de 10^{14} Hz, mais la fréquence Doppler induite par le déplacement de la cible n'est que de l'ordre de 10^6 Hz;
2. il n'y a aucune différence entre la fréquence Doppler induite par un solide qui se déplace à la vitesse $+V$ et la fréquence Doppler induite par un solide qui se déplace à la vitesse $-V$: il n'est donc pas possible de prévoir le signe de la vitesse.

Pour ces deux raisons, le faisceau incident est séparé en deux:

- ◇ un faisceau est dirigé tel quel sur la cible;
- ◇ un faisceau reste à l'intérieur de l'appareil, où il subit une modulation (ou hétérodynage), c'est-à-dire un décalage en fréquence provoqué artificiellement soit par un dispositif mécanique (disque à fentes, comme ce fut le cas encore très récemment chez Brüel et Kjær sur le modèle 3544), soit, et c'est le plus fréquent, par un dispositif électrique (cellule de Bragg, cas du modèle OFV 2200 de RMP-Polytec par exemple); ce décalage est de l'ordre de $40 \cdot 10^6$ Hz.

Ainsi, le faisceau rétrodiffusé peut-il être comparé au faisceau modulé conservé dans l'appareil (voir figure H.5). La fréquence obtenue est alors:

$$f = f_{Bragg} + f_{Doppler}$$

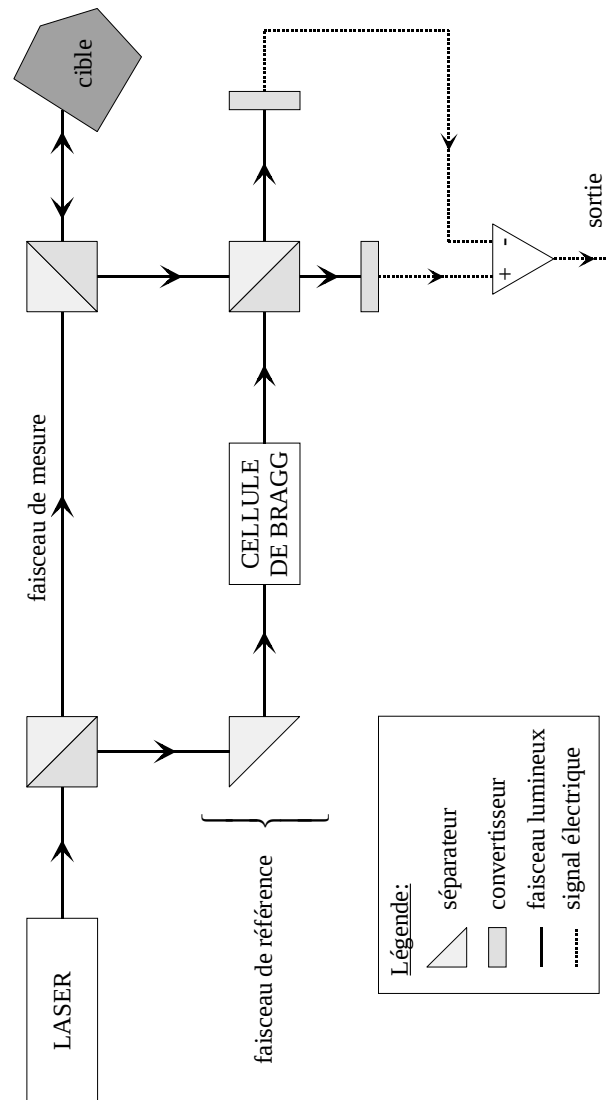


Figure H.5 : Schéma général d'un vélocimètre

Extensions

Les vélocimètres laser offrent souvent la possibilité de fournir des résultats relatifs au déplacement de la cible. Pour cela, au dispositif général de mesure de vitesse vient s'ajouter un petit dispositif indépendant de comptage de franges d'interférence obtenues

en croisant faisceau de référence et faisceau rétrodiffusé. La valeur de la frange multipliée par le nombre de franges fournit une estimation du déplacement.

Depuis quelques années, la société Brüel et Kjær notamment poursuit des études sur le développement d'un appareil de mesure laser permettant d'accéder directement à l'accélération de la cible [ROT96].

Appareil utilisé

L'appareil qui a servi aux mesures est le velocimètre OFV 2200 de la société Allemande RMP-Polytec. Le milieu amplificateur est du gaz (Hélium-Néon), et le faisceau possède une longueur d'onde de 633 nm (ce qui correspond au rouge dans le spectre visible). La puissance de l'appareil est de 0,8 mW (classe II). La gamme des vitesses s'étend jusqu'à 1,25 m/s, avec une dynamique de 90 dB. La gamme de mesure des déplacements va jusqu'à +/- 41 cm, avec une dynamique de 72 dB.

Annexe I

Rappels sur la théorie des plans d'expérience

Cette annexe rappelle quelques résultats de base nécessaires à la mise en œuvre des plans d'expérience. Pour des informations plus complètes, les deux ouvrages suivants seront très utiles: [GOU88] [BEN94b].

Introduction

Soit un système dont la réponse à une sollicitation donnée dépend d'un nombre n de paramètres. Pour étudier l'influence de ces paramètres sur la réponse, une première solution consiste à fixer $n - 1$ paramètres, puis à faire varier le dernier autant de fois que nécessaire. Cette méthode est sûre, mais si chaque paramètre doit prendre p valeurs, il faudra, pour mener à bien l'ensemble de l'étude, réaliser p^n essais. Comme p vaut au moins 2, p^n croît très vite: cette méthode est donc coûteuse.

La méthode des plans d'expérience permet de faire varier plusieurs facteurs en même temps, donc de limiter le nombre d'essais à réaliser, et, par l'analyse des résultats, d'en déduire les influences des paramètres ainsi que leurs interactions respectives sur la réponse.

Notions élémentaires

DÉFINITION

L'étude d'un système pour lequel chacun des n paramètres peut adopter deux valeurs donne lieu à la mise en place d'un plan dit *factoriel à deux niveaux et à n facteurs*: on parle alors de plan 2^n .

Ces facteurs (appelés également variables) seront notés X_i , et les deux valeurs prises par chacun des facteurs seront notées X_i^- et X_i^+ . En fait, la notion de facteur est très large: en particulier, les X_i ne correspondent pas systématiquement à des valeurs numériques, mais peuvent être associés à des états. Ainsi, dans une étude d'un système tournant, X_i peut correspondre à une valeur de vitesse, auquel cas X_i^- et X_i^+ pourraient être les bornes du domaine de variation de la vitesse, mais X_i pourrait également correspondre à un sens

de rotation, auquel cas X_i^- et X_i^+ correspondraient respectivement à un sens positif ou négatif de rotation.

VARIABLES CENTRÉES RÉDUITES

Tous les paramètres étudiés sont ramenés par affinité à l'intervalle $[-1; 1]$: aux facteurs X_i évoluant chacun dans l'intervalle $[X_i^-; X_i^+]$ sont donc substituées les valeurs x_i , dites *centrées réduites*, appartenant à $[-1; 1]$ (appelé domaine de la variable).

MATRICE D'ESSAI

Les essais réalisés dans le cadre des plan factoriels d'expérience à deux niveaux sont, *a priori*, tous les essais possibles à partir des deux niveaux de chacun des facteurs. Ainsi, dans le cas d'un plan à deux facteurs, il y a $2^2 = 4$ essais à réaliser:

Numéro essai	Niveau facteur 1	Niveau facteur 2	Réponse
1	+1	+1	r_1
2	+1	-1	r_2
3	-1	+1	r_3
4	-1	-1	r_4

Tableau I.1: Matrice d'expérience (plan à 2 facteurs)

Il est fréquent de trouver la matrice des essais (et, voir plus loin, la matrice des effets) sous la forme d'une matrice ne contenant que des signes + et -.

Ici, la réponse d'une seule grandeur est observée, mais il est possible que plusieurs grandeurs soient observées.

A partir des résultats r_i , certaines notions sont définies:

Notion de moyenne

Par définition, *la réponse moyenne* I d'un système est la valeur prise par la réponse de ce système lorsque les variables sont au centre du domaine (tous les x_i mis à 0). La notion de "0" pour une variable n'étant pas toujours bien définie (cf le sens de rotation évoqué plus haut), la moyenne est obtenue en calculant effectivement la moyenne des réponses du plan, soit, dans le cas présent:

$$I = \frac{r_1 + r_2 + r_3 + r_4}{4}$$

Notion d'effet

L'effet d'un facteur correspond à la quantité dont est modifié ce facteur lorsque la variable centrée réduite augmente d'une unité, les autres facteurs étant bloqués. *L'effet moyen*

correspond à la moyenne des effets d'un facteur. Là encore, la notion "d'augmentation d'une unité" n'est pas toujours simple à définir: c'est pourquoi les valeurs extrêmes sont utilisées pour calculer les effets. Dans le cas présent, l'effet du facteur 1 lorsque le facteur 2 est bloqué au niveau -1 est:

$$E_1^- = \frac{r_2 - r_4}{2}$$

et l'effet du facteur 1 quand le facteur 2 est bloqué au niveau $+1$ est:

$$E_1^+ = \frac{r_1 - r_3}{2}$$

Enfin, l'effet moyen du facteur 1 est:

$$E_1 = \frac{E_1^- + E_1^+}{2}$$

Notion d'interaction

L'*interaction* des autres facteurs sur un facteur donné est définie comme la variation de l'effet de ce facteur quand il augmente d'une unité. Dans le cas présent, l'interaction du facteur 2 avec le facteur 1 est:

$$E_{12} = \frac{E_1^+ - E_1^-}{2}$$

Si le plan compte trois facteurs, il sera possible de définir des interactions:

- ◇ d'ordre 2: entre les facteurs 1 et 2, puis 2 et 3, puis 1 et 3;
- ◇ d'ordre 3: entre les facteurs 1, 2 et 3.

Pour des plans comprenant un nombre plus important de facteurs, des interactions d'ordre supérieur apparaîtront. De manière générale, plus une interaction est d'ordre élevé, plus sa valeur est faible.

MATRICE DES EFFETS

Les différentes notions présentées ci-dessus peuvent être rassemblées sous forme matricielle de la façon suivante (les grandeurs présentées s'appuient sur l'exemple du plan à deux facteurs présenté). En notant:

$\mathcal{R}$ le vecteur des réponses tel que: $\mathcal{R}^t = \{r_1 \ r_2 \ r_3 \ r_4\}$;

$\mathcal{E}$ le vecteur des effets tel que: $\mathcal{E}^t = \{I \ E_1 \ E_2 \ E_{12}\}$;

$\mathbb{M}$ la matrice des effets telle que:

$$\mathbb{M} = \begin{bmatrix} +1 & +1 & +1 & +1 \\ +1 & +1 & -1 & -1 \\ +1 & -1 & +1 & -1 \\ +1 & -1 & -1 & +1 \end{bmatrix}$$

il vient:

$$\mathbb{M}^t \mathcal{R} = 4\mathcal{E}$$

et de manière générale, si le plan à deux niveaux comprend n facteurs, alors, en posant $N = 2^n$:

$$\mathbb{M}^t \mathcal{R} = N\mathcal{E}$$

Dans le cas où plusieurs réponses du système, par exemple q , sont étudiées, $\mathcal{R}$ et $\mathcal{E}$ deviennent des matrices de taille $N \times q$.

Il convient de remarquer que:

- ◇ la première colonne de $\mathbb{M}$, qui correspond à la moyenne, ne contient que des +1;
- ◇ la quatrième colonne de $\mathbb{M}$, qui correspond à l'interaction 1/2, contient des valeurs qui sont le produit des colonnes associées à l'effet 1 et à l'effet 2.

Ces remarques sont généralisables pour des plans comprenant davantage de facteurs: la matrice des effets contient alors une colonne pour la moyenne, une colonne pour chacun des facteurs, une colonne pour chacune des interactions d'ordre 2, puis pour les interactions d'ordre 3, etc.

Plans factoriels fractionnaires

PRINCIPE

Reprenons l'exemple du cas du plan à deux facteurs: on suppose qu'il apparaît, après analyse, qu'un troisième facteur doit être pris en compte pour étudier le système. Une première possibilité consiste à construire le plan factoriel à trois facteurs; une seconde possibilité consiste à exploiter le fait que les interactions d'ordre élevé sont, en général, faibles. Sous cette hypothèse, et en considérant le plan à deux facteurs seulement, l'interaction 1/2 est utilisée comme troisième facteur:

Numéro essai	Facteur 1	Facteur 2	Interaction 1/2 = facteur 3
1	+1	+1	+1
2	+1	-1	-1
3	-1	+1	-1
4	-1	-1	+1

Tableau I.2: Matrice d'expérience (plan à 3 facteurs fractionné)

Dans ce cas, l'effet associé à la colonne du facteur 3 est en fait somme de l'effet du facteur 3 seul et de l'interaction 1/2, ce qui se note:

$$l_3 = E_3 + E_{12}$$

Si l'interaction 1/2 est vraiment négligeable, la valeur de l_3 sera bien très proche de E_3 .

Dans une telle configuration, on dit que E_3 et E_{12} sont *aliasés*, et l_3 est alors appelé *contraste*. Cette démarche permet donc de réduire le nombre d'essais (4 seulement au lieu de 8 dans le cas présent), le prix à payer étant la pureté des résultats (TOUS les effets sont devenus impurs).

Il est bien sûr possible d'aliaser plusieurs interactions et d'ajouter, à un plan complet donné, autant de facteurs qu'il y a d'interactions. Ainsi, dans un plan de base comprenant n facteurs principaux, il est possible d'aliaser jusqu'à $2^n - n - 1$ interactions, le nombre maximum de facteurs qu'il est possible de prendre en compte étant alors $2^n - 1$. Lorsque les $2^n - n - 1$ interactions sont aliasées, le plan est dit *saturé*.

Une étude qui compte n facteurs dont k sont aliasés donne naissance à un plan noté 2^{n-k} . Dans le cas ci-dessus, l'étude comprend 3 facteurs, dont 1 seul est aliasé: nous avons créé un plan 2^{3-1} . Dans ce cas, le plan de base 2^3 a été coupé en deux demi-plans de taille 2^2 . De même, un plan 2^{5-2} est un plan 2^5 coupé en quatre quarts de plans, et un plan 2^{n-k} est un plan 2^n coupé en 2^k hyper-plans.

GÉNÉRATION DES ALIASES

Notation de Box

Il est commode d'introduire la notation de Box: un facteur est désigné par son numéro et une interaction par l'assemblage des numéros des facteurs auxquels elle fait référence. Ainsi, dans l'exemple utilisé, l'alias du facteur 3 avec l'interaction 1/2 s'écrit: $3 = 12$, et se traduit par $l_3 = E_3 + E_{12}$. Cette commodité est en fait bien plus qu'une simple notation, car elle permet de mettre en place une algèbre sur les colonnes de la matrice des effets: $3 = 12$ signifie que la colonne associée au facteur 3 s'obtient en multipliant terme à terme les colonnes des facteurs 1 et 2.

Il en découle une propriété très utilisée:

$$I = 1^2 = 2^2 = (12)^2$$

Groupe des générateurs d'aliasés

Dans l'exemple utilisé, on a:

$$3 = 12$$

soit, après multiplication par 3 de cette égalité:

$$I = 123$$

Cette relation définit le *groupe des générateurs d'aliasés indépendants* (noté GGAI en abrégé): elle permet en effet de retrouver tous les liens entre les différents facteurs et leurs interactions, et donc d'obtenir les expressions des contrastes. Par exemple, l'interaction

$2/3$ s'obtient en multipliant le générateur d'alias par 23, soit:

$$\begin{aligned} I \ 23 &= 1 \ 23 \ 23 \\ 23 &= 1 \ (23)^2 \\ 23 &= 1 \end{aligned}$$

Il est alors simple d'écrire le contraste associé:

$$l_1 = E_1 + E_{23}$$

Ceci montre que dans un plan 2^{3-1} où le seul générateur d'alias est $3 = 12$, les trois contrastes s'écrivent:

$$\begin{aligned} l_1 &= E_1 + E_{23} \\ l_2 &= E_2 + E_{31} \\ l_3 &= E_3 + E_{12} \end{aligned}$$

ce qui confirme bien que l'alias d'une interaction et d'un facteur n'affecte pas seulement ces deux grandeurs, mais tous les facteurs du plan.

Dans le cas où plusieurs alias ont été réalisées, le GGAI comprend plus de deux termes, et il est alors possible d'obtenir le groupe des générateurs d'alias *dépendants*, noté GGAD, par multiplication des générateurs indépendants entre eux, tout d'abord 2 à 2 puis 3 à 3, etc. Le groupe des générateurs d'alias (noté GGA) est la réunion du GGAI et du GGAD.

Supposons par exemple qu'un plan ait été aliasé de la façon suivante:

$$\begin{aligned} 4 &= 12 \\ 5 &= 23 \end{aligned}$$

Il s'agit donc d'un plan 2^{5-2} , dont le GGAI est: $\{I \ 124 \ 235\}$ et le GGAD est $\{1345\}$; le GGA est donc: $\{I \ 124 \ 235 \ 1345\}$. Un contraste s'obtient en multipliant tous les termes du GGA par l'effet recherché et en additionnant le résultat. Par exemple, le contraste associé à l'effet 2 est:

$$l_2 = E_2 + E_{14} + E_{35} + E_{12345}$$

Inversement, la connaissance d'un contraste permet de remonter rapidement au GGA.

Interprétations

La recherche du GGA permet d'écrire les contrastes, dont les valeurs sont ensuite calculées à partir des résultats d'essai: il reste donc à interpréter les valeurs de ces contrastes. Pour cela, un certain nombre d'hypothèses sont retenues:

1. les interactions d'ordre supérieur ou égal à 3 sont supposées négligeables;
2. si un contraste est nul (ou très faible devant les autres contrastes), cela peut signifier:

- ◇ soit que les effets aliasés se compensent (phénomène très rare);
 - ◇ soit que les effets aliasés sont tous négligeables (cas le plus fréquent);
3. si deux effets sont faibles, leur interaction le sera également et pourra être négligée;
 4. si une interaction d'ordre 2 comprend un effet fort et un effet faible, elle peut être négligée;
 5. si deux effets sont forts, leur interaction *peut éventuellement* être importante (mais ce n'est pas toujours le cas).

En appliquant ces règles au contraste l_2 donné dans le paragraphe précédent, il vient:

$$\begin{aligned} l_2 &= E_2 + E_{14} + E_{35} + E_{12345} \\ &\approx E_2 + E_{14} + E_{35} \end{aligned}$$

Si les résultats ont conduit par exemple à montrer que les facteurs 1 et 4 étaient influents, et que les facteurs 3 et 5 étaient peu influents, il en résulte:

$$l_2 \approx E_2 + E_{14}$$

DÉCOUPLAGE

Dans certains cas, les contrastes, après simplifications, comptent encore plusieurs termes (souvent 2) sur lesquels des doutes subsistent: en effet, deux effets influents peuvent très bien avoir une interaction négligeable. Dans l'exemple ci-dessus, bien que 1 et 4 soient influents, leur interaction n'a peut-être aucune incidence sur le résultat.

Pour lever le doute, un plan supplémentaire doit être mis en place: il faut *désaliaser*. Toujours à partir de l'exemple ci-dessus, l'idée consiste à trouver un plan qui fournirait un contraste l'_2 tel que:

$$l'_2 \approx E_2 - E_{14}$$

Il serait alors possible de séparer les contributions de 2 et de 14:

$$\begin{aligned} 2 &= \frac{l_2 + l'_2}{2} \\ 14 &= \frac{l_2 - l'_2}{2} \end{aligned}$$

Un tel plan est généré par $2 = -14$, ou $I = -124$. Comme il s'agit d'un plan 2^{5-2} , il faut fournir un second générateur, donné par exemple par $I = 1235$. Le plan complémentaire sera donc créé en aliasant 4 et -12, puis 5 et 123.

Modèle mathématique associé aux plans factoriels

On introduit une matrice $\mathbb{X}$ qui relie le vecteur $\mathcal{R}$ des réponses et le vecteur $\mathcal{E}$ des effets:

$$\mathcal{R} = \mathbb{X}.\mathcal{E}$$

avec (dans le cas de l'exemple du plan à deux facteurs):

$$\mathbb{X} = \begin{bmatrix} 1 & x_1 & x_2 & x_1x_2 \\ 1 & x_1 & x_2 & x_1x_2 \\ 1 & x_1 & x_2 & x_1x_2 \\ 1 & x_1 & x_2 & x_1x_2 \end{bmatrix}$$

Ce modèle suppose donc une dépendance linéaire par rapport à chacune des variables. Les courbes d'isoréponse (courbes $r_i = C^{te}$) sont donc des droites s'il n'y a pas d'interaction, et des hyperboles s'il y en a.