



Identification et commande sous-optimale d'un four à diffusion

Xabier Ruiz del Portal Anso

► To cite this version:

Xabier Ruiz del Portal Anso. Identification et commande sous-optimale d'un four à diffusion. Automatique / Robotique. Université Paul Sabatier - Toulouse III, 1976. Français. NNT: . tel-00176771

HAL Id: tel-00176771

<https://theses.hal.science/tel-00176771>

Submitted on 4 Oct 2007

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

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

THÈSE

présentée

devant L'UNIVERSITÉ PAUL SABATIER DE TOULOUSE (SCIENCES)

en vue de l'obtention

du Grade de DOCTEUR de Spécialité E.E.A. - Option "Automatique"

par

Xabier RUIZ DEL PORTAL ANSO

Maitre ès Sciences

IDENTIFICATION ET COMMANDE SOUS-OPTIMALE D'UN FOUR A DIFFUSION

Soutenue le 1^{er} Juillet 1976 devant la Commission d'Examen :

MM. Y. SEVELY

Président

J.-L. ABATUT

V. ALEIXANDRE

J.-P. BABARY

}
Examineurs

Astarketa honek biltzen du urte luzeetan egindako lana gure Herritik at, elburu bezela jarrita teknika ta jakintza hauek askatasun giro egoki batean eman dezaiola gure Aberriari aurrerakoi indar bat gizalari aberastasuna lortzeko eta bere nortasuna beste ainbat Herrirekin bereizteko

E R R A T A

- P. 15 ligne 7 lire : où q est le flux de chaleur
 ligne 12 lire : q_i est composé
- P. 18 ligne 16 : remplacer K_g par K_f
- P. 33 formule (2.3) lire : $L_n^* (x_K^*) \triangleq L_n^*$
 formule (2.5) lire : $\max \{ x_2, 1 - x_1 \}$
- P. 34 6ème ligne avant la fin, lire : Nous écrirons : $L_{n-1}^* = 2 L_n^* - \epsilon$
- P. 35 formule (2.9) lire : $L_{n-2}^* = L_{n-1}^* + L_n^*$
 3ème ligne avant la fin, lire $F_0 = F_1 = 1$
 formules (2.12) et (2.13) lire $K = 1, 2, \dots, n-1$
- P. 46 dernière ligne lire : $(W_j^{n+1})^4 =$
- P. 51 3ème ligne lire : où $\underline{z}(t)$ est un vecteur
 4ème ligne lire : $z_i(t) = T_{w_i}(t) - T_{m_i}(t)$
- P. 71 9ème ligne lire : où S et Q
- P. 76 13ème ligne lire : alors les équations (3.37) s'écrivent
 formule (3.39) remplacer $\underline{K}_2$ par $\underline{k}_2$
- P. 78 et P. 80 remplacer K_2 et K_3 par $\underline{k}_2$ et $\underline{k}_3$

AVANT - P R O P O S

Les marques de sympathie et les encouragements qui m'ont accueilli à mon arrivée à l'Université Paul Sabatier sont à l'origine de la réalisation de ce travail. Que Messieurs GIRALT, LACOSTE, LAGASSE et MASCART trouvent ici l'expression de ma gratitude la plus sincère.

Je suis profondément reconnaissant à Monsieur MARTINOT de m'avoir permis d'effectuer mes recherches au Laboratoire d'Automatique et d'Analyse des Systèmes.

Je tiens à exprimer toute ma reconnaissance à Monsieur le Professeur SEVELY pour l'intérêt qu'il a toujours manifesté à mon égard et l'honneur qu'il m'a fait en présidant le jury de cette thèse.

Je remercie vivement Monsieur BABARY de m'avoir accepté dans l'équipe "Systèmes à Paramètres Répartis-Unités Pilotes" qu'il dirige et pour les conseils dont il m'a entouré.

Que Monsieur ABATUT Maître de Conférences à l'Université Paul Sabatier et Monsieur ALEIXANDRE Professeur à l'Université de Valladolid qui ont bien voulu siéger à cette commission d'examen en soient sincèrement remerciés.

Je suis gré à Messieurs AMOUROUX, DOUCET, QUEMENER et ZABALA de l'aide qu'ils m'ont apporté.

Qu'il me soit permis de témoigner toute ma sympathie aux techniciens, secrétaires du laboratoire et aux membres de l'équipe "Systèmes à Paramètres Répartis-Unités Pilotes" dont le soutien m'a été particulièrement précieux.

Enfin, je ne saurais oublier le personnel de l'imprimerie et de la Société Monique CASTA sans qui ce mémoire n'aurait pas existé.

SYMBOLES UTILISES

X	: coordonnée axiale du tube de quartz
T_w	: température de la surface interne du tube de quartz
T_g	: température moyenne du gaz sur une section
$q(x,t)$	: flux de chaleur traversant la paroi interne du tube de quartz
$q_g(x,t)$	: flux de chaleur arrivant sur le matériau fictif
$q_s(x,t)$	: flux de chaleur latéral sortant de la section de tube de longueur dX
$q_e(x,t)$	: flux de chaleur latéral entrant dans cette section
q_i	: rayonnement total incident par unité de surface
q_o	: rayonnement total émergent par unité de surface
$F(X)$	: facteur de configuration géométrique relatif au rayonnement de la paroi à travers les orifices du tube
$K(Z)$	: facteur de configuration géométrique entre deux éléments de la surface interne du tube
ρ	: densité du gaz
σ	: constante de Stefan-Boltzman
u_m	: vitesse du gaz
h	: coefficient de transfert par convection
ϵ	: coefficient d'émissivité de la surface
C_p	: chaleur spécifique du gaz
ρ_f	: masse volumique du matériau fictif
C_{p_f}	: chaleur spécifique du matériau fictif
S_e	: surface extérieure du tube de matériau fictif
S_i	: surface intérieure du tube de quartz

V	: volume du tube de matériau fictif
L	: longueur du tube de quartz
D	: diamètre intérieur du tube de quartz
e	: épaisseur du tube de matériau fictif
β	: coefficient qui tient compte des pertes calorifiques et du changement d'unités
W	: température de la paroi (à une constante près)
G	: température du gaz (à une constante près).

VALEURS NUMERIQUES

$$\sigma = 4,96 \cdot 10^{-8} \text{ Kcal/h} \cdot \text{m}^2 \cdot ^\circ \text{K}^4$$

$$h = 7 \text{ Kcal/m}^2 \cdot \text{h} \cdot ^\circ \text{K}$$

$$\varepsilon = 0,41$$

$$u_m = 36 \text{ m/h}$$

$$\rho = 0,277 \text{ kg/m}^3$$

$$C_p = 0,28 \text{ kcal/kg} \cdot ^\circ \text{K}$$

Remarque: Nous conserverons dans un souci d'homogénéisation avec les précédentes études réalisées au L.A.A.S. les unités M.K.S.A.- De plus, le passage aux unités conventionnelles actuelles nécessiterait la modification de tous les coefficients intervenant au niveau de la modélisation.

I N T R O D U C T I O N

L'étude et la réalisation de structures de commandes numériques pour des processus de type industriel constitue à l'heure actuelle un thème de recherche très intéressant car il fait intervenir de nombreux domaines de l'automatique tels que :

- Modélisation de systèmes complexes
- Identification paramétrique
- Commande optimale et commande en temps réel ...

Un des thèmes principaux de l'équipe "Systèmes à Paramètres Répartis - Unités Pilotes " du L.A.A.S est constitué par la commande numérique d'un four de fabrication de composants électroniques.

Le problème peut être posé dans sa généralité de la façon suivante :

Déterminer une loi de commande et l'implanter sur le processus de façon à ce qu'un critère de coût soit optimisé.

Pour faire une telle étude, il faut disposer d'un modèle mathématique du processus. Un premier modèle a été établi en 1972 par M.AMOUROUX, J.P.BABARY et F.ROUBELLAT [I.4], [I.5]. Il est constitué de deux équations aux dérivées partielles non linéaires. Ces équations ont été étudiées à partir d'une linéarisation autour d'un profil de fonctionnement et de l'application d'une méthode d'approximation (Galerkin).

Ce sont ces équations qui ont été utilisées par A.BERGEAUD et L.SARMIENTO pour faire une étude de régulation autour du profil précédent. Le modèle linéaire se prêtait assez bien à l'identification de certains coefficients mal connus. Cette identification a été effectuée hors ligne en utilisant la technique du filtrage non linéaire continu discret [I.6]

La détermination de la loi de commande à partir d'un critère quadratique sur l'état et la commande, et d'un observateur permettant de reconstituer certaines des variables d'état [I.7], ont terminé cette première phase d'étude à caractère théorique.

Parallèlement à cela, toute la partie électronique nécessaire à l'acquisition de mesure et au couplage du calculateur (Cii.10010) avec le processus a été réalisée. Ceci a permis une vérification des résultats précédents.

Le travail que nous présentons dans ce mémoire consiste en une généralisation des résultats précédents. Dans certaines applications, l'état du système varie dans de grandes proportions. Ceci se vérifie, par exemple, dans le cas d'un problème de poursuite, d'atteinte d'un certain état désiré ou d'évolutions temporelles notables. Dans tous les cas, il est intéressant d'avoir un modèle mathématique global le plus représentatif possible du fonctionnement du processus.

C'est pour cela que nous avons conservé le caractère non linéaire du modèle initial.

Notre travail a consisté alors en une identification paramétrique de coefficients et en la détermination d'une loi de commande optimale. Le critère que nous avons choisi comporte en plus des termes quadratiques classiques sur l'écart entre l'état désiré et l'état courant, une fonction non linéaire de cet écart. Ceci a été fait essentiellement dans un but de simplification de l'étude et nous a conduit à une structure de commande sous-optimale, mais de calcul et de réalisation relativement simples.

Dans ce mémoire, nous avons adopté la présentation suivante :

- dans le premier chapitre, nous avons complété la modélisation du fonctionnement du four à diffusion. Pour cela,

nous l'avons considéré comme étant constitué d'un matériau homogène entre le tube de quartz et l'enroulement chauffant, ses paramètres physiques étant constants et mal connus. D'autre part, nous avons introduit un coefficient de perte au niveau des enroulements de chauffe.

Un bilan thermodynamique nous a conduit à un modèle global sous la forme d'un système d'équations aux dérivées partielles non linéaires, qui permet de connaître la température en chaque point du four et à chaque instant en fonction de la répartition de la puissance de chauffe le long du tube. Une discrétisation spatiale du modèle a été effectuée pour pouvoir appliquer des techniques d'optimisation plus connues.

- au deuxième chapitre, nous présentons l'algorithme d'identification et les résultats obtenus. Certaines définitions étant rappelées, le principe de la méthode de Fibonacci est alors présenté : c'est un algorithme de recherche unidimensionnel qu'il est intéressant d'utiliser au niveau de la méthode d'identification de Sphéra. Cette identification effectuée à partir de mesures de températures dans le processus et de la connaissance de la répartition des puissances calorifiques appliquées, nous permet de déterminer les paramètres inconnus du modèle.

- le dernier chapitre est consacré à la recherche d'une loi de commande et aux simulations correspondantes. Le problème est de déterminer la meilleure commande possible, permettant de minimiser un certain critère d'optimisation.

Pour le cas du problème de l'atteinte d'un état désiré, dans un temps donné, à partir d'un état initial quelconque, la forme du critère est celle proposée auparavant.

La résolution du problème nous a conduit à une loi de

commande sous-optimale en boucle fermée. Les résultats obtenus sont présentés à la fin du chapitre.

C H A P I T R E I

Introduction

I-1 - Description du processus

I-2 - Modélisation

I-2.1 - Modélisation de la partie interne du four

I-2.2 - Etablissement d'un système aux dérivées partielles

I-2.3 - Modèle global

I-2.4 - Modèle à paramètres localisés

Conclusion

Bibliographie

INTRODUCTION

La recherche d'une structure de modèle mathématique du fonctionnement d'un processus (par exemple, le four à diffusion qui fait l'objet de notre étude) doit toujours être fonction de son utilisation ultérieure. Cette structure est différente suivant que le but fixé est une simulation du fonctionnement du processus ou que l'objectif à atteindre est l'établissement d'un algorithme de commande ou encore celui d'une simple régulation autour d'un point de fonctionnement désiré. L'aspect "hardware" peut aussi intervenir selon que la boucle doit être analogique ou numérique.

Suivant le cas, le modèle obtenu peut être très complexe (cas d'une simulation fine) ou relativement simple (cas d'une régulation classique).

En ce qui nous concerne, notre problème est de déterminer une loi de commande du four à diffusion modélisé sous la forme d'un système non linéaire. Par conséquent, le modèle de son comportement global devrait être le plus précis possible. Mais dans ces conditions, la détermination de la loi de commande devient vite très complexe. Aussi avons-nous tenté d'obtenir un compromis entre une bonne précision du modèle et la relative simplicité de son utilisation.

Dans ce chapitre nous traitons le problème de la modélisation du four à diffusion de la manière suivante : après avoir décrit le processus utilisé, nous effectuons un bilan thermodynamique des transferts de chaleur en considérant les phénomènes de rayonnement, conduction dans les solides et convection dans le gaz. Ceci nous conduit à un système d'équations aux dérivées partielles non linéaires. Pour pouvoir identifier certains paramètres inconnus de ce système, une discrétisation spatiale est réalisée aux points de mesure des températures, le modèle ainsi obtenu est un système

différentiel ordinaire du douzième ordre. C'est à partir de ce modèle définitif que sera calculé l'algorithme de commande.

I -1. DESCRIPTION DU PROCESSUS

Le four qui fait l'objet de notre étude est un four de marque Electrogas qui est utilisé dans certaines usines de fabrication de composants électroniques.

Il est formé de plusieurs tubes concentriques et horizontaux, autour desquels sont bobinés les enroulements chauffants. L'ensemble est renfermé dans un caisson bourré de matériau isolant. Figure 1.1.

Le tube central est en quartz et débouche des deux côtés dans l'air ambiant. Il est traversé par un mélange gazeux nécessaire au dopage des composants situés au coeur du four.

Les trois éléments de chauffe sont bobinés sur les rainures d'un support réfractaire. Leur longueur totale est légèrement inférieure à la longueur totale du four.

La zone centrale de chauffe, de longueur deux fois et demi celle d'une zone latérale, correspond approximativement à la partie du four dans laquelle on veut obtenir un palier de température. Un tube de mulite placé entre le tube de quartz et l'enroulement chauffant;

- . crée un écran aux particules métalliques émises par l'enroulement chauffant, qui pourraient diffuser à travers le tube de quartz,
- . évite les gradients thermiques trop importants,
- . homogénéise la température.

La température du four est mesurée en six points, au moyen de six thermocouples en Nickel-Chrome-Nickel-Allié, placés sur la paroi extérieure du tube de quartz. Figure 1.2.

Principales grandeurs du four :

	$\varnothing$ intérieur	$\varnothing$ extérieur
- tube de quartz	51 mm	55 mm
- tube de mulite	70 mm	77 mm
- enveloppe réfractaire	100 mm	150 mm
- Longueur de four :	1365 mm	
- Longueur de la zone chauffée :	1170 mm	
- Longueur de la zone centrale :	650 mm	
- Longueur d'une zone latérale :	260 mm	
- Longueur d'une zone morte :	97,5 mm	

1 -2. MODELISATION

Cette étude peut être divisée en deux parties :

- La modélisation de la partie interne du four avec obtention d'un système aux dérivées partielles, qui a déjà fait l'objet de travaux antérieurs [1.4], [1.5], [1.6], [1.7].
- L'obtention du modèle global non linéaire qui nous est propre et augmente le champ des applications possibles.

1 -2.1 - Modélisation de la partie interne du four

Nous allons considérer comme seuls échanges d'énergie au coeur du four le rayonnement de la paroi du tube et la

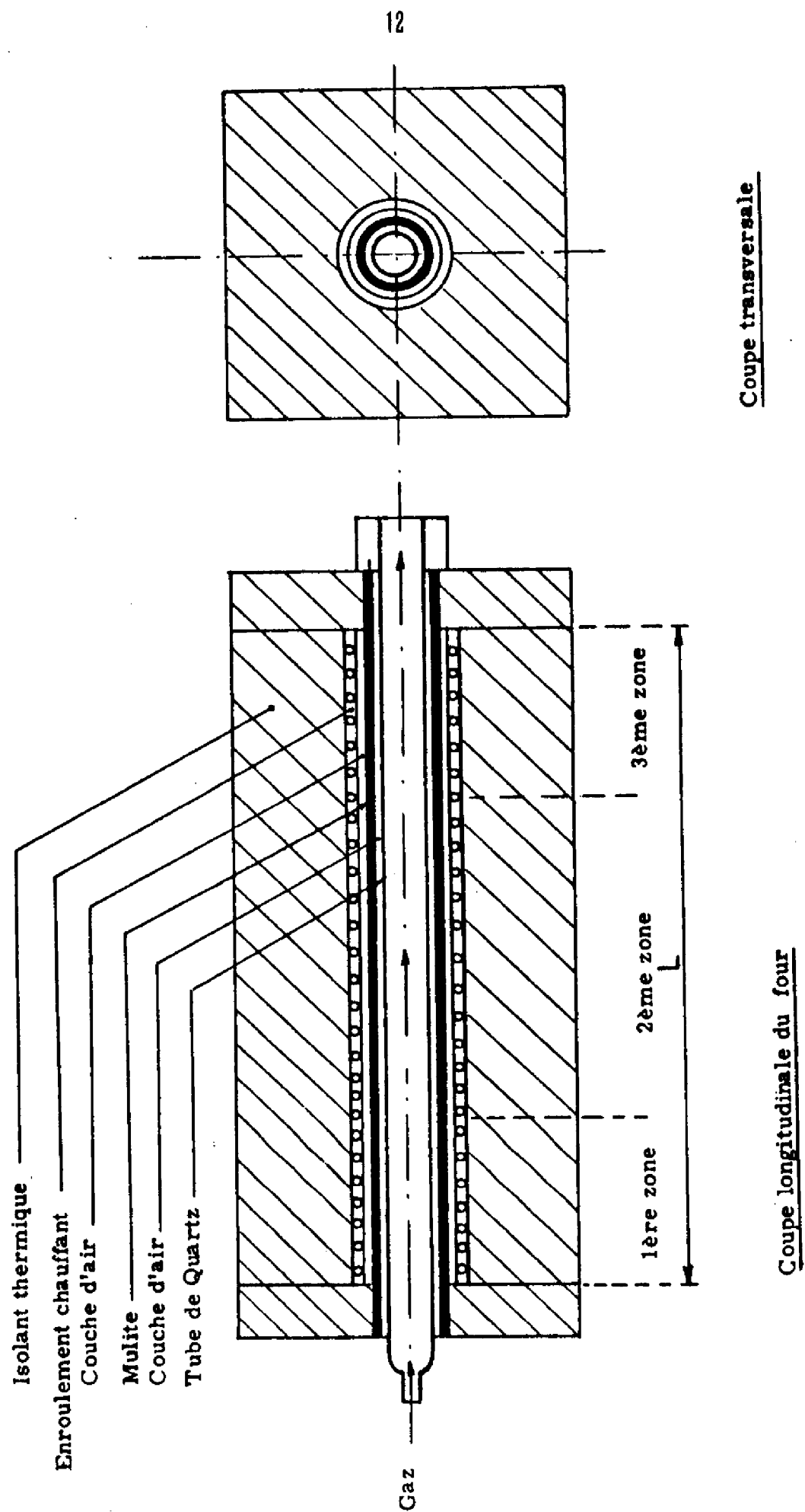


Figure 1.1

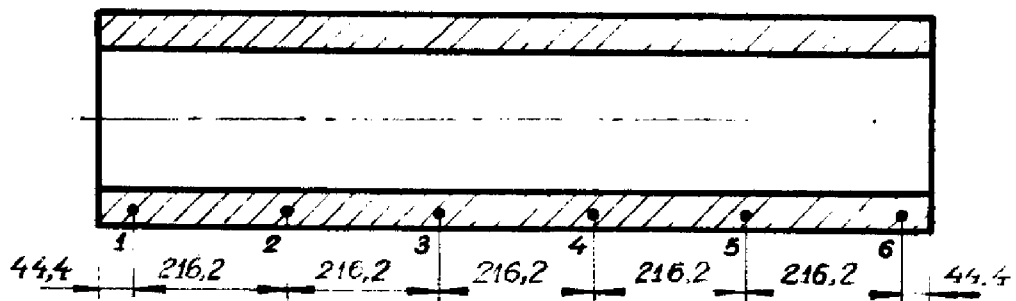


Figure 1.2

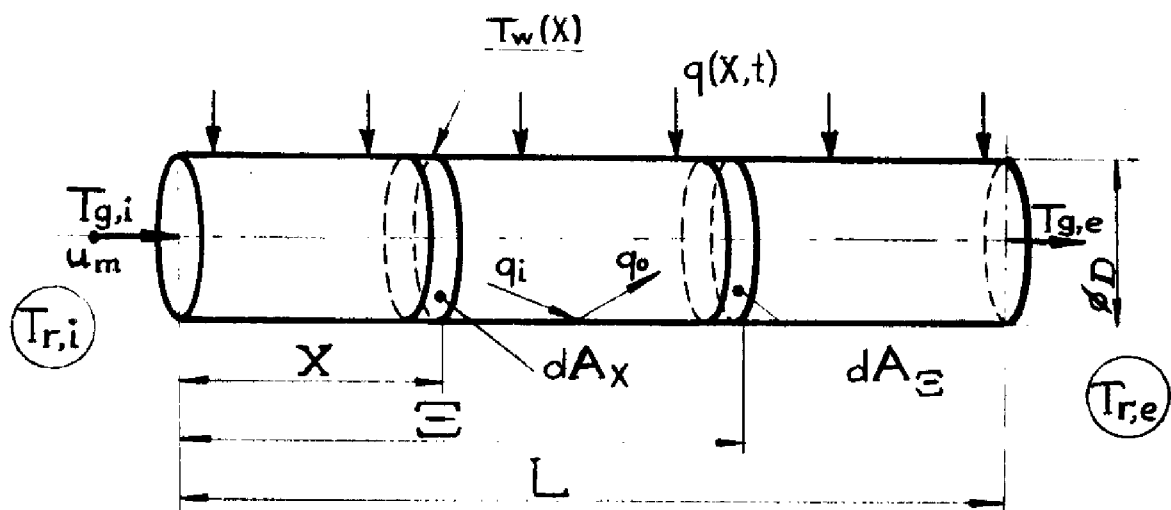


Figure 1.3

convection forcée due au débit de gaz.

Hypothèses

- On impose à la paroi interne tube un flux de chaleur à l'aide d'un dispositif extérieur (chauffage électrique)
- Le gaz qui s'écoule dans le tube laisse passer les rayonnement et ne modifie pas, par suite, les échanges par rayonnement entre les éléments de la paroi interne de la conduite.
- On suppose que l'intérieur de la conduite est une surface grise diffuse (les facteurs d'émission et d'absorption monochromatique sont pratiquement identiques pour toutes les longueurs d'ondes).
- La conduction dans le gaz est négligée. La vitesse du gaz u_m étant relativement faible, le régime est laminaire.
- Le coefficient de convection h entre la paroi et le gaz est supposé constant (indépendant de la température). Les études ultérieures de sensibilité ont montré a posteriori que la solution du modèle est très peu sensible par rapport aux coefficients u_m et h .

Le système que nous allons analyser est schématisé sur la figure 1.3. Un gaz transparent au rayonnement arrive à l'entrée du tube avec une température $T_{g,i}$, le traverse et ressort à la température $T_{g,e}$. Un flux de chaleur $q(x,t)$ est appliqué à la paroi du tube. Chaque extrémité du tube est placée dans un milieu ambiant de température $T_{r,i}$ à l'entrée et $T_{r,e}$ à la sortie.

Le bilan thermodynamique qui tient compte des échanges par rayonnement est obtenu à partir de la méthode de Poljak décrite en [1.2.7].

On désignera par q_0 l'émittance sortant de la surface

de la paroi, composée de son émission directe et de la réflexion du rayonnement incident sur cette surface. L'intensité du rayonnement incident, issu d'un élément de la surface est désignée par q_i . En considérant un élément de surface de la paroi du tube on peut écrire [1.1] :

$$q + q_i = q_0 + h [T_w(X,t) - T_g(X,t)] \quad (1.1)$$

où q est le flux de chaleur imposé et $T_w - T_g$ est la différence locale entre la température de la paroi et la température du volume gazeux. Or

$$q_0 = \epsilon \sigma T_w^4 + (1 - \epsilon) q_i \quad (1.2)$$

Le rayonnement total incident par unité de surface est composé de deux termes différents : Le rayonnement venant du milieu ambiant aux extrémités du tube plus le rayonnement venant du reste de la paroi.

$$q_i = \sigma T_{r,i}^4 F(X') + \sigma T_{r,e}^4 F(L' - X') + \int_0^{X'} q_0(\xi) K(X - \xi) d\xi + \int_{X'}^{L'} q_0(\xi) K(\xi - X') d\xi \quad (1.3)$$

$$\text{avec } X' = \frac{X}{D} \quad L' = \frac{L}{D} \quad \xi = \frac{\theta}{D}$$

La fonction $F(X')$ est le facteur de configuration géométrique pour le rayonnement d'un élément de la paroi du tube au point X' vers l'ouverture circulaire de l'entrée de ce tube. Ce facteur est donné par [1.3] :

$$F(X') = \frac{X'^2 + 0,5}{(X'^2 + 1)^{3/2}} - X' \quad X' \geq 0 \quad (1.4)$$

La fonction $K(z)$ est le facteur de configuration géométrique entre deux anneaux de la paroi distants de z [1.3]

$$K(z) = 1 - \frac{z^3 + 1,5z}{(z^2 + 1)^{3/2}} \quad z \geq 0 \quad (1.5)$$

Dans les articles cités en référence [1.10] et [1.11] les auteurs ont proposé une approximation des expressions (1.4) et (1.5) :

$$F(x') \cong \frac{1}{2} e^{-2x'} \quad K(z) \cong e^{-2z}$$

Par substitution de (1.1) et (1.2) en (1.3) on élimine q_1

$$q = \sigma \epsilon T_w^4 - \epsilon \sigma T_{r,i}^4 \left(\frac{1}{2} e^{-2x'} \right) - \sigma \epsilon T_{r,e}^4 \left(\frac{1}{2} e^{-2(L'-x')} \right) - \epsilon \int_0^{x'} q_0(\xi) e^{-2(x'-\xi)} d\xi - \epsilon \int_0^{L'-x'} q_0(\xi) e^{-2(\xi-x')} d\xi + h [T_w(x,t) - T_g(x,t)] \quad (1.6)$$

Pour déterminer T_g et T_w une équation complémentaire est établie au niveau du volume gazeux. En négligeant la conduction dans le gaz par rapport à la convection et puisque le gaz a été défini transparent au rayonnement, la seule quantité de chaleur transférée l'est par convection de la paroi. Pour l'élément de volume cylindrique de longueur dX et de diamètre D , la quantité d'énergie transférée de la paroi au gaz s'écrit :

$$h \pi D (dX) [T_w(X,t) - T_g(X,t)] dt$$

Celle-ci est égale à la quantité d'énergie emportée par le gaz augmentée de l'accroissement d'énergie interne du gaz :

$$\rho c_p \frac{\pi D^2}{4} u_m [T_g(X+dX,t) - T_g(X,t)] + \rho c_p \frac{\pi D^2}{4} dX [T_g(X,t+dt) - T_g(X,t)]$$

Ces deux quantités sont égales et le résultat est mis sous la forme:

$$4h [T_w(X,t) - T_g(X,t)] = \rho c_p D \left[\frac{\partial T_g(X,t)}{\partial t} + u_m \frac{\partial T_g(X,t)}{\partial X} \right] \quad (1.7)$$

I.2.2.-Etablissement d'un système aux dérivées partielles.

Les intégrales de l'expression (1.6) peuvent être éliminées si on dérive deux fois celle-ci par rapport à X' et si on la soustrait quatre fois au résultat de la dérivation.

On obtient :

$$\frac{\partial^2 q}{\partial X'^2} - 4q = \varepsilon \sigma \frac{\partial^2 T_w^4}{\partial X'^2} - 4\sigma \varepsilon T_w^4 + 4\varepsilon q_0(X') - 4h(T_w - T_g) + h \frac{\partial^2 (T_w - T_g)}{\partial X'^2} \quad (1.8)$$

Les équations (1.1) et (1.2) peuvent être combinées afin d'éliminer q_i

$$\varepsilon q_0 = \varepsilon \sigma T_w^4 - (1 - \varepsilon)q + h(1 - \varepsilon)(T_w - T_g) \quad (1.9)$$

En reportant (1.9) dans (1.8) on obtient :

$$\frac{\partial^2}{\partial X'^2} (\varepsilon \sigma T_w^4 + h(T_w - T_g)) = 4\varepsilon h(T_w - T_g) - 4\varepsilon q + \frac{\partial^2 q}{\partial X'^2} \quad (1.10)$$

Faisons le changement de variables

$$\begin{cases} X = Lx \\ X' = \frac{L}{D}x \end{cases} \quad 0 \leq x \leq 1$$

Nous pouvons alors exprimer les relations (1.7) et (1.10) en fonction de la variable réduite x .

$$\begin{cases} \frac{\partial^2}{\partial x^2} (\varepsilon \sigma T_w^4 + h(T_w - T_g)) = 4\varepsilon h \frac{L^2}{D^2} (T_w - T_g) - 4\varepsilon \frac{L^2}{D^2} q + \frac{\partial^2 q}{\partial x^2} \\ \frac{\partial T_g}{\partial t} + \frac{u_m}{L} \frac{\partial T_g}{\partial x} = \frac{4h}{\rho c_p D} (T_w - T_g) \end{cases} \quad (1.11)$$

On obtient ainsi un système de deux équations différentielles aux dérivées partielles, non linéaires, représentant l'évolution des températures dans la partie interne du four.

1.2.3. - Modèle global

Il existe entre le tube de quartz et l'enroulement de chauffe deux couches d'air séparées par le tube de mulite. Une modélisation fine de l'ensemble serait difficile et le modèle trop complexe.

Dans le but de surmonter cette difficulté, nous allons

énoncer deux hypothèses :

- Tous les matériaux existant entre la paroi interne du tube de quartz et l'enroulement chauffant sont assimilés à un seul corps, homogène, de masse spécifique ρ_g et de chaleur spécifique C_{pg} .
- Etant donné que toute la puissance électrique n'ira pas chauffer le tube de quartz, on définira un coefficient de pertes β , constant par rapport à l'espace et aux températures de fonctionnement.

On considère une section de tube de longueur dx .

Figure I.4

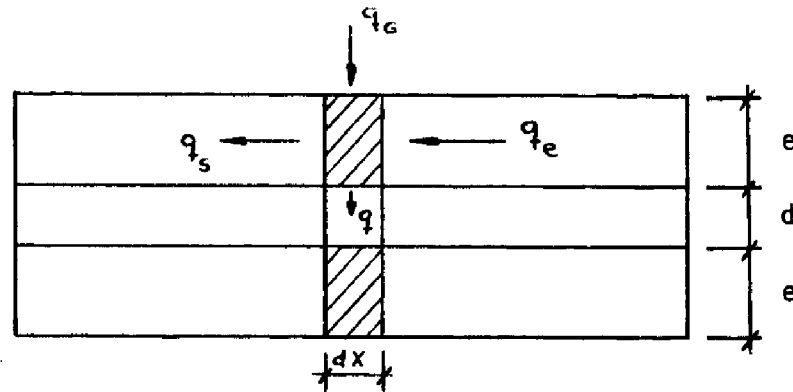


Figure I.4

On introduit:

- $q_g(x,t)$ flux de chaleur arrivant sur la face externe
- $q_s(x,t)$ flux de chaleur latéral sortant de la section de tube de longueur dx
- $q_e(x,t)$ flux de la chaleur latéral entrant dans cette section
- K_g conductibilité thermique du matériel
- S_e surface extérieure du tube
- S_i surface intérieure du tube
- V volume du tube

- Bilan énergétique :

$$\text{Variation d'énergie interne} = q_e - q_s - \text{conduction}$$

soit :

$$V \rho_f c_{Pf} \frac{\partial T_w}{\partial t} = S_e (\beta q_a(x,t) - S_i q(x,t) + K_f V \frac{\partial^2 T_w}{\partial x^2} \quad (1.12)$$

Sur la zone centrale du four, le gradient de température par rapport à x est faible. On suppose alors que la conduction latérale est négligeable sur cette partie où la régulation sera la meilleure.

En explicitant les expressions de V (volume du matériau, S_e (surface de la paroi externe) et S_i (surface de la paroi interne) nous avons :

$$e(e+D) \rho_f c_{Pf} \frac{\partial T_w}{\partial t} = (D+2e) \beta q_a(x,t) - D q(x,t) \quad (1.13)$$

Le modèle se compose alors des équations (1.11) et de (1.13). Au moins sur la partie centrale du four, le terme $\frac{\partial^2 q}{\partial x^2}$ (donc la courbure du profil du flux de chaleur) est négligeable devant $\frac{4\varepsilon}{D^2} q$. Par élimination du terme $q(x,t)$ nous obtenons le modèle global du four :

$$\left\{ \begin{array}{l} \frac{\partial^2}{\partial x^2} \left[\varepsilon \sigma T_w^4 + h(T_w - T_g) \right] = 4\varepsilon h \frac{L^2}{D^2} (T_w - T_g) - \frac{4\varepsilon L^2}{D^3} (D+2e) \beta q_a(x,t) + \\ \quad + \frac{4\varepsilon L^2}{D^3} e(D+e) \rho_f c_{Pf} \frac{\partial T_w}{\partial t} \\ \frac{\partial T_g}{\partial t} + \frac{u_m}{L} \frac{\partial T_g}{\partial x} = \frac{4h}{\rho c_p D} (T_w - T_g) \end{array} \right. \quad (1.14)$$

Conditions aux limites :

On considère les températures du gaz et de la paroi égales et constantes au point d'abscisse zéro.

$$\left. \begin{array}{l} T_w(0,t) = T_{w0} \\ T_g(0,t) = T_{g0} \end{array} \right\} T_{w0} = T_{g0}$$

En absence d'un thermocouple au point d'abscisse un, les températures correspondantes sont mal connues et les conditions aux limites difficiles à déterminer.

Néanmoins, la température de sortie ne va pas varier que de quelques degrés seulement et les travaux de simulation de Sarmiento [1.6] et Bergeaud [1.7] montrent qu'on commet une erreur relativement faible si on considère :

$$T_w(1,t) = T_{wo} = T_g(1,t)$$

Conditions initiales :

Elles sont choisies pour chaque manipulation :

$$T_w(x,0) = T_{wc}(x)$$

$$T_g(x,0) = T_{gc}(x)$$

Faisons le changement de variables :

$$T_w(x,t) = W(x,t) + T_{wo}$$

$$T_g(x,t) = G(x,t) + T_{go}$$

Le système d'équations (1.14) devient en définitive :

$$\left(\begin{aligned} \frac{\partial^2}{\partial x^2} \left[\epsilon \sigma W^4 + h(W-G) \right] &= 4\epsilon h \frac{L^2}{D^2} (W-G) - \frac{4\epsilon L^2}{D^3} (D+2e) \beta q_g(x,t) + \\ &+ \frac{4\epsilon L^2}{D^3} e(e+D) \rho_f c_{pf} \frac{\partial W}{\partial t} \\ \frac{\partial G}{\partial t} + \frac{u_m}{L} \frac{\partial G}{\partial x} &= \frac{4h}{\rho c_p D} (W-G) \end{aligned} \right. \quad (1.15)$$

avec les conditions aux limites :

$$W(0,t) = 0$$

$$G(0,t) = 0$$

$$W(1,t) = 0$$

$$G(1,t) = 0$$

et les conditions initiales :

$$W(x,0) = T_{wc}(x) - T_{wo}$$

$$G(x,0) = T_{gc}(x) - T_{go}$$

1.2.4 - Modèle à paramètres localisés :

Le modèle dynamique que nous venons d'obtenir, est une bonne approximation des phénomènes thermodynamiques qui se produisent dans le four. Néanmoins, son utilisation est assez gênante parce qu'il contient à la fois des dérivées partielles et un terme non linéaire important.

Une méthode classique de résolution d'un tel système consiste à discrétiser par rapport à la variable d'espace. Le résultat est un ensemble d'équations différentielles ordinaires, lesquelles peuvent être traitées au moyen des techniques plus connues (Principe du maximum entre autres).

Etant donné que la température de la paroi du tube est mesurée en six points au moyen de thermocouples placés à cet effet, nous garderons ces points comme points de discrétisation.

Nous appellerons $\Delta x_j = x_j - x_{j-1}$ la distance séparant deux points d'intégration. Les approximations des dérivés étant :

$$\left(\frac{\partial^2 W}{\partial x^2} \right)_j = \frac{Z}{\Delta x_j + \Delta x_{j+1}} \left(\frac{W_{j+1} - W_j}{\Delta x_{j+1}} - \frac{W_j - W_{j-1}}{\Delta x_j} \right)$$

Pour l'expression de $\left(\frac{\partial^2 G}{\partial x^2} \right)_j$ on obtient la même équation où les W sont remplacés par les G.

$$\left(\frac{\partial^2 W^4}{\partial x^2} \right)_j = \frac{Z}{\Delta x_j + \Delta x_{j+1}} \left(\frac{W_{j+1}^4 - W_j^4}{\Delta x_{j+1}} - \frac{W_j^4 - W_{j-1}^4}{\Delta x_j} \right)$$

$$\left(\frac{\partial G}{\partial x} \right)_j = \frac{G_{j+1} - G_{j-1}}{\Delta x_j + \Delta x_{j+1}}$$

Les équations (1.15) s'écrivent maintenant :

$$\left\{ \begin{aligned} \frac{dW_j}{dt} &= \frac{2A}{\Delta x_j + \Delta x_{j+1}} \left[\frac{W_{j+1}^4 - W_j^4}{\Delta x_{j+1}} - \frac{W_j^4 - W_{j-1}^4}{\Delta x_j} \right] + \frac{2B}{\Delta x_j + \Delta x_{j+1}} \left[\frac{W_{j+1} - W_j}{\Delta x_{j+1}} - \frac{W_j - W_{j-1}}{\Delta x_j} \right] \\ &\quad - \frac{2B}{\Delta x_j + \Delta x_{j+1}} \left[\frac{G_{j+1} - G_j}{\Delta x_{j+1}} - \frac{G_j - G_{j-1}}{\Delta x_j} \right] - c(W_j - G) + E q_G(x_j, t) \\ \frac{dG_j}{dt} &= M(W_j - G) - F \frac{G_{j+1} - G_{j-1}}{\Delta x_j + \Delta x_{j+1}} \end{aligned} \right. \quad (1.16)$$

$$j = 1, 2, \dots, 6$$

où

$$A = \frac{\sigma D^3}{4L^2 e(e+D) \rho_f c_{pf}}$$

$$B = \frac{h D^3}{4\epsilon L^2 e(e+D) \rho_f c_{pf}}$$

$$C = \frac{D h}{e(e+D) \rho_f c_{pf}}$$

$$M = \frac{4 h}{\rho c_p D}$$

$$F = \frac{U_m}{L}$$

$$E = \frac{(D + 2e) \beta}{e(e+D) \rho_f c_{pf}}$$

soit en regroupant les W_j et les G_j :

$$\left\{ \begin{aligned} \frac{dW}{dt} &= \frac{2B}{(\Delta x_j + \Delta x_{j+1}) \Delta x_j} W_{j-1} - \left(\frac{2B}{\Delta x_{j+1} \Delta x_j} + c \right) W_j + \frac{2B}{(\Delta x_j + \Delta x_{j+1}) \Delta x_{j+1}} W_{j+1} \\ &\quad - \frac{2B}{(\Delta x_j + \Delta x_{j+1}) \Delta x_j} G_{j-1} + \left(\frac{2B}{\Delta x_{j+1} \Delta x_j} + c \right) G_j + \frac{2B}{(\Delta x_j + \Delta x_{j+1}) \Delta x_{j+1}} G_{j+1} + \\ &\quad + \frac{2A}{(\Delta x_j + \Delta x_{j+1}) \Delta x_j} W_{j-1}^4 - \frac{2A}{\Delta x_{j+1} \cdot \Delta x_j} W_j^4 + \frac{2A}{(\Delta x_j + \Delta x_{j+1}) \Delta x_{j+1}} W_{j+1}^4 + \\ &\quad + E \cdot q_c(x_j, t) \\ \frac{dG_j}{dt} &= MW_j + \frac{F}{\Delta x_j + \Delta x_{j+1}} G_{j-1} - DG_j - \frac{F}{\Delta x_j + \Delta x_{j+1}} G_{j+1} \\ j &= 1, 2, \dots, 6 \end{aligned} \right. \quad (1.17)$$

CONCLUSION

La description du fonctionnement d'un four à diffusion et l'établissement d'un bilan thermodynamique au niveau des transferts de chaleur nous ont permis de modéliser le comportement thermique de ce four sous la forme d'un système d'équations aux dérivées partielles non linéaires.

Ce modèle fournit la distribution à chaque instant des températures du gaz et de paroi, laquelle est traversée par un flux de chaleur, en fonction de celui-ci et par là-même de

la répartition des puissances des trois zones de chauffage.

Pour des besoins d'identification paramétrique et d'optimisation, ce modèle a été transformé par une méthode de discretisation en un système différentiel ordinaire du douzième ordre non linéaire.

B I B L I O G R A P H I E

- [1.1] SIEGEL, M. PERLMUTTER
Convective and radiant heat transfert for flow of a transparent gas in a tube with a gray wall. Int. J. Heat Mass. Transfert - Vol.5, p.p.639-660 (1962).
- [1.2] M. JACOB
Heat transfert. Vol.11, p.p.21-24. Wiley, New-York (1957).
- [1.3] H.C. HOTTEL
Geometrical problems in radiant heat transfert Heat Transfert Lectures, Vol.II, NEPA-970, I.E.R.-13 Fairchild Corp. (1949).
- [1.4] J.P. BABARY, F. ROUBELLAT
Etablissement d'un modèle mathématique d'un four à diffusion.
Note Interne L.A.A.S. 72.I.11 - juillet 1972
- [1.5] M. AMOUREUX, J.P. BABARY
Etude d'un modèle mathématique de four à diffusion.
Note Interne L.A.A.S. 72.I.12 - juillet 1972
- [1.6] L.H. SARMIENTO SUESCUN
Contribution à l'étude d'un four à diffusion: Modélisation-Identification.
Thèse Doctorat Université, Toulouse 1973.
- [1.7] A. BERGEAUD
Etude d'une commande numérique d'un four à diffusion. Thèse Doctorat de Spécialité, Toulouse 30 septembre 1974.

- [1.8] A.P. SAGE
Optimum Systems Control;p.p. 150-153
Prentice-Hall, London 1968.
- [1.9] A.C.ROBINSON
A Survey of Optimal Control of Distributed
Parameter systems
Automatica,Vol.7,p.p.371-388 - (1971)
- [1.10] M.PERLMUTTER,R.SIEGEL
Effet of thermal radiation exchange in a tube
on convective heat transfer to a transparent gas.
Amer.Soc.Mech.Engrs;Paper 61-WA-169 (1961).
- [1.11] C.M. USISKIN,R.SIEGEL
Thermal radiation from a cylindrical enclosure
with specified wall heat flux.
Trans.A.S.M.E.;J.Heat Transfer,82,n°4,p.p.369-
374 (1960).
-

C H A P I T R E I I

Introduction

II-1 - Méthode d'identification

II-1.1 - Introduction

II-1.2 - Méthode de Fibonacci

II-1.3 - Méthode Sphéra

II-2 - Discrétisation totale du modèle.

II-3 - Application de la méthode Sphéra et résultats.

Conclusion

Bibliographie

INTRODUCTION

Dans le chapitre précédent nous avons obtenu une représentation mathématique du four à diffusion, qui a été possible en considérant certaines hypothèses simplificatrices et certains paramètres physiques dont les valeurs numériques étaient mal connues.

Il est alors indispensable de déterminer les valeurs de ces coefficients conduisant à un modèle dont la solution est la plus voisine possible du comportement du processus réel.

La similitude entre modèle et processus ne pouvant pratiquement pas être parfaite on est conduit à choisir un critère de minimisation de la "distance" entre le comportement réel et celui obtenu par simulation.

Cette identification paramétrique a été effectuée de la façon suivante. Dans un premier temps, du fait de l'utilisation d'un calculateur numérique, une discrétisation spatiale et temporelle du modèle constitué par un système d'équations aux dérivées partielles a été nécessaire.

Puis, l'application d'un échelon de puissance sur chaque zone du four a permis un enregistrement numérique de l'évolution des températures correspondantes.

Le critère à minimiser étant une forme quadratique sur la différence des températures obtenues au niveau du modèle d'une part et du processus d'autre part et la complexité du modèle étant relativement importante, nous avons choisi un algorithme de minimisation le plus simple possible, nous assurant néanmoins une bonne convergence. Nous avons choisi la méthode Sphéra [II.1] qui est une généralisation de celle de Fibonacci.

Nous rappelons dans la première partie de ce chapitre ce que sont les méthodes de Fibonacci et de Sphéra, puis

dans la deuxième partie nous appliquerons ces méthodes à notre problème d'identification.

II.1 - METHODE D'IDENTIFICATION

II.1.1.- Introduction

Quand on dispose d'une caractérisation du système physique qui donne une structure paramétrable (linéaire, non linéaire, stationnaire ou non, différentielle, discrète, etc ...), on est amené à comparer les comportements des deux systèmes : Physique et mathématique (ou modèle).

La théorie a établi des méthodes d'approche efficaces qui donnent entière satisfaction pour des processus simulés sur calculateurs analogiques, numériques ou hybrides.

Il s'agit d'une méthode itérative capable d'améliorations à chaque pas d'itération dont le principe est représenté sur la figure II.1.

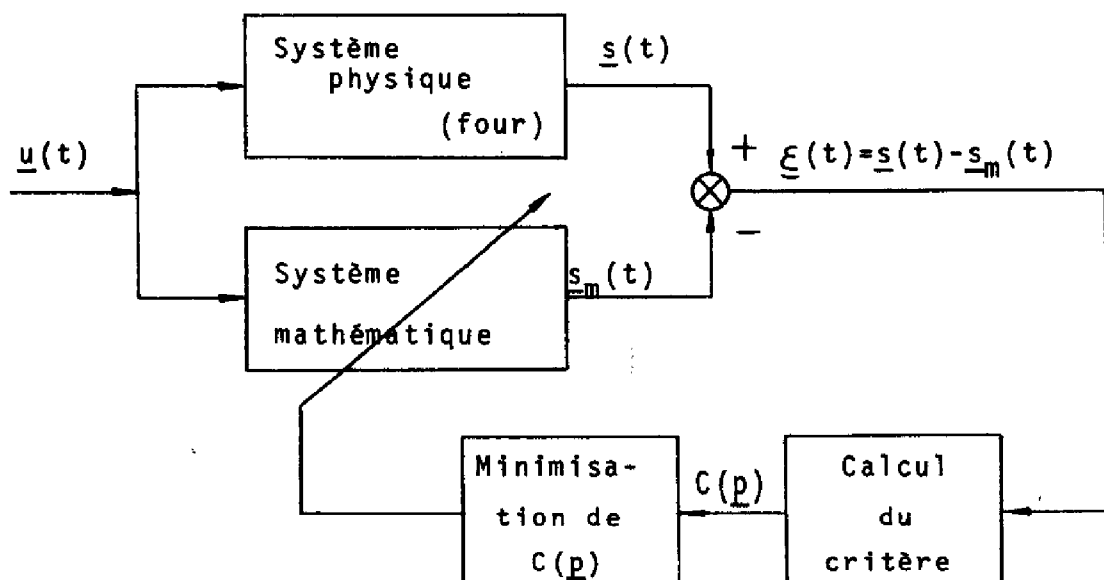


Figure II.1

Soit $\underline{p}$ le vecteur des paramètres inconnus et $c(\underline{p})$ la valeur du critère de performance associé.

Les méthodes heuristiques ont l'avantage d'une mise en oeuvre facile et ne nécessitent aucune connaissance analytique de la fonction $c(\underline{p})$. Elles consistent à explorer l'espace paramétrique par essais successifs et convergent assez rapidement vers la vallée, chaque méthode ne différant que par la stratégie plus ou moins efficace adoptée pour la suivre.

Les méthodes analytiques utilisent par contre, les propriétés analytiques des fonctions iso- c ; c'est-à-dire, les dérivées premières ou deuxièmes ou la réduction des iso- c à leur quadrique osculatrice au voisinage du minimum.

Les méthodes analytiques du second ordre sont déconseillées dans le cas où la valeur initiale du critère à minimiser est très différente de la valeur optimale.

N'ayant que deux paramètres à identifier et sachant que les méthodes heuristiques sont les moins sensibles aux bruits, nous avons utilisé ces dernières.

II.1.2.- Méthode de Fibonacci [II.1] - [II.3] -

Elle ne sera pas utilisée comme méthode d'identification du système mais servira, par la suite, au niveau du développement de la méthode d'identification Sphéra [II.1]

Nous présentons le principe de cette méthode sur l'exemple suivant, lequel consiste à déterminer le minimum d'une fonction unimodale $y = f(x)$ lorsqu'on dispose de n expériences pour explorer séquentiellement un intervalle normalisé $[x_0, x_{n+1}] = [0, 1]$.

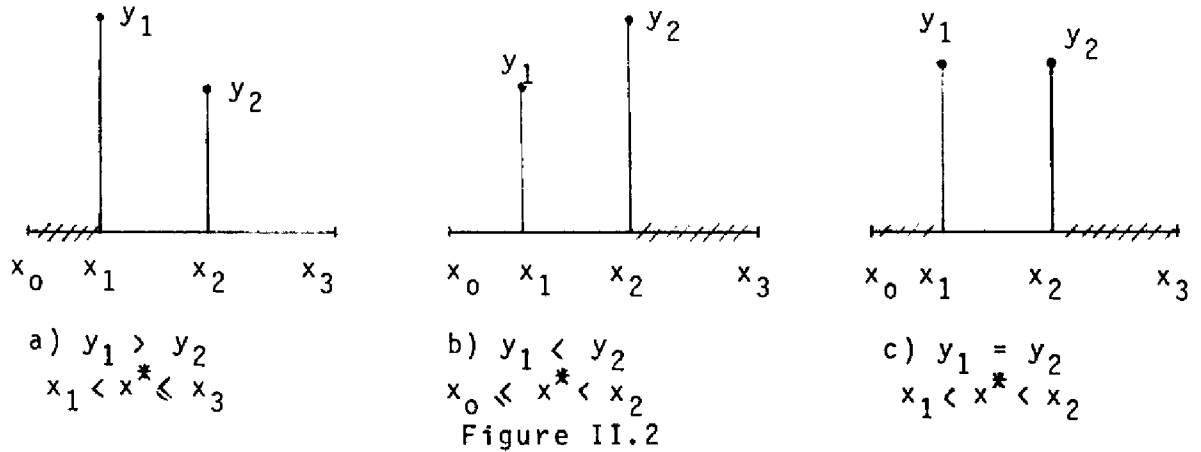
L'unimodalité, lorsqu'elle est assurée, permet d'affirmer après examen des résultats d'un couple quelconque d'expériences, que l'optimum x^* appartient à un intervalle plus petit que l'intervalle initial.

Intervalle d'incertitude

Considérons deux expériences x_1 et x_2 avec $x_1 < x_2$

Les valeurs correspondantes de la fonction seront y_1 et y_2

Trois cas peuvent se présenter comme le montre la figure II.2.



Pour $y_1 > y_2$ (figure II.2.a) le minimum est dans l'intervalle $\underline{[x_1, x_3]}$. Par contre si $y_1 < y_2$ (figure II.2.b) le minimum se trouve dans l'intervalle $\underline{[x_0, x_2]}$. Le cas de la figure II.2.c peut se ramener à l'un ou l'autre des cas précédents.

On appelle intervalle d'incertitude L_2 après deux expériences le plus grand de ces intervalles.

$$L_2 \triangleq \max_{k=1,2} \{ l_2(x_k, k) \} \quad (2.1)$$

où l_2 est la longueur des intervalles

$$l_2(x_k, k) \triangleq x_{k+1} - x_{k-1} \quad k = 1, 2$$

$$l_2(x_1, 1) \triangleq x_2 - x_0 \quad l_2(x_2, 2) \triangleq x_3 - x_1$$

Par la suite, nous verrons que le nombre n d'expériences nécessaires est calculé à partir du choix de la précision désirée au niveau de l'identification.

Pour cet ensemble de n expériences on écrira d'une façon générale :

$$L_n(x_k) = \max_{1 \leq k \leq n} \{l_n(x_k, k)\} \quad (2.2)$$

Si x_k^* représente la suite optimale des expériences et L_n^* l'intervalle d'incertitude final correspondant, on peut écrire

$$L_n^* \triangleq \min_{x_k} \{L_n(x_k)\} \quad \text{et} \quad L_n(x_k^*) \triangleq L_n^* \quad (2.3)$$

soit

$$L_n^* \triangleq \min_{x_k} \max_{1 \leq k \leq n} \{l_n(x_k, k)\} \quad (2.4)$$

Détermination des deux premières expériences

Considérons deux expériences telles que

$$0 \leq x_1 \leq x_2 \leq 1$$

Nous pouvons définir l'intervalle d'incertitude L_2 après ces deux expériences [II.3] par

$$L_2 = \max \{(x_2 - x_0), (x_3 - x_1)\} = \max \{x_2, 1 - x_1\} \quad (2.5)$$

Considérons $x_2 = x_1 + \epsilon$; ϵ étant au moins la plus petite valeur qui permet de différencier pratiquement les valeurs de y_1 et y_2 .

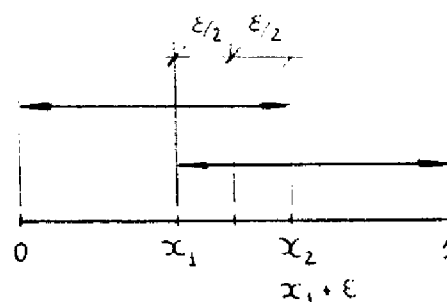
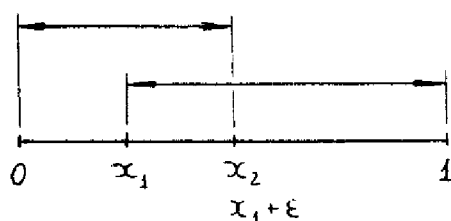
$$\text{Alors} \quad L_2 = \max \{x_1 + \epsilon, 1 - x_1\} \quad (2.6)$$

La valeur optimale de L_2 est obtenue, si les deux segments sont égaux, c'est-à-dire si $x_1 + \epsilon = 1 - x_1$ d'où

$$x_1 = \frac{1}{2} - \frac{\varepsilon}{2} \quad \text{et} \quad x_2 = \frac{1}{2} + \frac{\varepsilon}{2}$$

La valeur de L_2^* est alors :

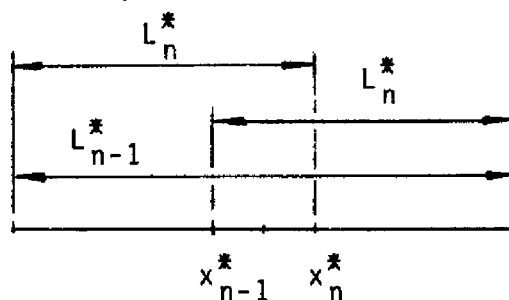
$$L_2^* = \frac{1}{2} + \frac{\varepsilon}{2}$$



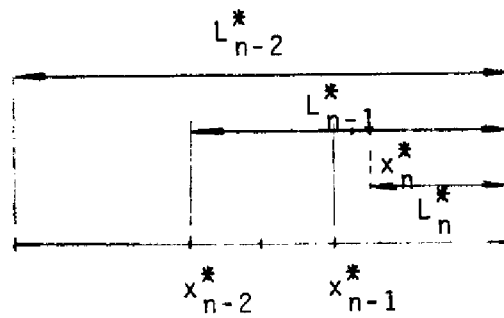
II.1.2.1.- Méthode de recherche

Supposons que $n-1$ expériences aient été effectuées de façon optimale. Il nous reste à explorer un intervalle d'incertitude de valeur L_{n-1}^* . Dans cet intervalle se trouve l'expérience x_{n-1}^* qui a donné à y la plus petite valeur des $n-1$ expériences.

Le point x_n^* distant ε de x_{n-1}^* est, comme on a vu plus haut, le symétrique de x_{n-1}^* par rapport au milieu de l'intervalle L_{n-1}^* représentant par la suite la précision désirée au niveau des paramètres à identifier. Nous écrirons :



Considérons maintenant l'intervalle L_{n-2}^* dans lequel se trouve le point x_{n-2}^* qui a donné la plus petite valeur des $n-2$ premiers essais. La meilleure façon de placer x_{n-1}^* est à nouveau la symétrique de x_{n-2}^* , ce qui nous permet d'écrire



$$L_{n-2}^* = L_{n-1}^* + L_n^* \quad (2.9)$$

ou encore $L_{n-2}^* = 3L_n^* - \epsilon \quad (2.10)$

On peut généraliser la relation (2.9) au cas de trois intervalles d'incertitude consécutifs

$$L_{j-1}^* = L_j^* + L_{j+1}^* \quad (2.11)$$

pour tout $2 < j < n-1$

Pour $j = n-2$ il vient

$$L_{n-3}^* = L_{n-2}^* + L_{n-1}^* = (3L_n^* - \epsilon) + (2L_n^* - \epsilon) = 5L_n^* - 2\epsilon$$

$j = n-3$

$$L_{n-4}^* = 8L_n^* - 3\epsilon$$

$j = n-4$

$$L_{n-5}^* = 13L_n^* - 5\epsilon$$

et d'une façon générale

$$L_{n-k}^* = F_{k+1} L_n^* - F_{k-1} \epsilon \quad (2.12)$$

$$K = 1, 2, \dots, n$$

$$\text{où } F_0 = F_j = 1$$

$$F_{k+1} = F_k + F_{k-1}$$

$$K = 1, 2, \dots, n \quad (2.13)$$

Les F_k forment une suite où F_k est appelé k-ième nombre de Fibonacci.

Soit L_1 la longueur de l'intervalle initial, il vient

$$L_1 = F_n L_n^* - F_{n-2} \cdot \varepsilon \quad (2.14)$$

La fraction L_n^* de l'intervalle original subsistant après n expériences séquentielles d'une recherche optimale s'obtient ainsi :

$$L_n^* = \frac{L_1}{F_n} + \frac{F_{n-2}}{F_n} \cdot \varepsilon \quad (2.15)$$

II.1.2.2. - Mise en oeuvre de la méthode [II.1]

Nous nous proposons donc de trouver le minimum d'une fonction dans un intervalle donné L_1 avec une précision donnée P . L_n est l'intervalle d'incertitude après n expériences et contient le minimum de la fonction. On peut poser :

$$P = \frac{L_n^*}{2} \quad (2.16)$$

ce que nous vérifierons par la suite.

En effet si ε est la distance entre deux expériences, l'intervalle d'incertitude L_n^* minimum que l'on peut espérer est $L_n^* = 2\varepsilon$. Sur la figure II.3 on a tracé L_{n-1} et quatre expériences : Deux qui fixent ces bornes et deux autres qui nous serviront pour choisir L_n .

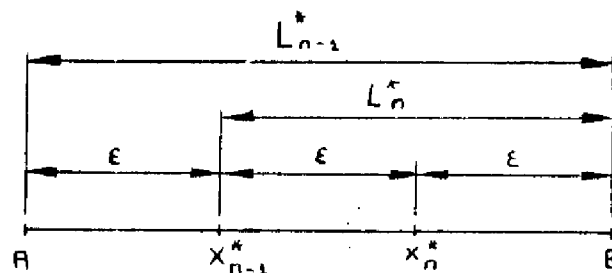


Figure II.3.

L'intervalle L_n est le plus petit possible car les expériences X_{n-1}^* , X_n^* et B sont au moins distantes de ϵ les unes des autres.

L'expression (2.14) devient

$$L_1 = \frac{L_n^*}{2} [2 F_n - F_{n-2}] \quad (2.17)$$

or, des égalités

$$F_{n+1} = F_n + F_{n-1}$$

et $F_n = F_{n-1} + F_{n-2}$

on déduit

$$F_{n+1} = 2F_n - F_{n-2} \quad (2.18)$$

La formule (2.17) s'écrit maintenant :

$$L_1 = \frac{L_n^*}{2} \cdot F_{n+1} = P \cdot F_{n+1} \quad (2.19)$$

11.1.2.3.- Détermination du nombre de mesures.

A fin de déterminer l'indice $n+1$ de l'expression (2.19), il faut calculer le produit $P \cdot F_i$, pour $i = 1, 2, \dots, n+1$ jusqu'à ce que $P \cdot F_{n+1} \geq L_1$.

Si la valeur de n est telle que $P \cdot F_{n+1} > L_1$, la précision effective ne sera plus P mais :

$$P_1 = \frac{L_1}{F_{n+1}} \quad \text{avec } P_1 < P \quad (2.20)$$

Remarque: Le nombre d'expériences qu'il faut effectuer est n .

II.1.2.4 - Détermination des points de mesure.

Considérons un intervalle de recherche $L_1 [A, B]$ où on a placé les deux premières expériences x_1 et x_2 , qui déterminent les intervalles L_2^* et L_3^* tels que $L_1 = L_2^* + L_3^*$. Figure II.4.

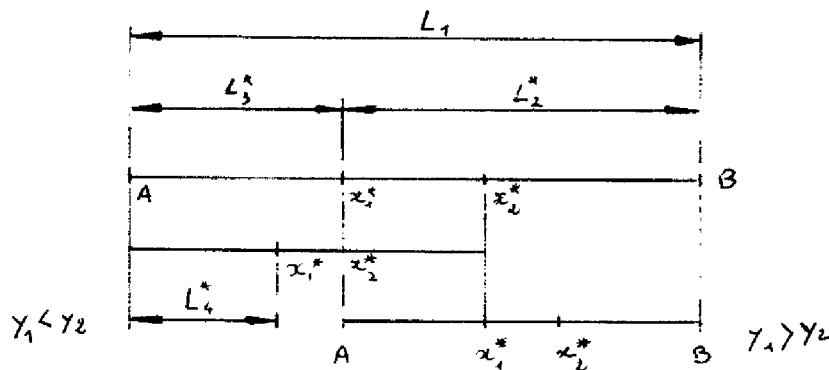


Figure II.4.

L_2^* est l'intervalle borné par une extrémité de l'intervalle initial L_1 et une expérience (x_1^* ou x_2^*), il reste donc $n-2$ mesures à faire dans cet intervalle.

A partir des expressions (2.12), (2.16) et (2.18) on obtient :

$$L_{n-k}^* = F_{k+1} \cdot 2P - F_{k-1} \cdot P = F_{k+2}P \quad (2.21)$$

En faisant $n-k=2$ on obtient la longueur de l'intervalle $L_2^* = F_n \cdot P$

de meme pour $n-k=3$ on a :

$$L_3^* = F_{n-1} \cdot P$$

Donc si a est l'abscisse de A , les abscisses de x_1^* et x_2^* sont :

$$x_1^* = a + F_{n-1} \cdot P$$

$$x_2^* = a + F_n \cdot P$$

Soient y_1 et y_2 les valeurs de $y(x)$ aux points x_1^* et x_2^* . Deux cas peuvent se présenter :

a) Si $y_1 < y_2$, la fonction $y(x)$ étant unimodale, on élimine l'intervalle $[x_2^*, B]$. Il reste l'intervalle $L_2 = [A, x_2^*]$ avec l'ancien point x_1^* qui va jouer le rôle de x_2^* .

Le nouveau point x_1^* est :

$$x_1^* = a + L_4^* = a + F_{n-2} \cdot P$$

b) Si $y_1 > y_2$, on élimine $[A, x_1^*]$, on pose $a = x_1^*$,

$x_1^* = x_2^*$ et on calcule le nouvel x_2^* pour la formule :

$$x_2^* = a + L_3^* = a + F_{n-1} \cdot P$$

II.1.2.5.- Organigramme.

Il est représenté sur la figure II.5.

Le langage Fortran exigeant des indices différents de zéro, on a initialisé la suite de Fibonacci comme suit :

$$F_1 = 1 \qquad F_2 = 1$$

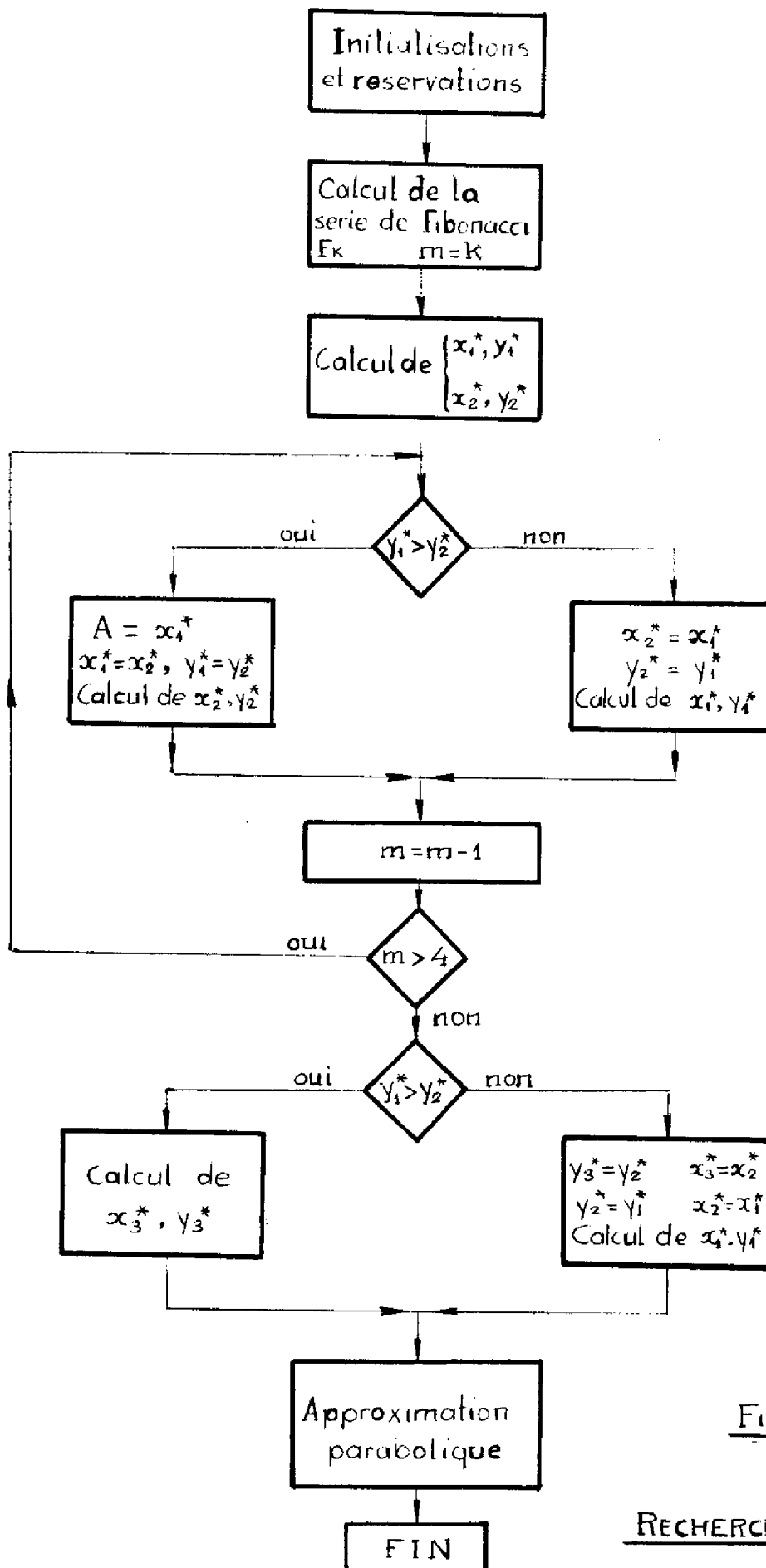


Figure II.5

RECHERCHE DE FIBONACCI

Si K est l'indice tel que

$$P.F_k \geq L_1$$

le nombre de mesures est $K-2$.

Initialisons m à la valeur K (dernier indice de la série de Fibonacci); décroissant de une unité à partir de la troisième mesure, le calcul sera terminé pour :

$$m = K - (K-2) + 2 = 4.$$

Approximation parabolique finale .

Si la fonction $y = f(x)$ a une allure parabolique autour du minimum, on peut accroître la précision de la méthode en remplaçant le meilleur point obtenu à l'aide de la méthode précédente par le minimum ($\hat{y}$) de la parabole qui passe par les trois derniers points (Figure II.6).

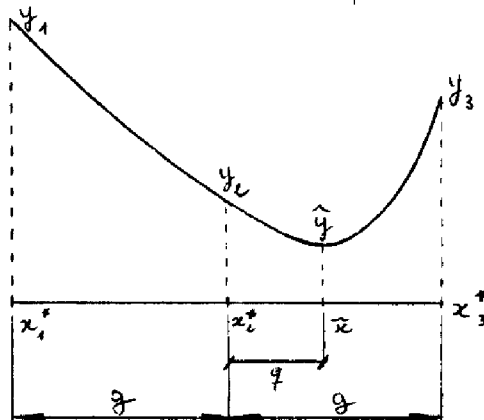


Figure II.6.

Soient y_1 , y_2 et y_3 les valeurs de $f(x)$ aux points de mesure. Le développement à l'ordre deux de $y = f(x)$ en série de Taylor au voisinage de x_2^* s'écrit :

$$y_1 = y_2 - g \cdot \dot{y}_2 + \frac{g^2}{2} \cdot \ddot{y}_2$$

$$\text{de même } y_3 = y_2 + g \cdot \dot{y}_2 + \frac{g^2}{2} \ddot{y}_2$$

$\dot{\hat{y}} = \dot{y}_2 + g \cdot \ddot{y}_2 = 0$ sera vérifié au point $\hat{y} = y(\hat{x})$ minimum de la parabole, soit :

$$q = -\frac{\dot{\ddot{y}}_2}{\ddot{y}_2} = \frac{g}{Z} \cdot \frac{y_1 - y_3}{y_1 + y_3 + 2y_2}$$

$$\text{et } \hat{x} = x_2 + q$$

Dans tous les exemples que nous avons traités par cette méthode nous avons toujours trouvé un point qui améliorerait la valeur du critère. Son utilisation devient très intéressante dans le cas où le minimum de la fonction se trouve à l'extérieur de l'intervalle initial de recherche et relativement éloigné de celui-ci.

II.1.3 - Méthode Sphéra

Cette méthode d'identification que nous avons utilisée est une généralisation de la méthode de Fibonacci. Elle présente l'avantage d'une mise en oeuvre facile et possède une convergence rapide. Nous exposerons son principe sur un exemple à deux dimensions avant de la généraliser au cas de problèmes à n dimensions.

Considérons le cas de la figure II.7, où sont représentées les courbes iso - $C_{(p)}$.

Détermination de la direction de recherche.

Soit $A(x_i, y_i)$ le point de départ. Traçons un cercle de centre A et de rayon R . Un point M du cercle a pour équations :

$$\begin{aligned} & \begin{cases} x = x_i + R \cos \theta \\ y = y_i + R \sin \theta \end{cases} \\ \text{soit} & \begin{cases} x = x_i (1+r \cos \theta) \\ y = y_i (1+r \sin \theta) \end{cases} \end{aligned} \quad (2.21)$$

où $R = x_i \cdot r$.

En faisant varier θ entre zéro et 2π , nous pourrions trouver l'angle θ_{opt} qui donnera la valeur la plus petite du critère $C(p)$ sur le cercle. Cette recherche sera effectuée à l'aide de la méthode de Fibonacci proposée au paragraphe II.1.2.

Détermination du meilleur point sur la direction.

Une fois l'angle θ_{opt} trouvé, il est intéressant de connaître le point B qui se trouve dans cette direction et le plus près de la vallée.

Les coordonnées de ce point s'écrivent :

$$\begin{cases} X_b = X_i(1 + \lambda \cos \theta_{opt}) \\ Y_b = Y_i(1 + \lambda \sin \theta_{opt}) \end{cases} \quad (2.22)$$

On applique une fois de plus la méthode de Fibonacci avec λ comme paramètre, qui nous conduira à ce point. On recommence ces deux procédures à partir du point B obtenu. On peut donc alterner les recherches des minimums par rapport à l'angle θ (R étant constant) et par rapport aux points situés dans une direction θ_{opt} .

Procédure de sécurité :

Si le point B qu'on vient de trouver est proche de la vallée, la recherche sur le cercle n'est pas nécessairement unimodale. Les intersections θ_1 et θ_2 de la vallée avec le cercle minimisent localement le critère $C(p)$, bien que le vrai minimum soit θ_1 . (Figure II.7.)

Si B est voisin de la vallée et le rayon R du cercle est petit on peut écrire :

$$\theta_1 \approx \theta_2 + \pi$$

ce qui donne l'angle θ_1 dans le cas où on aurait trouvé θ_2 . (En ce point le critère a une valeur supérieure à celle qu'il a au point B).

Procédure_finale :

Si les deux points θ_1 & θ_2 du cercle donnent une valeur du critère supérieure à celle du centre, on peut assurer que le minimum est à l'intérieur de ce cercle. Si nous considérons que la précision n'est pas suffisante, on peut diviser par deux le rayon et continuer la procédure.

Généralisation_à_n_paramètres :

La meilleure direction sera détectée sur une hypersphère à n dimensions de rayon R . Un point de cette hypersphère aura pour coordonnées :

$$\left\{ \begin{array}{l} X_1 = X_{10} + R \cos \alpha_1 \cos \alpha_2 \dots \cos \alpha_{n-1} \\ X_2 = X_{20} + R \cos \alpha_1 \cos \alpha_2 \dots \sin \alpha_{n-1} \\ \vdots \\ X_n = X_{n0} + R \sin \alpha_1 \end{array} \right. \quad (2.23)$$

On effectuera donc une recherche à $n - 1$ dimensions sur les paramètres $\alpha_1, \alpha_2, \dots, \alpha_{n-1}$. La détermination du meilleur point sur cette direction sera effectuée comme précédemment.

La méthode Sphéra peut être assimilée à une méthode du gradient mais où l'on calculerait le gradient sur l'hypersphère de rayon R non infiniment petit. La méthode est plus globale, moins sensible au bruit et donc plus efficace.

II.2 - DISCRETISATION TOTALE DU MODELE

L'application numérique de la méthode précédente entraîne la nécessité de discrétiser le modèle mathématique défini au chapitre I.

sa discrétisation dans le temps et dans l'espace conduit à des équations récurrentes relativement simples et

d'utilisation aisée au niveau d'une simulation [II.6] ou de l'identification paramétrique du modèle.

Pour des raisons logiques nous choisirons comme points de discrétisation ceux qui correspondent à l'emplacement des thermocouples.

Considérons le système régi par les équations (1.15).

W_j^n représente la valeur de la température moyenne de la paroi au point j , à l'instant $n \Delta t$. (avec n entier $0 \leq n \leq N$).

On pose pour tout j

$$\alpha_j = W_j^{n+1} - W_j^n$$

$\Delta x_j = x_j - x_{j-1}$ distance séparant deux points de discrétisation. Les équations discrétisées s'obtiennent [II.6], [II.7] en remplaçant :

$$\left[\frac{\partial^2 W}{\partial x^2} \right]_j^n \text{ par } \frac{1}{\Delta x_j + \Delta x_{j+1}} \left[\theta \left(\frac{\alpha_{j+1} - \alpha_j}{\Delta x_{j+1}} - \frac{\alpha_j - \alpha_{j-1}}{\Delta x_j} \right) + \frac{W_{j+1}^n - W_j^n}{\Delta x_{j+1}} - \frac{W_j^n - W_{j-1}^n}{\Delta x_j} \right]$$

L'expression de $\left[\frac{\partial^2 G}{\partial x^2} \right]_j^n$ est analogue à la précédente mais les α_j et W sont remplacés respectivement par les $\beta_j = G_j^{n+1} - G_j^n$ et G .

Pour $(W_j^{n+1})^4$ on admet l'approximation suivante :

$$(W_j^{n+1})^4 = (W_j^n + \alpha_j)^4 \simeq (W_j^n)^4 + 4 (W_j^n)^3 \alpha_j.$$

alors

$$\left(\frac{\partial^2 W^4}{\partial x^2}\right)_j^n \simeq \frac{2}{\Delta x_j + \Delta x_{j+1}} \left[\theta \frac{4(W_{j+1}^n)^3 \alpha_{j+1} - 4(W_j^n)^3 \alpha_j}{\Delta x_{j+1}} - \theta \frac{4(W_j^n)^4 \alpha_j - 4(W_{j-1}^n)^4 \alpha_{j-1}}{\Delta x_j} + \right. \\ \left. + \frac{(W_{j+1}^n)^4 - (W_j^n)^4}{\Delta x_{j+1}} - \frac{(W_j^n)^4 - (W_{j-1}^n)^4}{\Delta x_j} \right]$$

et

$$\left(\frac{\partial G}{\partial x}\right)_j^n \simeq \frac{1}{\Delta x_j + \Delta x_{j+1}} \left[\theta (\beta_{j+1} - \beta_{j-1}) + G_{j+1}^n - G_{j-1}^n \right]$$

Les équations (1.15) s'écrivent alors, après regroupement des α_j et β_j

$$X_j \alpha_{j+1} + Y_j \alpha_j + Z \alpha_{j-1} = S_j \quad 1 \leq j \leq N-1 \quad (2.24)$$

$$L_j \beta_{j+1} + K_j \beta_j + N_j \beta_{j-1} = V_j \quad (2.25)$$

avec

$$X_j = - \frac{8A\theta (W_{j+1}^n)^3 + 2B\theta}{\Delta x_{j+1} (\Delta x_j + \Delta x_{j+1})}$$

$$Y_j = \frac{8A\theta (W_j^n)^3 + 2B\theta \left(\frac{1}{\Delta x_j} + \frac{1}{\Delta x_{j+1}} \right)}{\Delta x_j + \Delta x_{j+1}} + \frac{1}{\Delta t}$$

$$Z_j = - \frac{8A\theta (W_{j-1}^n)^3 + 2B\theta}{\Delta x_j (\Delta x_j + \Delta x_{j+1})}$$

$$\begin{aligned} S_j = & \frac{2A}{\Delta x_j + \Delta x_{j+1}} \left(\frac{(W_{j+1}^n)^4 - (W_j^n)^4}{\Delta x_{j+1}} - \frac{(W_j^n)^4 - (W_{j-1}^n)^4}{\Delta x_j} \right) + \\ & + \frac{2B}{\Delta x_j + \Delta x_{j+1}} \left(\frac{W_{j+1}^n + W_j^n}{\Delta x_{j+1}} - \frac{W_j^n - W_{j-1}^n}{\Delta x_j} \right) - \frac{2B\theta\beta_{j+1}}{\Delta x_{j+1} (\Delta x_j + \Delta x_{j+1})} + \\ & + \frac{2B\theta\beta_j}{\Delta x_j + \Delta x_{j+1}} \left(\frac{1}{\Delta x_j} + \frac{1}{\Delta x_{j+1}} \right) - \frac{2B\theta\beta_{j-1}}{\Delta x_j (\Delta x_j + \Delta x_{j+1})} - \\ & - \frac{2B}{\Delta x_j + \Delta x_{j+1}} \left(\frac{G_{j+1}^n - G_j^n}{\Delta x_{j+1}} - \frac{G_j^n - G_{j-1}^n}{\Delta x_j} \right) - c(W_j - G_j) - E.q_a(x_j, t) \end{aligned}$$

$$L_j = \frac{F \theta}{\Delta x_j + \Delta x_{j+1}}$$

$$K_j = \frac{1}{\Delta t}$$

$$N_j = - \frac{F \theta}{\Delta x_j + \Delta x_{j+1}}$$

$$V_j = M (W_j^n - G_j^n) - \frac{G_{j+1}^n - G_{j-1}^n}{\Delta x_j + \Delta x_{j+1}}$$

Les conditions aux limites s'introduisent dans les équations par l'intermédiaire des variables portant l'indice zero ou N.

Des expressions α et β on déduit:

$$W_j^{n+1} = W_j^n + \alpha_j$$

$$G_j^{n+1} = G_j^n + \beta_j$$

Du système (2.25), on déduit les valeurs des β_j , $1 \leq j \leq N - 1$. β_0 et β_N sont donnés par les conditions aux limites.

La résolution du système (2.25) nous fournit les valeurs des β_j , donc des G_j^{n+1} . Les α_j ($1 \leq j \leq N - 1$) sont alors solutions du système (2.24), α_0 et α_N étant calculés à partir des conditions aux limites.

Les simulations que nous avons effectuées ont montré que la solution était à peu près la même pour des valeurs de θ comprises entre 0,5 et l'unité.

Plus délicat était le choix de Δt . Ne pouvant pas le calculer de façon théorique, plusieurs essais ont été

nécessaires, qui ont montré que pour $\Delta t \geq 2,7 \times 10^{-4}$ le système divergeait toujours, ce qui nous a conduit à choisir un pas d'intégration $\Delta t = 2,5 \times 10^{-4}$ h.

II.3. - APPLICATION DE LA METHODE SPHERA ET RESULTATS

L'identification du four est effectuée, pour avoir plus de précision, dans la gamme de température normale d'utilisation.

Pour la plupart des processus thermiques les modèles caractérisant la montée et la descente en température sont différents. En particulier, les constantes de temps correspondantes sont assez inégales.

Puisque, pratiquement, compte tenu des perturbations, il est toujours nécessaire de chauffer, l'identification du modèle a été effectuée en régime de chauffe du four.

Nous avons donc déterminé un modèle valable entre 750°K et 1050°K. Pour cela, un mini-ordinateur en ligne avec le processus permet d'appliquer simultanément sur les trois zones de chauffe, un échelon de puissance calorifique pendant une durée T nettement supérieure au temps de réponse du système et d'amplitude suffisante pour se situer dans la gamme de température désirée. Il enregistre séquentiellement la température de la paroi du tube de quartz aux six points où sont placés les thermo-couples et délivre ces valeurs sur un ruban : $T_{w_i}(j)$; $i = 1 \text{ à } 6$; $j = 1 \text{ à } M$ avec $M \cdot \Delta T = T$, ΔT étant la période d'échantillonnage. De cette façon on a la possibilité d'identifier le four "hors ligné", à l'aide d'un calculateur numérique scientifique.

L'ajustement des paramètres se fait en minimisant un critère d'écart entre les comportements du système réel et du modèle. La température moyenne du gaz est sensiblement identique à celle de la paroi, de ce fait nous ne tiendrons compte que de cette dernière. Nous avons choisi

le critère suivant :

$$J = \int_0^T \underline{z}^T(t) Q \underline{z}(t) dt$$

où $\underline{z}(t)$ est un vecteur de dimension six de composantes

$$z_i(t) = T_{wi}(t) - T_{mi}(t) \quad i = 1, 2, \dots, 6$$

$T_{mi}(t)$ est la valeur de la température calculée à l'aide du modèle.

Q est une matrice de pondération diagonale, définie positive.

$T_{wi}(t)$ n'est connu qu'aux instants d'échantillonnage. Alors nous approcherons le calcul de l'intégrale au moyen de la formule de Hardy

$$J = \frac{\Delta T}{140} \sum_{j=0, M-6, 6} \left\{ 41.f(j+6) + 216.f(j+5) + 27.f(j+4) + \right. \\ \left. + 272.f(j+3) + 27.f(j+2) + 216.f(j+1) + 41.f(j) \right\}$$

$$\text{où } f(j) = \underline{z}^T(t_j) Q \underline{z}(t_j).$$

Après une initialisation des paramètres à identifier, l'ordinateur effectue le calcul itératif qui conduit à la minimalisation du critère d'écart choisi et par conséquent aux valeurs des paramètres recherchés.

Etant donné que P_f et C_{pf} interviennent toujours sous forme de produits nous allons les considérer comme un seul paramètre, le deuxième étant le coefficient β .

La durée T pendant laquelle on a laissé évoluer le four en boucle ouverte est 20 heures, la période d'échantillonnage étant de trois minutes, on obtient ainsi quatre cents valeurs par thermo-couple.

La matrice Q du critère a été choisie sous la forme

suivante :

$$Q = \text{diag } [0,1 \quad 1 \quad 1 \quad 1 \quad 1 \quad 0,1]$$

Cela vient du fait que les hypothèses formulées lors de la modélisation étaient surtout vérifiées dans la zone centrale, mais moins sur les zones latérales pour lesquelles les gradients thermiques sont plus importants.

L'application directe de la méthode Sphera, avec le critère que nous venons de définir, aux équations du four donne les valeurs suivantes des paramètres :

$$\rho_g \cdot C_{pg} = 56,6 \quad \text{Kcal/m}^3 \cdot ^\circ\text{K}$$

$$\text{et } \beta = 0,17$$

L'erreur quadratique relative

$$e = \frac{\int_0^T \underline{z}^T(t) Q \underline{z}(t) dt}{\int_0^T \underline{I}_w^T(t) Q \underline{I}_w(t) dt}$$

$$\underline{T}_w(t) = [T_{w1}(t) \quad T_{w2}(t) \quad \dots \quad T_{w6}(t)]^T$$

que nous avons ainsi obtenue est $e = 0,00138$.

Les figures II.8, II.9 et II.10 nous montrent l'évolution des paramètres et du critère pendant l'identification.

En II.11 on donne les équations discrétisées du système après identification.

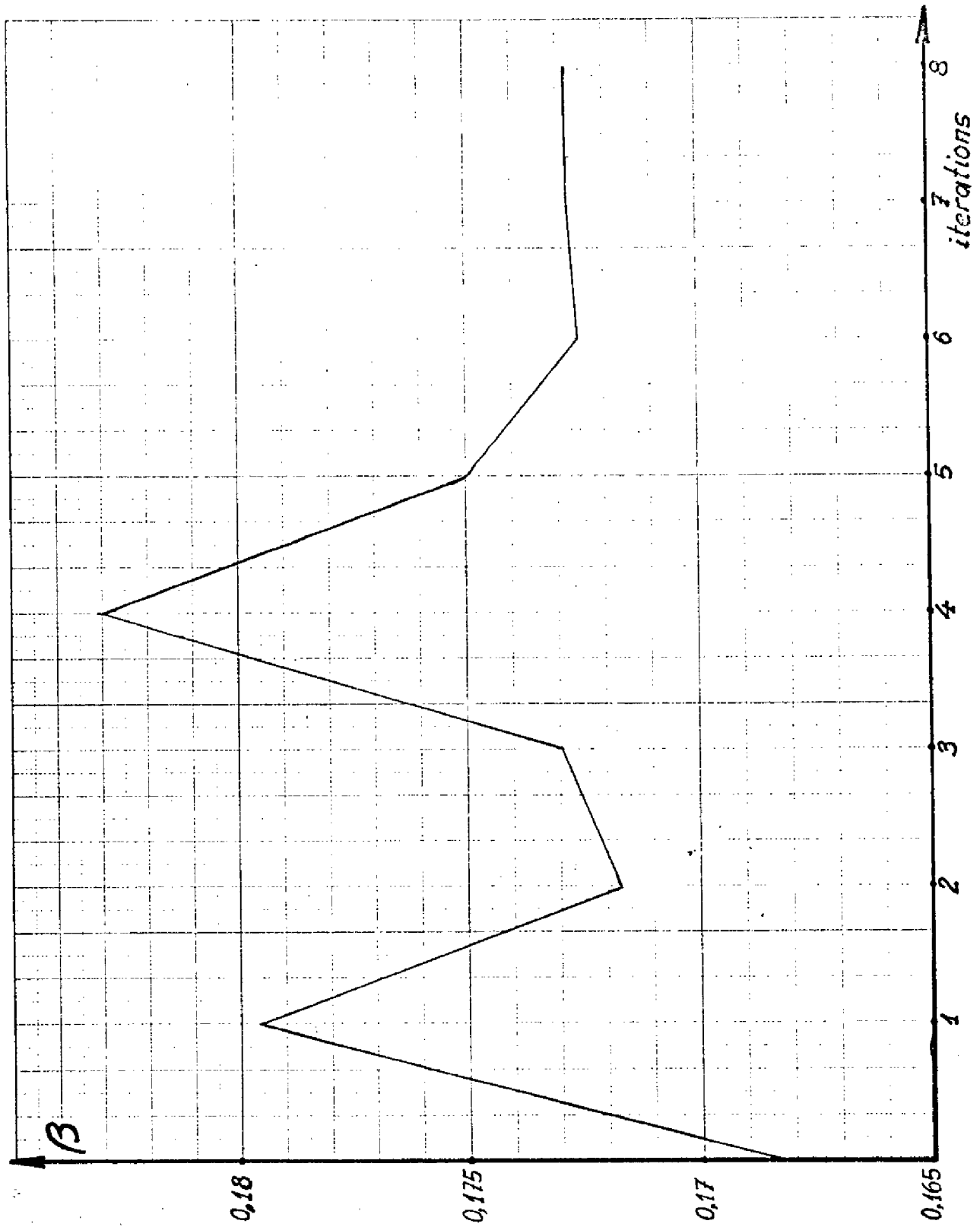


FIGURE 11-8.- Variation de β

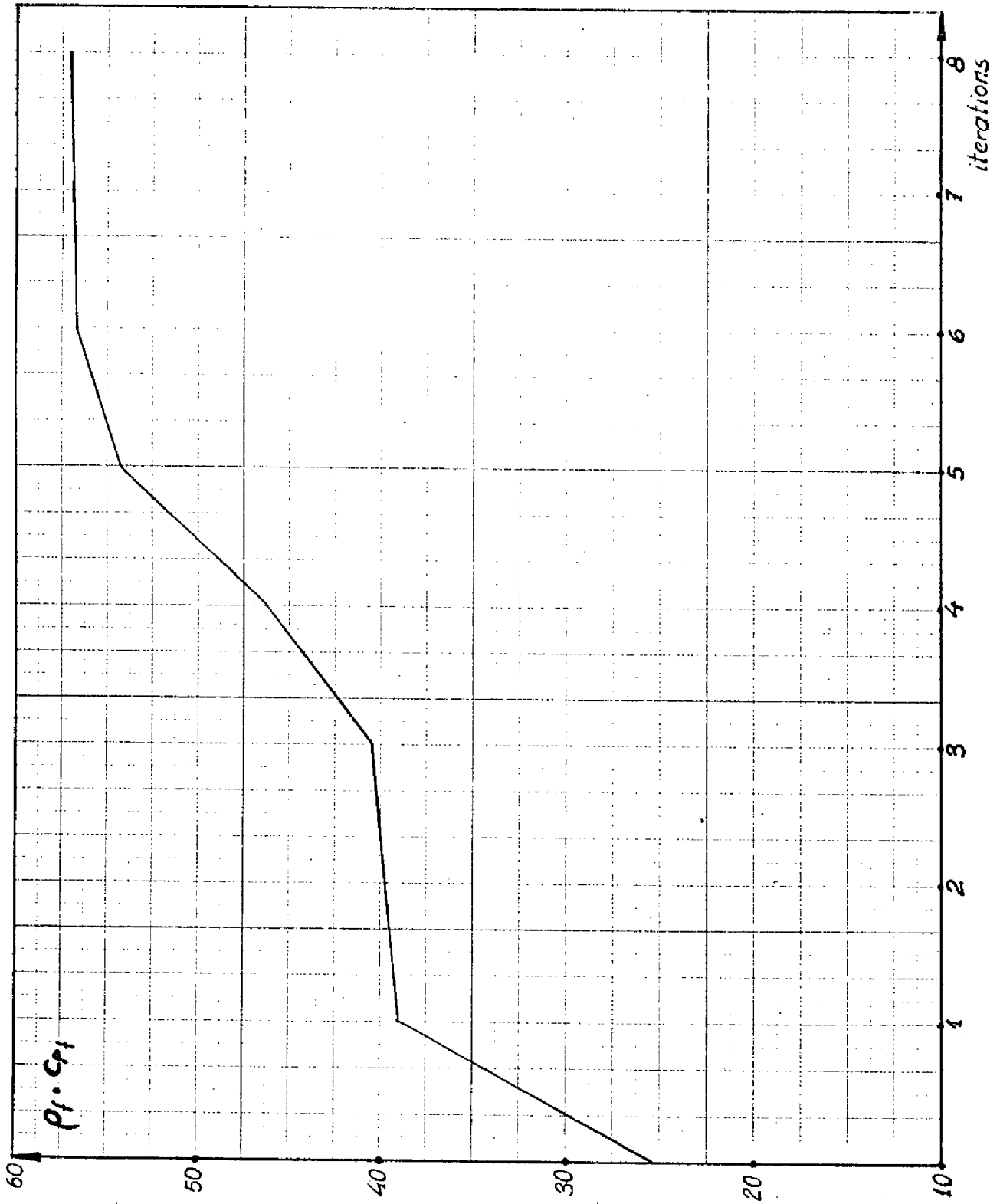


FIGURE II.9.- Variation de $p_1 \cdot C_{p_1}$

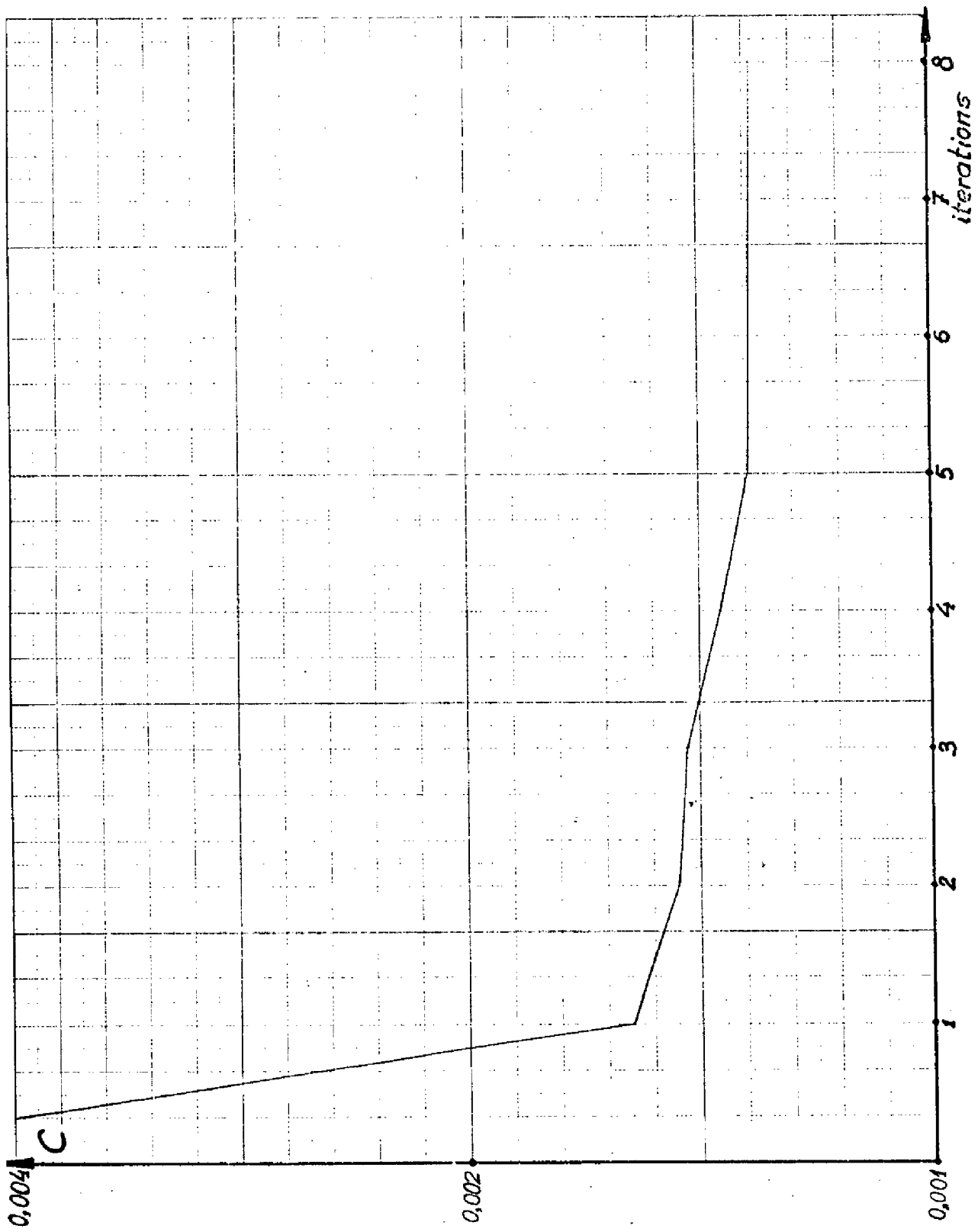


FIGURE II-10 .- Evolution du Critère

$$\begin{bmatrix} \dot{W}_1 \\ \dot{W}_2 \\ \dot{W}_3 \\ \dot{W}_4 \\ \dot{W}_5 \\ \dot{W}_6 \\ \dot{G}_1 \\ \dot{G}_2 \\ \dot{G}_3 \\ \dot{G}_4 \\ \dot{G}_5 \\ \dot{G}_6 \end{bmatrix} = \begin{bmatrix} -3,5941 & 0,0744 & & & & & 3,5941 & -0,0744 & & & & & -5,9829 W_1^* + 2,1621 W_2^* & 0,165 & 0 & 0 \\ 0,0621 & -3,5331 & 0,0621 & & & & -0,0621 & 3,5331 & -0,0619 & & & & 1,806 W_1^* - 3,610 W_2^* + 1,804 W_3^* & 0,095 & 0,038 & 0 \\ 0,0620 & -3,5328 & 0,0619 & & & & -0,062 & 3,5328 & -0,0619 & & & & 1,802 W_2^* - 3,603 W_3^* + 1,8 W_4^* & 0 & 0,066 & 0 \\ 0,0619 & -3,5327 & 0,0619 & & & & -0,0619 & 3,5327 & -0,0619 & & & & 1,799 W_3^* - 3,593 W_4^* + 1,797 W_5^* & 0 & 0,066 & 0 \\ 0,0619 & -3,5329 & 0,0621 & & & & -0,0619 & 3,5329 & -0,0621 & & & & 1,801 W_4^* - 3,606 W_5^* + 1,801 W_6^* & 0 & 0,038 & 0,095 \\ 0,0744 & -3,5941 & & & & & -0,0744 & 3,5941 & & & & & 2,162 W_5^* - 5,382 W_6^* & 0 & 0 & 0,165 \\ 7078,64 & & & & & & -7078,64 & -73,057 & & & & & & & & & \\ 7078,64 & & & & & & 61,023 & -7078,64 & -61,023 & & & & & & & & \\ 7078,64 & & & & & & 60,959 & -7078,64 & -60,959 & & & & & & & & \\ 7078,64 & & & & & & 60,923 & -7078,64 & -60,923 & & & & & & & & \\ 7078,64 & & & & & & 60,985 & -7078,64 & -60,98 & & & & & & & & \\ 7078,64 & & & & & & 73,057 & -7078,64 & & & & & & & & & \end{bmatrix} + 10^{-10} \begin{bmatrix} P_1 \\ P_2 \\ P_3 \end{bmatrix}$$

II.11 Equations d'état du système

C O N C L U S I O N

Une structure d'un modèle dynamique du four à diffusion étant établie, nous nous sommes intéressés à l'identification des paramètres inconnus, pour une plage de températures donnée. Pour ceci, nous avons défini un critère de "ressemblance" ou d'écart, dont la minimisation a été effectuée à l'aide de la méthode Sphéra.

C'est à partir de ce modèle ainsi établi, que nous déterminerons dans le prochain chapitre un algorithme de commande en boucle fermée réalisant la minimisation d'un critère quadratique sur l'état, la commande et une fonction semi-définie positive sur l'état du processus.

B I B L I O G R A P H I E

- [~II.1_] J.RICHALET, A.RAULT,R.POULIQUEN
Identification des processus par la méthode
du modèle Gordon & Breach.
- [~II.2_] Daniel GRAUPE
Identification of Systems
Van Nostrand Reinhold Company (New York),1972.
- [~II.3_] D.J.WILDE
Méthodes de recherche d'un optimum
Dunod 1966.
- [~II.4_] A.LAVI, T.P. WOGL
Recent advances in optimisation techniques.
J.Wiley & Sons,1966.
- [~II.5_] A.J.FOSSARD, M.GAUVRIT, M.GEEGUEN, E.TOUMIRE
Structures identifiables des systèmes multiva-
riables et leurs applications à la commande
des processus.
Rapport E.N.S.A.E.-C.E.R.A. n°66.00.227,
décembre 1967.
- [~II.6_] J.E. DOUCET
Programme de simulation d'un four à diffusion.
Note Interne L.A.A.S.- S.I.S. 75.I.01.
- [~II.7_] R.D. RICHTMYER, K.W. MORTON
Difference Methods for initial value problems.
John Wiley 1 & Sons - 1967.
- [~II.8_] M.GAUVRIT, J.F. LEMAITRE
Analyse et commande d'un four de rechauffage
Rapport E.N.S.A.E.- C.E.R.A. n° 68.01.483,
avril 1970.
-

CHAPITRE III

Introduction

- III.1. - Position du problème
- III.2. - Le problème de poursuite pour un système linéaire
- III.3. - Obtention d'une loi de commande en boucle fermée
sous-optimale
- III.4. - Détermination des éléments de l'algorithme de
régulation

Conclusion

Bibliographie

INTRODUCTION

Dans les chapitres précédents nous avons donné une représentation mathématique du fonctionnement du four. La méthode Sphéra nous a permis de déterminer les coefficients inconnus du modèle que nous avons ensuite transformé à l'aide d'une méthode de discrétisation spatiale en vue de pouvoir appliquer des techniques d'optimisation déjà établies pour les systèmes à paramètres localisés.

Cette discrétisation nous a conduit à un modèle dynamique non linéaire sous la forme d'un système d'équations différentielles ordinaires d'ordre douze.

La détermination d'une loi de commande optimale en boucle fermée d'un tel système conduit généralement à une mise en oeuvre relativement complexe.

A l'exception de quelques systèmes très particuliers, des techniques d'approximation permettent de déterminer une commande sous-optimale minimisant un certain critère de performance.

Après avoir présenté les principaux résultats liés à la résolution d'un problème de poursuite dans le cas de systèmes linéaires, nous appliquerons ces résultats au cas de systèmes non linéaires, pour lesquels on supposera constant l'état que l'on désire atteindre.

III.1. - POSITION DU PROBLEME

Ecrivons les équations de la figure II.11 sous la forme plus générale suivante :

$$\dot{\underline{x}}(t) = A \underline{x}(t) + \eta \cdot f(\underline{x}) + B \underline{u}(t) \quad (3.1)$$

où $\underline{x}(t)$ est le vecteur d'état de dimension n (pour notre

problème : $\underline{x}(t) = [W_1(t) \dots W_6(t) G_1(t) \dots G_6(t)]^T$,
 $A_{n \times n}$ et $B_{n \times m}$ sont des matrices constantes, $f(\underline{x})$ est un
 vecteur de fonctions qui contient la partie non linéaire
 du système, η est un petit scalaire constant, et $\underline{u}(t)$
 est le vecteur de commande de dimension m .

D'un point de vue physique, nous nous proposons, en
 partant d'un état initial $\underline{x}(0) = \underline{x}_0$, d'arriver à un état
 désiré $\underline{x}_d$ constant, au bout d'un temps T , tout en mi-
 nimisant un critère de la forme

$$J(\underline{x}_0, \underline{u}, \underline{x}_d, \eta) = Q[\underline{e}(t_f), t_f] + \int_{t_0}^{t_f} L[\underline{e}(t), \underline{u}(t), t, \eta] dt \quad (3.2.)$$

$$\text{où } \underline{e}(t) = \underline{x}_d - \underline{x}(t) \quad (3.3.)$$

$$\text{pour } t_0 \leq t \leq t_f \quad ; \quad t_f - t_0 = T$$

φ est une fonction non linéaire (souvent quadratique)
 de l'erreur à l'instant t_f .

L est une fonction non linéaire de l'erreur et de la
 commande.

Le problème qui se pose maintenant est la recherche de
 la commande $\underline{u}(t)$ pour le système qui satisfait les con-
 ditions énoncées précédemment et minimise l'indice de
 performance (3.3).

III.2. - LE PROBLEME DE POURSUITE POUR UN SYSTEME LINEAIRE

Dans un premier temps rappelons quelques résultats con-
 cernant les systèmes linéaires et qui nous serviront
 par la suite.

Considérons un système linéaire observable et comman-
 dable décrit par les équations :

$$\left. \begin{aligned} \dot{\underline{x}}(t) &= A(t) \cdot \underline{x}(t) + B(t) \cdot \underline{u}(t) \\ \underline{y}(t) &= C(t) \cdot \underline{x}(t) \end{aligned} \right\} \quad (3.4)$$

$$\underline{x}(t_0) = \underline{x}_0$$

où $\underline{x}(t)$ est le vecteur d'état, $\underline{u}(t)$ est le vecteur de commande et $\underline{y}(t)$ représente le vecteur de sortie [III.1].

Désignons par $\underline{z}(t)$ la sortie désirée et $\underline{e}(t)$ sa différence avec la sortie réelle.

$$\underline{e}(t) = \underline{z}(t) - \underline{y}(t) = \underline{z}(t) - C(t) \cdot \underline{x}(t) \quad (3.5)$$

Remarque : Le cas où $\underline{z}(t)$ est constant et $C(t) = \mathbb{1}$ correspond à celui envisagé au paragraphe III.1.

D'une façon générale, considérons le problème de la minimisation d'un indice de performance tel que :

$$J(\underline{\hat{e}}, \underline{\hat{u}}) = \underline{e}^T(t_f) \cdot S(t_f) \underline{e}(t_f) + \int_{t_0}^{t_f} [\underline{e}^T(t) Q(t) \underline{e}(t) + \underline{u}^T(t) \cdot R(t) \cdot \underline{u}(t)] dt \quad (3.6)$$

où $S(t_f)$ et $Q(t_f)$ sont des matrices symétriques semi-définies positives et $R(t)$ est une matrice symétrique définie positive.

Pour simplifier les notations, nous ne préciserons plus les indices de la variable temporelle en argument des différentes fonctions du temps du système.

L'hamiltonien associé à (3.4) et (3.6) est défini par :

$$H = (\underline{z} - C\underline{x})^T Q (\underline{z} - C\underline{x}) + \underline{u}^T R \underline{u} + \underline{\lambda}^T (A\underline{x} + B\underline{u}) \quad (3.7)$$

L'équation d'Hamilton-Jacobi-Bellman est définie par :

$$\frac{\partial V(\underline{x}, t)}{\partial t} + H(\underline{x}, \underline{\lambda}, \underline{u}) = 0 \quad (3.8)$$

$$\text{où } \underline{\lambda} = \frac{\partial V(\underline{x}, t)}{\partial \underline{x}} \quad (3.9)$$

$$\text{avec la condition finale } V[\underline{x}(t_f), t_f] = \underline{e}^T(t_f) S(t_f) \underline{e}(t_f) \quad (3.10)$$

La résolution de (3.8) est en général très difficile.

Une simplification importante sera apportée dans le cas où on choisit $\underline{u}$ comme étant une fonction de $\underline{x}$. La minimisation de l'Hamiltonien par rapport à $\underline{u}$ donne :

$$\frac{\partial H}{\partial \underline{u}} = 0 = 2.R.\underline{u} + B^T \underline{\lambda}$$

soit $\underline{u} = -\frac{1}{2} R^{-1} . B^T . \underline{\lambda}$ (3.11)

L'hamiltonien s'écrit alors :

$$H = (\underline{z} - C\underline{x})^T Q(\underline{z} - C\underline{x}) + \underline{\lambda}^T A \underline{x} - \frac{1}{4} \underline{\lambda}^T B . R^{-1} . B^T \underline{\lambda} \quad (3.12)$$

Tenant compte de la condition (3.9) et de la relation (3.12) nous écrirons l'équation d'Hamilton-Jacobi-Bellman sous la forme suivante :

$$\begin{aligned} \frac{\partial V(\underline{x}, t)}{\partial t} + \underline{z}^T . Q . \underline{z} - 2\underline{x}^T . C^T . Q . \underline{z} + \underline{x}^T . C^T . Q . C . \underline{x} + \frac{\partial V(\underline{x}, t)^T}{\partial \underline{x}} A \underline{x} \\ - \frac{1}{4} \frac{\partial V(\underline{x}, t)^T}{\partial \underline{x}} . B . R^{-1} . B^T . \frac{\partial V(\underline{x}, t)}{\partial \underline{x}} = 0 \end{aligned}$$

(3.13)

Des relations (3.5) et (3.10) on déduit :

$$\begin{aligned} V[\underline{x}(t_f), t_f] = \underline{z}^T(t_f) . S(t_f) . \underline{z}(t_f) - 2\underline{x}^T(t_f) . C^T(t_f) . \\ S(t_f) . \underline{z}(t_f) + \underline{x}^T(t_f) C^T(t_f) S(t_f) C(t_f) . \\ \underline{x}(t_f) \end{aligned} \quad (3.14)$$

Nous allons chercher une solution de $V(\underline{x}, t)$ qui satisfait l'expression (3.13) et qui est semblable à la relation (3.14), ce qui nous permet d'écrire :

$$V(\underline{x}, t) = h(t) + 2.\underline{x}^T . \underline{k}(t) + \underline{x}^T . M(t) . \underline{x} \quad (3.15)$$

où $M(t)$ est une matrice $n \times n$ symétrique, $\underline{k}(t)$ est un vecteur de dimension n et $h(t)$ est un scalaire.

En reportant la relation (3.15) dans la relation (3.13) on obtient :

$$\begin{aligned} \underline{x}^T \dot{M} \underline{x} + 2 \underline{x}^T \underline{k} + \dot{h} + \underline{z}^T Q \underline{z} - 2 \underline{x}^T C^T Q \underline{z} + \underline{x}^T C^T Q C \underline{x} + \\ + (2M \underline{x} + 2\underline{k})^T A \underline{x} - (M \underline{x} + \underline{k})^T B R^{-1} B^T (M \underline{x} + \underline{k}) = 0 \end{aligned} \quad (3.16)$$

Les conditions suffisantes pour que (3.16) soit toujours vérifiée quels que soient $\underline{x}(t)$ et $\underline{z}(t)$ sont donc :

$$\dot{M} + C^T Q C + 2 M A - M B R^{-1} B^T M = 0 \quad (3.17)$$

$$\dot{\underline{k}} = C^T Q \underline{z} + (M B R^{-1} B^T - A^T) \underline{k} \quad (3.18)$$

$$\dot{h} + \underline{z}^T Q \underline{z} - \underline{k}^T B R^{-1} B^T \underline{k} = 0 \quad (3.19)$$

(3.17) est une équation de Riccati et il a été démontré par Kalman [III.9_] qu'elle possède une solution unique définie positive.

Nous pouvons déduire à partir des relations (3.14) et (3.15) les conditions finales pour les équations (3.17) (3.18) et (3.19). On obtient alors par identification

$$M(t_f) = C^T(t_f) S(t_f) C(t_f) \quad (3.20)$$

$$\underline{k}(t_f) = - C^T(t_f) S(t_f) \underline{z}(t_f) \quad (3.21)$$

$$h(t_f) = \underline{z}(t_f) S(t_f) \underline{z}(t_f) \quad (3.22)$$

La loi de commande optimale s'écrit donc :

$$\underline{u}(t) = - \frac{1}{2} R^{-1}(t) B^T(t) \frac{\partial V(\underline{x}, t)}{\partial \underline{x}(t)} = - R^{-1}(t) B^T(t) \cdot$$

$$[M(t) \cdot \underline{x}(t) + \underline{k}(t)]$$

(3.23)

La fonction $h(t)$ ne figurant pas dans l'expression de la commande optimale $\underline{u}(t)$ on ne tient donc pas compte de la condition (3.19).

Le schéma bloc de simulation du problème optimal de poursuite et représenté sur la figure III.1.

Si l'horizon d'intégration $t_f - t_0$ est suffisamment grand, au bout d'un temps assez long (par rapport aux constantes de temps) on peut considérer $\dot{M}(t) \neq 0$ et résoudre l'équation de Riccati en régime statique. Si $\underline{z}(t) = \underline{z}$ (vecteur constant) on aura également $\underline{k}(t) \neq 0$ et $\underline{k}$ sera approximé par l'expression :

$$\underline{k} \cong - [M^{-1}B^T - A^T]^{-1} C^T Q \underline{z} \quad (3.24)$$

III.3. - OBTENTION D'UNE LOI DE COMMANDE EN BOUCLE FERMEE SOUS-OPTIMALE

Ayant rappelé les principaux résultats concernant la résolution du problème d'optimisation formulé en (3.6) d'un système linéaire, nous exposons dans ce paragraphe la manière dont nous avons utilisé ces résultats pour résoudre le même problème dans le cas d'un système non linéaire.

En réalité, l'étude d'une commande sous-optimale pour un système non linéaire a été traitée par de nombreux auteurs de façons différentes.

Sannuti et Kokotović (III.2) ont proposé une "méthode des perturbations singulières". Pour une catégorie de systèmes non linéaires, ils ont établi une condition suffisante pour que la commande optimale soit continue et dérivable par rapport à des perturbations qui modifient l'ordre du système, les "perturbations singulières".

Nishikawa, Sannomiya et Itakura (III.3) ont considéré la non linéarité comme une perturbation du système et

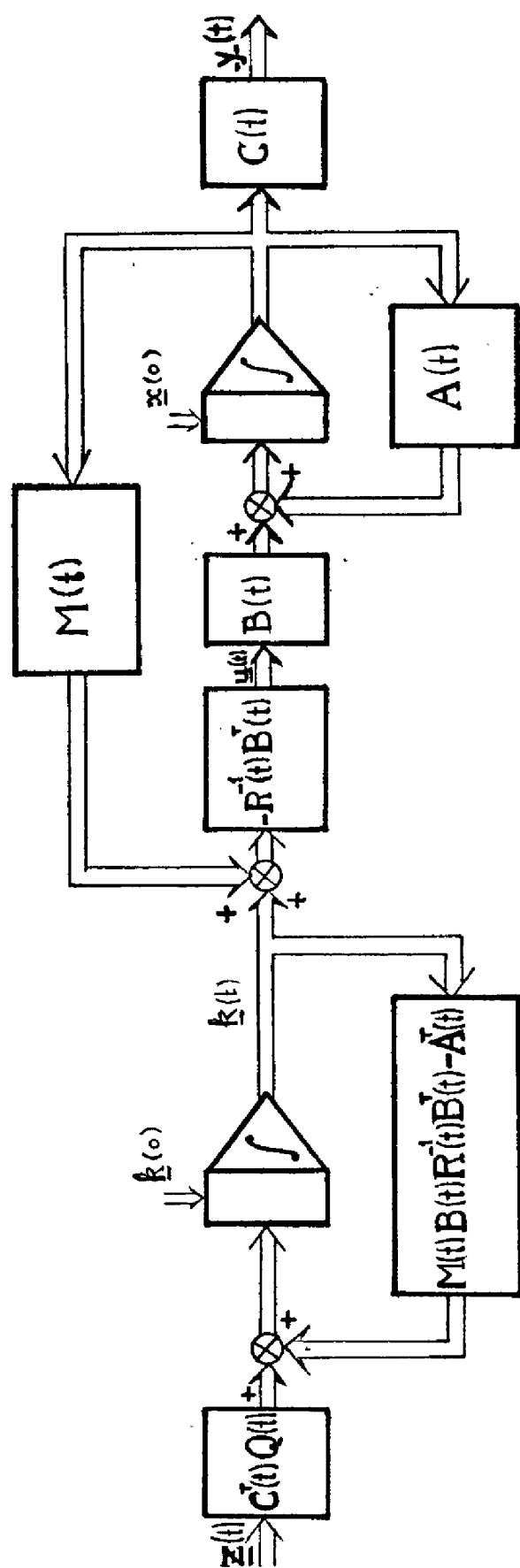


FIGURE III-1.- Schema de simulation.

ont introduit un paramètre ϵ pour la représenter. Un développement en série selon les puissances de ϵ , donne un système d'équations différentielles dont les solutions fournissent une loi de commande sous-optimale.

Pearson (III.4) optimise un système non stationnaire par rapport à un indice de performance, en considérant le système comme linéaire et stationnaire à un instant donné.

Garrard et ses co-auteurs (III.5), (III.6), (III.7) ont utilisé différentes méthodes pour trouver une solution approchée de l'équation d'Hamilton-Jacobi-Bellman.

Burghart (III.8) développe une loi de commande sous-optimale en développant en série de Taylor la matrice des gains de retour d'état pour le système linéarisé à chaque instant.

Pour ce qui concerne notre problème d'optimisation, nous avons préféré utiliser la méthode développée par Garrard (III.5).

Le type d'algorithme de commande sous-optimale obtenu par cette méthode, tient compte à la fois des conditions d'optimisation à satisfaire et des conditions de mise en oeuvre sur le calculateur.

Description de la méthode

Considérons un système dynamique commandable régi par les équations

$$\dot{\underline{x}}(t) = A.\underline{x}(t) + \eta . \underline{f} [\underline{x}(t)] + B.\underline{u}(t) \quad (3.25)$$

ou $\underline{x}(t)$ est le vecteur d'état, $\underline{x} \in R^n$; $\underline{u}(t)$ est le vecteur de commande, $\underline{u} \in R^m$; η est un scalaire constant et $\underline{f} [\underline{x}(t)]$ est un vecteur de fonctions non linéaires de l'état vérifiant la condition $\underline{f}(\underline{0}) = \underline{0}$. A de dimension $n \times n$ et B de dimension $n \times m$ sont deux matrices constantes.

Dans ce qui suit, le vecteur de sortie est identifié au vecteur d'état : $\underline{y}(t) = \underline{x}(t)$ et la sortie désirée $\underline{z}(t)$ est supposée constante $\underline{z}(t) = \underline{x}_d$.

Nous nous proposons de trouver la commande sous-optimale $\underline{u}(t)$ qui permet de transférer l'état du système d'une valeur initiale $\underline{x}(0) = \underline{x}_0$ à une valeur finale $\underline{x}(t_f) = \underline{x}_d$ en minimisant l'indice de performance J .

$$J(\underline{x}_0, \underline{u}, \underline{x}_d, \eta) = \underline{e}^T(t_f) S \underline{e}(t_f) + \int_{t_0}^{t_f} [\underline{e}^T(t) Q \underline{e}(t) + \eta g(\underline{e}) + \underline{u}^T(t) R \underline{u}(t)] dt \quad (3.26)$$

où F et Q sont deux matrices symétriques $n \times n$ semi-définies positives et R est une matrice symétrique $m \times m$ définie positive.

$$\underline{e}(t) = \underline{x}_d - \underline{x}(t)$$

$g[\underline{e}(t)]$ est une fonction de $\underline{e}(t)$ tel que :

$$g(\underline{e}) \geq 0$$

$$g(\underline{0}) = 0$$

Le paramètre η est le même que celui défini précédemment.

Ecrivons l'hamiltonien H de la façon suivante :

$$H(\underline{x}, \underline{\lambda}, \underline{u}, \eta) = (\underline{x}_d - \underline{x})^T Q (\underline{x}_d - \underline{x}) + \eta \cdot g(\underline{e}) + \underline{u}^T R \underline{u} + \underline{\lambda}^T [\underline{A} \underline{x} + \eta \cdot \underline{f}(\underline{x}) + B \underline{u}] \quad (3.27)$$

L'équation d'Hamilton-Jacobi-Bellman correspondante est (III.10)

$$\frac{\partial V(\underline{x}, t)}{\partial t} + H(\underline{x}, \underline{\lambda}, \underline{u}, \eta) = 0 \quad (3.28)$$

$$\text{ou } \underline{\lambda} = \frac{\partial V(\underline{x}, t)}{\partial \underline{x}}$$

avec la condition finale : $V(\underline{x}(t_f), t_f) = \underline{e}^T(t_f) \cdot S \cdot \underline{e}(t_f)$

La minimisation de l'Hamiltonien par rapport à $\underline{u}$ donne :

$$\underline{u}(t) = -\frac{1}{2} R^{-1} B^T \underline{\lambda} = -\frac{1}{2} R^{-1} B^T \cdot \frac{\partial V(\underline{x}, t)}{\partial \underline{x}} \quad (3.29)$$

Des équations (3.27), (3.28), et (3.29) on déduit :

$$\begin{aligned} \frac{\partial V}{\partial t} + \underline{x}_d^T Q \underline{x}_d - 2 \underline{x}_d^T Q \underline{x} + \underline{x}^T Q \underline{x} + \frac{\partial V^T}{\partial \underline{x}} [-A \underline{x} + \eta \cdot \underline{f}(\underline{x})] \\ + \eta \cdot g(\underline{e}) - \frac{1}{4} \frac{\partial V^T}{\partial \underline{x}} B R^{-1} B^T \frac{\partial V}{\partial \underline{x}} = 0 \end{aligned}$$

(3.30)

Developpons maintenant $V(\underline{x}, t)$ suivant les puissances de η (III.5)

$$V(\underline{x}, t) = \sum_{n=2}^{\infty} \eta^{n-2} v_n(\underline{x}, t) \quad (3.31)$$

En remplaçant $V(\underline{x}, t)$ dans (3.30) par son développement (3.31) et en identifiant par rapport aux puissances de η on obtient le système différentiel suivant :

$$\begin{aligned} \frac{\partial v_2}{\partial t} + \underline{x}_d^T Q \underline{x}_d - 2 \underline{x}_d^T Q \underline{x} + \frac{\partial v_2^T}{\partial \underline{x}} A \underline{x} - \frac{1}{4} \frac{\partial v_2^T}{\partial \underline{x}} B \cdot R^{-1} \cdot B^T \frac{\partial v_2}{\partial \underline{x}} + \\ + \underline{x}^T Q \underline{x} = 0 \end{aligned}$$

$$\begin{aligned}
& \frac{\partial V_3}{\partial t} + \frac{\partial V_3^T}{\partial \underline{x}} A \underline{x} + \frac{\partial V_2^T}{\partial \underline{x}} \underline{f}(\underline{x}) + g(\underline{e}) + \frac{1}{4} \frac{\partial V_2^T}{\partial \underline{x}} B R^{-1} B^T \frac{\partial V_3}{\partial \underline{x}} - \frac{1}{4} \frac{\partial V_3^T}{\partial \underline{x}} B R^{-1} B^T \frac{\partial V_2}{\partial \underline{x}} = 0 \\
& \vdots \\
& \frac{\partial V_i}{\partial t} + \frac{\partial V_i^T}{\partial \underline{x}} A \underline{x} + \frac{\partial V_{i-1}^T}{\partial \underline{x}} \underline{f}(\underline{x}) - \frac{1}{4} \sum_{\substack{k \geq 2 \\ l \geq 2}}^{i+2} \frac{\partial V_k^T}{\partial \underline{x}} B R^{-1} B^T \frac{\partial V_l}{\partial \underline{x}} = 0
\end{aligned} \tag{3.32}$$

avec $K + 1 = i + 2$

Si nous écrivons $\underline{f}(\underline{x})$ sous la forme $F(\underline{x})\underline{x}$ et si nous choisissons $g(\underline{e}) = \underline{e}^T D(\underline{e}) \underline{e} + \underline{b}^T(\underline{e}) \underline{e}$, où $D(\underline{e})$ est une matrice symétrique et $\underline{b}(\underline{e})$ est un vecteur de même dimension que $\underline{e}$.

Les équations précédentes s'écrivent alors :

$$\begin{aligned}
& \frac{\partial V_2}{\partial t} + \underline{x}_d^T Q \underline{x}_d - 2 \underline{x}_d^T Q \underline{x} + \underline{x}^T Q \underline{x} + \frac{\partial V_2^T}{\partial \underline{x}} A \underline{x} - \frac{1}{4} \frac{\partial V_2^T}{\partial \underline{x}} B R^{-1} B^T \frac{\partial V_2}{\partial \underline{x}} = 0 \\
& \frac{\partial V_3}{\partial t} + \frac{\partial V_3^T}{\partial \underline{x}} A \underline{x} + \frac{\partial V_2^T}{\partial \underline{x}} F(\underline{x}) \underline{x} + \underline{e}^T D(\underline{e}) \underline{e} + \underline{b}^T(\underline{e}) \underline{e} - \frac{1}{4} \frac{\partial V_2^T}{\partial \underline{x}} B R^{-1} B^T \frac{\partial V_3}{\partial \underline{x}} - \\
& \quad - \frac{1}{4} \frac{\partial V_3^T}{\partial \underline{x}} B R^{-1} B^T \frac{\partial V_2}{\partial \underline{x}} = 0 \\
& \vdots \\
& \frac{\partial V_i}{\partial t} + \frac{\partial V_i^T}{\partial \underline{x}} A \underline{x} + \frac{\partial V_{i-1}^T}{\partial \underline{x}} F(\underline{x}) \underline{x} - \frac{1}{4} \sum_{\substack{k \geq 2 \\ l \geq 2 \\ K+1=l+2}}^{i+2} \frac{\partial V_k^T}{\partial \underline{x}} B R^{-1} B^T \frac{\partial V_l}{\partial \underline{x}} = 0
\end{aligned} \tag{3.33}$$

Choisissons V_i de la même façon que dans le cas du

problème de poursuite associé à un système linéaire, soit,

$$V_i(\underline{x}, t) = \underline{x}^T M_i(t) \underline{x} + 2 \underline{x}^T \underline{k}_i(t) + h_i(t) \quad (3.34)$$

$$i = 2, 3, 4, \dots$$

où $M_i(t)$ est une matrice $n \times n$ symétrique

$\underline{k}_i(t)$ est un vecteur de dimension n

et $h_i(t)$ est un scalaire

En combinant (3.33) et (3.34) on obtient le nouveau système :

$$\begin{aligned} & \underline{x}^T \dot{M}_2 \underline{x} + 2 \underline{x}^T \dot{\underline{k}}_2 + \dot{h}_2 + \underline{x}_d^T Q \underline{x}_d - 2 \underline{x}_d^T Q \underline{x} + \underline{x}^T Q \underline{x} + \\ & + (2M_2 \underline{x} + 2\underline{k}_2)^T A \underline{x} - (M_2 \underline{x} + \underline{k}_2)^T B R^{-1} B^T (M_2 \underline{x} + \underline{k}_2) = 0 \end{aligned} \quad (3.35 - a)$$

$$\begin{aligned} & \underline{x}^T \dot{M}_3 \underline{x} + 2 \underline{x}^T \dot{\underline{k}}_3 + \dot{h}_3 + 2(M_3 \underline{x} + \underline{k}_3)^T A \underline{x} + 2(M_2 \underline{x} + \underline{k}_2)^T F(\underline{x}) \underline{x} + \\ & + (\underline{x}_d - \underline{x})^T D(\underline{e})(\underline{x}_d - \underline{x}) + \underline{b}^T (\underline{x}_d - \underline{x}) - 2(M_2 \underline{x} + \underline{k}_2)^T B R^{-1} B^T (M_3 \underline{x} + \underline{k}_3) = 0 \end{aligned} \quad (3.35 - b)$$

.

.

.

.

$$\begin{aligned} & \underline{x}^T \dot{M}_i \underline{x} + 2 \underline{x}^T \dot{\underline{k}}_i + \dot{h}_i + 2(M_i \underline{x} + \underline{k}_i)^T A \underline{x} + 2(M_{i-1} \underline{x} + \underline{k}_{i-1})^T F(\underline{x}) \underline{x} - \\ & - \sum_{\substack{k \geq 2 \\ l \geq 2 \\ k+l=i+2}}^{i+2} (M_k \underline{x} + \underline{k}_k)^T B R^{-1} B^T (M_l \underline{x} + \underline{k}_l) = 0 \end{aligned} \quad (3.35 - i)$$

.

.

.

La commande optimale peut être obtenue à partir des équations (3.35). Il suffit de résoudre successivement

(3.35 -a), (3.35 -b), ..., (3.35 - i) et reporter les résultats en (3.29).

Néanmoins, si η est suffisamment petit on pourra ne tenir compte que des N premières équations du système (3.35).

Calcul des deux premiers termes d'une commande de sous-optimale.

L'équation (3.35 - a) étant satisfaite pour tout $\underline{x}$ nous permet d'écrire :

$$\dot{\underline{M}}_2 + 2\underline{M}_2\underline{A} - \underline{M}_2\underline{B}\underline{R}^{-1}\underline{B}^T\underline{M}_2 + \underline{Q} = \underline{0} \quad (3.35 - a)$$

$$\dot{\underline{k}}_2 - \underline{Q}\underline{x}_d + (\underline{A}^T - \underline{M}_2\underline{B}\underline{R}^{-1}\underline{B}^T)\underline{k}_2 = \underline{0} \quad (3.36 - b)$$

$$\dot{h} + \underline{x}_d^T \underline{Q}\underline{x}_d + \underline{k}_2^T \underline{B}\underline{R}^{-1}\underline{B}^T \underline{k}_2 = 0 \quad (3.36 - c)$$

Ce système est exactement celui obtenu dans le cas linéaire (3.17), (3.18), (3.19).

$\underline{M}_2(t)$ est la solution d'une équation de Riccati, matrice définie positive.

$\underline{k}_2(t)$ est la solution d'un ensemble d'équations différentielles ordinaires.

La résolution de l'équation (3.36 - c) est inutile puisque dans la commande n'intervient que le terme

$$\frac{\partial V}{\partial \underline{x}} \quad \text{lequel est indépendant de } h(t).$$

De la même façon, nous écrivons pour (3.35 - b) :

$$\dot{\underline{M}}_3 + 2\underline{M}_3\underline{A} + 2\underline{M}_2\underline{F}(\underline{x}) + \underline{D}(\underline{e}) - 2\underline{M}_2\underline{B}\underline{R}^{-1}\underline{B}^T\underline{M}_3 = \underline{0} \quad (3.37 - a)$$

$$\dot{\underline{k}}_3 + \underline{A}^T \underline{k}_3 + \underline{F}^T(\underline{x}) \underline{k}_2 - 2\underline{D}(\underline{e}) \underline{x}_d - 2\underline{M}_2\underline{B}\underline{R}^{-1}\underline{B}^T \underline{k}_3 - 2\underline{M}_3\underline{B}\underline{R}^{-1}\underline{B}^T \underline{k}_2 - \underline{b} = \underline{0} \quad (3.37 - b)$$

$$\dot{h}_3 + \underline{x}_d^T \underline{D}(\underline{e}) \underline{x}_d + \underline{b}^T(\underline{e}) \underline{x}_d - 2\underline{k}_2^T \underline{B}\underline{R}^{-1}\underline{B}^T \underline{k}_3 = 0 \quad (3.37 - c)$$

Dans ce système apparaissent la matrice $D(\underline{e})$ et le vecteur $\underline{b}(\underline{e})$ intervenant dans l'expression de $g(\underline{e})$.

La solution la plus commode est de déterminer $D(\underline{e})$ et $\underline{b}(\underline{e})$ a posteriori de telle façon que les calculs et l'expression de la commande soient simplifiés, en essayant de diminuer au maximum l'influence de $\underline{f}(\underline{x})$.

(3.37 = a) devient un ensemble d'équation différentielles ordinaires si on choisit $D(\underline{e})$ de telle sorte que :

$$2M_2 F(\underline{x}) + D(\underline{e}) \equiv W \quad (3.38)$$

où W est une matrice symétrique et constante.

De même si on choisit $\underline{b}$ tel que :

$$F^T(\underline{x}) \underline{k}_2 - 2D(\underline{e}) \underline{x}_d - \underline{b} \equiv \underline{0} \quad (3.39)$$

alors les équations (3.38) s'écrivent

$$\begin{cases} \dot{M}_3 + 2M_3(A - BR^{-1}B^T M_2) + W = 0 \\ \dot{\underline{k}}_3 + (A^T - M_2 BR^{-1}B^T) \underline{k}_3 - M_3 BR^{-1}B^T \underline{k}_2 = 0 \\ \dot{\underline{h}}_3 + \underline{x}_d^T D(\underline{e}) \underline{x}_d + \underline{b}^T(\underline{e}) \underline{x}_d - 2\underline{k}_2^T BR^{-1}B^T \underline{k}_3 = 0 \end{cases} \quad (3.40)$$

Une simulation numérique de cette commande, nous permettra de tester la validité de ces approches.

Cette approche du problème conduit à une loi de commande sous-optimale de la forme :

$$\underline{u}(t) = -R^{-1}B^T \{ (M_2(t) + \eta M_3(t)) \underline{x} + \underline{k}_2(t) + \eta \underline{k}_3(t) \}$$

Le schéma bloc de simulation du fonctionnement du système commandé est représenté sur la figure III.2.

Détermination des valeurs finales

Comme on a vu au début de l'étude, la valeur finale de $V(\underline{x}, t)$ est donnée par :



FIGURE III-2 .- Schéma de simulation.

$$V(\underline{x}(t_f), t_f) = (\underline{x}_d - \underline{x}(t_f))^T S (\underline{x}_d - \underline{x}(t_f)) = \underline{x}_d^T S \underline{x}_d - 2 \underline{x}^T(t_f) S \underline{x}_d + \underline{x}^T(t_f) S \underline{x}(t_f)$$

mais on remarque aussi que :

$$\begin{aligned} V(\underline{x}(t_f), t_f) &= \sum_{n=2}^{\infty} \eta^{n-2} V_n(\underline{x}(t_f), t_f) = \\ &= h_2(t_f) + 2 \underline{x}^T(t_f) \underline{k}_2(t_f) + \underline{x}^T(t_f) M_2(t_f) \underline{x}(t_f) + \\ &+ \eta [h_3(t_f) + 2 \underline{x}^T(t_f) \underline{k}_3(t_f) + \underline{x}^T(t_f) M_3(t_f) \underline{x}(t_f)] + \dots \end{aligned}$$

$\underline{x}(t_f)$ devant être identique à $\underline{x}_d$, on en déduit :

$$\left. \begin{aligned} M_2(t_f) &= S \\ K_2(t_f) &= - S \underline{x}_d \\ h_2(t_f) &= \underline{x}_d^T S \underline{x}_d \end{aligned} \right\} \begin{aligned} M_3(t_f) &= \mathbb{O} \\ K_3(t_f) &= \underline{0} \\ h_3(t_f) &= 0 \end{aligned} \quad (3.41)$$

Ne connaissant que les valeurs finales, le calcul de $M_2(t)$, $K_2(t)$, $M_3(t)$, $K_3(t)$, ..., sera effectué hors ligne pour une mise en oeuvre ultérieure sur mini-calculateur.

Pour simplifier le calcul de la commande, Garrard [III.5_] a proposé de remplacer les matrices $M_k(t)$ par leurs moyennes en utilisant la formule :

$$\bar{M}_k = \frac{1}{t_f - t_0} \int_{t_0}^{t_f} M_k(t) dt \quad k = 2, 3, \dots \quad (3.42)$$

III.4. - DETERMINATION DES ELEMENTS DE L'ALGORITHME DE REGULATION

Un bilan thermodynamique des phénomènes de transfert de chaleur dans un four à diffusion (chapitre I) nous

a conduit à l'établissement d'un modèle mathématique de celui-ci composé de deux équations différentielles aux dérivées partielles couplées non linéaires (1.15).

L'objectif recherché étant la conduite numérique du four, nous avons identifié les paramètres inconnus à l'aide de la méthode Sphéra, puis nous avons discrétisé le modèle, figure II.11, pour pouvoir appliquer des techniques d'optimisation établies pour les systèmes à paramètres localisés.

La résolution du problème de poursuite ou d'atteinte d'un état désiré pour le système non linéaire

$$\begin{cases} \dot{\underline{x}}(t) = A \underline{x}(t) + \eta \int \underline{x}(t) + B \underline{u}(t) \\ \underline{y}(t) = C \underline{x}(t) \end{cases} \quad (3.43)$$

$$\underline{x}(0) = \underline{x}_0$$

avec minimisation d'un critère J , quadratique sur l'état et la commande augmenté d'une certaine fonction de l'état semi-définie positive, donné par l'expression :

$$J(\underline{x}_0, \underline{u}(t), \underline{x}_d, \eta) = \underline{e}^T(t_f) S \underline{e}(t_f) + \int_t^{t_f} [\underline{e}^T(t) Q \underline{e}(t) + \eta g(\underline{e}) + \underline{u}^T(t) R \underline{u}(t)] dt \quad (3.44)$$

$$\underline{e}(t) = \underline{x}_d - \underline{x}(t)$$

($\underline{x}_d$ est l'état désiré)

conduit à une loi de commande sous optimale en boucle fermé de la forme :

$$\underline{u}(t) = -R^{-1} B^T \left[(M_2(t) + \eta M_3(t)) \underline{x}(t) + \underline{k}_2(t) + \eta \underline{k}_3(t) \right] \quad (3.45)$$

où $M_2(t)$ est une matrice carrée $n \times n$, symétrique et définie

positive, solution de l'équation de Riccati :

$$\dot{M}_2(t) + 2M_2(t).A - M_2(t)BR^{-1}B^T M_2(t) + Q = 0 \quad (3.46)$$

$M_3(t)$ est une matrice carrée $n \times n$, symétrique, solution de l'équation de Lyapounov dynamique :

$$\dot{M}_3(t) = 2M_3(t)[A - BR^{-1}B^T M_2(t)] + W \quad (3.47)$$

et $\underline{k}_2(t)$ et $\underline{k}_3(t)$ sont deux vecteurs de dimension n solutions des équations :

$$\dot{\underline{k}}_2(t) = [M_2(t)BR^{-1}B^T - A^T] \underline{k}_2(t) + Q \cdot \underline{x}_d \quad (3.48)$$

$$\dot{\underline{k}}_3(t) = [M_2(t)BR^{-1}B^T - A^T] \underline{k}_3(t) + M_3(t)BR^{-1}B^T \underline{k}_2(t) \quad (3.49)$$

Les valeurs finales sont :

$$\begin{aligned} M_2(t_f) &= S & M_3(t_f) &= 0 \\ K_2(t_f) &= -S\underline{x}_d & K_3(t_f) &= 0 \end{aligned} \quad (3.50)$$

III.4.1. - Résolution de l'équation de Riccati

Quand le temps t_f tend vers l'infini, la matrice M_2 , solution de l'équation de Riccati, atteint une valeur constante. Il résulte que cette matrice est invariante et que $\dot{M}_2(t_f) = 0$. L'équation de Riccati est alors constituée par un système d'équations algébriques dont la résolution fournit la matrice des gains de retour d'état.

Dans le cas contraire, M_2 est une fonction du temps et sa résolution en ligne sur mini-ordinateur pose des problèmes difficiles à résoudre, si on tient compte de l'emplacement mémoire et du temps nécessaires pour effectuer les calculs.

Nous avons jugé intéressant dans un but de simplification

d'initialiser $M_2(t_f) = S$ à sa valeur obtenue en résolvant l'équation de Riccati algébrique [III.12].

Le choix des matrices de pondération

$$Q = 30 \times \mathbb{1} \quad \text{et} \quad R = \mathbb{1}$$

conduit à une matrice M_1 donnée par la figure III.3.

III.4.2. - Résolution de l'équation de Lyapounov

Soit :

$$\dot{M}_3(t) = 2M_3(t) \left[A - BR^{-1}B^T M_2(t) \right] + W \quad (3.51)$$

où $W = 2M_2 \cdot F(x) + D(e)$ est donné par (3.38)

$D(e)$ fait partie du critère à minimiser (3.26) et peut aussi être choisi tel que les calculs soient simplifiés ; par exemple tel que la matrice $W = 2M_2 \cdot F(x) + D(e)$ soit une matrice diagonale.

Le vecteur $\underline{k}_3(t)$ est une fonction de la matrice $M_3(t)$ elle même fonction de la matrice diagonale W . Des considérations pratiques nous permettent de choisir W de telle façon que les coefficients du vecteur $\underline{k}_3(t)$ prennent des valeurs satisfaisantes (ni trop grandes, ni trop petites). Par exemple si $W = \text{diag} [10^9 \ 10^9 \dots \dots \dots 10^9 \ 7]$

les coefficients de la matrice $M_3(t)$ calculés par leur moyenne temporelle sur un intervalle de temps fini, d'après (3.42), sont donnés par la figure III.4.

III.4.3. - Calcul des vecteurs $\underline{k}_2(t)$ et $\underline{k}_3(t)$

La résolution des équations (3.48) et (3.49) est effectuée simultanément à condition de définir le vecteur augmenté de dimension vingt-quatre :

$$\underline{k}(t) = \begin{bmatrix} \underline{k}_2(t) \\ \dots \\ \underline{k}_3(t) \end{bmatrix}$$

413,0	- 669,1	380,9	13,6	- 10,2	6,11	0,21	- 0,34	0,19	0,001	- 0,005	0,003
- 669,1	1 242,4	- 659,0	- 59,4	18,6	- 10,9	- 0,34	0,63	- 0,34	- 0,02	0,01	- 0,005
380,9	- 659,0	1801,5	- 1429,5	85,3	- 35,8	0,19	- 0,34	0,91	- 0,73	0,06	- 0,02
13,6	- 59,4	- 1429,5	2338,1	- 1327,3	777,0	0,007	- 0,02	- 0,72	1,19	- 0,69	0,41
- 10,2	18,6	85,3	- 1327,3	2244,1	- 1325,4	- 0,005	- 0,009	0,048	- 0,67	1,14	- 0,7
6,11	- 10,9	- 35,8	777,0	- 1325,4	838,1	0,003	- 0,005	- 0,02	0,39	- 0,67	0,44

0,21	- 0,34	0,19	0,007	- 0,005	0,003	$2,2 \times 10^{-3}$	$-1,7 \times 10^{-4}$	$1,02 \times 10^{-4}$	$1,03 \times 10^{-6}$	$-2,7 \times 10^{-6}$	$1,6 \times 10^{-6}$
- 0,34	0,63	- 0,34	- 0,02	- 0,009	- 0,005	$-1,7 \times 10^{-4}$	$2,4 \times 10^{-3}$	$-1,7 \times 10^{-4}$	$-7,2 \times 10^{-6}$	$4,8 \times 10^{-6}$	$-2,8 \times 10^{-6}$
0,19	- 0,34	0,91	- 0,72	0,048	- 0,02	$1,02 \times 10^{-4}$	$1,7 \times 10^{-4}$	$2,5 \times 10^{-3}$	$-3,7 \times 10^{-4}$	$3,3 \times 10^{-5}$	$-1,1 \times 10^{-5}$
0,001	- 0,02	- 0,73	1,19	- 0,67	0,39	$1,03 \times 10^{-6}$	$-7,2 \times 10^{-4}$	$-3,7 \times 10^{-4}$	$2,7 \times 10^{-3}$	$-3,5 \times 10^{-4}$	$2,09 \times 10^{-4}$
- 0,005	0,01	0,06	- 0,69	1,14	- 0,67	$-2,7 \times 10^{-6}$	$4,8 \times 10^{-6}$	$3,3 \times 10^{-4}$	$-3,5 \times 10^{-4}$	$2,7 \times 10^{-3}$	$-3,5 \times 10^{-4}$
0,003	- 0,005	- 0,02	0,41	- 0,7	0,44	$1,6 \times 10^{-6}$	$-2,8 \times 10^{-6}$	$-1,1 \times 10^{-5}$	$2,09 \times 10^{-4}$	$-3,5 \times 10^{-4}$	$2,3 \times 10^{-3}$

FIGURE III.3 - SOLUTION DE L'EQUATION DE RICCATI

155	-140,9	9,85	49,4	5,48	- 3,51	0,08	0,07	0,006	0,02	0,002	- 0,002
-140,9	257,8	- 18,9	- 58,5	- 18,4	10,3	- 0,07	0,13	- 0,01	- 0,03	- 0,008	- 0,005
9,85	- 18,9	260,1	- 75,7	- 99,6	42,2	0,005	- 0,01	0,13	- 0,04	- 0,04	0,02
49,4	- 58,5	- 75,7	233,6	- 85,5	35,9	0,02	- 0,03	- 0,04	0,12	- 0,04	0,02
5,48	- 18,4	- 99,6	- 85,5	407,3	-156,2	0,003	- 0,009	- 0,05	- 0,04	0,21	- 0,08
- 3,51	10,3	42,2	35,9	-156,2	-121,5	- 0,002	0,005	0,02	0,018	- 0,08	0,07
0,08	- 0,07	0,005	0,02	0,003	- 0,002	$7,1 \times 10^{-3}$	$-4,3 \times 10^{-5}$	$3,4 \times 10^{-5}$	$1,2 \times 10^{-5}$	$1,1 \times 10^{-6}$	$9,6 \times 10^{-7}$
- 0,07	0,13	- 0,01	- 0,03	-0,009	0,005	$-4,3 \times 10^{-5}$	$7,1 \times 10^{-3}$	7×10^{-6}	$1,4 \times 10^{-5}$	$- 4 \times 10^{-6}$	$2,6 \times 10^{-6}$
0,006	- 0,01	0,13	- 0,04	- 0,05	- 0,02	$3,4 \times 10^{-6}$	$- 7 \times 10^{-6}$	$7,1 \times 10^{-3}$	-2×10^{-5}	$-2,4 \times 10^{-5}$	$1,1 \times 10^{-5}$
0,02	- 0,03	- 0,04	0,12	- 0,04	- 0,018	$1,2 \times 10^{-5}$	$1,4 \times 10^{-5}$	$- 2 \times 10^{-5}$	$7,1 \times 10^{-3}$	$-2,3 \times 10^{-5}$	$9,7 \times 10^{-6}$
0,002	- 0,008	- 0,04	- 0,04	0,21	- 0,08	$1,1 \times 10^{-6}$	$- 4 \times 10^{-6}$	$-2,4 \times 10^{-5}$	$-2,4 \times 10^{-5}$	$7,1 \times 10^{-3}$	$-3,6 \times 10^{-5}$
- 0,002	- 0,005	0,02	0,02	- 0,08	0,07	$9,6 \times 10^{-7}$	$2,6 \times 10^{-6}$	$1,1 \times 10^{-5}$	$9,7 \times 10^{-6}$	$-3,6 \times 10^{-5}$	$7,1 \times 10^{-3}$

FIGURE III . 4 - SOLUTION DE L'EQUATION DE LYAPOUNOV

De cette façon on aboutit au système d'équations différentielles :

$$\dot{\underline{k}}(t) = \begin{pmatrix} \dot{\underline{k}}_2(t) \\ \dot{\underline{k}}_3(t) \end{pmatrix} = \begin{pmatrix} M_2 B R^{-1} B^T - A^T & \textcircled{0} \\ M_3 B R^{-1} B^T & M_2 B R^{-1} B^T - A^T \end{pmatrix} \begin{pmatrix} \underline{k}_2(t) \\ \underline{k}_3(t) \end{pmatrix} + \begin{pmatrix} Q \underline{x}_d \\ \underline{0} \end{pmatrix} \quad (3.52)$$

avec la condition finale :

$$\underline{k}(t_f) = \begin{pmatrix} \underline{k}_2(t_f) \\ \underline{k}_3(t_f) \end{pmatrix} = \begin{pmatrix} -F \underline{x}_d \\ \underline{0} \end{pmatrix}$$

La solution a été obtenue à partir de la méthode prédictive-corrective d'Hamming modifiée [III.13]. Il s'agit d'une procédure d'intégration du quatrième ordre, qui exige l'évaluation du membre de droite de l'équation matricielle (3.52) seulement deux fois à chaque pas d'itération, contrairement à d'autres méthodes du même ordre de précision, comme la méthode de Runge-Kutta, qui exigent le calcul de cette quantité quatre fois à chaque pas. Un autre avantage est qu'à chaque pas on obtient une estimation de l'erreur et, en fonction de celle-ci, on peut modifier ce pas d'intégration.

Cependant, la méthode d'Hamming ne peut pas être initialisée d'elle-même, car il est nécessaire de connaître les valeurs de la fonction et de ses dérivées en quatre points équidistants au pas précédent. Le calcul de ces points a été possible grâce à l'utilisation d'une méthode de Runge-Kutta modifiée par Ralston [III.14].

L'intégration a été effectuée sur un horizon

temporel de trois heures.

Les courbes des figures III.5 et III.6 représentent respectivement la variation temporelle des coefficients $-R^{-1}B^T \underline{k}_2(t)$ ($\eta = 0$) et $-R^{-1}B^T [\underline{k}_2(t) + \eta \underline{k}_3(t)]$ ($\eta \neq 0$) intervenant dans le calcul de la loi de commande (3.45).

III.4.4 - Commentaires sur les résultats obtenus

A partir du modèle dynamique du système (1.15) et de la loi de commande sous-optimale (3.45)

$$\underline{u}(t) = -R^{-1}B^T [\underline{M}_2(t) + \underline{k}_2(t) + \eta (\underline{M}_3(t) + \underline{k}_3(t))]$$

nous avons réalisé une série de simulations de fonctionnement du système commandé.

Ces expériences ont été effectuées dans un premier temps pour $\eta = 0$. Les résultats sont donnés sur les figures III.7 et III.8.

On remarque que les variations de la commande tendent assez bien vers zéro et permettent de ramener l'état du système à la valeur désirée dans l'intervalle de temps fixé.

Les résultats que nous avons obtenus pour $\eta \neq 0$ sont donnés sur les figures III.9 et III.10.

Les performances du système de régulation restent sensiblement les mêmes, pour un certain choix de pondération au niveau de la matrice W . Si on augmente les valeurs des éléments de cette matrice on obtient une commande d'amplitude maximale plus faible et répartie de façon plus égale sur l'intervalle de temps considéré.

CONCLUSION

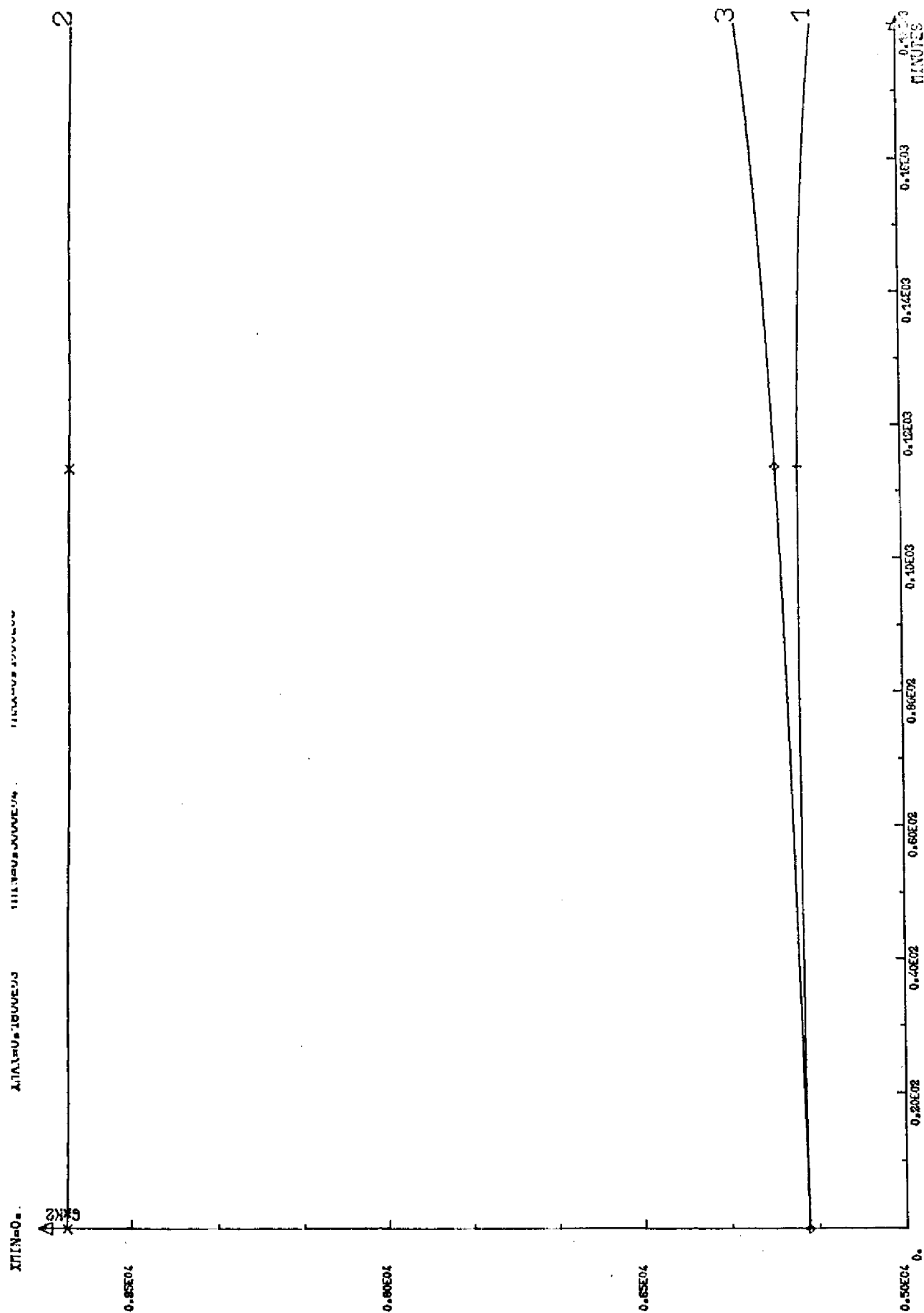
Dans ce chapitre nous avons développé un algorithme de commande pour le problème de l'atteinte d'un état

fixe pour un système non linéaire.

Nous avons choisi un critère quadratique classique augmenté d'un terme non linéaire dans le but de simplifier les calculs. L'addition de ce terme revient à considérer une contrainte supplémentaire sur l'état du système.

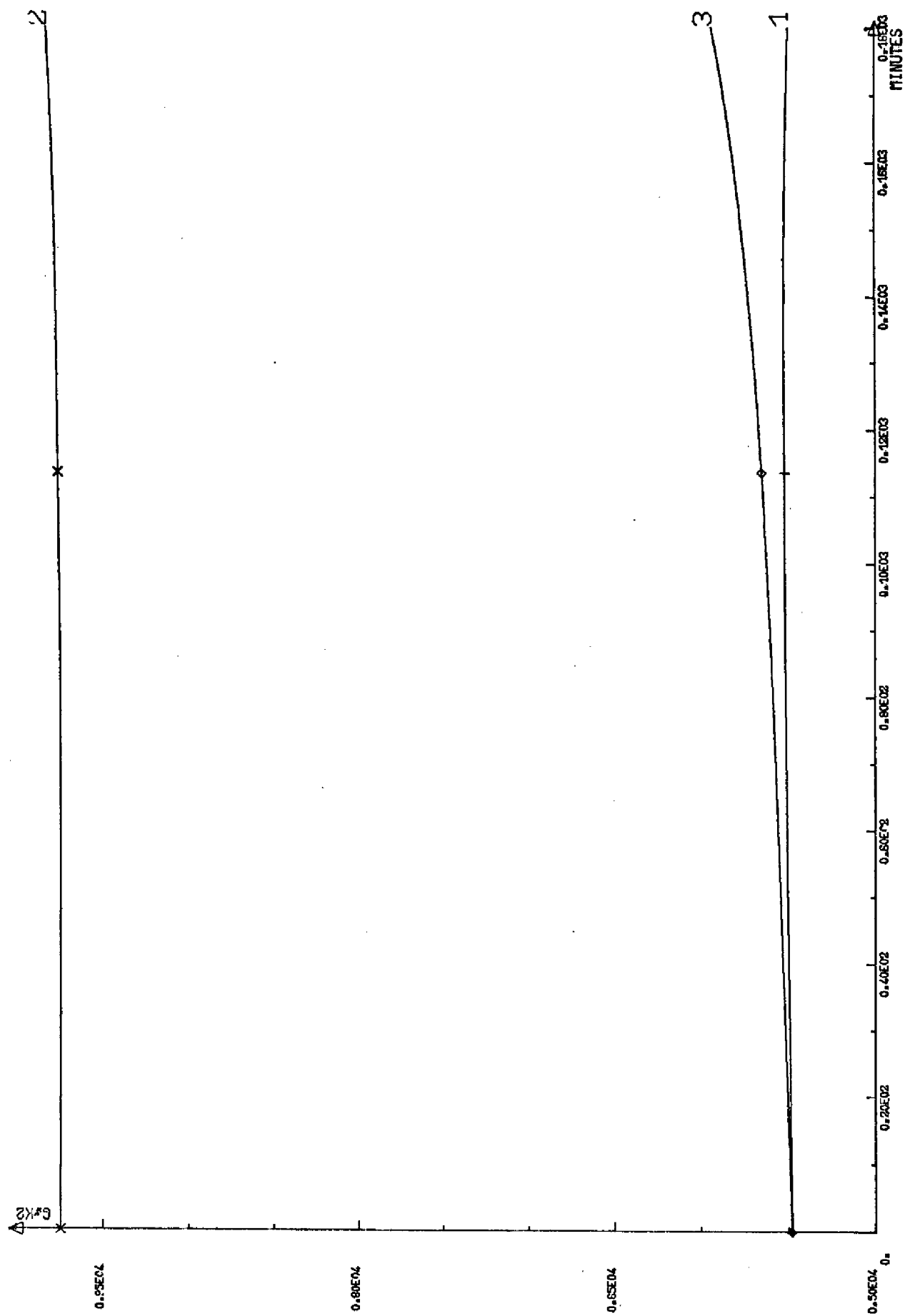
La résolution de l'équation d'Hamilton-Jacobi-Bellman pour ce problème est simplifiée d'une façon notable au moyen d'un développement en série. Ne pouvant pas tenir compte pratiquement de tous les termes, nous nous sommes arrêtés au deuxième ordre.

Les résultats que nous avons obtenus par simulation nous ont permis de vérifier la validité de cet algorithme et du choix des approximations effectuées.



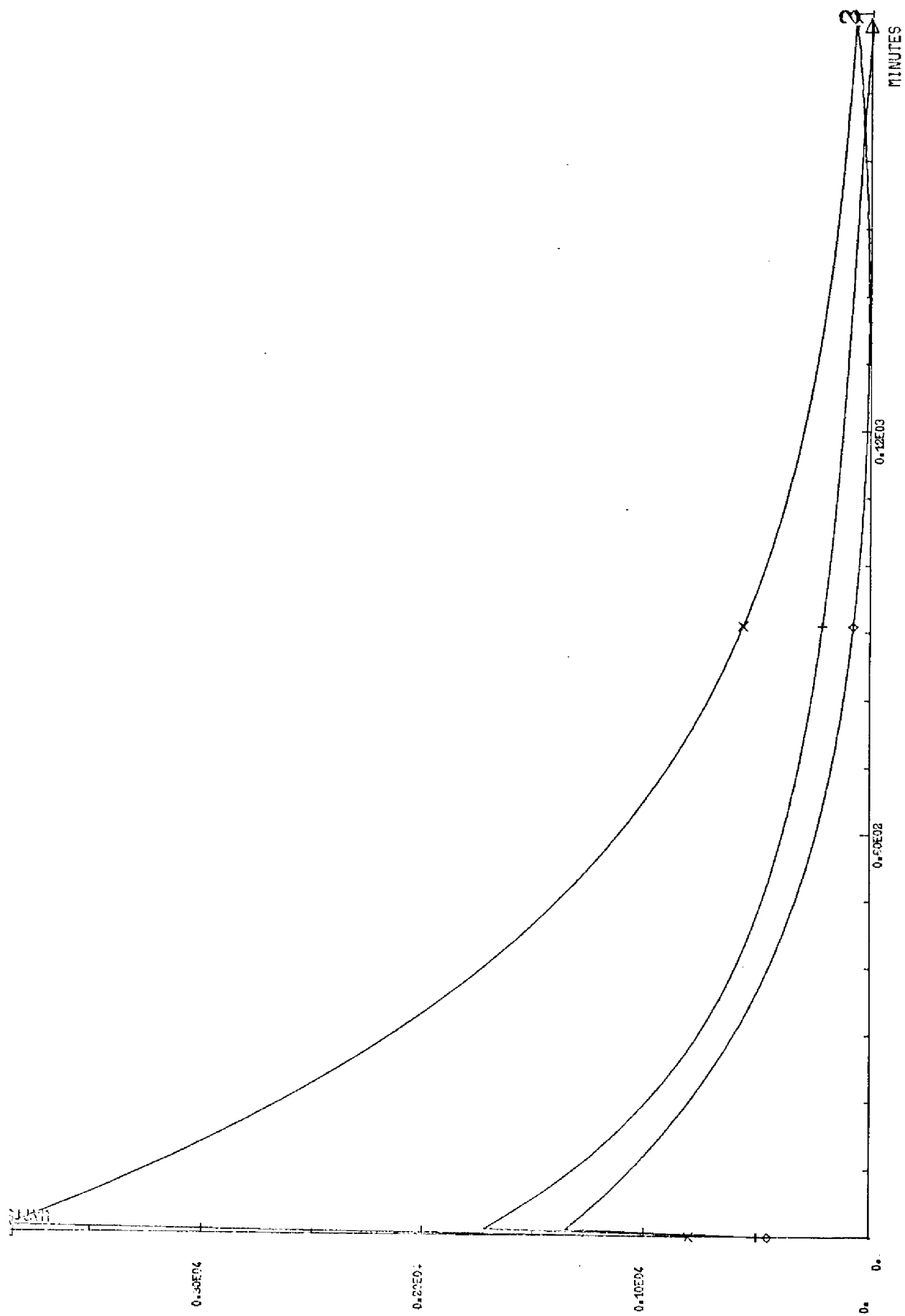
EVOLUTION TEMPORELLE DE G*K2

Figure III.5



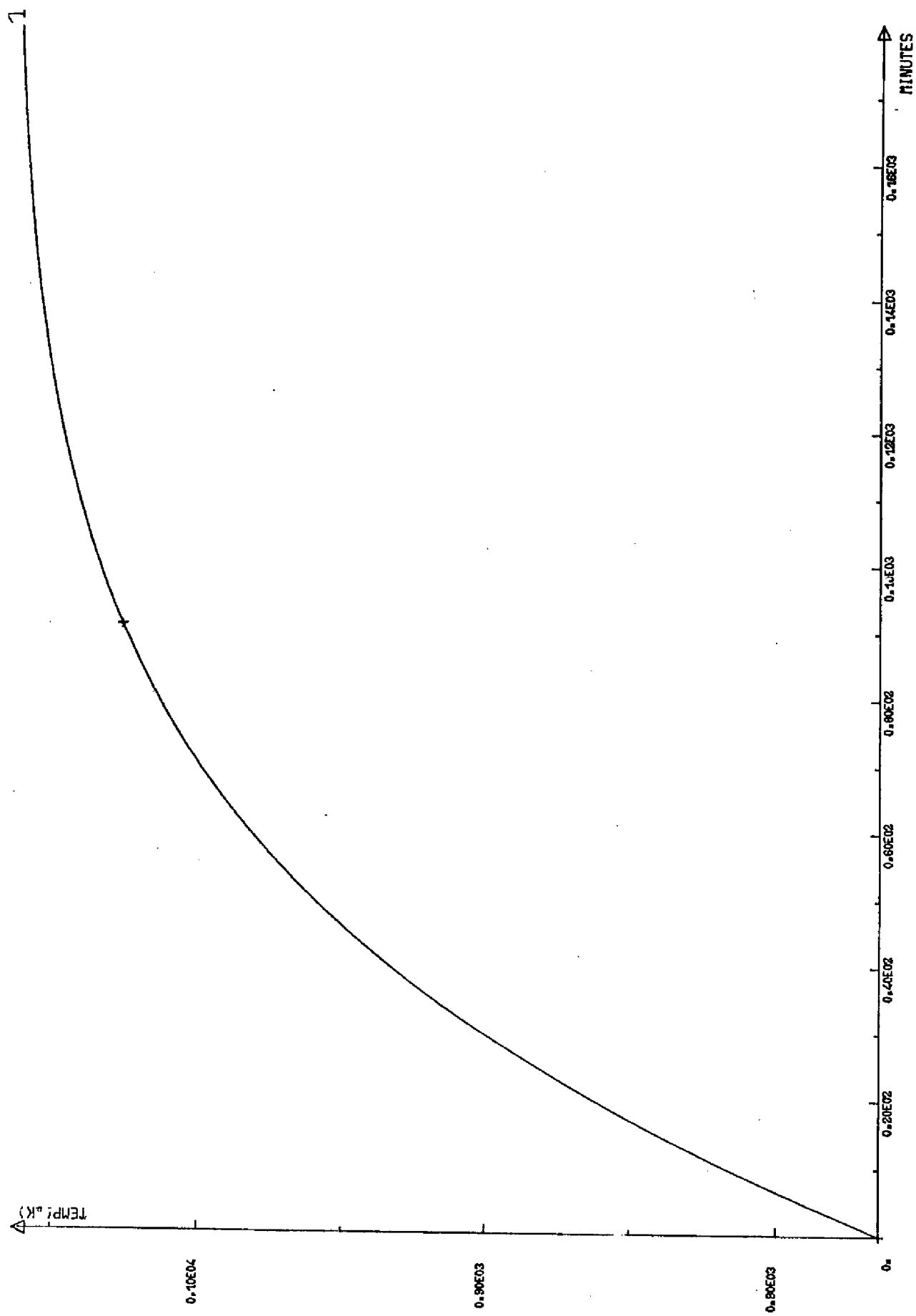
EVOLUTION TEMPORELLE DE
 $G^*(K_2 + \eta * K_3)$

Figure III.6



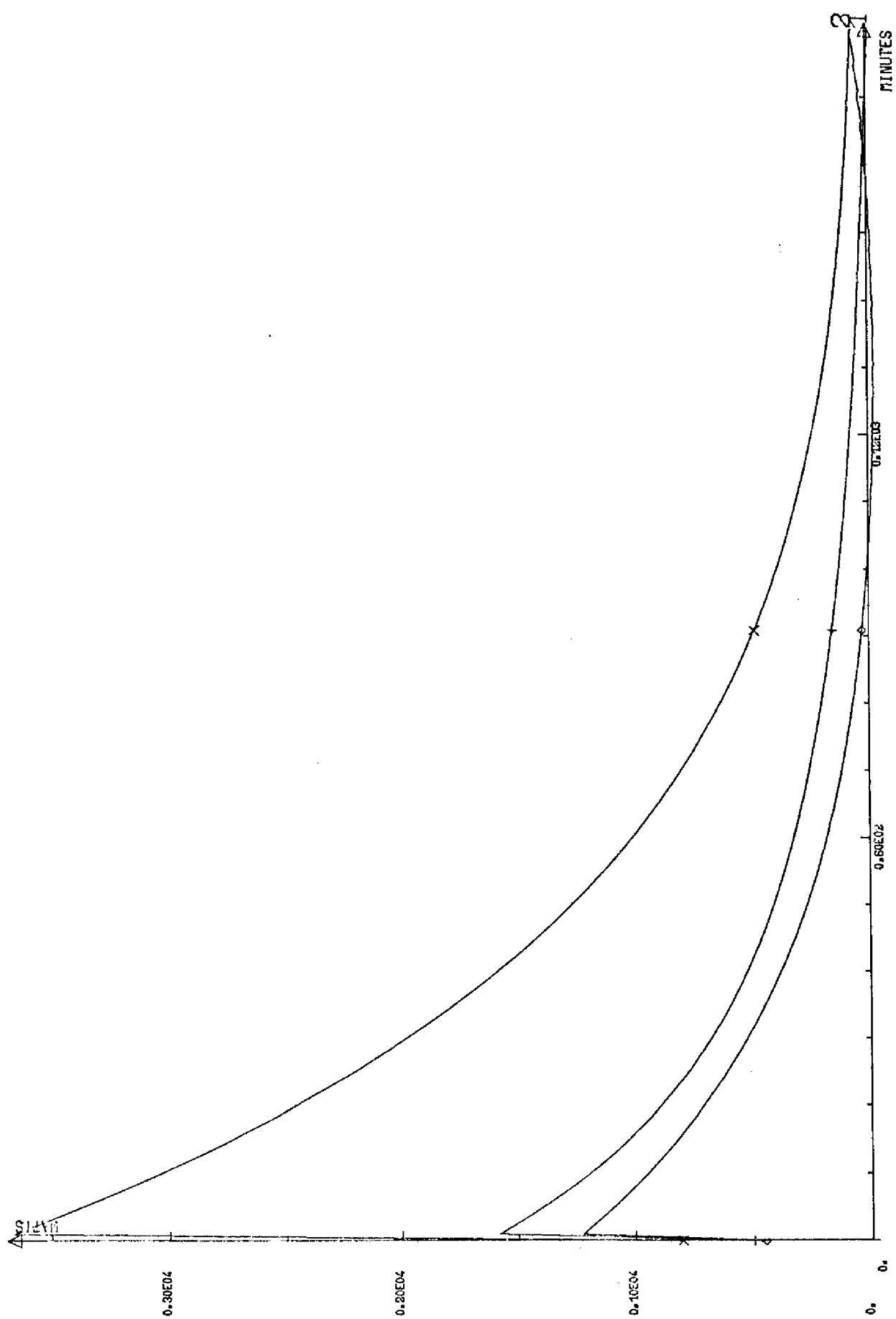
EVOLUTION DES PUISSANCES
APPROX. DU PREMIER ORDRE

Figure III.7



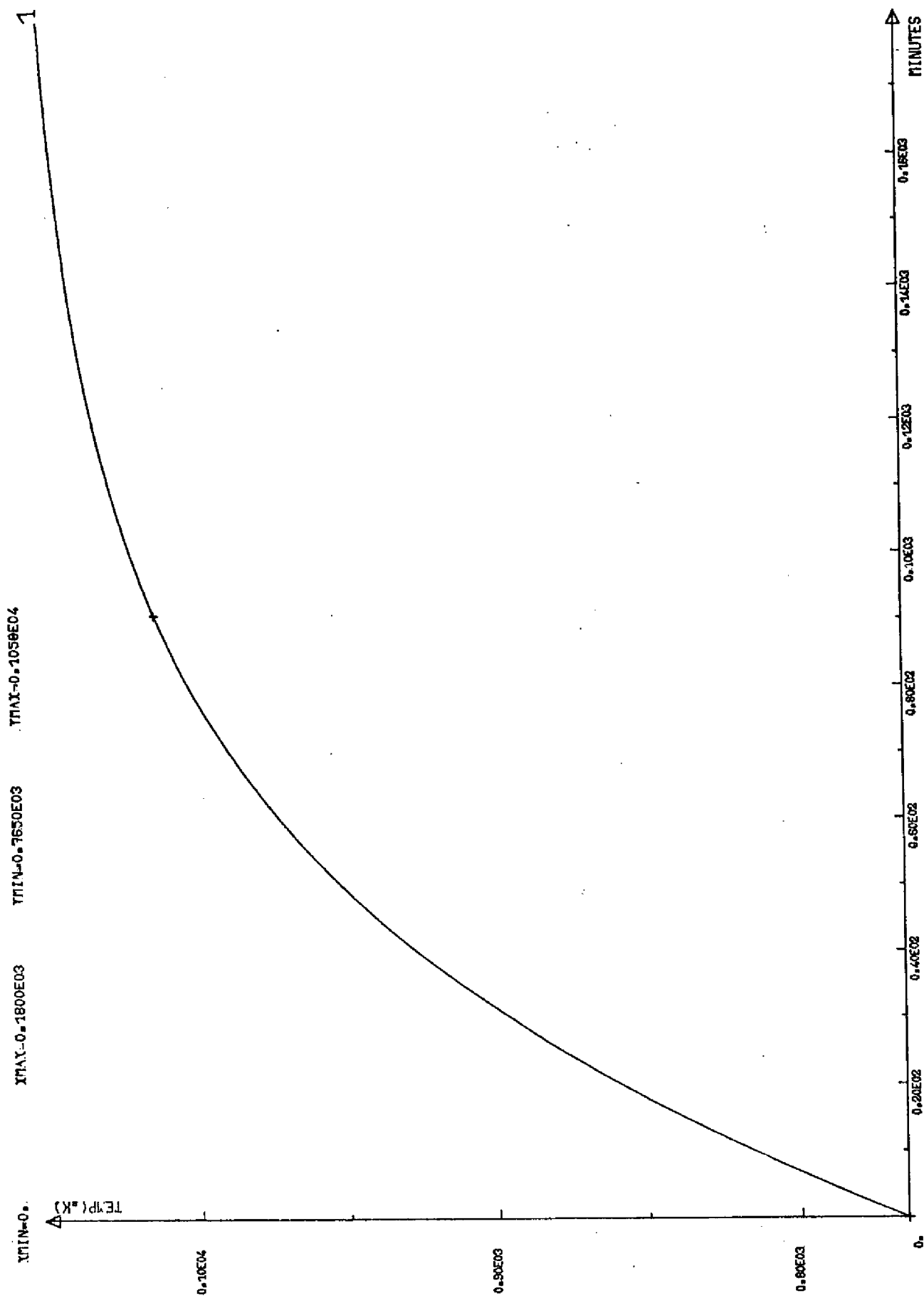
REGIME TRANSITOIRE
THERMOCOUPLE NO 3

Figure III.8



EVOLUTION DES PUISSANCES
APPROX. DU SECOND ORDRE

Figure III.9



REGIME TRANSITOIRE

THERMOCOUPLE N° 2

Figure III.10

B I B L I O G R A P H I E

- $\angle$ III.1_7 M.ATHANS, P.L. FALB
 Optimal Control
 Mc Graw-Hill Book Company (1966)
- $\angle$ III.2_7 P. SANNUTI, P. KOKOTOVIC
 Singular Perturbation Method for Near Op-
 timum
 Design of High-Order Non-Linear Systems
 Automatica, Vol.5, pp. 773-779 (1969)
- $\angle$ III.3_7 Y. NISHIKAWA, N. SANNOMIYA, H. ITAKURA
 A Method for Suboptimal Design of Non-
 Linear
 Feedback Systems
 Automatica, Vol.7, pp. 703-712 (1972)
- $\angle$ III.4_7 J.D. PEARSON
 Approximation Methods in Optimal Control
 J.Electron. Control 13, pp. 453-469 (1962)
- $\angle$ III.5_7 W.L. GARRARD, N.H. Mc. CLAMROCH, L.G. CLARK
 An Approach to Sub-optimal Feedback Control
 of Non-Linear Systems
 International Journal of Control, Vol.5,
 n° 5, pp. 425-435 (1967)
- $\angle$ III.6_7 W.L. GARRARD
 Additional results on sub-optimal feed-
 back control of non-linear systems
 International Journal of Control
 Vol.10, n°6, pp.657-663 (1969)

- [~III.7_] W.L. GARRARD
Sub-optimal Feedback Control for Non-Linear
Systems
Automatica, Vol.8, pp. 219-221 (1972)
- [~III.8_] J.H. BURGHART
A Technique for Sub-optimal Feedback Control
of Non-Linear Systems
I.E.E.E. Transactions on Automatic Control
pp. 530 -533 (1969)
- [~III.9_] R.E. KALMAN
Contributions to the Theory of Optimal
Control
Boln. Soc. Mat. Mex., Vol.5, pp. 102-119
(1960)
- [~III.10_] R. BELLMAN
Introduction to Matrix Analysis
Mc Graw-Hill (1960)
- [~III.11_] M. AMOUROUX, J.P. BABARY, Y.QUEMENER,
X. RUIZ del PORTAL
Contrôle par ordinateur numérique de la
température d'un four d'élaboration de
matériaux et composants électroniques
3° Congreso Nacional de Informatica y
Automatica Madrid 29/31 octobre 1975
- [~III.12_] M. COURDESSES, G. ALENGRIN
Programmes de résolution de l'équation al-
gébrique matricielle de Riccati
Note Interne L.A.A.S. 75.I.34 Novembre 1975

- [III.13_] A. RALSTON, H.S. WILF
Mathematical Methods for Digital Computers
Wiley (1960) pp. 95-109
- [III.14_] A. RALSTON
Runge-Kutta Methods with Minimum Errors
Bounds
M.T.A.C. Vol.16, n° 80 (1962) pp. 431-437
-

C O N C L U S I O N

L'objet des études menées dans le cadre de cette thèse concernait l'obtention d'un modèle dynamique du fonctionnement d'un four à diffusion, dans le but de concevoir une loi de commande sous-optimale en boucle fermée du profil thermique.

Compte tenu de quelques hypothèses simplificatrices réalistes, nous avons modélisé le comportement du four sous la forme d'un système de deux équations différentielles aux dérivées partielles non linéaires, permettant de connaître la température en chaque point du four, à chaque instant, en fonction de la répartition de la puissance de chauffe le long des trois zones de chauffe.

Le modèle tient compte de paramètres susceptibles de varier, en fonction des expériences, tels que la nature des gaz, sa vitesse de propagation, les caractéristiques géométriques et thermodynamiques du tube de quartz ...

Certains coefficients de ce modèle, étant inconnus, ont été identifiés à partir de mesures de température effectuées en six points du four, alors que la puissance de chauffe était constante, en utilisant la méthode heuristique d'identification Sphéra, qui est une extension de la méthode de Fibonacci.

Une loi de commande sous-optimale, basée sur l'approximation de la solution d'Hamilton-Jacobi-Bellman, a été obtenue à partir de la considération du modèle discrétisé et d'un critère de type quadratique modifié sur la commande et l'écart entre l'état réel et l'état désiré du système.

La loi de commande en boucle fermée obtenue est relativement simple à mettre en oeuvre sur un calculateur numérique temps réel.

Les simulations numériques effectuées ont montré l'efficacité de la méthode proposée.

Le type d'approche que nous avons utilisé est facilement applicable à des problèmes non linéaires de type polynomial. Dans le cas où l'état désiré varie en fonction du temps, la méthode présentée dans ce mémoire est facilement généralisable.

T A B L E D E S M A T I E R E S

SYMBOLES UTILISES

INTRODUCTION

CHAPITRE 1 :

Introduction	9
I.1. - Description du processus	10
I.2. - Modélisation	11
I.2.1. - Modélisation de la partie interne du four	11
I.2.2. - Etablissement d'un système aux dérivées partielles	16
I.2.3. - Modèle global	17
I.2.4. - Modèle à paramètres localisés	21
Conclusion	23
Bibliographie	25

CHAPITRE II :

Introduction	29
II.1. - Méthode d'identification	30
II.1.1. - Introduction	30
II.1.2. - Méthode de Fibonacci	31
II.1.2.1. - Méthode de Recherche	34
II.1.2.2. - Mise en oeuvre de la méthode	36
II.1.2.3. - Détermination du nom- bre de mesures	37
II.1.2.4. - Détermination des points de mesure	38
II.1.2.5. - Organigramme	39
II.1.3. - Méthode Sphéra	42
II.2. - Discrétisation totale du modèle	45

II.3. - Application de la méthode Sphéra et résultats	50
Conclusion	57
Bibliographie	58

CHAPITRE III :

Introduction	63
III.1. - Position du problème	63
III.2. - Le problème de poursuite pour un système linéaire	64
III.3. - Obtention d'une loi de commande en boucle fermée sous-optimale	68
III.4. - Détermination des éléments de l'algorithme de régulation	78
III.4.1. - Résolution de l'équation de Riccati	80
III.4.2. - Résolution de l'équation de Lyapounov	81
III.4.3. - Calcul des vecteurs $\underline{k}_2(t)$ et $\underline{k}_3(t)$	81
III.4.4. - Commentaires sur les résultats obtenus	85
Conclusion	85
Bibliographie	93

CONCLUSION