

Support de cours officiel du cours de *Modélisation : Représentations et analyse des modèles* de première année de CentraleSupélec écrit par les personnes suivantes, mentionnées ici en ordre alphabétique : *Thomas Chevet* (Département Automatique, Laboratoire des Signaux et Systèmes), *Marie-Anne Lefebvre* (Institut d'Électronique et de Télécommunications de Rennes), *Véronique Letort - Le Chevalier* (Département de Mathématiques, Laboratoire MICS), *Cristina Maniu* (Département Automatique, Laboratoire des Signaux et Systèmes), *Guillaume Sandou* (Département Automatique, Laboratoire des Signaux et Systèmes) et *Cristina Vlad* (Département Automatique, Laboratoire des Signaux et Systèmes). Merci de signaler toute erreur que vous constateriez à modelisation.cours@centralesupelec.fr.

Deuxième édition, 20 novembre 2019

Remerciements

La première partie de ce document a comme base de départ entre autres le polycopié du cours de *Signaux et Systèmes 2* (Antoine et al., 2014) du *cursus Ingénieur Supélec* et reprend une grande partie du texte d'origine pour sa première partie. Ainsi, nous tenons à remercier nos collègues Jacques Antoine, Jean-Luc Collette, Gilles Duc, Hervé Guéguen et Daniel Poulton pour ce partage.

La modélisation par réseaux de Petri a été largement inspirée du polycopié du cours de *Grafcet et réseaux de Petri* (Chlique, 2000), *cursus Ingénieur Supélec*. Nous remercions Pierre Chlique pour ce partage.

Nous remercions Gilles Duc et Anne-Laure Castelier pour la relecture et leurs suggestions pour l'amélioration du document.

Notations

a	Notation pour un nombre scalaire
$\mathbf{a}$	Notation pour un vecteur
$\mathbf{A}$	Notation pour une matrice
$\mathbf{A}^T$	Transposée de la matrice $\mathbf{A}$
$\text{Tr}(\mathbf{A})$	Trace de la matrice $\mathbf{A}$
$\det(\mathbf{A})$	Déterminant de la matrice $\mathbf{A}$
$\ \mathbf{a}\ $	Norme euclidienne du vecteur $\mathbf{a}$
$\mathbb{N}$	Ensemble des nombres entiers naturels
$\mathbb{Z}$	Ensemble des nombres entiers
$\mathbb{R}$	Ensemble des nombres réels
$\mathbb{R}_+$	Ensemble des nombres réels positifs
$\mathbb{R}_+^*$	Ensemble des nombres réels strictement positifs
$L^1(\mathbb{R})$	Espace des fonctions sommables
$\mathcal{D}(I)$	Ensemble des fonctions à support compact de classe C^∞ dans I
$\text{Re}(a)$	Partie réelle du nombre complexe a
$\text{Im}(a)$	Partie imaginaire du nombre complexe a
$\text{Pf}(a)$	Partie finie de a
VP	Valeur principale de Cauchy
$X(f) = \mathcal{F}\{x(t)\}(f)$	Transformée de Fourier du signal $x(t)$
p	Variable de Laplace
$X(p) = \mathcal{L}\{x(t)\}(p)$	Transformée de Laplace du signal $x(t)$
$H(p)$	Fonction de transfert d'un système linéaire invariant dans le temps, à temps continu
$H(z)$	Fonction de transfert d'un système linéaire invariant dans le temps, à temps discret
$p(n) = \mathbb{P}[X = n]$	Fonction de masse de la variable aléatoire X discrète
$\mathbb{P}[X < x]$	Probabilité que la variable aléatoire X soit inférieure à x
$\mathbb{P}[X Y]$	Probabilité conditionnelle de X sachant Y
$\mathbb{E}[X]$	Espérance de la variable aléatoire X
$\text{Var}[X]$	Variance de la variable aléatoire X
$\text{Cov}[X, Y]$	Covariance des variables aléatoires X et Y
$X \sim \mathcal{N}(\mu, \sigma^2)$	Variable aléatoire X gaussienne (ou normale), avec espérance μ et variance σ^2
δ	Impulsion de Dirac
$\mathbb{1}$	Echelon unitaire
$x * y$	Produit de convolution des signaux x et y
$R = (P, T, \text{Arc})$	Réseau de Petri avec les ensembles des places $P = \{P_1, \dots, P_{n_P}\}$, des transitions $T = \{tr_1, \dots, tr_{n_T}\}$ et des arcs orientés Arc
$\mathcal{A}$	Notation générale pour un automate
θ	Vecteur des paramètres
$\hat{\theta}$	Vecteur des paramètres estimés à partir de données expérimentales

Table des matières

Remerciements	iii
Notations	v
1 Introduction	1
1.1 Objectifs et liens avec d'autres cours	1
1.2 Organisation du polycopié	2
2 Du système à la formalisation mathématique	5
2.1 Introduction	5
2.1.1 Notion de système	5
2.1.2 Modèles et Modélisation	6
2.1.3 Représentations des systèmes : différentes approches	7
2.2 Représentations par schéma blocs	9
2.3 Représentation par variables d'état	11
2.4 Modélisation à événements discrets	13
2.5 Classification des systèmes	14
2.5.1 Nombre d'entrées/sorties	14
2.5.2 Critère temporel	15
2.5.3 Linéarité	15
2.5.4 Invariance dans le temps	16
2.5.5 Causalité	16
2.5.6 Évolution au cours du temps	16
2.6 Modèles de signaux certains	16
2.6.1 Echelon unité ou échelon de Heaviside	17
2.6.2 Impulsion de Dirac	17
2.6.3 Signal sinusoïdal	17
I Modélisation des systèmes à état continu au sens d'un système dynamique	19
3 Modélisation sous forme de représentation d'état	21
3.1 Définition	21
3.1.1 Équation d'état et équation de sortie	22
3.2 Stabilité d'un point d'équilibre ou d'une trajectoire d'équilibre	24
3.3 Linéarisation autour d'un point d'équilibre	26
3.3.1 Théorème de superposition	26
3.3.2 Linéarisation	27
3.4 Résolution des équations d'état d'un système linéaire	29
3.4.1 Régime libre	29
3.4.2 Régime forcé	29
3.4.3 Cas particulier des systèmes linéaires et invariants	30
3.4.4 Méthodes de calcul de l'exponentielle matricielle	30
3.5 Stabilité asymptotique	33

3.6	Analyse des systèmes discrets	34
3.6.1	Modélisation des systèmes linéaires discrets par variables d'état	34
3.6.2	Résolution des équations d'état	34
3.6.3	Stabilité asymptotique	35
4	Modélisation sous forme de fonction de transfert	37
4.1	Introduction	37
4.2	Notions préliminaires : transformation de Fourier, de Laplace, en z	37
4.2.1	Notions de base sur la transformation de Fourier	37
4.2.2	Transformation de Laplace	39
4.2.3	Transformation en z	43
4.3	Représentation temporelle des systèmes	46
4.3.1	Réponses impulsionnelle et indicielle des systèmes linéaires invariants à temps discret	46
4.3.2	Relation de convolution dans les systèmes linéaires invariants	47
4.4	Etude des systèmes linéaires et invariants par fonction de transfert	48
4.4.1	Fonction de transfert	48
4.4.2	Stabilité EBSB	49
4.4.3	Réponse en fréquence	51
4.4.4	Passage état – transfert. Formes compagnons	53
5	Analyse des performances des systèmes dynamiques	59
5.1	Introduction	59
5.2	Diagramme de Bode des systèmes à temps continu	59
5.2.1	Définition	59
5.2.2	Systèmes du 1 ^{er} ordre	60
5.2.3	Systèmes du 2 ^{ème} ordre	61
5.2.4	Méthode pratique de tracé asymptotique	62
5.3	Diagramme de Bode des systèmes numériques	64
5.3.1	1 ^{er} ordre tout pôle	64
5.3.2	2 ^{ème} ordre tout pôle	65
5.4	Diagramme de Bode des systèmes passe-tout, à déphasage minimal ou avec retard pur	66
5.4.1	Systèmes passe-tout	66
5.4.2	Systèmes à déphasage minimal	67
5.4.3	Systèmes avec retard pur	68
5.5	Analyse temporelle	69
5.5.1	Systèmes analogiques	69
5.6	Conclusion de la partie I	73
II	Modélisation des systèmes à état discret	75
6	Introduction sur la modélisation des systèmes à état discret	77
6.1	Problématique	77
6.2	Motivation et structure de la partie II	77
6.3	Définition des systèmes à événements discrets	80
7	Modélisation des systèmes à événements discrets non temporisés	81
7.1	Modélisation par automate	81
7.1.1	Définition et exemple d'automate	81
7.1.2	Propriétés des automates	83
7.1.3	Composition des automates	85
7.1.4	Analyse d'un système à événements discrets non temporisé représenté par un automate	88
7.1.5	Utilisation de la notion de langage par les automates	89
7.2	Modélisation par réseau de Petri non temporisés	90

7.2.1	Réseaux de Petri autonomes	91
7.2.2	Propriétés des modèles représentés par réseaux de Petri	95
7.2.3	Analyse par réduction	97
7.2.4	Analyse linéaire	99
7.2.5	Réseaux non autonomes synchronisés	102
8	Modélisation des systèmes à évènements discrets temporisés	103
8.1	Introduction	103
8.2	Structure d'horloge	103
8.3	Automates temporisés	105
8.3.1	Définition et simulation d'un automate temporisé	105
8.4	Automates temporisés avec gardes	106
8.4.1	Introduction	106
8.4.2	Ajout d'invariants	108
8.4.3	Formalisation	108
8.4.4	Parallélisation d'automates temporisés avec gardes	109
8.5	Systèmes hybrides	110
8.5.1	Présentation informelle	110
8.5.2	Description formelle	111
8.6	Réseaux de Petri temporisés	112
8.6.1	Définition	113
9	Automates stochastiques	115
9.1	Introduction	115
9.2	Structure d'horloge stochastique	115
9.3	Automate temporisé stochastique	116
9.4	Processus de Poisson	116
9.4.1	Définition de la densité de probabilité	116
9.4.2	Moyenne et variance du processus de Poisson	117
9.4.3	Propriétés du processus de Poisson	118
9.4.4	Intérêt du processus de Poisson	119
9.5	Autres types de processus	119
9.5.1	Loi d'Erlang ou processus de renouvellement	119
9.5.2	Distribution de Rayleigh	119
9.6	Conclusion de la partie II	120
III	Méthodes pour l'identification et l'évaluation des modèles	121
10	Identification	123
10.1	Introduction	123
10.1.1	Démarche générale pour l'identification	123
10.1.2	Classification sommaire des méthodes	123
10.2	Analyse indicielle	124
10.2.1	Procédure expérimentale	124
10.2.2	Identification par une fonction de transfert du premier ordre	125
10.2.3	Identification par la méthode de Strejé	126
10.2.4	Identification par une fonction de transfert du second ordre	128
10.3	Analyse harmonique	129
10.3.1	Procédure expérimentale	129
10.4	Méthode générale statistique	130
10.4.1	Principe général de la méthode statistique d'identification de modèle	130
10.4.2	Estimation paramétrique par la méthode des moindres carrés	131
10.4.3	Interprétation statistique	132
10.4.4	Notions d'optimisation	136
10.5	Évaluation et sélection de modèles	139

10.5.1 Quelques critères d'évaluation	140
10.5.2 Sélection de modèle	142
11 Analyse d'incertitude et de sensibilité	147
11.1 Contexte et objectifs	147
11.1.1 Notations	147
11.1.2 Analyse d'incertitude : définition et objectifs	147
11.1.3 Analyse de sensibilité : définition et objectifs	148
11.2 Modélisation de l'incertitude	149
11.2.1 Terminologie	149
11.2.2 Sources d'incertitude sur les facteurs	150
11.2.3 Définition générale de l'incertitude	151
11.2.4 Modélisation des incertitudes de type A	152
11.2.5 Modélisation des incertitudes de type B	157
11.3 Méthodes de propagation des incertitudes	160
11.3.1 Méthode des intervalles ou théorie des erreurs maximales ou approche min-max	161
11.3.2 Approximation de la variance par méthode de Taylor : méthode dite de la combinaison des variances	162
11.3.3 Propagation des incertitudes par méthode de Monte-Carlo	164
11.4 Analyse de sensibilité	166
11.4.1 Méthodes locales	166
11.4.2 Méthode d'analyse de sensibilité globale pour un modèle linéaire ou quasi- linéaire : indices SRC	167
11.4.3 Méthode d'analyse de sensibilité globale pour un modèle quelconque : indices de Sobol	169
11.5 Conclusion de la partie III	172
Annexe A Transformée de Laplace de signaux usuels	173
Annexe B Transformée en z de signaux usuels	175
Annexe C Termes techniques en anglais	177
Annexe D Convolution	179
Annexe E Rappels de probabilités	183
Annexe F Rappels de calcul différentiel	191
Annexe G Algorithme d'estimation des indices de Sobol	193

Chapitre 1

Introduction

Les progrès technologiques et scientifiques actuels n'auraient pas pu être possibles sans la compréhension et l'évolution des systèmes complexes faisant intervenir des domaines variés d'application comme l'énergie, les télécommunications, le transport, l'aéronautique et le spatial, la robotique, l'économie et la finance, la santé etc. Afin de pouvoir maîtriser le comportement des systèmes complexes, il s'avère nécessaire de modéliser leur comportement de façon quantitative ou qualitative, en tenant compte d'une bonne connaissance du système réel. La modélisation de systèmes a un rôle essentiel pour la commande des systèmes, ainsi que pour l'analyse des interactions entre les divers composants d'un système ou entre différents systèmes. En fonction des hypothèses faites en lien avec le but de la modélisation, des modèles plus simples (utilisés pour la synthèse des lois de commande) ou plus complets (utilisés pour l'analyse du comportement du système réel) peuvent être élaborés. Ce polycopié contient les outils nécessaires pour représenter différents types de systèmes en choisissant une modélisation adéquate (continue ou discrète), pour analyser leur comportement vis-à-vis des divers signaux appliqués en entrée du système, pour faire une analyse de sensibilité et pour valider les modèles choisis selon les hypothèses prises en compte.

1.1 Objectifs et liens avec d'autres cours

A l'issue de ce cours, vous serez capables de représenter et analyser l'évolution d'un système au moyen d'un modèle exploitable analytiquement ou numériquement, adapté à l'objectif de modélisation, déterminé en termes d'hypothèses de modélisation, de représentativité et de niveau de complexité, et d'en déterminer le domaine de validité. Pour cela, vous serez capables de choisir et justifier l'échelle temporelle et spatiale d'intérêt, ainsi que la représentation la plus opportune. Puis, à partir de données expérimentales, vous serez capables de concevoir une structure de modèle et d'en identifier les paramètres, malgré les bruits de mesure inhérents, et finalement d'évaluer les modèles proposés.

Ce cours présente une vision générale de différentes approches de modélisation et permet également de voir comment les différentes étapes de la démarche s'articulent entre elles. Il s'appuie sur des notions enseignées dans le cours de *Convergence, intégration, probabilités, équations aux dérivées partielles*, ainsi que sur des notions d'analyse et des bases d'algorithmique et physique. Ce cours constitue le principal pré-requis du cours de tronc commun d'*Automatique* en deuxième année, qui permettra de synthétiser des lois de commande à partir d'un modèle dynamique. De plus, certaines notions enseignées pendant ce cours seront approfondies lors de prochains cours de tronc commun, de séquences thématiques, de cours électifs ou lors de divers projets dans le cadre des pôles projets associés. On peut citer notamment les cours d'*Optimisation* ou de *Statistiques et Machine Learning* dont on verra en partie III de ce cours toute l'utilité dans la démarche de modélisation. Des renvois à d'autres cours émailleront ainsi ce document, qui constitue à ce titre une sorte de fil-rouge reliant entre elles et donnant du sens à différentes activités pédagogiques du cursus.

1.2 Organisation du polycopié

Chapitre 2 - Du système à la formalisation mathématique

Ce chapitre introduit d'abord la notion de système et ensuite les différentes approches de modélisation des systèmes qui seront proposées dans ce polycopié. Ensuite, la taxonomie des modèles est abordée. Une démarche de modélisation des systèmes est proposée, permettant de représenter différentes classes de systèmes, partant d'un système simple (composé de sous-systèmes) et allant jusqu'à la modélisation des systèmes complexes (systèmes multi-agents ou systèmes de systèmes).

Partie I - Modélisation des systèmes à état continu au sens d'un système dynamique

Chapitre 3 - Modélisation sous forme de représentation d'état

En revenant sur la problématique de la représentation des systèmes dynamiques à temps continu ou à temps discret, ce chapitre introduit le concept de variables d'état qui joue un rôle fondamental pour l'étude des systèmes dynamiques, qu'ils soient linéaires ou non. Nous allons également voir comment le concept de modèle linéaire tangent (autour d'un point de fonctionnement ou d'une trajectoire) permet d'approcher par un système linéaire le comportement d'un système non linéaire. La résolution explicite de l'équation d'état linéaire, ainsi que l'analyse de stabilité par étude des valeurs propres sont proposées.

Chapitre 4 - Modélisation sous forme de fonction de transfert

Les principaux outils permettant l'étude des systèmes linéaires et des signaux qui leur sont associés sont présentés dans le chapitre 4. Les transformées de Fourier (pour les signaux à temps continu ou discret), de Laplace (pour les signaux à temps continu) et la transformée en z (pour les signaux à temps discret) sont brièvement introduites. Les transformées de Laplace et en z permettent de définir la représentation d'un système dynamique par fonction de transfert, qui ne s'applique qu'aux systèmes linéaires et qui joue un rôle fondamental aussi bien pour le traitement des signaux que pour l'analyse des systèmes. Un premier aperçu de l'utilisation de la fonction de transfert (qui ne s'intéresse qu'à la relation entre l'entrée et la sortie d'un système) est donné dès ce chapitre. Celui-ci s'achève par l'analyse de stabilité des systèmes linéaires donnés sous la forme d'une fonction de transfert (à temps continu ou discret). Ensuite, le lien avec la représentation par variables d'état est proposé.

Chapitre 5 - Analyse des performances des systèmes dynamiques

Ce chapitre est principalement dédié à l'étude fréquentielle des systèmes élémentaires d'ordre 1 ou 2 (c'est-à-dire décrits par une équation différentielle linéaire d'ordre 1 ou 2, ou encore par une fonction de transfert rationnelle dont le numérateur et le dénominateur sont des polynômes de degré 1 ou 2), d'où l'on va déduire plus généralement le diagramme du Bode des systèmes d'ordres plus élevés. Le cas des systèmes numériques d'ordre 1 et 2 est également proposé, puis sont présentés les systèmes passe-tout et à déphasage minimal. Enfin, on étudie les systèmes avec retard pur, qui présentent une fonction de transfert non rationnelle.

Après avoir rappelé comment représenter la réponse en fréquence d'un système linéaire dans le plan de Bode, nous étudions l'influence de la position des pôles et des zéros sur les réponses temporelles. La réponse impulsionnelle, la réponse indicielle, ainsi que la réponse à une rampe seront analysées. Là encore, les systèmes à temps continu et discret sont traités en parallèle.

Partie II - Modélisation des systèmes à état discret

Chapitre 6 - Introduction à la modélisation des systèmes à état discret

Nous présentons dans la deuxième partie la modélisation des systèmes à événements discrets. Ce type de système est caractérisé par de brusques transitions du système dans l'espace d'état en réponse à la survenue d'événements. De nombreux systèmes physiques travaillant dans des modes opératoires différents, partageant des ressources ou reposant sur un ordonnancement des tâches (réseaux informatiques, systèmes de production, réseaux de transport, etc.) peuvent ainsi se modéliser en adoptant ce point de vue.

Chapitre 7 - Modélisation des systèmes à événements discrets non temporisés

Après un chapitre introductif présentant de manière informelle les types de modèles inhérents à ces systèmes, nous aborderons tout d'abord dans le chapitre 7 les systèmes à événements discrets non temporisés. Pour ces systèmes, le temps n'apparaît pas de façon explicite dans la modélisation, et seule la séquence d'événements agissant sur le système est considérée pour en décrire l'évolution. Les principaux outils de modélisation de tels systèmes, à savoir les automates et les réseaux de Pétri sont successivement présentés ainsi que les propriétés à analyser pour ces modèles.

Chapitre 8 - Modélisation des systèmes à événements discrets temporisés

Nous traitons ensuite au chapitre 8 le cas des systèmes à événements discrets temporisés où la dimension temporelle est ajoutée au formalisme de modélisation pour décrire le fonctionnement du système. En considérant le cas des automates temporisés, une horloge est ainsi associée à chaque événement pour orchestrer la séquence d'événements. Le cas des automates avec gardes et invariants est ensuite présenté, permettant d'associer des conditions à chaque état et chaque transition. L'extension aux systèmes hybrides où chaque état discret est caractérisé par une dynamique continue différente est enfin exposée, permettant d'englober une large classe de systèmes et faisant le lien avec les développements des systèmes à état continu de la première partie.

Chapitre 9 - Automates stochastiques

Cette deuxième partie de modélisation des systèmes à événements discrets se conclut par une brève présentation de processus stochastiques nécessaires à la modélisation de ces systèmes au chapitre 9. En effet, l'arrivée d'événements n'est que rarement connue à l'avance et correspond, du point de vue de la modélisation du système le recevant, à un processus aléatoire. L'accent est mis dans ce chapitre sur le processus de Poisson qui est de prime importance dans ce contexte.

Partie III - Méthodes pour l'identification et l'évaluation des modèles

Dans cette partie, l'objectif est de confronter le modèle à la réalité, soit via un comportement attendu, soit de manière quantitative via des observations collectées sur le système réel.

Chapitre 10 - Identification

Ce chapitre développe la problématique de l'identification, c'est-à-dire la question de savoir comment spécifier un modèle mathématique à partir de données expérimentales et éventuellement de connaissances a priori. La question posée est de savoir comment fixer la structure d'un modèle et les valeurs de ses paramètres de sorte qu'il permette de reproduire 'au mieux' un comportement attendu ou des valeurs observées expérimentalement sur le système réel. Nous introduirons tout d'abord des approches simples, comme l'analyse indicielle ou l'analyse harmonique, qui peuvent être mises en œuvre dans des cas particuliers. Puis nous présenterons la méthode générale d'identification de modèle paramétrique à partir de données expérimentales. Cette méthode inclut des étapes concernant l'évaluation et la sélection de modèles. Des méthodes pour tester de manière fiable les propriétés du modèle en terme de reproduction fidèle de la réalité d'une part et prédictives d'autre part, ainsi que des critères permettant de choisir entre deux (ou plusieurs) modèles candidats seront discutés dans ce chapitre.

Chapitre 11 - Analyse d'incertitude et de sensibilité

Dans la phase d'analyse et d'évaluation du modèle, ainsi qu'en préalable à son utilisation comme outil de prédiction ou comme support à la prise de décision, il est important de pouvoir quantifier le degré de fiabilité des sorties du modèle, que l'on caractérise notamment par leur incertitude. Ce chapitre présente les méthodes de modélisation de l'incertitude (comment représenter le manque d'information lié à un paramètre ou à une grandeur d'entrée?) et les méthodes de propagation de l'incertitude qui permettent de quantifier l'incertitude obtenue sur une sortie du modèle en fonction des incertitudes sur ses paramètres ou entrées. Lorsque l'incertitude sur une sortie est trop importante, cela peut pénaliser fortement l'utilisation possible du modèle : il faut alors trouver des moyens d'y remédier. Pour cela, on présentera des méthodes d'analyse de sensibilité globale qui permettent notamment de quantifier la contribution des différents facteurs à l'incertitude de la sortie afin de pouvoir mieux cibler les actions de réduction de l'incertitude et revaloriser ainsi l'utilité du modèle développé.

Il est important de noter que, dans la pratique, les étapes présentées ici (identification, évaluation et sélection de modèles, analyse d'incertitude et de sensibilité) ne s'effectuent pas de manière linéaire, mais de manière cyclique dans le processus de développement du modèle. Par exemple, l'analyse de sensibilité peut aussi répondre à la question de sélection des paramètres à identifier dans un modèle présentant des problèmes d'identifiabilité (i.e. dans lequel il n'est pas possible d'estimer de manière fiable tous les paramètres à la fois).

Chapitre 2

Du système à la formalisation mathématique

2.1 Introduction

2.1.1 Notion de système

2.1.1.1 Définitions

Un *système* représente un ensemble d'éléments en interaction entre eux et éventuellement avec l'environnement, qui se coordonnent pour concourir à un résultat. Parmi les caractéristiques d'un système, on peut citer l'interdépendance de ses éléments, ainsi que l'influence du et vers le milieu extérieur via les entrées et les sorties.

Le mot *système* est d'origine grecque *σύνστημα* signifiant combinaison, assemblage, composition. Parmi les synonymes du mot *système* se trouve : *procédé*, *processus* (avec leurs équivalents anglais *process*, *plant*).

La notion de système est utilisée dans divers domaines. Selon Larousse, un *système informatique* est vu comme un ensemble de moyens d'acquisition, de traitement, de stockage et de restitution de données, et de moyens de télécommunication mis en œuvre pour une application ou un ensemble d'applications spécifiés. Un *système social* est un modèle théorique qui décrit et analyse l'interaction sociale humaine. Dans le domaine de l'*Automatique*, le concept de *système dynamique* représente tout système évoluant avec le temps, selon des lois que l'on cherchera à identifier et qui peuvent être déterministes ou stochastiques.

2.1.1.2 Niveaux hiérarchiques

Les systèmes utilisés en pratique étant souvent complexes, une décomposition d'un système en sous-systèmes peut être envisagée dans la phase de modélisation d'un système. La démarche consiste à décomposer le système en sous-systèmes élémentaires effectuant une seule tâche simple, à étudier le comportement de chaque sous-système et le comportement global tenant compte des interactions entre les différents sous-systèmes.

Quand plusieurs systèmes indépendants (appelés *agents*¹ dans ce contexte) interagissent dans un environnement afin de répondre à un but commun ils forment ce que l'on appelle un *système multi-agents*. La notion de système multi-agents a été utilisée au départ dans le domaine de l'*Intelligence Artificielle*, étant caractérisée par le nombre de ses agents, l'environnement dans lequel les agents interagissent et le niveau d'autonomie des agents. Les caractéristiques principales de tels systèmes sont l'autonomie de chaque agent, la communication entre les agents et la coordination pour un but commun. La notion de systèmes (dynamiques) multi-agents est également utilisée en *Automatique* afin de modéliser et de piloter des applications complexes comme le vol

1. Un agent est une entité matérielle ou logicielle qui se trouve dans un certain environnement, et qui est capable d'actions autonomes dans l'environnement afin de répondre à ses objectifs.

en formation de plusieurs drones, le comportement des piétons dans une foule, la fluidité de la circulation automobile etc.

Une autre structure de hiérarchisation des systèmes est ce que l'on appelle les *systèmes de systèmes*, définis comme des systèmes de grande taille, intégrant plusieurs sous-systèmes, possiblement multidisciplinaires, autonomes, mais inter-connectés, afin de satisfaire un besoin global. Parmi les caractéristiques des systèmes de systèmes on trouve l'indépendance opérationnelle, l'autonomie de gestion et de répartition géographique des systèmes constitutifs, le développement évolutif et le comportement émergent. Ces types de systèmes correspondent par exemple au réseau national/européen de distribution d'énergie électrique, au réseau de contrôle du trafic aérien au niveau européen, aux constellations de satellites etc.

2.1.2 Modèles et Modélisation

Au sens premier, le mot *modèle*, issu du latin *modulus*, désigne 'une œuvre ou un objet que l'on imite, que l'on reproduit ; ce qui sert d'exemple, de type, de norme' (dictionnaire de l'Académie Française). Par un glissement sémantique de l'objet imité à l'objet imitant, un modèle désigne également une représentation, pouvant être par exemple de type physique, graphique ou mathématique, qui formalise les relations unissant les différents éléments d'un système, d'un processus, d'une structure, en vue de faciliter la compréhension de certains mécanismes, d'en prévoir le comportement, ou de permettre la validation d'une hypothèse. C'est de manière générale une représentation simplifiée de la réalité, qui peut, ou pas, être accompagnée d'une écriture mathématique. La notion de modèle peut prendre des sens très différents selon les domaines : par exemple en biologie, on parlera de 'modèle' pour désigner un organisme (plante, animal).

La *modélisation* est l'activité visant à représenter un objet ou un phénomène du monde réel par une instance du système formel choisi. Le terme *simulation* désignera ici la méthode qui permet de mettre en œuvre un modèle, généralement via un outil informatique permettant d'interpréter la représentation choisie pour le modèle. La simulation permet notamment d'obtenir des valeurs *calculées* ou *prédites* qui pourront être confrontées à des valeurs *observées* sur le système réel afin d'en évaluer la concordance.

Enfin, on peut terminer en insistant sur le fait qu'un modèle est toujours une simplification du fonctionnement du système réel. Dans ce cours, nous nous intéresserons principalement aux modèles correspondant à des représentations des systèmes sous forme mathématique (d'équations) ou sous forme de graphes.

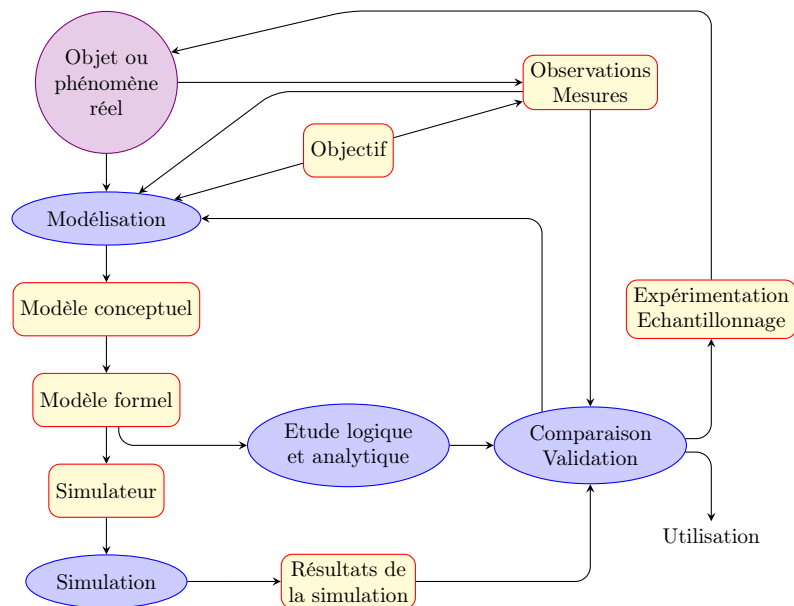


FIGURE 2.1 – Illustration de la démarche de modélisation

La figure 2.1 illustre bien la démarche de modélisation que nous avons pour objectif de présenter

dans ce cours : en fonction de l'objectif formulé, l'objet ou le phénomène d'intérêt est traduit par un modèle conceptuel ou formel (parties I et II : représentations de systèmes à états continus ou discrets) dont on peut, d'une part, étudier les propriétés de manière analytique, d'autre part, simuler le comportement de manière généralement numérique. Les résultats de l'étude analytique ainsi que ceux de la simulation peuvent alors être confrontés à des informations recueillies sur le système réel (partie III), ce qui conduira soit à la mise en utilisation du modèle, soit à le remettre en cause ou bien à indiquer un besoin de collecte de données expérimentales supplémentaires.

2.1.2.1 Définitions des composantes d'un modèle : vocabulaire

Nous avons vu qu'un système (ou processus) est une partie d'un ensemble que nous appréhendons comme un tout et qui interagit avec son environnement. L'activité de modélisation va consister à identifier sur ce système, parfois de manière arbitraire ou subjective (biais de vision du modélisateur) et, en tout cas, toujours en fonction de l'objectif recherché pour la modélisation, un certain nombre de composantes :

- des grandeurs caractérisant le système lui-même, c'est-à-dire permettant de décrire son activité interne de manière aussi exhaustive que possible, les *variables d'état*.
- des variables dont le système subit l'action, les *entrées*, parfois aussi appelées co-variables. Si les entrées peuvent être maîtrisées, c'est-à-dire imposées par un utilisateur externe (par exemple un thermostat), on parlera d'*entrée de commande*. Il existe également des *entrées de perturbation*, vues comme des actions de l'environnement qui ne sont pas manipulées par l'utilisateur (par exemple des conditions météorologiques agissant sur un système). Une étude plus approfondie de ces notions sera réalisée dans le cours d'*Automatique*.
- des variables auxquelles le modélisateur s'intéresse, les *sorties*. Les sorties fournissent la réponse du système aux entrées et dépendent généralement de son état interne, c'est-à-dire de ses variables d'état. Elles peuvent être mesurées (elles sont alors dites observables) ou non. La fonction définissant une sortie observable sera appelée *fonction d'observation*.
- les différentes fonctions ou représentations liant les variables d'état, d'entrée et de sortie peuvent faire apparaître un certain nombre de constantes, que l'on appellera des *paramètres*² du modèle. Ces paramètres sont spécifiques au système considéré et peuvent, ou pas, être mesurés. Ce peut être par exemple un taux de naissance, un coefficient de frottement, une masse, la constante de gravitation, etc. Certains paramètres pourront être déterminés avec précision et d'autres, même mesurables, feront l'objet d'incertitudes dont il faudra déterminer les conséquences sur les sorties du modèle.

2.1.3 Représentations des systèmes : différentes approches

Deux approches, qu'on pourrait presque qualifier de deux philosophies, se présentent pour la démarche de modélisation d'un système.

L'approche dite *mécaniste* ou par *modèle de connaissance* consiste à s'appuyer sur toutes les connaissances disponibles sur les mécanismes de fonctionnement du système afin d'en rendre compte dans le modèle. C'est généralement la démarche adoptée en modélisation en physique où les équations décrivant le système à identifier sont généralement bien connues et issues de lois génériques souvent très précises, au sens où elles permettent de représenter le phénomène réel de manière très fiable (penser à un moteur électrique, par exemple, ou aux lois de la gravitation). Ces modèles restent néanmoins valables seulement dans un certain domaine de validité, c'est-à-dire sous certaines conditions (par exemple les lois de la mécanique générale versus celles de la mécanique quantique). Cette approche peut aussi être mise en œuvre même si le système et les lois qui le régissent sont mal connues, dès lors qu'on s'attache à proposer une vision explicative du fonctionnement du système. Les paramètres du modèle ont alors une interprétation physique. Ce type d'approche peut être réalisé même en l'absence de données.

L'approche *empirique* ou par *modèle de comportement* intervient en particulier lorsque les équations décrivant le système sont inconnues ou bien trop complexes pour une simulation détaillée

2. Parfois ces paramètres peuvent varier dans le temps (par exemple la masse d'un réservoir à essence au fur et à mesure de la consommation lors du déplacement de la voiture).

(par exemple des phénomènes d'échange de chaleur dans un four). Le modélisateur a alors pour objectif de déterminer un modèle simple, qui se comporte comme une bonne approximation du système, même s'il ne représente pas tous les mécanismes internes : le modèle doit simplement être en mesure de reproduire le comportement du système, c'est-à-dire la façon dont il transforme ses entrées en sorties. Ce type d'approche repose donc fortement sur la disponibilité de données : on parle aussi dans ce cas d'approche '*data-driven*' (guidée par les données). Les méthodes de type apprentissage automatique (machine learning) s'apparentent à cette approche.

Il faut noter que de nombreux modèles sont construits selon un mélange des deux approches. Par exemple, on pourra avoir une description mécaniste de la déformation d'un matériau en réponse à une contrainte, dont le coefficient module de Young aura une variation en réponse à la température qui aura été déterminée par une relation issue d'une démarche empirique (en utilisant un jeu de données acquises sur des expériences systématiques d'étude de l'effet de la température).

Exemple 2.1.1. *On peut illustrer ces deux approches sur un exemple, dans lequel elles conduisent toutes deux au même type de modèle. En 1934, deux faisans mâles et six femelles ont été introduits sur une île sur la côte près de Washington. Au cours des 5 années suivantes, la population a connu une croissance d'allure exponentielle (Lack 1954, figure 2.2). La méthode empirique consisterait ici à choisir une forme de modèle, par exemple en introduisant une fonction de type exponentiel $f(t) = ae^{bx}$, où a et b sont des paramètres, et à en estimer les paramètres à partir des données. Si les données étaient beaucoup plus nombreuses ou à évolution complexe, on pourrait utiliser un réseau de neurones récurrent, par exemple.*

La méthode mécaniste consiste à essayer de formaliser le processus de reproduction des faisans. Sous un certain nombre d'hypothèses, entre autres qu'on ne distingue pas entre mâle et femelle et que l'on peut considérer la population comme continue, on va écrire un modèle sous la forme d'une équation différentielle : $\frac{dx(t)}{dt} = r \cdot x(t)$, où r représente le taux de reproduction. La résolution de cette équation conduit également à une fonction de forme exponentielle dont on estimera les paramètres r et x_0 (qui, cette fois, ont des interprétations biologiques) à partir des données. L'intérêt de l'approche mécaniste est que la construction du modèle, en elle-même, fait progresser vers une meilleure compréhension du processus ou du système étudié. Elle possède en revanche un certain nombre de limites (complexité, processus mal représentés lorsque mal connus) et de nombreux modèles dits '*mécanistes*' comportent en réalité souvent un certain nombre de composantes empiriques pour certains sous-processus modélisés.

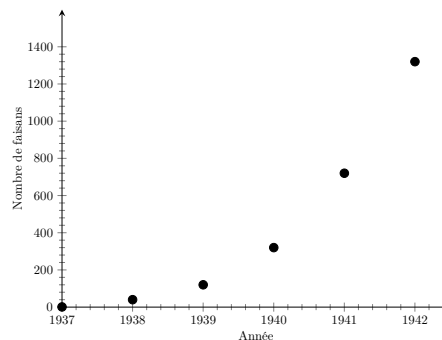


FIGURE 2.2 – Evolution des faisans

2.1.3.1 Sur l'approche par modèle de comportement (empirique)

L'approche par modèle de comportement pour la représentation des systèmes est essentiellement *fonctionnelle*, c'est-à-dire qu'elle permet de représenter les systèmes en tenant compte des échanges des flux dans une architecture de fonctions. En *Automatique* et *Cybernétique*, les *modèles de comportement* intègrent des aspects continus et/ou événementiels, i.e. impliquant des évolutions fondées sur une succession d'états et de transitions. En *Informatique*, des *modèles sémantiques* (plus précisément des modèles de données) sont utilisés pour répondre aux problématiques liées à l'abstraction du logiciel.

Comme on le verra dans la suite de ce chapitre il existe plusieurs types de représentations du comportement d'un système qui ont bien sûr toutes leurs avantages et leurs limitations en fonction de l'utilisation que l'on veut en faire et dont l'équivalence est à étudier.

Une première façon d'aborder l'étude des systèmes dans le cadre d'une modélisation comportementale est de considérer qu'il s'agit d'entités qui reçoivent des signaux d'entrée et qui les transforment en de nouveaux signaux de sortie. La transformation effectuée par le système peut alors être assimilée à la manière dont il réagit à des stimuli et donc à son *comportement*. Pour certains systèmes (dits *systèmes dynamiques*) cette description doit prendre en compte la composante temporelle des différents signaux ; c'est-à-dire qu'à un instant donné la valeur d'un signal de sortie ne dépend pas uniquement de la valeur des signaux d'entrées à cet instant mais également de leurs valeurs aux instants précédents³.

Un problème fondamental que pose alors leur étude est de pouvoir décrire cette transformation que réalise le système. Depuis Descartes, l'approche occidentale pour prendre en compte la complexité est de décomposer les systèmes dans une démarche descendante, d'étudier chacun des sous-systèmes puis de revenir au système global. Pour les systèmes qui nous intéressent ici cette démarche se traduit par la composition des représentations comportementales des sous-systèmes afin d'en déduire le comportement global. Comme nous le verrons, cette approche demande de prendre quelques précautions pour certains systèmes. Il est donc important de comprendre le processus de modélisation qui permet de passer d'un système à sa représentation comportementale. Enfin nous verrons comment le comportement des systèmes peut être caractérisé par certaines propriétés structurelles c'est-à-dire inhérentes au système lui-même.

Cette question de la composition des modèles comportementaux des systèmes physiques (choix d'équations implicites ou explicites, définition de ports d'interaction et affectation en entrées et sorties, prise en compte des contraintes algébriques, etc.) et les problèmes qu'elle induit par exemple au niveau de la simulation est un problème fondamental pour l'ingénieur qui doit construire des modèles de systèmes complexes en prenant en compte les contraintes économiques de temps et de coût de développement de ce modèle.

Cette activité de modélisation ne peut pas être parfaite. La modélisation est donc effectuée pour répondre à un objectif particulier et repose sur un ensemble d'hypothèses simplificatrices. Lorsqu'on utilise le modèle pour mener certaines études et que l'on en tire des conclusions il faut s'assurer que l'on est bien dans un cadre où les hypothèses sont respectées et que les conclusions restent valables pour le système réel.

2.2 Représentations par schéma blocs

Lorsque le système dont on veut représenter le comportement est complexe, il est souvent intéressant de faire apparaître sa structure fonctionnelle et ses sous-systèmes, de représenter ceux-ci puis de les composer pour obtenir une représentation du système global.

Dans ce but, il est possible de représenter les sous-systèmes par des blocs possédant des points d'entrée et de sortie correspondant aux signaux d'entrée et de sortie du système (figure 2.3).

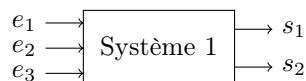


FIGURE 2.3 – Représentation par bloc

La composition est traduite dans cette représentation en reliant les points d'entrée et de sortie des blocs pour indiquer l'égalité des signaux. Ainsi, sur la figure 2.4, il est possible de représenter la composition des systèmes "Système1" et "Système2" par la connexion de leurs entrées et sorties, comme par exemple Système1.s1 et Système2.e1, qui assure l'égalité des signaux associés.

3. Dans ce document on considérera des systèmes dont la composante 'temporelle' correspond à la notion usuelle du temps et est représentée par l'ensemble des réels (cas des systèmes continus) ou des entiers (cas des systèmes discrets) même s'il s'agit d'un cas particulier puisque d'un point de vue théorique la notion de systèmes dynamiques ne demande pour la composante temporelle que l'existence d'un ensemble muni d'une relation d'ordre.

Cette représentation graphique par schémas blocs permet également de prendre en compte des niveaux hiérarchiques d'abstraction de la représentation comportementale puisqu'on peut englober le système issu de la composition dans un nouveau bloc qui possède ses propres entrées et sorties. Les règles d'abstraction qui assurent la cohérence de la représentation sont relativement simples et consistent principalement à assurer que les entrées (respectivement les sorties) d'un bloc sont connectées aux entrées (respectivement aux sorties) des blocs de niveaux inférieurs.

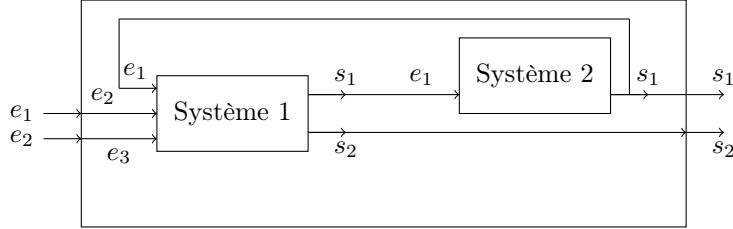


FIGURE 2.4 – Composition de blocs

Exemple 2.2.1 (Échangeur de chaleur). *On s'intéresse aux températures en sortie T_1 et T_2 d'un échangeur de chaleur en fonction des débits d'entrée Q_1 et Q_2 et des températures d'entrée θ_1 et θ_2 . Les équations suivantes permettent de modéliser le système d'échangeur :*

$$C_1 \dot{T}_1(t) = k_1 Q_1(t) (\theta_1(t) - T_1(t)) + \frac{1}{R} (T_2(t) - T_1(t)) \quad (2.2.1)$$

$$C_2 \dot{T}_2(t) = k_2 Q_2(t) (\theta_2(t) - T_2(t)) + \frac{1}{R} (T_1(t) - T_2(t)) \quad (2.2.2)$$

avec C_1, C_2 les capacités thermiques du fluide, R la résistance thermique et k_1, k_2 les coefficients d'échange du fluide.

En linéarisant ce modèle autour d'un point de fonctionnement⁴ défini par Q_{10} , Q_{20} , T_{10} et T_{20} , ce modèle peut se mettre sous la forme du schéma bloc donné dans la figure suivante, avec $N_1(p) = C_1 p + \left(k_1 Q_{10} + \frac{1}{R}\right)$, $N_2(p) = C_2 p + \left(k_2 Q_{20} + \frac{1}{R}\right)$, $a = k_1 Q_{10}$, $b = k_2 Q_{20}$, $c = 1/R$, $\Delta_1 = k_1 (\theta_{10} - T_{10})$ and $\Delta_2 = k_2 (\theta_{20} - T_{20})$.

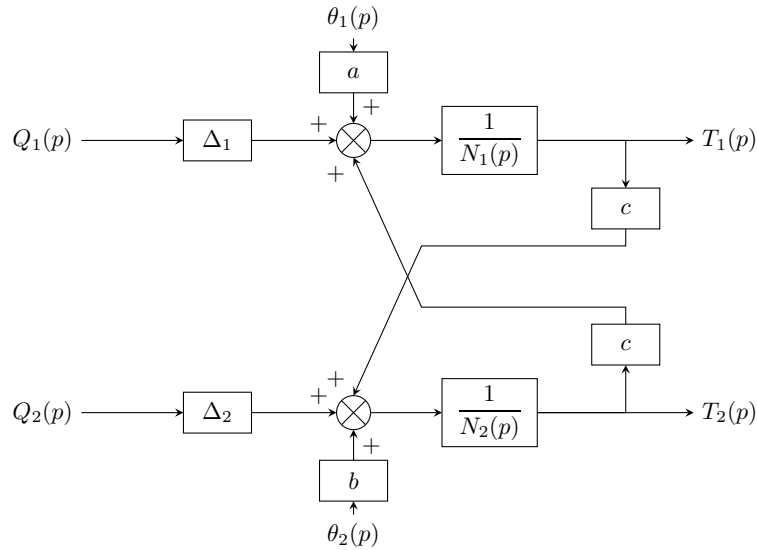


FIGURE 2.5 – Schéma bloc d'un échangeur de chaleur

4. Cette notion est abordée au Chapitre 3.

Cet exemple a permis de mettre en évidence un schéma bloc à base de fonctions de transfert qui seront étudiées dans le Chapitre 4.

2.3 Représentation par variables d'état

Lorsqu'on représente le comportement d'un système par la relation qui relie ses signaux de sortie à ses signaux d'entrée on peut formellement considérer le système comme une application de l'espace des signaux d'entrée dans l'espace des signaux de sortie. Si on veut connaître la réponse du système il est nécessaire de connaître l'ensemble des signaux d'entrée depuis l'origine des temps pour pouvoir calculer leur image par l'application. Ceci reste vrai même si on ne s'intéresse à la réponse du système qu'à partir d'un certain instant et que les entrées avant cet instant restent inconnues. Il est donc intéressant de disposer d'une représentation qui permet de résumer tout le passé du système afin de pouvoir déterminer quelles seront ses réponses à des signaux d'entrée à partir d'un certain instant sans connaître toute l'histoire de ces signaux.

Dans ce but, nous allons introduire la *notion d'état d'un système* qui est définie comme un ensemble de grandeurs qu'il faut connaître à un instant donné pour pouvoir déterminer l'évolution du système à partir de cet instant, connaissant un modèle mathématique et les signaux appliqués à l'entrée.

L'équation d'état est un ensemble d'équations différentielles (voir le cours *Convergence, intégration, probabilités, équations aux dérivées partielles*) du premier ordre décrivant le système à temps continu :

$$\begin{aligned}\dot{\mathbf{x}}(t) &= \mathbf{f}(\mathbf{x}(t), \mathbf{u}(t), t) \\ \mathbf{y}(t) &= \mathbf{g}(\mathbf{x}(t), \mathbf{u}(t), t)\end{aligned}\quad (2.3.1)$$

ou un ensemble d'équations de récurrence d'ordre 1 décrivant un système à temps discret :

$$\begin{aligned}\mathbf{x}_{k+1} &= \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k, k) \\ \mathbf{y}_k &= \mathbf{g}(\mathbf{x}_k, \mathbf{u}_k, k)\end{aligned}\quad (2.3.2)$$

Exemple 2.3.1 (Modèle d'un satellite). *La trajectoire d'un satellite autour de la terre peut être considérée comme plane. Les équations de la mécanique qui définissent cette trajectoire en coordonnées polaires dans ce plan sont les suivantes :*

$$\ddot{r}(t) = r(t)\dot{\theta}^2(t) - \frac{GM}{r^2(t)} + \frac{1}{m}u_1(t) \quad (2.3.3)$$

$$\ddot{\theta}(t) = -\frac{2\dot{r}(t)\dot{\theta}(t)}{r(t)} + \frac{1}{mr(t)}u_2(t) \quad (2.3.4)$$

où G représente la constante de gravitation, M la masse de la terre, m la masse du satellite et u_1 et u_2 correspondent aux composantes radiale et tangentielle de la force délivrée par les moteurs du satellite.

Si on ne s'intéresse qu'à la composante radiale de la position il est possible, en différenciant les équations (2.3.3) et (2.3.4), d'en éliminer la composante tangentielle et d'explicitier par une équation différentielle la relation entre le signal de sortie ($r(t)$) et les signaux d'entrée ($u_1(t)$ et $u_2(t)$).

Si on introduit le vecteur $\mathbf{x}$ défini par $\mathbf{x} = (r \quad \dot{r} \quad \dot{\theta})^T$, les équations (2.3.3) et (2.3.4) peuvent se réécrire :

$$\dot{\mathbf{x}}(t) = \begin{pmatrix} x_2(t) \\ x_1(t)x_3^2(t) - \frac{GM}{x_1^2(t)} + \frac{1}{m}u_1(t) \\ -\frac{2x_2(t)x_3(t)}{x_1(t)} + \frac{1}{mx_1(t)}u_2(t) \end{pmatrix} \quad (2.3.5)$$

qui peut s'écrire sous la forme suivante, où $\mathbf{f}$ est une fonction vectorielle non différentielle :

$$\dot{\mathbf{x}}(t) = \mathbf{f}(\mathbf{x}(t), \mathbf{u}(t))$$

Si on s'intéresse à la composante radiale de la position, considérée comme la sortie du système, la valeur de cette sortie est donnée par :

$$y(t) = \begin{pmatrix} 1 & 0 & 0 \end{pmatrix} x(t) \quad (2.3.6)$$

Exemple 2.3.2 (Modèle d'un circuit électrique). Considérons le circuit électrique de la figure 2.6 pour lequel il est possible de commander la tension $e(t)$ qui est donc l'entrée du système.

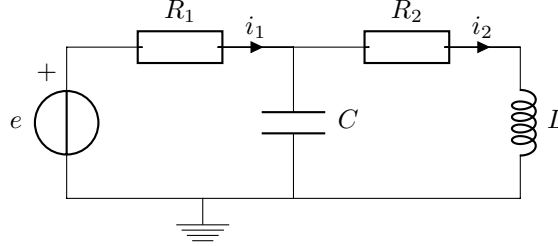


FIGURE 2.6 – Exemple de circuit électrique

Si on s'intéresse au courant qui circule dans R_1 qui est donc considéré comme une sortie du système, la relation qui relie les signaux d'entrée et de sortie est définie par l'équation (2.3.7), dont la résolution pour un signal d'entrée donné permet de connaître le signal de sortie :

$$R_1 LC \ddot{i}_1(t) + (R_1 R_2 C + L) \dot{i}_1(t) + (R_1 + R_2) i_1(t) = LC \ddot{e}(t) + R_2 C \dot{e}(t) + e(t) \quad (2.3.7)$$

Si on s'intéresse au courant dans R_2 considéré comme une autre sortie du système la relation entre l'entrée et cette deuxième sortie est définie par l'équation (2.3.8) :

$$R_1 LC \ddot{i}_2(t) + (R_1 R_2 C + L) \dot{i}_2(t) + (R_1 + R_2) i_2(t) = e(t) \quad (2.3.8)$$

On obtient donc autant d'équations différentielles à résoudre que de sorties. Il est cependant possible de remarquer que l'évolution des grandeurs électriques ne dépend pas des grandeurs observées et que pour un signal d'entrée donné la valeur de la tension aux bornes de la capacité, par exemple, est la même quelle que soit la sortie. En considérant la tension aux bornes de la capacité et le courant dans l'inductance il est possible de modéliser le comportement du circuit par les équations (2.3.9) et (2.3.10) :

$$C \frac{dV_C(t)}{dt} = \frac{e(t) - V_C(t)}{R_1} - i_2(t) \quad (2.3.9)$$

$$L \frac{di_2(t)}{dt} = V_C(t) - R_2 i_2(t) \quad (2.3.10)$$

En considérant le vecteur $\mathbf{x} = (V_C \ i_2)^T$ cette équation peut se réécrire en l'équation (2.3.11) qui est toujours valide quels que soient les signaux observés :

$$\dot{\mathbf{x}}(t) = \begin{pmatrix} -\frac{1}{R_1 C} & -\frac{1}{C} \\ \frac{1}{L} & -\frac{R_2}{L} \end{pmatrix} \mathbf{x}(t) + \begin{pmatrix} \frac{1}{R_1 C} \\ 0 \end{pmatrix} e(t) \quad (2.3.11)$$

Si on veut connaître les signaux correspondant aux courants i_1 et i_2 il suffit d'introduire le vecteur $\mathbf{y} = (i_1 \ i_2)^T$ dont la valeur est définie par l'équation (2.3.12) :

$$\mathbf{y}(t) = \begin{pmatrix} -\frac{1}{R_1} & 0 \\ 0 & 1 \end{pmatrix} \mathbf{x}(t) + \begin{pmatrix} \frac{1}{R_1} & 0 \end{pmatrix} e(t) \quad (2.3.12)$$

L'introduction du vecteur $\mathbf{x}$ permet ici encore de séparer dans la représentation comportementale une relation dynamique (équation (2.3.11)) qui est valable quelles que soient les grandeurs observées et une relation statique qui définit les grandeurs observées (équation (2.3.12)).

Cette forme de représentation d'état est très générale. Néanmoins, pour certains systèmes⁵ il est nécessaire de prendre en compte un vecteur d'état hybride comportant une composante homogène à $\mathbb{R}^n$ dont l'évolution peut être prise en compte par une équation de type (2.3.1) et une autre composante logique dont l'évolution est spécifiée par une machine d'état.

2.4 Modélisation à événements discrets

Les modèles de systèmes physiques sont souvent des modèles à base d'équations différentielles (temps continu) ou des équations aux différences (temps discret) tels que nous venons de les introduire. Pour un certain nombre de systèmes, ce type de représentation n'est pas approprié. Ainsi, il existe par exemple des systèmes réels (par exemple des convertisseurs de puissance) qui sont composés de sous-processus continus qui sont démarrés, reconfigurés et arrêtés par une commande logique à état discrets. Un *système à événements discrets* est un système dynamique défini par un espace d'états discrets et par des évolutions (appelées trajectoires) fondées sur une succession d'états et de transitions. Les transitions sont étiquetées par des symboles (appelés *événements*).

Il existe plusieurs outils pour la modélisation des systèmes à événements discrets : les *automates* (à états finis⁶ ou temporisés⁷), le *grafcet* (utilisé pour la modélisation des machines évoluant en parallèle), les *réseaux de Petri* (qui sont des graphes formés de deux types de noeuds : des places et des transitions, reliées par des arcs orientés) et les *StateCharts* (formalisme de description graphique pour les systèmes réactifs complexes, caractérisés habituellement par les mots clés suivants : automates, profondeur, orthogonalité et communication par diffusion).

Les automates ou *machines à états* permettent de modéliser explicitement l'espace d'état d'un système. Les automates sont particulièrement bien adaptés à la modélisation de la logique de commande, notamment pour les systèmes critiques. En effet, les automates se prêtent à des méthodes formelles d'analyse qui permettent de valider certaines propriétés. Enfin, il est facile d'implémenter un automate aussi bien en logiciel qu'en matériel.

Un inconvénient des automates est leur faible expressivité : il s'agit d'une modélisation d'assez bas niveau et de petites modifications de la spécification d'un système peuvent entraîner de grands changements dans la structure d'un automate. Un autre inconvénient est que le nombre d'états d'un automate peut devenir très grand, même pour un système de complexité modérée.

En combinant les automates (qui sont purement séquentiels) avec des modèles qui supportent le parallélisme, on peut modéliser des systèmes complexes avec des automates ayant un nombre gérable d'états. Par exemple, les *StateCharts*, qui sont maintenant utilisés par UML pour modéliser les aspects dynamiques d'un objet, combinent les automates avec un modèle réactif synchrone. Lorsque l'on combine des automates avec des équations différentielles, on obtient un modèle de *système dynamique hybride*.

Exemple 2.4.1 (Exemple d'automate). La figure 2.7 montre un automate qui modélise le fonctionnement d'un répondeur téléphonique. Les entrées du système sont : *sonnerie* (signal d'appel), *décroché* (le combiné est décroché), *fin_annonce* (la lecture de l'annonce est terminée), *fin_message* (la fin du message est détectée – présence de la tonalité par exemple).

Les sorties sont : *répondre* (répondre à l'appel et commencer à lire l'annonce), *enregistrer* (commencer à enregistrer le message), *messages* (un ou plusieurs messages ont été enregistrés).

Les états du système sont : *attente* (état de repos du système), *1 sonn.* (on a compté une sonnerie), *2 sonn.* (on a compté 2 sonneries), *annonce* (on a compté 3 sonneries et on a lancé la lecture de l'annonce), *enregistre* (l'annonce est terminée, on a lancé l'enregistrement du message).

Les transitions marquées par le symbole $\perp$ sont prises lorsqu'aucune entrée n'est présente. Les transitions marquées par *sinon* sont prises lorsqu'une entrée autre que celles pour lesquelles existe une transition explicite est présente.

Les transitions $\perp$ qui bouclent sur un état sans produire de sortie ne sont en général pas indiquées sur le diagramme d'un automate. Ces transitions ne sont utiles que lorsqu'on compose

5. C'est par exemple le cas d'un radiateur électrique dans une pièce à qui on fixe une température à assurer θ_e signal d'entrée et qui pour ce faire met en marche sa résistance lorsque la température est inférieure à $(\theta_e - h)$ et l'éteint lorsque la température est supérieure à $(\theta_e + h)$.

6. Un automate à états finis est un automate avec un nombre fini d'états.

7. Les automates temporisés sont des automates à états discrets finis auxquels on a associé une horloge.

un automate avec un autre, l'un des automates étant parfois amené à ne rien faire (*to stutter* : bégayer, balbutier) pendant que l'autre réagit à une entrée.

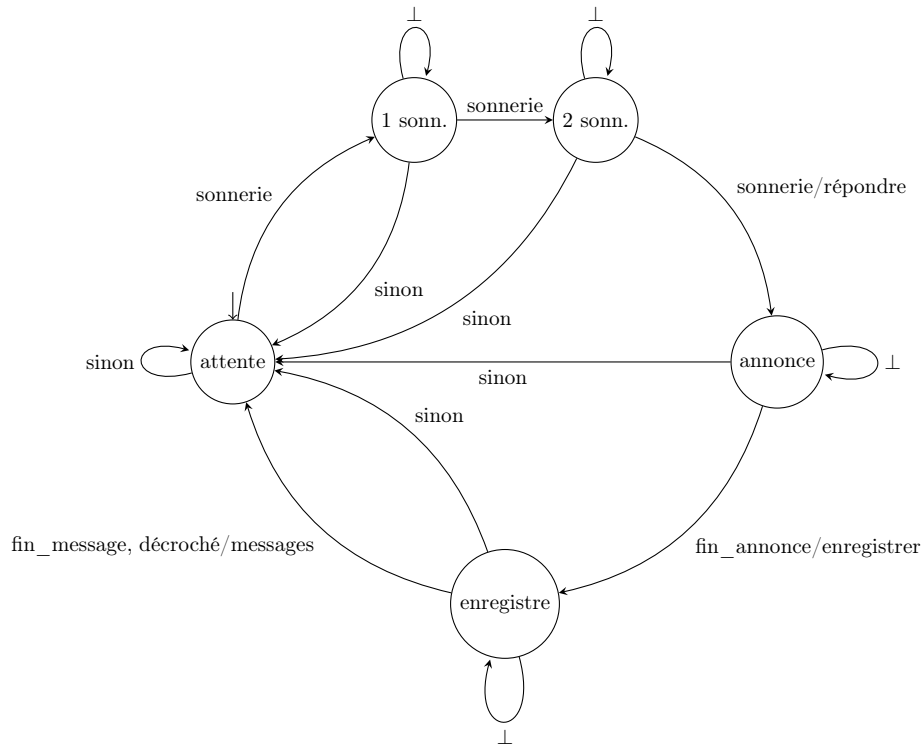


FIGURE 2.7 – Automate d'un répondeur téléphonique

2.5 Classification des systèmes

Cette partie propose une classification des systèmes selon divers critères.

2.5.1 Nombre d'entrées/sorties

Un système est un dispositif présentant un certain nombre d'entrées (*signaux d'entrée*) et de sorties (*signaux de sortie*) ; il réalise une application de l'espace des signaux d'entrée dans celui des signaux de sortie (figure 2.8).

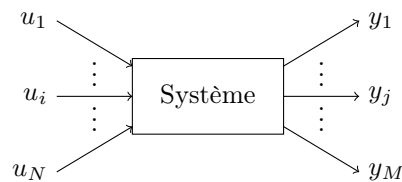


FIGURE 2.8 – Système

2.5.1.1 Systèmes monovariables

Définition 2.5.1

Un système avec une seule entrée et une seule sortie est appelé *système monovariabile* ou système SISO (*Single-Input Single-Output*).

2.5.1.2 Systèmes multivariables

Définition 2.5.2

Un système avec plusieurs entrées et plusieurs sorties est appelé *système multivariable* ou système MIMO (*Multiple-Input Multiple-Output*).

Un système avec plusieurs entrées et une seule sortie est appelé système MISO (*Multiple-Input Single-Output*). Un système avec une seule entrée et plusieurs sorties est appelé système SIMO (*Single-Input Multiple-Output*).

Exemple 2.5.1 (Systèmes multivariables). *Robots, drones etc.*

Dans la suite, sauf indication contraire, on s'intéresse aux systèmes à une entrée et une sortie.

2.5.2 Critère temporel

2.5.2.1 Système continu

Définition 2.5.3

Un système est dit *continu* (ou *analogique*) si les signaux d'entrée et de sortie sont des signaux à temps continu.

Exemple 2.5.2. *Filtres analogiques, moteurs, gyroscopes etc.*

2.5.2.2 Système discret

Définition 2.5.4

Un système est dit *discret* (ou *numérique*) si les signaux d'entrées et de sorties sont des signaux à temps discret.

Exemple 2.5.3. *Filtres numériques, systèmes de gestion de production etc.*

2.5.2.3 Système hybride entrée/sortie

Définition 2.5.5

Un système est dit *hybride* (ou *mixte*) si les entrées sont des signaux à temps continu (respectivement discret) et les sorties sont des signaux à temps discret (respectivement continu).

Exemple 2.5.4. *Convertisseurs de puissance, systèmes en commutation etc.*

2.5.3 Linéarité

Définition 2.5.6

Un système est dit *linéaire* si à toute combinaison linéaire des signaux d'entrée correspond la même combinaison linéaire des signaux de sortie.

Exemple 2.5.5. *Soit un système linéaire à une entrée u et une sortie y . Si y_1 et y_2 sont les réponses aux actions u_1 et u_2 , la réponse à $\lambda u_1 + \mu u_2$ sera $\lambda y_1 + \mu y_2$.*

Tout système qui ne respecte pas cette propriété est un *système non linéaire*.

2.5.4 Invariance dans le temps

Définition 2.5.7

Soit $y(t)$ la réponse du système à l'action $u(t)$. Le système est dit *invariant dans le temps* (ou *stationnaire* ou *permanent*) si, quel que soit le décalage τ , $y(t - \tau)$ est la réponse à l'action $u(t - \tau)$.

Un système réel n'est quasiment jamais stationnaire : même les systèmes conçus par l'homme ne le sont pas à cause du vieillissement des composants qui modifie la relation entrée/sortie. On considère en pratique qu'un système est invariant dans le temps si cette modification est imperceptible dans la fenêtre d'observation du système.

Exemple 2.5.6. *À l'échelle de la minute, un filtre LC est un système invariant contrairement à l'appareil de production de la parole qui est non stationnaire. Ce dernier peut être considéré comme stationnaire à l'échelle de la dizaine de millisecondes.*

Une classe particulière de systèmes est représentée par les systèmes linéaires et invariants dans le temps.

2.5.5 Causalité

Définition 2.5.8

Un système est dit *causal* si sa réponse à un instant donné ne dépend que des valeurs de l'entrée aux instants précédents.

Autrement dit, si on considère un système au repos auquel on applique une action $u(t)$ nulle pour t inférieur à t_0 , le système sera dit causal si la sortie $y(t)$ est nulle pour t inférieur à t_0 . Bien que la plupart des systèmes réels soient causaux, il peut être intéressant dans certains cas (filtre numérique fonctionnant à temps différé par exemple) d'étudier des systèmes non causaux.

Exemple 2.5.7. *Un retard pur est un système causal : $y(t) = u(t - \tau)$. Le filtre numérique de la figure 2.9 est un système non causal puisque $y(k) = u(k - 1) + u(k) + u(k + 1)$.*

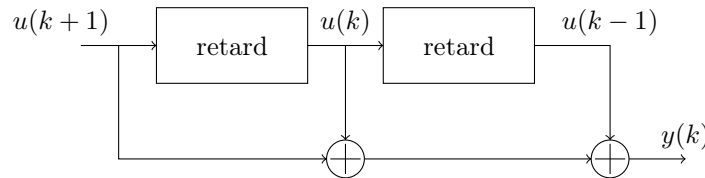


FIGURE 2.9 – Filtre numérique non causal

2.5.6 Évolution au cours du temps

Définition 2.5.9

Un système est dit *statique* si ses sorties ne dépendent que des valeurs présentes des entrées. Sinon il est *dynamique*.

Exemple 2.5.8 (Système dynamique). *Système avec une sortie décrite par exemple par $y(k) = f(u(k), u(k - 1))$, où $u(k)$ représente le signal d'entrée.*

2.6 Modèles de signaux certains

Enfin, pour terminer ce chapitre d'introduction, on définit un certain nombre de signaux de base qui seront utilisés dans le cours : l'échelon, l'impulsion et le signal sinusoïdal.

De manière générale, un signal à temps continu est une fonction du temps et un signal à temps discret est une suite infinie de réels.

2.6.1 Echelon unité ou échelon de Heaviside

La figure 2.10 présente l'échelon unité ou échelon de Heaviside (noté $\mathbb{1}$ dans ce polycopié). Il s'agit d'une modélisation très simple de la mise en marche d'un dispositif.



FIGURE 2.10 – Echelon de Heaviside

2.6.2 Impulsion de Dirac

Formellement, la distribution de Dirac en 0 est la forme linéaire définie sur $\mathcal{D}(I)$, l'ensemble des fonctions à support compact de classe C^∞ dans I ou ensemble des fonctions-test, par :

$$\forall \Phi \in \mathcal{D}(I), \quad \Phi(0) = \int_{-\infty}^{+\infty} \delta(t) \Phi(t) dt \quad (2.6.1)$$

L'impulsion de Dirac peut être vue comme la « dérivée », au sens des distributions, de l'échelon de Heaviside. Elle correspond physiquement à un phénomène infiniment bref et infiniment intense. C'est un signal théorique très utile pour caractériser des systèmes, bien qu'elle ne soit pas réalisable physiquement car la fonction qui lui serait associée (par exemple si la distribution était régulière) devrait vérifier des conditions du type :

$$\delta(t) = 0 \text{ si } t \neq 0 \quad (2.6.2)$$

$$\delta(0) = +\infty \quad (2.6.3)$$

$$\text{et } \int_{-\infty}^{+\infty} \delta(t) dt = 1 \quad (2.6.4)$$

ce qui n'est pas possible. Du fait de son utilité importante dans ce cours, on s'autorisera cependant la notation abusive de l'impulsion de Dirac sous forme d'une fonction.

Dans le cas discret, l'impulsion discrète $\delta(k)$ est la suite nulle partout sauf au point $k = 0$ où elle vaut 1.

2.6.3 Signal sinusoïdal

Dans le cas continu, un signal sinusoïdal est défini par :

$$x(t) = A \sin(2\pi f_0 t + \varphi) \quad (2.6.5)$$

La fréquence du signal est f_0 , sa pulsation est $\omega_0 = 2\pi f_0$, son amplitude est A et sa phase est φ . Dans le cas de signaux discrets, un signal sinusoïdal peut être défini par :

$$x(k) = A \sin(2\pi \nu_0 k + \varphi) \quad (2.6.6)$$

où ν_0 est la fréquence du signal sinusoïdal. Il s'agit ici d'une grandeur sans dimension appartenant à l'intervalle $\left[-\frac{1}{2}, \frac{1}{2}\right]$. En effet, les signaux de fréquence $\nu_0 + n$, avec n entier, définissent la même suite de nombres, avec ν_0 appelée la *fréquence réduite*.

Première partie

Modélisation des systèmes à état continu au sens d'un système dynamique

Chapitre 3

Modélisation sous forme de représentation d'état

3.1 Définition

Soit S un système à temps continu ayant pour entrée le signal vectoriel $\mathbf{u}(t)$.

Définition 3.1.1

Une *représentation d'état* du comportement dynamique de S est la donnée d'un vecteur $\mathbf{x}$ et d'une équation différentielle du type :

$$\dot{\mathbf{x}}(t) = f(\mathbf{x}(t), \mathbf{u}(t), t) \quad (3.1.1)$$

avec $\mathbf{u}(t) \in \mathbb{R}^m$ le vecteur d'entrée et $\mathbf{x} \in \mathbb{R}^n$ le vecteur d'état formé par les variables d'état $x_i(t)$ qui représentent un ensemble de n variables indépendantes $x_1(t), x_2(t), \dots, x_n(t)$ choisies telles que leur seule donnée à un instant $t = t_0$ suffit à décrire, de façon exhaustive, l'état du système à l'instant t_0 .

C'est-à-dire que toute grandeur du système peut être déterminée, à l'instant t , à partir des $x_i(t_0)$ et des $u_i(\tau)$, pour $\tau \in [t_0, t]$.

Les équations (2.3.5) et (2.3.11) sont du type de l'équation (3.1.1) pour les exemples de satellite et de circuit électrique.

L'équation (3.1.1) est appelée *équation dynamique* (*équation d'évolution* ou *première équation d'état*) du système. Il est possible de lui adjoindre une *équation d'observation* (ou *équation de sortie*) qui permet de définir les sorties du système :

$$\mathbf{y}(t) = g(\mathbf{x}(t), \mathbf{u}(t), t) \quad (3.1.2)$$

avec $\mathbf{y}(t) \in \mathbb{R}^p$ le vecteur de sortie du système.

Cette définition s'étend au cas des systèmes à temps discret par la donnée d'un vecteur $\mathbf{x}$, d'une équation dynamique récurrente et d'une équation d'observation.

Définition 3.1.2

Un système à temps discret sous forme d'état est donné par :

$$\begin{aligned} \mathbf{x}(k+1) &= f(\mathbf{x}(k), \mathbf{u}(k), k) \\ \mathbf{y}(k) &= g(\mathbf{x}(k), \mathbf{u}(k), k) \end{aligned} \quad (3.1.3)$$

Remarque 3.1.1. Pour un système donné, la représentation d'état n'est pas unique.

L'intérêt de disposer d'une telle représentation est qu'elle va en général permettre le calcul des réponses du système à des signaux d'entrée à partir d'un instant particulier si la valeur du vecteur

d'état à cet instant est connue. Son intérêt est également qu'à partir de l'étude de f il est possible de donner des propriétés du système sans calculer les réponses à des signaux particuliers.

Une première propriété qu'il est ainsi possible d'étudier est l'existence et l'unicité de la solution de l'équation (3.1.1) dans le cas continu¹ (problème de Cauchy). Il est ainsi possible de montrer que si la fonction f est continûment dérivable, c'est-à-dire si sa dérivée selon chacun de ses arguments existe et est continue, alors pour tout instant t_0 , tout état initial $\mathbf{x}_0$ et tout signal $\mathbf{u}(t)$ il existe un unique signal vectoriel $\mathbf{x}(t)$ défini sur un intervalle I contenant t_0 , solution de (3.1.1) et tel que $\mathbf{x}(t_0) = \mathbf{x}_0$. Ce résultat classique de la théorie des équations différentielles est un résultat local puisque l'unicité de la solution n'est garantie que pour un intervalle contenant t_0 . Cependant dans la pratique des systèmes étudiés dans ce cours, cet intervalle s'étendra sur l'ensemble des réels considéré comme la composante temporelle des différents signaux.

Définition 3.1.3

Un système *invariant* est un système dont la dynamique ne dépend pas du temps mais uniquement du vecteur d'état et de la commande.

Ainsi, un système *non linéaire invariant dans le temps* s'écrit sous la forme suivante :

— A temps continu :

$$\begin{aligned}\dot{\mathbf{x}}(t) &= f(\mathbf{x}(t), \mathbf{u}(t)) \\ \mathbf{y}(t) &= g(\mathbf{x}(t), \mathbf{u}(t))\end{aligned}\tag{3.1.4}$$

— A temps discret :

$$\begin{aligned}\mathbf{x}(k+1) &= f(\mathbf{x}(k), \mathbf{u}(k)) \\ \mathbf{y}(k) &= g(\mathbf{x}(k), \mathbf{u}(k))\end{aligned}\tag{3.1.5}$$

3.1.1 Équation d'état et équation de sortie

L'ensemble des valeurs $x_i(t_0)$ résumant le passé du système, son évolution pour $t > t_0$ est uniquement fonction de ces valeurs $x_i(t_0)$ et des actions $u_i(t)$ (pour $t > t_0$).

Pour chaque variable d'état $x_i(t)$ il existe donc une fonction f_i telle que :

$$\frac{dx_i(t)}{dt} = f_i(x_1(t), x_2(t), \dots, x_n(t), u_1(t), u_2(t), \dots, u_m(t))$$

Si le système est linéaire², les fonctions f_i ont une forme linéaire :

$$\frac{dx_i(t)}{dt} = \sum_{k=1}^n a_{ik}(t)x_k(t) + \sum_{k=1}^m b_{ik}(t)u_k(t)$$

On obtient alors la première équation d'état d'un système linéaire :

$$\boxed{\frac{d\mathbf{x}(t)}{dt} = \mathbf{A}(t)\mathbf{x}(t) + \mathbf{B}(t)\mathbf{u}(t)}\tag{3.1.8}$$

1. Les conditions d'existence de solutions à l'équation (3.1.3) sont simples à déterminer et sont liées au domaine de définition de f .

2. Un système est dit à *dynamique linéaire* si et seulement si :

$$\begin{aligned}\forall \mathbf{x}_1, \mathbf{x}_2, \mathbf{u}_1, \mathbf{u}_2, t \quad & f(\mathbf{x}_1 + \mathbf{x}_2, \mathbf{u}_1 + \mathbf{u}_2, t) = f(\mathbf{x}_1, \mathbf{u}_1, t) + f(\mathbf{x}_2, \mathbf{u}_2, t) \\ \forall \mathbf{x}, \mathbf{u}, \lambda \in \mathbb{R}, t \quad & f(\lambda \mathbf{x}, \lambda \mathbf{u}, t) = \lambda f(\mathbf{x}, \mathbf{u}, t)\end{aligned}\tag{3.1.6}$$

Il est dit à *observation linéaire* si et seulement si :

$$\begin{aligned}\forall \mathbf{x}_1, \mathbf{x}_2, \mathbf{u}_1, \mathbf{u}_2, t \quad & g(\mathbf{x}_1 + \mathbf{x}_2, \mathbf{u}_1 + \mathbf{u}_2, t) = g(\mathbf{x}_1, \mathbf{u}_1, t) + g(\mathbf{x}_2, \mathbf{u}_2, t) \\ \forall \mathbf{x}, \mathbf{u}, \lambda \in \mathbb{R}, t \quad & g(\lambda \mathbf{x}, \lambda \mathbf{u}, t) = \lambda g(\mathbf{x}, \mathbf{u}, t)\end{aligned}\tag{3.1.7}$$

Un système sous forme d'état est un *système linéaire* s'il respecte les propriétés (3.1.6) et (3.1.7).

où $\mathbf{x}(t)$ est le vecteur d'état $\mathbf{x}(t) = \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_n(t) \end{pmatrix}$ et $\mathbf{u}(t)$ est le vecteur source $\mathbf{u}(t) = \begin{pmatrix} u_1(t) \\ u_2(t) \\ \vdots \\ u_m(t) \end{pmatrix}$.

$\mathbf{A}$ s'appelle la *matrice d'évolution* ; elle est de dimension $n \times n$. $\mathbf{B}$ s'appelle la *matrice d'entrée* (ou de commande) ; elle est de dimension $n \times m$.

De même, les variables d'état suffisant à décrire l'état du système les sorties $y_j(t)$ s'expriment sous la forme :

$$y_j(t) = h_j(x_1(t), x_2(t), \dots, x_n(t), u_1(t), u_2(t), \dots, u_m(t))$$

Si le système est linéaire, on obtient la deuxième équation d'état :

$$\boxed{\mathbf{y}(t) = \mathbf{C}(t)\mathbf{x}(t) + \mathbf{D}(t)\mathbf{u}(t)} \quad (3.1.9)$$

où $\mathbf{y}(t)$ est le vecteur de sortie $\mathbf{y}(t) = \begin{pmatrix} y_1(t) \\ y_2(t) \\ \vdots \\ y_p(t) \end{pmatrix}$.

$\mathbf{C}$ s'appelle la *matrice d'observation* ; elle est de dimension $p \times n$. $\mathbf{D}$ s'appelle la *matrice de transmission directe* ; elle est de dimension $p \times m$.

Si l'on récapitule, un *système à temps continu, linéaire variant dans le temps*, s'écrit sous la forme :

$$\begin{aligned} \dot{\mathbf{x}}(t) &= \mathbf{A}(t)\mathbf{x}(t) + \mathbf{B}(t)\mathbf{u}(t) \\ \mathbf{y}(t) &= \mathbf{C}(t)\mathbf{x}(t) + \mathbf{D}(t)\mathbf{u}(t) \end{aligned} \quad (3.1.10)$$

Remarque 3.1.2. (i) Si le système est linéaire ET stationnaire (ou invariant dans le temps), les matrices $\mathbf{A}$, $\mathbf{B}$, $\mathbf{C}$ et $\mathbf{D}$ sont indépendantes du temps : $\mathbf{A}(t) = \mathbf{A}$, $\mathbf{B}(t) = \mathbf{B}$, $\mathbf{C}(t) = \mathbf{C}$, $\mathbf{D}(t) = \mathbf{D}$.

(ii) Pour un système donné, plusieurs vecteurs d'état sont possibles pour représenter le système. On passe d'un vecteur d'état à un autre par une matrice $\mathbf{M}$ inversible appelée *matrice de transformation*. Soit $\mathbf{x}(t)$ un vecteur d'état ; $\mathbf{z}(t) = \mathbf{M}\mathbf{x}(t)$ satisfait la nouvelle équation d'état :

$$\frac{d\mathbf{z}(t)}{dt} = \mathbf{M}\mathbf{A}\mathbf{M}^{-1}\mathbf{z}(t) + \mathbf{M}\mathbf{B}\mathbf{u}(t) = \mathbf{A}'\mathbf{z}(t) + \mathbf{B}'\mathbf{u}(t)$$

La représentation d'un système par variables d'état n'est donc pas unique. En toute rigueur on devrait parler d'*une* représentation d'état et non de *la* représentation d'état.

(iii) Dans certains ouvrages, la première équation d'état s'appelle l'équation d'état et la deuxième équation d'état porte le nom d'équation de sortie ou équation de mesure.

Ces propriétés sont données pour des systèmes continus, mais elles peuvent s'étendre sans difficulté au cas des *systèmes à temps discret, linéaire et invariant dans le temps* modélisés par les équations (3.1.3). Les équations d'état du système sont alors exprimées par les équations (3.1.11) :

$$\begin{aligned} \mathbf{x}(k+1) &= \mathbf{A}(k)\mathbf{x}(k) + \mathbf{B}(k)\mathbf{u}(k) \\ \mathbf{y}(k) &= \mathbf{C}(k)\mathbf{x}(k) + \mathbf{D}(k)\mathbf{u}(k) \end{aligned} \quad (3.1.11)$$

Comme on le verra dans la suite de ce cours qui se concentre sur les systèmes linéaires invariants, ceux-ci possèdent un certain nombre de caractéristiques qui facilitent leur étude. En particulier la linéarité assure que les propriétés (stabilité, etc.) sont des propriétés structurelles du système qui ne dépendent donc pas d'une utilisation particulière de celui-ci.

3.2 Stabilité d'un point d'équilibre ou d'une trajectoire d'équilibre

La *stabilité* d'un système correspond intuitivement au fait que le système "ne va pas diverger et garder un comportement maîtrisable".

Considérons un système à temps continu dont le comportement est décrit par l'équation d'état³ dynamique la plus générale (3.1.1) et tel que l'entrée $\mathbf{u}(t)$ soit fixée à une valeur quelconque $\mathbf{u}$. Cette équation peut alors s'écrire :

$$\dot{\mathbf{x}}(t) = \mathbf{f}_{\mathbf{u}}(\mathbf{x}(t), t) \quad (3.2.1)$$

Définition 3.2.1

Le point $\bar{\mathbf{x}}$ de l'espace d'état est un *point d'équilibre d'un système à temps continu* si $\mathbf{f}_{\mathbf{u}}(\bar{\mathbf{x}}, t) = 0, \forall t$.

Définition 3.2.2

Un point d'équilibre est *stable au sens de Lyapunov* si lorsque l'état du système est déplacé dans un voisinage de $\bar{\mathbf{x}}$ il reste dans un voisinage du point d'équilibre, ce qui peut formellement s'énoncer par l'équation (3.2.2) où $\mathbf{x}(t_0)$ est l'état à l'instant t_0 et $\mathbf{x}(t)$ est la solution de (3.2.1) définie par ces conditions initiales :

$$\forall \varepsilon > 0, \exists \delta(\varepsilon, t_0) \text{ tq } \|\bar{\mathbf{x}} - \mathbf{x}(t_0)\| < \delta(\varepsilon, t_0) \Rightarrow \|\mathbf{x}(t) - \bar{\mathbf{x}}\| < \varepsilon, \forall t > t_0 \quad (3.2.2)$$

Définition 3.2.3

Le point d'équilibre est *asymptotiquement stable* s'il est stable et si l'état tend vers ce point d'équilibre lorsqu'il a été déplacé. Ceci revient à ajouter la condition (3.2.3) à la condition (3.2.2) :

$$\lim_{t \rightarrow \infty} \|\mathbf{x}(t) - \bar{\mathbf{x}}\| = 0 \quad (3.2.3)$$

Si le point d'équilibre est asymptotiquement stable il est parfois possible de caractériser la convergence de l'état vers ce point.

Définition 3.2.4

Le point d'équilibre sera dit *exponentiellement stable* s'il est possible de majorer l'évolution de la distance entre l'état et le point d'équilibre par une exponentielle décroissante.

Cette stabilité exponentielle demande l'existence de deux réels positifs M et α indépendants de t_0 et de $\mathbf{x}(t_0)$ tels que l'on ait :

$$\|\mathbf{x}(t) - \bar{\mathbf{x}}\| \leq M \|\mathbf{x}(t_0) - \bar{\mathbf{x}}\| e^{-\alpha(t-t_0)} \quad (3.2.4)$$

Ces différentes définitions sont données ici dans le cas des systèmes continus et se transposent facilement au cas des systèmes discrets.

Définition 3.2.5

Un *point d'équilibre* (ou *point fixe*⁴) d'un système à temps discret décrit par $\mathbf{x}_{k+1} = \mathbf{f}_{\mathbf{u}}(\mathbf{x}_k, k)$ est un point qui satisfait $\bar{\mathbf{x}} = \mathbf{f}_{\mathbf{u}}(\bar{\mathbf{x}}, k) \forall k$.

Ces considérations sur la stabilité sont locales et ne préjugent pas du comportement du système sur l'ensemble de son domaine de fonctionnement. Elles correspondent d'autre part à une vue statique du système et il est parfois plus pertinent d'en considérer une vue dynamique. Par exemple si on s'intéresse au comportement d'une fusée en vol la notion de point d'équilibre n'est pas très utile. Il est alors plus intéressant d'introduire une notion de *stabilité de trajectoire* qui va caractériser la capacité du système à suivre cette trajectoire.

3. Pour cette partie l'équation d'observation n'est pas nécessaire.

4. Le terme de 'point fixe' est utilisé principalement dans le cas des systèmes discrets mais peut également être utilisé pour un système continu.

Reprenons donc un système défini par l'équation d'état (3.1.1) et considérons que l'entrée du système est un signal vectoriel prédéfini $\mathbf{u}_1(t)$. Le choix d'une valeur initiale⁵ pour le vecteur d'état définit complètement la *trajectoire* suivie par la solution de (3.1.1). Notons $\Phi_{\mathbf{x}}(t)$ cette trajectoire solution de (3.1.1) en réponse à l'entrée $\mathbf{u}_1(t)$ définie par la valeur initiale $\mathbf{x}$ du vecteur d'état.

Définition 3.2.6

La trajectoire $\Phi_{\mathbf{x}_0}(t)$ définie par le point initial $\mathbf{x}_0$ est *stable* si et seulement si :

$$\forall \varepsilon, \exists \delta(\varepsilon) \text{ tq } \|\mathbf{x}_0 - \mathbf{x}_1\| < \delta(\varepsilon) \Rightarrow \|\Phi_{\mathbf{x}_0}(t) - \Phi_{\mathbf{x}_1}(t)\| < \varepsilon, \forall t \quad (3.2.5)$$

Cette propriété peut s'illustrer par l'exemple en dimension 1 de la figure 3.1. La donnée du point initial $\mathbf{x}_0$ définit une trajectoire qui est stable si, lorsqu'on choisit un écart ε , il est possible de trouver un voisinage de $\mathbf{x}_0$ tel que les trajectoires définies par les points de ce voisinage soient toutes dans le 'tube' généré par la trajectoire initiale et l'écart.

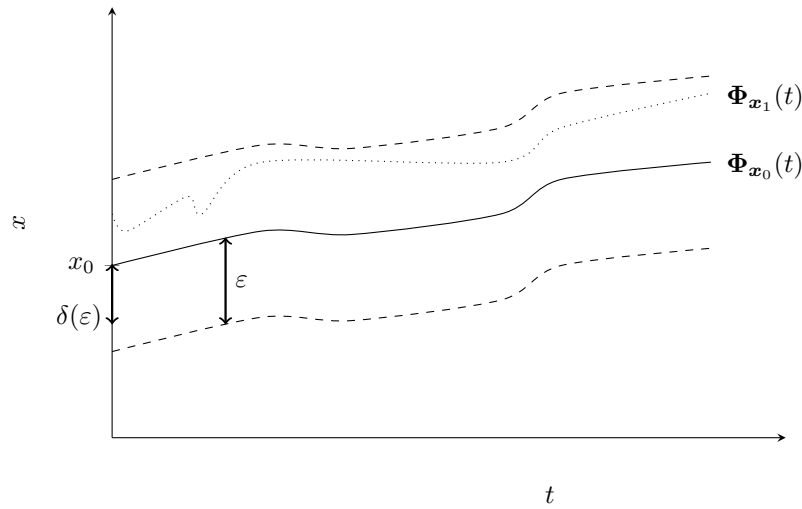


FIGURE 3.1 – Stabilité d'une trajectoire

En considérant $\mathbf{z}(t) = \Phi_{\mathbf{x}}(t) - \Phi_{\mathbf{x}_0}(t)$ et en dérivant cette expression il est possible, étant donnée la définition des trajectoires $\Phi_{\mathbf{x}}$, d'écrire l'équation (3.2.6), et donc en réutilisant la définition de $\mathbf{z}(t)$, l'équation (3.2.7). Cette équation peut alors s'écrire sous la forme (3.2.8) où la fonction w dépend de $\mathbf{x}_0$ et de $\mathbf{u}_1(t)$:

$$\dot{\mathbf{z}}(t) = f(\Phi_{\mathbf{x}}(t), \mathbf{u}_1(t), t) - f(\Phi_{\mathbf{x}_0}(t), \mathbf{u}_1(t), t) \quad (3.2.6)$$

$$\dot{\mathbf{z}}(t) = f(\mathbf{z}(t) + \Phi_{\mathbf{x}_0}(t), \mathbf{u}_1(t), t) - f(\Phi_{\mathbf{x}_0}(t), \mathbf{u}_1(t), t) \quad (3.2.7)$$

$$\dot{\mathbf{z}}(t) = w(\mathbf{z}(t), t) \quad (3.2.8)$$

Il est aisé de vérifier que le point $\mathbf{z} = 0$ est un point d'équilibre de ce nouveau système. Il est alors possible de voir que la condition (3.2.5) pour la trajectoire du système initial est équivalente à la condition (3.2.2) pour ce point d'équilibre du nouveau système dynamique. Il sera donc possible d'étudier la stabilité d'une trajectoire en se ramenant à l'étude de la stabilité d'un point d'équilibre.

Une dernière notion de stabilité, utile en traitement du signal, est la notion de *stabilité entrée bornée sortie bornée (EBSB)* qui exprime le fait que la norme de la sortie du système est bornée pour tout signal d'entrée de norme bornée.

Définition 3.2.7

Un système défini par :

$$\dot{\mathbf{x}}(t) = f(\mathbf{x}(t), \mathbf{u}(t), t)$$

$$\mathbf{y}(t) = g(\mathbf{x}(t), \mathbf{u}(t), t)$$

5. Dans cette partie nous considérons, sans perte de généralité, que l'instant initial est l'instant $t = 0$.

est *stable entrée bornée sortie bornée (EBSB)* par le fait que quel que soit le couple état-instant initial $(\mathbf{x}_0, t_0)$, la condition (3.2.9), où N est un nombre fini ne dépendant pas de t_0 , est vérifiée :

$$\|\mathbf{u}(t)\| < M < \infty, \forall t > t_0 \Rightarrow \|\mathbf{y}(t)\| < N(M, \mathbf{x}_0) \quad (3.2.9)$$

Comme on pourra le voir par la suite ces différentes notions de stabilité peuvent parfois se rejoindre lorsqu'on s'intéresse à des classes particulières de systèmes mais dans le cas le plus général elles correspondent à des propriétés différentes.

3.3 Linéarisation autour d'un point d'équilibre

3.3.1 Théorème de superposition

Une propriété fondamentale issue de la linéarité du système est le théorème de superposition qui permet de considérer la réaction du système à un signal d'entrée comme la somme de réactions élémentaires.

Considérons un système décrit par les équations (3.1.10)⁶ dont l'état initial à l'instant t_0 est considéré comme étant le vecteur nul ($\mathbf{x}(t_0) = 0$). Soient $\Phi_1(t)$ la solution de l'équation dynamique définie par ces conditions initiales et le signal d'entrée $\mathbf{u}_1(t)$ et $\Phi_2(t)$ celle définie par ces mêmes conditions et le signal $\mathbf{u}_2(t)$.

Soit $\Phi(t)$ la solution de l'équation dynamique définie par ces conditions initiales et le signal d'entrée $\mathbf{u}(t) = \mathbf{u}_1(t) + \mathbf{u}_2(t)$. On a $\Phi(t) = \Phi_1(t) + \Phi_2(t)$ puisque la solution est unique et que $\Phi_1(t) + \Phi_2(t)$ est une solution. Si on s'intéresse aux signaux de sortie, il est alors facile de voir que la sortie correspondant à $\mathbf{u}(t) = \mathbf{u}_1(t) + \mathbf{u}_2(t)$ est la somme de la sortie correspondant à $\mathbf{u}_1(t)$ et de celle correspondant à $\mathbf{u}_2(t)$.

Par un raisonnement similaire, si nous considérons le même système mais que nous considérons maintenant que le signal d'entrée est toujours le vecteur nul il est aisé de vérifier que la réponse du système définie par la condition initiale $\mathbf{x}(t_0) = \mathbf{x}_1 + \mathbf{x}_2$ est la somme des réponses définies par $\mathbf{x}(t_0) = \mathbf{x}_1$ et $\mathbf{x}(t_0) = \mathbf{x}_2$.

Enfin si $\Phi_1(t)$ est la solution définie par $\mathbf{x}(t_0) = 0$ et $\mathbf{u}(t) = \mathbf{u}_1(t)$ et $\Phi_2(t)$ celle définie par $\mathbf{x}(t_0) = \mathbf{x}_0$ et $\mathbf{u}(t) = 0$ il est possible de vérifier que $\Phi(t) = \Phi_1(t) + \Phi_2(t)$ est la solution définie par $\mathbf{x}(t_0) = \mathbf{x}_0$ et $\mathbf{u}(t) = \mathbf{u}_1(t)$. La réponse du système à une condition initiale et un signal d'entrée est donc la somme de la réponse aux conditions initiales seules (entrée nulle) et de la réponse au signal d'entrée seul (condition initiale nulle).

Une conséquence importante de ce théorème de superposition est la possibilité d'étudier les réactions des systèmes à plusieurs entrées comme la somme de systèmes à une entrée (figure 3.2). En effet tout signal vectoriel à k composantes peut être considéré, par décomposition dans une base de l'espace de dimension k , comme la somme de k signaux dont une seule composante est non nulle. Les réactions à ces signaux à une composante définissent alors autant de systèmes à une entrée.

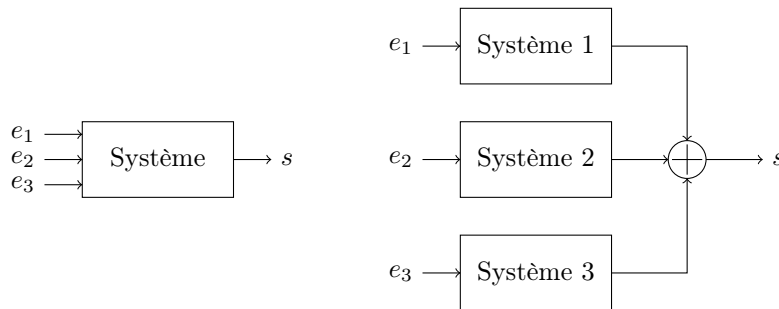


FIGURE 3.2 – Système multi-entrées

6. Encore une fois le cas des systèmes discrets sera considéré par adaptation du cas continu.

Dans la suite de ce cours les systèmes considérés seront en général des systèmes à une entrée puisque, pour les systèmes linéaires, les propriétés des systèmes à plusieurs entrées se déduisent facilement de celles des systèmes à une entrée définis par chacune des entrées.

3.3.2 Linéarisation

Lorsque le système ne respecte pas les contraintes (3.1.6) et (3.1.7), il n'est pas linéaire, son étude est alors plus compliquée et demande l'utilisation d'outils mathématiques plus complexes. Cependant si on n'a pas besoin d'étudier le système sur l'ensemble de l'espace d'état mais seulement au voisinage d'une position d'équilibre il est possible d'utiliser un modèle linéaire tangent au modèle non-linéaire en ce point d'équilibre.

Par exemple si on considère un satellite géo-stationnaire on n'a pas besoin d'étudier le comportement du système, défini par l'équation (2.3.5), dans tout l'espace d'état mais seulement dans un voisinage du point d'équilibre

$$\left(\left(\sqrt[3]{\frac{GM}{\bar{\omega}^2}} \quad 0 \quad \bar{\omega} \right)^T \quad (0 \quad 0)^T \right)$$

qui définit sa trajectoire nominale à la vitesse orbitale $\bar{\omega}$.

Pour étudier les systèmes non linéaires on effectuera donc souvent une linéarisation autour d'un point d'équilibre en considérant le modèle linéaire tangent au point d'équilibre. L'idée de base est très simple à illustrer pour une fonction d'une variable (figure 3.3), la tangente à la courbe $(x, f(x))$ au point $(\bar{x}, f(\bar{x}))$ est alors la droite d'équation :

$$y - f(\bar{x}) = f'(\bar{x})(x - \bar{x})$$

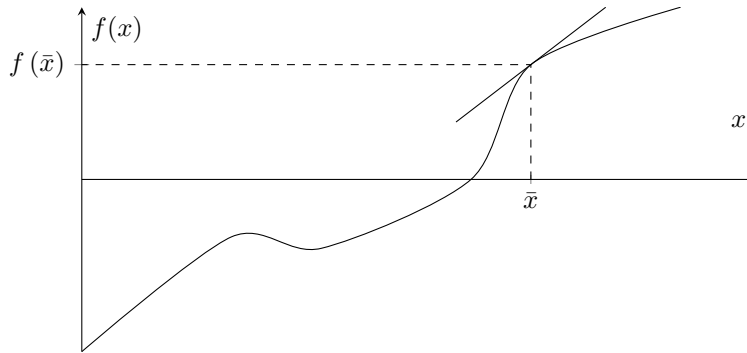


FIGURE 3.3 – Linéarisé tangent

Pour une fonction réelle de plusieurs variables, c'est-à-dire de $\mathbb{R}^n$ dans $\mathbb{R}$, ceci se généralise par la formule suivante (voir aussi la section de rappels de calcul différentiel F) :

$$\begin{aligned} y - f(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n) &= \frac{\partial f}{\partial x_1}(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n) \cdot (x_1 - \bar{x}_1) \\ &\quad + \frac{\partial f}{\partial x_2}(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n) \cdot (x_2 - \bar{x}_2) \\ &\quad + \dots + \frac{\partial f}{\partial x_n}(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n) \cdot (x_n - \bar{x}_n) \end{aligned} \quad (3.3.1)$$

Si on s'intéresse à un système dynamique défini par les équations (3.1.1) et (3.1.2) au voisinage du point d'équilibre $(\bar{x}, \bar{u})$, il est possible d'appliquer la formule (3.3.1) pour chacune des composantes du vecteur dérivé et du vecteur de sortie et on peut écrire, en faisant intervenir les *matrices*

jacobiennes :

$$\dot{\mathbf{x}}(t) = \begin{pmatrix} \frac{\partial f_1}{\partial x_1}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) & \cdots & \frac{\partial f_1}{\partial x_n}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) \\ \vdots & \ddots & \vdots \\ \frac{\partial f_n}{\partial x_1}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) & \cdots & \frac{\partial f_n}{\partial x_n}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) \end{pmatrix} (\mathbf{x}(t) - \bar{\mathbf{x}}) \quad (3.3.2)$$

$$+ \begin{pmatrix} \frac{\partial f_1}{\partial u_1}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) & \cdots & \frac{\partial f_1}{\partial u_m}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) \\ \vdots & \ddots & \vdots \\ \frac{\partial f_n}{\partial u_1}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) & \cdots & \frac{\partial f_n}{\partial u_m}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) \end{pmatrix} (\mathbf{u}(t) - \bar{\mathbf{u}})$$

$$\mathbf{y}(t) - \bar{\mathbf{y}} = \begin{pmatrix} \frac{\partial g_1}{\partial x_1}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) & \cdots & \frac{\partial g_1}{\partial x_n}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) \\ \vdots & \ddots & \vdots \\ \frac{\partial g_p}{\partial x_1}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) & \cdots & \frac{\partial g_p}{\partial x_n}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) \end{pmatrix} (\mathbf{x}(t) - \bar{\mathbf{x}}) \quad (3.3.3)$$

$$+ \begin{pmatrix} \frac{\partial g_1}{\partial u_1}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) & \cdots & \frac{\partial g_1}{\partial u_m}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) \\ \vdots & \ddots & \vdots \\ \frac{\partial g_p}{\partial u_1}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) & \cdots & \frac{\partial g_p}{\partial u_m}(\bar{\mathbf{x}}, \bar{\mathbf{u}}, t) \end{pmatrix} (\mathbf{u}(t) - \bar{\mathbf{u}})$$

Si on introduit les nouvelles variables $\tilde{\mathbf{x}}(t) = \mathbf{x}(t) - \bar{\mathbf{x}}$, $\tilde{\mathbf{u}}(t) = \mathbf{u}(t) - \bar{\mathbf{u}}$, $\tilde{\mathbf{y}}(t) = \mathbf{y}(t) - \bar{\mathbf{y}}$, on voit que ces équations s'écrivent :

$$\begin{aligned} \dot{\tilde{\mathbf{x}}}(t) &= \mathbf{A}(t)\tilde{\mathbf{x}}(t) + \mathbf{B}(t)\tilde{\mathbf{u}}(t) \\ \tilde{\mathbf{y}}(t) &= \mathbf{C}(t)\tilde{\mathbf{x}}(t) + \mathbf{D}(t)\tilde{\mathbf{u}}(t) \end{aligned} \quad (3.3.4)$$

Ces équations définissent bien un système linéaire tangent au système d'origine dans un voisinage du point d'équilibre. L'étude de ce système permettra la mise en évidence de propriétés locales du système d'origine. On a notamment le résultat suivant.

Remarque 3.3.1. *Le point d'équilibre $(\mathbf{x}(t), \mathbf{u}(t)) = (\bar{\mathbf{x}}, \bar{\mathbf{u}})$ du système non linéaire est asymptotiquement (respectivement exponentiellement) stable si et seulement si le point d'équilibre $(\tilde{\mathbf{x}}(t), \tilde{\mathbf{u}}(t)) = (\mathbf{0}, \mathbf{0})$ du système linéaire tangent en $(\bar{\mathbf{x}}, \bar{\mathbf{u}})$ est asymptotiquement (respectivement exponentiellement) stable.*

L'intérêt de ce résultat est qu'on dispose de méthodes simples et efficaces pour tester la stabilité dans le cas d'un système linéaire.

Le cas des systèmes discrets se déduit facilement de cette étude en procédant de façon similaire, puisque les différenciations qui sont effectuées pour obtenir le modèle linéaire sont faites par rapport aux composantes des vecteurs d'état et d'entrée et non par rapport au temps.

Exemple 3.3.1 (Modèle de satellite (suite)). *Reprenons l'exemple du satellite modélisé par l'équation (2.3.5) et intéressons nous à son comportement autour de la trajectoire définie par le point d'équilibre vu précédemment $\left((\bar{\alpha} \ 0 \ \bar{\omega})^T \ (0 \ 0)^T \right)$, en notant $\bar{\alpha} = \sqrt[3]{\frac{GM}{\bar{\omega}^2}}$:*

$$\dot{\mathbf{x}}(t) = \begin{pmatrix} x_2(t) \\ x_1(t)x_3(t)^2 - \frac{GM}{x_1(t)^2} + \frac{1}{m}u_1(t) \\ -\frac{2x_2(t)x_3(t)}{x_1(t)} + \frac{1}{mx_1(t)}u_2(t) \end{pmatrix} \quad (3.3.5)$$

En utilisant les conventions précédentes le modèle tangent à ce point d'équilibre est alors le suivant :

$$\dot{\tilde{\mathbf{x}}}(t) = \begin{pmatrix} 0 & 1 & 0 \\ 3\bar{\omega}^2 & 0 & 2\bar{\alpha}\bar{\omega} \\ 0 & -\frac{2\bar{\omega}}{\bar{\alpha}} & 0 \end{pmatrix} \tilde{\mathbf{x}}(t) + \begin{pmatrix} 0 & 0 \\ \frac{1}{m} & 0 \\ 0 & \frac{1}{m\bar{\alpha}} \end{pmatrix} \tilde{\mathbf{u}}(t) \quad (3.3.6)$$

3.4 Résolution des équations d'état d'un système linéaire

On se propose de résoudre la première équation d'état, qui est une équation différentielle matricielle du premier ordre, en deux temps. On étudie tout d'abord l'équation sans second membre (régime libre) puis on cherche une solution particulière de l'équation avec second membre (régime forcé).

3.4.1 Régime libre

Il s'agit de résoudre :

$$\frac{d\mathbf{x}(t)}{dt} = \mathbf{A}(t)\mathbf{x}(t)$$

avec la condition initiale $\mathbf{x}(t_0) = \mathbf{x}_0$. Cette équation étant linéaire ainsi que le système, et étant donnée la définition des variables d'état, la solution s'exprime linéairement en fonction de $\mathbf{x}_0$ sous la forme :

$$\boxed{\mathbf{x}(t) = \Phi(t, t_0) \mathbf{x}_0} \quad (3.4.1)$$

$\Phi(t, t_0)$ s'appelle la *matrice de transition* du système ; elle est de dimension $n \times n$.

Cette matrice de transition possède les propriétés suivantes :

$$\Phi(t, t) = \mathbf{I} \quad (3.4.2)$$

$$\Phi(t_2, t_0) = \Phi(t_2, t_1) \Phi(t_1, t_0) \quad (3.4.3)$$

$$\Phi(t_2, t_1) = (\Phi(t_1, t_2))^{-1} \quad (3.4.4)$$

$$\frac{d\Phi}{dt}(t, t_0) = \mathbf{A}(t)\Phi(t, t_0) \quad (3.4.5)$$

3.4.2 Régime forcé

Il s'agit de résoudre :

$$\frac{d\mathbf{x}(t)}{dt} = \mathbf{A}(t)\mathbf{x}(t) + \mathbf{B}(t)\mathbf{u}(t) \quad (3.4.6)$$

connaissant la solution générale de l'équation homogène établie ci-dessus. On cherche une solution particulière $\mathbf{x}_p(t)$ avec, par exemple, la méthode de variation de la constante :

$$\begin{aligned} \mathbf{x}_p(t) = \Phi(t, t_0) \mathbf{x}_0(t) &\Rightarrow \frac{d\mathbf{x}_p(t)}{dt} = \frac{d\Phi}{dt}(t, t_0) \mathbf{x}_0(t) + \Phi(t, t_0) \frac{d\mathbf{x}_0}{dt}(t) \\ &= \mathbf{A}(t)\Phi(t, t_0) \mathbf{x}_0(t) + \Phi(t, t_0) \frac{d\mathbf{x}_0}{dt}(t) \end{aligned}$$

d'où en utilisant (3.4.6) et (3.4.5) :

$$\Phi(t, t_0) \frac{d\mathbf{x}_0}{dt}(t) = \mathbf{B}(t)\mathbf{u}(t) \Leftrightarrow \frac{d\mathbf{x}_0}{dt}(t) = (\Phi(t, t_0))^{-1} \mathbf{B}(t)\mathbf{u}(t)$$

Finalement :

$$\mathbf{x}_0(t) = \int_{t_0}^t (\Phi(\tau, t_0))^{-1} \mathbf{B}(\tau) \mathbf{u}(\tau) d\tau$$

et, compte tenu des propriétés de Φ on a :

$$\mathbf{x}_p(t) = \Phi(t, t_0) \mathbf{x}_0(t) = \Phi(t, t_0) \int_{t_0}^t (\Phi(\tau, t_0))^{-1} \mathbf{B}(\tau) \mathbf{u}(\tau) d\tau = \int_{t_0}^t \Phi(t, \tau) \mathbf{B}(\tau) \mathbf{u}(\tau) d\tau$$

La solution de la première équation d'état est donc :

$$\boxed{\mathbf{x}(t) = \Phi(t, t_0) \mathbf{x}(t_0) + \int_{t_0}^t \Phi(t, \tau) \mathbf{B}(\tau) \mathbf{u}(\tau) d\tau} \quad (3.4.7)$$

La solution de la deuxième équation d'état est donc :

$$\boxed{\mathbf{y}(t) = \mathbf{C}(t) \left[\Phi(t, t_0) \mathbf{x}(t_0) + \int_{t_0}^t \Phi(t, \tau) \mathbf{B}(\tau) \mathbf{u}(\tau) d\tau \right] + \mathbf{D}(t) \mathbf{u}(t)} \quad (3.4.8)$$

La difficulté de cette résolution réside dans le calcul de la matrice de transition $\Phi(t, \tau)$ qui fait appel à des techniques de résolution numérique sauf dans le cas où le système est linéaire ET invariant auquel cas cette matrice peut se calculer analytiquement. Nous allons voir plusieurs méthodes de calcul de $\Phi(t, \tau)$ dans ce cas.

3.4.3 Cas particulier des systèmes linéaires et invariants

Lorsque la matrice $\mathbf{A}$ est indépendante du temps, l'équation :

$$\frac{d\Phi}{dt}(t, t_0) = \mathbf{A}\Phi(t, t_0)$$

se résout par analogie avec les cas scalaire en recherchant une solution sous la forme de la somme d'une série entière matricielle :

$$\begin{aligned} \Phi(t, t_0) &= \mathbf{A}_0 + \mathbf{A}_1(t - t_0) + \mathbf{A}_2(t - t_0)^2 + \dots + \mathbf{A}_k(t - t_0)^k + \dots \\ \Rightarrow \frac{d\Phi}{dt}(t, t_0) &= \mathbf{0} + \mathbf{A}_1 + 2\mathbf{A}_2(t - t_0) + \dots + k\mathbf{A}_k(t - t_0)^{k-1} + \dots \end{aligned}$$

Or :

$$\frac{d\Phi}{dt}(t, t_0) = \mathbf{A}\Phi(t, t_0) = \mathbf{A} \left[\mathbf{A}_0 + \mathbf{A}_1(t - t_0) + \mathbf{A}_2(t - t_0)^2 + \dots + \mathbf{A}_k(t - t_0)^k + \dots \right]$$

d'où en identifiant les deux développements on obtient :

$$\mathbf{A}_1 = \mathbf{A}\mathbf{A}_0; \mathbf{A}_2 = \frac{1}{2}\mathbf{A}\mathbf{A}_1; \dots; \mathbf{A}_k = \frac{1}{k}\mathbf{A}\mathbf{A}_{k-1}$$

La propriété $\Phi(t, t) = \mathbf{I}$ nous donne $\mathbf{A}_0 = \mathbf{I}$. On aboutit alors au développement suivant pour $\Phi(t, t_0)$:

$$\Phi(t, t_0) = \mathbf{I} + \mathbf{A}(t - t_0) + \frac{1}{2!}\mathbf{A}^2(t - t_0)^2 + \dots + \frac{1}{k!}\mathbf{A}^k(t - t_0)^k + \dots$$

Ce développement matriciel rappelle celui de l'exponentielle dans le cas scalaire ; pour cette raison on note de façon symbolique :

$$\Phi(t, t_0) = e^{\mathbf{A}(t-t_0)}$$

La solution de la première équation d'état s'écrit alors :

$$\boxed{\mathbf{x}(t) = e^{\mathbf{A}(t-t_0)} \mathbf{x}(t_0) + \int_{t_0}^t e^{\mathbf{A}(t-\tau)} \mathbf{B}(\tau) \mathbf{u}(\tau) d\tau} \quad (3.4.9)$$

3.4.4 Méthodes de calcul de l'exponentielle matricielle

3.4.4.1 Utilisation du développement en série

Cette méthode consiste à calculer directement le développement en série vu ci-dessus en calculant les différentes puissances de $\mathbf{A}$. On ne l'utilise que si le calcul peut se faire simplement.

Exemple 3.4.1.

$$\mathbf{A} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \Rightarrow \mathbf{A}^2 = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} = \mathbf{A}^k \text{ pour } k \geq 2 \Rightarrow e^{\mathbf{A}(t-t_0)} = \mathbf{I} + \mathbf{A}(t-t_0) = \begin{pmatrix} 1 & t-t_0 \\ 0 & 1 \end{pmatrix}$$

$$\mathbf{A} = \begin{pmatrix} 0 & 1 \\ 0 & -a \end{pmatrix} \Rightarrow \mathbf{A}^2 = \begin{pmatrix} 0 & -a \\ 0 & a^2 \end{pmatrix} \\ \Rightarrow \dots \mathbf{A}^k = \begin{pmatrix} 0 & (-a)^{k-1} \\ 0 & (-a)^k \end{pmatrix} \Rightarrow e^{\mathbf{A}(t-t_0)} = \begin{pmatrix} 1 & \frac{1 - e^{-a(t-t_0)}}{a} \\ 0 & e^{-a(t-t_0)} \end{pmatrix}$$

3.4.4.2 Méthode des modes

Cette méthode consiste à effectuer un changement de vecteur d'état (changement de base) de sorte que la nouvelle matrice d'évolution $\mathbf{A}'$ soit diagonale ; en d'autres termes cette méthode consiste à diagonaliser la matrice $\mathbf{A}$.

$$\mathbf{A}' = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix} \Rightarrow \Phi(t, t_0) = e^{\mathbf{A}'(t-t_0)} = \begin{pmatrix} e^{\lambda_1(t-t_0)} & 0 & \dots & 0 \\ 0 & e^{\lambda_2(t-t_0)} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & e^{\lambda_n(t-t_0)} \end{pmatrix}$$

Exemple 3.4.2. $\mathbf{A} = \begin{pmatrix} 0 & 1 \\ 0 & -a \end{pmatrix}$ de valeurs propres $\lambda_1 = 0$ et $\lambda_2 = -a$ d'où :

$$\mathbf{A}' = \begin{pmatrix} 0 & 0 \\ 0 & -a \end{pmatrix} \Rightarrow e^{\mathbf{A}'(t-t_0)} = \begin{pmatrix} 1 & 0 \\ 0 & e^{-a(t-t_0)} \end{pmatrix}$$

On cherche maintenant à déterminer $\mathbf{M}$. On a par exemple :

$$\mathbf{A}\mathbf{x}_1 = \lambda_1\mathbf{x}_1 = 0 \Rightarrow \mathbf{x}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$$

et :

$$\mathbf{A}\mathbf{x}_2 = \lambda_2\mathbf{x}_2 = -a\mathbf{x}_2 \Rightarrow \mathbf{x}_2 = \begin{pmatrix} 1 \\ -a \end{pmatrix}$$

d'où :

$$\mathbf{M} = \begin{pmatrix} 1 & 1 \\ 0 & -a \end{pmatrix} \text{ et } \mathbf{M}^{-1} = \begin{pmatrix} 1 & \frac{1}{a} \\ 0 & -\frac{1}{a} \end{pmatrix}$$

soit finalement :

$$e^{\mathbf{A}(t-t_0)} = \begin{pmatrix} 1 & \frac{1 - e^{-a(t-t_0)}}{a} \\ 0 & e^{-a(t-t_0)} \end{pmatrix}$$

3.4.4.3 Formule de Sylvester

Cette méthode s'applique lorsque les valeurs propres de la matrice d'évolution $\mathbf{A}$ sont toutes **distinctes**. On a alors :

$$e^{\mathbf{A}(t-t_0)} = \sum_{i=1}^n e^{\lambda_i(t-t_0)} \left[\prod_{j=1, j \neq i}^n \frac{\mathbf{A} - \lambda_j \mathbf{I}}{\lambda_i - \lambda_j} \right] \quad (3.4.10)$$

Exemple 3.4.3. $\mathbf{A} = \begin{pmatrix} 0 & 1 \\ 0 & -a \end{pmatrix}$ de valeurs propres $\lambda_1 = 0$ et $\lambda_2 = -a$:

$$e^{\mathbf{A}(t-t_0)} = \frac{\mathbf{A} + a\mathbf{I}}{a} + e^{-a(t-t_0)} \frac{\mathbf{A}}{-a} = \frac{1 - e^{-a(t-t_0)}}{a} \mathbf{A} + \mathbf{I}$$

d'où :

$$e^{\mathbf{A}(t-t_0)} = \begin{pmatrix} 1 & \frac{1 - e^{-a(t-t_0)}}{a} \\ 0 & e^{-a(t-t_0)} \end{pmatrix}$$

3.4.4.4 Interpolation de Sylvester

On utilise cette méthode notamment lorsque les valeurs propres de $\mathbf{A}$ ne sont pas toutes distinctes. Le processus consiste à remplir un tableau de n lignes et $n + 1$ colonnes comme suit :

— à chaque valeur propre simple correspond une ligne :

$$1 \quad \lambda_1 \quad \lambda_1^2 \quad \dots \quad \lambda_1^{n-1} \quad e^{\lambda_1(t-t_0)}$$

— à chaque valeur propre double correspond deux lignes ; la première est identique à celle des valeurs propres simples, la deuxième est la dérivée de la première :

$$\begin{array}{cccccc} 1 & \lambda_2 & \lambda_2^2 & \dots & \lambda_2^{n-1} & e^{\lambda_2(t-t_0)} \\ 0 & 1 & 2\lambda_2 & \dots & (n-1)\lambda_2^{n-2} & (t-t_0)e^{\lambda_2(t-t_0)} \end{array}$$

— à chaque valeur propre triple correspond trois lignes ; la première est identique à celle des valeurs propres simples, la deuxième est la dérivée de la première et la troisième est la dérivée de la deuxième :

$$\begin{array}{cccccc} 1 & \lambda_3 & \lambda_3^2 & \dots & \lambda_3^{n-1} & e^{\lambda_3(t-t_0)} \\ 0 & 1 & 2\lambda_3 & \dots & (n-1)\lambda_3^{n-2} & (t-t_0)e^{\lambda_3(t-t_0)} \\ 0 & 0 & 2 & \dots & (n-1)(n-2)\lambda_3^{n-3} & (t-t_0)^2 e^{\lambda_3(t-t_0)} \end{array}$$

et ainsi de suite pour les valeurs propres d'ordre supérieur à 3. On complète alors le tableau par une $(n+1)^{\text{ième}}$ ligne égale à :

$$\mathbf{I} \quad \mathbf{A} \quad \mathbf{A}^2 \quad \dots \quad \mathbf{A}^{n-1} \quad e^{\mathbf{A}(t-t_0)}$$

La méthode de Sylvester consiste à calculer formellement le déterminant de la matrice obtenue ; on déduit alors la matrice de transition en écrivant que ce déterminant est nul.

Exemple 3.4.4. $\mathbf{A} = \begin{pmatrix} 0 & 1 \\ 0 & -a \end{pmatrix}$ de valeurs propres $\lambda_1 = 0$ et $\lambda_2 = -a$. Le tableau est ici égal à :

$$\begin{array}{ccc} 1 & 0 & 1 \\ 1 & -a & e^{-a(t-t_0)} \\ \mathbf{I} & \mathbf{A} & e^{\mathbf{A}(t-t_0)} \end{array}$$

Son déterminant Δ développé par rapport à la dernière ligne est :

$$\begin{aligned} \Delta &= \mathbf{I} \begin{vmatrix} 0 & 1 \\ -a & e^{-a(t-t_0)} \end{vmatrix} - \mathbf{A} \begin{vmatrix} 1 & 1 \\ 1 & e^{-a(t-t_0)} \end{vmatrix} + e^{\mathbf{A}(t-t_0)} \begin{vmatrix} 1 & 0 \\ 1 & -a \end{vmatrix} \\ &= a\mathbf{I} - \mathbf{A} \left(e^{-a(t-t_0)} - 1 \right) - ae^{\mathbf{A}(t-t_0)} \end{aligned}$$

Ainsi :

$$\Delta = 0 \Rightarrow e^{\mathbf{A}(t-t_0)} = \mathbf{I} - \mathbf{A} \frac{e^{-a(t-t_0)} - 1}{a} = \begin{pmatrix} 1 & \frac{1 - e^{-a(t-t_0)}}{a} \\ 0 & e^{-a(t-t_0)} \end{pmatrix}$$

3.4.4.5 Utilisation de la transformée de Laplace

Cette sous-section fait appel à la notion de transformée de Laplace qui sera expliquée dans le chapitre suivant.

La première équation d'état est une équation différentielle qui peut se résoudre en appliquant la transformée de Laplace. On se place ici dans le cas des systèmes invariants.

$$\mathcal{L}\left\{\frac{d\mathbf{x}(t)}{dt}\right\}(p) = \mathcal{L}\{\mathbf{A}\mathbf{x}(t) + \mathbf{B}\mathbf{u}(t)\}(p) \Leftrightarrow pX(p) - \mathbf{x}(t_0)e^{-pt_0} = \mathbf{A}X(p) + \mathbf{B}U(p)$$

d'où :

$$\boxed{X(p) = (p\mathbf{I} - \mathbf{A})^{-1}(\mathbf{x}(t_0)e^{-pt_0} + \mathbf{B}U(p))} \quad (3.4.11)$$

En identifiant cette équation avec l'équation temporelle :

$$\mathbf{x}(t) = e^{\mathbf{A}(t-t_0)}\mathbf{x}(t_0) + \int_{t_0}^t e^{\mathbf{A}(t-\tau)}\mathbf{B}\mathbf{u}(\tau)d\tau$$

on déduit que $(p\mathbf{I} - \mathbf{A})^{-1}e^{-pt_0}$ est la transformée de Laplace de $e^{\mathbf{A}(t-t_0)}$.

Exemple 3.4.5.

$$\mathbf{A} = \begin{pmatrix} 0 & 1 \\ 0 & -a \end{pmatrix} \Rightarrow p\mathbf{I} - \mathbf{A} = \begin{pmatrix} p & -1 \\ 0 & p+a \end{pmatrix} \Rightarrow (p\mathbf{I} - \mathbf{A})^{-1}e^{-pt_0} = \begin{pmatrix} \frac{e^{-pt_0}}{p} & \frac{e^{-pt_0}}{p(p+a)} \\ 0 & \frac{e^{-pt_0}}{p+a} \end{pmatrix}$$

En utilisant les transformées usuelles, on en déduit :

$$e^{\mathbf{A}(t-t_0)} = \begin{pmatrix} 1 & \frac{1 - e^{-a(t-t_0)}}{a} \\ 0 & e^{-a(t-t_0)} \end{pmatrix}$$

3.5 Stabilité asymptotique

Pour un système linéaire, toute trajectoire est stable si et seulement si le point d'équilibre $x = 0$ de $\dot{x}(t) = \mathbf{A}x(t)$ est stable. D'après le tableau de Sylvester, chaque terme de $e^{\mathbf{A}t}$ est de la forme :

$$\sum_{k=1}^n p_k(t)e^{\lambda_k t} \quad (3.5.1)$$

où $p_k(t) = \sum_{i=0}^{n_k-1} \alpha_i t^i$ est un polynôme de degré $n_k - 1$, avec n_k l'ordre de multiplicité de λ_k .

On en déduit les conditions pour :

1. la *stabilité asymptotique* : $x(t) = e^{\mathbf{A}(t-t_0)}x(t_0)$ tend vers 0, $\forall x(t_0)$ si et seulement si $\forall k$, $e^{\lambda_k t} \rightarrow 0$, ce qui est équivalent à $\forall k$, $\text{Re}(\lambda_k) < 0$
2. la *stabilité au sens de Lyapunov* : $x(t) = e^{\mathbf{A}(t-t_0)}x(t_0)$ reste borné $\forall x(t_0)$ si $\forall k$, $\text{Re}(\lambda_k) < 0$ ou $\text{Re}(\lambda_k) = 0$ avec $n_k = 1$. En effet, pour une multiplicité $n_k = 2$ avec $\lambda_k = j\omega_k$, le polynôme $p_k(t)e^{\lambda_k t} = (\alpha_0 + \alpha_1 t)e^{j\omega_k t}$ ne sera pas borné.

On remarque aussi que s'il y a convergence asymptotique, elle est aussi exponentielle.

Remarque 3.5.1. En résumé, un système continu linéaire invariant dans le temps est :

1. *stable asymptotiquement et exponentiellement si et seulement si les valeurs propres de la matrice d'évolution $\mathbf{A}$ ont toutes une partie réelle strictement négative ;*
2. *stable au sens de Lyapunov si la matrice d'évolution $\mathbf{A}$ a uniquement*
 - *des valeurs propres avec la partie réelle strictement négative ;*
 - *les valeurs propres à partie réelle nulle ont une multiplicité égale à 1 ;*
3. *instable dans les autres cas.*

En pratique, cette condition pourra être testée entre autres en calculant le polynôme caractéristique $\det(\lambda\mathbf{I} - \mathbf{A})$ de la matrice $\mathbf{A}$, puis en appliquant le critère de Routh (paragraphe 4.4.2.3.1).

3.6 Analyse des systèmes discrets

3.6.1 Modélisation des systèmes linéaires discrets par variables d'état

Reprenons l'équation donnant l'évolution de l'état d'un système continu entre deux instants t_0 et t :

$$\mathbf{x}(t) = \Phi(t, t_0) \mathbf{x}(t_0) + \int_{t_0}^t \Phi(t, \tau) \mathbf{B}(\tau) \mathbf{u}(\tau) d\tau$$

et appliquons cette équation entre deux instants t_k et t_{k+1} en supposant l'entrée $u(t)$ constante (égale à $u(t_k)$) entre ces deux instants. L'évolution du système entre t_k et t_{k+1} s'exprime alors par :

$$\mathbf{x}(t_{k+1}) = \Phi(t_{k+1}, t_k) \mathbf{x}(t_k) + \left(\int_{t_k}^{t_{k+1}} \Phi(t_{k+1}, \tau) \mathbf{B}(\tau) d\tau \right) \mathbf{u}(t_k)$$

En posant :

$$\mathbf{F}(k) = \Phi(t_{k+1}, t_k) \text{ et } \mathbf{G}(k) = \int_{t_k}^{t_{k+1}} \Phi(t_{k+1}, \tau) \mathbf{B}(\tau) d\tau$$

on obtient :

$$\mathbf{x}(t_{k+1}) = \mathbf{F}(k) \mathbf{x}(t_k) + \mathbf{G}(k) \mathbf{u}(t_k)$$

Si le système est permanent et que les instants t_k sont multiples d'une période d'échantillonnage T , on obtient :

$$\boxed{\mathbf{F}(k) = \Phi(t_{k+1}, t_k) = e^{\mathbf{A}(t_{k+1}-t_k)} = e^{\mathbf{A}T}} \quad (3.6.1)$$

$$\boxed{\begin{aligned} \mathbf{G}(k) &= \int_{t_k}^{t_{k+1}} \Phi(t_{k+1}, \tau) \mathbf{B}(\tau) d\tau = \int_0^T e^{\mathbf{A}(T-\theta)} \mathbf{B} d\theta = \int_0^T e^{\mathbf{A}u} \mathbf{B} du \\ \mathbf{G}(k) &= \mathbf{A}^{-1} (e^{\mathbf{A}T} - \mathbf{I}) \mathbf{B} \text{ si } \mathbf{A} \text{ est inversible} \end{aligned}} \quad (3.6.2)$$

$\mathbf{F}(k) = \mathbf{F}$ et $\mathbf{G}(k) = \mathbf{G}$ sont alors deux matrices constantes et les équations d'état du système sont :

$$\begin{cases} \mathbf{x}(k+1) &= \mathbf{F} \mathbf{x}(k) + \mathbf{G} \mathbf{u}(k) \\ \mathbf{y}(k) &= \mathbf{C} \mathbf{x}(k) + \mathbf{D} \mathbf{u}(k) \end{cases} \quad (3.6.3)$$

REMARQUE : Comme dans le cas scalaire, si T est "petit", on peut faire l'approximation :

$$e^{\mathbf{A}T} \approx \mathbf{I} + \mathbf{A}T$$

On obtient alors comme première équation d'état :

$$\mathbf{x}(k+1) \approx (\mathbf{I} + \mathbf{A}T) \mathbf{x}(k) + \mathbf{B}T \mathbf{u}(k)$$

c'est-à-dire :

$$\frac{\mathbf{x}(k+1) - \mathbf{x}(k)}{T} \approx \mathbf{A} \mathbf{x}(k) + \mathbf{B} \mathbf{u}(k)$$

ce qui revient à dire que l'on a fait l'approximation classique (d'autant plus justifiée que T est petit) suivante :

$$\left(\frac{d\mathbf{x}}{dt} \right)_{t=kT} \approx \frac{\mathbf{x}(k+1) - \mathbf{x}(k)}{T}$$

3.6.2 Résolution des équations d'état

La résolution de la première équation dans le domaine temporel se fait en procédant par récurrence :

$$\mathbf{x}(n+1) = \mathbf{F}^{n+1} \mathbf{x}(0) + \sum_{k=0}^n \mathbf{F}^{n-k} \mathbf{G} \mathbf{u}(k) \quad (3.6.4)$$

Il est en général plus intéressant d'utiliser la *transformée en z* (qui sera détaillée dans le chapitre suivant) pour résoudre la première équation d'état afin d'éviter les multiplications de matrices :

$$\mathcal{Z}\{\mathbf{x}(k+1)\}(z) = \mathcal{Z}\{\mathbf{F}\mathbf{x}(k) + \mathbf{G}\mathbf{u}(k)\}(z) \Leftrightarrow zX(z) = \mathbf{F}X(z) + \mathbf{G}U(z)$$

d'où :

$$\boxed{X(z) = (z\mathbf{I} - \mathbf{F})^{-1} \mathbf{G}U(z)} \quad (3.6.5)$$

Il suffit ensuite de revenir dans l'espace des temps pour obtenir $\mathbf{x}(n)$ afin d'avoir accès à $\mathbf{y}(n)$ par la deuxième équation d'état.

Remarque 3.6.1. *Le calcul de $(z\mathbf{I} - \mathbf{F})^{-1}$ peut se faire en utilisant l'algorithme de Leverrier-Souriau.*

3.6.3 Stabilité asymptotique

En considérant la relation (3.6.4) pour une entrée $\mathbf{u}(k)$ nulle, on établit facilement que $\mathbf{x}(n)$ tend vers 0 pour toute condition initiale $\mathbf{x}(0)$ si et seulement si toutes les valeurs propres λ_i de $\mathbf{F}$ vérifient $|\lambda_i| < 1$.

Remarque 3.6.2. *En résumé, un système discret linéaire invariant dans le temps est :*

1. *stable asymptotiquement et exponentiellement si et seulement si les valeurs propres de la matrice d'évolution $\mathbf{F}$ ont toutes un module strictement inférieur à 1 ;*
2. *stable au sens de Lyapunov si la matrice d'évolution $\mathbf{F}$ a uniquement*
 - *des valeurs propres de module strictement inférieur à 1 ;*
 - *les valeurs propres de module 1 ont une multiplicité égale à 1 ;*
3. *instable dans les autres cas.*

Cette considération peut être testée en appliquant le *critère de Jury* (paragraphe 4.4.2.3.2) au polynôme caractéristique de $\mathbf{F}$.

Chapitre 4

Modélisation sous forme de fonction de transfert

4.1 Introduction

Ce chapitre présente les méthodes qui peuvent être mises en œuvre dans le cas de l'analyse de systèmes linéaires invariants dans le temps (stationnaires). Nous allons voir comment mobiliser les outils mathématiques (que nous introduirons) que sont la transformée de Laplace (pour les systèmes à temps continu) et la transformée en z (pour les systèmes à temps discret) pour analyser un système linéaire et stationnaire à partir d'une représentation externe de type entrée/sortie. Des approches permettant le passage entre les deux représentations externe (fonction de transfert) et interne (représentation d'état) sont également proposées.

Une partie des méthodes que nous allons présenter dans ce cours s'appuient sur deux notions mathématiques importantes :

- la transformée de Laplace, qui sera utile pour les systèmes à temps continus ;
- la transformée en z , utile pour les systèmes à temps discrets.

Ces deux opérateurs héritent de formalismes et propriétés définis pour la transformée de Fourier et sa version discrète, que l'on commencera donc par introduire.

4.2 Notions préliminaires : transformation de Fourier, de Laplace, en z

4.2.1 Notions de base sur la transformation de Fourier

La transformation de Fourier est un outil incontournable pour caractériser un signal en faisant le lien entre le domaine temporel et domaine fréquentiel. En vue de son utilisation dans ce cours, nous en introduisons ici la définition sur l'espace $L^1(\mathbb{R})$ des fonctions sommables et quelques unes de ses propriétés, qui seront explicitées dans le cours de *Convergence, intégration, probabilités, équations aux dérivées partielles* et étendues à d'autres espaces fonctionnels, ce qui permettra de traiter des classes plus larges de signaux.

Définition 4.2.1

Dans ce cours, nous définissons la *transformée de Fourier* d'une fonction sommable¹ par :

$$X(f) = \mathcal{F}\{x(t)\}(f) = \int_{-\infty}^{+\infty} x(t)e^{-2\pi jft} dt \quad (4.2.1)$$

1. $x(t)$ absolument sommable vérifie $\int_{-\infty}^{+\infty} |x(t)| dt < +\infty$.

La transformation de Fourier associe à une fonction x de la variable temporelle t une fonction X d'une variable pouvant s'interpréter comme une variable fréquentielle f . La transformée de Fourier d'une fonction sommable sur $\mathbb{R}$ est continue et bornée sur $\mathbb{R}$ et tend vers 0 en $\pm\infty$ (théorème de Riemann-Lebesgue). Si x est la valeur réelle, le module de la transformée de Fourier est une fonction paire et sa phase est une fonction impaire. La connaissance de $X(f)$ dans le domaine positif est donc suffisante dans le cas où $x(t)$ est à valeurs réelles.

Le tableau 4.1 contient la transformée de Fourier des signaux usuels. On note que ce tableau inclut la transformée de Fourier de l'impulsion de Dirac, δ , qui repose sur la définition de la transformation de Fourier sur l'espace des distributions (voir cours de *Convergence, intégration, probabilités, équations aux dérivées partielles*).

$x(t)$	$X(f)$
1	$\delta_0(f)$
$e^{2\pi j f_0 t}$	$\delta_{f_0}(f)$
$\delta_0(t)$	1
$\delta_{f_0}(t)$	$e^{-2\pi j f_0 t}$
$\sin(2\pi f_0 t)$	$\frac{1}{2j}\delta_{f_0}(f) - \frac{1}{2j}\delta_{-f_0}(f)$
$\cos(2\pi f_0 t)$	$\frac{1}{2}\delta_{f_0}(f) + \frac{1}{2}\delta_{-f_0}(f)$
$x(t) = 1$ si $t \in [0, T]$ $x(t) = 0$ sinon	$T e^{-\pi j f T} \frac{\sin(\pi f T)}{\pi f T}$

TABLEAU 4.1 – Transformées de Fourier des fonctions usuelles (cas continu)

Définition 4.2.2

Dans le cas discret, la *transformée de Fourier discrète* d'un signal $x(k)$ est définie de façon similaire :

$$X(\nu) = \mathcal{F}\{x(k)\}(\nu) = \sum_{k=-\infty}^{k=+\infty} x(k) e^{-2\pi j \nu k} \quad (4.2.2)$$

La transformée de Fourier discrète est définie pour tout *signal* $x(k)$ *absolument sommable*, i.e. :

$$\sum_{k=-\infty}^{k=+\infty} |x(k)| < +\infty \quad (4.2.3)$$

De plus, la transformée de Fourier discrète $X(\nu)$, lorsqu'elle est définie, est une fonction périodique de période 1. La connaissance de $X(\nu)$ sur l'intervalle $\left[0, \frac{1}{2}\right]$ est donc suffisante pour les signaux à valeur réelle.

Le tableau 4.2 contient la transformée de Fourier des signaux usuels dans le cas discret, où δ est l'impulsion discrète.

On peut noter que la transformée de Fourier n'est pas définie pour des signaux laissant apparaître des divergences exponentielles (de la forme $t^k e^{at}$ avec $a > 0$) qui apparaissent couramment dans les systèmes réels étudiés. Pour cette raison, les transformées de Laplace (pour les signaux à temps continu) et en z (pour les signaux à temps discret) ont été introduites afin de pouvoir étudier ces phénomènes d'instabilité, caractérisés par des signaux à divergence exponentielle, dans les systèmes linéaires.

$x(k)$	$X(\nu)$
1	$\delta_0(\nu)$
$e^{2\pi j a k}$	$\delta_a(\nu)$
$a^k, k \geq 0, a < 1$	$\frac{1}{1 - ae^{-2\pi j \nu}}$
$\delta(k - a)$	$e^{-2\pi j \nu a}$
$\sin(2\pi \nu_0 k)$	$\frac{1}{2j} \delta_{\nu_0}(\nu) - \frac{1}{2j} \delta_{-\nu_0}(\nu)$
$\cos(2\pi \nu_0 k)$	$\frac{1}{2} \delta_{\nu_0}(\nu) + \frac{1}{2} \delta_{-\nu_0}(\nu)$
$x(k) = 1$ si $k \in [0, N - 1]$ $x(k) = 0$ sinon	$e^{-\pi j \nu (N-1)} \frac{\sin(\pi \nu N)}{\sin(\pi \nu)}$

TABLEAU 4.2 – Transformées de Fourier des fonctions usuelles (cas discret)

4.2.2 Transformation de Laplace

La *transformée de Laplace* permet l'étude des phénomènes d'instabilité dans le cas des signaux à temps continu. L'idée est d'ajouter à l'argument imaginaire pur $j2\pi f$ une partie réelle notée σ que l'on choisira de façon à faire converger l'intégrale suivante appelée transformée de Laplace bilatérale du signal $x(t)$:

$$X(p) = \mathcal{L}\{x(t)\}(p) = \int_{-\infty}^{+\infty} x(t)e^{-pt} dt \text{ en posant } p = \sigma + j2\pi f \quad (4.2.4)$$

$X(p)$ est une fonction de la variable complexe p .

REMARQUE : Dans la littérature anglo-saxonne, p est noté s .

4.2.2.1 Existence de la transformée de Laplace

Soit $X(p)$ la transformée de Laplace du signal $x(t)$. Alors :

$$|X(p)| = \left| \int_{-\infty}^{+\infty} x(t)e^{-pt} dt \right| \leq \int_{-\infty}^{+\infty} |x(t)| e^{-\sigma t} dt$$

Donc la transformée de Laplace existe si $\int_{-\infty}^{+\infty} |x(t)| e^{-\sigma t} dt$ existe.

On va montrer que $X(p)$ existe si $x(t)$ est localement sommable et a une croissance dite "au plus exponentielle".

Supposons que :

$$\begin{cases} \exists (A > 0, \alpha, t_1 \geq 0) & \text{tel que } |x(t)| \leq Ae^{\alpha t}, \forall t \geq t_1 \\ \exists (B > 0, \beta, t_2 \leq 0) & \text{tel que } |x(t)| \leq Be^{\beta t}, \forall t \leq t_2 \end{cases}$$

Alors :

$$|X(p)| \leq \int_{-\infty}^{+\infty} |x(t)| e^{-\sigma t} dt = \int_{-\infty}^{t_2} |x(t)| e^{-\sigma t} dt + \int_{t_2}^{t_1} |x(t)| e^{-\sigma t} dt + \int_{t_1}^{+\infty} |x(t)| e^{-\sigma t} dt$$

donc :

$$|X(p)| \leq B \int_{-\infty}^{t_2} e^{(\beta - \sigma)t} dt + \int_{t_2}^{t_1} |x(t)| e^{-\sigma t} dt + A \int_{t_1}^{+\infty} e^{(\alpha - \sigma)t} dt$$

Or :

$$\left\{ \begin{array}{ll} \int_{-\infty}^{t_2} e^{(\beta-\sigma)t} dt & \text{existe si } \sigma < \beta \\ \int_{t_1}^{+\infty} |x(t)| e^{-\sigma t} dt & \text{existe puisqu'on a supposé } x(t) \text{ localement sommable} \\ \int_{t_1}^{t_2} e^{(\alpha-\sigma)t} dt & \text{existe si } \sigma > \alpha \end{array} \right.$$

La transformée de Laplace existe donc pour $\alpha < \operatorname{Re}(p) < \beta$.

Soit σ_1 et σ_2 la plus petite (respectivement la plus grande) des valeurs possibles de α (respectivement β) ; on appelle domaine de convergence (ou bande de convergence) de $X(p)$ l'intervalle $]\sigma_1, \sigma_2[$.

Exemple 4.2.1. Soit le signal $x(t) = \begin{cases} e^{at} & \text{si } t \geq 0, a > 0 \\ 0 & \text{si } t < 0 \end{cases}$. Sa transformée de Laplace est :

$$X(p) = \int_0^{+\infty} e^{at} e^{-pt} dt = \left[\frac{e^{(a-p)t}}{a-p} \right]_0^{+\infty} = \frac{1}{p-a} \text{ si } \operatorname{Re}(p) = \sigma > a$$

Son domaine de convergence est $]a, +\infty[$.

4.2.2.2 Propriétés de la transformée de Laplace

Etant donné la ressemblance entre l'expression de la transformée de Laplace ($X(p)$) et celle de la transformée de Fourier ($\mathcal{F}\{x(t)\}$), les propriétés de $X(p)$ sont analogues à celles de $\mathcal{F}\{x(t)\}$.

4.2.2.2.1 Linéarité L'intégrale étant une fonction linéaire, on vérifie aisément que $\mathcal{L}\{\lambda x(t) + \mu y(t)\}(p) = \lambda \mathcal{L}\{x(t)\}(p) + \mu \mathcal{L}\{y(t)\}(p)$. Le domaine de convergence de $\mathcal{L}\{\lambda x(t) + \mu y(t)\}(p)$ est l'intersection des domaines de convergence de $\mathcal{L}\{x(t)\}(p)$ et de $\mathcal{L}\{y(t)\}(p)$.

4.2.2.2.2 Décalage temporel, décalage en p Soient $x(t)$ et $y(t) = x(t - \tau)$:

$$Y(p) = \int_{-\infty}^{+\infty} y(t) e^{-pt} dt = \int_{-\infty}^{+\infty} x(t - \tau) e^{-pt} dt = \int_{-\infty}^{+\infty} x(u) e^{-p(u+\tau)} du$$

$$Y(p) = e^{-p\tau} \int_{-\infty}^{+\infty} x(u) e^{-pu} du = e^{-p\tau} X(p)$$

$$\boxed{y(t) = x(t - \tau) \Leftrightarrow Y(p) = e^{-p\tau} X(p)} \quad (4.2.5)$$

Un décalage temporel de $x(t)$ se traduit par une multiplication de $X(p)$ par une exponentielle complexe.

$$\mathcal{L}\{x(t)\}(p) = \int_{-\infty}^{+\infty} x(t) e^{-pt} dt = X(p)$$

$$\mathcal{L}\{x(t) e^{-at}\}(p) = \int_{-\infty}^{+\infty} x(t) e^{-at} e^{-pt} dt = \int_{-\infty}^{+\infty} x(t) e^{-(p+a)t} dt = X(p+a)$$

$$\boxed{\mathcal{L}\{x(t) e^{-at}\}(p) = X(p+a)} \quad (4.2.6)$$

Un décalage de $X(p)$ se traduit par une multiplication de $x(t)$ par une exponentielle complexe.

4.2.2.2.3 Transformée d'un produit de convolution Soient deux signaux $x(t)$ et $y(t)$. Leur produit de convolution est donné par :

$$\begin{aligned}\mathcal{L}\{(x * y)(t)\}(p) &= \int_{-\infty}^{+\infty} \left(\int_{-\infty}^{+\infty} x(u)y(t-u)du \right) e^{-pt}dt = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} x(u)y(t-u)e^{-pt}dtdu \\ \mathcal{L}\{(x * y)(t)\}(p) &= \underbrace{\int_{-\infty}^{+\infty} x(u)e^{-pu}du}_{X(p)} \underbrace{\int_{-\infty}^{+\infty} y(t-u)e^{-p(t-u)}dt}_{Y(p)} \\ \boxed{\mathcal{L}\{(x * y)(t)\}(p) = X(p)Y(p)} &\quad (4.2.7)\end{aligned}$$

La transformée de Laplace transforme un produit de convolution en un produit.

Les deux transformées $X(p)$ et $Y(p)$ devant exister, le domaine de convergence de $\mathcal{L}\{(x * y)(t)\}(p)$ est l'intersection des domaines de convergence de $\mathcal{L}\{x(t)\}(p)$ et de $\mathcal{L}\{y(t)\}(p)$.

4.2.2.2.4 Intégration, dérivation Soit un signal $x(t)$ continu et $X(p)$ sa transformée de Laplace.

$$\mathcal{L}\left\{\frac{dx}{dt}\right\}(p) = \int_{-\infty}^{+\infty} \frac{dx}{dt} e^{-pt} dt = [x(t)e^{-pt}]_{-\infty}^{+\infty} + p \int_{-\infty}^{+\infty} x(t)e^{-pt} dt = [x(t)e^{-pt}]_{-\infty}^{+\infty} + pX(p)$$

$x(t)$ étant à croissance au plus exponentielle, on a :

$$\begin{cases} |x(t)e^{-pt}| \leq \exp((\sigma_1 - \operatorname{Re}(p))t) & \text{quand } t \rightarrow +\infty \\ |x(t)e^{-pt}| \leq \exp((\sigma_2 - \operatorname{Re}(p))t) & \text{quand } t \rightarrow -\infty \end{cases}$$

d'où :

$$\mathcal{L}\left\{\frac{dx}{dt}\right\}(p) = pX(p); \sigma_1 < \operatorname{Re}(p) < \sigma_2$$

Remarque 4.2.1. Le domaine de convergence de $\mathcal{L}\left\{\frac{dx}{dt}\right\}(p)$ est égal à celui de $\mathcal{L}\{x(t)\}(p)$.

Remarque 4.2.2. L'expression trouvée ci-dessus conduit à considérer la variable p au sens d'un calcul symbolique comme un opérateur de dérivation. La transformée de Laplace devient alors un outil puissant pour résoudre certaines équations différentielles.

$$\begin{array}{ccc} x(t) & \xrightarrow{\quad} & \boxed{\frac{d}{dt}} \longrightarrow y(t) = \frac{dx}{dt} \\ X(p) & \xrightarrow{\quad} & \boxed{p} \longrightarrow Y(p) = pX(p) \end{array}$$

Si $x(t)$ présente une discontinuité de première espèce en $t = t_0$, en scindant l'intégrale en deux parties on trouve :

$$\mathcal{L}\left\{\frac{dx}{dt}\right\}(p) = pX(p) + (x(t_0^-) - x(t_0^+))e^{-t_0 p}; \sigma_1 < \operatorname{Re}(p) < \sigma_2$$

Par exemple, si $x(t)$ est un signal causal :

$$\boxed{\mathcal{L}\left\{\frac{dx}{dt}u(t)\right\}(p) = pX(p) - x(0^+); \sigma_1 < \operatorname{Re}(p) < \sigma_2} \quad (4.2.8)$$

Remarque 4.2.3. De même, on démontre :

$$\boxed{\mathcal{L}\left\{\int_0^t x(u)du\right\}(p) = \frac{1}{p}X(p)} \quad (4.2.9)$$

4.2.2.2.5 Multiplication par la variable d'évolution Dans les domaines où elle est définie, la convergence de la transformée de Laplace est absolue ; on peut donc écrire :

$$\begin{aligned}\frac{dX(p)}{dp} &= \frac{d}{dp} \left(\int_{-\infty}^{+\infty} x(t)e^{-pt} dt \right) = \int_{-\infty}^{+\infty} \frac{d}{dp} (x(t)e^{-pt}) dt \\ &= \int_{-\infty}^{+\infty} -tx(t)e^{-pt} dt \text{ pour } \sigma_1 < \operatorname{Re}(p) < \sigma_2\end{aligned}$$

soit :

$$\mathcal{L}\{tx(t)\}(p) = -\frac{dX(p)}{dp}; \sigma_1 < \operatorname{Re}(p) < \sigma_2$$

En réitérant ce processus, on obtient, de façon générale :

$$\boxed{\mathcal{L}\{t^n x(t)\}(p) = (-1)^n \frac{d^n X(p)}{dp^n}; \sigma_1 < \operatorname{Re}(p) < \sigma_2} \quad (4.2.10)$$

4.2.2.2.6 Valeur initiale, valeur finale Soit $x(t)$ un signal causal de transformée $X(p)$. Lorsque les deux limites suivantes existent, on a les relations :

$$\boxed{\begin{cases} \lim_{t \rightarrow 0^+} x(t) = \lim_{\operatorname{Re}(p) \rightarrow +\infty} pX(p) & \text{théorème de la valeur initiale} \\ \lim_{t \rightarrow +\infty} x(t) = \lim_{p \rightarrow 0} pX(p) & \text{théorème de la valeur finale} \end{cases}} \quad (4.2.11)$$

4.2.2.3 Transformée de Laplace inverse

Le calcul de la transformée de Laplace inverse est donné par l'expression suivante (formule de Bromwich-Wagner) :

$$\boxed{x(t) = \mathcal{L}^{-1}\{X(p)\}(t) = \frac{1}{2\pi j} \operatorname{VP} \left(\int_{\sigma_0 - j\infty}^{\sigma_0 + j\infty} X(p)e^{pt} dp \right)} \quad (4.2.12)$$

où $j = \sqrt{-1}$, $\sigma_1 < \sigma_0 < \sigma_2$ et VP représente la valeur principale de Cauchy.

Le calcul de cette intégrale de la variable complexe fait appel au théorème des résidus, théorème important du domaine de l'analyse fonctionnelle complexe mais qui sort du cadre de ce cours. En conséquence, dans ce cours, on se ramènera systématiquement à des formes de fonctions connues telles que présentées dans le paragraphe 4.2.2.4 suivant.

4.2.2.4 Transformées usuelles

Le tableau 4.3 donne la transformée de Laplace des principales fonctions **causales** (pour lesquelles on a donc $\sigma_2 = +\infty$) rencontrées en traitement du signal. Le calcul de $X(p)$ en fonction de $x(t)$ s'obtient facilement en utilisant les propriétés données en 4.2.2.2. Ce tableau sera en revanche très pratique pour trouver $x(t)$ connaissant $X(p)$ sans utiliser la transformée de Laplace inverse vue en 4.2.2.3.

Dans le tableau 4.3, $\delta_0(t)$ représente l'impulsion de Dirac et $\mathbb{1}(t)$ représente la fonction échelon de Heaviside.

Exemple 4.2.2 (Résolution d'une équation différentielle). *On recherche les signaux $x(t)$ causaux vérifiant l'équation différentielle suivante :*

$$\frac{d^2 x}{dt^2} + \frac{dx}{dt} + x = 0$$

avec les conditions limites $x(t=0) = x(0)$ et $\frac{dx}{dt}(t=0) = x'(0)$.

$x(t)$	$X(p)$	σ_1
$\delta_0(t)$	1	$-\infty$
$\delta_0^{(n)}(t)$	p^n	$-\infty$
$\mathbb{1}(t)$	$\frac{1}{p}$	0
$\mathbb{1}(t)e^{-at}$	$\frac{1}{p+a}$	$-a$
$\mathbb{1}(t)t^n e^{-at}$	$\frac{1}{(p+a)^{n+1}}$	$-a$
$\mathbb{1}(t) \sin(\omega t) e^{-at}$	$\frac{\omega}{(p+a)^2 + \omega^2}$	$-a$
$\mathbb{1}(t) \cos(\omega t) e^{-at}$	$\frac{p+a}{(p+a)^2 + \omega^2}$	$-a$
$x_m(t) * \sum_{k=0}^{+\infty} \delta(t - kT)$ (signal périodique de motif période $x_m(t)$)	$\frac{X_m(p)}{1 - e^{-Tp}}$	σ_{m_1}

TABLEAU 4.3 – Transformées de Laplace des fonctions usuelles

En appliquant la transformée de Laplace à cette équation, on obtient :

$$\begin{aligned} (p^2 + p + 1) X(p) &= (1 + p)x(0) + x'(0) \\ \Leftrightarrow X(p) &= \frac{p+1}{p^2 + p + 1}x(0) + \frac{1}{p^2 + p + 1}x'(0) \end{aligned}$$

soit :

$$\begin{aligned} X(p) = & \left(\frac{1}{2} \frac{2}{\sqrt{3}} \frac{\frac{\sqrt{3}}{2}}{\left(p + \frac{1}{2}\right)^2 + \left(\frac{\sqrt{3}}{2}\right)^2} + \frac{p + \frac{1}{2}}{\left(p + \frac{1}{2}\right)^2 + \left(\frac{\sqrt{3}}{2}\right)^2} \right) x(0) \\ & + \frac{2}{\sqrt{3}} \frac{\frac{\sqrt{3}}{2}}{\left(p + \frac{1}{2}\right)^2 + \left(\frac{\sqrt{3}}{2}\right)^2} x'(0) \end{aligned}$$

D'où, en utilisant les transformées usuelles :

$$x(t) = e^{-\frac{1}{2}t} \left(\frac{2}{\sqrt{3}} \left(\frac{1}{2}x(0) + x'(0) \right) \sin\left(\frac{\sqrt{3}}{2}t\right) + \cos\left(\frac{\sqrt{3}}{2}t\right) x(0) \right) u(t)$$

4.2.3 Transformation en z

La transformée en z permet l'étude des phénomènes d'instabilité dans le cas des signaux à temps discret. Définie à partir de la transformée de Fourier discrète, l'idée est d'ajouter à l'argument imaginaire pur $j2\pi\nu$ une partie réelle notée σ que l'on choisira de façon à faire converger la somme suivante appelée transformée en z du signal $x(k)$:

$$X(z) = \mathcal{Z}\{x(k)\}(z) = \sum_{k=-\infty}^{+\infty} x(k)z^{-k} \text{ en posant } z = e^{\sigma + j2\pi\nu} \quad (4.2.13)$$

$X(z)$ est une fonction de la variable complexe z .

4.2.3.1 Existence de la transformée en z

Soit $X(z)$ la transformée en z du signal $x(k)$. Alors :

$$|X(z)| = \left| \sum_{k=-\infty}^{+\infty} x(k)z^{-k} \right| \leq \sum_{k=-\infty}^{+\infty} |x(k)| \rho^{-k} \text{ en posant } z = \rho e^{j\theta}$$

Donc la transformée en z existe si $\sum_{k=-\infty}^{+\infty} |x(k)| \rho^{-k}$ existe.

On va montrer que $X(z)$ existe si $x(k)$ a une croissance dite "au plus exponentielle". Supposons que :

$$\begin{cases} \exists (A > 0, \alpha \geq 0, k_1 \geq 0) & \text{tel que } |x(k)| \leq A\alpha^k, \forall k \geq k_1 \geq 0 \\ \exists (B > 0, \beta \geq 0, k_2 \leq 0) & \text{tel que } |x(k)| \leq B\beta^k, \forall k \leq k_2 \leq 0 \end{cases}$$

Alors :

$$|X(z)| \leq \sum_{k=-\infty}^{+\infty} |x(k)| \rho^{-k} = \sum_{k=-\infty}^{k_2} |x(k)| \rho^{-k} + \sum_{k=k_2+1}^{k_1-1} |x(k)| \rho^{-k} + \sum_{k=k_1}^{+\infty} |x(k)| \rho^{-k}$$

donc :

$$|X(z)| \leq B \sum_{k=-\infty}^{k_2} \left(\frac{\beta}{\rho}\right)^k + \sum_{k=k_2+1}^{k_1-1} |x(k)| \rho^{-k} + A \sum_{k=k_1}^{+\infty} \left(\frac{\alpha}{\rho}\right)^k$$

Or :

$$\begin{cases} \sum_{k=-\infty}^{k_2} \left(\frac{\beta}{\rho}\right)^k & \text{existe si } \rho < \beta \\ \sum_{k=k_1}^{+\infty} \left(\frac{\alpha}{\rho}\right)^k & \text{existe si } \rho > \alpha \end{cases}$$

La transformée en z existe donc pour $\alpha < |z| < \beta$.

Soit R_1 et R_2 la plus petite (respectivement la plus grande) des valeurs possibles de α (respectivement β) ; on appelle domaine de convergence ou couronne de convergence de $X(z)$ la surface comprise entre les cercles de rayon R_1 et R_2 .

Exemple 4.2.3. Soit le signal $x(k) = \begin{cases} a^k & \text{si } k \geq 0, a > 0 \\ 0 & \text{si } k < 0 \end{cases}$. Sa transformée en z est :

$$X(z) = \sum_{k=-\infty}^{+\infty} a^k z^{-k} = \sum_{k=0}^{+\infty} (az^{-1})^k = \left[\frac{1}{1 - az^{-1}} \right]_0^{+\infty} \text{ si } |az^{-1}| < 1 \Leftrightarrow |z| > a$$

Sa couronne de convergence est l'extérieur du disque de rayon a .

4.2.3.2 Propriétés de la transformée en z

Etant donné la ressemblance entre l'expression de la transformée de Laplace ($X(p)$) et celle de la transformée en z, les propriétés de $X(z)$ sont analogues à celles de $X(p)$.

4.2.3.2.1 Linéarité L'opérateur somme étant une fonction linéaire, on vérifie que :

$$\mathcal{Z}\{\lambda x(k) + \mu y(k)\}(z) = \lambda \mathcal{Z}\{x(k)\}(z) + \mu \mathcal{Z}\{y(k)\}(z).$$

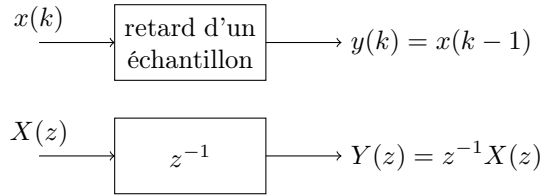
Le domaine de convergence de $\mathcal{Z}\{\lambda x(k) + \mu y(k)\}(z)$ est l'intersection des domaines de convergence de $\mathcal{Z}\{x(k)\}(z)$ et de $\mathcal{Z}\{y(k)\}(z)$.

4.2.3.2.2 Décalage temporel Soient $x(k)$ et $y(k) = x(k - k_0)$. Alors :

$$Y(z) = \sum_{k=-\infty}^{+\infty} y(k)z^{-k} = \sum_{k=-\infty}^{+\infty} x(k - k_0)z^{-k} = \sum_{k=-\infty}^{+\infty} x(k)z^{-(k+k_0)}$$

$$\boxed{y(k) = x(k - k_0) \Leftrightarrow Y(z) = z^{-k_0} X(z)} \quad (4.2.14)$$

Un décalage temporel de $x(k)$ de k_0 échantillons se traduit par une multiplication de $X(z)$ par z^{-k_0} . Cette propriété conduit à considérer la variable z^{-1} au sens du calcul symbolique comme un opérateur retard de 1 échantillon.



4.2.3.2.3 Multiplication par une exponentielle Soit $y(k) = a^k x(k)$:

$$Y(z) = \sum_{k=-\infty}^{+\infty} a^k x(k)z^{-k} = \sum_{k=-\infty}^{+\infty} x(k) \left(\frac{z}{a}\right)^{-k} = X\left(\frac{z}{a}\right) ; |a| R_1 < |z| < |a| R_2$$

$$\boxed{\mathcal{Z}\{a^k x(k)\}(z) = X\left(\frac{z}{a}\right)} \quad (4.2.15)$$

4.2.3.2.4 Transformée d'un produit de convolution Soient deux signaux $x(k)$ et $y(k)$. Le produit de convolution de deux signaux discrets est donné par :

$$\mathcal{Z}\{(x * y)(k)\}(z) = \sum_{k=-\infty}^{+\infty} \left[\sum_{l=-\infty}^{+\infty} x(k-l)y(l) \right] z^{-k} = \sum_{k=-\infty}^{+\infty} \sum_{l=-\infty}^{+\infty} x(k-l)y(l)z^{-(k-l)}z^{-l}$$

$$\mathcal{Z}\{(x * y)(k)\}(z) = \underbrace{\sum_{l=-\infty}^{+\infty} y(l)z^{-l}}_{Y(z)} \underbrace{\sum_{u=-\infty}^{+\infty} x(u)z^{-u}}_{X(z)}$$

$$\boxed{\mathcal{Z}\{(x * y)(k)\}(z) = X(z)Y(z)} \quad (4.2.16)$$

La transformée en z transforme un produit de convolution en un produit.

Les deux transformées $X(z)$ et $Y(z)$ devant exister, le domaine de convergence de $\mathcal{Z}\{(x * y)(k)\}(z)$ est l'intersection des domaines de convergence de $\mathcal{Z}\{x(k)\}(z)$ et de $\mathcal{Z}\{y(k)\}(z)$.

4.2.3.2.5 Multiplication par la variable d'évolution Dans les domaines où elle est définie, la convergence de la transformée en z est absolue ; on peut donc écrire :

$$\frac{dX(z)}{dz} = \frac{d}{dz} \left(\sum_{k=-\infty}^{+\infty} x(k)z^{-k} \right) = \sum_{k=-\infty}^{+\infty} \frac{d}{dz} (x(k)z^{-k})$$

$$= - \sum_{k=-\infty}^{+\infty} kx(k)z^{-k-1} \text{ pour } R_1 < |z| < R_2$$

soit :

$$\mathcal{Z}\{kx(k)\}(z) = -z \frac{dX(z)}{dz} ; R_1 < |z| < R_2$$

En réitérant ce processus, on obtient de façon générale :

$$\mathcal{Z}\{k^n x(k)\}(z) = \left(-z \frac{d}{dz}\right)^n X(z); R_1 < |z| < R_2 \quad (4.2.17)$$

où $\left(-z \frac{d}{dz}\right)^n$ signifie que l'on applique n fois à $X(z)$ l'opérateur $-z \frac{d}{dz}$.

4.2.3.2.6 Valeur initiale, valeur finale Soit $x(k)$ un signal causal de transformée $X(z)$. Lorsque les deux limites suivantes existent, on a les relations :

$$\begin{cases} x(0) = \lim_{|z| \rightarrow +\infty} X(z) & \text{théorème de la valeur initiale} \\ \lim_{n \rightarrow +\infty} x(k) = \lim_{z \rightarrow 1} (1 - z^{-1}) X(z) & \text{théorème de la valeur finale} \end{cases} \quad (4.2.18)$$

4.2.3.3 Transformée en z inverse

Le calcul de la transformée en z inverse est donné par l'expression suivante :

$$x(k) = \mathcal{Z}^{-1}\{X(z)\}(k) = \frac{1}{2\pi j} \oint_{\Gamma} X(z) z^{k-1} dz \quad (4.2.19)$$

où Γ est un cercle centré sur l'origine et appartenant à la couronne de convergence de $X(z)$.

Comme pour la transformée de Laplace inverse, le calcul de cette intégrale de la variable complexe sort du cadre de ce cours et on se ramènera aux formes des transformées en z de signaux usuels telles que listées dans le paragraphe 4.2.3.4 suivant.

4.2.3.4 Transformées en z usuelles

De même que pour la transformée de Laplace, le tableau 4.4 donne la transformée en z pour des signaux discrets usuels **causaux** (pour lesquels on a donc $R_2 = \infty$). Il sera très pratique pour trouver $x(k)$ connaissant $X(z)$ sans utiliser la transformée en z inverse vue en 4.2.3.3.

Dans le tableau 4.4, $\delta_0(k)$ représente l'impulsion discrète et $\mathbb{1}(k)$ représente la fonction échelon de Heaviside.

$x(k)$	$X(z)$	R_1
$\delta_0(k)$	1	0
$\mathbb{1}(k)$	$\frac{1}{1 - z^{-1}}$	1
$k\mathbb{1}(k)$	$\frac{z^{-1}}{(1 - z^{-1})^2}$	1
$a^k \mathbb{1}(k)$	$\frac{1}{1 - az^{-1}}$	$ a $
$a^k \cos(k2\pi\nu) \mathbb{1}(k)$	$\frac{1 - z^{-1}a \cos(2\pi\nu)}{1 - 2z^{-1}a \cos(2\pi\nu) + a^2 z^{-2}}$	$ a $
$a^k \sin(k2\pi\nu) \mathbb{1}(k)$	$\frac{z^{-1}a \sin(2\pi\nu)}{1 - 2z^{-1}a \cos(2\pi\nu) + a^2 z^{-2}}$	$ a $

TABLEAU 4.4 – Transformées en z des fonctions usuelles

4.3 Représentation temporelle des systèmes

4.3.1 Réponses impulsionnelle et indicielle des systèmes linéaires invariants à temps discret

Définition 4.3.1

La *réponse impulsionnelle* d'un système discret est la sortie $y(k) = h(k)$ en réponse à l'entrée $u(k) = \delta(k)$, où $\delta(k)$ est l'impulsion discrète (voir figure 4.1).

Dans le cas de systèmes causaux, cette réponse impulsionnelle est nulle pour les temps négatifs.

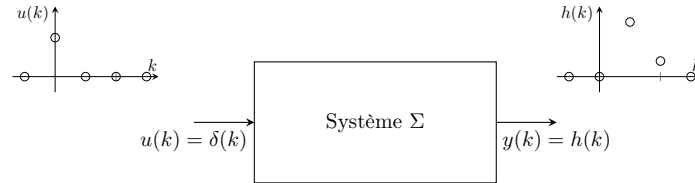


FIGURE 4.1 – Réponse impulsionnelle d'un système discret

Définition 4.3.2

La *réponse indicielle* d'un système est la sortie en réponse à une entrée en échelon. Dans le cas de systèmes linéaires et invariants, la réponse indicielle est la somme cumulée (intégrale numérique) de la réponse impulsionnelle.

4.3.2 Relation de convolution dans les systèmes linéaires invariants

4.3.2.1 Cas discret

Dans le cas des systèmes linéaires et invariants à temps discret, la sortie y peut se calculer à partir de la réponse impulsionnelle du système.

En effet, le signal d'entrée u peut être vu comme une suite infinie d'impulsions décalées dans le temps :

$$u(k) = \sum_{n=-\infty}^{+\infty} u(n)\delta(k-n) \quad (4.3.1)$$

Comme le système est linéaire et invariant, la sortie est la somme infinie des réponses à ces impulsions. La réponse à l'impulsion $u(n)\delta(k-n)$ est égale à la suite $u(n)h(k-n)$ (système invariant). Il vient donc :

$$y(k) = \sum_{n=-\infty}^{+\infty} u(n)h(k-n) \quad (4.3.2)$$

Finalement, pour une entrée u quelconque et dans le cas d'un système linéaire et invariant, la sortie y peut donc être calculée par convolution de l'entrée par la réponse impulsionnelle :

$$y(k) = h * u(k) = \sum_{n=-\infty}^{+\infty} h(n)u(k-n) \quad (4.3.3)$$

Plus de détails sur la convolution sont donnés en annexe D.

4.3.2.2 Cas continu

Comme pour le cas discret, la réponse impulsionnelle d'un système est la sortie $y(t) = h(t)$ en réponse à l'entrée $u(t) = \delta(t)$, où $\delta(t)$ est l'impulsion de Dirac (voir figure 4.2). La réponse indicielle d'un système est la sortie en réponse à une entrée en échelon. Dans le cas de systèmes linéaires et invariants, la réponse indicielle est l'intégrale de la réponse impulsionnelle.

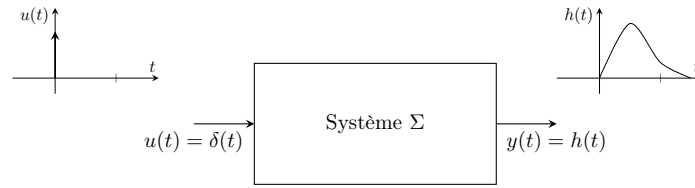
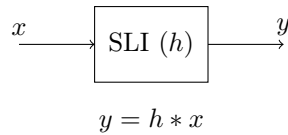


FIGURE 4.2 – Réponse impulsionnelle d'un système continu

4.4 Etude des systèmes linéaires et invariants par fonction de transfert

4.4.1 Fonction de transfert

La relation entrée-sortie d'un système linéaire et invariant (SLI) de réponse impulsionnelle $h(t)$ (en continu) ou $h(k)$ (en discret) est une équation de convolution dans le domaine temporel (figure 4.3).

FIGURE 4.3 – Produit de convolution $y = h * x$

Si l'on applique la transformée adéquate (Laplace dans le cas continu ou en z dans le cas discret), de par les propriétés des transformées de Laplace et en z , on obtient les relations suivantes :

$$\boxed{Y(p) = H(p)X(p) \text{ ou } Y(z) = H(z)X(z)} \quad (4.4.1)$$

où Y , H et X sont les transformées (de Laplace ou en z) des signaux y , h et x .

Définition 4.4.1

La transformée de la réponse impulsionnelle du SLI :

$$H(p) = \frac{Y(p)}{X(p)} \text{ ou } H(z) = \frac{Y(z)}{X(z)}$$

est appelée *fonction de transfert* ou *transmittance* du SLI.

L'intérêt de l'utilisation des transmittances vient du fait que la plupart des systèmes étudiés sont constitués par la mise en cascade de sous-systèmes élémentaires (figure 4.4).

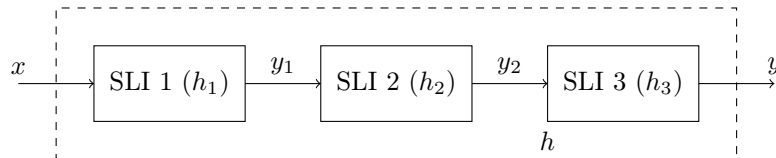


FIGURE 4.4 – Mise en cascade de systèmes

L'étude de ce système dans le domaine temporel conduit à l'expression, en général compliquée : $y = h_3 * h_2 * h_1 * x$. L'utilisation des transmittances conduit à l'expression relativement simple : $Y = H_3 H_2 H_1 X$.

Une autre structure importante où l'utilisation des transmittances est indispensable est celle des systèmes bouclés (figure 4.5).

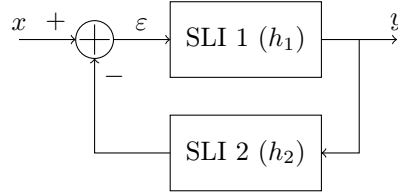


FIGURE 4.5 – Système bouclé

L'étude de ce système dans le domaine temporel donne :

$$\left. \begin{array}{l} \varepsilon = x - h_2 * y \\ y = h_1 * \varepsilon \end{array} \right\} \Rightarrow y = h_1 * x - h_1 * h_2 * y \Leftrightarrow (\delta + h_1 * h_2) * y = h_1 * x$$

Cette expression ne permet pas d'expliciter y à partir de x . L'utilisation des transmittances conduit à une expression plus simple et permet d'expliciter Y à partir de X :

$$Y = H_1 (X - H_2 Y) \Leftrightarrow Y = \frac{H_1}{1 + H_1 H_2} X$$

4.4.2 Stabilité EBSB

4.4.2.1 Définition

Définition 4.4.2

Un SLI sera dit *stable EBSB* (*Entrée Bornée Sortie Bornée*) si sa réponse à toute entrée bornée en amplitude est bornée.

Cette définition de la stabilité est la plus souvent utilisée en pratique lorsqu'on représente un système par fonction de transfert.

Une condition nécessaire et suffisante de la stabilité EBSB est :

$$\boxed{\int_{-\infty}^{+\infty} |h(u)| du < \infty \text{ ou } \sum_{k=-\infty}^{+\infty} |h(k)| < \infty} \quad (4.4.2)$$

c'est-à-dire que $h(t)$ (en continu) ou $h(k)$ (en discret) est absolument sommable.

Démonstration. Cette démonstration est seulement présentée dans le cas continu.

(i) Condition suffisante de stabilité : soit l'entrée bornée $x(t)$ ($|x(t)| < M$). Alors :

$$|y(t)| = \left| \int_{-\infty}^{+\infty} h(u)x(t-u)du \right| \leq \int_{-\infty}^{+\infty} |h(u)| |x(t-u)| du \leq M \int_{-\infty}^{+\infty} |h(u)| du$$

donc $y(t)$ est bornée ; la condition est suffisante.

(ii) Condition nécessaire de stabilité : soit l'entrée $x(t) = \text{sgn}(h(-t))$:

$$y(0) = \int_{-\infty}^{+\infty} h(u)x(-u)du = \int_{-\infty}^{+\infty} |h(u)| du$$

$y(0)$ doit être borné donc la condition est nécessaire.

□

Remarque 4.4.1. Cette démonstration se transpose sans difficulté dans le cas discret.

4.4.2.2 Conditions de stabilité et transmittance

On a vu ci-dessus une condition nécessaire et suffisante de stabilité EBSB sur $h(t)$ ou $h(k)$. On va chercher ici un critère de stabilité sur la transmittance. On va montrer que ce critère dépend des propriétés de causalité du SLI (causal ou anti-causal), ce qui n'est pas le cas de la condition de stabilité donnée dans le domaine temporel.

Considérons tout d'abord le cas d'un système continu causal dont la fonction de transfert $H(p)$ a pour abscisse de convergence σ_1 :

$$H(p) = \int_0^{+\infty} h(t)e^{-pt} dt \text{ pour } \operatorname{Re}(p) > \sigma_1$$

$$|H(0)| = \left| \int_0^{+\infty} h(t) dt \right| \leq \int_0^{+\infty} |h(t)| dt$$

Si le système est stable EBSB cette dernière intégrale est bornée, donc $H(0)$ existe, ce qui signifie que 0 est dans le domaine de convergence de $H(p)$, c'est-à-dire que σ_1 est négative et que $H(p)$ existe pour tout p à partie réelle positive ou nulle.

Dans la plupart des cas étudiés, la transmittance sera une fraction rationnelle :

$$H(p) = \frac{N(p)}{D(p)}$$

où $N(p)$ et $D(p)$ sont des polynômes. Dans ce cas les pôles de $H(p)$ (valeurs de p rendant $H(p)$ infinie) étant par définition à l'extérieur du domaine de convergence sont à partie réelle négative. Pour conclure :

Théorème 4.4.1

Un SLI continu causal dont la transmittance est une fraction rationnelle sera stable EBSB si et seulement si les pôles de sa transmittance sont à partie réelle strictement négative (c'est-à-dire situés dans le demi-plan complexe gauche).

Dans le cas discret, on montre par une démonstration analogue en calculant $H(z)$ pour $z = 1$ que 1 est dans le domaine de convergence de $H(z)$, d'où l'on déduit :

Théorème 4.4.2

Un SLI discret causal dont la transmittance est une fraction rationnelle est stable EBSB si et seulement si les pôles de sa transmittance sont de module strictement inférieur à 1 (c'est-à-dire à l'intérieur du cercle unité).

Dans le cas des systèmes anti-causaux les conclusions précédentes sont inversées ; c'est-à-dire qu'un SLI anti-causal sera stable si et seulement si les pôles sont à partie réelle positive dans le cas continu ou de module supérieur à 1 dans le cas discret.

4.4.2.3 Critères de Routh et de Jury

Ces deux critères permettent de savoir si un système continu (Routh) ou discret (Jury) est stable EBSB. Leur démonstration étant longue et de peu d'intérêt ici, seul le résultat est donné.

4.4.2.3.1 Critère de Routh Le *critère de Routh* donne une condition nécessaire et suffisante pour que les racines d'un polynôme à coefficient réels aient toutes une partie réelle strictement négative.

Soit le polynôme $D(p) = a_n p^n + a_{n-1} p^{n-1} + \dots + a_1 p + a_0$. Les racines de $D(p)$ ont toutes une partie réelle strictement négative si et seulement si :

- (i) $a_i > 0$; $0 \leq i \leq n$
- (ii) les coefficients de la première colonne du tableau de Routh (tableau 4.5, sont tous positifs.

1	a_n	a_{n-2}	a_{n-4}	$\cdots$
2	a_{n-1}	a_{n-3}	a_{n-5}	$\cdots$
3	$c_1 = \frac{a_{n-1}a_{n-2} - a_n a_{n-3}}{a_{n-1}}$	$c_2 = \frac{a_{n-1}a_{n-4} - a_n a_{n-5}}{a_{n-1}}$	$\cdots$	
4	$d_1 = \frac{c_1 a_{n-3} - c_2 a_{n-1}}{c_1}$	$d_2 = \frac{c_1 a_{n-5} - c_3 a_{n-1}}{c_1}$	$\cdots$	
$\vdots$	$\vdots$	$\vdots$	$\vdots$	
n	m_1	m_2		
$n+1$	n_1			

TABLEAU 4.5 – Tableau de Routh

De plus, le nombre de changements de signe dans la première colonne est égal au nombre de racines de $D(p)$ à partie réelle positive.

L'application du critère suppose évidemment qu'on puisse aller jusqu'au bout de la construction du tableau, c'est-à-dire qu'aucun des coefficients de la première colonne ne soit nul (dans le cas contraire, on conclut que le critère n'est pas vérifié).

Remarque 4.4.2. Lorsque dans les formules du tableau 4.5 un coefficient est absent (par exemple si a_{n-5} n'existe pas), on le remplace par un 0.

4.4.2.3.2 Critère de Jury Le critère de Jury donne une condition nécessaire et suffisante pour que les racines d'un polynôme à coefficients réels aient toutes un module strictement inférieur à 1.

Soit le polynôme $D(z) = a_{0,0} + a_{0,1}z + a_{0,2}z^2 + \cdots + a_{0,n}z^n$, avec $a_{0,n} > 0$. On construit le tableau $(n-1) \times (n+1)$ suivant :

$$\begin{pmatrix} a_{0,0} & a_{0,1} & a_{0,2} & a_{0,3} & \cdots & a_{0,n-1} & a_{0,n} \\ a_{1,0} & a_{1,1} & a_{1,2} & a_{1,3} & \cdots & a_{1,n-1} & 0 \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ a_{n-2,0} & a_{n-2,1} & a_{n-2,2} & 0 & \cdots & 0 & 0 \end{pmatrix}$$

avec :

$$a_{j+1,k} = \begin{cases} \begin{vmatrix} a_{j,0} & a_{j,n-j-k} \\ a_{j,n-j} & a_{j,k} \end{vmatrix} & \text{pour } 0 \leq k \leq n-j-1 \\ 0 & \text{pour } k > n-j-1 \end{cases}$$

$D(z)$ n'a pas de racine de module supérieur ou égal à 1 si et seulement si :

1. $\sum_{i=0}^n a_{0,i} = D(1) > 0$
2. $(-1)^n \sum_{i=0}^n (-1)^i a_{0,i} = (-1)^n D(-1) > 0$
3. $|a_{0,0}| - a_{0,n} < 0$
4. $|a_{j,0}| - |a_{j,n-j}| > 0$, pour $j = 1, 2, \dots, n-2$

4.4.3 Réponse en fréquence

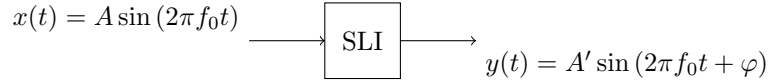
4.4.3.1 Définition

La réponse en fréquence $\hat{h}(f)$ (notion qui sera approfondie dans le cours *Traitement de signaux*) d'un système linéaire invariant peut être déterminée expérimentalement par une analyse harmonique du système en injectant à l'entrée un signal sinusoïdal. Les données expérimentales $|\hat{h}(f)|$ et $\arg(\hat{h}(f))$ sont reportées dans un système de coordonnées, par exemple dans le plan de Bode.

$$|\hat{h}(f_0)| = \frac{A'}{A} \text{ et } \arg(\hat{h}(f_0)) = \varphi$$

Cas où $n = 2$ $D(z) = a_{0,2}z^2 + a_{0,1}z + a_{0,0}$, $a_{0,2} > 0$	Cas où $n = 3$ $D(z) = a_{0,3}z^3 + a_{0,2}z^2 + a_{0,1}z + a_{0,0}$, $a_{0,3} > 0$
Conditions de stabilité :	
1. $a_{0,0} + a_{0,1} + a_{0,2} > 0$	1. $a_{0,0} + a_{0,1} + a_{0,2} + a_{0,3} > 0$
2. $a_{0,0} - a_{0,1} + a_{0,2} > 0$	2. $-a_{0,0} + a_{0,1} - a_{0,2} + a_{0,3} > 0$
3. $ a_{0,0} - a_{0,2} < 0$	3. $ a_{0,0} - a_{0,3} < 0$
4. ne s'applique pas puisque $n - 2 = 0$	4. $ a_{1,0} - a_{1,2} > 0$ soit $ a_{0,0}^2 - a_{0,3}^2 - a_{0,0}a_{0,2} - a_{0,1}a_{0,3} > 0$

TABLEAU 4.6 – Conditions de Jury dans les cas de polynômes d'ordre 2 et 3



4.4.3.2 Fonction de transfert et réponse harmonique

Cas continu

On considère un système mono-variable continu, le cas des systèmes multi-variables s'en déduisant immédiatement. On applique à l'entrée du système linéaire invariant un signal sinusoïdal :

$$u(t) = u_0 \sin(\omega_0 t + \varphi) = \text{Im}(u_0 e^{j\varphi} e^{j\omega_0 t}) \quad (4.4.3)$$

La sortie est égale à la convolution de ce signal par la réponse impulsionnelle h :

$$\begin{aligned} y(t) &= h * \text{Im}(u_0 e^{j\varphi} e^{j\omega_0 t}) = \text{Im}(h * u_0 e^{j\varphi} e^{j\omega_0 t}) \\ &= u_0 \text{Im} \left(\int_{-\infty}^{+\infty} h(\tau) e^{j\varphi} e^{j\omega_0(t-\tau)} d\tau \right) \\ &= u_0 \text{Im} \left(\left(\int_{-\infty}^{+\infty} h(\tau) e^{-j\omega_0 \tau} d\tau \right) e^{j\varphi} e^{j\omega_0 t} \right) \end{aligned} \quad (4.4.4)$$

On reconnaît la transformée de Laplace de h calculée au point $j\omega_0$, notée $H(j\omega_0)$. Finalement :

$$y(t) = u_0 |H(j\omega_0)| \sin(\omega_0 t + \varphi + \arg(H(j\omega_0))) \quad (4.4.5)$$

La sortie est donc un signal sinusoïdal de même pulsation, dont l'amplitude est multipliée par le module de la fonction de transfert en $j\omega_0$, i.e. $H(j\omega_0)$, avec un déphasage supplémentaire égal à l'argument de $H(j\omega_0)$. La réponse en fréquence, ou réponse harmonique peut donc se déduire de la transformée de Laplace en posant formellement $p = j\omega$. Cependant, pour que cette opération ait un sens, il est nécessaire que la transformée de Laplace soit définie sur l'axe imaginaire. C'est en particulier le cas pour les systèmes causaux stables qui ont tous leurs pôles dans le demi-plan gauche.

Cas discret

De la même façon, pour un système discret attaqué par une entrée sinusoïdale du type :

$$u(k) = u_0 \sin(2\pi k\nu + \varphi) \quad (4.4.6)$$

la sortie se calcule à partir de la transformée en z , notée $H(z)$, de la réponse en fréquence calculée au point $e^{2\pi j\nu}$:

$$y(k) = u_0 |H(e^{2\pi j\nu})| \sin(2\pi k\nu + \varphi + \arg(H(e^{2\pi j\nu}))) \quad (4.4.7)$$

Comme dans le cas continu, pour que cette opération ait un sens, il est nécessaire que la transformée en z soit définie sur le cercle unité. C'est en particulier le cas pour les systèmes causaux stables, qui ont tous leurs pôles dans le cercle unité.

4.4.3.3 Conditions d'existence de la réponse en fréquence

La réponse en fréquence est la transformée de Fourier de la réponse impulsionnelle. En particulier, cette transformée existe si la réponse impulsionnelle est absolument sommable. Cette condition étant également la condition de stabilité des SLI on peut en déduire que pour les systèmes stables (mais **pour les systèmes stables seulement**) on peut passer des transmittances à variables complexes à la réponse en fréquence par le changement de variable adéquat suivant :

$$p = j2\pi f \text{ ou } z = e^{j2\pi\nu}.$$

Exemple 4.4.1 (Système du premier ordre). Soit $H(p) = \frac{1}{p+1}$. Le pôle de $H(p)$ ($p = -1$) est à partie réelle négative donc ce système est stable, et :

$$\hat{h}(f) = H(p) \Big|_{p=j2\pi f} = \frac{1}{j2\pi f + 1}$$

Exemple 4.4.2 (Système intégrateur $h(t) = u(t)$ (échelon)). Soit $H(p) = \frac{1}{p}$. Le pôle de $H(p)$ ($p = 0$) n'est pas à partie réelle strictement négative donc ce système est instable et :

$$\hat{h}(f) = \mathcal{F}\{h(t)\}(f) = \frac{1}{2}\delta(f) + \frac{1}{2\pi j} \text{Pf} \left(\frac{1}{f} \right) \neq H(p) \Big|_{p=j2\pi f} = \frac{1}{j2\pi f}$$

La différence provient du fait que $H(p)$ est définie au sens des fonctions alors que $\hat{h}(f)$ est définie au sens des distributions.

4.4.4 Passage état – transfert. Formes compagnons

Nous avons vu dans les chapitres précédents deux méthodes d'analyse des systèmes linéaires permanents : une représentation externe caractérisée par une fonction de transfert $H(p)$ et une représentation interne caractérisée par la représentation d'état. Nous allons voir ici comment passer d'une représentation interne à une représentation externe et inversement.

Dans le cas général où le système a m entrées et p sorties, la transmittance $\mathbf{H}(p)$ est une matrice (appelée matrice de transfert) de dimension $p \times m$:

$$\begin{pmatrix} Y_1(p) \\ Y_2(p) \\ \vdots \\ Y_i(p) \\ \vdots \\ Y_n(p) \end{pmatrix} = \begin{pmatrix} H_{11}(p) & H_{12}(p) & \cdots & H_{1j}(p) & \cdots & H_{1m}(p) \\ H_{21}(p) & H_{22}(p) & \cdots & H_{2j}(p) & \cdots & H_{2m}(p) \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots \\ H_{i1}(p) & H_{i2}(p) & \cdots & H_{ij}(p) & \cdots & H_{im}(p) \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots \\ H_{n1}(p) & H_{n2}(p) & \cdots & H_{nj}(p) & \cdots & H_{nm}(p) \end{pmatrix} \begin{pmatrix} U_1(p) \\ U_2(p) \\ \vdots \\ U_j(p) \\ \vdots \\ U_m(p) \end{pmatrix}$$

$$\text{avec } H_{ij}(p) = \left(\frac{Y_i}{U_j}(p) \right)_{U_k(p)=0 \text{ pour } k \neq j}.$$

4.4.4.1 Passage état $\rightarrow$ transfert

Systèmes à temps continu Connaissant les équations d'état du système linéaire invariant, on cherche la matrice de transfert $\mathbf{H}(p)$. Pour cela on applique la transformée de Laplace aux deux équations d'état. En ne considérant que la partie causale des signaux on obtient :

$$\begin{cases} \mathcal{L} \left\{ \frac{d\mathbf{x}(t)}{dt} \right\} (p) &= \mathcal{L} \{ \mathbf{A}\mathbf{x}(t) + \mathbf{B}\mathbf{u}(t) \} (p) \\ \mathcal{L} \{ \mathbf{y}(t) \} (p) &= \mathcal{L} \{ \mathbf{C}\mathbf{x}(t) + \mathbf{D}\mathbf{u}(t) \} (p) \end{cases} \Leftrightarrow \begin{cases} p\mathbf{X}(p) - \mathbf{x}(t_0)e^{-pt_0} &= \mathbf{A}\mathbf{X}(p) + \mathbf{B}\mathbf{U}(p) \\ \mathbf{Y}(p) &= \mathbf{C}\mathbf{X}(p) + \mathbf{D}\mathbf{U}(p) \end{cases}$$

d'où :

$$X(p) = (p\mathbf{I} - \mathbf{A})^{-1} (\mathbf{x}(t_0) e^{-pt_0} + \mathbf{B}U(p))$$

et :

$$Y(p) = \left(\mathbf{C} (p\mathbf{I} - \mathbf{A})^{-1} \mathbf{B} + \mathbf{D} \right) U(p) + \mathbf{C} (p\mathbf{I} - \mathbf{A})^{-1} \mathbf{x}(t_0) e^{-pt_0}$$

Par définition, la transmittance caractérise la relation entre les variations de la (ou des) sortie(s) en fonction des variations du (ou des) signaux d'entrée(s) ; cela revient à définir $\mathbf{H}(p)$ comme $Y(p) = \mathbf{H}(p)U(p)$ pour $\mathbf{x}(t_0) = 0$.

On a donc ici :

$$\boxed{\mathbf{H}(p) = \mathbf{C} (p\mathbf{I} - \mathbf{A})^{-1} \mathbf{B} + \mathbf{D}} \quad (4.4.8)$$

Remarque 4.4.3. Calcul de $(p\mathbf{I} - \mathbf{A})^{-1}$:

Afin de pouvoir revenir dans le domaine temporel, il peut être intéressant de mettre $(p\mathbf{I} - \mathbf{A})^{-1}$ sous la forme :

$$(p\mathbf{I} - \mathbf{A})^{-1} = \frac{\mathbf{B}_0 p^{n-1} + \mathbf{B}_1 p^{n-2} + \dots + \mathbf{B}_{n-1}}{p^n + d_1 p^{n-1} + \dots + d_n}$$

où les $\mathbf{B}_i$ sont des matrices carrées $n \times n$ et les scalaires d_i sont les coefficients du polynôme caractéristique de $\mathbf{A}$:

$$\det(p\mathbf{I} - \mathbf{A}) = p^n + d_1 p^{n-1} + \dots + d_n$$

Il existe pour cela des méthodes systématiques et récursives comme l'algorithme de Leverrier-Souriau décrit ci-dessous :

$$\begin{array}{ll} \mathbf{B}_0 = \mathbf{I} & d_1 = -\text{Tr}(\mathbf{B}_0 \mathbf{A}) \\ \mathbf{B}_1 = \mathbf{B}_0 \mathbf{A} + d_1 \mathbf{I} & d_2 = -\frac{1}{2} \text{Tr}(\mathbf{B}_1 \mathbf{A}) \\ \mathbf{B}_2 = \mathbf{B}_1 \mathbf{A} + d_2 \mathbf{I} & d_3 = -\frac{1}{3} \text{Tr}(\mathbf{B}_2 \mathbf{A}) \\ \vdots & \vdots \\ \mathbf{B}_k = \mathbf{B}_{k-1} \mathbf{A} + d_k \mathbf{I} & d_{k+1} = -\frac{1}{k+1} \text{Tr}(\mathbf{B}_k \mathbf{A}) \\ \vdots & \vdots \\ \mathbf{B}_{n-1} = \mathbf{B}_{n-2} \mathbf{A} + d_{n-1} \mathbf{I} & d_n = -\frac{1}{n} \text{Tr}(\mathbf{B}_{n-1} \mathbf{A}) \end{array}$$

La dernière équation $\mathbf{B}_n = \mathbf{B}_{n-1} \mathbf{A} + d_n \mathbf{I} = \mathbf{0}$ sert de vérification au calcul.

Remarque 4.4.4. Changement de base de la représentation interne :

Nous avons déjà vu que la représentation d'état d'un système n'est pas unique. Si on effectue le changement de base $\mathbf{x}'(t) = \mathbf{M}\mathbf{x}(t)$ on obtient comme équations d'état :

$$\begin{aligned} \frac{d\mathbf{x}'(t)}{dt} &= \mathbf{A}' \mathbf{x}'(t) + \mathbf{B}' \mathbf{u}(t) \quad \text{avec} \quad \mathbf{A}' = \mathbf{M} \mathbf{A} \mathbf{M}^{-1} \text{ et } \mathbf{B}' = \mathbf{M} \mathbf{B} \\ \mathbf{y}(t) &= \mathbf{C}' \mathbf{x}'(t) + \mathbf{D}' \mathbf{u}(t) \quad \text{avec} \quad \mathbf{C}' = \mathbf{C} \mathbf{M}^{-1} \text{ et } \mathbf{D}' = \mathbf{D} \end{aligned}$$

d'où :

$$\begin{aligned} \mathbf{H}'(p) &= \mathbf{C}' (p\mathbf{I} - \mathbf{A}')^{-1} \mathbf{B}' + \mathbf{D}' \\ &= \mathbf{C} \mathbf{M}^{-1} (p\mathbf{I} - \mathbf{M} \mathbf{A} \mathbf{M}^{-1})^{-1} \mathbf{M} \mathbf{B} + \mathbf{D} \\ &= \mathbf{C} (\mathbf{M}^{-1} (p\mathbf{I} - \mathbf{M} \mathbf{A} \mathbf{M}^{-1}) \mathbf{M})^{-1} \mathbf{B} + \mathbf{D} \\ &= \mathbf{C} (p\mathbf{I} - \mathbf{A})^{-1} \mathbf{B} + \mathbf{D} \\ &= \mathbf{H}(p) \end{aligned}$$

Le changement de base interne n'a bien évidemment pas modifié la matrice de transfert du système.

Remarque 4.4.5. *Stabilité du système :*

La relation liant la fonction de transfert $H(p)$ et la matrice d'évolution $\mathbf{A}$:

$$\boxed{H(p) = \mathbf{C} (p\mathbf{I} - \mathbf{A})^{-1} \mathbf{B} + \mathbf{D}} \quad (4.4.9)$$

implique que les pôles de $H(p)$ sont inclus dans l'ensemble des valeurs propres de $\mathbf{A}$. On peut en déduire la relation suivante :

$$\boxed{\text{stabilité asymptotique} \Rightarrow \text{stabilité EBSB}}$$

Systèmes à temps discret Les relations entre les équations d'état et la fonction de transfert d'un SLI à temps discret se déduisent par simple transposition du cas continu. En particulier, en appliquant la transformée en z aux équations (3.6.3), on obtient la fonction de transfert :

$$\boxed{H(z) = \mathbf{C} (z\mathbf{I} - \mathbf{F})^{-1} \mathbf{G} + \mathbf{D}} \quad (4.4.10)$$

soit une relation analogue à (4.4.9), en remplaçant p par z , et les matrices $\mathbf{A}$ et $\mathbf{B}$ par $\mathbf{F}$ et $\mathbf{G}$ respectivement.

4.4.4.2 Passage transfert $\rightarrow$ état

Cette section présente diverses approches pour passer d'une fonction de transfert à une représentation d'état d'un système à temps continu. Les passages transfert $\rightarrow$ état pour des systèmes à temps discret se font suivant une démarche analogue au cas continu.

Nous allons voir ici comment passer d'une représentation externe caractérisée par la matrice de transfert $\mathbf{H}(p)$ à une représentation interne. Afin de simplifier les calculs, on se limitera ici aux systèmes à une entrée et une sortie. La généralisation au cas multi-entrées et/ou multi-sorties peut se faire sur chaque fonction de transfert $H_{ij}(p)$ au prix d'une complexité accrue.

Nous allons voir deux représentations d'état appelées forme compagnon pour la commande et forme compagnon pour l'observation.

4.4.4.2.1 Forme compagnon pour la commande Posons :

$$H(p) = \frac{b_0 + b_1p + \dots + b_m p^m}{a_0 + a_1p + \dots + p^n} = \frac{N(p)}{D(p)} = \frac{Y(p)}{U(p)} = \frac{Y(p)W(p)}{W(p)U(p)}$$

On se place dans un premier temps dans le cas où m est strictement inférieur à n .

En identifiant $\frac{Y(p)}{W(p)}$ avec $N(p)$ et $\frac{U(p)}{W(p)}$ avec $D(p)$ on trouve :

$$\begin{cases} Y(p) &= (b_0 + b_1p + \dots + b_m p^m) W(p) \\ U(p) &= (a_0 + a_1p + \dots + p^n) W(p) \end{cases}$$

soit dans le domaine temporel :

$$\begin{cases} y(t) &= b_0 w(t) + b_1 \frac{dw(t)}{dt} + \dots + b_m \frac{d^m w(t)}{dt^m} \\ \frac{d^n w(t)}{dt^n} &= -a_0 w(t) - a_1 \frac{dw(t)}{dt} - \dots - a_{n-1} \frac{d^{n-1} w(t)}{dt^{n-1}} + u(t) \end{cases}$$

ce qui conduit à choisir comme variables d'état du système :

$$\begin{aligned} x_1(t) &= w(t) \\ x_2(t) &= \frac{dx_1(t)}{dt} = \frac{dw(t)}{dt} \\ x_3(t) &= \frac{dx_2(t)}{dt} = \frac{d^2 w(t)}{dt^2} \\ &\vdots \\ x_n(t) &= \frac{dx_{n-1}(t)}{dt} = \frac{d^{n-1} w(t)}{dt^{n-1}} \end{aligned}$$

ce qui donne sous forme matricielle :

$$\frac{d}{dt} \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_n(t) \end{pmatrix} = \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 \\ 0 & 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & 1 \\ -a_0 & -a_1 & -a_2 & \cdots & -a_{n-1} \end{pmatrix} \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_n(t) \end{pmatrix} + \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix} u(t)$$

c'est-à-dire :

$$\frac{d\mathbf{x}(t)}{dt} = \mathbf{A}\mathbf{x}(t) + \mathbf{B}u(t) \text{ (première équation d'état)}$$

et :

$$y(t) = (b_0 \quad b_1 \quad \cdots \quad b_m \quad 0 \quad \cdots \quad 0) \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_n(t) \end{pmatrix}$$

c'est-à-dire :

$$y(t) = \mathbf{C}\mathbf{x}(t) \text{ (deuxième équation d'état avec } D = 0)$$

Cette représentation d'état est appelée **forme compagnon pour la commande**. L'implantation de cette forme correspond à la figure 4.6.

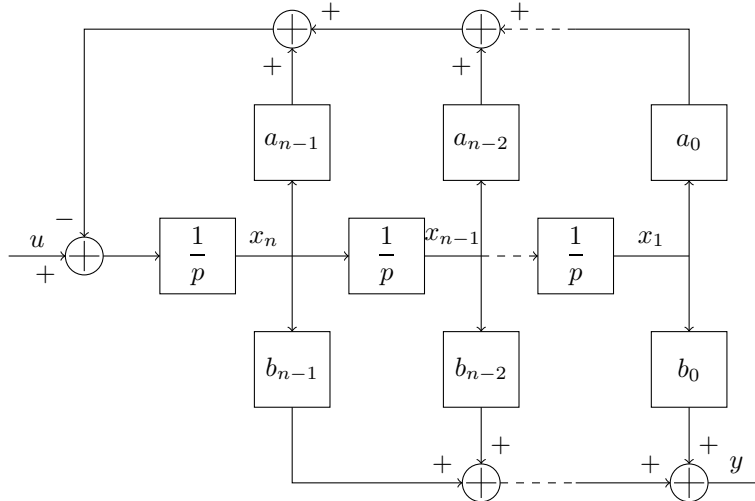


FIGURE 4.6 – Forme compagnon pour la commande

Remarque 4.4.6. Cas où $m = n$:

Si $m = n$ (degré du numérateur de $H(p)$ égal au degré du dénominateur), on peut écrire $H(p)$ sous la forme $H(p) = H_1(p) + K$ avec $H_1(p)$ nulle à l'infini.

On peut alors appliquer la méthode précédente à $H_1(p)$; il suffit d'ajouter au signal de sortie dans la deuxième équation du système d'équations d'état, le terme $Ku(t)$.

4.4.4.2.2 Forme compagnon pour l'observation On pose :

$$H(p) = \frac{b_0 + b_1p + \cdots + b_{n-1}p^{n-1}}{a_0 + a_1p + \cdots + p^n} = \frac{N(p)}{D(p)} = \frac{Y(p)}{U(p)}$$

en divisant $N(p)$ et $D(p)$ par p^n , on obtient :

$$Y(p) = \left(b_{n-1} \frac{1}{p} + \cdots + b_1 \frac{1}{p^{n-1}} + b_0 \frac{1}{p^n} \right) U(p) - \left(a_{n-1} \frac{1}{p} + \cdots + a_1 \frac{1}{p^{n-1}} + a_0 \frac{1}{p^n} \right) Y(p)$$

que l'on peut écrire sous la forme :

$$Y(p) = \frac{1}{p} \left(\zeta_{n-1} + \frac{1}{p} \left(\zeta_{n-2} + \frac{1}{p} \left(\zeta_{n-3} + \frac{1}{p} \left(\cdots + \frac{1}{p} \zeta_0 \right) \right) \right) \right) \text{ avec } \zeta_j = b_j U(p) - a_j Y(p)$$

en posant $X_n(p) = \mathcal{L}\{x_n(t)\}(p) = Y(p)$, on obtient les relations suivantes :

$$\begin{cases} X_n(p) &= \frac{1}{p} (b_{n-1}U(p) - a_{n-1}X_n(p) + X_{n-1}(p)) \\ X_{n-1}(p) &= \frac{1}{p} (b_{n-2}U(p) - a_{n-2}X_n(p) + X_{n-2}(p)) \\ \vdots & \\ X_1(p) &= \frac{1}{p} (b_0U(p) - a_0X_n(p)) \end{cases}$$

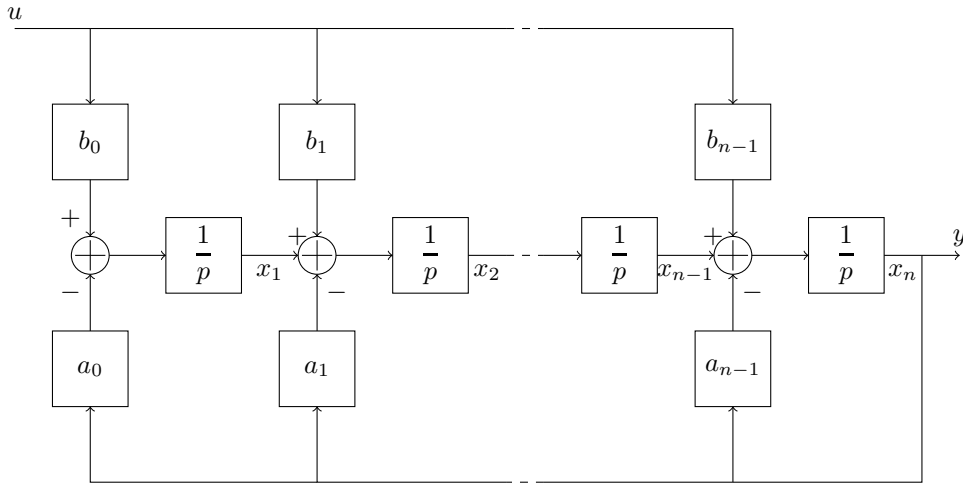


FIGURE 4.7 – Forme compagnon pour l'observation

Ces n équations donnent dans le domaine temporel :

$$\begin{cases} \frac{dx_1(t)}{dt} &= b_0 u(t) - a_0 x_n(t) \\ \frac{dx_2(t)}{dt} &= b_1 u(t) - a_1 x_n(t) + x_1(t) \\ \frac{dx_3(t)}{dt} &= b_2 u(t) - a_2 x_n(t) + x_2(t) \\ \vdots & \\ \frac{dx_n(t)}{dt} &= b_{n-1} u(t) - a_{n-1} x_n(t) + x_{n-1}(t) \end{cases}$$

ce qui s'écrit sous forme matricielle :

$$\frac{d}{dt} \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_{n-1}(t) \\ x_n(t) \end{pmatrix} = \begin{pmatrix} 0 & 0 & \cdots & 0 & -a_0 \\ 1 & 0 & \cdots & 0 & -a_1 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \cdots & 0 & -a_{n-2} \\ 0 & 0 & \cdots & 1 & -a_{n-1} \end{pmatrix} \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_{n-1}(t) \\ x_n(t) \end{pmatrix} + \begin{pmatrix} b_0 \\ b_1 \\ \vdots \\ b_{n-2} \\ b_{n-1} \end{pmatrix} u(t)$$

c'est-à-dire :

$$\frac{d\mathbf{x}}{dt} = \mathbf{A}\mathbf{x}(t) + \mathbf{B}u(t) \text{ (première équation d'état)}$$

et :

$$y(t) = \begin{pmatrix} 0 & 0 & \cdots & 0 & 1 \end{pmatrix} \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_{n-1}(t) \\ x_n(t) \end{pmatrix}$$

c'est-à-dire :

$$y(t) = \mathbf{C}\mathbf{x}(t) \text{ (deuxième équation d'état avec } \mathbf{D} = 0)$$

Cette représentation d'état est appelée **forme compagnon pour l'observation** et est implantée par le schéma bloc de la figure 4.7.

Chapitre 5

Analyse des performances des systèmes dynamiques

5.1 Introduction

Dans ce chapitre sont étudiés les comportements temporel et fréquentiel de systèmes linéaires invariants élémentaires du 1^{er} ou du 2^{ème} ordre, analogiques ou numériques. Ils interviennent dans la décomposition des fonctions de transfert des systèmes d'ordre plus élevé, facilitant ainsi l'étude de celles-ci, comme le montre par exemple la représentation de Bode pour les systèmes analogiques. Des exemples d'utilisation de cette représentation sont proposés. Le diagramme de Bode d'un système numérique avec un pôle du premier ordre ainsi que le diagramme de Bode d'un système du second ordre numérique sont analysés. Puis, les diagrammes de Bode des systèmes passe-tout, des systèmes à déphasage minimal et des systèmes avec retard pur sont présentés. On peut plus généralement envisager une approche des comportements fréquentiel et temporel des systèmes analogiques ou numériques en considérant la position des pôles et des zéros de la fonction de transfert. Une analyse temporelle des systèmes du premier ordre et du second ordre est également proposée dans ce chapitre.

5.2 Diagramme de Bode des systèmes à temps continu

5.2.1 Définition

Une fonction de transfert, dite à constantes localisées, peut se décomposer sous la forme :

$$H(p) = K p^q e^{-\tau p} \frac{\prod_{k=1}^{n_k} \left(1 + \frac{p}{\omega_k}\right) \prod_{l=1}^{n_l} \left(1 + 2\zeta_l \frac{p}{\omega_l} + \left(\frac{p}{\omega_l}\right)^2\right)}{\prod_{m=1}^{n_m} \left(1 + \frac{p}{\omega_m}\right) \prod_{n=1}^{n_n} \left(1 + 2\zeta_n \frac{p}{\omega_n} + \left(\frac{p}{\omega_n}\right)^2\right)} \quad (5.2.1)$$

avec le gain $K \in \mathbb{R}$, le comportement basse fréquence donné par $q \in \mathbb{Z}$, le retard pur $\tau \in \mathbb{R}$, les pulsations de brisure pour les termes du premier ordre $\omega_k, \omega_m \in \mathbb{R}^*$, les pulsations de brisure pour les termes du deuxième ordre $\omega_l, \omega_n \in \mathbb{R}_+^*$ et les amortissements $|\zeta_l|, |\zeta_n| < 1$.

La représentation de Bode donne les variations de $20 \log(|H(j\omega)|)$ (logarithme décimal et module exprimé en dB) et de $\arg(H(j\omega))$ (phase exprimée en degrés) en fonction de la pulsation ω , l'échelle des pulsations étant logarithmique. L'étude des asymptotes de ces deux fonctions donne le tracé asymptotique de Bode, utilisé notamment en Automatique.

Le module en décibel (respectivement la phase) de $H(j\omega)$ est la somme des modules en décibel (respectivement des phases) des termes élémentaires (du 1^{er}, du 2^{ème} ordre et de la forme $K p^q e^{-\tau p}$ avec q entier relatif et τ un nombre réel) intervenant dans la décomposition de $H(p)$. Dans la suite nous allons d'abord traiter le cas des systèmes sans retard (avec $\tau = 0$) et ensuite le cas des systèmes à retard pur sera abordé dans la section 5.4.3.

Il convient donc de donner dans un premier temps le comportement fréquentiel des systèmes du 1^{er} et du 2^{ème} ordre afin de proposer ensuite une méthode permettant d'obtenir le tracé asymptotique de Bode de $H(p)$.

5.2.2 Systèmes du 1^{er} ordre

5.2.2.1 Systèmes avec un zéro

La fonction de transfert :

$$H(p) = 1 + \frac{p}{\omega_0} \quad (5.2.2)$$

possède la représentation de Bode tracée en figure 5.1 pour $\omega_0 > 0$. On a représenté en trait fin le tracé réel et en trait gras le tracé asymptotique. Pour $\omega_0 < 0$, le module est identique et la phase est l'opposée de celle représentée ici.

Par convention, la pente de 20 dB/décade est qualifiée de pente +1.

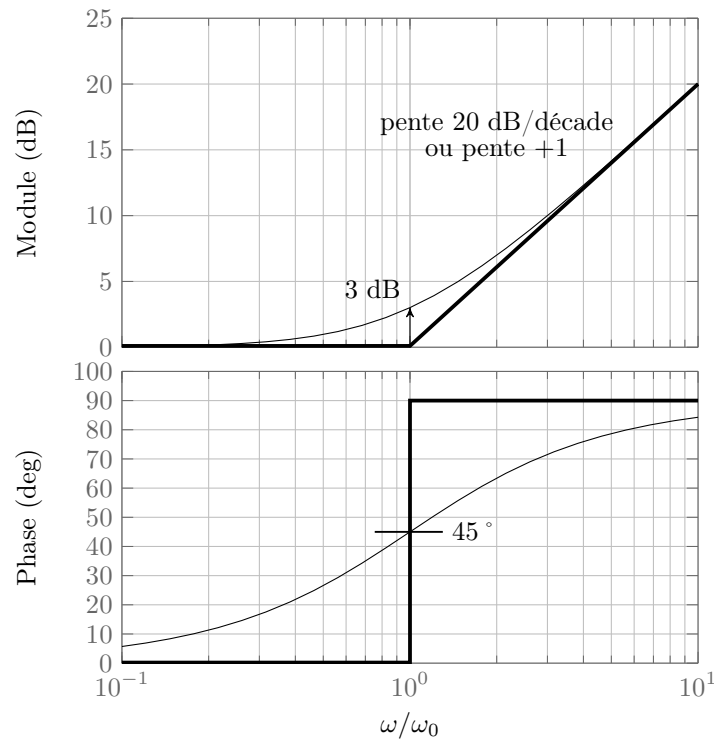


FIGURE 5.1 – Tracé de Bode d'un système avec un zéro

5.2.2.2 Système avec un pôle

La fonction de transfert :

$$H(p) = \frac{1}{1 + \frac{p}{\omega_0}} = \frac{\omega_0}{p + \omega_0} \quad (5.2.3)$$

possède la représentation de Bode tracée en figure 5.2 pour $\omega_0 > 0$.

On a représenté en trait fin le tracé réel et en trait gras le tracé asymptotique. Pour $\omega_0 < 0$, le module est identique et la phase est l'opposée de celle représentée ici.

Par convention, la pente de -20 dB/décade est qualifiée de pente -1.

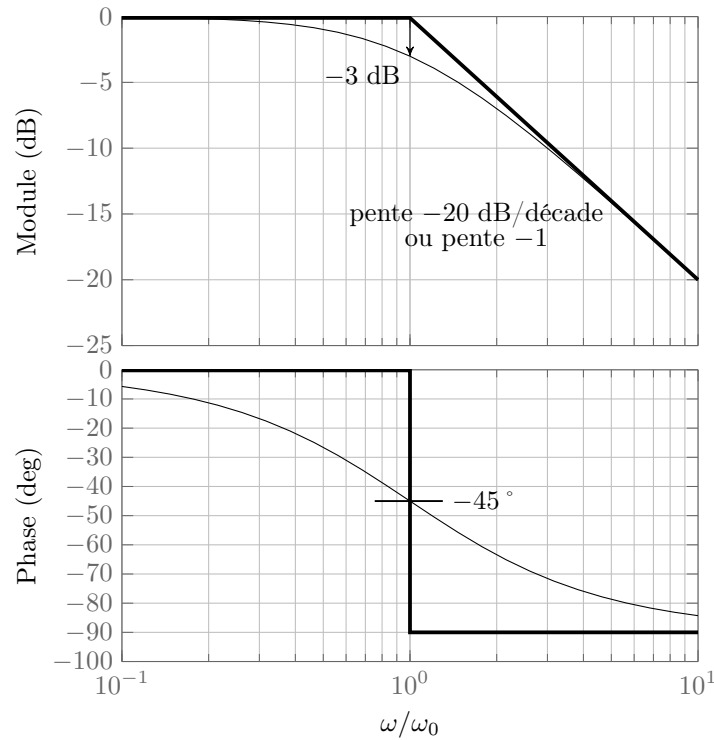


FIGURE 5.2 – Tracé de Bode d'un système avec un pôle

5.2.3 Systèmes du 2^{ème} ordre

La fonction de transfert :

$$H(p) = \frac{1}{1 + 2\zeta \frac{p}{\omega_0} + \left(\frac{p}{\omega_0}\right)^2} = \frac{\omega_0^2}{p^2 + 2\zeta\omega_0 p + \omega_0^2} \quad (5.2.4)$$

possède la représentation de Bode tracée en figure 5.3 pour $\omega_0 > 0$ et :

$$\zeta \in \{0, 1; 0, 2; 0, 3; 0, 4; 0, 5; 0, 6; 0, 7; 0, 8; 0, 9; 1; 2; 3; 5; 7; 10; 20\}$$

Remarque 5.2.1. *Le lecteur pourra remarquer que pour $\zeta > 1$, le système du "second ordre" se factorise en deux systèmes du premier ordre.*

La pulsation ω_0 est appelée pulsation propre et ζ est l'amortissement. Lorsque $\zeta < \frac{\sqrt{2}}{2}$, le module présente un maximum :

$$|H(j\omega)|_{\max} = \frac{1}{2\zeta\sqrt{1-\zeta^2}} \text{ pour } \frac{\omega}{\omega_0} = \sqrt{1-2\zeta^2} \quad (5.2.5)$$

On a représenté en trait fin le tracé réel et en trait gras le tracé asymptotique. Par convention, la pente de -40 dB/décade est qualifiée de pente -2 . Pour des valeurs de ζ opposées, le module est identique et la phase est l'opposée de celle représentée ici.

Pour obtenir le tracé de :

$$H(p) = 1 + 2\zeta \frac{p}{\omega_0} + \left(\frac{p}{\omega_0}\right)^2$$

il suffit de prendre l'opposé du module en dB et de la phase de la fonction de transfert précédente.

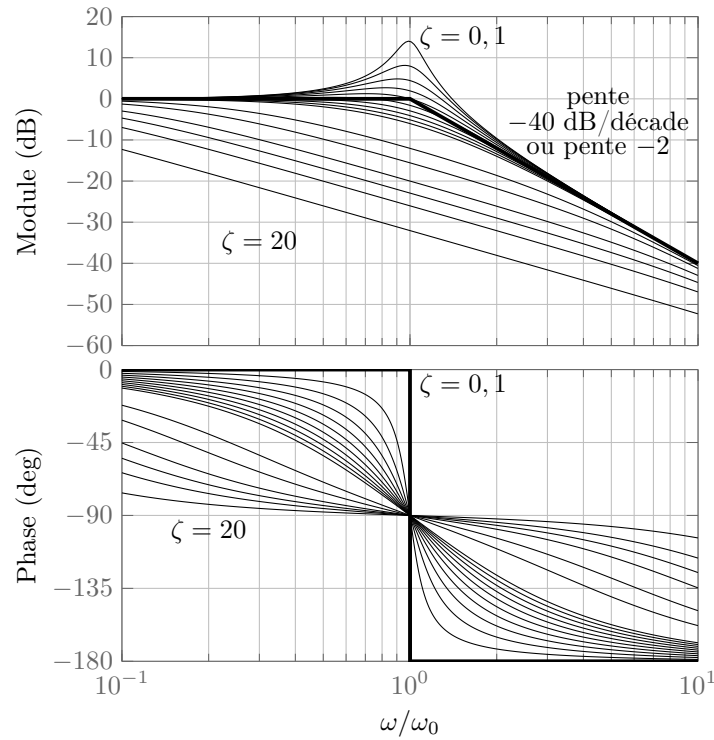


FIGURE 5.3 – Tracé de Bode d'un système du second ordre

5.2.4 Méthode pratique de tracé asymptotique

La fonction de transfert $H(p)$ sans retard ($\tau = 0$) présentée en 5.2.1 est décomposée en termes élémentaires décrits en 5.2.2 et 5.2.3. Plutôt que de faire la somme des tracés asymptotiques de ces termes pour obtenir le tracé de Bode associé à $H(p)$, on prend en compte leur contribution au fur et à mesure en faisant croître la pulsation ω de zéro jusqu'à l'infini.

- Pour $\omega \rightarrow 0$, $H(j\omega) \rightarrow Kj^q\omega^q$. En module, l'asymptote est donc une droite de pente $q \times 20$ dB/décade (ou $+q$) coupant l'axe 0 dB à la pulsation $\omega = \left| \frac{1}{K} \right|^{\frac{1}{q}}$. En phase, l'asymptote est une droite horizontale d'ordonnée $q \times 90^\circ$ si $K > 0$ ou $q \times 90^\circ + 180^\circ$ si $K < 0$.

Ensuite on fait croître ω . Quand une des pulsations de brisure $|\omega_k|$, ω_l , $|\omega_m|$, ω_n est rencontrée, on fait évoluer le tracé en adoptant les règles suivantes :

- Pour $\omega = |\omega_k|$, le terme à prendre en compte est $1 + \frac{p}{\omega_k}$ au numérateur. En module, le tracé asymptotique subit un changement de pente de $+20$ dB/décade (ou $+1$). En phase, l'asymptote subit un saut de $90^\circ \times \text{sgn}(\omega_k)$.
- Pour $\omega = |\omega_m|$, le terme à prendre en compte est $1 + \frac{p}{\omega_m}$ au dénominateur. En module, le tracé asymptotique subit un changement de pente de -20 dB/décade (ou -1). En phase, l'asymptote subit un saut de $-90^\circ \times \text{sgn}(\omega_m)$.
- Pour $\omega = \omega_l$, le terme à prendre en compte est $1 + 2\zeta_l \frac{p}{\omega_l} + \left(\frac{p}{\omega_l} \right)^2$ au numérateur. En module, le tracé asymptotique subit un changement de pente de $+40$ dB/décade (ou $+2$). En phase, l'asymptote subit un saut de $180^\circ \times \text{sgn}(\zeta_l)$.
- Pour $\omega = \omega_n$, le terme à prendre en compte est $1 + 2\zeta_n \frac{p}{\omega_n} + \left(\frac{p}{\omega_n} \right)^2$ au dénominateur. En module, le tracé asymptotique subit un changement de pente de -40 dB/décade (ou -2). En phase, l'asymptote subit un saut de $-180^\circ \times \text{sgn}(\zeta_n)$.

Exemple 5.2.1 (Diagramme de Bode). *On s'intéresse au tracé de Bode de la fonction de transfert :*

$$H(p) = \frac{(1+p)(1+50p)(1+5p)}{(1+0,05p)(1+8p+100p^2)} \quad (5.2.6)$$

Pour faire apparaître les pulsations où vont intervenir les évolutions du tracé (pulsations de brisure), on peut réécrire la fonction de transfert :

$$H(p) = \frac{(1+p) \left(1 + \frac{p}{0,02}\right) \left(1 + \frac{p}{0,2}\right)}{\left(1 + \frac{p}{20}\right) \left(1 + 2 \cdot 0,4 \cdot \frac{p}{0,1} + \left(\frac{p}{0,1}\right)^2\right)} \quad (5.2.7)$$

On obtient alors le tracé de la figure 5.4 après application des règles décrites en 5.2.4.

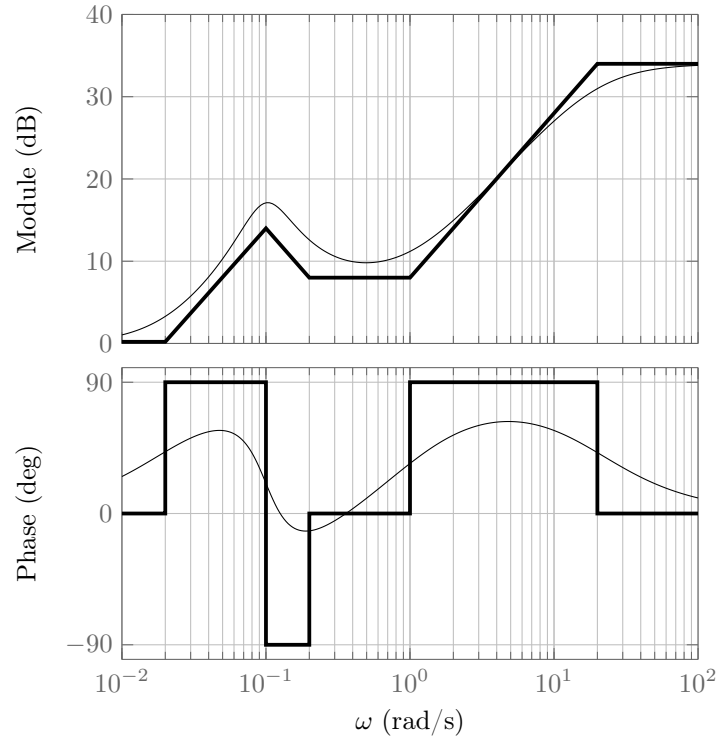


FIGURE 5.4 – Exemple de tracé de Bode

On peut constater que le tracé réel s'éloigne d'autant plus du tracé asymptotique que les pulsations de brisure sont proches.

Exemple 5.2.2 (Exemple d'application). *En Automatique, on est fréquemment amené à modéliser un système physique en le considérant comme linéaire invariant. Pour un moteur à courant continu entraînant une charge mécanique, il est possible par exemple de déterminer une relation simplifiée entre la tension d'alimentation du moteur et sa vitesse de rotation :*

$$\frac{\Omega_m(p)}{U_m(p)} = \frac{K_v}{1 + \tau p}$$

Les paramètres de cette fonction de transfert peuvent être déterminés en procédant à une analyse harmonique (voir 10.3.1 : mesure de l'amplitude et du déphasage de la sortie pour une tension $u_m(t) = A \cos(\omega t)$ avec plusieurs valeurs de la pulsation ω). Le modèle est aussi validé par l'analyse des équations différentielles régissant le système. On peut en déduire la relation entre la tension d'alimentation et la position angulaire :

$$H(p) = \frac{\Theta_m(p)}{U_m(p)} = \frac{K_v}{p(1 + \tau p)}$$

puisque $\Omega_m(p) = p\Theta_m(p)$. Un schéma simple d'asservissement où on souhaite que la position $\theta_m(t)$ soit aussi proche que possible d'une commande $e(t)$ peut être proposé en figure 5.5. Un dispositif de mesure de position angulaire (potentiomètre) délivre une tension égale à cette position. Le gain K est réglable et on se propose de lui donner une valeur convenable. La fonction de transfert de l'ensemble vaut :

$$H_{ensemble}(p) = \frac{\Theta_m(p)}{E(p)} = \frac{KH(p)}{1 + KH(p)} = \frac{1}{1 + \frac{p}{KK_v} + \frac{\tau p^2}{KK_v}} \quad (5.2.8)$$

(une analyse sommaire du schéma de la figure 5.5 permet d'obtenir cette fonction de transfert).

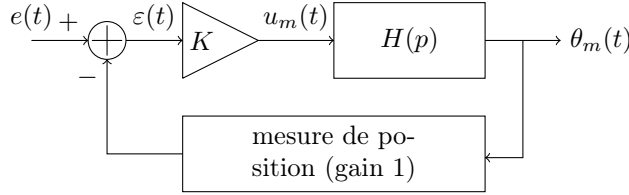


FIGURE 5.5 – Schéma d'un asservissement de position

Pour se ramener à l'étude du paragraphe 5.2.3, on peut l'écrire sous la forme :

$$H_{ensemble}(p) = \frac{1}{1 + 2\zeta \frac{p}{\omega_0} + \left(\frac{p}{\omega_0}\right)^2} \text{ avec } \omega_0 = \sqrt{\frac{KK_v}{\tau}} \text{ et } \zeta = \frac{1}{2\sqrt{KK_v\tau}}$$

Pour faire en sorte que la sortie suive au mieux la commande, il faudrait que $H_{ensemble}(j\omega)$ reste proche de 1 pour des pulsations les plus élevées possibles, soit une valeur de ω_0 la plus élevée possible. En effet, le tracé asymptotique nous indique que $H_{ensemble}(j\omega) \approx 1$ pour $\omega < \omega_0$. On pourrait pour cela augmenter arbitrairement le gain K . Cependant, si K est trop grand, l'amortissement ζ devient trop faible et on s'éloigne trop de la condition $H_{ensemble}(j\omega) \approx 1$ pour $\omega < \omega_0$ car une résonance apparaît dans la réponse en fréquence.

Une valeur convenable de l'amortissement (correspondant à un tracé réel proche du tracé asymptotique) est par exemple $\zeta = \frac{\sqrt{2}}{2}$ (figure 5.11), d'où la valeur :

$$K = \frac{1}{2K_v\tau}$$

La bande passante du système vaut alors $\omega_0 = \frac{1}{\tau\sqrt{2}}$.

5.3 Diagramme de Bode des systèmes numériques

5.3.1 1^{er} ordre tout pôle

Un système du premier ordre numérique tout pôle possède une fonction de transfert :

$$H(z) = \frac{1}{1 - az^{-1}} \quad (5.3.1)$$

L'équation aux différences qui lui est associée est :

$$y(n) = u(n) + ay(n-1)$$

La réponse en fréquence vaut :

$$H(e^{2\pi j\nu}) = \frac{1}{1 - ae^{-2\pi j\nu}}$$

On peut représenter en figure 5.6 le module de cette réponse pour différentes valeurs de a (a varie de $-0,9$ à $0,9$ par pas de $0,1$) et faire une interprétation géométrique de cette représentation en fonction de la position du pôle a de la fonction $H(z)$. En effet :

$$H(z) = \frac{z}{z-a} \text{ soit } |H(e^{2\pi j\nu})| = \frac{1}{|e^{2\pi j\nu} - a|} = \frac{1}{MP}$$

où P est le point dans le plan complexe correspondant au pôle et M un point se déplaçant sur le cercle unité. On observe directement ainsi un comportement passe-bas pour $a > 0$ et passe haut pour $a < 0$.

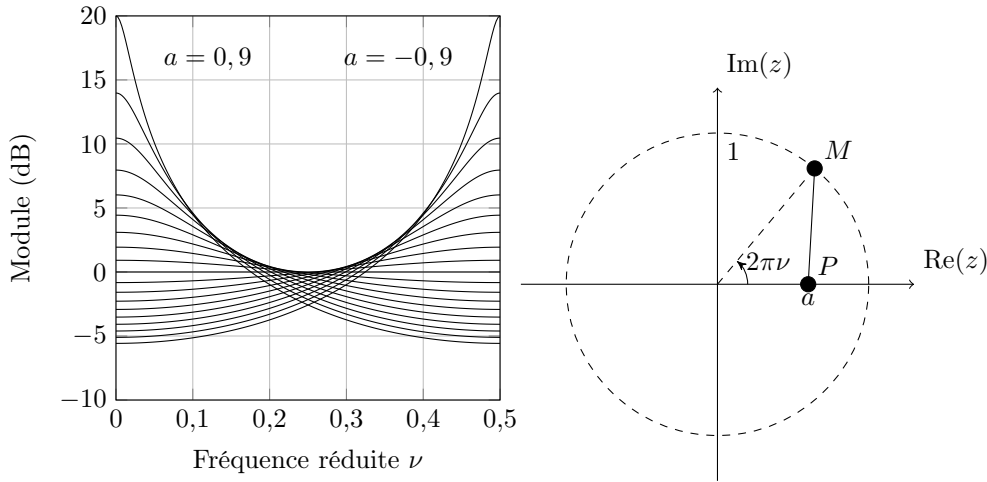


FIGURE 5.6 – Réponse en fréquence d'un système numérique du premier ordre

5.3.2 2^{ème} ordre tout pôle

Un système du second ordre numérique tout pôle possède une fonction de transfert :

$$H(z) = \frac{1}{1 + a_1 z^{-1} + a_2 z^{-2}} = \frac{1}{(1 - r e^{j\theta} z^{-1})(1 - r e^{-j\theta} z^{-1})} = \frac{1}{1 - 2r \cos \theta z^{-1} + r^2 z^{-2}} \quad (5.3.2)$$

(on n'envisage pas le cas où les deux pôles sont réels, ce qui reviendrait à une décomposition en deux systèmes du premier ordre).

L'équation aux différences qui lui est associée est :

$$y(n) = u(n) - a_1 y(n-1) - a_2 y(n-2)$$

La réponse en fréquence vaut :

$$H(e^{2\pi j\nu}) = \frac{1}{1 + a_1 e^{-2\pi j\nu} + a_2 e^{-4\pi j\nu}}$$

On a représenté en figure 5.7 le module des réponses en fréquences associées à des valeurs de $\theta = 2\pi\nu_0$ avec les valeurs de ν_0 et r suivantes :

$$\nu_0 \in \{0, 1; 0, 15; 0, 2; 0, 25; 0, 3; 0, 35; 0, 4\}$$

$$r \in \{0, 65; 0, 7; 0, 75; 0, 8; 0, 85; 0, 9; 0, 95\}$$

On constate que pour des valeurs de r proches de 1 et une valeur de $\theta = 2\pi\nu_0$, la réponse présente une résonance pour une fréquence proche de ν_0 .

En effet, les deux fréquences $\nu = 0$ et $\nu = \frac{1}{2}$ correspondent à des extremums de la réponse en fréquence et il y a une troisième fréquence extrême lorsque :

$$\cos \theta < \frac{2r}{1 + r^2}$$

Cette fréquence $\nu_{\max}$ vérifie alors :

$$\cos(2\pi\nu_{\max}) = \cos\theta \frac{1+r^2}{2r}$$

et le module de la réponse à cette fréquence vaut :

$$|H(e^{2\pi j\nu_{\max}})| = \frac{1}{(1-r^2)\sin\theta}$$

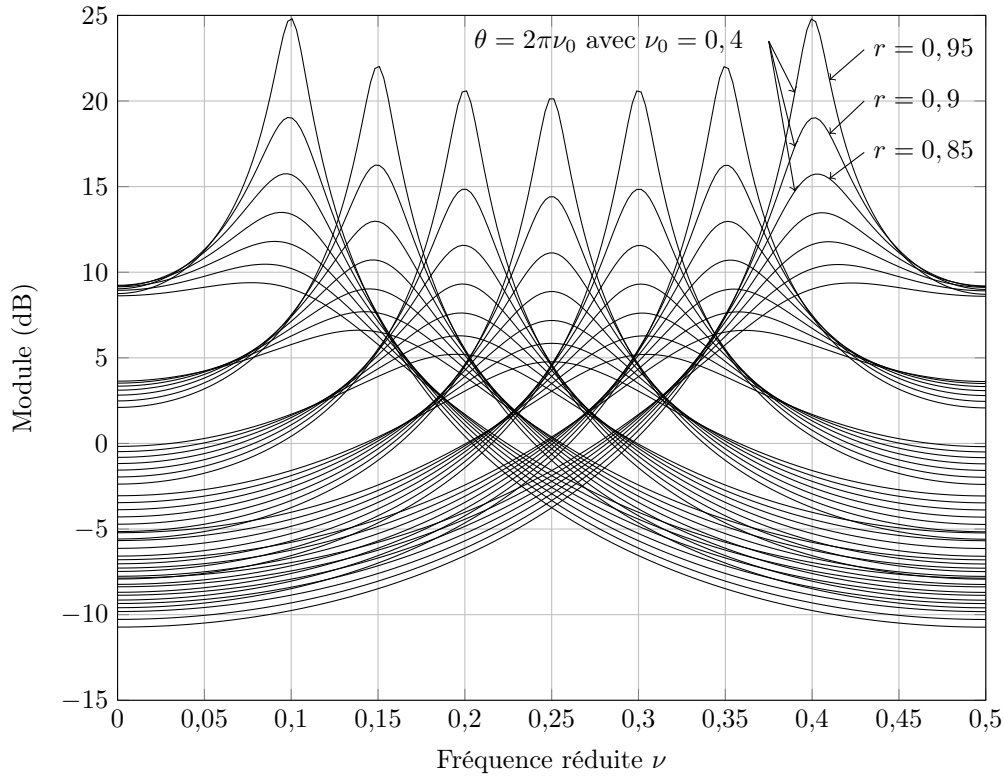


FIGURE 5.7 – Réponse en fréquence d'un système numérique du second ordre

On peut par conséquent connaître le comportement fréquentiel à partir des pôles de la fonction de transfert (figure 5.8) et faire une interprétation géométrique de la réponse en fréquence que l'on peut exprimer sous la forme :

$$|H(e^{2\pi j\nu})| = \frac{1}{MP_1 \cdot MP_2}$$

5.4 Diagramme de Bode des systèmes passe-tout, à déphasage minimal ou avec retard pur

5.4.1 Systèmes passe-tout

Un *système passe-tout* analogique vérifie :

$$|H(j\omega)| = 1 \text{ pour tout } \omega$$

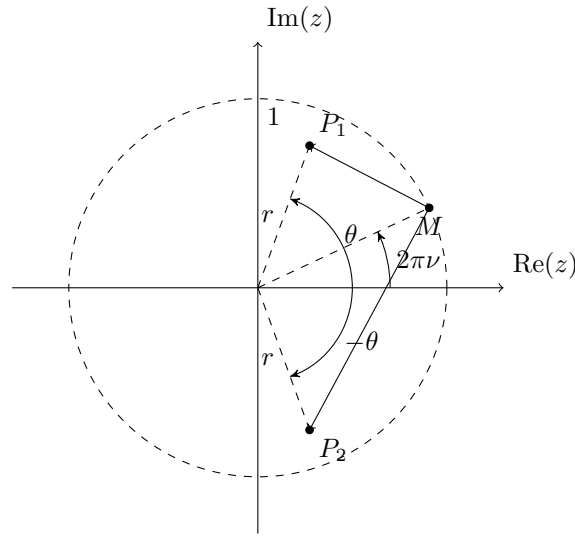


FIGURE 5.8 – Position des pôles de la fonction de transfert 5.3.2

sans contrainte sur la phase $\arg(H(j\omega))$. Cette situation survient lorsqu'un pôle P est compensé par un zéro $-\overline{P}$. On obtient alors des filtres de la forme :

$$H(p) = \frac{\prod_{k=1}^N (p + \overline{P}_k)}{\prod_{k=1}^N (p - P_k)} \quad (5.4.1)$$

avec la contrainte de stabilité $\text{Re}(P_k) < 0$, appelés aussi déphaseurs analogiques purs.

Un système passe-tout numérique vérifie :

$$|H(e^{2\pi j\nu})| = 1 \text{ pour tout } \nu$$

sans contrainte sur la phase $\arg(H(e^{2\pi j\nu}))$. Cette situation survient lorsqu'un pôle P est compensé par un zéro $Z = \frac{1}{\overline{P}}$. On obtient alors des filtres de la forme :

$$H(z) = \frac{\prod_{k=1}^N (\overline{P}_k - z^{-1})}{\prod_{k=1}^N (1 - P_k z^{-1})} \quad (5.4.2)$$

avec la contrainte de stabilité $|P_k| < 1$, appelés aussi déphaseurs numériques purs.

5.4.2 Systèmes à déphasage minimal

Considérons par exemple la fonction de transfert :

$$H(p) = \frac{(1+p) \left(1 + \frac{p}{10}\right) (1+10p)}{\left(1 + \frac{p}{5}\right) (1+2p) (1+20p)} \quad (5.4.3)$$

On représente en figure 5.9 le module et la phase de $H(j\omega)$ (tracé en gras pour la phase). Toutes les autres fonctions de la forme :

$$H(p) = \frac{(1 \pm p) \left(1 \pm \frac{p}{10}\right) (1 \pm 10p)}{\left(1 + \frac{p}{5}\right) (1+2p) (1+20p)} \quad (5.4.4)$$

ont le même module (figure 5.9).

On les obtient en effet en multipliant la fonction de transfert initiale par un déphaseur pur étudié au paragraphe 5.4.1. Les phases associées sont représentées en trait fin sur le graphe. On constate que la phase est maximale lorsque tous les zéros sont à partie réelle négative (en considérant que le déphasage est l'opposé de la phase étudiée, on a bien un *déphasage minimal*).

Définition 5.4.1

Plus généralement, on appelle *filtre analogique à déphasage minimal* un filtre dont tous les zéros sont à partie réelle négative.

Pour un filtre quelconque, on peut toujours se ramener à un tel filtre en multipliant la fonction de transfert initiale par le déphaseur analogique pur approprié.

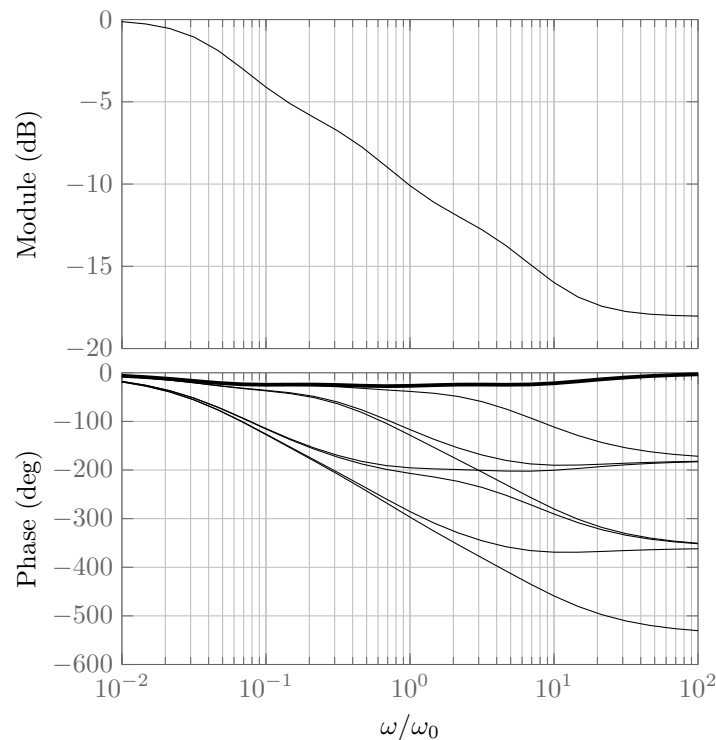


FIGURE 5.9 – Réponse en fréquence d'une famille de filtres ayant le même module

Dans le domaine numérique, on pourrait faire la même constatation. On appelle alors filtre numérique à déphasage minimal un filtre dont tous les zéros sont de module inférieur à 1. Pour un filtre quelconque, on peut toujours se ramener à un tel filtre en multipliant la fonction de transfert initiale par le déphaseur numérique pur approprié.

5.4.3 Systèmes avec retard pur

La sortie d'un système peut ne pas évoluer instantanément lorsque l'on fait évoluer son entrée. Il est alors nécessaire de considérer des fonctions de transfert de la forme :

$$G(p) = e^{-\tau p} H(p) \quad (5.4.5)$$

où $H(p)$ est une fonction de transfert à constantes localisées (paragraphe 5.2.1) et où τ est un *retard pur*. Si $y(t)$ est la réponse indicielle du système $H(p)$ seul, la réponse indicielle du système $G(p)$ est alors $y(t - \tau)$. On peut ainsi rendre compte d'un comportement temporel, difficilement modélisable par une simple fonction de transfert à constantes localisées (voir identification par la méthode de Strejc, paragraphe 10.2.3).

Pour le tracé de Bode d'un tel système, la règle est la suivante :

- le module est le même que $|H(j\omega)|$ car $|e^{-j\omega\tau}| = 1$;
- la phase du système est $\arg(H(j\omega)) - 180\frac{\omega\tau}{\pi}$ exprimée en degré.

En revanche, la notion de retard pur ne présente aucune difficulté pour un système numérique ; on obtiendrait une fonction de transfert $z^{-\tau_d}H(z)$ où le retard τ_d s'exprime en nombre d'échantillons.

5.5 Analyse temporelle

5.5.1 Systèmes analogiques

5.5.1.1 1^{er} ordre tout pôle

L'étude du comportement fréquentiel de :

$$H(p) = \frac{1}{1 + \frac{p}{\omega_0}} = \frac{\omega_0}{p + \omega_0} \quad (5.5.1)$$

a été faite au 5.2.2. On peut s'intéresser aussi à la réponse temporelle d'un tel système, en particulier la réponse à un échelon $\mathbb{1}(t)$. Il s'agit donc de déterminer le signal temporel associé à une transformée de Laplace égale à :

$$\frac{H(p)}{p} = \frac{\omega_0}{p(p + \omega_0)}$$

La décomposition en éléments simples donne accès au résultat :

$$Y(p) = \frac{H(p)}{p} = \frac{1}{p} - \frac{1}{p + \omega_0}$$

L'expression temporelle de la réponse indicielle est alors :

$$y(t) = (1 - e^{-\omega_0 t}) \mathbb{1}(t) \quad (5.5.2)$$

On constate ainsi que le pôle $P_0 = -\omega_0$ détermine directement le comportement temporel du système.

Si l'on note la *constante de temps* $T = \frac{1}{\omega_0}$, la fonction de transfert (5.5.1) devient :

$$H(p) = \frac{1}{1 + Tp} \quad (5.5.3)$$

La *réponse indicielle* (5.5.2) (représentée figure 5.10) peut s'écrire en fonction de T comme suit :

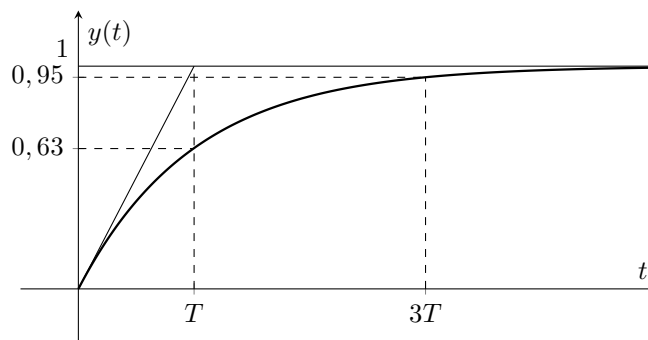


FIGURE 5.10 – Réponse indicielle d'une fonction de transfert d'ordre 1 sans zéro avec gain unitaire

$$y(t) = \left(1 - e^{-\frac{t}{T}}\right) \mathbb{1}(t) \quad (5.5.4)$$

Le *temps de réponse* à 5% de ce système à est de $3T$ (voir figure 5.10).

5.5.1.2 2^{ème} ordre tout pôle

L'étude du comportement fréquentiel de :

$$H(p) = \frac{1}{1 + 2\zeta \frac{p}{\omega_0} + \left(\frac{p}{\omega_0}\right)^2} \quad (5.5.5)$$

a été faite au paragraphe 5.2.3. On peut s'intéresser aussi à la réponse temporelle d'un tel système, en particulier la réponse à un échelon unitaire $\mathbb{1}(t)$. Il s'agit donc de déterminer le signal temporel associé à une transformée de Laplace égale à :

$$\frac{H(p)}{p} = \frac{\omega_0^2}{p(p^2 + 2\zeta\omega_0 p + \omega_0^2)} = \frac{\omega_0^2}{p(p - P_1)(p - P_2)}$$

La décomposition en éléments simples donne accès au résultat :

$$Y(p) = \frac{H(p)}{p} = \frac{1}{p} + \frac{\omega_0^2}{P_1 - P_2} \left(\frac{1}{P_1(p - P_1)} - \frac{1}{P_2(p - P_2)} \right)$$

L'expression temporelle de la *réponse indicielle* est alors :

$$y(t) = \left(1 + \frac{\omega_0^2}{P_1 - P_2} \left(\frac{1}{P_1} e^{P_1 t} - \frac{1}{P_2} e^{P_2 t} \right) \right) \mathbb{1}(t) \quad (5.5.6)$$

Cette réponse est tracée sur la figure 5.11 pour $\omega_0 > 0$ et :

$$\zeta \in \{0, 1; 0, 2; 0, 3; 0, 4; 0, 5; 0, 6; 0, 7; 0, 8; 0, 9; 1; 2; 3; 5; 7; 10; 20\}$$

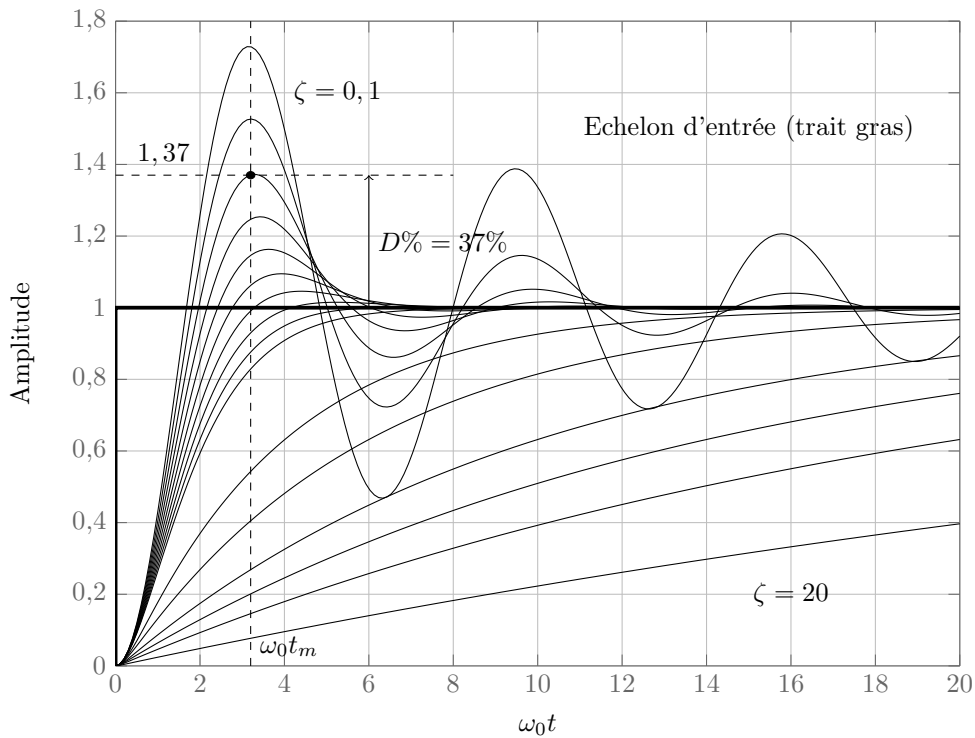


FIGURE 5.11 – Réponse indicielle d'un système du second ordre

On peut remarquer que pour des amortissements faibles ($0 < \zeta < 0,5$), le *temps du premier maximum* noté t_m vérifie approximativement :

$$\omega_0 t_m \approx 3$$

la relation exacte étant en fait :

$$t_m = \frac{\pi}{\omega_0 \sqrt{1 - \zeta^2}} \quad (5.5.7)$$

Pour $0 < \zeta < 1$, une autre caractéristique de la réponse indicielle est le dépassement, qui est l'écart (amplitude maximale de la sortie – amplitude de la sortie en régime permanent) divisé par l'amplitude de la sortie en régime permanent. On montre qu'il est égal à :

$$D = \exp \left(-\frac{\pi \zeta}{\sqrt{1 - \zeta^2}} \right) \quad (5.5.8)$$

Lorsque $0 < \zeta < 1$, les deux pôles P_1 et P_2 sont complexes conjugués et valent :

$$P_1 = \omega_0 \left(-\zeta + j\sqrt{1 - \zeta^2} \right) \text{ et } P_2 = \omega_0 \left(-\zeta - j\sqrt{1 - \zeta^2} \right)$$

La position de ces pôles peut être représentée dans le plan complexe (figure 5.12). On peut noter que l'angle θ qui apparaît sur la figure vérifie $\sin \theta = \zeta$.

Ainsi, pour des amortissements inférieurs à 1, le module des pôles (ω_0) détermine approximativement le temps du premier maximum t_m de la réponse indicielle (donné par la relation (5.5.7)), et la partie réelle détermine le dépassement D (donné par (5.5.8)).

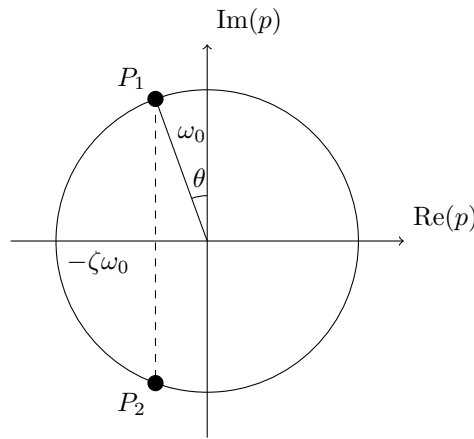


FIGURE 5.12 – Position des pôles d'un système du 2^{ème} ordre, pour $\zeta < 1$

5.5.1.3 2^{ème} ordre avec pôles et zéros

Sans considérer un cas général, observons l'incidence de la position des pôles et des zéros sur les comportements fréquentiel et temporel du système de fonction de transfert :

$$H(p) = \frac{p^2 + \omega_0^2}{p^2 + 2\zeta\omega_0 p + \omega_0^2} = \frac{(p - Z_1)(p - Z_2)}{(p - P_1)(p - P_2)} \text{ avec } 0 < \zeta < 1 \quad (5.5.9)$$

La position des pôles et des zéros dans le plan complexe est représentée en figure 5.13.

On observe deux pôles $\omega_0 \left(-\zeta \pm j\sqrt{1 - \zeta^2} \right)$ et deux zéros $\pm j\omega_0$. La réponse en fréquence est tracée en figure 5.14 pour :

$$\zeta \in \{0, 1; 0, 2; 0, 3; 0, 4; 0, 5; 0, 6; 0, 7; 0, 8; 0, 9; 1\}$$

On obtient une famille de filtres réjecteurs d'un signal harmonique de pulsation ω_0 . La présence des deux zéros a pour effet d'annuler la réponse à la pulsation ω_0 soit $H(j\omega_0) = 0$ quel que soit l'amortissement. Ils sont tous l'approximation d'un filtre idéal de gain 1 à toutes pulsations sauf à ω_0 où le gain vaut 0.

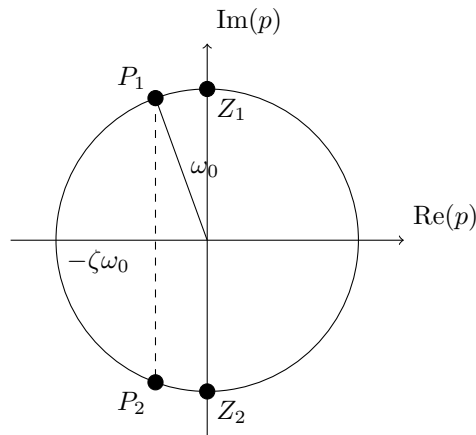


FIGURE 5.13 – Position des pôles et des zéros de la fonction de transfert 5.5.9

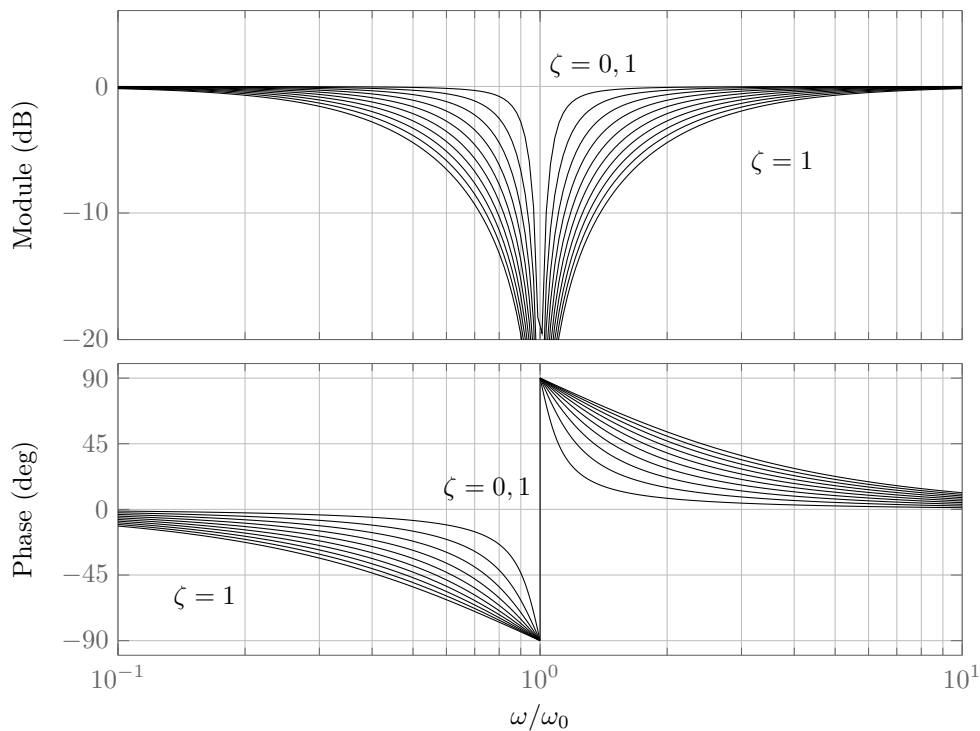


FIGURE 5.14 – Réponse d'un filtre réjecteur

On peut montrer que l'intervalle de largeur $\Delta\omega$ dans lequel :

$$|H(j\omega)| < \frac{\sqrt{2}}{2} \text{ (soit } -3 \text{ dB)}$$

vaut $\Delta\omega \approx 2\omega_0\zeta$; l'approximation est d'autant meilleure que l'amortissement est faible.

La réponse indicielle d'un tel système n'est pas la plus intéressante. En revanche, la réponse du système à une entrée $u(t) = \sin(\omega_0 t) \mathbb{1}(t)$ donnera une information sur le temps que met le système à tendre vers le régime permanent qui est ici un signal nul (puisque $H(j\omega_0) = 0$).

On a :

$$U(p) = \frac{\omega_0}{p^2 + \omega_0^2} \text{ d'où } Y(p) = \frac{\omega_0}{p^2 + 2\zeta\omega_0 p + \omega_0^2}$$

La sortie vaut alors :

$$y(t) = \frac{1}{\sqrt{1-\zeta^2}} e^{-\zeta\omega_0 t} \sin\left(\sqrt{1-\zeta^2}\omega_0 t\right) \mathbb{1}(t) \quad (5.5.10)$$

La sortie tend d'autant plus lentement vers le régime permanent que l'amortissement est faible. On obtient une constante de temps $\tau = \frac{1}{\omega_0 \zeta}$ associée au terme en exponentielle, qui traduit la vitesse de convergence vers cet état permanent que l'on souhaite atteindre le plus rapidement possible.

Globalement, on observe un compromis entre :

- la précision sur la réalisation du filtre (largeur de bande atténuée $\Delta\omega = 2\omega_0\zeta$)
- le temps de réponse du système (constante de temps $\tau = \frac{1}{\omega_0\zeta}$)

Le produit $\tau\Delta\omega$ est constant et vaut 2, ce qui permet d'énoncer pour cet exemple une sorte de principe d'incertitude : "on ne peut avoir simultanément une bonne précision sur la réponse fréquentielle et temporelle".

5.6 Conclusion de la partie I

Ici s'achève la partie I sur la représentation et l'analyse de modèles de systèmes à états continus. En synthèse, cette partie a permis de présenter, pour des systèmes à temps continu ou discret respectivement :

- la représentation d'un système sous forme d'état, constituée d'une équation d'évolution donnant la variation des variables d'état au cours du temps et d'une équation de sortie. Même si elle est loin de couvrir tous les modèles possibles (par exemple les équations à dérivées partielles ne s'y retrouvent pas), cette représentation offre un cadre de modélisation assez général puisqu'elle peut englober des modèles non-linéaires et variants dans le temps.
- l'étude de propriétés des points d'équilibre de systèmes non linéaires sous forme de représentation d'état : au-delà de la définition formelle des différents types de stabilité, on a vu que leur étude passe généralement par une linéarisation du système non linéaire autour d'un point d'équilibre afin de se ramener au cas d'un système linéaire.
- le cas particulier des systèmes linéaires et invariants dans le temps, traité de manière approfondie car de nombreux résultats peuvent être obtenus analytiquement dans ce cas, alors que la résolution ne peut généralement être traitée que de manière numérique lorsque cette hypothèse n'est pas vérifiée. On a vu que les SLI peuvent se représenter sous forme d'une fonction de transfert dont les propriétés permettent ainsi des conclusions sur la stabilité du point d'équilibre.
- l'analyse fréquentielle (via le diagramme de Bode) ou temporelle (via la réponse indicielle) d'un système linéaire invariant dans le temps.

Ces propriétés, qui caractérisent le comportement de tels systèmes, seront exploitées en partie III lorsqu'il s'agira de traiter le problème d'identification. En effet, cela nous permettra, face à un système réel que l'on cherchera à représenter par un modèle, d'utiliser les propriétés observées sur le système réel pour en déduire une forme adéquate de modèle à proposer.

Nous allons à présent passer à la description de la représentation d'une autre classe de systèmes, les systèmes à états discrets.

Deuxième partie

Modélisation des systèmes à état discret

Chapitre 6

Introduction sur la modélisation des systèmes à état discret

6.1 Problématique

Dans de nombreux systèmes complexes, l'enchaînement des tâches effectuées par le système provient d'*événements* tels que l'arrivée d'un signal, la fin d'une tâche précédente, la libération d'une ressource, le changement d'un mode opératoire, etc. La survenue d'un tel événement entraîne ainsi des phénomènes tels que la concurrence pour une même ressource, la synchronisation de tâches ou encore la saturation de certaines variables. De tels systèmes sont particulièrement présents dans les réseaux informatiques, les systèmes de production et de logistique, les réseaux de transport, les réseaux de communication, les réseaux de régulation biologiques et nous en verrons de nombreux exemples dans ce cours.

De tels systèmes ne peuvent généralement pas se modéliser de façon simple à l'aide des outils abordés dans les chapitres précédents traitant de la modélisation des systèmes à état continu. Il est en particulier difficile de représenter le système par un jeu d'équations différentielles de la forme $\dot{\mathbf{x}}(t) = f(\mathbf{x}(t), \mathbf{u}(t))$ où $\mathbf{x}(t)$ représente l'état du système et $\mathbf{u}(t)$ les entrées, car l'état $\mathbf{x}(t)$ subit des variations instantanées en réponse aux événements.

Nous allons ainsi voir dans cette partie certaines méthodes permettant d'appréhender la modélisation des systèmes à événements discrets

6.2 Motivation et structure de la partie II

Afin d'introduire de manière informelle les concepts et notions traités dans ce chapitre, nous considérons une marche aléatoire dans le plan sous la forme d'un jeu à 4 joueurs. Un mobile, supposé ponctuel et initialement placé à l'origine, peut ainsi se déplacer dans un espace à deux dimensions $(x_1(t), x_2(t))$ sous l'action de 4 joueurs N, E, S, O ayant pour objectif d'emmener le mobile en un point particulier de l'espace, différent pour chaque joueur et tenu secret. Chaque joueur dispose d'un bouton :

- lorsque le joueur N presse son bouton, le mobile se déplace d'une unité dans le sens des $x_2(t)$ croissants ;
- lorsque le joueur E presse son bouton, le mobile se déplace d'une unité dans le sens des $x_1(t)$ croissants ;
- lorsque le joueur S presse son bouton, le mobile se déplace d'une unité dans le sens des $x_2(t)$ décroissants ;
- lorsque le joueur O presse son bouton, le mobile se déplace d'une unité dans le sens des $x_1(t)$ décroissants.

L'état du jeu est donné par la position du mobile dans le plan $(x_1(t), x_2(t))$. L'évolution du jeu est entièrement définie par la *séquence d'événements*. La figure 6.1 présente ainsi l'évolution du jeu en réponse à la séquence $N \rightarrow E \rightarrow S \rightarrow S \rightarrow O \rightarrow S \rightarrow O$.

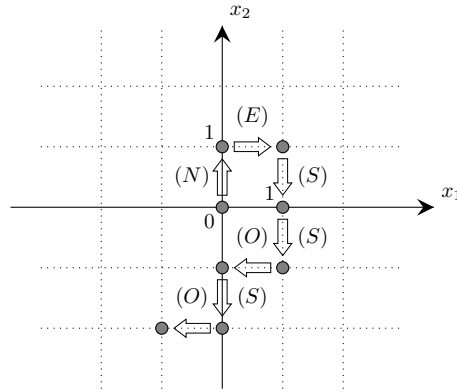
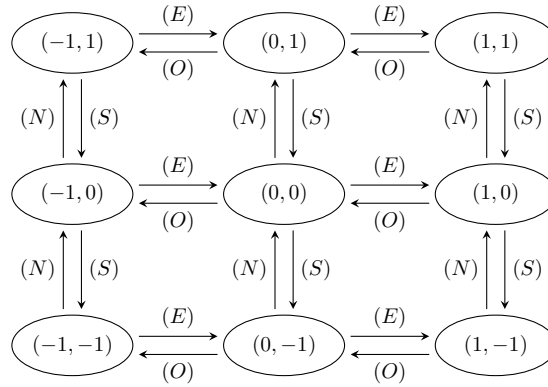


FIGURE 6.1 – Exemple de marche aléatoire

Nous voyons dans cet exemple que le temps n'apparaît pas de façon explicite. La date à laquelle se produit l'évènement n'a pas d'influence sur l'évolution du mobile. On parlera ici d'un système à évènement discret *non temporisé*. La taille de l'espace d'état est ici a priori infinie (c'est-à-dire que $x_1(t)$ et $x_2(t)$ ne sont a priori pas bornés). Si l'on ajoute au jeu la règle que le mobile doit rester dans le domaine du plan défini par $-1 \leq x_1(t) \leq 1$ et $-1 \leq x_2(t) \leq 1$ (les joueurs n'ont pas le droit de provoquer un mouvement qui ferait sortir le mobile de ce domaine), on peut représenter l'évolution du système sous la forme d'un *automate* représenté sur la figure 6.2, dans laquelle, pour simplifier, on fait l'hypothèse que deux joueurs ne peuvent appuyer sur leur bouton exactement au même instant. L'automate présenté dans cette figure contient tous les mouvements possibles dans les intervalles spécifiés, et non seulement la séquence indiquée figure 6.1. Nous présenterons formellement les systèmes à évènements discrets non temporisés, et en particulier ce type de modélisation par automate fini dans le chapitre 7.

FIGURE 6.2 – Système à évènement discret *non temporisé*. Exemple d'automate pour modéliser la marche aléatoire

Supposons désormais que l'on rajoute une nouvelle règle au jeu, qui doit se jouer en un temps limité, avec un délai maximum de réflexion autorisé pour chaque joueur durant la partie. Dans ce cas, les dates auxquelles se produisent les évènements deviennent importantes et vont influencer sur l'évolution future du jeu. On peut par exemple représenter l'évolution du jeu sous la forme d'un chronogramme comme sur la figure 6.3. Dans cette figure, à partir de l'état initial $(0, 0)$, à l'instant t_1 , le mobile fait un saut d'un pas vers le Nord, arrivant ainsi en $(0, 1)$ jusqu'au moment t_2 , où il fait un saut vers l'Est jusqu'à $(1, 1)$ et ainsi de suite. Nous aborderons les systèmes à évènements discrets dits *temporisés* dans le chapitre 8.

Les joueurs confirmés décident alors de complexifier leur jeu. La pression sur leur bouton d'action n'engendre plus alors un saut dans l'espace, mais une translation rectiligne uniforme (par exemple) dans leur direction d'action. Ainsi, si le joueur N appuie sur son bouton à l'instant t_0 ,

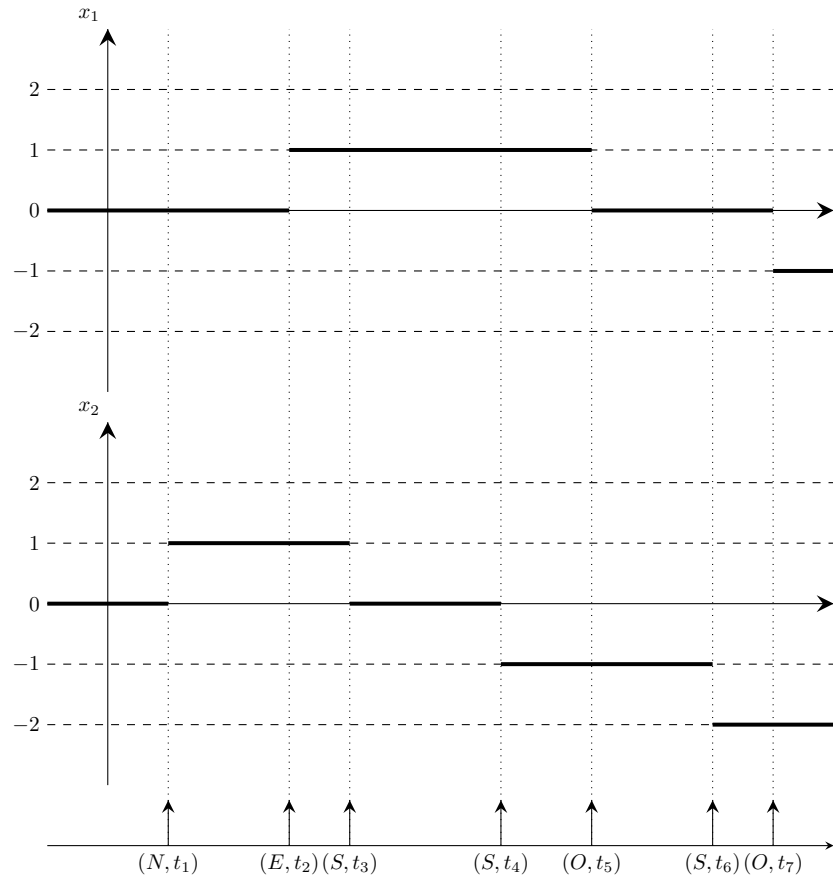


FIGURE 6.3 – Systèmes à événements discrets *temporisés*. Représentation par chronogramme de la marche aléatoire

l'évolution du mobile devient ($a > 0$ et fixé à l'avance) :

$$\begin{aligned} \dot{x}_1(t) &= 0 \rightarrow x_1(t) = x_1(t_0) \\ \dot{x}_2(t) &= +a \rightarrow x_2(t) = x_2(t_0) + a(t - t_0) \end{aligned} \quad (6.2.1)$$

et ce jusqu'à ce qu'un autre joueur décide d'appuyer sur son bouton. Un tel système est dit *système hybride*. L'arrivée d'un événement entraîne une modification brusque de la dynamique continue du système. La modélisation de tels systèmes sera abordée à la section 8.5 du chapitre 8.

Enfin, le jeu présenté ici est un jeu autonome, dans le sens où aucun signal externe ne vient agir sur le mouvement du mobile. On peut tout à fait imaginer qu'un événement extérieur (sous la forme de l'apparition d'un signal exogène par exemple) vienne modifier l'évolution du système : à une date aléatoire un mauvais génie viendra décaler la position du mobile d'une unité, modifier la valeur du paramètre a , etc. La simulation d'événements discrets nécessite l'introduction de distributions d'événements discrets telles que les lois de Rayleigh ou de Poisson que nous verrons au chapitre 9.

Remarque 6.2.1. *Attention, il ne faut absolument pas confondre la notion de systèmes à événements discrets avec, par exemple, les systèmes échantillonnés vus dans la partie de modélisation des systèmes à état continu, même si dans cet exemple introductif très simple il serait possible de modéliser le système de la marche aléatoire par une équation d'état discrète de la forme $\mathbf{x}_{k+1} = f(\mathbf{x}_k, \mathbf{u}_k)$ en considérant que l'indice k ne fait pas référence au temps, mais à la $k^{\text{ième}}$ évolution du système. Pour les systèmes à événements discrets, c'est bien la survenue d'événements qui provoque les changements d'état.*

6.3 Définition des systèmes à évènements discrets

A partir de l'exemple précédent, nous pouvons définir un système à évènements discrets comme suit.

Définition 6.3.1

Un *Système à Evénements Discrets* (SED) est un système défini par un *espace d'états discrets* et des *évolutions*, i.e. des trajectoires, fondées sur une succession d'états et de transitions.

Plusieurs points sont ici à souligner :

- La notion d'espace d'état est la même que celle qui a été utilisée dans la modélisation des systèmes à états continus par l'intermédiaire de la représentation d'état : la connaissance de l'état du système à un instant t_0 est suffisante pour en prédire le comportement futur dès lors que l'on connaît les entrées qui sont appliquées au système.
- Cet espace d'état est discret, ce qui veut dire que le nombre des états possibles est dénombrable ; il peut être fini ou infini.
- Les évènements sont dits discrets s'ils apparaissent à des instants discrets du temps, c'est-à-dire que leur apparition est instantanée et déclenche une transition instantanée de l'état.
- Des symboles, définis par les éléments d'un alphabet et appelés évènements, sont associés aux transitions.

Si l'on s'intéresse seulement à l'ordre de récurrence tout en ignorant le temps, des SED non-temporisés peuvent être élaborés à base d'automates à états finis, réseaux de Petri non-temporisés, Grafcet, etc.

Des informations temporelles cruciales peuvent être intégrées au système afin de spécifier les trajectoires d'évènements et leurs instants. Ces systèmes peuvent être représentés par :

- Modèles temporisés non-stochastiques où le temps est connu a priori : automates temporisés, réseaux de Petri temporisés, algèbre min-max etc.
- Modèles temporisés stochastiques fondés sur des hypothèses statistiques : chaînes de Markov, réseaux de files d'attente, processus semi-markoviens généralisés, réseaux de Petri stochastiques.

Chapitre 7

Modélisation des systèmes à évènements discrets non temporisés

Ce chapitre permet de formaliser ce qu'est un modèle non temporisé de système à évènements discrets et d'en donner les principales caractéristiques. Les éléments de base de la modélisation de ce type de systèmes par des automates ou par des réseaux de Petri non-temporisés seront abordés par la suite.

Définition 7.0.1

Un système à évènements discrets est dit *non temporisé* lorsque seule la séquence des évènements est considérée, les dates d'occurrences n'ayant pas d'effet.

7.1 Modélisation par automate

7.1.1 Définition et exemple d'automate

Nous formalisons ici la notion d'automate introduite dans la section 6.2 à l'aide de l'exemple de la marche aléatoire.

Définition 7.1.1

Un *automate à états finis*, également appelé *automate fini* $\mathcal{A}$ peut être représenté par un quadruplet $\mathcal{A} = (Q, \Sigma, \phi, Q_0)$, avec :

- Q un ensemble fini de sommets représentant des états discrets ;
- Σ un ensemble de symboles (évènements) appelés alphabet ;
- ϕ la relation de transition définie de $Q \times \Sigma$ dans Q . Cette relation ϕ peut être également définie comme un ensemble de transitions entre les états, définie alors par un triplet (*sommet source, évènement, sommet but*) ;
- $Q_0 \subset Q$ un ensemble d'états initiaux.

Ainsi, si le système se trouve dans un état q_1 et qu'un évènement représenté par le symbole σ se produit, il évoluera vers l'état q_2 tel que $q_2 = \phi(q_1, \sigma)$. La fonction ϕ peut n'être que *partielle* si certaines transitions ne sont pas définies. Il peut donc être plus rigoureux de définir la transition ϕ non pas comme une fonction, mais comme une partie de $Q \times \Sigma \times Q$, c'est-à-dire $\phi \subset Q \times \Sigma \times Q$, avec la convention suivante : si $(q_1, \sigma, q_2) \in \phi$, alors le symbole σ fait passer l'automate de l'état q_1 à l'état q_2 . Cette représentation est nécessaire dans le cas d'automates non déterministes (cf. définition 7.1.2).

Un automate se représentera de façon très naturelle par un graphe dont les noeuds correspondront aux états et les transitions correspondront aux arcs. Les états initiaux sont généralement représentés par une flèche entrante sur le graphe. On parlera d'automate *fini* si l'ensemble Q est fini, d'automate infini dans le cas contraire.

Exemple 7.1.1. Nous considérons ici le cas d'un (petit !) loueur de véhicules électriques. Le loueur possède 2 véhicules. Nous supposons pour simplifier qu'un véhicule loué nécessite forcément une recharge avant de pouvoir être loué à nouveau. L'état de l'ensemble de véhicules est donc caractérisé par le nombre de véhicules disponibles chargés n_c et le nombre de véhicules déchargés n_d . Nous noterons l'état défini par le vecteur $X = (n_c, n_d)$. Les événements pouvant faire évoluer le système sont la location d'un véhicule, notée *loc* ; le retour d'un véhicule, noté *ret* ou la recharge d'un véhicule, notée *ch*. Nous supposons que : 1. initialement le loueur possède deux véhicules chargés prêt à la location ; 2. pendant la journée, les voitures sont louées une par une et non pas les deux voitures en même temps ; 3. toutes les voitures sont retournées le soir. La modélisation du système peut se faire à l'aide de l'automate représenté sur la figure 7.1.

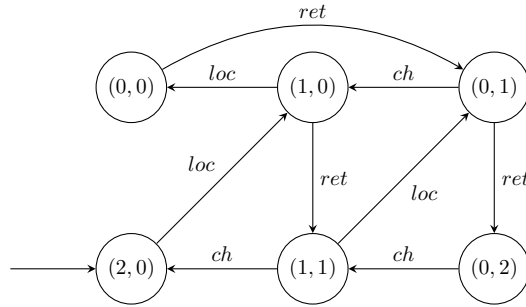


FIGURE 7.1 – Automate pour la modélisation de l'ensemble de véhicules

Pour reprendre les notations de la définition 7.1.1, nous avons donc :

- $Q = \{(0,0), (0,1), (1,0), (1,1), (2,0), (0,2)\}$;
- $\Sigma = \{loc, ret, ch\}$;
- ϕ peut être définie par simple énumération, par exemple $\phi((1,1), ch) = (2,0)$. Si l'on considère la définition comme ensemble de transitions, nous avons :

$$\begin{aligned} \phi = \{ & ((2,0), loc, (1,0)), ((1,0), loc, (0,0)), ((0,0), ret, (0,1)), \\ & ((0,1), ch, (1,0)), ((1,0), ret, (1,1)), ((1,1), ch, (2,0)), \\ & ((0,1), ret, (0,2)), ((0,2), ch, (1,1)), ((1,1), loc, (0,1)) \} \end{aligned}$$

- $Q_0 = \{(2,0)\}$.

Selon les applications, on pourra ajouter un cinquième élément $Q_f \subset Q$ à la définition de l'automate qui correspondra aux états finaux de l'automate. On parle d'*états marqués* pour ces états. Ces états seront représentés sur l'automate par un double cercle. Dans notre exemple, on peut imaginer que le loueur souhaite que ses deux véhicules soient de retour en fin de journée pour pouvoir les recharger dans la nuit. L'ensemble des états marqués est alors $Q_f = \{(2,0), (1,1), (0,2)\}$.

Nous pouvons remarquer que la notion d'évènement est très générale, puisqu'elle peut correspondre à l'arrivée d'un signal ou d'une grandeur extérieure au système (l'utilisateur a appuyé sur un bouton par exemple) ou à une condition satisfaite par un signal interne au système ("le niveau d'eau a atteint la hauteur désirée").

Il existe différentes variantes à la définition de l'automate telle que donnée à la définition 7.1.1 lorsque l'on souhaite ajouter explicitement des entrées et des sorties. On trouve ainsi les extensions suivantes :

- *Automate de Moore* : dans ce cas, on trouve une fonction qui associe à chaque état une valeur pour la (ou les) sortie(s) du système. La sortie ne dépend ainsi que de l'état (l'évènement de sortie est émis lorsque l'automate rentre dans l'état). Conventionnellement, la sortie est indiquée en gras au dessus de l'état ; elle peut correspondre par exemple à une information délivrée par un capteur à l'arrivée dans cet état.
- *Automate de Mealy* : la sortie dépend d'un état et d'un évènement d'entrée reçu dans cet état. Ainsi, dans un état x si l'on reçoit un évènement d'entrée e_i , il y aura une transition

vers un état y et émission d'un événement de sortie e_o lors de la transition. Les transitions sont étiquetées avec des événements de la forme 'événement d'entrée e_i / événement de sortie e_o '.

7.1.2 Propriétés des automates

7.1.2.1 Caractère déterministe / non déterministe

Définition 7.1.2

- Un automate sera dit *déterministe* si les deux conditions suivantes sont remplies :
- il n'y a qu'un seul état initial, c'est-à-dire $Q_0 = \{q_0\}$;
 - à partir d'un état, un même symbole ne peut conduire que dans un seul état, c'est à dire :

$$\forall q \in Q, \forall \sigma \in \Sigma, ((q, \sigma, q_1) \in \phi \ \& \ (q, \sigma, q_2) \in \phi) \Rightarrow q_1 = q_2 \quad (7.1.1)$$

Dans le cas contraire, on parlera d'automates non déterministes ou stochastiques. Ainsi, il est possible que pour certains automates un même symbole puisse potentiellement entraîner plusieurs transitions (bien sûr une seule transition sera effectuée), ou bien que certaines transitions soient spontanées (la transition a lieu sans qu'aucun événement ne soit apparu). Du point de vue de la modélisation, ce type d'automate peut correspondre à des incertitudes de modélisation (on ne sait pas par exemple avec certitude quelle va être la conséquence exacte d'un événement, un événement peut ne pas être mesuré, etc.). Du point de vue de la théorie des automates, les applications des automates non déterministes sont nombreuses, mais leur étude complète dépasse le cadre de ce cours.

7.1.2.2 Accessibilité, co-accessibilité, complément

Reprenons tout d'abord la notion de fonction de transition, donnée à la définition 7.1.1. Nous pouvons étendre cette définition de la transition pour une suite de symboles reçus à partir d'un état q_0 . En effet, si l'automate reçoit une suite de symboles valides (au sens où toutes les transitions correspondantes sont définies) $\sigma_1\sigma_2...\sigma_n$ à partir d'un état q_0 , le système évoluera jusqu'à l'état q_n défini par $\phi(\phi(...\phi(\phi(q_0, \sigma_1), \sigma_2), ...), \sigma_n))$. On notera de façon abrégée la *lecture* de cette suite de symboles $\sigma_1\sigma_2...\sigma_n$ par $\phi(q_0, \sigma_1\sigma_2...\sigma_n) = q_n$.

Définition 7.1.3

Un état q sera dit *accessible* depuis un état initial q_0 s'il existe une suite de symboles $\sigma_1\sigma_2...\sigma_n$ valide telle que $\phi(q_0, \sigma_1\sigma_2...\sigma_n) = q$.

Un automate dont tous les états sont accessibles sera dit accessible. Clairement, on peut retirer de l'automate de départ la partie non accessible et obtenir un automate dont le comportement sera identique. Du point de vue de la modélisation d'un système à événements discrets, trouver des états non accessibles est souvent le signe d'une erreur dans le processus qui a conduit à ce modèle.

Définition 7.1.4

Un état q sera dit *co-accessible* s'il existe une suite de symboles $\sigma_1\sigma_2...\sigma_n$ valide permettant d'aller dans l'un des états marqués, c'est à dire que $\phi(q, \sigma_1\sigma_2...\sigma_n) = q_m$ avec $q_m \in Q_f$.

Si l'automate obtenu en éliminant la partie non co-accessible est identique à l'automate de départ, alors l'automate est dit *co-accessible*.

Remarque 7.1.1. *Il ne suffit pas que tous les états soient co-accessibles pour que l'automate le soit. En effet, supposons qu'un automate ait deux états marqués q_{f1} et q_{f2} . Tous les états peuvent être co-accessibles, mais avoir seulement des chemins conduisant à q_{f1} , tandis que q_{f2} n'est pas accessible.*

Si un automate n'est pas co-accessible, cela peut souvent être le signe d'un système mal conçu. En effet, les états marqués sont souvent ceux vers lesquels on souhaite que le système évolue. Dire qu'un état n'est pas co-accessible signifie qu'il est possible que le système soit arrivé dans un état à partir duquel il n'est pas possible de le faire évoluer vers un état marqué.

Définition 7.1.5

Un automate sera dit *élagué* s'il est à la fois accessible et co-accessible.

L'élaguage d'un automate consistera ainsi à prendre successivement la partie accessible et co-accessible d'un automate, ces deux opérations pouvant être faites dans un ordre quelconque.

Exemple 7.1.2. Afin d'illustrer ces notions, considérons l'automate initial représenté sur la figure 7.2.

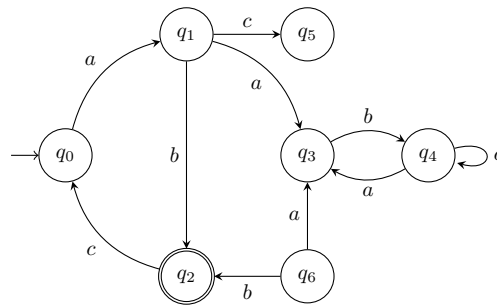


FIGURE 7.2 – Automate initial

L'état 6 n'est pas accessible à partir de l'état initial q_0 . Pour obtenir un automate accessible (représenté sur la figure 7.3), on le supprime ainsi que les transitions sortant de cet état.

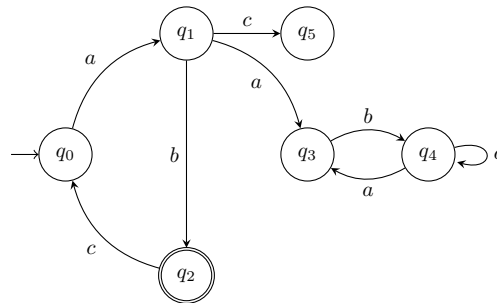


FIGURE 7.3 – Automate accessible

Si par ailleurs on repart de l'automate de la figure 7.2, la recherche de la partie co-accessible conduit à l'automate représenté sur la figure 7.4 où on trouve tous les états permettant d'aller dans l'état marqué q_2 . Dans cet exemple, les états q_3 , q_4 et q_5 ne permettent pas des transitions vers l'état marqué.

Enfin, si on applique successivement les deux opérations, on obtient l'automate élagué de la figure 7.5.

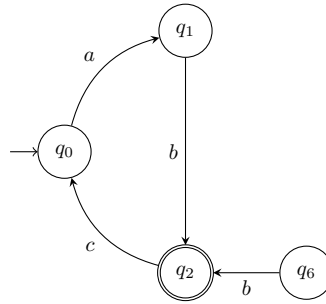


FIGURE 7.4 – Automate co-accessible

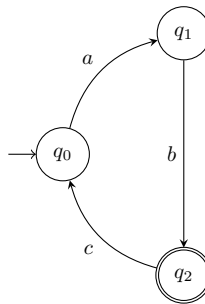


FIGURE 7.5 – Automate élagué

7.1.3 Composition des automates

Nous avons vu dans les chapitres précédents qu'un système peut être vu comme une architecture de sous-systèmes en interactions. Il est donc nécessaire de pouvoir hiérarchiser un modèle (dans le sens où un automate peut être décomposé en sous-automates) ou bien de combiner plusieurs automates pour les faire évoluer ensemble.

7.1.3.1 Automates hiérarchiques

Définition 7.1.6

Un automate est dit hiérarchique lorsque certains de ses états sont eux-mêmes des automates.

Nous ne rentrerons pas ici dans une présentation détaillée des machines d'état ou statecharts, notions par ailleurs reprises dans UML et qui seront abordées dans le cours de Modélisation Systèmes de deuxième année. De manière très intuitive, nous pouvons étendre la notion des automates, en considérant la représentation d'un état donnée sur la figure 7.6.

Il apparaît assez évident que ce qui est repris par les termes "actions" et "activités" à l'intérieur d'un état peuvent renvoyer à une description du comportement de cet état sous forme d'un automate.

Exemple 7.1.3. La figure 7.7 donne un exemple d'automate hiérarchique pour le fonctionnement d'une machine.

Au niveau supérieur, la machine peut être dans l'un des trois états suivants $Q = \{a, n, d\}$:

- a : la machine est à l'arrêt ;
- n : la machine est dans le mode de fonctionnement normal ;
- d : la machine est en défaillance.

Entre ces états, on trouve les transitions associées aux événements suivants $\Sigma = \{\mu, \alpha, \pi, \rho\}$:

- μ : démarrage de la machine ;
- α : demande d'arrêt de la machine ;

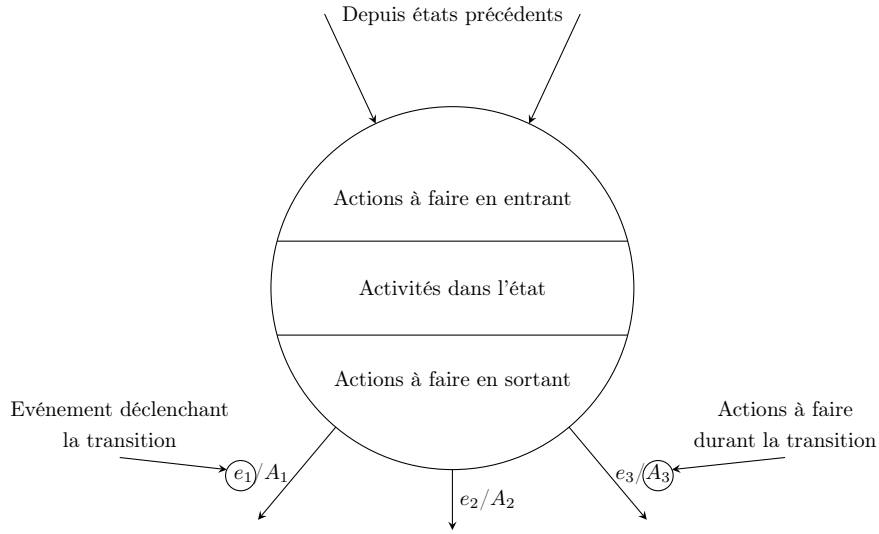


FIGURE 7.6 – Description générale d'un état

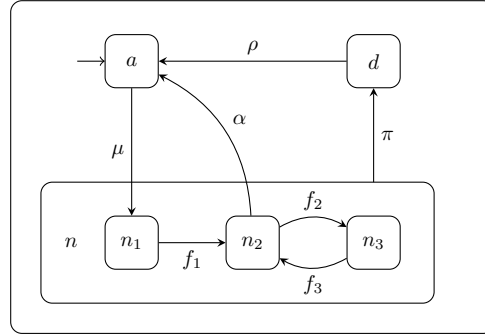


FIGURE 7.7 – Exemple d'automate hiérarchisé

- π : apparition d'un défaut ;
- ρ : réparation de la machine.

Un exemple de transition associée à un événement dans le niveau supérieur est le suivant :

$$\phi_{sup}(a, \mu) = n.$$

L'état initial est $q_0 = \{a\}$. Au niveau inférieur, les états n_1, n_2, n_3 et les transitions associées aux événements f_1, f_2 et f_3 permettent de décrire le fonctionnement de la machine dans le mode normal (par exemple $\phi_{inf}(n_1, f_1) = n_2$).

7.1.3.2 Produit synchrone d'automates

Nous considérons ici deux automates (incluant des états marqués) notés :

$$\mathcal{A}_1 = (Q_1, \Sigma_1, \phi_1, Q_{01}, Q_{f1}) \text{ et } \mathcal{A}_2 = (Q_2, \Sigma_2, \phi_2, Q_{02}, Q_{f2}).$$

Définition 7.1.7

Le produit synchrone des automates $\mathcal{A}_1$ et $\mathcal{A}_2$, noté $\mathcal{A}_1 \times \mathcal{A}_2$ est défini par : $\mathcal{A}_1 \times \mathcal{A}_2 = Ac(Q_1 \times Q_2, \Sigma_1 \cap \Sigma_2, \phi_{prod}, Q_{01} \times Q_{02}, Q_{f1} \times Q_{f2})$, avec :

$$\phi_{prod}((q_1, q_2), e) = \begin{cases} (\phi_1(q_1, e), \phi_2(q_2, e)) & \text{si les deux transitions sont définies} \\ \text{non définie} & \text{sinon} \end{cases} \quad (7.1.2)$$

La notation $Ac(.)$ signifie que l'on ne garde que la partie accessible de l'automate résultant. Plusieurs faits sont en outre à remarquer :

- l'opération $\times$ ne prend en compte que les événements qui appartiennent à l'intersection des ensembles considérés $\Sigma_1 \cap \Sigma_2$.
- le produit est nécessairement synchrone dans le sens où un événement arrive exactement au même instant au niveau des deux automates entraînant une transition simultanée.
- tous les événements privés (c'est-à-dire n'appartenant qu'à l'un des deux automates) sont absents du produit.

Exemple 7.1.4. On considère les deux automates donnés sur la figure 7.8.

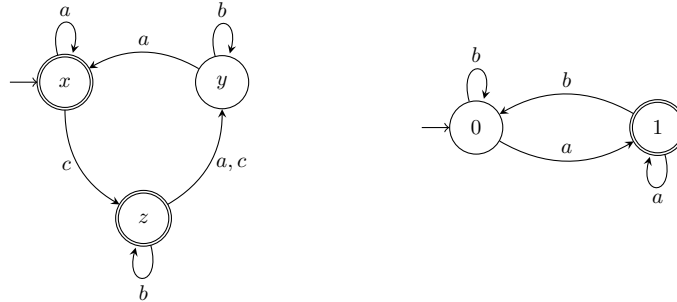


FIGURE 7.8 – Automates utilisés pour illustrer le produit synchrone

L'état initial du produit de ces deux automates (également appelé automate produit) est $(x, 0)$ correspondant au seul couple possible d'états initiaux pour les deux automates. A partir de cet état, on voit que la seule transition simultanément possible pour les deux automates est a , qui conduit à l'état $(x, 1)$. Une fois dans cet état, de nouveau une seule transition est possible simultanément sur les deux automates : l'événement a fait ainsi rester l'automate dans l'état $(x, 1)$. Il n'y a plus de nouvel état à analyser, le produit synchrone obtenu est représenté sur la figure 7.9. Les états x pour le premier automate et 1 pour le deuxième étant des états marqués, l'état $(x, 1)$ est marqué dans l'automate produit.

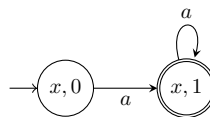


FIGURE 7.9 – Produit synchrone des deux automates

7.1.3.3 Parallélisation d'automates

On reprend ici les deux automates $\mathcal{A}_1 = (Q_1, \Sigma_1, \phi_1, Q_{01}, Q_{f1})$ et $\mathcal{A}_2 = (Q_2, \Sigma_2, \phi_2, Q_{02}, Q_{f2})$. La parallélisation consiste à faire évoluer les deux automates de façon partiellement synchrone en considérant tous les événements qui peuvent se produire pour l'un ou l'autre des automates.

Définition 7.1.8

La parallélisation des automates $\mathcal{A}_1$ et $\mathcal{A}_2$, notée $\mathcal{A}_1 || \mathcal{A}_2$ est définie par :

$$\mathcal{A}_1 || \mathcal{A}_2 = Ac(Q_1 \times Q_2, \Sigma_1 \cup \Sigma_2, \phi_{par}, Q_{01} \times Q_{02}, Q_{f1} \times Q_{f2}),$$

avec :

$$\phi_{par}((q_1, q_2), e) = \begin{cases} (\phi_1(q_1, e), \phi_2(q_2, e)) & \text{si les deux transitions sont définies} \\ (\phi_1(q_1, e), q_2) & \text{si la transition n'est définie que pour } \mathcal{A}_1 \\ (q_1, \phi_2(q_2, e)) & \text{si la transition n'est définie que pour } \mathcal{A}_2 \\ \text{non définie} & \text{sinon} \end{cases} \quad (7.1.3)$$

On notera que la parallélisation des deux automates est partiellement synchronisée dans le sens où un événement présent simultanément sur les deux automates doit se déclencher exactement au même instant pour que ce type de transition ait lieu dans l'automate résultant.

Exemple 7.1.5. La figure 7.10 présente le résultat de la parallélisation des deux automates représentés sur la figure 7.8 (avec $Q_1 = \{x, y, z\}$ et $Q_2 = \{0, 1\}$) :

$$Q = \{(x, 0), (x, 1), (y, 0), (y, 1), (z, 0), (z, 1)\}.$$

Les états $(x, 1)$ et $(z, 1)$ correspondent aux états marqués. En parallélisant ces deux automates, on peut voir qu'à partir de l'état $(x, 0)$ si l'on applique la transition associée à l'événement a on arrive dans l'état $(x, 1)$, en revanche la transition associée à l'événement b permet de rester dans l'état $(x, 0)$. Retrouver l'ensemble de l'automate de la figure 7.10 est laissé comme exercice au lecteur.

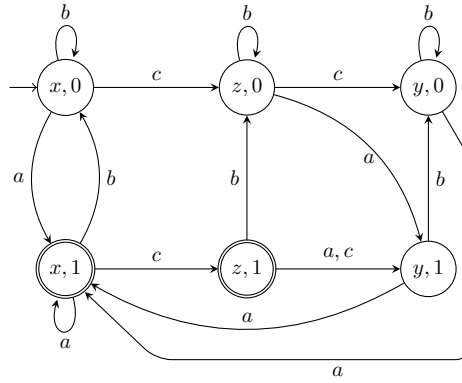


FIGURE 7.10 – Parallélisation des deux automates

7.1.4 Analyse d'un système à événements discrets non temporisé représenté par un automate

7.1.4.1 Simulation et bisimulation d'automates

Considérons deux automates $\mathcal{A}_1$ et $\mathcal{A}_2$ définis sur un même ensemble d'événements Σ . De manière informelle, on dira que $\mathcal{A}_1$ *simule* $\mathcal{A}_2$ si toute suite d'événements exécutée par $\mathcal{A}_2$ peut également être exécutée par $\mathcal{A}_1$. Si réciproquement $\mathcal{A}_2$ simule $\mathcal{A}_1$ (c'est-à-dire que toute suite d'événements exécutée par $\mathcal{A}_1$ peut l'être également par $\mathcal{A}_2$), les deux automates sont dits *bisimilaires*.

Plus formellement, on a la définition suivante de la bisimilarité, en se limitant au cas d'automates déterministes :

Définition 7.1.9

Soient deux automates $\mathcal{A}_1 = (Q_1, \Sigma, \phi_1, Q_{01}, Q_{f1})$ et $\mathcal{A}_2 = (Q_2, \Sigma, \phi_2, Q_{02}, Q_{f2})$. Ces deux automates sont *bisimilaires* s'il existe une relation $\mathcal{R}$ binaire de $Q_1 \times Q_2$, au sens où $(q_1, q_2) \in \mathcal{R}$ signifie que les deux états q_1 et q_2 sont en relation $\mathcal{R}$, ce que nous noterons $q_1 \mathcal{R} q_2$, telle que :

- les états initiaux Q_{01} et Q_{02} sont en relation (pour un automate déterministe il n'y a qu'un seul état initial) ;
- pour tout état $q_1 \in Q_1$, il existe $q_2 \in Q_2$ tel que $q_1 \mathcal{R} q_2$;
- pour tout état $q_2 \in Q_2$, il existe $q_1 \in Q_1$ tel que $q_1 \mathcal{R} q_2$;
- si $q_1 \mathcal{R} q_2$ et $\sigma \in \Sigma$, si $q_3 = \phi_1(q_1, \sigma)$ est défini, alors $q_4 = \phi_2(q_2, \sigma)$ est défini et $q_3 \mathcal{R} q_4$;
- si $q_1 \mathcal{R} q_2$ et $\sigma \in \Sigma$, si $q_4 = \phi_2(q_2, \sigma)$ est défini, alors $q_3 = \phi_1(q_1, \sigma)$ est défini et $q_3 \mathcal{R} q_4$;
- pour les états marqués, si $q_1 \mathcal{R} q_2$, alors $(q_1 \in Q_{f1} \Leftrightarrow q_2 \in Q_{f2})$.

Ces notions sont particulièrement intéressantes pour aborder des problèmes de simplification ou réduction de modèle.

7.1.4.2 Deadlock et livelock

Définition 7.1.10

Le *deadlock* ou *blocage fatal* consiste en une situation où il n'y a plus d'évolution possible alors que l'automate est arrivé dans un état non marqué.

Définition 7.1.11

Le *livelock* ou *blocage vivant* caractérise une situation où un sous-ensemble d'états non marqués devient le seul accessible.

L'évolution future de l'automate reste alors confinée dans ce seul sous-ensemble. Le blocage n'est pas fatal car l'état peut encore évoluer sous l'effet des événements, mais l'automate ne peut plus atteindre un état marqué.

Exemple 7.1.6. La figure 7.11 présente l'exemple d'un automate dans lequel :

- l'état q_4 constitue un *deadlock* ;
- les sous-ensembles formés par les états q_1 et q_3 constituent un *livelock*.

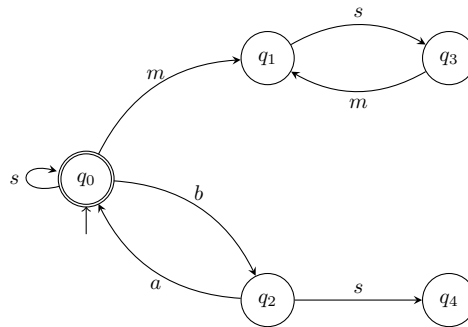


FIGURE 7.11 – Exemple d'automate présentant un deadlock et un livelock

La détermination de ces situations permet d'analyser le bon fonctionnement d'un système. Ces notions sont bien sûr liées aux notions d'atteignabilité ou d'accessibilité vues précédemment. Elles trouvent également un large champ d'application dans les problématiques liées au diagnostic et à la sûreté de fonctionnement.

7.1.5 Utilisation de la notion de langage par les automates

Bien qu'elle ne soit pas primordiale dans une optique de modélisation, la notion de *langage* est d'importance majeure dans la théorie des automates à laquelle elle est souvent associée. Nous en présentons ici quelques aspects et les liens naturels qui existent avec les automates.

Définition 7.1.12

On appelle *alphabet*, noté Σ dans la suite, un ensemble, fini ou infini, de lettres ou de symboles.

Définition 7.1.13

Un *mot* u est un ensemble fini de lettres de Σ . Sa longueur est notée $|u|$. On note Σ^* l'ensemble des mots finis qui peuvent être formés à partir des lettres de Σ .

Le mot vide est très souvent noté ε . De nombreuses opérations peuvent être définies sur les mots. L'une des plus importantes est la *concaténation* de deux mots u et v qui permet d'obtenir un nouveau mot $w = u.v$, tel que :

$$w(i) = \begin{pmatrix} u(i), & \text{si } i \leq |u| \\ v(i - |u|), & \text{si } |u| + 1 \leq i \leq |u| + |v| \end{pmatrix} \quad (7.1.4)$$

Définition 7.1.14

Un langage L est un ensemble fini ou infini de mots, c'est-à-dire que $L \subset \Sigma^*$.

Exemple 7.1.7. Si $\Sigma = \{\varepsilon, a, b, c\}$, $L_1 = \{\varepsilon, a, aab, ca, bc\}$ est un langage (fini).
 $L_2 = \{\text{combinaisons de } a \text{ et } b\} = \{\varepsilon, a, b, aa, bb, ab, ba, bba, aab, aba, \dots\}$ est un langage (infini).

De nombreuses opérations existent également sur les langages (union, intersection, par exemple, dès lors qu'il s'agit d'opérations d'union ou d'intersection sur des ensembles). En particulier, la concaténation de langages revêt une grande importance. La *concaténation de langage* L_1 et L_2 est ainsi donnée par :

$$L_1 \subset \Sigma^*, L_2 \subset \Sigma^*, L_1 L_2 = \{w = u.v, u \in L_1, v \in L_2\} \quad (7.1.5)$$

Définition 7.1.15

La *fermeture par préfixe* $\overline{L}$ d'un langage L est définie de la façon suivante :

$$L \subset \Sigma^*, \overline{L} = \{s \in \Sigma^* | \exists t \in \Sigma^*, st \in L\} \quad (7.1.6)$$

Il s'agit de l'ensemble des préfixes des mots de l'alphabet L . Généralement $L \subset \overline{L}$ (il suffit que $\varepsilon \in L$). Si $L = \overline{L}$, l'alphabet L est dit *fermé* (ou *clos*) *par préfixe*. On définit enfin la *fermeture de Kleene* de L (aussi nommée *fermeture itérative*) par :

$$L^* = \{\varepsilon\} \cup L \cup LL \cup LLL \cup \dots \quad (7.1.7)$$

On remarquera que l'on a utilisé la même notation, Σ , pour l'ensemble des symboles entrant dans la définition d'un automate et l'alphabet permettant de définir un langage. En effet, si l'on considère les symboles d'un automate comme un alphabet, les suites de transitions permises par un automate vont constituer des mots *reconnus* par l'automate. Ainsi, la question "comment créer un automate reproduisant un langage donné?" renverra-t-elle à la phase de modélisation d'un système, tandis que la question "quel est le langage parlé par cet automate?" renverra à la phase d'analyse des propriétés du modèle.

Considérons un automate $\mathcal{A} = (Q, \Sigma, \phi, Q_0, Q_f)$. Le *langage reconnu par l'automate* est défini par :

$$\mathcal{L}(\mathcal{A}) = \{s \in \Sigma^* | \exists q_0 \in Q_0 \text{ pour lequel } \phi(q_0, s) \text{ est défini}\} \quad (7.1.8)$$

On rappelle que pour un automate déterministe l'état initial est unique. Le *langage marqué* par l'automate est quant à lui donné par :

$$\mathcal{L}_m(\mathcal{A}) = \{s \in \Sigma^* | \exists q_0 \in Q_0 \text{ pour lequel } \phi(q_0, s) \in Q_f\} \quad (7.1.9)$$

De nombreuses propriétés des automates peuvent être analysées en utilisant cette notion de langage. Ces approches sortent néanmoins du cadre de ce cours traitant des aspects de *modélisation* des systèmes à événements discrets.

7.2 Modélisation par réseau de Petri non temporisés

Cette partie est très largement inspirée du polycopié (Pierre Chlique, 2000), dont elle reprend de très larges extraits.

Les réseaux de Petri constituent une alternative aux automates à états finis pour la représentation des systèmes à événements discrets. Ils sont en particulier bien adaptés à la modélisation des processus ayant des activités parallèles.

Leur définition dans les années soixante répondait à un besoin d'analyse et de validation des programmes conçus pour les machines multiprocesseurs. Ils restent actuellement largement utilisés pour la modélisation des activités des systèmes temps réel.

Nous définirons tout d'abord dans ce chapitre les réseaux de Petri simples autonomes. Ceux-ci permettent de modéliser le comportement de systèmes isolés et présentent surtout un ensemble de propriétés vérifiables.

Notons que l'on peut montrer que tout automate peut se représenter sous la forme d'un réseau de Petri, alors que la réciproque est fausse. En ce sens, les réseaux de Petri offrent un cadre de modélisation plus général que les automates vus dans la section précédente.

7.2.1 Réseaux de Petri autonomes

7.2.1.1 Structure statique des graphes

Définition 7.2.1

Un graphe R correspondant à un *réseau de Petri autonome* représente les relations entre trois ensembles d'éléments :

$$R = (P, T, Arc)$$

- des *places* P symbolisées par des cercles ;
- des *transitions* T symbolisées par des traits ;
- des *arcs* Arc orientés.

Les places et les transitions constituent les nœuds du graphe, reliés par des arcs orientés.

7.2.1.2 Contrainte

Dans un réseau de Petri, un arc doit relier deux nœuds de natures différentes. Le réseau de la figure 7.12 respecte cette contrainte. Le réseau de Petri proposé dans cette figure est caractérisé par $P = \{P_1, P_2, \dots, P_6\}$, $T = \{tr_1, tr_2, tr_3, tr_4\}$ et par un ensemble Arc formé de 12 arcs orientés.

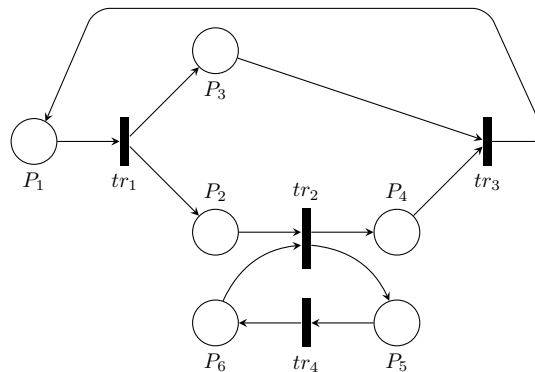


FIGURE 7.12 – Exemple de réseau de Petri, sans marquage initial

Il n'y a aucune limitation pour le nombre d'arcs entrant ou sortant d'un nœud.

7.2.1.3 Modélisation de la dynamique des systèmes

Définition 7.2.2

La notion de *marquage* permet de modéliser l'état d'un système. Il s'agit d'une application m qui associe un nombre entier à toute place du réseau :

$$m : P \rightarrow \mathbb{N}$$

Le marquage est symbolisé sur le graphe par des jetons.

Le *marquage initial* représente l'état initial du système modélisé. Un réseau de Petri est ainsi défini par la donnée de son graphe ET d'un marquage initial associé.

En notant $m(P_i)$ le marquage de la place P_i , on a pour le réseau de la figure 7.13 : $m(P_1) = 2$, $m(P_6) = 1$, $m(P_i) = 0$ pour $i \in \{2, 3, 4, 5\}$.

La règle suivante permet de modéliser l'évolution des systèmes discrets.

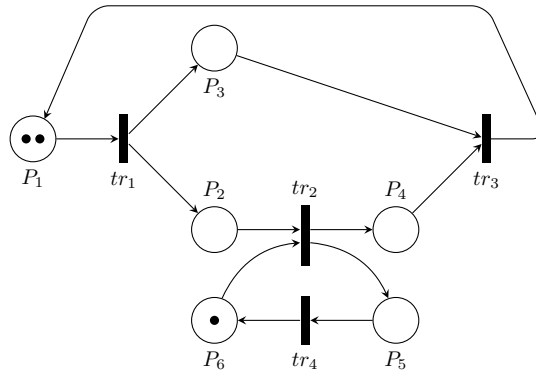


FIGURE 7.13 – Exemple de réseau de Petri, avec marquage initial

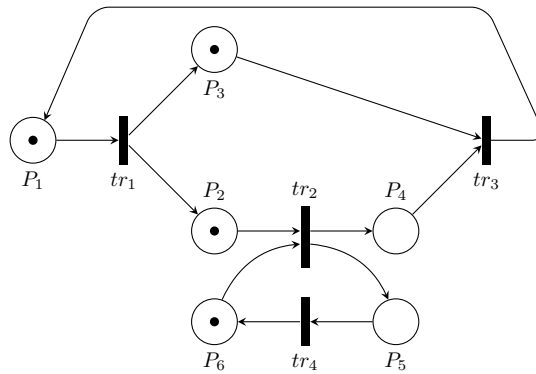
Définition 7.2.3

Une transition est dite *validée* si toutes ses places d'entrée sont marquées (c'est-à-dire contiennent au moins un jeton).

Définition 7.2.4

Une transition d'un réseau de Petri autonome est *franchie* si elle est validée en :

- retirant un jeton de chaque place d'entrée,
- ajoutant un jeton dans chaque place de sortie.

FIGURE 7.14 – Exemple de réseau de Petri, après franchissement de la transition t_1

La durée de franchissement d'une transition n'est pas nécessairement nulle.

Après avoir introduit ces aspects, une définition plus rigoureuse est donnée dans la suite.

Définition 7.2.5

Un *réseau de Petri discret autonome marqué* (i.e. le temps n'intervient pas) est défini par le quintuplet $R = (P, T, Pre, Post, m_0)$ avec :

- $P = \{P_1, P_2, \dots, P_n\}$ un ensemble fini, non vide de places ;
- $T = \{tr_1, tr_2, \dots, tr_m\}$ un ensemble fini, non vide de transitions, avec $P \cap T = \emptyset$;
- $Pre : P \times T \rightarrow \mathbb{N}$ l'application d'incidence avant, où $Pre(P_i, tr_j)$ contient la valeur entière associée à l'arc allant de P_i à tr_j ;
- $Post : P \times T \rightarrow \mathbb{N}$ l'application d'incidence arrière, où $Post(P_i, tr_j)$ contient la valeur entière associée à l'arc allant de tr_j à P_i ;
- $m_0 : P \rightarrow \mathbb{N}$ le marquage initial du réseau.

7.2.1.4 Séquence de franchissements

Plusieurs transitions franchissables sont franchies successivement suivant une *séquence de franchissement*. Il n'y a pas de franchissement simultané des transitions franchissables. Ainsi, pour interpréter un réseau de Petri autonome, on applique l'algorithme suivant :

```

répéter
    déterminer les transitions franchissables ;
    choisir une séquence de franchissement ;
    franchir la première transition de cette séquence ;
fin_répéter

```

Ainsi, dans l'exemple de la figure 7.13, la transition tr_1 est seule validée au départ. Après son franchissement, le marquage suivant est celui de la figure 7.14. Pour ce marquage, les deux transitions tr_1 et tr_2 sont validées et une séquence de franchissement doit être choisie.

Le marquage du réseau est représenté par un vecteur colonne $m(P)$ ayant pour dimension le nombre de places du réseau.

Le **graphe des marquages** représente l'ensemble des évolutions du marquage et correspond à une "remise à plat" de la description faisant apparaître les états globaux du système modélisé.

Un nœud de ce graphe correspond à une valeur du vecteur $m(P)$ et un arc de sortie à une transition franchissable pour ce marquage (figure 7.15).

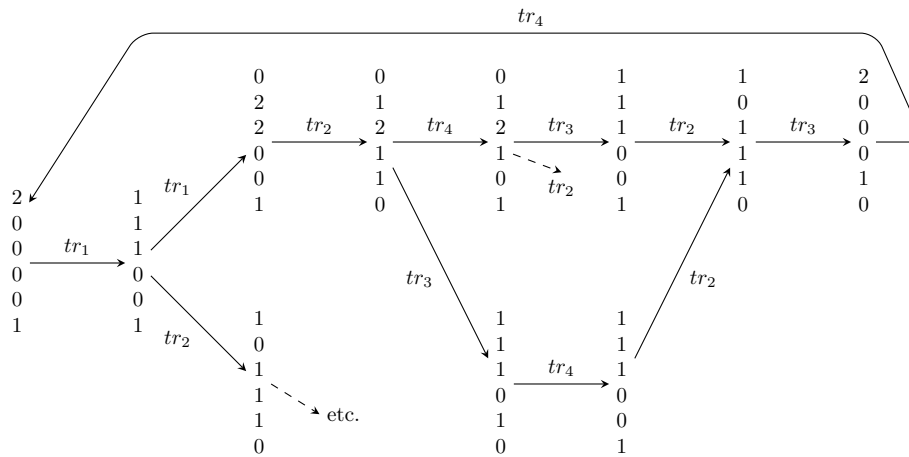


FIGURE 7.15 – Exemple de graphe de marquage

7.2.1.5 Applications

Les réseaux de Petri permettent de représenter l'évolution de systèmes parallèles. Des structures courantes sont présentées en figure 7.16.

Remarque 7.2.1. Les transitions source n'ont pas de place d'entrée et sont toujours franchissables. A contrario, les transitions puits n'ont pas de place de sortie et font donc disparaître des marques. Les représentations correspondantes sont données sur la figure 7.17.

Exemple 7.2.1. Validation d'une application de commande répartie.

Soit un ensemble de tâches réparties sur un système multi-processeurs dont la spécification des communications est fournie en figure 7.18. Les tâches communiquent par des requêtes d'envoi de messages dans des boîtes à lettres et d'attente de message sur ces boîtes à lettres.

Le problème global est d'analyser les synchronisations des tâches de cet ensemble indépendamment des traitements séquentiels qu'elles effectuent. Leur comportement est modélisé par le réseau de Petri de la figure 7.19 que l'on pourra essayer de retrouver à titre d'exercice.

En simulant l'évolution de celui-ci, on peut remarquer qu'il existe des séquences de franchissement de transition conduisant à un blocage du réseau (les retrouver à titre d'exercice)

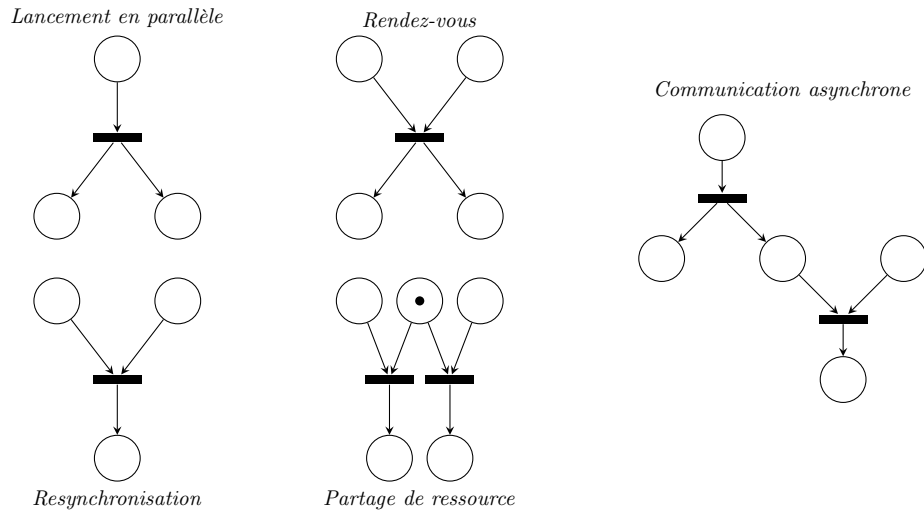


FIGURE 7.16 – Exemple de réseaux de Petri permettant la représentation de systèmes parallèles

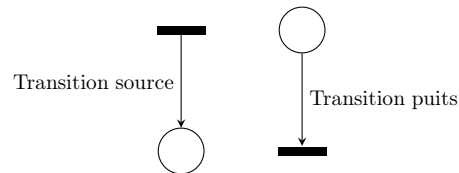


FIGURE 7.17 – Transition source et transition puits

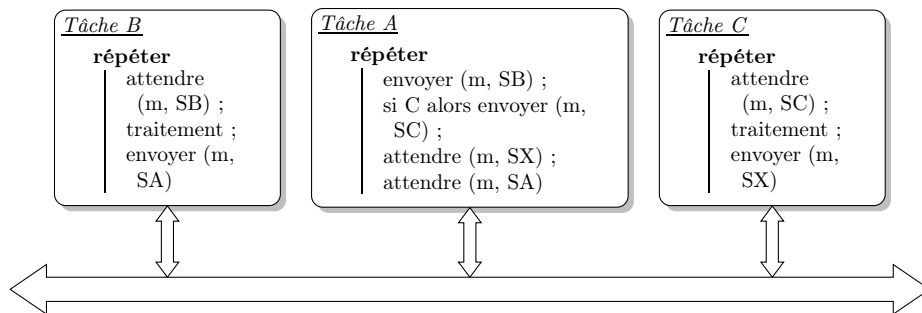


FIGURE 7.18 – Définition des tâches d'un système multi-processeurs

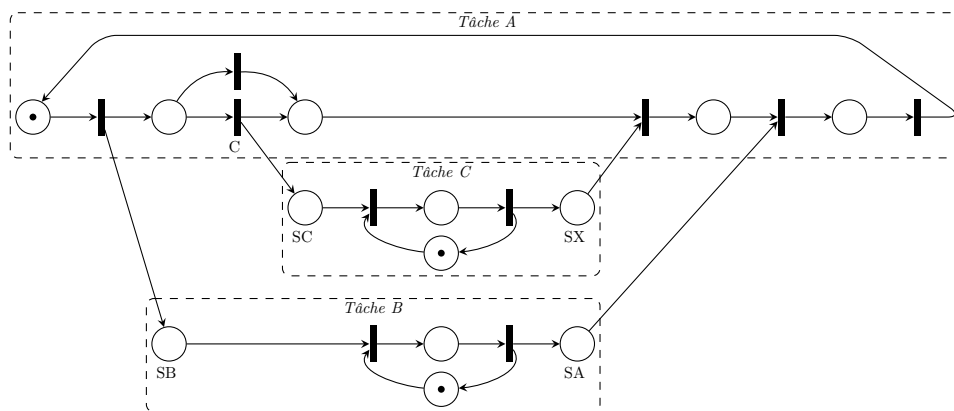


FIGURE 7.19 – Validation d'une application de commande répartie

Exemple 7.2.2. *Spécification du rendez-vous du langage Ada.*

Le mécanisme de synchronisation du langage Ada met en jeu une tâche serveur et une (ou plusieurs) tâche(s) cliente(s).

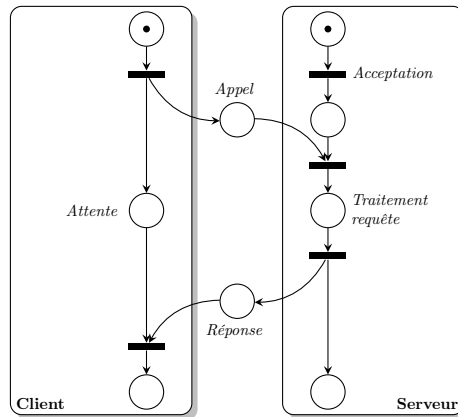


FIGURE 7.20 – Tâche client et tâche serveur

Une tâche "cliente" présente des requêtes à la tâche "serveur" qui ne les satisfait que lorsqu'elle est arrivée en un point d'acceptation. Elle doit donc attendre que la tâche "serveur" soit prête.

Inversement, lorsqu'une tâche serveur n'a reçu aucune demande, elle attend en son point d'acceptation. Une spécification de ce mécanisme est donnée en figure 7.20.

7.2.2 Propriétés des modèles représentés par réseaux de Petri

Lors de la phase de spécification d'un système, plusieurs types d'erreurs peuvent être commises telles que représentées sur la figure 7.21 :

- des erreurs de connaissance du système modélisé,
- des erreurs de modélisation dues à une mauvaise utilisation de l'outil utilisé, ici les réseaux de Petri.

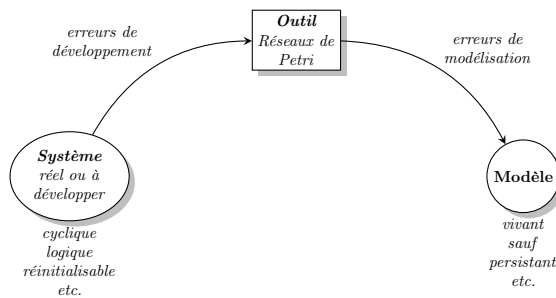


FIGURE 7.21 – Source d'erreur dans la modélisation par réseau de Petri

Le système modélisé possède en général un certain nombre de priorités identifiables. Il est intéressant de pouvoir vérifier que le modèle obtenu possède des propriétés similaires ou compatibles.

7.2.2.1 Propriétés**Définition 7.2.6**

Un *réseau vivant* est tel qu'à partir du marquage initial et de tout marquage consécutif, toute transition peut être incluse dans une séquence de franchissement.

Le réseau de la figure 7.22.a n'est pas vivant. La transition tr de ce réseau ne peut jamais être franchie. Ceci traduit une erreur de modélisation. Ce réseau fonctionne cependant sans blocage et se réinitialise.

Le réseau de la figure 7.22.b correspond à un cycle courant d'automatisme avec une phase d'initialisation. Bien que représentant exactement le comportement attendu, il n'est pas vivant au sens de la définition précédente.

Définition 7.2.7

Un réseau de Petri est *pseudo-vivant* si pour tous ses marquages accessibles, au moins une transition peut être franchie.

Cette propriété garantit l'absence de blocage mais pas l'absence de transition inutile ne pouvant jamais être franchie.

Les réseaux des figure 7.22.a et b sont pseudo-vivants alors que celui de la figure 7.22.c comporte une situation de blocage et n'est ni vivant, ni pseudo-vivant. Ce dernier ne peut modéliser un système temps réel correct.

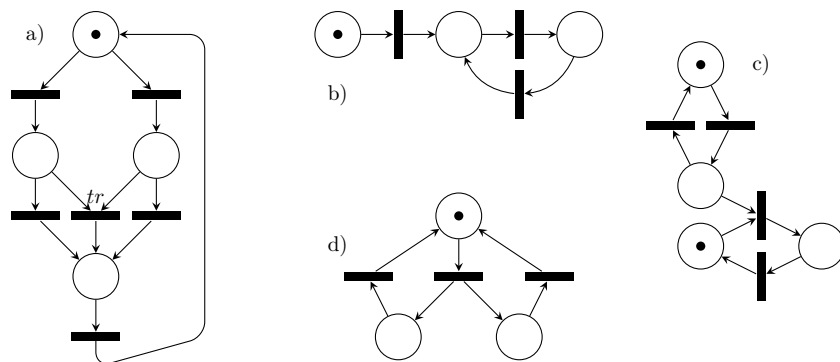


FIGURE 7.22 – Exemples de réseaux de Petri présentant des erreurs

Définition 7.2.8

Un réseau de Petri est *borné* si pour tout marquage accessible, le nombre de marques contenues dans chaque place du réseau reste inférieur à un nombre k .

Le réseau de la figure 7.22.d n'est pas borné. Il ne possède pas de fonctionnement cyclique et n'est pas réinitialisable.

Définition 7.2.9

Un réseau vivant dont la borne est égale à 1 est dit *sauf*.

C'est une propriété qui est souhaitable pour les modèles de systèmes logiques.

Il y a **conflit** lorsque plusieurs transitions sont franchissables et que le franchissement d'une d'entre elles invalide le franchissement d'une autre. La figure 7.23.a présente une telle situation. Il existe ainsi un conflit lorsque le nombre de marques dans une place est inférieur au nombre de transitions validées en sortie de cette place.

Le conflit de la figure 7.23.b n'est pas effectif car le marquage de la place partagée est égal au nombre de transitions validées en sortie.

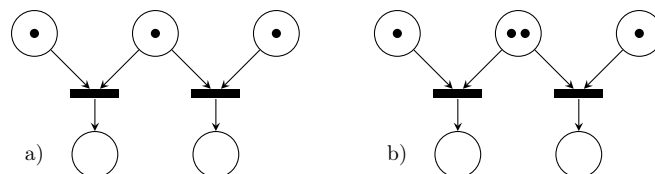


FIGURE 7.23 – Exemples de conflit

Définition 7.2.10

Un réseau de Petri est dit *persistant* s'il ne contient pas de situation de conflit.

L'analyse de ces principales propriétés des réseaux peut se faire avec les approches ci-après.

7.2.2.2 Graphe des marquages

A partir du marquage initial d'un réseau, il est possible de tracer un graphe dont les nœuds sont les marquages successifs et les arcs sont les transitions faisant passer d'un marquage à un autre. Le graphe de la figure 7.15 en est ainsi un exemple pour le réseau de Petri de la figure 7.13. L'analyse de ce graphe des marquages par des méthodes issues de la théorie des graphes permet de mettre en évidence des propriétés du réseau (cycles, marquages atteignables, etc.). Lorsque le réseau n'est pas borné, ce graphe a une structure d'arbre infini.

7.2.2.3 Recherche de structures particulières

Il s'agit ici d'isoler des structures de graphes ne vérifiant pas les propriétés fondamentales illustrées par la figure 7.24.

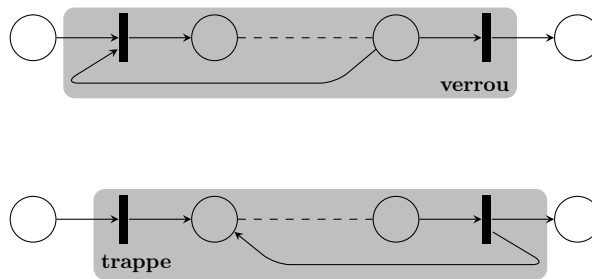


FIGURE 7.24 – Verrou et trappe dans un réseau de Petri

Définition 7.2.11

Un *verrou* est un ensemble de places tel que l'ensemble de ses transitions d'entrée est inclus dans l'ensemble de ses transitions de sortie.

Un verrou n'ayant aucune place marquée initialement ne pourra pas être activé.

Définition 7.2.12

Une *trappe* est un ensemble de places tel que l'ensemble de ses transitions de sortie est inclus dans l'ensemble de ses transitions d'entrée.

Un réseau de Petri contenant des trappes n'est pas borné.

7.2.3 Analyse par réduction

Afin de détecter des situations particulières, une méthode d'analyse procède par réduction du réseau, dont le principe est le suivant. La méthode est fondée sur la réduction du modèle par des transformations conservant les propriétés "vivants et bornés" des réseaux de Petri.

L'ensemble de ces transformations est présenté en figure 7.25.

Pour la **transformation 1** concernant une transition simple, la place de départ est origine d'un seul chemin. Il n'y a donc pas d'autre arc sortant que celui du chemin à simplifier.

Remarque 7.2.2. Lors de l'application de cette transformation, les jetons doivent être conservés.

Pour la **transformation 2**, la seconde transition n'a qu'une seule place d'entrée, la place devant être simplifiée. Il n'y a donc pas d'autre arc entrant sur cette transition.

La **transformation 3** permet de simplifier deux places qui sont toujours marquées simultanément et la **transformation 4** deux chemins redondants entre deux mêmes places.

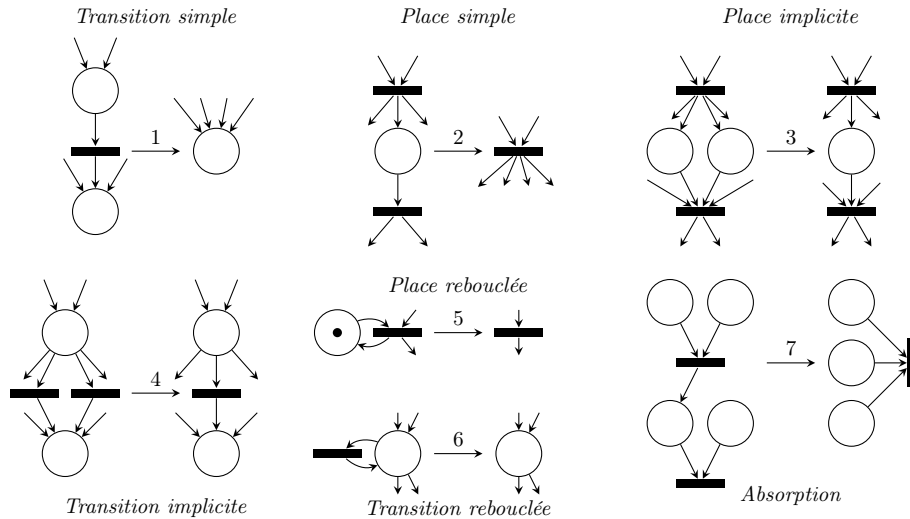


FIGURE 7.25 – Opérations de réduction d'un réseau de Petri

La **transformation 5** permet de simplifier une place toujours marquée qui ne conditionne pas le franchissement d'une transition et la **transformation 6** simplifie un chemin rebouclé. La **transformation 7** d'absorption est comparable à la **transformation 2**.

Exemple 7.2.3. Soit à modéliser l'accès concurrent à un magasin de pièces constitué d'un ensemble vertical de cellules organisées matriciellement et d'un transtockeur manipulant les éléments stockés (figure 7.26).

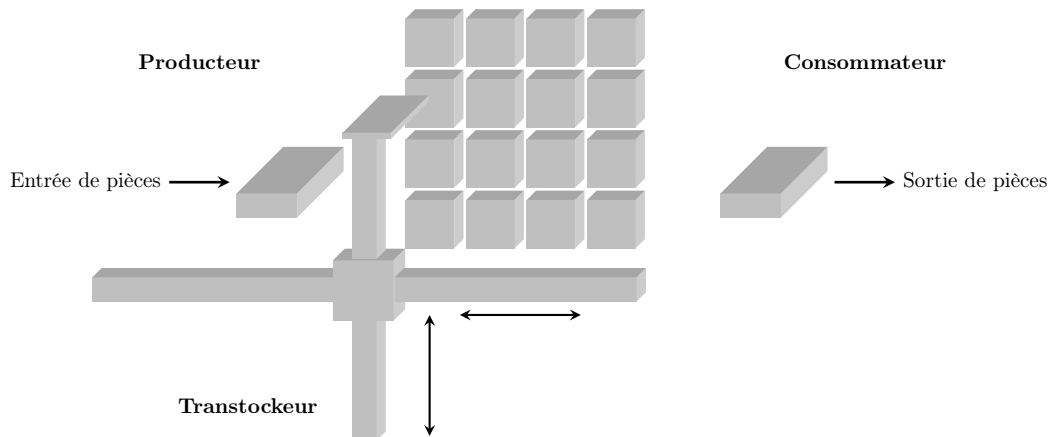


FIGURE 7.26 – Exemple d'un magasin de pièces

Le producteur crée un groupe de pièces, puis réserve le transtockeur pour le ranger dans le magasin s'il y a des emplacements libres.

S'il y a des pièces dans le magasin, le consommateur réserve le transtockeur lorsqu'il est libre afin de reprendre des pièces pour les traiter.

Une spécification des synchronisations entre le producteur et le consommateur est fournie en figure 7.27.a).

En appliquant les réductions 2 d'une place simple et 7 d'absorption, on obtient le réseau b). Ensuite, en utilisant la réduction d'une transition simple 1 et en conservant le marquage, on obtient le réseau réduit c). Ce réseau comporte plusieurs places marquées rebouclées que l'on peut supprimer pour obtenir d). Ensuite, on peut encore simplifier en supprimant la place simple et la transition rebouclée.

Le réseau e) qui est évidemment vivant et borné a les mêmes propriétés que le réseau a).

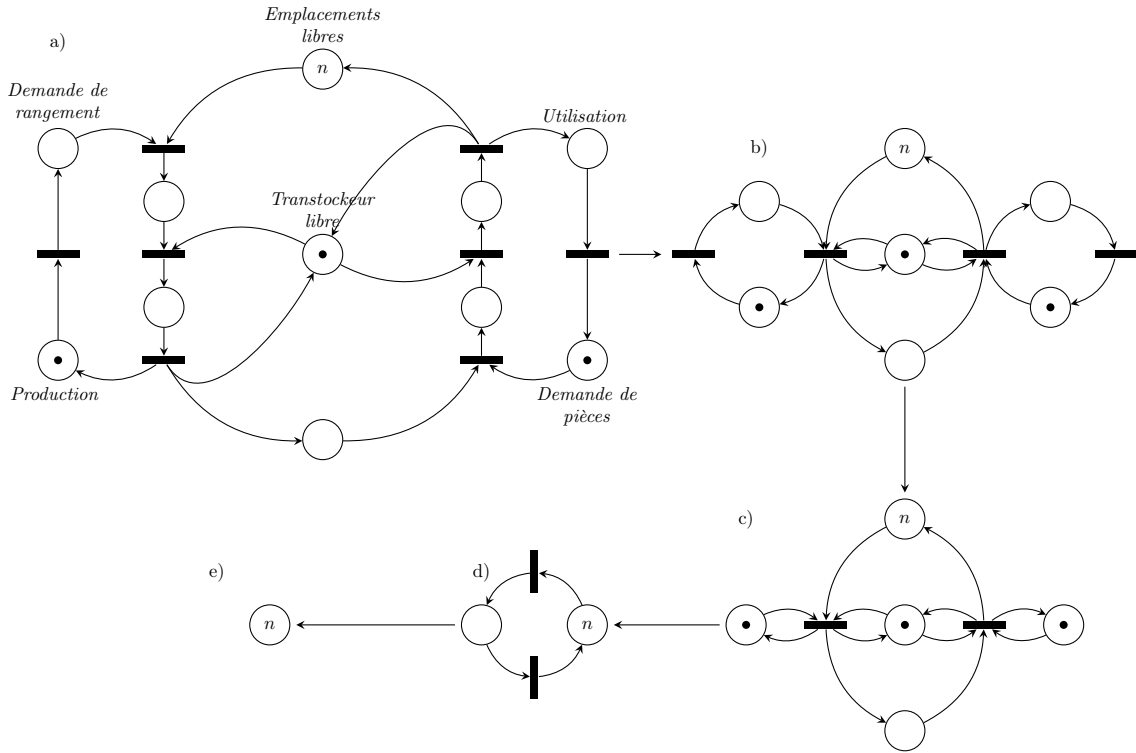


FIGURE 7.27 – Modélisation par réseau de Petri d'un magasin de pièces, et analyse par réduction du modèle

L'analyse des modèles n'est pas toujours aussi simple car certains réseaux restent compliqués à analyser, même après réduction. Cette méthode doit donc être complétée par une autre approche.

7.2.4 Analyse linéaire

7.2.4.1 Principe

Le modèle en réseaux de Petri est représenté sous une forme matricielle et les principes de l'analyse linéaire sont appliqués pour montrer certaines propriétés. La structure du modèle en réseaux de Petri est représentée par une **matrice d'incidence** :

$$W = [w_{ij}] \quad (7.2.1)$$

avec

$$w_{ij} = Post(P_i, tr_j) - Pre(P_i, tr_j), \quad (7.2.2)$$

pour $i \in \{1, \dots, n\}$, $j \in \{1, \dots, m\}$ (voir définition 7.2.5). Les lignes de la matrice d'incidence correspondent aux places et les colonnes aux transitions. Les éléments de la matrice $Pre(P, t)$ représentent les relations des transitions avec leurs places d'entrée. Ainsi, si la place d'indice P_i est place d'entrée de la transition d'indice tr_j , on a :

$$Pre(P_i, tr_j) = 1 \quad (7.2.3)$$

et dans le cas contraire :

$$Pre(P_i, tr_j) = 0 \quad (7.2.4)$$

De la même façon, les éléments de la matrice $Post(P, t)$ représentent les relations des transitions avec leurs places de sortie. Chaque colonne de la matrice d'incidence indique la modification du marquage lors du franchissement de la transition correspondante. La matrice d'incidence représente bien la structure du modèle à condition qu'il n'y ait pas de *boucle élémentaire* dans laquelle une même place est en entrée et en sortie d'une transition.

Dans ce cas, il faut introduire une place intermédiaire fictive.
Considérons à titre d'exemple le réseau représenté sur la figure 7.28.

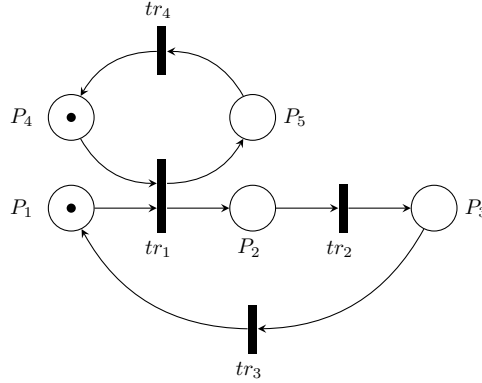


FIGURE 7.28 – Exemple de réseau pour illustrer l'analyse linéaire

Le marquage du réseau est représenté par un vecteur colonne $m(P)$ de taille égale au nombre de places du réseau. La validation d'une transition tr_i est testée par le critère suivant :

$$m(P) \geq -W(P, tr_i) \quad (7.2.5)$$

Le franchissement d'une transition tr_i validée correspond alors à l'opération matricielle suivante :

$$m_2(P) = m_1(P) + W(P, tr_i) \quad (7.2.6)$$

avec :

- $m_2(P)$: marquage résultant ;
- $m_1(P)$: marquage de départ ;
- $W(P, tr_i)$: colonne de la matrice correspondant à la transition tr_i .

Dans notre exemple, nous pouvons vérifier que la transition tr_1 est validée pour le marquage initial :

$$m(P) \geq -W(P, tr_1) \Rightarrow \begin{pmatrix} 1 \\ 0 \\ 0 \\ 1 \\ 0 \end{pmatrix} \geq \begin{pmatrix} 1 \\ -1 \\ 0 \\ 1 \\ -1 \end{pmatrix} \quad (7.2.7)$$

Le franchissement de cette transition tr_1 conduit au marquage suivant :

$$m_2(P) = m_1(P) + W(P, tr_1) = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 1 \\ 0 \end{pmatrix} + \begin{pmatrix} -1 \\ 1 \\ 0 \\ -1 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \\ 1 \end{pmatrix} \quad (7.2.8)$$

Le vecteur séquence $S(t)$ d'une dimension égale au nombre de transitions du réseau représente un ensemble des séquences de franchissement. Ainsi, le vecteur $S(t)$ transposé suivant :

$$S^T(t) = (1 \quad 2 \quad 1 \quad 0) \quad (7.2.9)$$

représente l'ensemble des séquences contenant une fois tr_1 , deux fois tr_2 , une fois tr_3 et pas tr_4 . Il est maintenant possible d'introduire la relation fondamentale qui, à partir d'un marquage initial m_i , détermine le marquage final m_f d'un réseau après franchissement d'une séquence $S(t)$:

$$m_f = m_i + W(P, t) \cdot S(t) \quad (7.2.10)$$

Remarque 7.2.3. L'équation (7.2.10) est similaire à la représentation d'état des systèmes à états continus présentée dans la partie I, avec la différence que le 'vecteur d'état' m (i.e., le marquage du réseau de Petri) ne peut pas être négatif et doit être un entier.

Ainsi dans notre exemple, pour le franchissement d'une séquence contenant tr_2, tr_3, tr_4 à partir du marquage $m_2 = (0 \ 1 \ 0 \ 0 \ 1)^T$, on a :

$$m_f(P) = \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \\ 1 \end{pmatrix} + \begin{pmatrix} -1 & 0 & 1 & 0 \\ 1 & -1 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ -1 & 0 & 0 & 1 \\ 1 & 0 & 0 & -1 \end{pmatrix} \cdot \begin{pmatrix} 0 \\ 1 \\ 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 1 \\ 0 \end{pmatrix} \quad (7.2.11)$$

7.2.4.2 Application à la recherche de cycle

L'ensemble des cycles vérifie la relation :

$$m_f(P) = m_i(P) \quad (7.2.12)$$

Rechercher les cycles consiste donc à trouver les solutions $X(t)$ à composantes entières à l'équation :

$$W(P, t) \cdot X(t) = \mathbf{0} \quad (7.2.13)$$

En reprenant notre exemple, on a :

$$W(P, t) \cdot X(t) = \begin{pmatrix} -1 & 0 & 1 & 0 \\ 1 & -1 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ -1 & 0 & 0 & 1 \\ 1 & 0 & 0 & -1 \end{pmatrix} \cdot \begin{pmatrix} a \\ b \\ c \\ d \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix} \Rightarrow X(t) = \begin{pmatrix} \alpha \\ \alpha \\ \alpha \\ \alpha \end{pmatrix}, \text{ avec } \alpha \in \mathbb{Z} \quad (7.2.14)$$

Les solutions sont donc toutes les séquences pour lesquelles les transitions du réseau sont franchies un même nombre de fois. Sur le même principe, on peut vérifier que certains marquages sont atteignables.

7.2.4.3 Application à l'amélioration de la sûreté de fonctionnement

Dans ce cas, on recherche un vecteur pondération $\tilde{P}(P)$ qui s'il existe est tel que :

$$\tilde{P}^T(P)m_f = \tilde{P}^T(P)m_i = k \quad (7.2.15)$$

et ceci quels que soient les marquages initiaux m_i et finaux m_f donc, quelle que soit la séquence de franchissement $S(t)$ effectuée.

On a par conséquent :

$$\tilde{P}^T(P)m_f = \tilde{P}^T(P)m_i + \tilde{P}^T(P)W(P, t) \cdot S(t) \quad (7.2.16)$$

Trouver $\tilde{P}(P)$ revient donc à chercher les solutions de :

$$\tilde{P}^T(P)W(P, t) = \mathbf{0} \quad (7.2.17)$$

Dans notre exemple, on a :

$$\tilde{P}^T(P)W(P, t) = \begin{pmatrix} a \\ b \\ c \\ d \\ e \end{pmatrix}^T \cdot \begin{pmatrix} -1 & 0 & 1 & 0 \\ 1 & -1 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ -1 & 0 & 0 & 1 \\ 1 & 0 & 0 & -1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix} \Rightarrow \tilde{P}(P) = \begin{pmatrix} \lambda \\ \lambda \\ \lambda \\ \gamma \\ \gamma \end{pmatrix} \quad (7.2.18)$$

avec $\alpha, \gamma \in \mathbb{Z}$ et donc pour tout marquage valide :

$$\tilde{P}^T(P)m(P) = \lambda + \gamma \quad (7.2.19)$$

L'existence de ce vecteur de pondération permet ainsi, dans un système supposé non fiable, de vérifier cette condition pour tout nouveau marquage calculé.

7.2.5 Réseaux non autonomes synchronisés

Définition 7.2.13

Un *réseau de Petri synchronisé* est un réseau de Petri autonome muni d'une application d'un ensemble d'événements sur l'ensemble de ses transitions.

Parmi ces événements, on distingue l'événement e toujours vrai. A toute transition tr_i du réseau est donc associé un événement e_i .

Définition 7.2.14

La transition tr_i est dite *réceptive à l'événement e_i* si elle est validée, c'est-à-dire si toutes ses places d'entrée sont marquées. Sur une occurrence de l'événement e_i , la transition devient *franchissable*.

L'algorithme d'interprétation d'un réseau synchronisé est le suivant :

répéter

- déterminer les événements vrais ;
- déterminer les transitions franchissables ;
- effectuer une séquence de franchissement contenant ces transitions ;

fin_répéter

Ainsi, un événement se produit à un instant déterminé et il n'entraîne qu'un seul franchissement de transition dans les cas présentés en figure 7.29.

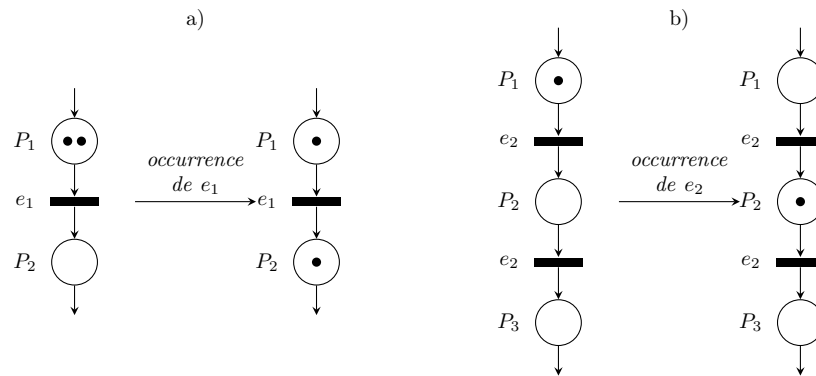


FIGURE 7.29 – Exemple de franchissement de transition dans le cas d'un réseau de Petri synchronisé

Chapitre 8

Modélisation des systèmes à événements discrets temporisés

8.1 Introduction

Nous avons vu dans le chapitre précédent la modélisation de systèmes à événements discrets pour lesquels seule l'occurrence d'événements était à prendre en compte pour en décrire l'évolution indépendamment des instants d'apparition de ces événements. Nous nous intéressons ici à des systèmes à événements discrets pour lesquels le temps exerce une influence sur le comportement du système. Ainsi, la trajectoire d'un tel système n'est plus simplement caractérisée par la suite des événements $\{\sigma_1, \sigma_2, \dots\}$ mais par la suite des événements et de leurs instants d'occurrence $\{(\sigma_1, t_1), (\sigma_2, t_2), \dots\}$. L'utilisation de tels modèles de systèmes à événements discrets permettra de répondre à des questions telles que :

- Combien de temps le système va-t-il / doit-il rester dans un état donné ?
- Combien peut-il / doit-il y avoir d'événements entre deux instants donnés ?
- Quel est le temps d'attente entre deux événements ?
- ...

Avant de définir les automates temporisés à la section 8.3 et leurs extensions sous forme d'automates temporisés avec garde à la section 8.4, nous devons définir formellement la notion d'horloge, ce que nous ferons à la section 8.2. De façon assez naturelle, nous présenterons ensuite la modélisation de systèmes dits hybrides à la section 8.5, où les événements discrets entraînent des transitions entre états, états dans lesquels le système évolue avec une dynamique continue particulière. Enfin, nous reprendrons brièvement le formalisme des réseaux de Petri pour l'étendre au cas des réseaux de Petri temporisés à la section 8.6.

8.2 Structure d'horloge

Afin d'introduire simplement la notion de structure d'horloge, nous allons l'illustrer sur un exemple simple. Considérons donc un automate tel que celui défini au chapitre 7, sous la forme $\mathcal{A} = (Q, \Sigma, \phi, Q_0)$. Pour simplifier, nous considérons ici que $\Sigma = (\alpha, \beta)$ et que la fonction de transition est définie pour tout état, c'est-à-dire que les fonctions $\phi(q, \alpha)$ et $\phi(q, \beta)$ sont définies $\forall q \in Q$.

Définition 8.2.1

Un événement σ est actif dans l'état q si la transition $\phi(q, \sigma)$ est définie.

Clairement, les événements α et β sont tout le temps actifs dans notre exemple. Dans le cas général, les événements vont devenir successivement actifs ou inactifs au fur et à mesure de l'évolution du système. On attribue à l'événement σ une *durée de vie* notée $v_{\sigma, m}$ la $m^{\text{ième}}$ fois où celui-ci devient actif. Une fois qu'un événement devient actif, une horloge décompte le temps à partir de

$v_{\sigma,m}$ jusqu'à atteindre 0. Plusieurs événements pouvant être actifs, l'évènement dont l'horloge arrive à 0 la première doit se produire. L'évènement en question se produit et fait évoluer le système dans un nouvel état q' . Dans cet état q' :

- Si σ s'est produit et σ est de nouveau actif dans q' , alors une nouvelle durée de vie $v_{\sigma,m+1}$ lui est attribuée ;
- Si σ s'est produit et σ n'est pas actif dans q' , alors aucune durée de vie ne lui est attribuée jusqu'à ce que le système soit de nouveau dans un état où σ est actif ;
- Si σ ne s'est pas produit et σ est toujours actif dans q' , alors son horloge continue de décompter le temps ;
- Si σ ne s'est pas produit et σ n'est pas actif dans q' , alors son horloge est arrêtée et une nouvelle durée de vie lui sera attribuée lorsque le système sera de nouveau dans un état où σ est actif.

La figure 8.1 présente un exemple d'évolution du système. Les états sont notés x_i sur l'abscisse. On ne précise pas plus avant les états dans cet exemple, notre but étant seulement de décrire une structure d'horloge.

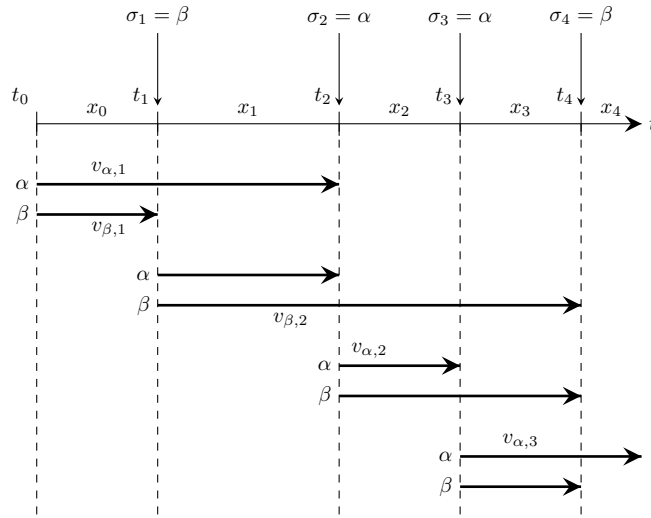


FIGURE 8.1 – Illustration de la structure d'horloge avec deux événements toujours actifs

A l'instant initial t_0 , le système est dans l'état x_0 et les deux événements sont actifs et se voient attribuer une durée de vie $v_{\alpha,1}$ et $v_{\beta,1}$. Comme $v_{\beta,1} < v_{\alpha,1}$, à $t = t_1 = t_0 + v_{\beta,1}$, l'évènement $\sigma_1 = \beta$ se produit provoquant la transition du système dans l'état x_1 . Les deux événements étant toujours actifs dans cet état, une nouvelle durée de vie $v_{\beta,2}$ est attribuée à l'évènement β , tandis que l'horloge correspondant à α continue de décompter le temps. A $t = t_2$, l'horloge associée à α arrive à 0 la première. L'évènement α doit donc se produire, le système évolue vers l'état x_2 , une nouvelle durée de vie $v_{\alpha,2}$ est attribuée à α , de nouveau actif quand le système est en x_2 , tandis que l'horloge associée à β continue de décompter le temps ; β étant toujours actif en x_2 . Le système continue ensuite d'évoluer selon les mêmes règles.

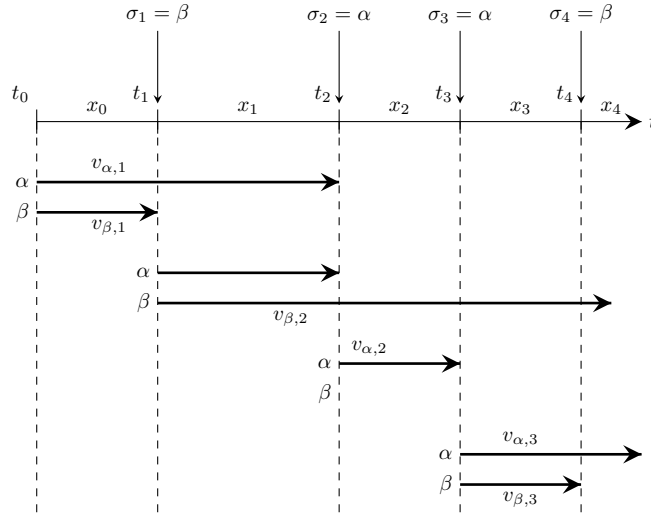
Supposons maintenant que l'évènement β ne soit pas actif dans l'état x_2 (c'est à dire que la transition $\phi(x_2, \beta)$ n'est pas définie). L'évolution du système se fait alors selon la figure 8.2.

Dans l'état x_2 , seul l'évènement α est actif. L'horloge correspondant à β est supprimée, même si elle n'est pas arrivée à 0. A $t = t_3$, l'évènement α se produit, provoquant une transition vers x_3 . Dans cet état, les deux événements sont actifs et deux nouvelles durées de vie $v_{\alpha,3}$ et $v_{\beta,3}$ sont attribuées à α et β .

Nous pouvons désormais définir une structure d'horloge dans le cas général.

Définition 8.2.2

Une *structure d'horloge* associée à un ensemble d'évènements Σ est la donnée d'un ensemble $\mathbf{V} = \{ \mathbf{v}_\sigma, \sigma \in \Sigma \}$ de séquences d'horloges $\mathbf{v}_\sigma = \{v_{\sigma,1}, v_{\sigma,2}, \dots\}$.

FIGURE 8.2 – Illustration de la structure d'horloge avec un évènement β pas actif dans l'état x_2

8.3 Automates temporisés

8.3.1 Définition et simulation d'un automate temporisé

On définit également le score d'un évènement :

Définition 8.3.1

Le score d'un évènement σ après la $k^{\text{ième}}$ transition, noté $N_{\sigma,k}$ est le nombre de fois que l'évènement σ a été activé depuis l'instant initial.

Ainsi, pour l'exemple de la figure 8.1, à l'instant initial, on a $N_{\alpha,0} = 1$ et $N_{\beta,0} = 1$ puisque les deux évènements sont activés à $t = t_0$. Comme β se produit à l'instant $t = t_1$ on a $N_{\beta,1} = 2$ et comme les deux évènements sont possibles dans l'état x_1 , on a $N_{\alpha,1} = 1$. Ce score permet de repérer la nouvelle valeur de la durée de vie à attribuer dans la séquence de la structure d'horloge lorsqu'un évènement devient actif.

Dans la suite, nous noterons $c_{\sigma}(t)$ la valeur courante de l'horloge de l'évènement σ et $\Delta\tau_k = t_{k+1} - t_k$ le temps passé dans l'état x_k . Ainsi, supposons que le système soit dans l'état x_k correspondant à l'état après les k premières transitions, chaque état ayant un score $N_{\sigma,k}$. L'évènement σ_1 se produit occasionnant la $k + 1^{\text{ième}}$ transition vers l'état x_{k+1} . Nous avons les évolutions suivantes pour les scores des évènements :

- $N_{\sigma_1,k+1} = N_{\sigma_1,k} + 1$ si σ_1 est de nouveau actif dans l'état x_{k+1} ;
- $N_{\sigma_1,k+1} = N_{\sigma_1,k}$ si σ_1 n'est pas actif dans l'état x_{k+1} ;
- $N_{\sigma_2,k+1} = N_{\sigma_2,k} + 1$ si σ_2 n'était pas actif dans x_k et est de nouveau actif dans l'état x_{k+1} ;
- $N_{\sigma_2,k+1} = N_{\sigma_2,k}$ si σ_2 n'est pas actif dans l'état x_{k+1} ;

Pour les valeurs des horloges, à l'instant t_{k+1} :

- $c_{\sigma_1}(t_{k+1}) = v_{\sigma,N_{\sigma_1,k+1}}$ si σ_1 est actif dans l'état x_{k+1} ;
- $c_{\sigma_2}(t_{k+1}) = v_{\sigma,N_{\sigma_2,k+1}}$ si σ_2 est actif dans l'état x_{k+1} , mais ne l'était pas dans l'état x_k ;
- $c_{\sigma_2}(t_{k+1}) = v_{\sigma,N_{\sigma_2,k}} - \Delta\tau_k$ si σ_2 est actif dans l'état x_{k+1} et l'était aussi dans l'état x_k .

Ces évolutions des scores et des horloges associés à chaque évènement permettent alors de définir un automate temporisé.

Définition 8.3.2

Un automate temporisé fini $\mathcal{A}$ peut être représenté par un quintuplet $\mathcal{A} = (Q, \Sigma, \phi, Q_0, \mathbf{V})$, avec :

- Q un ensemble d'états ;
- Σ un ensemble de symboles ;
- ϕ la relation de transition définie de $Q \times \Sigma$ dans Q ;
- $Q_0 \subset Q$ un ensemble d'états initiaux ;
- $\mathbf{V}$ la structure d'horloge associée.

Dans cette définition, Q, Σ, ϕ, Q_0 ont bien sur la même interprétation que dans le chapitre 7. La structure d'horloge $\mathbf{V}$ correspond bien entendu à ce qui a été présenté à la section 8.2.

La simulation d'un automate temporisé revient ainsi à suivre l'algorithme suivant lorsque le système est dans l'état x_k pour déterminer la transition suivante, l'instant de transition et le futur état :

- **Etape 1.** Déterminer le prochain événement par l'équation :

$$\sigma' = \arg \min_{\sigma \in \Sigma} \{c_\sigma(t_k)\} \quad (8.3.1)$$

sous contrainte que σ est actif dans l'état x_k

- **Etape 2.** Déterminer l'instant du prochain événement :

$$t_{k+1} = t_k + c_{\sigma'}(t_k) \quad (8.3.2)$$

- **Etape 3.** Déterminer l'état suivant :

$$x_{k+1} = \phi(x_k, \sigma') \quad (8.3.3)$$

- **Etape 4.** Remettre à jour les scores $N_{\sigma, k+1}$ et les horloges $c_\sigma(t_{k+1}), \forall \sigma \in \Sigma$.

Exemple 8.3.1. *Modélisation d'un système de file d'attente sous la forme d'un automate temporisé. On considère ici un système de file d'attente dans lequel 2 événements peuvent se produire, à savoir l'arrivée d'un client, notée a et le départ d'un client (après traitement) d . L'état du système peut de façon naturelle être caractérisé par la longueur de la file d'attente. Le système peut se représenter sous la forme d'un automate non temporisé conformément à la figure 8.3 ci-dessous, dans lequel la file d'attente est vide à l'instant initial.*

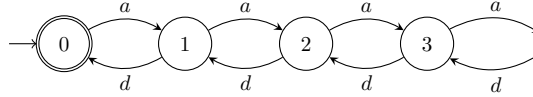


FIGURE 8.3 – Modélisation d'une file d'attente sous la forme d'un automate non temporisé

Si ce modèle permet de décrire les évolutions possibles de l'état, il est insuffisant pour prédire réellement ce qu'il va advenir du système. Pour ce faire il est nécessaire de connaître le temps de traitement d'un client, ainsi que les dates d'arrivée de ces clients. Ces données correspondent à des structures d'horloges pour les événements a et d : $v_{a,k}$ correspond à l'intervalle de temps entre les clients $k-1$ et k , tandis que $v_{d,k}$ correspond au temps de traitement du client k . Hormis dans l'état 0, les deux événements sont actifs dans tous les états. Ce type de système pourra se modéliser de manière très simple en utilisant un chronogramme comme sur la figure 8.4.

Notons pour finir que nous avons supposé ici que les dates d'arrivée des événements et les temps de traitement étaient ici connus à l'avance. Ceci n'est que très rarement le cas, et il sera nécessaire pour caractériser le comportement de nombreux systèmes, d'introduire des distributions de probabilités pour ces grandeurs. C'est ce que nous ferons au chapitre 9.

8.4 Automates temporisés avec gardes

8.4.1 Introduction

Les automates temporisés avec gardes sont une alternative des automates temporisés de la section 8.3. Nous considérons ainsi qu'un certain nombre d'horloges c_i sont définies et nécessaires

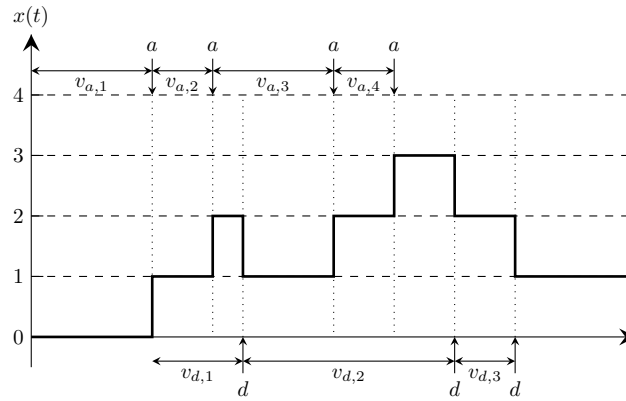


FIGURE 8.4 – Chronogramme représentant l'évolution du système de file d'attente

pour la modélisation du système. Dans chaque état, ces horloges obéissent à une évolution de la forme :

$$\frac{dc_i}{dt}(t) = 1 \quad (8.4.1)$$

Dans ce formalisme, chaque transition est en fait associée à un triplet (*garde*, *evenement*, *reset*) dans lequel :

- *garde* correspond à un ensemble de conditions que doivent satisfaire les horloges décrivant l'automate temporisé pour que la transition soit autorisée ;
- *evenement* a la signification habituelle et correspond à l'évènement permettant la transition ;
- *reset* rassemble l'ensemble des horloges qui doivent être remises à zéro au moment de la transition.

Exemple 8.4.1. Nous considérons l'exemple d'un distributeur de boisson. Lorsque le client introduit une pièce, il peut entrer son choix après une seconde. Il doit avoir fait son choix avant 10 secondes, sinon, l'automate lui rend sa pièce et il faut recommencer. Lorsque le choix est fait, le distributeur doit délivrer la boisson dans un intervalle de temps compris entre 2 et 5 secondes, sinon il rend la pièce. La modélisation du système peut se faire selon la représentation de la figure 8.5. Les ensembles suivantes sont définis : $\text{garde} = \{1 \leq c_{\text{client}} < 10; 2 \leq c_{\text{distributeur}} < 5; c_{\text{distributeur}} \geq 5\}$, $\text{evenement} = \{\text{pièce}, \text{choix}, \text{livraison}, \text{erreur}, \text{indécision}\}$, $\text{reset} = \{c_{\text{client}}, c_{\text{distributeur}}\}$.

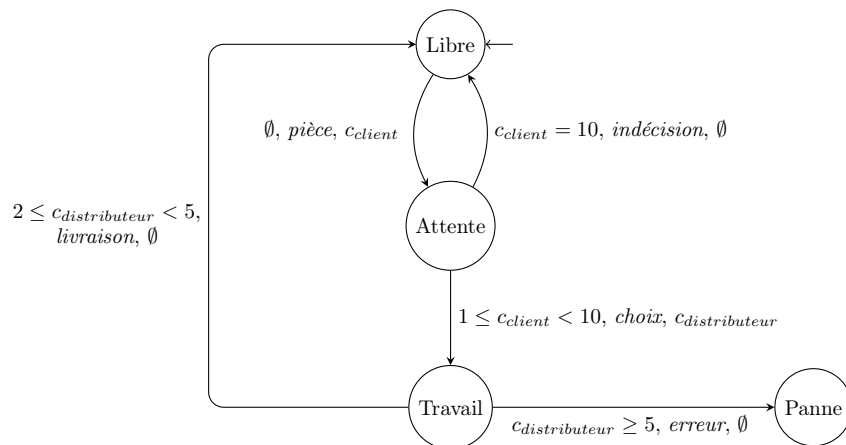


FIGURE 8.5 – Modélisation d'un distributeur de boissons

L'introduction peut se faire à n'importe quel instant (ce qui est traduit sur la transition de Libre à Attente par le symbole $\emptyset$) lorsque l'automate est dans l'état Libre, ce qui produit la transition

vers l'état *Attente* et la remise à 0 de l'horloge c_{client} . Si le client fait son choix, l'automate passe dans le mode *Travail*. L'horloge $c_{distributeur}$ est remise à zéro. Si la livraison est faite en moins de 5 secondes le distributeur revient à l'état *Libre*, sinon, il passe à l'état *Panne*.

8.4.2 Ajout d'invariants

L'exemple précédent montre que la modélisation précédente est incomplète :

- Dans l'état *Attente*, si aucun événement n'a lieu avant que $c_{client} > 10$, le système restera bloqué dans cet état.
- Dans l'état *Travail*, rien n'oblige le système à aller dans l'état *Panne* si $c_{distributeur}$ devient strictement supérieur à 5.

En effet, les gardes associées aux transitions indiquent les conditions nécessaires pour qu'une transition *puisse* se produire, mais n'obligent pas à ce que cette transition se produise effectivement. Dans le formalisme précédent, il est donc nécessaire d'ajouter des conditions pour forcer les transitions. Ces conditions, associées aux états, sont appelées des invariants. Ces conditions doivent être satisfaites pour pouvoir rester dans un état. Ainsi, nous ajoutons à l'état *Attente* l'invariant $c_{client} \leq 10$ et à l'état *Travail* l'invariant $c_{distributeur} \leq 5$. L'automate correspondant est représenté sur la figure 8.6.

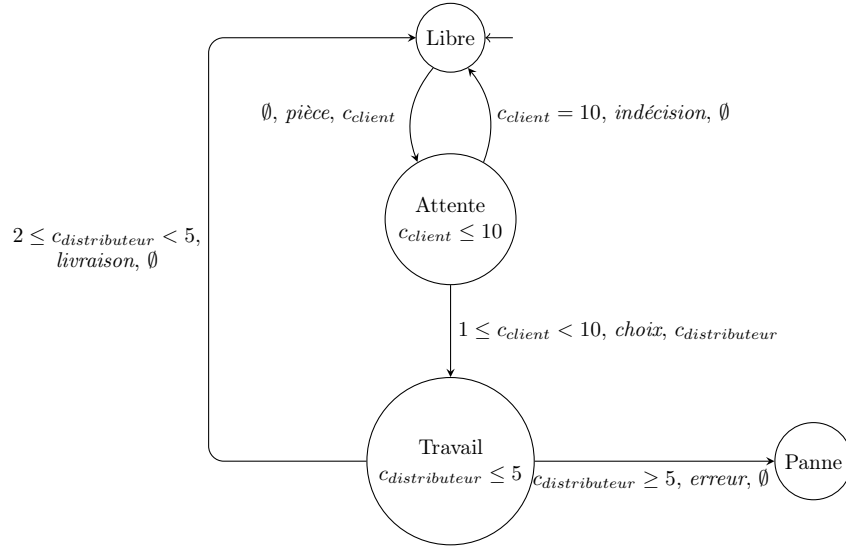


FIGURE 8.6 – Modélisation d'un distributeur de boissons avec invariants

8.4.3 Formalisation

L'exemple précédent peut se généraliser pour obtenir la définition suivante d'un automate temporisé avec gardes et invariants.

Définition 8.4.1

Un automate temporisé avec garde est un sextuplet $\mathcal{A} = (Q, \Sigma, C, Tra, Inv, q_0)$, avec :

- Q l'ensemble des états ;
- Σ l'ensemble (fini) des événements ;
- C l'ensemble des horloges, $\{c_1(t), \dots, c_n(t)\}$, avec $c_i(t) \in \mathbb{R}_+$ et $t \in \mathbb{R}_+$;
- Tra l'ensemble des transitions temporisées de l'automate, c'est à dire que :

$$Tra \subseteq Q \times \mathcal{C}(C) \times \Sigma \times 2^C \times Q \quad (8.4.2)$$

où $\mathcal{C}(C)$ représente l'ensemble des conditions admissibles pour les horloges ;

- Inv est l'ensemble des invariants pour les états :

$$Inv : Q \mapsto \mathcal{C}(C) \quad (8.4.3)$$

- q_0 l'état initial.

L'ensemble Tra doit se comprendre de la façon suivante : si $(q_{in}, garde, evenement, reset, q_{out}) \in Tra$, alors il existe une transition temporisée $(garde, evenement, reset)$ permettant d'aller de q_{in} à q_{out} .

8.4.4 Parallélisation d'automates temporisés avec gardes

Beaucoup de notions existant pour les automates non temporisés se transposent au cas des automates temporisés. C'est par exemple le cas de la parallélisation d'automates qui peut se définir de la façon suivante.

Définition 8.4.2

La parallélisation de deux automates temporisés avec gardes $\mathcal{A}_1 = (Q_1, \Sigma_1, C_1, Tra_1, Inv_1, q_{0,1})$ et $\mathcal{A}_2 = (Q_2, \Sigma_2, C_2, Tra_2, Inv_2, q_{0,2})$, est définie par un automate temporisé avec gardes :

$$\mathcal{A}_1 || \mathcal{A}_2 = Ac(Q_1 \times Q_2, \Sigma_1 \cup \Sigma_2, C_1 \cup C_2, Tra_{1,2}, Inv_{1,2}, (q_{0,1}, q_{0,2})) \quad (8.4.4)$$

avec :

- les conditions d'invariance obtenues par les intersections des conditions d'invariance des automates :

$$\begin{aligned} Inv_{1,2} : Q_1 \times Q_2 &\mapsto \mathcal{C}(C)_{1,2} = \mathcal{C}(C)_1 \wedge \mathcal{C}(C)_2 \\ Inv_{1,2}(q_1, q_2) &= Inv_1(q_1) \wedge Inv_2(q_2) \end{aligned} \quad (8.4.5)$$

- les conditions de garde obtenues comme :

$$Tra_{1,2} \subseteq (Q_1 \times Q_2) \times \mathcal{C}(C)_{1,2} \times (\Sigma_1 \cup \Sigma_2) \times 2^{(C_1 \cup C_2)} \times (\Sigma_1 \cup \Sigma_2) \quad (8.4.6)$$

et construites de la façon suivante :

- Pour tout $\sigma_c \in \Sigma_1 \cap \Sigma_2$,
si $(q_{in,1}, garde_1, \sigma_c, reset_1, q_{out,1}) \in Tra_1$ et $(q_{in,2}, garde_2, \sigma_c, reset_2, q_{out,2}) \in Tra_2$,
alors $((q_{in,1}, q_{in,2}), garde_1 \wedge garde_2, \sigma_c, reset_1 \cup reset_2, (q_{out,1}, q_{out,2})) \in Tra_{1,2}$;
- Pour tout $\sigma_1 \in \Sigma_1 \setminus \Sigma_2$ et $q_2 \in Q_2$, si $(q_{in,1}, garde_1, \sigma_1, reset_1, q_{out,1}) \in Tra_1$ alors
 $((q_{in,1}, q_2), garde_1, \sigma_1, reset_1, (q_{out,1}, q_2)) \in Tra_{1,2}$;
- Pour tout $\sigma_2 \in \Sigma_2 \setminus \Sigma_1$ et $q_1 \in Q_1$, si $(q_{in,2}, garde_2, \sigma_2, reset_2, q_{out,2}) \in Tra_2$ alors
 $((q_1, q_{in,2}), garde_2, \sigma_2, reset_2, (q_1, q_{out,2})) \in Tra_{1,2}$.

Rappelons que cette notion est d'importance dans le cas de la modélisation de systèmes complexes puisqu'il s'agit d'un des outils permettant la décomposition d'un système en sous-systèmes, le comportement global s'en déduisant par parallélisation.

Exemple 8.4.2. Nous reprenons ici un exemple tiré de (Cassandras and Lafortune, 2008). La figure 8.7 présente la modélisation d'un système recevant des messages. Les messages doivent être reçus avec un intervalle supérieur à une seconde, sinon le système émettra une alarme au bout d'une seconde au maximum et sera bloqué.

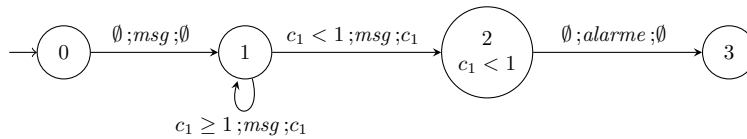


FIGURE 8.7 – Automate temporisé modélisant la réception de messages

Un deuxième système, représenté sur la figure 8.8, obéit au même principe de fonctionnement, mais avec des constantes de temps différentes : il doit recevoir un message avec un intervalle de temps inférieur à 5 secondes, sinon une deuxième alarme est émise.

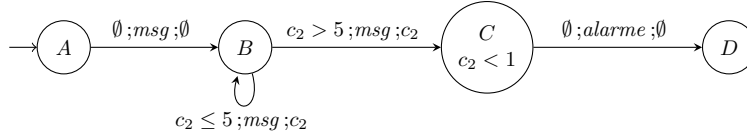


FIGURE 8.8 – Automate temporisé modélisant la réception de messages avec des constantes de temps différentes

La figure 8.9 présente alors l'automate obtenu à partir de la définition de la parallélisation des deux automates. De manière intuitive, le système global fonctionne si les messages sont reçus avec un intervalle de temps compris entre 1 et 5 secondes, sinon l'une ou l'autre des alarmes sera émise.

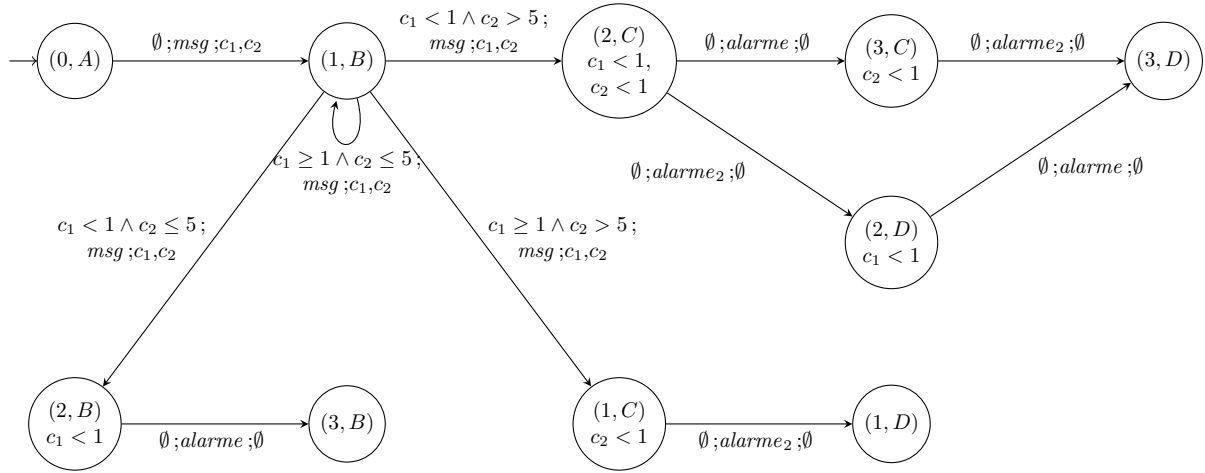


FIGURE 8.9 – Automate temporisé modélisant la réception de messages avec des constantes de temps différentes

Notons que le système global s'arrête lorsque l'un des deux systèmes s'arrête, car l'évènement msg appartient à Σ_1 et Σ_2 . Dans ce cas, selon la définition, pour qu'une transition existe sur l'automate parallèle, elle doit exister sur les deux automates initiaux.

8.5 Systèmes hybrides

8.5.1 Présentation informelle

Les systèmes décrits par des automates temporisés avec gardes peuvent être de façon très naturelle étendus au cas de systèmes hybrides. Les systèmes hybrides sont en effet modélisés par des états discrets q (au sens utilisé jusqu'à présent) et par des états continus que nous noterons x . Dans chacun des états discrets (on parle parfois de *mode de fonctionnement du système*), l'état continu évolue en suivant une description similaire à ce qui a été défini dans la partie I correspondant à la modélisation des systèmes à états continus (par exemple sous la forme d'une équation d'état continue). Notons que ces dynamiques continues peuvent dépendre d'entrées exogènes comme des variables de commande. Les conditions d'invariance et de garde sont alors constituées de contraintes sur les états continus ou les horloges.

Ainsi, la modélisation de systèmes sous la forme d'automates temporisés avec garde est un cas particulier de système hybride dans lequel les états continus sont représentés par les horloges et les dynamiques continues associées à chaque état sont très simples et décrites par $dc_i(t)/dt = 1$.

8.5.2 Description formelle

De manière plus rigoureuse, un système hybride décrit par un automate peut se définir de la façon suivante.

Définition 8.5.1

Un *automate hybride* est représenté par la donnée de :

$$(\mathcal{A})_h = (Q, X, \Sigma, U, f, \phi, Inv, Garde, \rho, q_0, x_0) \quad (8.5.1)$$

Avec :

- Q un ensemble d'états discrets ou modes de fonctionnement ;
- X un espace de variables d'états continues ;
- Σ un ensemble d'évènements discrets ;
- U un ensemble de signaux extérieurs, appelés souvent entrées de commande ;
- f un champ de vecteur, c'est à dire une application de la forme $(Q \times X \times U) \mapsto X$;
- ϕ une fonction de transition discrète (partielle), c'est à dire une application de la forme $\phi : (Q \times X \times \Sigma) \mapsto Q$;
- Inv un ensemble décrivant les invariants : $Inv \subseteq (Q \times X)$;
- $Garde$ un ensemble décrivant les conditions de garde : $Garde \subseteq (Q \times Q \times X)$;
- ρ une fonction de reset $\rho : (Q \times Q \times X \times \Sigma) \mapsto X$;
- q_0 un état discret initial ;
- x_0 un état continu initial.

Ainsi, lorsque le système est dans un mode de fonctionnement q_i , l'état continu obéit à une équation différentielle continue de la forme :

$$\frac{dx(t)}{dt} = f(q_i, x(t), u(t), t) \quad (8.5.2)$$

Le système peut rester dans cet état q_i tant que la condition d'invariant associée est vérifiée, c'est à dire tant que :

$$(q_i, x(t)) \in Inv \quad (8.5.3)$$

Si un évènement σ se produit à un instant t_0 , alors que l'état continu a la valeur $x(t)$, le système opère une transition vers l'état q_j si :

$$(q_i, q_j, x(t)) \in Garde \quad (8.5.4)$$

Lorsque la transition est effectuée, certains états sont remis à 0 (ou *réinitialisés*) selon la fonction ρ :

$$x_{ini}(t_0^+) = \rho(q_i, q_j, x(t_0), \sigma) \quad (8.5.5)$$

Illustrons cette définition sur l'exemple ci-dessous.

Exemple 8.5.1. On considère un système de deux cuves tel que représenté sur la figure 8.10. On dispose deux vannes V_1 et V_2 qui peuvent être en position ouverte O ou fermée F . Le système a donc 4 états discrets (ou 4 modes de fonctionnement) qui sont $(F, F), (F, O), (O, F), (O, O)$. Le système a également deux états continus qui sont les niveaux d'eau dans les cuves, notés x_1 et x_2 . Lorsque la vanne 1 est ouverte, son débit est constant et noté D . Lorsque la vanne 2 est ouverte, le débit qui va de la cuve 1 à la cuve 2 est proportionnel à la racine carrée du niveau dans la cuve 1. La vanne 2 possède toujours un débit de fuite qui est proportionnel à la racine carrée du niveau dans la cuve 2. On suppose pour simplifier que les sections des cuves sont égales à $1m^2$.

Il est facile de déterminer les équations d'évolution du système dans chacun des modes de fonctionnement :

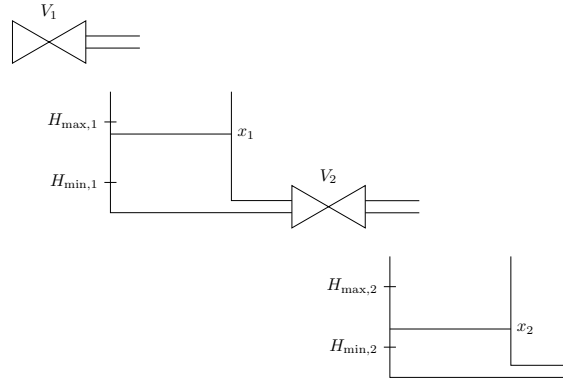


FIGURE 8.10 – Système de deux cuves

$$(F, F) : \begin{cases} \frac{dx_1(t)}{dt} = 0 \\ \frac{dx_2(t)}{dt} = -k_2 \sqrt{x_2(t)} \end{cases} \quad (8.5.6)$$

$$(F, O) : \begin{cases} \frac{dx_1(t)}{dt} = -k_1 \sqrt{x_1(t)} \\ \frac{dx_2(t)}{dt} = k_1 \sqrt{x_1(t)} - k_2 \sqrt{x_2(t)} \end{cases} \quad (8.5.7)$$

$$(O, F) : \begin{cases} \frac{dx_1(t)}{dt} = D \\ \frac{dx_2(t)}{dt} = -k_2 \sqrt{x_2(t)} \end{cases} \quad (8.5.8)$$

$$(O, O) : \begin{cases} \frac{dx_1(t)}{dt} = D - k_1 \sqrt{x_1(t)} \\ \frac{dx_2(t)}{dt} = k_1 \sqrt{x_1(t)} - k_2 \sqrt{x_2(t)} \end{cases} \quad (8.5.9)$$

Les cuves ont des niveaux minimums et maximums autorisés, notés $H_{\min,1}$, $H_{\min,2}$, $H_{\max,1}$, $H_{\max,2}$. Les niveaux doivent rester dans ces bornes autorisées ; si ce n'est pas le cas, le système se met dans un cinquième état discret labélisé Défaut. Afin d'essayer d'éviter une mise en défaut, on définit une tolérance ϵ . Si le niveau dans la cuve 1 atteint $H_{\max,1} - \epsilon$, il est interdit d'ouvrir la vanne 1, s'il atteint $H_{\min,1} + \epsilon$, il est interdit d'ouvrir la vanne 2 ou de fermer la vanne 1. Si le niveau dans la cuve 2 atteint $H_{\max,2} - \epsilon$, il est interdit d'ouvrir la vanne 2 ; s'il atteint $H_{\min,2} + \epsilon$, il est interdit de fermer la vanne 2.

La représentation sous forme d'automate hybride est alors celle donnée sur la figure 8.11. On suppose pour simplifier qu'une seule action peut être faite au même instant. Pour alléger la figure on n'a pas reproduit les équations d'état continues à l'intérieur des états discrets. De la même façon, seule la garde est indiquée sur les transitions. Dans le système considéré, il n'y a pas de fonction de reset des états lors des transitions. Les transitions vers l'état Défaut ne sont pas indiquées sur la figure. Les gardes vers cet état sont toutes identiques et valent :

$$(x_1 = H_{\min,1}) \vee (x_1 = H_{\max,1}) \vee (x_2 = H_{\min,2}) \vee (x_2 = H_{\max,2}) \quad (8.5.10)$$

8.6 Réseaux de Petri temporisés

Nous avons définis les réseaux de Petri non temporisés dans le chapitre 7. Ces réseaux s'étendent très naturellement au cas des systèmes temporisés. Cette catégorie de réseaux a une grande importance dans l'analyse et la validation de systèmes temps réel. Nous en donnons un bref aperçu ici.

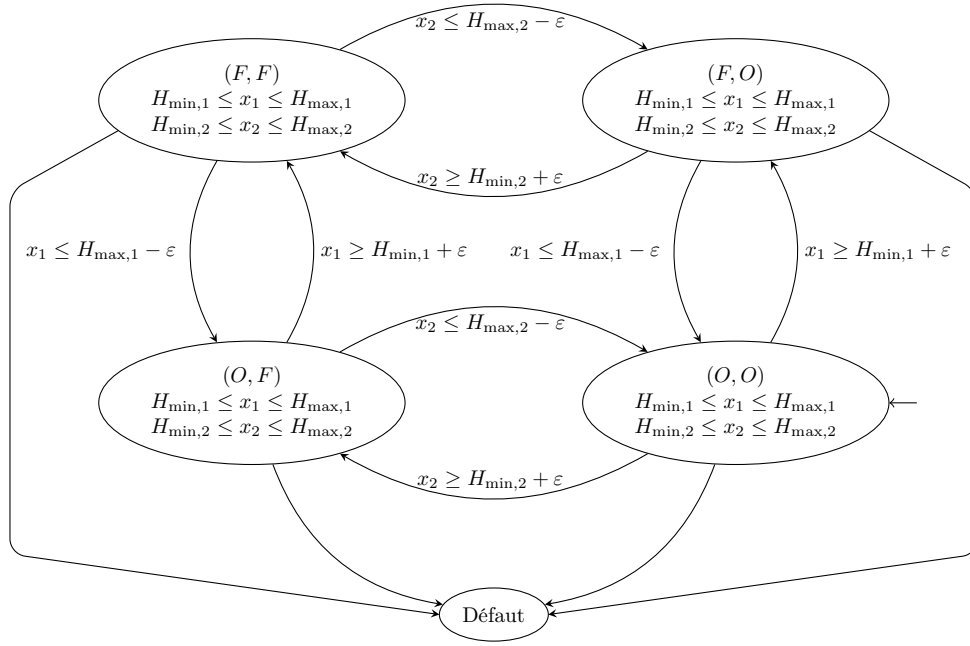


FIGURE 8.11 – Modélisation sous forme d'automate hybride du système de deux cuves

8.6.1 Définition

Les *réseaux de Petri p-temporisés* sont des réseaux autonomes pour lesquels une temporisation fixe (éventuellement nulle) T_i est associée à chaque place P_i .

Lorsqu'une place devient marquée, la marque ajoutée reste indisponible pendant la durée t_i de la temporisation. Le marquage global $m(P)$ est donc constitué de la somme des marques non disponibles $m_n(P)$ et des marques disponibles $m_d(P)$:

$$m(P) = m_d(P) + m_n(P)$$

Les règles d'évolution habituelles des réseaux de Petri sont appliquées au marquage disponible.

On considère que le temps de franchissement des transitions est nul. Le marquage initial est supposé être constitué de marques disponibles.

Le fonctionnement d'un réseau temporisé comporte en général une phase transitoire suivie d'une phase de fonctionnement stationnaire. Il est possible dans cette seconde phase de caractériser les modes de fonctionnement du réseau : vitesse maximale, fréquence de franchissement des transitions, etc.

Remarque 8.6.1. Une autre variante est constituée des réseaux de Petri *t-temporisés* qui associent des temporisations aux franchissements des transitions.

Lorsqu'une transition est franchie, les marques des places d'entrée sont réservées pendant la durée de la temporisation. Après écoulement de celle-ci, elles sont rendues disponibles dans les places de sortie.

Les réseaux *p-temporisés* et les réseaux *t-temporisés* sont équivalents.

Chapitre 9

Automates stochastiques

9.1 Introduction

Nous avons vu dans les chapitres précédents des techniques permettant de modéliser des systèmes à événements discrets. Dans tout ce qui a été présenté, les phénomènes ont été considérés comme parfaitement connus à l'avance. Il apparaît que cette vision déterministe du monde est incomplète pour analyser le fonctionnement de certains systèmes. Ainsi, si nous reprenons l'exemple classique de la file d'attente, les modélisations précédentes considéraient que les dates d'arrivées des clients et les temps de traitement étaient parfaitement connus. Ces hypothèses ne sont que peu réalistes au regard du fonctionnement d'un tel système. Afin de caractériser le comportement des systèmes à événements discrets, il est donc nécessaire de caractériser le caractère aléatoire de l'arrivée des événements. Nous commençons ici par définir les structures d'horloges stochastiques à la section 9.2, puis les automates temporisés stochastiques à la section 9.3. Nous détaillerons plus spécifiquement les automates dont les horloges sont caractérisées par des processus de Poisson à la section 9.4. Nous finirons enfin par donner brièvement quelques autres exemples de types de processus entrant dans la modélisation des systèmes à événements discrets stochastiques à la section 9.5.

9.2 Structure d'horloge stochastique

Nous avons défini au chapitre 8 les structures d'horloges permettant d'obtenir des automates temporisés. Il est nécessaire ici de les étendre au cas stochastique en considérant que les séquences d'horloges sont des grandeurs aléatoires caractérisées par des densités de probabilités. Ainsi, une structure d'horloge stochastique est donnée par la définition suivante :

Définition 9.2.1

Une structure d'horloge stochastique associée à un ensemble d'événements Σ est la donnée d'un ensemble de fonctions de distribution :

$$G = \{G_\sigma, \sigma \in \Sigma\} \quad (9.2.1)$$

caractérisant des séquences stochastiques d'horloges :

$$\{V_{\sigma,k}\} = \{V_{\sigma,1}, V_{\sigma,2}, \dots\}, \text{ avec } V_{\sigma,k} \in \mathbb{R}_+ \quad (9.2.2)$$

c'est-à-dire que $V_{\sigma,k}$ est une réalisation d'une variable aléatoire ayant pour fonction de distribution G_σ .

Nous nous limiterons ici au cas où les séquences d'horloges sont constituées de variables indépendantes et identiquement distribuées. Ainsi, chaque valeur de la séquence d'horloges $\{V_{\sigma,k}\}$ est complètement caractérisée par la fonction de répartition $G_\sigma(t) = \mathbb{P}[V_\sigma \leq t]$ qui ne dépend ni de k ni des valeurs précédentes de la séquence d'horloge.

9.3 Automate temporisé stochastique

Nous présentons ici la définition d'un automate temporisé stochastique. Nous ne considérons que la version *discrète*, l'extension au cas de systèmes hybrides étant faisable.

Définition 9.3.1

Un automate temporisé stochastique est caractérisé par un sextuplet $\mathcal{A} = (Q, \Sigma, \phi, p_{tr}, p_0, G)$ avec :

- Q un ensemble dénombrable d'états ;
- Σ un ensemble dénombrable d'événements ;
- ϕ un ensemble de transitions possibles, avec $\phi \subseteq Q \times \Sigma \times Q$;
- p_{tr} une probabilité de transition ;
- p_0 une fonction de répartition pour l'état initial ;
- G une structure d'horloge stochastique.

Dans cette définition, $(q_1, \sigma, q_2) \in \phi$ ne signifie pas que la transition de q_1 à q_2 sera effectuée si l'événement σ se produit alors que le système est dans l'état q_1 . Cela signifie que la transition correspondante est possible, et qu'elle sera caractérisée par une variable aléatoire de loi de probabilité p_{tr} . Bien entendu, $p_{tr}(q_1, \sigma, q_2) = 0$ si $(q_1, \sigma, q_2) \notin \phi$.

L'automate génère ainsi une suite aléatoire d'états $(Q_0, Q_1, \dots)$ de la façon suivante :

- L'état initial Q_0 est tirée selon la loi p_0 .
- Lorsque l'automate est dans l'état Q_k , les événements actifs (c'est-à-dire ceux pour lesquels une transition depuis Q_k est possible si cet événement se produit) possèdent une horloge. Celle-ci est tirée aléatoirement selon la structure d'horloge stochastique si l'événement est redevenu actif en Q_k . Si l'événement était déjà actif avant d'arriver en Q_k , l'horloge continue de décompter le temps.
- L'horloge qui la première arrive à 0 entraîne le déclenchement de l'événement correspondant.
- Le système opère alors une transition vers l'état Q_{k+1} obtenu en le tirant aléatoirement parmi les différents états suivants possibles, avec les probabilités p_{tr} .
- Lorsque le système arrive dans l'état Q_{k+1} , les scores des événements sont remis à jour, et des horloges sont de nouveau tirées pour les états nouvellement actifs.

La description proposée ici est très générale, puisque les transitions et les horloges ont un caractère aléatoire. Il est bien entendu possible d'avoir des modélisations de systèmes dans lesquelles les transitions gardent leur caractère déterministe, et les horloges gardent une structure stochastique. Les développements présentés ici permettent également d'aborder les modélisations de système sous la forme de chaînes de Markov, mais cela dépasse le cadre de ce cours.

Une classe particulière d'automates stochastiques est constituée par les automates dont l'arrivée des événements est donnée par une loi de probabilité correspondant à un processus de Poisson. Nous détaillons ce processus dans la section suivante.

9.4 Processus de Poisson

9.4.1 Définition de la densité de probabilité

Nous considérons un système à événements discrets dans lequel un seul événement peut se produire et nous notons $N(t)$ le nombre de fois où cet événement s'est produit sur l'intervalle de temps $]0, t]$ (nous supposons $N(0) = 0$). De la même façon, nous notons $N(t_{k-1}, t_k)$ le nombre d'événements survenus durant l'intervalle de temps $]t_{k-1}, t_k]$. De manière évidente, nous avons donc :

$$\begin{aligned} 0 \leq N(t_1) \leq N(t_2) \leq \dots \leq N(t_{k-1}) \leq N(t_k) \\ N(t_{k-1}, t_k) = N(t_k) - N(t_{k-1}) \end{aligned} \tag{9.4.1}$$

Nous faisons les hypothèses suivantes :

- **Hypothèse 1** : deux évènements ne peuvent se produire exactement au même instant.
- **Hypothèse 2** : les variables aléatoires $N(t), N(t, t_1), N(t_1, t_2), \dots, N(t_{k-1}, t_k)$ sont mutuellement indépendantes.
- **Hypothèse 3** : $\mathbb{P}[N(t_{k-1}, t_k) = n]$, avec $n = 0, 1, \dots$ peut dépendre de la longueur de l'intervalle, c'est-à-dire de $t_k - t_{k-1}$, mais pas des instants t_k et t_{k-1} .

Définition 9.4.1

Un processus $\{N(t)\}$ vérifiant les deux premières hypothèses est un *processus à incréments indépendants*. Si l'hypothèse 3 est également satisfaite, le processus est à *incrémentés indépendants et stationnaires*.

Du fait de l'hypothèse 3, nous avons :

$$\mathbb{P}[N(t_{k-1}, t_k) = n] = \mathbb{P}[N(s) = n], \text{ avec } s = t_k - t_{k-1} \quad (9.4.2)$$

Il est donc suffisant pour caractériser ce processus de déterminer les probabilités $\mathbb{P}[N(t) = n]$. Considérons tout d'abord deux instants t et $t + \tau$. Nous avons :

$$\mathbb{P}[N(t + \tau) = 0] = \mathbb{P}[N(t) = 0 \cap N(t, t + \tau) = 0] \quad (9.4.3)$$

Du fait de l'hypothèse 2 :

$$\mathbb{P}[N(t + \tau) = 0] = \mathbb{P}[N(t) = 0] \cdot \mathbb{P}[N(t, t + \tau) = 0] = \mathbb{P}[N(t) = 0] \cdot \mathbb{P}[N(\tau) = 0] \quad (9.4.4)$$

En notant la probabilité $\mathbb{P}_0(t) = \mathbb{P}[N(t) = 0]$, cette probabilité obéit donc à une équation de la forme :

$$\mathbb{P}_0(t + \tau) = \mathbb{P}_0(t) \cdot \mathbb{P}_0(\tau) \quad (9.4.5)$$

On peut montrer que les seules fonctions solutions d'une telle équation sont les fonctions exponentielles. On obtient donc :

$$\mathbb{P}_0(t) = e^{-\lambda t} \quad (9.4.6)$$

En calculant les probabilités successives (la démonstration n'est pas donnée ici mais peut être facilement trouvée dans la littérature) :

$$\mathbb{P}[N(t) = n] = \mathbb{P}_n(t) = \frac{(\lambda t)^n}{n!} e^{-\lambda t} \quad (9.4.7)$$

Le processus $\{N(t)\}$ caractérisé par cette loi de probabilité est un appelé *Processus de Poisson*. On parle également de *distribution exponentielle*.

9.4.2 Moyenne et variance du processus de Poisson

Nous pouvons calculer analytiquement la moyenne et la variance du processus de Poisson :

$$\begin{aligned} \mathbb{E}[N(t)] &= \sum_{n=0}^{+\infty} n \mathbb{P}_n(t) = \sum_{n=0}^{+\infty} n \frac{(\lambda t)^n}{n!} e^{-\lambda t} = \sum_{n=1}^{+\infty} n \frac{(\lambda t)^n}{n!} e^{-\lambda t} \\ &= \sum_{n=1}^{+\infty} \frac{(\lambda t)^{n-1}}{(n-1)!} e^{-\lambda t} \lambda t = \lambda t e^{-\lambda t} \sum_{n=0}^{+\infty} \frac{(\lambda t)^n}{n!} \\ &= \lambda t \end{aligned} \quad (9.4.8)$$

Le scalaire λ est souvent appelé l'*intensité* ou le *taux du processus de Poisson*, car $\lambda = \mathbb{E}[N(t)]/t$ correspond au taux moyen d'apparition des évènements par unité de temps.

La variance se calcule de manière similaire :

$$\begin{aligned}
\text{Var}[N(t)] &= \mathbb{E}[N^2(t)] - \mathbb{E}[N(t)]^2 = \sum_{n=0}^{+\infty} n^2 \mathbb{P}_n(t) - (\lambda t)^2 = e^{-\lambda t} \sum_{n=0}^{+\infty} n^2 \frac{(\lambda t)^n}{n!} - (\lambda t)^2 \\
&= e^{-\lambda t} \sum_{n=1}^{+\infty} n^2 \frac{(\lambda t)^n}{n!} - (\lambda t)^2 = e^{-\lambda t} \sum_{n=1}^{+\infty} (n-1+1) \frac{(\lambda t)^n}{(n-1)!} - (\lambda t)^2 \\
&= e^{-\lambda t} \sum_{n=1}^{+\infty} (n-1) \frac{(\lambda t)^n}{(n-1)!} + e^{-\lambda t} \sum_{n=1}^{+\infty} \frac{(\lambda t)^n}{(n-1)!} - (\lambda t)^2 \\
&= e^{-\lambda t} \sum_{n=2}^{+\infty} (n-2) \frac{(\lambda t)^n}{(n-2)!} + e^{-\lambda t} \sum_{n=1}^{+\infty} \frac{(\lambda t)^{n-1}}{(n-1)!} \lambda t - (\lambda t)^2 \\
&= (\lambda t)^2 e^{-\lambda t} \sum_{n=0}^{+\infty} \frac{(\lambda t)^n}{(n)!} + (\lambda t) e^{-\lambda t} \sum_{n=0}^{+\infty} \frac{(\lambda t)^n}{(n)!} - (\lambda t)^2 \\
&= (\lambda t)^2 + (\lambda t) - (\lambda t)^2 = \lambda t
\end{aligned} \tag{9.4.9}$$

9.4.3 Propriétés du processus de Poisson

9.4.3.1 Temps inter-événements

Nous pouvons tout d'abord évaluer la loi de probabilité du temps inter-événement. Si nous notons t_k l'instant d'apparition de l'évènement k , le temps inter-événement est défini par :

$$V_k = t_k - t_{k-1} \tag{9.4.10}$$

Ceci est une variable aléatoire (nous avons pris la notation V_k car ce temps inter-événement est proche d'une horloge stochastique). La fonction de répartition de cette variable aléatoire est donnée par :

$$\mathbb{P}[V_k \leq t] = 1 - \mathbb{P}[V_k > t] = 1 - \mathbb{P}[N(t_{k-1}, t_{k-1} + t) = 0] = 1 - e^{-\lambda t} \tag{9.4.11}$$

Nous remarquons que ce temps inter-événement est indépendant de k . La densité de probabilité correspondante du temps inter-événement est obtenue en différenciant cette équation et vaut donc $\lambda e^{-\lambda t}$.

9.4.3.2 Processus sans mémoire

Supposons qu'aucun évènement ne se soit produit depuis un temps θ , nous cherchons à déterminer la loi de probabilité du prochain évènement. Cette propriété s'écrit ainsi :

$$\mathbb{P}[V_k \leq \theta + t | V_k > \theta] = \frac{\mathbb{P}[V_k \leq \theta + t \cap V_k > \theta]}{\mathbb{P}[V_k > \theta]} = \frac{\mathbb{P}[V_k \leq \theta + t \cap V_k > \theta]}{1 - \mathbb{P}[V_k \leq \theta]} \tag{9.4.12}$$

En utilisant la propriété de l'équation 9.4.11, nous obtenons ainsi :

$$\mathbb{P}[V_k \leq \theta + t | V_k > \theta] = \frac{(1 - e^{-\lambda(\theta+t)}) - (1 - e^{-\lambda\theta})}{e^{-\lambda\theta}} = 1 - e^{-\lambda t} \tag{9.4.13}$$

Ceci montre que quelque soit le temps passé depuis l'apparition d'un évènement dans le processus de Poisson, la fonction de répartition de la prochaine occurrence est toujours donné par $1 - e^{-\lambda t}$, ce qui justifie le terme de *processus sans mémoire*.

9.4.3.3 Superposition de processus de Poisson

Nous considérons ici que m évènements peuvent désormais se produire, chacun correspondant à un processus de Poisson de paramètre λ_i . Si ces processus sont indépendants, nous pouvons de nouveau évaluer le temps inter-événements, noté V_k , en se rappelant qu'un évènement se produit lorsque la première horloge arrive à 0.

Ainsi, l'expression suivante est obtenue :

$$\mathbb{P}[V_k \leq t] = 1 - \mathbb{P}[V_k > t] = 1 - \mathbb{P}\left[\min_{i=1, \dots, m} V_{i,k} > t\right] = 1 - \prod_{i=1}^m \mathbb{P}[V_{i,k} > t] = 1 - \prod_{i=1}^m e^{-\lambda_i t} \quad (9.4.14)$$

En notant $\tilde{\lambda} = \sum_{i=1}^m \lambda_i$, nous obtenons :

$$\mathbb{P}[V_k \leq t] = 1 - e^{-\tilde{\lambda}t} \quad (9.4.15)$$

ce qui montre que la superposition de m processus de Poisson est également un processus de Poisson de paramètre $\tilde{\lambda}$.

Ce résultat est particulièrement d'importance pour les automates dont les structures d'horloge sont des processus de Poisson.

9.4.4 Intérêt du processus de Poisson

Le processus de Poisson est extrêmement utilisé dans la modélisation des systèmes à événements discrets. Outre son intérêt théorique illustré par les différentes propriétés que nous avons définies précédemment, de nombreux processus physiques ont, naturellement, des distributions de Poisson (physique des particules, problèmes de trafic, de réseaux de communication, etc.). En effet, ce type de processus permet de compter des événements arrivant de façon aléatoire, mais de façon invariante par rapport au temps.

9.5 Autres types de processus

De nombreux autres types de processus aléatoires entrent dans la représentation des arrivées d'événements dans le cadre des systèmes à événements discrets. Nous en donnons quelques exemples ici.

9.5.1 Loi d'Erlang ou processus de renouvellement

La *densité de probabilité de la loi d'Erlang* est donnée par :

$$p(x) = \frac{\lambda^k x^{k-1} e^{-\lambda x}}{(k-1)!} \quad (9.5.1)$$

Elle dépend de deux paramètres : un scalaire entier k qui est appelé *facteur de forme* et un scalaire réel λ qui est l'*intensité*. Sa fonction de répartition est donnée par :

$$\mathbb{P}[X \leq x] = 1 - \sum_{n=0}^{k-1} e^{-\lambda x} \frac{(\lambda x)^n}{n!} \quad (9.5.2)$$

On peut interpréter cette loi comme la distribution de la somme de k variables aléatoires, indépendantes et identiquement distribuées, selon une loi exponentielle de paramètre λ .

9.5.2 Distribution de Rayleigh

La *densité de probabilité d'une loi de Rayleigh* est donnée par :

$$p(x) = \frac{x}{\sigma^2} \exp\left(\frac{-x^2}{2\sigma^2}\right) \quad (9.5.3)$$

Si U est une loi uniforme sur $]0; 1[$, alors $X = \sigma \sqrt{-2 \ln(U)}$ suit une loi de Rayleigh. Si on note D_n la distance à l'origine dans le cas d'une marche aléatoire dans le plan après n pas au hasard, on montre que $D_n/\sqrt{n}$ converge en loi vers la loi de Rayleigh.

9.6 Conclusion de la partie II

Nous avons vu dans cette partie différents formalismes pour représenter des systèmes à états discrets, déterministes ou stochastiques, temporisés ou non temporisés.

En particulier, les réseaux de Pétri constituent des outils génériques applicables dans ce cadre, avec des propriétés que l'on peut déduire de leur étude par réduction ou analyse linéaire. Leurs intérêts sont un formalisme rigoureux et compact, une sémantique précise et une représentation visuelle graphique. Ils offrent la possibilité de traiter facilement des processus se déroulant en parallèle ou selon les concepts de coopération, compétition (partage de ressources) ou synchronisation. De nombreux outils industriels s'appuient sur leur formalisme pour de nombreuses applications (ingénieur de production, modélisation de services, tâches d'un système multi-processeur, systèmes de planification de production, réseaux de transport, etc.). Il existe beaucoup d'extensions : réseaux de Pétri colorés (avec des jetons de couleurs pour représenter des processus parallèles), généralisés (fonction de valuation ou poids sur les arcs), à arcs inhibiteurs, à priorités, etc. Leur étude dépasse le cadre de ce cours.

Troisième partie

Méthodes pour l'identification et l'évaluation des modèles

Chapitre 10

Identification

10.1 Introduction

10.1.1 Démarche générale pour l'identification

Définition 10.1.1

Le problème de l'*identification* consiste à déterminer un modèle mathématique d'un système réel à partir d'observations, en s'appuyant éventuellement sur des connaissances a priori.

Il s'agit donc de déterminer ou d'ajuster à la fois la structure du modèle et ses paramètres, afin de reproduire au mieux le comportement du système.

La démarche classique pour identifier un modèle consiste à :

- *Etape 1* : formaliser les connaissances disponibles a priori ;
- *Etape 2* : recueillir des données expérimentales ;
- *Etape 3* : identifier la structure, estimer les paramètres (et les incertitudes associées) ;
- *Etape 4* : évaluer le modèle ainsi obtenu en analysant ses propriétés et en le confrontant à des données n'ayant pas servi dans les étapes précédentes, par exemple des données acquises en soumettant le système et le modèle à des entrées nouvelles (on parle de *validation croisée*).

Dans les cas de systèmes complexes, cette procédure s'opère généralement de façon cyclique, les résultats d'une étape pouvant conduire à une remise en cause du modèle. Nous pouvons par exemple obtenir en étape 3 des problèmes d'identifiabilité pratique ou structurelle qui rendent impossible l'estimation des paramètres (voir partie 10.4.2). Ou bien, en étape 4, il est possible de s'apercevoir que les résidus, c'est-à-dire l'écart entre les sorties du système réel et du modèle, contiennent encore de l'information explicable (corrélation avec les signaux d'entrée), ce qui peut amener à proposer une nouvelle structure de modèle et à refaire l'estimation paramétrique.

La procédure d'identification se termine lorsque le modèle répond à l'objectif pour lequel il a été établi.

10.1.2 Classification sommaire des méthodes

Selon les systèmes considérés et les domaines d'applications, certaines des étapes citées ci-dessus peuvent ne pas être réalisables ou sans intérêt. Par exemple, lorsque l'on considère un modèle de connaissance, la structure du modèle peut être considérée comme donnée et l'identification a alors seulement pour objectif de déterminer la valeur des paramètres (ou d'une partie d'entre eux) qui interviennent dans ces équations. Dans d'autres cas, au contraire, on se place dans une démarche empirique et l'identification a pour objectif de déterminer un modèle mathématique, simple, qui se comporte comme une bonne approximation du système. Enfin, la phase d'évaluation peut être très restreinte si l'on ne dispose que de peu de données expérimentales ou pour un système n'ayant pas d'entrées manipulables par un expérimentateur. On peut penser par exemple à un modèle de croissance d'une tumeur cancéreuse.

Ces caractéristiques des modèles et des protocoles expérimentaux applicables à un système déterminent le choix des méthodes d'identification. Certaines s'appliqueront plus spécifiquement à la détermination d'un modèle à temps continu, d'autres à la détermination d'un modèle à temps discret. Nous en présentons dans ce cours les plus usuelles, qu'on peut répartir en deux catégories :

- les méthodes graphiques qui s'appuient sur les réponses à des signaux usuels (échelon, impulsion, sinusoïde) : on présente ici les principes de l'analyse indicielle en paragraphe 10.2 et de l'analyse harmonique en paragraphe 10.3.
- les méthodes statistiques, plus générales (paragraphe 10.4.2) ; elles s'expriment sous la forme d'un problème d'optimisation dont la résolution peut se révéler plus ou moins compliquée et faire parfois appel à un algorithme récursif (paragraphe 10.4.4).

Exemple 10.1.1. La figure 10.1 illustre la procédure d'identification pour un système admettant des entrées : on soumet le système à un signal d'entrée test (ou à plusieurs signaux-tests consécutivement) ; on enregistre la sortie du système ; on en déduit un modèle mathématique (structure et/ou paramètres et/ou incertitudes) ; on valide le modèle en comparant ses réponses et celles du système pour différents signaux d'entrée (dernière étape non représentée sur le schéma).

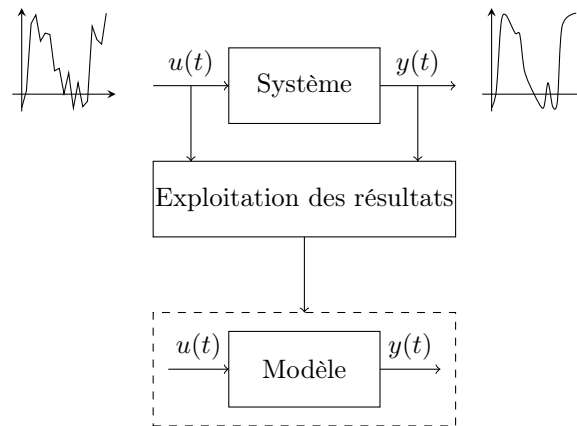


FIGURE 10.1 – Identification

Le choix du signal-test appliqué lors de la procédure expérimentale est évidemment important dans ce cas, et conditionne la qualité du résultat de l'identification. Par exemple, on pourra notamment envisager les signaux suivants :

- impulsion, signal intéressant puisqu'il donne accès à la réponse impulsionnelle, qui caractérise complètement un système linéaire ; mais ce signal est impossible à réaliser physiquement à temps continu ;
- échelon, plus facile à réaliser expérimentalement ;
- sinusoïde.

Remarque 10.1.1. Le vocabulaire employé pour l'identification est remarquablement divers selon les domaines d'applications : on rencontrera ainsi, avec parfois des différences subtiles d'interprétation, les termes *identification*, *inférence*, *inversion du modèle* ou *problème inverse*, *estimation*, *ajustement*, *calibration*, *'fitting'*, *calage*, etc.

10.2 Analyse indicielle

10.2.1 Procédure expérimentale

Le principe de cette approche consiste à déduire un modèle du système physique à partir du relevé de la réponse obtenue lorsqu'on applique en entrée un échelon. Ce signal est l'un des plus facile à mettre en œuvre. La procédure expérimentale consiste, à un instant t_0 où le système se trouve en régime permanent, à changer instantanément la valeur de l'entrée (figure 10.2). On

enregistre la sortie, et on en déduit la réponse indicielle (réponse à un échelon de Heaviside, valant 0 pour $t < 0$ et 1 pour $t \geq 0$) par recalage par rapport à 0, et proportionnalité.

Remarque 10.2.1. *Si l'hypothèse que le système est linéaire et invariant dans le temps est valide, la réponse indicielle obtenue sera la même quelles que soient les valeurs de t_0 , E_1 et E_2 . En pratique, il convient de faire des essais pour différentes valeurs, pour vérifier cette hypothèse ou l'infirmier (il arrive souvent par exemple que l'hypothèse de linéarité ne soit plus vérifiée pour des valeurs trop faibles, ou au contraire trop fortes de la différence $E_2 - E_1$).*

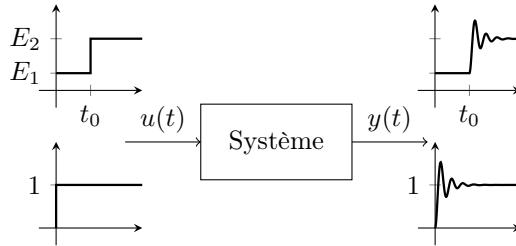


FIGURE 10.2 – Analyse indicielle

L'intérêt théorique de cette procédure est que la transformée de Laplace de la réponse indicielle est $\frac{H(p)}{p}$, où $H(p)$ est la fonction de transfert du système. Par 'multiplication par p ', c'est-à-dire dans le domaine temporel par dérivation, on en déduit théoriquement $H(p)$ (ou encore la réponse impulsionnelle dans le domaine temporel). Cette opération est bien sûr impossible à réaliser en pratique, mais on peut en retenir que la réponse indicielle présente théoriquement toute l'information nécessaire à la détermination de la fonction de transfert.

Ayant obtenu la réponse indicielle, un premier examen visuel permet de déduire quelques caractéristiques de la fonction de transfert. Ainsi par exemple :

- si la tangente à l'origine est non nulle, et la réponse monotone sans discontinuité en $t = 0$, un modèle de fonction de transfert peut être du premier ordre sans zéro (paragraphe 5.5.1.1) ;
- si la tangente à l'origine est nulle, un modèle de fonction de transfert peut être d'ordre deux au moins ;
- si la tangente à l'origine est nulle et que la réponse présente des oscillations, un modèle de fonction de transfert peut être du second ordre sans zéro avec $\zeta < 1$ (paragraphe 5.5.1.2) ;
- si la réponse reste nulle pour tout $0 < t < \tau$, la fonction de transfert présente un retard pur, elle est donc de la forme $e^{-\tau p} H(p)$ (paragraphe 5.4.3).

Le problème, ayant choisi a priori la forme de la fonction de transfert, consiste ensuite à déterminer ses paramètres. Nous présentons ci-dessous quelques cas typiques.

10.2.2 Identification par une fonction de transfert du premier ordre

Dans le cas où la réponse indicielle est monotone avec une tangente non nulle en $t = 0$, on peut rechercher un modèle simple sous la forme d'une fonction de transfert du premier ordre :

$$H(p) = \frac{K}{1 + Tp} = K \frac{1/T}{p + 1/T} \quad (10.2.1)$$

La réponse indicielle correspondante (paragraphe 5.5.1.1) s'écrit :

$$y(t) = K \left(1 - \exp\left(-\frac{t}{T}\right) \right) \mathbb{1}(t) \quad (10.2.2)$$

avec $\mathbb{1}(t)$ l'échelon unitaire. La tangente à l'origine a pour équation $y(t) = \frac{K}{T}t$. On en déduit les propriétés suivantes (figure 10.3) :

- la réponse tend asymptotiquement vers la droite d'équation $y = K$;
- pour $t = T$, la sortie vaut $y(T) \approx 0,63K$;
- pour $t = 3T$, la sortie vaut $y(3T) \approx 0,95K$;
- pour $t = T$, la tangente à l'origine coupe l'asymptote horizontale.

Ces caractéristiques permettent aisément de déterminer le gain K et la constante de temps T . En pratique, la réponse obtenue peut être bruitée (ou ne pas correspondre exactement à la réponse d'une fonction de transfert du 1^{er} ordre), et il est donc conseillé, notamment pour déterminer T , de se baser à la fois sur la tangente et sur plusieurs points de la réponse, et de prendre par exemple la moyenne des valeurs obtenues.

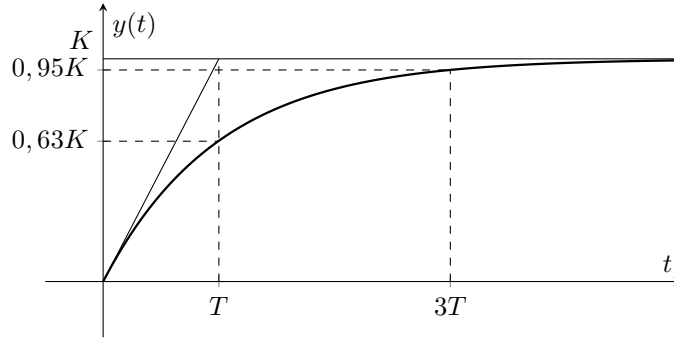


FIGURE 10.3 – Réponse indicielle (fonction de transfert d'ordre 1 sans zéro)

10.2.3 Identification par la méthode de Strejč

La méthode proposée par Vladimir Strejč s'applique à toute réponse indicielle continue en $t = 0$, monotone et présentant une asymptote horizontale. Le modèle recherché est une fonction de transfert à temps continu de la forme suivante :

$$H(p) = \frac{K e^{-\tau p}}{(1 + Tp)^n} \quad (10.2.3)$$

Il est à noter que l'on ne suppose pas que le système correspond exactement à une fonction de transfert de cette forme : on cherche des paramètres K , T , n et τ pour lesquelles la réponse indicielle déduite de cette fonction de transfert se rapproche au mieux de celle du système physique. L'expérience montre que ce modèle est largement justifié pour de nombreux systèmes, même si ceux-ci correspondent en fait à une fonction de transfert avec des constantes de temps très différentes.

10.2.3.1 Principe de la méthode

La réponse indicielle déduite de la fonction de transfert cherchée est :

$$y(t) = \mathcal{L}^{-1} \left\{ \frac{H(p)}{p} \right\} (t)$$

Supposons tout d'abord que le système ne présente pas de retard pur, soit $\tau = 0$. On calcule alors, en utilisant les propriétés de la transformée de Laplace :

$$\begin{aligned} \frac{dy(t)}{dt} &= \mathcal{L}^{-1} \left\{ p \cdot \frac{H(p)}{p} \right\} (t) = \mathcal{L}^{-1} \left\{ \frac{K}{(1 + Tp)^n} \right\} (t) = K e^{-\frac{t}{T}} \frac{t^{n-1}}{(n-1)!T^n} \mathbb{1}(t) \\ y(t) &= \int_0^t K e^{-\frac{\theta}{T}} \frac{\theta^{n-1}}{(n-1)!T^n} d\theta \\ &= K \left(1 - e^{-\frac{t}{T}} \left(1 + \frac{t}{T} + \frac{t^2}{2!T^2} + \cdots + \frac{t^{n-1}}{(n-1)!T^{n-1}} \right) \right) \mathbb{1}(t) \quad (10.2.4) \end{aligned}$$

La position du point d'inflexion I de cette réponse (figure 10.4, en trait plein) est caractéristique de l'ordre n de la fonction de transfert. En effet $\frac{d^2y(t)}{dt^2}$ s'annule pour $t_I = (n - 1)T$.

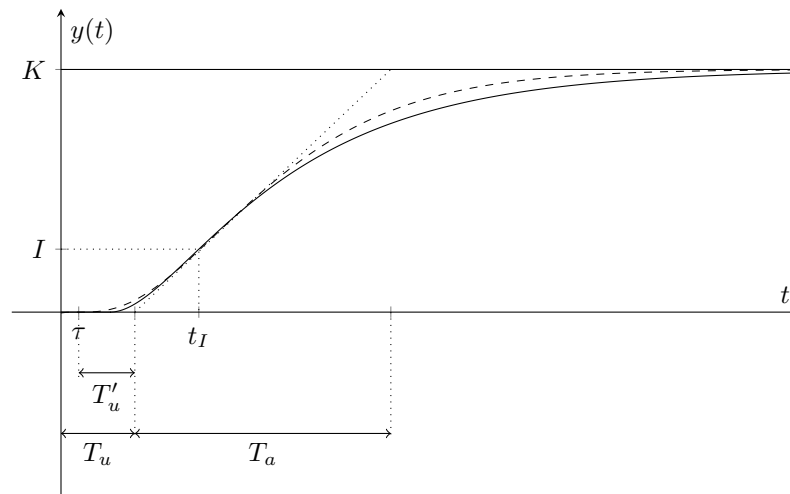


FIGURE 10.4 – Réponse indicielle (modèle de Strejč)

A partir de l'équation de la tangente au point d'inflexion :

$$y(t) = \left(\frac{dy(t)}{dt} \right)_{t=t_I} (t - t_I) + y(t_I) \quad (10.2.5)$$

on détermine les 2 intervalles de temps T_u et T_a représentés sur la figure 10.4. Les valeurs obtenues suivant l'ordre n sont reportées dans le tableau 10.1.

n	$\frac{T_a}{T}$	$\frac{T_u}{T}$	$\frac{T_u}{T_a}$
1	1	0	0
2	2,72	0,28	0,10
3	3,70	0,80	0,22
4	4,46	1,42	0,32
5	5,12	2,10	0,41
6	5,70	2,81	0,49
7	6,23	3,55	0,57

TABLEAU 10.1 – Caractéristiques du modèle de Strejč

10.2.3.2 Mise en œuvre

La valeur de K s'identifie directement à partir de la valeur asymptotique de la réponse indicielle. Le tracé de la tangente au point d'inflexion permet de déterminer le rapport $\frac{T_u}{T_a}$. Si celui-ci correspond à une valeur figurant dans le tableau 10.1, on en déduit immédiatement les valeurs de n et T .

Il est fréquent néanmoins que le rapport $\frac{T_u}{T_a}$ mesuré ne soit pas égal à l'une des valeurs du tableau, d'abord parce que la détermination pratique du point d'inflexion et le tracé de la tangente sont assez approximatifs, ensuite parce que la fonction de transfert du système ne correspond pas forcément à l'expression recherchée. Dans ce cas, on prend pour l'ordre n la valeur immédiatement inférieure, et on compense l'erreur commise par l'introduction d'un retard pur fictif τ .

En effet, la présence de ce retard ne modifie pas la valeur de T_a , mais change T_u en $T'_u = T_u - \tau$ (réponse en pointillés sur la figure 10.4). Soit donc $\frac{T'_u}{T_a}$ la valeur dans le tableau qui est immédiatement inférieure à $\frac{T_u}{T_a}$. Le tableau donne n et T , puis on calcule successivement T'_u et $\tau = T_u - T'_u$.

Exemple 10.2.1. Si $T_a = 10$ et $T_u = 3$, alors $\frac{T_u}{T_a} = 0,3$ d'où $\frac{T'_u}{T_a} = 0,22$. On en déduit successivement $n = 3$, $T = \frac{T_a}{3,7} = 2,7$, $T'_u = 0,22T_a = 2,2$ et $\tau = 3 - 2,20 = 0,8$.

Remarque 10.2.2. Lorsque la réponse indicielle présente un retard pur naturel θ , il convient d'ajouter τ à θ . La mesure de θ s'effectue par l'écart séparant l'instant d'application de l'échelon de l'instant où la réponse commence à évoluer, et c'est ce dernier instant qui est pris comme origine pour la mesure de T_u .

10.2.4 Identification par une fonction de transfert du second ordre

Dans le cas d'une réponse avec oscillations et tangente nulle à l'origine, on peut tenter de rechercher un modèle sous la forme d'une fonction de transfert du second ordre sans zéro (paragraphe 5.5.1.2), avec pôles complexes conjugués :

$$H(p) = \frac{K}{1 + 2\zeta \frac{p}{\omega_0} + \left(\frac{p}{\omega_0}\right)^2} = K \frac{\omega_0^2}{p^2 + 2\zeta\omega_0 p + \omega_0^2}, \text{ avec } |\zeta| < 1 \quad (10.2.6)$$

La détermination des trois paramètres est particulièrement simple (figure 10.5) : l'asymptote horizontale donne la valeur de K , et les coordonnées du premier maximum fournissent D et t_m , d'où on déduit ζ et ω_0 , respectivement :

$$\begin{cases} D = \exp\left(-\frac{\pi\zeta}{\sqrt{1-\zeta^2}}\right) \Leftrightarrow \zeta = \frac{1}{\sqrt{1 + \left(\frac{\pi}{\ln D}\right)^2}} \\ \omega_0 = \frac{\pi}{t_m \sqrt{1-\zeta^2}} \end{cases} \quad (10.2.7)$$

Remarque 10.2.3. L'approximation $\omega_0 t_m \approx 3$ est connue dans la littérature.

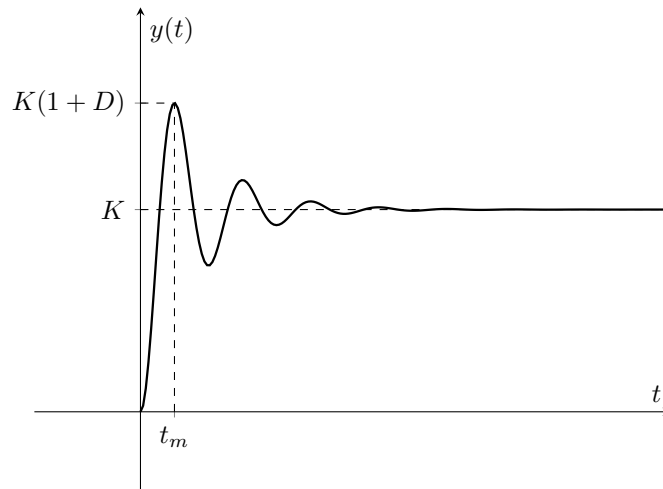


FIGURE 10.5 – Réponse indicielle (fonction de transfert d'ordre 2 sans zéro)

10.3 Analyse harmonique

10.3.1 Procédure expérimentale

Considérons un système linéaire de fonction de transfert $H(p)$, excité par un signal sinusoïdal $u(t) = U \sin(\omega t)$. Si le système est stable, la sortie du système est, en régime permanent, une sinusoïde de même pulsation (figure 10.6) :

$$y(t) = U |H(j\omega)| \sin(\omega t + \arg(H(j\omega))) \quad (10.3.1)$$

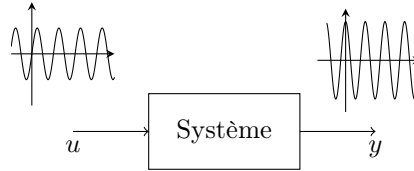


FIGURE 10.6 – Analyse harmonique

La procédure expérimentale consiste à relever $|H(j\omega)|$ et $\arg(H(j\omega))$ (appelés usuellement gain et phase) pour un ensemble de valeurs de ω . Cette opération peut se faire avec un oscilloscope, avec des appareils appelés transféromètres, ou en utilisant des logiciels spécialisés. Les courbes obtenues représentent la réponse harmonique du système, qui constitue un modèle non paramétrique de celui-ci (ce qui peut suffire dans certains cas).

L'ensemble des valeurs de ω choisies doit permettre de tracer un graphe suffisamment complet, et suffisamment précis. Il faut aussi s'assurer que la sortie a effectivement une allure sinusoïdale, et que les résultats sont relativement indépendants de la valeur de U , faute de quoi l'hypothèse de linéarité n'est pas valide. Il convient de remarquer que cette méthode nécessitant un nombre assez important de mesures, elle devient fastidieuse dans le cas de systèmes très lents.

En général, on souhaite aller plus loin et obtenir un modèle paramétrique, qui peut s'écrire sous la forme déjà rencontrée dans la section 5.2.1 :

$$H(p) = K p^q \frac{\prod_{k=1}^{n_k} \left(1 + \frac{p}{\omega_k}\right) \prod_{l=1}^{n_l} \left(1 + 2\zeta_l \frac{p}{\omega_l} + \left(\frac{p}{\omega_l}\right)^2\right)}{\prod_{m=1}^{n_m} \left(1 + \frac{p}{\omega_m}\right) \prod_{n=1}^{n_n} \left(1 + 2\zeta_n \frac{p}{\omega_n} + \left(\frac{p}{\omega_n}\right)^2\right)} \quad (10.3.2)$$

Il n'y a pas de méthode systématique pour effectuer cette opération, dont les principes généraux, déduits des résultats de la section 5.2.4 sont les suivants :

- le comportement asymptotique en basse fréquence fixe les valeurs du réel K et de l'entier relatif q ;
- la recherche des pulsations de brisure $\omega_k, \omega_l, \omega_m, \omega_n$ s'effectue principalement à partir des asymptotes de la courbe de gain tracée dans le plan de Bode, en veillant à ce que leurs positions soient cohérentes avec la courbe de phase ;
- les résonances éventuelles sur la courbe de gain permettent de déterminer les amortissements $\zeta_l < 1$ des paires de pôles ou zéros complexes, en veillant là aussi à la cohérence avec la courbe de phase ;
- s'il n'est pas possible de rendre la phase cohérente avec le gain, il peut être nécessaire de compléter la fonction de transfert ci-dessus en la multipliant par un retard pur $\exp(-\tau p)$ (section 5.4.3) ; ou par une fonction passe-tout $\frac{1 - \tau p}{1 + \tau p}$ (section 5.4.1) : celles-ci n'ont aucun effet sur la courbe de gain, et la valeur de τ peut donc être choisie pour ajuster au mieux la courbe de phase.

En pratique, on cherchera en général à minimiser le nombre de pôles et de zéros, de façon à obtenir le modèle le plus simple possible. Par ailleurs, on voit que la détermination du modèle se

base essentiellement sur les propriétés des réponses harmoniques des systèmes du premier et du second ordre (section 5.5.1), qui de plus sont eux-mêmes des modèles suffisants dans un grand nombre de cas. Il importe donc de bien connaître ces réponses.

Remarque 10.3.1. *Toute information a priori peut être utilisée pour faciliter l'obtention du modèle paramétrique ; si par exemple on connaît l'expression littérale de la fonction de transfert, voire la valeur de certains paramètres, le nombre de pôles et de zéros de la fonction de transfert (et même dans certains cas leur nature réelle ou complexe) est connu ; on connaît donc à l'avance l'allure des courbes de gain et de phase, ce qui simplifie notablement le travail.*

10.4 Méthode générale statistique

10.4.1 Principe général de la méthode statistique d'identification de modèle

Cette méthode est très générale et peut se décliner selon différentes modalités qui englobent de très nombreux cas d'applications. Elle s'applique notamment lorsqu'il n'est pas possible de mesurer directement les valeurs des paramètres d'un modèle ou lorsque l'on a plusieurs hypothèses sur la forme à donner aux équations d'un modèle (sa structure) et que l'on cherche à choisir la plus adaptée pour reproduire le comportement du système réel. A l'aide de données expérimentales, on va chercher (i) à identifier la *structure du modèle*, c'est-à-dire à déterminer quel est le 'meilleur' modèle parmi cet ensemble et (ii) en estimer les *paramètres*, c'est-à-dire à donner des valeurs numériques adéquates à ses paramètres.

La figure 10.7 illustre ce cadre dans le cas de modèles à temps continu d'ordre 1. On cherche, parmi un ensemble de *modèles candidats* $M(\theta)$, $\bar{M}(\theta)$, $\tilde{M}(\tilde{\theta})$, lequel permet de reproduire au mieux un jeu de données expérimentales une fois estimées les valeurs de ses paramètres, avec θ , $\bar{\theta}$ et $\tilde{\theta}$ des vecteurs contenant les paramètres à identifier.

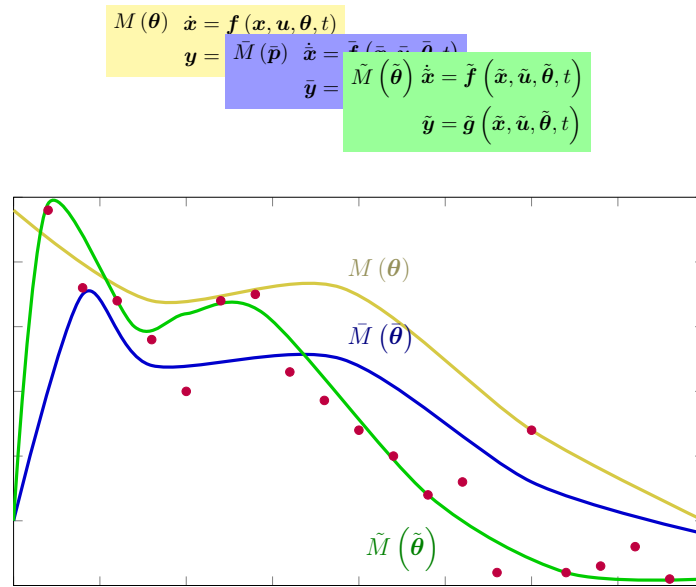


FIGURE 10.7 – Méthode générale (statistique) pour l'identification de modèle

Pour simplifier la présentation, notons ici un modèle par une fonction de réponse déterministe à une entrée u et caractérisée par un vecteur de paramètres θ . Les N_m modèles candidats¹ seront ainsi représentés par les fonctions f^i , avec $i \in \{1, \dots, N_m\}$:

$$y^i = f^i(u, \theta_i)$$

1. Dans la suite, le nombre des modèles candidats est noté N_m .

En réalité, chaque fonction f^i peut représenter par exemple un ensemble d'équations d'évolution et d'observation constituant un modèle dynamique continu ou discret, dont y constitue une sortie ou un ensemble de sorties pour lesquelles on peut disposer de mesures associées.

Le principe général de la méthode est le suivant :

1. On construit un ensemble de modèles candidats, tous supposés pouvoir convenir pour représenter le système réel. Il peut s'agir, indifféremment, de *modèles de connaissance ou de comportement*. Ils peuvent être construits à partir d'une analyse préalable du fonctionnement du système, d'une analyse de données ou d'expériences antérieures, d'informations a priori recueillies auprès d'experts, de la littérature scientifique, etc.
2. On recueille des données expérimentales (u_n, y_n^{obs}) , avec $n \in \{1, \dots, N\}$, sur le système réel, soit en le soumettant à une entrée-test lorsque l'expérimentateur peut agir sur les entrées, soit en les mesurant lorsqu'elles sont imposées (par exemple des conditions météorologiques). Dans toute la suite, N désignera le nombre de mesures.
3. Pour chacun des N_m modèles candidats :
 - On applique cette même entrée au modèle i .
 - On définit un critère quantifiant la distance entre la sortie du modèle y^i et les observations y_{obs}^i , c'est-à-dire qu'on évalue à quel point le modèle permet de reproduire les données. Ce critère dépend du vecteur de paramètres² $\theta^i = (\theta_1^i, \theta_2^i, \dots, \theta_{n_i}^i)^T$ puisque les équations du modèle en dépendent.
 - On cherche les valeurs à affecter aux paramètres afin d'optimiser ce critère, pour le modèle considéré.
4. On compare les modèles ainsi obtenus pour en sélectionner le meilleur.

Pour rendre cette procédure opérationnelle, on voit qu'il s'agit de spécifier différents éléments, que nous allons traiter dans les paragraphes suivants : en étape 3, la notion de distance entre une sortie de modèle et des observations (paragraphe 10.4.2), la façon d'optimiser ce critère (paragraphe 10.4.4) et enfin, en étape 4, le critère de comparaison permettant de sélectionner un modèle (paragraphe 10.5.2).

En dernier lieu, il s'agira d'évaluer le modèle ainsi obtenu afin d'en juger les performances et la fiabilité (paragraphe 10.5 et chapitre 11).

10.4.2 Estimation paramétrique par la méthode des moindres carrés

Nous nous intéressons ici à la procédure d'identification des paramètres d'un modèle, aussi appelée *estimation paramétrique*, qui intervient en étape 3 de la démarche générale présentée ci-dessus. La structure du modèle, c'est-à-dire la forme de ses équations, est ici considérée comme connue, puisqu'on va appliquer cette procédure à chacun de nos N_m modèles candidats. L'objectif consiste à trouver quelles valeurs affecter aux paramètres inconnus qui interviennent dans ces équations et ne peuvent pas être mesurés de manière directe.

De nombreuses méthodes coexistent, *bayésiennes* ou *fréquentistes*, qui seront vues dans le cours de Statistiques et Machine Learning. Afin de pouvoir compléter le traitement des problèmes de modélisation qui nous intéressent dans ce cours, nous allons introduire ici l'une des méthodes fréquentistes les plus classiques (avec celle du *maximum de vraisemblance* dont elle est en fait un cas particulier) : la *méthode des moindres carrés*.

Remarque 10.4.1. Pour faciliter la lecture du document, nous allons omettre les indices i de chaque modèle candidat et représenter le modèle simplement par $y = f(u, \theta)$ puisque sa structure est supposée connue.

Comme on s'en doute, les observations expérimentales ne peuvent jamais être reproduites de manière exacte par le modèle f . L'erreur, notée $\epsilon = y - f(u, \theta)$, est la différence entre la prédiction du modèle et la réponse mesurée pour une situation donnée. Ce terme d'erreur se décompose en deux composantes : l'*erreur de mesure* et l'*erreur liée à la non-adéquation exacte du modèle au système réel*.

2. L'indice supérieur est associé au modèle. L'indice inférieur permet de désigner l'élément correspondant d'un vecteur.

Avec la *méthode des moindres carrés*, le critère permettant de juger de l'adéquation entre une sortie de modèle et des observations se définit de la façon suivante.

Définition 10.4.1

Le *critère des moindres carrés ordinaires* $J(\boldsymbol{\theta})$ est la norme euclidienne des erreurs, qui correspond à une distance euclidienne entre les observations y_n^{obs} et les simulations $y(u_n, \boldsymbol{\theta})$:

$$J(\boldsymbol{\theta}) = \|\epsilon\|^2 = \sum_{n=1}^N (y_n^{obs} - y(u_n, \boldsymbol{\theta}))^2 \quad (10.4.1)$$

La procédure d'estimation paramétrique consiste à chercher les valeurs de paramètres minimisant ce terme :

$$\hat{\boldsymbol{\theta}} = \arg \min_{\boldsymbol{\theta}} J(\boldsymbol{\theta}) \quad (10.4.2)$$

Cette équation définit un *problème d'optimisation*. On va voir dans les deux paragraphes suivants que, lorsque le critère J est linéaire en les paramètres, on peut en exprimer la solution sous forme analytique. En revanche, lorsque le problème d'optimisation est non-linéaire (on parle alors parfois de *méthode des moindres carrés non-linéaires*), la solution ne pourra généralement être obtenue que de manière numérique, selon les méthodes qui seront décrit en paragraphe 10.4.4.

10.4.3 Interprétation statistique

Il est intéressant d'adopter ici un point de vue probabiliste. En effet, les observations ne peuvent généralement pas être mesurées avec une précision infinie et sont toujours entachées d'un bruit de mesure, c'est-à-dire que si on refaisait une nouvelle série de mesures, on obtiendrait un résultat légèrement différent, et donc un vecteur de paramètres optimal lui aussi légèrement différent. On va donc considérer les mesures y_n^{obs} comme des variables aléatoires Y_n^{obs} : la solution du problème d'optimisation (10.4.2), notée $\hat{\boldsymbol{\theta}}$, définit un *estimateur* de $\boldsymbol{\theta}$. On pourra donc chercher à caractériser la loi de probabilité de cet estimateur ou simplement certaines de ses propriétés comme sa *variance* ou son *espérance*.

Plaçons-nous dans le cas idéal, mais très étudié du fait du grand nombre de propriétés que l'on peut obtenir sur l'estimateur, où les équations du modèle sont supposées être les 'bonnes', c'est-à-dire qu'on a seulement du bruit de mesure. On note $\boldsymbol{\theta}^*$ le vecteur des 'vrais' paramètres et on suppose que les mesures se distribuent de manière *indépendante et identiquement distribuées* (i.i.d.), avec une distribution Gaussienne, centrée et de variance constante, autour des valeurs théoriques $y(u_n, \boldsymbol{\theta}^*)$. Cela revient à dire que les erreurs $\epsilon_n = Y_n^{obs} - y(u_n, \boldsymbol{\theta}^*)$ vérifient $\epsilon_n \sim \mathcal{N}(0, \sigma^2)$. Cela veut dire que l'on considère que les observations vont statistiquement se répartir de manière 'équilibrée' autour de la valeur prédite par le modèle ($\forall n, \mathbb{E}[\epsilon_n] = 0$) et avec une dispersion autour de cette valeur qui ne dépend pas de la valeur de l'entrée u , ou du temps pour un modèle dynamique ($\forall n, \text{Var}[\epsilon_n] = \sigma^2$). L'estimation paramétrique consiste :

- d'une part à obtenir des valeurs numériques $\boldsymbol{\theta}^{num}$, réalisations de la variable aléatoire $\boldsymbol{\theta}$, qui minimisent le critère des moindres carrés, c'est-à-dire la norme des résidus (ou erreurs résiduelles, qui sont les réalisations des variables aléatoires ϵ_n), c'est-à-dire encore l'écart entre les valeurs observées et la simulation obtenue par le modèle ;
- d'autre part à obtenir des propriétés statistiques sur l'estimateur $\boldsymbol{\theta}$. L'intérêt de ces propriétés est qu'elles nous permettent de savoir si nous pouvons être plus ou moins confiants sur le fait que les valeurs numériques que nous attribuerons aux paramètres via cette méthode soient proches des vrais paramètres. Autrement dit, que si l'on refaisait une nouvelle série de mesures, qui serait forcément un peu différente de la précédente, les paramètres estimés ne seraient pas trop éloignés des précédents.

Cette démarche est illustrée graphiquement en figure 10.8, où la variabilité des mesures est représentée par les barres d'incertitudes.

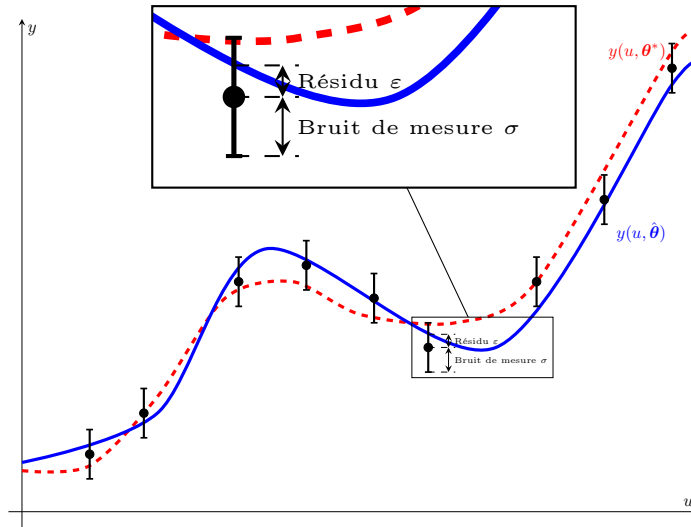


FIGURE 10.8 – Illustration de la méthode des moindres carrés

Définition 10.4.2

On appelle *biais d'un estimateur* la quantité $\mathbb{E}[\hat{\theta}] - \theta^*$.

On peut montrer que, sous ces hypothèses, l'estimateur des moindres carrés :

- converge en probabilité vers le vrai paramètre quand le nombre d'observation augmente ;
- est sans biais, c'est-à-dire que son espérance est égale au vrai paramètre ;
- est de variance minimale parmi tous les estimateurs sans biais ;
- l'incertitude sur les paramètres estimés se caractérise ainsi via la matrice de covariance de l'estimateur. Celle-ci s'exprime, de *manière exacte* dans le cas d'un modèle linéaire en ses paramètres ou de *manière approchée* dans le cas d'un modèle non-linéaire (de *manière asymptotique*, c'est-à-dire pour N suffisamment grand), par :

$$\text{Cov}[\hat{\theta}] \approx \sigma^2 (F^T F)^{-1}$$

où F est la matrice jacobienne du modèle (matrice des dérivées partielles par rapport aux paramètres) en $\hat{\theta}$ que l'on approxime généralement de manière numérique.

Exemple 10.4.1. Pour un modèle linéaire s'écrivant $\mathbf{Y} = \mathbf{U}\boldsymbol{\theta} + \boldsymbol{\epsilon}$, où $\mathbf{U}$ est la matrice définissant les variables du modèle et $\boldsymbol{\epsilon}$ le vecteur des erreurs avec $\epsilon_n \sim \mathcal{N}(0, \sigma^2)$ indépendante et identiquement distribuées, la solution du problème d'optimisation (10.4.2), i.e. $\min_{\boldsymbol{\theta}} J(\boldsymbol{\theta}) = \min_{\boldsymbol{\theta}} \sum_{n=1}^N \epsilon_n^2$ s'écrit :

$$\hat{\boldsymbol{\theta}} = (\mathbf{U}^T \mathbf{U})^{-1} \mathbf{U}^T \mathbf{Y}$$

sous réserve que $\mathbf{U}^T \mathbf{U}$ soit bien inversible (voir détails de résolution sur les applications dans les paragraphes suivants). Alors l'estimateur $\hat{\boldsymbol{\theta}}$ suit une loi normale dont l'espérance et la variance valent :

$$\begin{aligned} \mathbb{E}[\hat{\boldsymbol{\theta}}] &= \mathbb{E}[(\mathbf{U}^T \mathbf{U})^{-1} \mathbf{U}^T \mathbf{Y}] \\ &= (\mathbf{U}^T \mathbf{U})^{-1} \mathbf{U}^T \mathbb{E}[\mathbf{U}\boldsymbol{\theta} + \boldsymbol{\epsilon}] \end{aligned}$$

Comme $\mathbb{E}[\boldsymbol{\epsilon}] = \mathbf{0}$, on peut écrire :

$$\begin{aligned} \mathbb{E}[\hat{\boldsymbol{\theta}}] &= (\mathbf{U}^T \mathbf{U})^{-1} \mathbf{U}^T \mathbf{U} \boldsymbol{\theta} \\ &= \boldsymbol{\theta} \end{aligned}$$

$$\begin{aligned}
\text{Var}[\hat{\theta}] &= \text{Var}[(U^T U)^{-1} U^T Y] \\
&= (U^T U)^{-1} U^T \text{Var}[U\theta + \epsilon] U (U^T U)^{-1} \\
&= \sigma^2 (U^T U)^{-1} U^T U (U^T U)^{-1} \\
&= \sigma^2 (U^T U)^{-1}
\end{aligned}$$

car $\text{Var}[U\theta + \epsilon] = \text{Var}[\epsilon] = \sigma^2 I_N$. Le cas du modèle linéaire sera développé dans les paragraphes suivants sur des cas d'application.

Remarque 10.4.2. Lorsque l'hypothèse d'une variance constante ($\forall n, \text{Var}[\epsilon_n] = \sigma^2$) n'est pas vérifiée, ce qui arrive fréquemment en pratique, on généralise l'approche en pondérant chaque terme de la somme de l'équation (10.4.1) par la variance de l'erreur associée ; on parle alors de méthode des moindres carrés pondérés dont le critère devient :

$$J(\theta) = \|\epsilon\|^2 = \sum_{n=1}^N \frac{(y_n^{\text{obs}} - y(x_n, \theta))^2}{\sigma_n^2} \quad (10.4.3)$$

Intuitivement, une mesure aura alors d'autant plus de poids que son incertitude (la variance de la variable aléatoire dont elle est une réalisation) sera faible.

Remarque 10.4.3. Lorsque le problème d'optimisation défini en (10.4.2) admet plusieurs solutions, on parle de **non-identifiabilité**. Cela signifie que différentes valeurs de paramètres peuvent définir exactement la même relation entrée-sortie. Plus formellement, un paramètre θ_i est dit **structurellement globalement identifiable** si pour presque tout $\theta \in \Theta$, on a :

$$f(\theta) = f(\hat{\theta}) \implies \theta_i = \hat{\theta}_i$$

Nous allons à présent appliquer cette méthode sur deux exemples de modèles conduisant à des critères **linéaires** par rapport aux paramètres, en montrant comment on peut, dans ce cas, résoudre analytiquement le problème d'optimisation (10.4.3).

10.4.3.1 Utilisation de la méthode des moindres carrés linéaires : détermination de la fonction de transfert d'un système discret à réponse impulsionnelle finie

Rappelons tout d'abord que la réponse impulsionnelle $h(n)$ d'un système discret constitue un modèle paramétrique de ce système : en effet sa réponse à une entrée quelconque $u(n)$ se déduit par convolution de $u(n)$ et de $h(n)$.

Nous ferons l'hypothèse que le système est causal, c'est-à-dire que $h(n) = 0$ pour tout $n < 0$, et qu'il est stable, de sorte que $\lim_{n \rightarrow \infty} h(n) = 0$.

Si le système est à réponse impulsionnelle finie, il existe un instant N tel que $h(n) = 0$ pour tout³ $n > N$. Sa fonction de transfert s'écrit alors :

$$H(z) = \sum_{k=0}^N h(k) z^{-k} = a_0 + a_1 z^{-1} + \dots + a_N z^{-N}, \text{ avec } a_k = h(k) \quad (10.4.4)$$

Soit $u(n)$ le signal appliqué en entrée et $y(n)$ le signal observé en sortie du système. La sortie du modèle ayant pour paramètres $a_0, \dots, a_N$ s'écrit :

$$y_m(n) = a_0 u(n) + a_1 u(n-1) + \dots + a_N u(n-N) \quad (10.4.5)$$

En pratique, comme nous l'avons expliqué plus haut, la sortie $y(n)$ observée ne sera jamais rigoureusement égale à $y_m(n)$: en effet, même en supposant que l'hypothèse de linéarité et les hypothèses formulées plus haut sont valides, les paramètres $a_0, \dots, a_N$ que nous allons déterminer seront toujours entachés d'incertitudes, et les mesures effectuées en sortie entachées de bruit.

3. Plus généralement les développements proposés s'appliquent si on peut supposer $h(n) \approx 0$ pour tout $n > N$.

Considérons l'erreur de sortie $e(n) = y(n) - y_m(n)$ entre la sortie observée et la sortie du modèle, et le critère :

$$J = \sum_{n=N}^{N+M} e(n)^2 \quad (10.4.6)$$

L'erreur $e(n)$ dépendant linéairement des paramètres $a_0, \dots, a_N$, les valeurs de ceux-ci qui minimisent le critère J peuvent être obtenus de manière analytique. En effet, les différentes valeurs de $e(n)$ qui apparaissent dans le critère s'écrivent sous forme matricielle :

$$\begin{pmatrix} e(N) \\ e(N+1) \\ \vdots \\ e(N+M) \end{pmatrix} = \begin{pmatrix} y(N) \\ y(N+1) \\ \vdots \\ y(N+M) \end{pmatrix} - \begin{pmatrix} u(N) & u(N-1) & \cdots & u(0) \\ u(N+1) & u(N) & \cdots & u(1) \\ \vdots & \vdots & \ddots & \vdots \\ u(N+M) & u(N+M-1) & \cdots & u(M) \end{pmatrix} \begin{pmatrix} a_0 \\ a_1 \\ \vdots \\ a_N \end{pmatrix}$$

Cette relation peut s'écrire :

$$\mathbf{e} = \mathbf{y} - \mathbf{U}\boldsymbol{\theta}, \text{ avec } J = \|\mathbf{e}\|^2 = \mathbf{e}^T \mathbf{e} = \mathbf{y}^T \mathbf{y} - (2\mathbf{y}^T \mathbf{U}) \boldsymbol{\theta} + \boldsymbol{\theta}^T (\mathbf{U}^T \mathbf{U}) \boldsymbol{\theta} \quad (10.4.7)$$

Le gradient de J par rapport à $\boldsymbol{\theta}$ s'écrit :

$$\frac{\partial J}{\partial \boldsymbol{\theta}} = -2\mathbf{U}^T \mathbf{y} + 2(\mathbf{U}^T \mathbf{U}) \boldsymbol{\theta}$$

Le vecteur $\boldsymbol{\theta}$ qui annule le gradient (et donc qui minimise J) est par conséquent :

$$\boldsymbol{\theta} = (\mathbf{U}^T \mathbf{U})^{-1} \mathbf{U}^T \mathbf{y} \quad (10.4.8)$$

La procédure expérimentale demande donc de choisir :

- l'instant N à partir duquel on peut affirmer ou supposer que $h(n) \approx 0$;
- la longueur $N + M$ de la séquence d'entrée appliquée, sachant que pour que la matrice $\mathbf{U}^T \mathbf{U}$ soit inversible, où $\mathbf{U}$ est de dimension $(M+1) \times (N+1)$, il faut que $M \geq N$; théoriquement, le résultat est d'autant meilleur que M est grand, mais la matrice $\mathbf{U}^T \mathbf{U}$ peut être mal conditionnée et donc difficile à inverser, notamment pour de fortes valeurs de M ;
- la nature de l'entrée appliquée. En pratique, n'importe quelle entrée peut théoriquement être utilisée, mais celle-ci doit être suffisamment riche pour que les paramètres obtenus soient réutilisables dans d'autres conditions. La sinusoïde, signal au spectre beaucoup trop pauvre, est donc à proscrire ; un *échelon* est déjà plus riche, mais la *séquence binaire pseudo-aléatoire*⁴ (SBPA) est dans bien des cas le signal le plus approprié.

10.4.3.2 Utilisation de la méthode des moindres carrés linéaires : Détermination de la fonction de transfert d'un système discret à réponse impulsionnelle infinie

Si le système est à réponse impulsionnelle infinie, sa fonction de transfert en z s'écrit sous la forme plus générale :

$$H(z) = \frac{a_0 + a_1 z^{-1} + a_2 z^{-2} + \cdots + a_N z^{-N}}{1 + b_1 z^{-1} + b_2 z^{-2} + \cdots + b_N z^{-N}} \quad (10.4.9)$$

Soit $y(n)$ la sortie observée en réponse à une entrée appliquée $u(n)$. L'équation de récurrence qui correspond à cette fonction de transfert s'écrit :

$$y(n) = a_0 u(n) + a_1 u(n-1) + \cdots + a_N u(n-N) - b_1 y(n-1) - \cdots - b_N y(n-N) \quad (10.4.10)$$

En pratique, pour les mêmes raisons que dans le paragraphe précédent, les valeurs observées en sortie du système ne vérifient jamais rigoureusement cette équation. L'utilisation directe du critère des moindres carrés portant sur l'erreur de sortie conduirait ici à un problème d'optimisation non-linéaire, plus compliqué à résoudre. On va préférer se ramener à un critère linéaire en les

4. Une séquence binaire pseudo-aléatoire est un signal qui se présente comme une succession d'échelons d'amplitudes identiques, mais de durée variable.

paramètres, qui, si le modèle était parfait, donnerait des résultats équivalents. Pour cela, définissons l'erreur d'équation $\epsilon(n)$ comme la différence entre les deux membres de l'équation (10.4.10) : c'est la différence entre une valeur observée au temps n et la valeur prédite par le modèle en prenant pour valeurs des états précédents, non pas celles calculées par le modèle, mais directement les observations. Alors, comme cette erreur d'équation $\epsilon(n)$ dépend linéairement des paramètres $\{a_i, b_i\}$, le critère défini par :

$$J = \sum_{n=N}^{N+M} \epsilon(n)^2 \quad (10.4.11)$$

en dépend lui aussi linéairement.

En effet, les différentes valeurs de $\epsilon(n)$ qui apparaissent dans le critère s'écrivent sous forme matricielle :

$$\begin{pmatrix} \epsilon(N) \\ \epsilon(N+1) \\ \vdots \\ \epsilon(N+M) \end{pmatrix} = \begin{pmatrix} y(N) \\ y(N+1) \\ \vdots \\ y(N+M) \end{pmatrix} - \begin{pmatrix} u(N) & \cdots & u(0) & -y(N-1) & \cdots & -y(0) \\ u(N+1) & \cdots & u(1) & -y(N) & \cdots & -y(1) \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ u(N+M) & \cdots & u(M) & -y(N+M-1) & \cdots & -y(M) \end{pmatrix} \begin{pmatrix} a_0 \\ \vdots \\ a_N \\ b_1 \\ \vdots \\ b_N \end{pmatrix}$$

On peut donc écrire cette relation sous la forme :

$$\boldsymbol{\epsilon} = \mathbf{y} - \mathbf{A}\boldsymbol{\theta} \text{ avec } J = \|\boldsymbol{\epsilon}\|^2 = \boldsymbol{\epsilon}^T \boldsymbol{\epsilon} \quad (10.4.12)$$

Le vecteur des paramètres $\boldsymbol{\theta}$ qui minimise J est donc :

$$\boldsymbol{\theta} = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{y} \quad (10.4.13)$$

La procédure expérimentale demande donc de choisir :

- l'ordre N de la fonction de transfert ;
- la longueur $N + M$ de la séquence d'entrée appliquée, sachant que pour que la matrice $\mathbf{A}^T \mathbf{A}$ soit inversible, avec $\mathbf{A}$ de dimension $(M + 1) \times (2N + 1)$, il faut que $M \geq 2N$; là aussi, le résultat est théoriquement d'autant meilleur que M est grand, avec les mêmes problèmes liés à l'inversion de la matrice $\mathbf{A}^T \mathbf{A}$;
- la nature de l'entrée appliquée, ce qui renvoie là aussi à la discussion du paragraphe précédent.

Pour les cas **non-linéaires**, le problème d'optimisation (10.4.3) ne se résout pas aussi simplement et nécessite de faire appel à des algorithmes d'optimisation numérique. La section suivante aborde brièvement ces méthodes, qui seront vues de manière approfondie dans le cours d'*Optimisation*.

10.4.4 Notions d'optimisation

On donne ici une courte introduction à l'optimisation, dans le but de son utilisation en modélisation.

10.4.4.1 Généralités et formulation du problème

L'optimisation est l'activité de minimiser ou maximiser des fonctions en vérifiant des relations annexes. Dans le cadre de la modélisation, on peut en distinguer trois grands usages :

1. Le modèle peut, en lui-même, se présenter sous la forme d'un problème d'optimisation. C'est par exemple le cas des modèles dits *téléonomiques*. On peut citer en exemple :
 - le cas d'un modèle de croissance de plantes dans lequel la position des nouvelles pousses est déterminée de sorte à maximiser son interception lumineuse ;
 - le cas de la minimisation de l'énergie potentielle totale sert à trouver la position d'un système en mécanique (ex : la position d'une corde).

2. La phase d'estimation paramétrique des modèles à partir de données expérimentales se formule généralement comme un problème d'optimisation : il s'agit de maximiser une vraisemblance ou de minimiser un critère de moindres carrés. Dans les cas simples ce problème pourra être résolu *analytiquement* ; la plupart du temps *numériquement*.
3. Les problèmes d'optimisation se posent enfin également dans le cadre de l'utilisation des modèles pour la conception et la commande optimale. On va par exemple chercher à maximiser la performance d'un système sous contraintes de résistance, qualité, fabrication, coût, etc.

La formulation générale d'un problème d'optimisation est la suivante :

$$\begin{aligned}
 & \min_{x \in \mathbb{R}^n} f(x) \\
 & \text{sous contraintes} \\
 & g(x) \leq 0 \\
 & h(x) = 0
 \end{aligned} \tag{10.4.14}$$

où x désigne les variables d'optimisation. On notera x^* une solution de ce problème. Il s'agit ici d'un problème d'optimisation *continue* puisque les variables sont prises dans $\mathbb{R}^n$. La fonction f est appelée la *fonction coût*. Dans le problème d'optimisation (10.4.14), g désigne les *contraintes d'inégalité* et h les *contraintes d'égalité*. Il va de soi que la plupart des problèmes réels ou industriels ne sont pas initialement sous cette forme et la première étape du travail consiste en général à la formulation du problème sous forme standard.

Le problème défini précédemment est dit *global* car on en cherche les solutions sur tout l'espace de définition des variables. Ce problème peut être numériquement compliqué, typiquement quand il y a des solutions isolées comme illustré en figure 10.9.

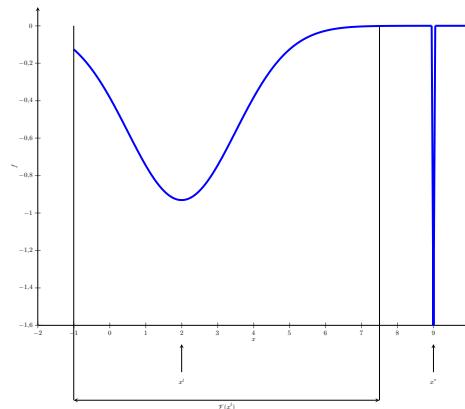


FIGURE 10.9 – Cas d'un problème d'optimisation globale difficile.

On se contentera donc parfois de résoudre le problème d'optimisation *local* qui consiste à trouver une solution sur un voisinage d'un point de référence.

Les problèmes d'optimisation peuvent parfois être résolus analytiquement, mais, le plus souvent, on approche leur(s) solution(s) avec des méthodes numériques itératives.

10.4.4.2 Principe d'une méthode de descente pour l'estimation paramétrique

Ce paragraphe fait appel au calcul différentiel. Un rappel sur les notions de base du calcul différentiel est disponible en annexe F. Nous présentons ici des résultats qui s'appliquent essentiellement dans le cas d'une optimisation sans contrainte. La prise en compte des contraintes requiert des adaptations qui seront vues dans le cours d'*Optimisation*.

Revenons à notre problème d'estimation paramétrique par minimisation du terme d'erreur $J(\theta)$, de $\mathbb{R}^N$ dans $\mathbb{R}$, en notant θ le vecteur des paramètres cherchés. Le gradient de J par rapport à θ est le vecteur colonne de dimension N :

$$\nabla J(\theta) = \left(\frac{\partial J}{\partial \theta_1} \quad \cdots \quad \frac{\partial J}{\partial \theta_N} \right)^T \tag{10.4.15}$$

Une procédure possible (*méthode du gradient*) est alors la suivante :

- *Etape 1* : On initialise le vecteur de paramètres à une valeur θ_0 , soit à partir de connaissances a priori (ordre de grandeur par exemple), soit de façon quelconque si aucune information n'est disponible ;
- *Etape 2* : On calcule le critère $J(\theta_k)$ correspondant à la valeur courante θ_k des paramètres ;
- *Etape 3* : A partir de θ_k et de $J(\theta_k)$, on calcule une nouvelle valeur θ_{k+1} ;
- *Etape 4* : Tant que le modèle n'est pas satisfaisant ou qu'un critère de convergence n'est pas satisfait, on revient en étape 2.

Le calcul effectué en étape 3 doit évidemment permettre de faire décroître globalement le critère, de façon à converger vers un minimum de la fonction $J(\theta)$. A l'itération k de la procédure, on dispose d'une valeur θ_k du vecteur θ . Le développement limité au 1^{er} ordre de $J(\theta)$ au voisinage de θ_k s'écrit :

$$J(\theta) \approx J(\theta_k) + (\theta - \theta_k)^T \nabla J(\theta_k)$$

On cherche une nouvelle valeur θ_{k+1} de façon que $J(\theta_{k+1})$ soit le plus petit possible. D'après la formule ci-dessus, il convient donc de choisir $\theta_{k+1} - \theta_k$ colinéaire à $\nabla J(\theta_k)$ et de sens contraire, soit :

$$\theta_{k+1} = \theta_k - \lambda_k \nabla J(\theta_k) \quad (10.4.16)$$

Le paramètre $\lambda_k > 0$ est appelé le *pas* d'optimisation. Il est en général adapté à chaque itération.

La figure 10.10 illustre le principe correspondant à cette formule dans le cas d'une fonction dépendant de 2 paramètres : les courbes portées sur la figure sont les courbes de niveau de la fonction $J(\theta)$, c'est-à-dire l'ensemble des vecteurs θ conduisant à une valeur donnée de $J(\theta)$. A l'itération k , le déplacement à partir de θ_k se fait suivant la ligne de plus grande pente. Cette démarche s'apparente à celle d'un marcheur qui veut rejoindre le plus rapidement possible le fond d'une vallée sans le voir.

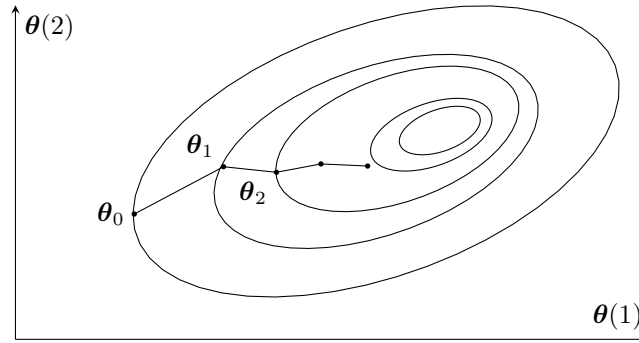


FIGURE 10.10 – Optimisation d'une fonction de 2 variables par la méthode du gradient

En pratique, si le gradient $\nabla J(\theta_k)$ n'est pas connu analytiquement, il est déterminé de façon approchée en calculant les valeurs de $J(\theta)$ pour un ensemble de points proches de θ_k . L'avantage des méthodes de gradient est qu'on observe en général une évolution rapide de θ_k au cours des itérations successives lorsqu'on est loin d'un minimum de $J(\theta)$. En revanche, cette technique peut conduire à des oscillations ou à une progression très lente au voisinage du minimum (figure 10.10), et sa convergence peut donc être relativement longue.

La *méthode de Newton* fournit une alternative intéressante. Elle utilise la matrice des dérivées secondes de $J(\theta)$, appelée *matrice hessienne*, qui est carrée symétrique de dimension $N \times N$:

$$\mathbf{H}(\theta) = \begin{pmatrix} \frac{\partial^2 J}{\partial \theta_1^2} & \cdots & \frac{\partial^2 J}{\partial \theta_1 \partial \theta_N} \\ \vdots & \ddots & \vdots \\ \frac{\partial^2 J}{\partial \theta_N \partial \theta_1} & \cdots & \frac{\partial^2 J}{\partial \theta_N^2} \end{pmatrix} \quad (10.4.17)$$

et s'appuie sur le développement limité au 2^{ème} ordre de $J(\boldsymbol{\theta})$ au voisinage de $\boldsymbol{\theta}_k$:

$$J(\boldsymbol{\theta}) \approx J(\boldsymbol{\theta}_k) + (\boldsymbol{\theta} - \boldsymbol{\theta}_k)^T \nabla J(\boldsymbol{\theta}_k) + \frac{1}{2} (\boldsymbol{\theta} - \boldsymbol{\theta}_k)^T \mathbf{H}(\boldsymbol{\theta}_k) (\boldsymbol{\theta} - \boldsymbol{\theta}_k)$$

Cherchons la nouvelle valeur $\boldsymbol{\theta}_{k+1}$ de façon que le second membre soit minimum. En annulant sa dérivée par rapport à $\boldsymbol{\theta}$, on obtient :

$$\nabla J(\boldsymbol{\theta}_k) + \mathbf{H}(\boldsymbol{\theta}_k) (\boldsymbol{\theta}_{k+1} - \boldsymbol{\theta}_k) = 0, \text{ d'où } \boldsymbol{\theta}_{k+1} = \boldsymbol{\theta}_k - \mathbf{H}(\boldsymbol{\theta}_k)^{-1} \nabla J(\boldsymbol{\theta}_k) \quad (10.4.18)$$

Cette méthode présente l'inconvénient de nécessiter le calcul ou l'évaluation des dérivées secondes de $J(\boldsymbol{\theta})$ et pose évidemment des problèmes si $\mathbf{H}(\boldsymbol{\theta}_k)$ devient singulière ou mal conditionnée. Elle est aussi plus sensible à la valeur $\boldsymbol{\theta}_0$ choisie. En revanche, sa convergence est rapide au voisinage d'un minimum, où en général l'approximation par un développement limité d'ordre 2 est bien justifiée. Il existe par ailleurs des variantes (introduire un pas comme dans la méthode du gradient, remplacer $\mathbf{H}(\boldsymbol{\theta}_k)$ par une matrice définie positive lorsqu'on est loin de l'optimum, etc.).

Il est donc intéressant de combiner ces deux approches en utilisant la première lorsqu'on est loin du minimum, et la seconde lorsqu'on s'en approche.

Dans le cas, fréquent en pratique, où la fonction $J(\boldsymbol{\theta})$ présente des minima locaux, on comprend facilement que le résultat obtenu dépend en général de l'initialisation choisie : comme l'illustration (en dimension 1) de la figure 10.11 permet de le comprendre, la procédure ci-dessus risque de converger vers un minimum local suivant la valeur de $\boldsymbol{\theta}_0$.

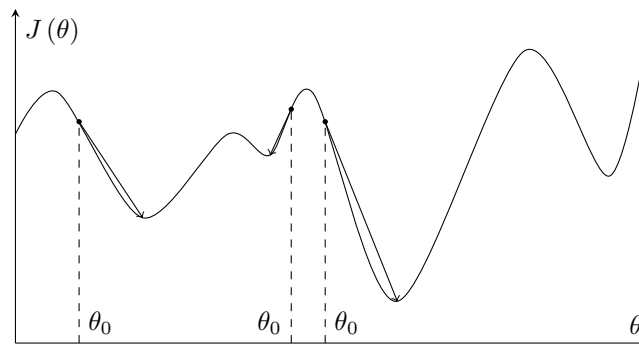


FIGURE 10.11 – Problème des minima locaux

C'est pourquoi toute information a priori sur la valeur des paramètres est très utile, et il convient de relancer la procédure ci-dessus avec des initialisations $\boldsymbol{\theta}_0$ différentes. Notons pour finir qu'il existe aussi des méthodes d'*optimisation stochastiques*, appelées *méta-heuristiques*, qui évitent de rester bloqué sur un minimum local : on peut citer notamment les *algorithmes évolutionnaires* parmi lesquels les *algorithmes génétiques*, l'*optimisation par essaim particulaire* (en anglais *Particle Swarm Optimisation*), les *méthodes de recuit simulé*. Une partie de ces méthodes sera abordée plus tard dans le cursus.

10.5 Évaluation et sélection de modèles

Les paragraphes précédents ont donné désormais des moyens d'effectuer l'estimation paramétrique des modèles proposés. La dernière étape est de les évaluer et de pouvoir les comparer.

Le but de la phase d'évaluation est de déterminer la qualité du modèle en fonction des objectifs proposés quant à son utilisation. On préférera ce terme à celui de 'validation' d'un modèle qui laisserait penser que l'on attend une réponse binaire (oui/non) à cette question alors que c'est en pratique rarement le cas. Une citation fameuse de Box (1979) indique que *tous les modèles sont faux, mais quelques uns sont utiles*⁵. Il ne s'agit bien sûr pas de remettre en cause la validité de la

5. Box (1979) – "All models are wrong but some are useful" : Now it would be very remarkable if any system existing in the real world could be exactly represented by any simple model. However, cunningly chosen parsimonious models often do provide remarkably useful approximations. (...) For such a model there is no need to ask the question "Is the model true?". If "truth" is to be the "whole truth" the answer must be "No". The only question of interest is "Is the model illuminating and useful?".

méthode de résolution des équations ou de l'exécution de l'algorithme définissant le modèle, mais bien l'adéquation de ce modèle au phénomène réel qu'il cherche à décrire. En ce sens, aucun modèle n'est bien évidemment le reflet exact du phénomène réel avec toute sa complexité. En revanche, un modèle est évalué en fonction de son utilité vis à vis de ses objectifs.

10.5.1 Quelques critères d'évaluation

L'évaluation d'un modèle est une étape importante du processus de modélisation car elle oblige, si ce n'est pas déjà fait, à se poser un certain nombre de questions :

- Quel est l'objectif de mon modèle ?
- Dans quels types de conditions mon modèle sera-t-il amené à être utilisé ?
- Quelle est ma définition d'une qualité acceptable du modèle ? avec quel(s) critère(s) la mesurer ?

Nous présentons ici différents critères d'évaluation des modèles pouvant être utilisés pour la classe d'objectifs la plus couramment utilisée, à savoir la *prédiction*. D'autres critères existent pour l'évaluation vis à vis d'objectifs d'aide à la décision ou de test d'hypothèses scientifiques par exemple. Une fois encore, il n'existe pas de critère unique d'évaluation : il est fréquent d'en utiliser plusieurs à la fois, chacun mettant en lumière différents aspects du comportement du modèle vis à vis des données.

On reprend les notations simplifiées précédentes et on considérera ici la sortie comme univariée (1D) :

$$\hat{y}_i = f(\mathbf{u}_i, \hat{\boldsymbol{\theta}})$$

où $y_i \in \mathbb{R}$ est la sortie, $\mathbf{u}_i \in \mathbb{R}^k$ est le vecteur des k variables d'entrée pour l'individu ou la situation i et $\hat{\boldsymbol{\theta}} \in \mathbb{R}^p$ représente le vecteur des valeurs de paramètres. On se donne par ailleurs un vecteur $(y_i^{obs})_{i \in \{1, \dots, N\}}$ de valeurs observées réelles que l'on a mesurées pour N individus ou situations données correspondant à une combinaison de valeurs des variables d'entrée $\mathbf{u}_i$ connues ou imposées.

Remarque 10.5.1. On rappelle que la notation $\hat{\cdot}$ indique que ces valeurs ont été estimées à partir de données expérimentales.

10.5.1.1 Evaluation graphique

Le premier type d'évaluation est l'**évaluation graphique**, sous différents formats possibles comme ceux présentés en figure 10.12 : en (a) le graphe des valeurs simulées (représentées par un trait continu) et observées (représentées par des cercles) en fonction d'une variable d'entrée u . On peut noter que l'interprétation visuelle de ce type de graphe peut être assez trompeuse car elle ne met pas en valeur l'erreur importante sur l'un des points, qui apparaît plus clairement sur le graphe des valeurs simulées vs observées en (b) ou bien, encore mieux, sur le graphe des résidus $y^{obs} - \hat{y}$ en (c). Le graphe des résidus (c) est particulièrement important dans le processus d'évaluation et devrait toujours être examiné. Il permet d'identifier les points présentant la plus grande erreur et l'amplitude relative de cette erreur par rapport à celle sur les autres points. Il permet aussi une inspection visuelle de la distribution des résidus, ce qui est un point important dans la procédure d'estimation paramétrique, surtout lorsque l'une des hypothèses sous-jacentes est que les résidus sont centrés en zéro et de même variance.

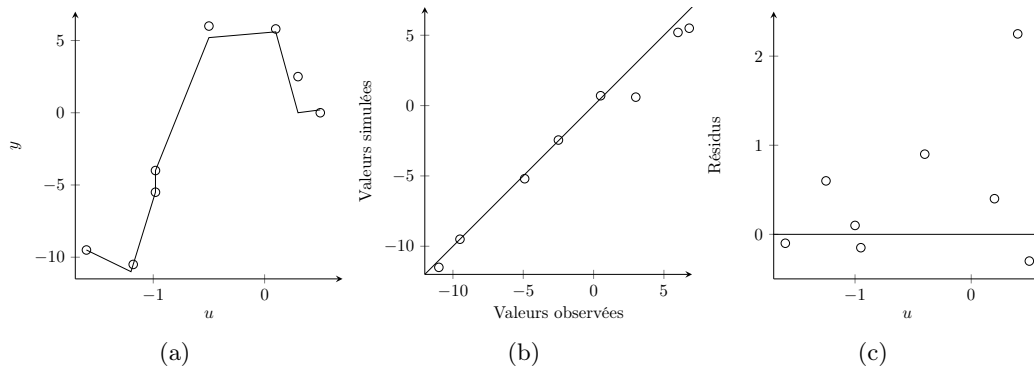


FIGURE 10.12 – (a) Graphe des valeurs simulées et observées en fonction de la variable explicative u (b) Graphe des valeurs simulées vs observées (c) Graphes des résidus $y^{obs} - \hat{y}$ en fonction de u

10.5.1.2 Evaluation via des critères quantitatifs

Plusieurs critères seront proposés et analysés dans ce paragraphe.

Définition 10.5.1

Le *biais* B est la moyenne des écarts entre les valeurs mesurées et les valeurs simulées :

$$B = \frac{1}{N} \sum_{i=1}^N (y_i^{obs} - \hat{y}_i),$$

avec N le nombre des mesures.

Un biais négatif indique une tendance du modèle à la sous-prédiction et réciproquement. Sur l'exemple de la figure 10.12, le biais vaut 0,48 ce qui est confirmé par la figure : on voit que le modèle a une tendance à systématiquement prédire des valeurs moins élevées que les données observées. Le biais seul reste un critère insuffisant car, si un 'bon' modèle doit avoir un biais faible, la réciproque n'est pas vraie : un biais faible peut être obtenu par un équilibre entre des erreurs de grande amplitude successivement négatives ou positives.

Cet inconvénient peut être surpassé en utilisant le critère suivant.

Définition 10.5.2

L'*erreur quadratique moyenne* (en anglais *Mean Squared Error*) est définie comme la moyenne des écarts au carré entre les valeurs mesurées et les valeurs simulées :

$$MSE = \frac{1}{N} \sum_{i=1}^N (y_i^{obs} - \hat{y}_i)^2$$

On utilise aussi souvent deux critères dérivés de l'erreur quadratique moyenne :

— sa racine carrée :

$$RMSE = \sqrt{MSE}$$

qui a l'avantage de s'exprimer dans la même unité que y ;

— sa valeur relative :

$$RRMSE = \frac{RMSE}{\bar{y}^{obs}}$$

qui correspond à la diviser par $\bar{y}^{obs} = \frac{1}{N} \sum_{i=1}^N y_i^{obs}$ et permet de comparer des erreurs fondées sur des jeux de données différents ou entre des variables de différents ordres de grandeur.

On peut enfin citer une mesure d'évaluation normalisée.

Définition 10.5.3

Le coefficient de détermination ou *efficience* se calcule de la manière suivante :

$$R^2 = 1 - \frac{\sum_{i=1}^N (y_i^{obs} - \hat{y}_i)^2}{\sum_{i=1}^N (y_i^{obs} - \bar{y}^{obs})^2}$$

Le coefficient R^2 prend une valeur maximale de 1 pour un modèle reproduisant parfaitement les observations. Il vaut 0 pour un modèle qui consisterait à utiliser la valeur moyenne des observations pour prédire toutes les situations. Il n'y a pas de borne inférieure a priori : un modèle peut être un pire prédicteur que la moyenne des données ($R^2 < 0$).

Dans de nombreux cas, il sera intéressant de fournir plusieurs des critères listés ci-dessus, afin de permettre un diagnostic fin des sources d'erreur. Il existe également plusieurs décompositions de l'erreur quadratique moyenne qui peuvent être employées dans ce but.

En cas de sortie multi-variée, différentes stratégies sont possibles, parmi lesquelles notamment : soit considérer les erreurs associées à chaque variable indépendamment, soit définir une erreur moyenne. On privilégie en ce cas les critères relatifs afin de prendre en compte les différences possibles d'ordre de grandeur. Si les sorties sont par exemple des distributions ou des nombres d'items dans différentes classes (par exemple nombre de fruits selon leur calibre), on pourra utiliser comme critère les statistiques utilisées pour les tests d'hypothèse de concordance entre distributions comme par exemple celui de Kolmogorov-Smirnov (voir cours de *Statistiques*).

Remarque 10.5.2. *Si certaines situations comportent beaucoup plus d'observations que d'autres, il peut être intéressant de faire d'abord une moyenne des erreurs par situations, afin d'éviter une influence trop prédominante de la situation comportant le plus d'observations.*

10.5.2 Sélection de modèle

On se pose à présent la question de choisir entre plusieurs modèles disponibles, pour lesquels on a au préalable estimé les paramètres et que l'on a évalués. Une première idée naturelle et naïve serait de se servir du critère des moindres carrés optimisé comme base de cette décision : plus il est faible, plus cela veut dire que les simulations 'collent' aux données et meilleur serait le modèle. Le problème de cette démarche est le risque de sur-apprentissage : l'ajout de paramètres supplémentaires, c'est-à-dire l'augmentation de la complexité du modèle, permet de manière mécanique un meilleur ajustement aux données, alors que confrontés à de nouvelles données ou à d'autres situations, il sera très peu performant. Une citation célèbre attribuée à John von Neumann (1903-1957) permet d'illustrer ce concept : "with four parameters I can fit an elephant and with five I can make him wiggle his trunk" ⁶.

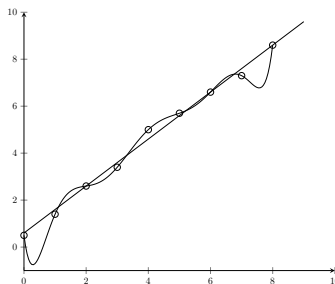


FIGURE 10.13 – Deux modèles candidats : un modèle linéaire et un modèle polynomial

La figure 10.13 illustre ce problème sur la comparaison entre un *modèle polynomial* et un *modèle linéaire* : le modèle polynomial permet de reproduire exactement les données disponibles. Mais avec

6. Certains s'y sont bien sûr essayés et cet exercice a même donné lieu à publications (voir par exemple Mayer et al. (2010)). En réalité, utiliser 4 paramètres complexes correspond à 8 paramètres réels !

des données issues d'une nouvelle collecte, qui seraient légèrement différentes du fait du bruit de mesures, on peut imaginer que les prédictions seraient bien moins bonnes. Le modèle linéaire, en revanche, sera un peu moins bon sur l'ajustement, mais plus robuste sur la prédiction en cas de nouvelles données ou de nouvelles situations.

Nous avons donc besoin de critères pour comparer des modèles de complexités différentes, c'est-à-dire ayant des nombres de paramètres différents. Nous décrivons ici deux approches différentes pour traiter ce problème.

10.5.2.1 Première approche : utilisation de critères pénalisés

La première approche consiste à pénaliser les modèles les plus complexes en introduisant un terme supplémentaire qui dépendra du nombre de paramètres à estimer. Différents critères existent. L'un des plus classiques, issu du domaine de la théorie de l'information, s'appuie sur la définition d'une notion de distance entre un 'vrai' modèle sous-jacent au système réel et un modèle candidat à le représenter. Il a été proposé en 1985 par H. Akaike qui l'avait d'abord nommé "An Information Criterion" (AIC) : l'acronyme est resté, mais sa signification actuelle est donnée dans la définition suivante.

Définition 10.5.4 (*Critère d'information d'Akaike (AIC)*)

Sous les hypothèses que les résidus sont indépendants et distribués de manière gaussienne $\mathcal{N}(0, \sigma^2)$, ce critère s'écrit à une constante près :

$$AIC = N \ln(J_{min}) + 2k$$

où k est le nombre de paramètres estimés du modèle et $J_{min} = \min_{\theta} \sum_{n=1}^N (y_n^{obs} - y(\mathbf{u}_n, \theta))^2$ est le minimum atteint par la somme des carrés des résidus, c'est-à-dire le critère (10.4.1) minimisé de la méthode des moindres carrés utilisée pour l'estimation.

Ce critère est utilisé pour classer des modèles, dès lors que leurs paramètres ont été estimés sur les **mêmes données expérimentales**. Le modèle choisi est celui qui aura la plus faible valeur d'AIC. La suppression d'une constante ne change donc rien à ce classement. En revanche, il ne dit rien de la qualité absolue du modèle : il ne rendrait ainsi pas compte du fait qu'aucun des modèles candidats ne produit de bons ajustements.⁷

Remarque 10.5.3. Lorsque le rapport du nombre de données au nombre de paramètres n'est pas suffisamment élevé, ce qui se traduit en pratique par la condition $N/k < 40$, il peut être démontré qu'il convient d'ajouter un terme de correction au critère AIC : $AICc = AIC + \frac{2k(k+1)}{N-k-1}$, sans quoi le critère tend à sélectionner des modèles trop complexes Hurvich and Tsai (1995). On voit que le terme de correction AICc tend vers AIC lorsque N devient grand, ce qui est attendu.

10.5.2.2 Deuxième approche : évaluer les capacités prédictives des modèles

Une autre façon d'éviter les problèmes de sur-ajustement est de tester la capacité du modèle à prédire une sortie pour des situations n'ayant pas servi à l'estimation de ses paramètres. Alors que les critères présentés en paragraphe 10.5.1 permettent d'évaluer l'ajustement du modèle aux données disponibles et utilisées pour son identification, on s'intéresse ici à essayer de prévoir ses performances face à de futures nouvelles données, non encore acquises.

7. Le cours de *Statistiques* permettra de découvrir l'utilisation de la fonction de vraisemblance $\mathcal{L}$ qui fournit une formulation plus générale de ce critère : $AIC = -2\ln(\mathcal{L}) + 2k$. La dérivation de l'expression donnée ici, sous les hypothèses de résidus gaussiens à variance constante, se fait assez simplement en utilisant l'expression de l'estimateur de la variance des résidus $\hat{\sigma}^2 = J/N$.

Définition 10.5.5

La *capacité prédictive* d'un modèle est quantifiée par le critère d'*erreur quadratique moyenne de prédiction* (en anglais *Mean Squared Error of Prediction*) :

$$MSEP(\hat{\theta}) = \mathbb{E}_{(U,Y)} \left[(Y^{obs} - f(U, \hat{\theta}))^2 \right] \quad (10.5.1)$$

Si la variable de sortie dépend du temps, on utilisera l'*erreur quadratique moyenne intégrée de prédiction* :

$$IMSEP(\hat{\theta}) = \mathbb{E}_{(U,Y)} \left[\int (Y^{obs}(t) - f(t, U, \hat{\theta}))^2 dt \right] \quad (10.5.2)$$

ou bien son équivalent discret en remplaçant l'intégrale par une somme.

Il s'agit de l'espérance du carré des différences entre les observations et les sorties simulées correspondantes, toutes deux vues comme des variables aléatoires. La distribution jointe de (U, Y) doit être spécifiée en fonction de :

- l'ensemble des situations possibles dans lesquelles on souhaite utiliser le modèle, c'est-à-dire sur une distribution de la variable d'entrée U qu'on se donne ;
- la variabilité des observations associées (bruit de mesure), pour la distribution de Y .

On note que $\hat{\theta}$ est ici vu comme fixé : il s'agit en toute rigueur d'une espérance conditionnelle⁸. L'évaluation de la capacité prédictive d'un même modèle pourra être fortement différente selon la distribution-cible choisie des situations pour U , elle-même en lien avec l'objectif du modèle, dont dépend intrinsèquement l'évaluation. Par exemple, si l'objectif du modèle est la prédiction de rendement d'une culture en situation irriguée, il ne devra pas être confronté à des situations de stress hydrique.

Remarque 10.5.4. Utiliser la racine carrée d'erreur quadratique moyenne de prédiction $RMSEP = \sqrt{MSEP}$ permet d'obtenir un critère dont l'unité est la même que celle de Y .

On voit que l'on se rapproche beaucoup de la définition de l'erreur quadratique moyenne⁹ (10.5.2) et il pourrait être tentant d'estimer l'erreur quadratique moyenne de prédiction MSEP par l'erreur quadratique moyenne MSE. Cependant, les N observations disponibles pour calculer la MSE pourraient ne pas être représentatives de la population de situations ciblées pour l'utilisation du modèle. Et, surtout, l'erreur quadratique moyenne va généralement sous-estimer l'erreur quadratique moyenne de prédiction MSEP : on conçoit bien que l'erreur entre une prédiction et des données utilisées pour développer le modèle va être la plupart du temps plus faible que l'erreur qui sera obtenue en confrontant le modèle à de nouvelles situations.

En pratique, deux cas peuvent se présenter pour estimer l'erreur quadratique moyenne de prédiction MSEP :

- On dispose d'un nouveau jeu de données, indépendant de celui utilisé pour l'identification du modèle. Dans ce cas, on obtient une approximation de la MSEP simplement par l'erreur quadratique moyenne du nouvel échantillon :

$$MSEP(\hat{\theta}) \approx MSE^{new} = \frac{1}{N} \sum_{i=1}^N \left(y_i^{new\ obs} - f(u_i^{new}, \hat{\theta}) \right)^2$$

où on a noté $(u_i^{new}, y_i^{new\ obs})_{i \in \{1, \dots, N\}}$ les N nouvelles données.

- Le cas le plus courant est celui où on ne dispose pas de données supplémentaires à celles utilisées pour l'identification du modèle. On va alors procéder à une *validation croisée*. Il

8. L'annexe E contient un rappel sur les probabilités.

9. Pour obtenir une approximation de l'espérance dans l'expression (10.5.1), on utilise classiquement l'estimateur de la moyenne empirique (qui sera abordé dans le cours de *Statistiques*) sur N situations :

$$\mathbb{E} \left[(Y - f(U, \hat{\theta}))^2 | \hat{\theta} \right] \approx \frac{1}{N} \sum_{i=1}^N (Y_i - f(U_i, \hat{\theta}))^2.$$

Le cours de *Statistiques* montrera que cet estimateur est sans biais, convergent et asymptotiquement normal.

en existe plusieurs protocoles, parmi lesquels celui du *leave-one-out* : pour $i \in \{1, \dots, N\}$, on estime les paramètres du modèle en écartant la i -ème observation, puis on calcule l'erreur de prédiction sur cette dernière. En faisant la moyenne de ces termes, on obtient ainsi l'estimation :

$$MSEP(\hat{\theta}) \approx MSEP_{cv} = \frac{1}{N} \sum_{i=1}^N \left(y_i^{obs} - f(u_i, \hat{\theta}_{-i}) \right)^2,$$

où $\hat{\theta}_{-i}$ désigne le vecteur de paramètres obtenus via une estimation sur toutes les données sauf la i -ème.

L'erreur quadratique moyenne de prédiction MSEP peut être utilisée comme critère de sélection de modèle car elle ne favorise pas forcément les modèles les plus complexes : le sur-apprentissage est mécaniquement pénalisé par l'évaluation de la prédiction.

Chapitre 11

Analyse d'incertitude et de sensibilité

L'annexe E présente quelques notions utiles pour ce chapitre.

11.1 Contexte et objectifs

11.1.1 Notations

On se donne un modèle déterministe défini par :

$$Y = f(X, \mathbf{d}) \quad (11.1.1)$$

où $Y \in \mathbb{R}$ est une variable de sortie, que l'on va ici supposer être de dimension 1 par souci de clarté : les résultats peuvent assez facilement être étendus à des dimensions supérieures. Si le système est dynamique, on en fera l'analyse à un point de temps bien choisi et on la réitérera pour d'autres points de temps si besoin. Cette expression signifie que l'on considère dans cette partie le modèle comme une relation reliant un certain nombre de facteurs à une sortie. Le terme générique de *facteur* est souvent employé dans les domaines de l'*analyse d'incertitude* et de *sensibilité* car il peut s'agir à la fois de paramètres ou de variables d'entrée moyennes. On décompose ces facteurs en deux catégories : $X = (X_1, \dots, X_k) \in \mathbb{R}^{1 \times k}$ est le vecteur ligne des facteurs sur lesquels on suppose une incertitude et le vecteur ligne $\mathbf{d} = (d_1, \dots, d_{n_d}) \in \mathbb{R}^{1 \times n_d}$ est le vecteur des facteurs fixés (i.e. pris constants). On note que l'on adopte dès à présent la notation de Y et X en majuscule car elles seront, dans la suite, considérées comme des variables aléatoires.

11.1.2 Analyse d'incertitude : définition et objectifs

Définition 11.1.1

L'*analyse d'incertitude* vise à caractériser l'incertitude¹ associée soit directement à une grandeur mesurée², soit à une sortie d'un modèle en fonction de l'incertitude sur certains de ses composants.

Actuellement, avec l'évolution des moyens de simulation et l'acquisition de nouvelles données expérimentales, des modèles de plus en plus complexes sont développés dans de nombreux domaines et utilisés en tant qu'outils d'intégration de connaissances, de prédiction, d'aide à la décision ou de développement de lois de commande³. Il n'est pas rare, par exemple en biologie, de rencontrer des modèles dont le nombre de paramètres dépasse la cinquantaine. Cette complexité pose des problèmes pour la détermination des valeurs des paramètres lorsqu'il s'agit de reproduire et prédire les propriétés d'un phénomène réel : tous les paramètres ne peuvent généralement pas être connus avec précision. Cette complexité et les incertitudes associées peuvent ainsi limiter la confiance que l'on peut avoir dans les réponses du modèle et réduire son utilisation.

1. Le sens de cette caractérisation reste à définir ultérieurement.

2. On est alors dans le champ de la *métrologie*.

3. Des méthodes de synthèse de lois de commande seront abordées dans le cours d'*Automatique*.

Bien que ce précepte ne soit pas toujours vérifié en pratique, notamment dans les médias destinés au grand public, toute prédiction réalisée à l'aide d'un modèle devrait systématiquement être accompagnée d'un indicateur de son incertitude. L'analyse d'incertitude est en effet cruciale avant toute prise de décision fondée sur des résultats de simulation et est parfois même une obligation réglementaire dans certains secteurs sensibles (par exemple centrales nucléaires, rejets de CO₂, etc.). Il s'agit par exemple d'être sûr qu'une affirmation fondée sur un modèle peut affronter la critique, c'est-à-dire que les conclusions restent valables même si certains facteurs d'entrée varient (i.e. il faut prouver qu'il n'y a pas d'exploration incomplète des hypothèses d'entrée). L'estimation de cette incertitude est un exercice délicat : (i) la sous-estimer peut être dangereux puisque des situations non anticipées seraient susceptibles de se produire ; (ii) un excès de précautions conduisant à la sur-estimer est également néfaste en ce qu'il limite fortement le degré de confiance porté au résultat et finalement l'utilité même du modèle.

Dans le cadre d'un problème de modélisation, les *incertitudes* ou *erreurs* peuvent être présentes à différents niveaux :

1. dans la formulation même du modèle, qu'elle soit mathématique ou algorithmique : provenant par exemple des approximations, des incertitudes structurales ;
2. dans l'implémentation numérique du modèle : ayant comme sources la précision des algorithmes, les seuils de convergence, l'aspect aléatoire rencontré dans les modèles stochastiques ;
3. dans la détermination des paramètres ou variables d'entrée du modèle. La modélisation d'objets complexes implique souvent de nombreux paramètres empiriques, c'est-à-dire s'appuyant sur l'expérience et non sur la théorie. Dès lors qu'il s'agit de paramètres découlant de mesures, ils sont incertains.

Dans cette partie, nous nous intéressons essentiellement aux incertitudes issues du point 3.

11.1.3 Analyse de sensibilité : définition et objectifs

Pour des modèles complexes, l'influence de chaque paramètre sur la variable de sortie n'est pas simple à déterminer et est souvent diluée au milieu de celles de l'ensemble des paramètres : il est donc difficile de trier les paramètres indispensables et ceux moins utiles. L'*analyse de sensibilité*, illustrée en figure 11.1, se définit comme une sorte de réciproque à l'analyse d'incertitude. Cette figure montre que : (i) dans la démarche d'analyse d'incertitude on cherche à caractériser l'incertitude sur la variable de sortie Y du modèle selon celle définie sur ses facteurs d'entrée ; (ii) dans l'analyse de sensibilité, il s'agit de quantifier la contribution des incertitudes sur les facteurs d'entrée à l'incertitude sur Y .

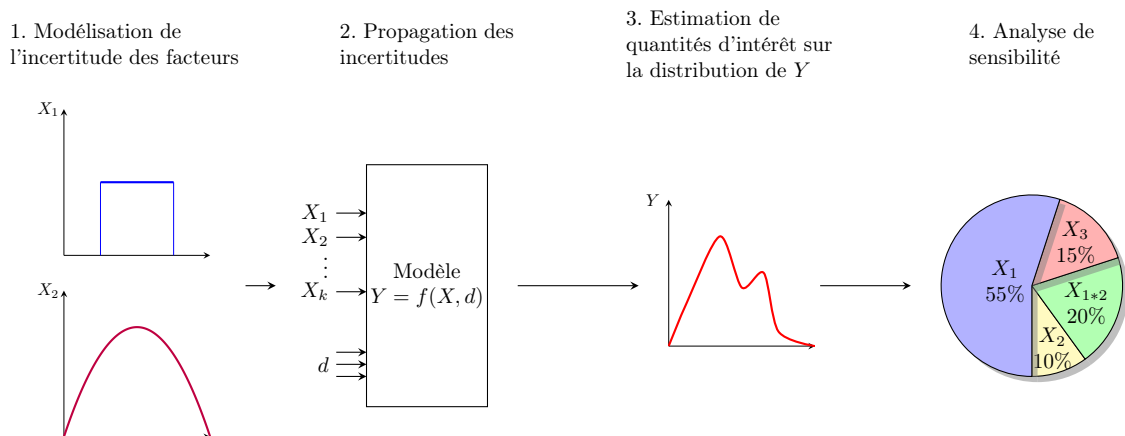


FIGURE 11.1 – Illustration de la démarche d'analyse d'incertitude et de sensibilité

Deux familles de méthodes seront abordées : les *méthodes locales* (les plus connues mais généralement moins informatives) et les *méthodes globales*.

Définition 11.1.2

L'*analyse de sensibilité globale* est l'étude de la façon dont les incertitudes dans les sorties d'un modèle (numériques ou autres) peuvent être attribuées à différentes sources d'incertitudes parmi les entrées du modèle.

Saltelli et al. (2000) définissent quatre types d'objectifs à l'analyse de sensibilité :

1. *Factors Prioritisation* (FP) : on veut faire un pari éclairé sur le facteur qu'il faudrait fixer pour avoir la plus grande réduction de l'incertitude de la sortie. Cela peut permettre de désigner des objectifs prioritaires en termes de recherche ou d'expérimentation sur les facteurs sur lesquels il est le plus important d'obtenir des informations plus précises.
2. *Factors Fixing* (FF) : on parcourt les paramètres pour identifier ceux qui n'ont aucune ou très peu d'influence, que l'on pourra fixer à n'importe quelle valeur dans leur intervalle d'incertitude sans effet significatif sur la sortie. Cela peut faciliter le processus d'estimation paramétrique, notamment en cas de non-identifiabilités. Cela peut permettre aussi de simplifier un modèle, si l'on décide de ne pas prendre en compte les processus ayant peu d'influence sur la sortie d'intérêt.
3. *Variance Cutting* (VC) : l'objectif consiste à trouver l'ensemble minimal de paramètres à fixer pour atteindre une réduction donnée de l'incertitude de sortie. Ce problème peut se poser lorsque le modèle doit être utilisé pour de la prise de décision ou en analyse de risque.
4. *Factors Mapping* (MP) : le but est de trouver les facteurs responsables de réalisations de la sortie Y dans une région donnée.

On peut y ajouter, pour un modélisateur, les objectifs de mettre à jour les interactions existantes entre des entrées dans leur effet sur une sortie et de faciliter l'estimation paramétrique en identifiant quels facteurs auront une influence réelle sur la fonction de coût utilisée (en faisant l'analyse de sensibilité sur cette fonction de coût). On peut également s'en servir pour juger de la cohérence d'un modèle et sa pertinence, en comparant ses réponses en terme de sensibilité à ce que l'on attend d'après nos connaissances du phénomène réel.

Certaines des méthodes d'analyse de sensibilité présentées ci-après peuvent s'appliquer indifféremment que l'on connaisse la structure du modèle ou pas. Elles peuvent donc s'appliquer soit pour tester un modèle que l'on est soi-même en train de développer, soit pour tester un programme de simulation transmis par quelqu'un d'autre, même si l'on n'a pas accès au code ou aux équations du modèle.

Qu'il s'agisse d'analyse d'incertitude ou de sensibilité, un ingrédient nécessaire en première étape est la caractérisation de l'incertitude des différents facteurs considérés : on parle de modélisation de l'incertitude.

11.2 Modélisation de l'incertitude

11.2.1 Terminologie

Lorsqu'on rend compte du résultat d'un mesurage d'une grandeur physique ou, dans le cadre d'un modèle, du résultat d'une simulation, il faut l'assortir d'une indication quantitative permettant à un utilisateur d'en estimer la fiabilité. Sans une telle indication, les résultats de mesure ne peuvent pas être comparés, soit entre eux, soit par rapport à des valeurs de référence données dans une spécification ou une norme. Cette indication est désignée par le terme *incertitude*, qui caractérise le doute raisonnable que l'on peut avoir sur le résultat.

On donne ici les définitions de quelques termes⁴ issus du champ de la *Métrologie*, qui sont utiles pour désigner de manière rigoureuse les concepts associés à l'*analyse d'incertitude* :

— *mesurage* : ensemble d'opérations ayant pour but de déterminer une valeur d'une grandeur ;

4. Cette terminologie est issue des références suivantes BIPM et al. (2008) et BIPM et al. (1993), publiées par sept organisations contributrices : le Bureau International des Poids et Mesures (BIPM), la Commission Electrotechnique Internationale (IEC), la Fédération Internationale de Chimie Clinique (FICC), l'Organisation internationale de normalisation (ISO), l'Union Internationale de Chimie Pure et Appliquée (IUPAC), l'Union Internationale de Physique Pure et Appliquée (UIPPA) et l'Organisation Internationale de Métrologie Légale (OIML).

- *mesurande* : grandeur particulière soumise à mesurage. La définition du mesurande peut nécessiter des indications relatives à des grandeurs telles que le temps, la température et la pression (par exemple la pression de vapeur d'un échantillon donné d'eau à 20°C) ;
- *valeur vraie* d'une grandeur mesurable : valeur, indéterminée par nature, que l'on obtiendrait par un mesurage parfait ;
- *erreur de mesure* : résultat d'un mesurage moins une valeur vraie du mesurande ;
- *exactitude de mesure* : étroitesse de l'accord entre le résultat d'un mesurage et une valeur vraie du mesurande ;
- *répétabilité* des résultats de mesurage : étroitesse de l'accord entre les résultats des mesurages successifs du même mesurande, mesurages effectués dans la totalité des mêmes conditions de mesure (même mode opératoire, même observateur, même instrument de mesure utilisé dans les mêmes conditions, même lieu, répétition durant une courte période de temps) ;
- *reproductibilité des résultats de mesure* : étroitesse de l'accord entre les résultats des mesurages du même mesurande, mesurages effectués en faisant varier les conditions de mesure (principe de mesure, méthode de mesure, observateur, instrument de mesure, étalon de référence, lieu, conditions d'utilisation, temps).

11.2.2 Sources d'incertitude sur les facteurs

Différentes raisons peuvent justifier de considérer un facteur comme incertain.

1. Il ne peut pas être mesuré et le modélisateur dispose seulement d'un certain degré de connaissances à son sujet.
2. Il peut être mesuré mais ces mesures sont entachées d'erreurs qui peuvent être de deux sortes :
 - *Erreur systématique*. C'est la composante de l'erreur de mesure qui, dans des mesurages répétés, demeure constante ou varie de façon prévisible. On peut la définir comme la moyenne qui résulterait d'un nombre infini de mesurages du même mesurande, effectués dans les conditions de répétabilité, moins une valeur vraie du mesurande. Comme la valeur vraie, les erreurs systématiques et leurs causes ne peuvent pas être connues complètement. Elles influencent l'exactitude de la mesure en créant un biais, par exemple par l'emploi d'appareils mal-étalonnés ou par un phénomène négligé à tort (par exemple l'influence du champ magnétique terrestre dans une mesure magnétique).
 - *Erreur aléatoire* ou *erreur accidentelle* ou *dispersion statistique*. Elle se révèle par des mesures répétées d'une même grandeur et peut être liée à des phénomènes perturbateurs (bruit de fond électronique, sensibilité d'un appareil à des changements infimes de températures, etc.) ou à la nature aléatoire du phénomène mesuré. Elle affecte la précision de la mesure. Elle peut être issue par exemple du fait de l'observateur (sûreté du geste ou de l'acuité visuelle par exemple) ou bien du fait de phénomènes aléatoires inhérents au processus de mesure.

Pour recenser de manière exhaustive les possibles origines des incertitudes, la *méthode dite des '5M'*, imaginée par Kaoru Ishikawa et classiquement employée aujourd'hui dans le domaine de l'analyse qualité, permet de rechercher les différentes causes possibles d'un problème et de les représenter de manière synthétique en un diagramme en forme d'arête de poisson pour matérialiser de manière structurée le lien entre les causes et leur effet (défaut, panne, dysfonctionnement, erreur de mesure, etc.). Ishikawa classe les différentes causes d'un problème en cinq grandes familles :

- *Matière* : différents consommables utilisés, matières premières, emballage, etc.
- *Milieu* : lieu de travail (disposition, localisation, etc.), son aspect, son organisation physique (éclairage, température, bruits divers), etc.
- *Méthode* : procédure, flux d'information, etc.
- *Matériel ou moyen* : équipements, machines, outillages, pièces de rechange, etc.
- *Main d'oeuvre* : ressources humaines, qualifications du personnel, etc.

Une version étendue ajoute deux items :

- *Management* : méthodes d'encadrement, style de commandement, délégation, organigramme imprécis, etc.
- *Moyens financiers* : budget alloué, coûts, etc.

La figure 11.2 permet d'illustrer l'application de la méthode des '5M' sur un exemple tiré de Pernot (2018) d'une analyse de plomb dans une eau. Les résultats de mesure dépendent donc des cinq familles, dont les éléments sont explicités sur la figure.

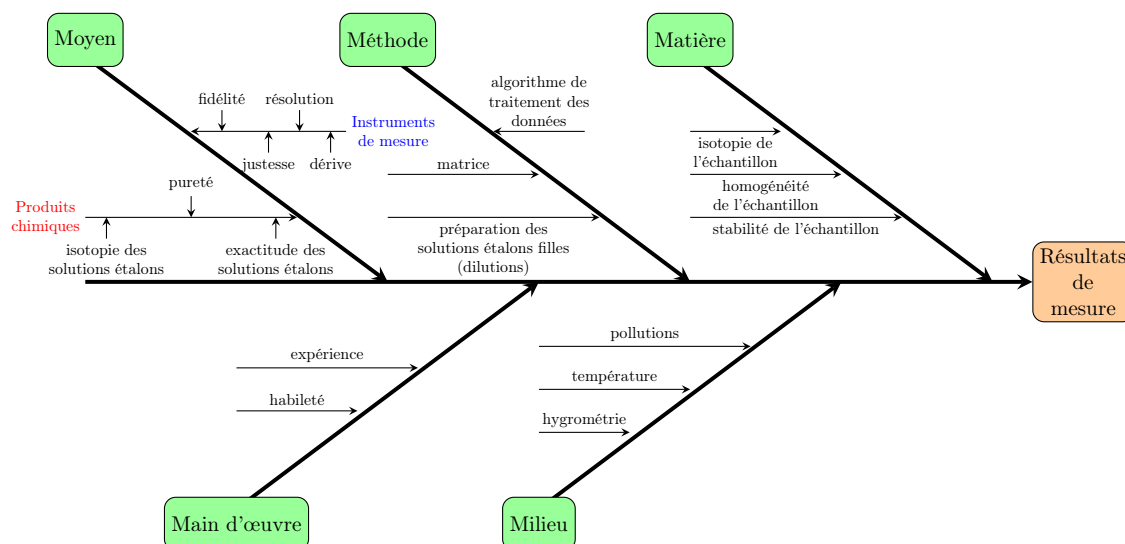


FIGURE 11.2 – Diagramme des '5M' sur l'exemple d'une analyse de plomb dans une eau

11.2.3 Définition générale de l'incertitude

Différentes théories de l'incertitude existent. On adopte ici la définition du 'Guide to the expression of uncertainty in measurement' (GUM) (voir BIPM et al. (2008)) :

Définition 11.2.1

L'incertitude de mesure représente le terme associé au résultat d'un mesurage, qui caractérise la dispersion des valeurs qui pourraient raisonnablement être attribuées au mesurande.

Remarque 11.2.1. *L'incertitude de mesure est une expression du fait que, pour un mesurande donné et un résultat de mesure donné de ce mesurande, il n'y a pas une valeur, mais un nombre infini de valeurs dispersées autour du résultat, qui sont compatibles avec toutes les observations et les données, avec la connaissance que l'on peut avoir du monde physique et qui peuvent être attribuées au mesurande avec des degrés de crédibilité divers.*

Pour estimer l'incertitude associée à un résultat de mesure, on suppose que le résultat d'un mesurage a été corrigé pour tous les effets systématiques reconnus comme significatifs et qu'on a fait tous les efforts pour leur identification. L'incertitude du résultat d'un mesurage reflète l'impossibilité de connaître exactement la valeur du mesurande. Le résultat d'un mesurage après correction des effets systématiques reconnus reste encore seulement une estimation de la valeur du mesurande en raison de l'incertitude provenant des effets aléatoires et de la correction imparfaite du résultat pour les effets systématiques.

Remarque 11.2.2. *L'incertitude peut être représentée par un paramètre ou une loi de probabilité. Un exemple de paramètre permettant de représenter l'incertitude est un écart-type⁵ (ou un multiple de celui-ci) ou la demi-largeur d'un intervalle de niveau de confiance déterminé.*

5. On parle alors d'*incertitude-type*.

Ces considérations nous amènent naturellement à adopter une démarche probabiliste dans laquelle des *facteurs* (pouvant être des entrées, des paramètres ou bien des conditions initiales ou aux limites, par exemple), sont considérés comme des variables aléatoires suivant des lois de probabilité à déterminer : on peut attribuer un degré de confiance $\mathbb{P}[X = x]$ à chaque valeur possible x d'une quantité incertaine X . L'état des connaissances correspondant est décrit au mieux par une distribution sur l'ensemble des valeurs possibles du mesurande. La figure 11.3 illustre cette procédure pour des variables continues : les densités de probabilité des facteurs X_1 , X_2 et X_3 sont représentées par les fonctions g_{X_1} , g_{X_2} , g_{X_3} respectivement, et on note g_Y la densité de probabilité de Y , dont la caractérisation est l'objectif de l'analyse d'incertitude.

Remarque 11.2.3. Des approches alternatives, que nous ne traiterons pas dans ce cours, peuvent être considérées pour représenter l'incertitude (par exemple l'analyse par intervalles ou par la théorie des ensembles flous (en anglais *fuzzy set theory*), etc.).

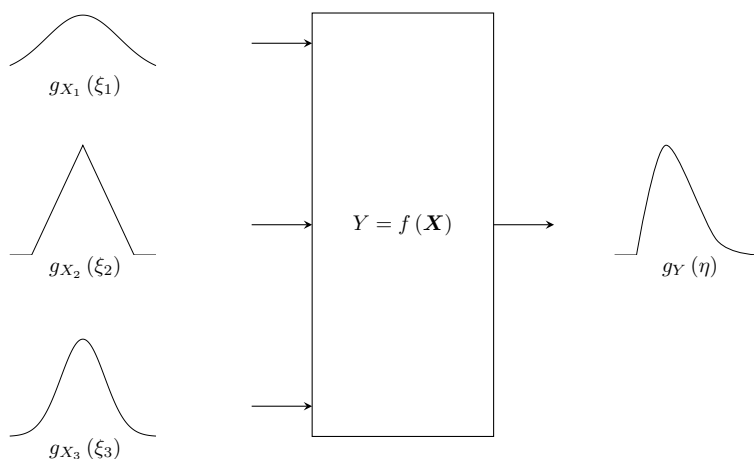


FIGURE 11.3 – Principe de l'analyse d'incertitude

On peut dès à présent souligner que la modélisation des incertitudes, c'est-à-dire ici le choix d'une distribution de probabilité (comme illustré par la figure 11.3) pour un facteur donné, est l'un des problèmes-clé lors de la mise en œuvre des méthodes d'analyse d'incertitude et de sensibilité.

Selon la typologie du BIPM et al. (2008), on peut distinguer deux familles d'approches, nommées A et B, qui seront présentées par la suite.

11.2.4 Modélisation des incertitudes de type A

Les *incertitudes de type A* portent sur des grandeurs qui ont pu être observées expérimentalement de manière répétée. On suppose disposer d'un n -échantillon *indépendamment et identiquement distribué* (i.i.d.) $X = (X_1, \dots, X_n)$ de loi marginale $\mathbb{P}_X$ sur $\mathbb{R}$ (on se place pour simplifier dans le cas d'une variable unidimensionnelle) et on suppose de plus que $\mathbb{P}_X$ admet une densité par rapport à la mesure de Lebesgue que l'on notera f , inconnue. Deux approches sont possibles :

- *approche directe*, dite *non-paramétrique* : on construit une densité de probabilité approchée à partir de l'échantillon : histogrammes, méthode des noyaux. Cette famille de méthodes est applicable lorsque l'on dispose d'un échantillon de grande taille.
- *approche inverse* ou *paramétrique* : on fait l'hypothèse d'une loi suivie par X et on estime les paramètres de sa densité de probabilité.

La figure 11.4 permet d'illustrer ces deux approches : on se donne une distribution de probabilité pour une variable aléatoire quelconque, grâce à laquelle on simule le tirage de 100 réalisations indépendantes qui vont jouer le rôle de données réelles. À partir de cet échantillon, les méthodes

non-paramétriques consistent à construire l'histogramme représenté, selon la méthodologie qui sera présentée en section 11.2.4.2, ou bien la densité continue estimée par la méthode des noyaux (courbe verte). L'application d'une méthode paramétrique consiste quant à elle à reconnaître une forme de densité gaussienne, par exemple, et à en estimer les moments afin de la caractériser (courbe bleue). On note que, du fait de la perte d'information liée à la taille finie de l'échantillon et à l'aspect stochastique des réalisations, toutes ces méthodes fournissent seulement des approximations de la 'vraie' densité ayant servi à générer les données (courbe rouge).

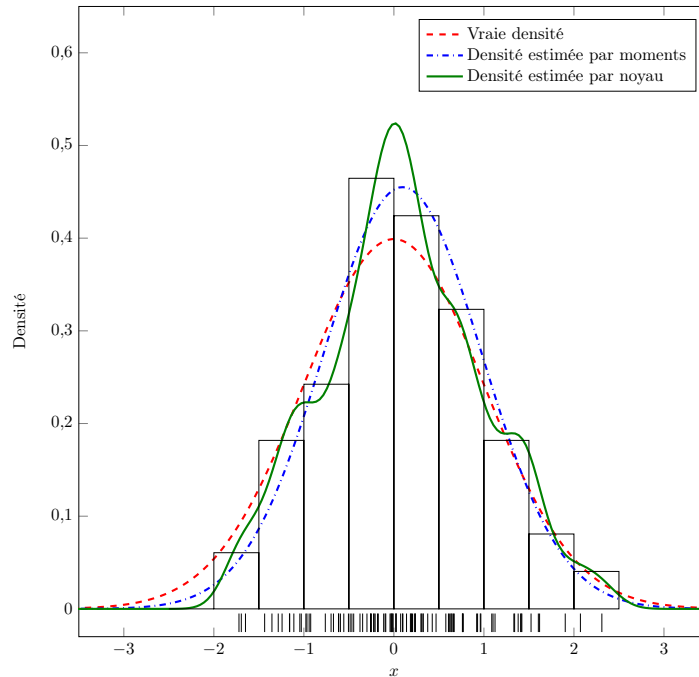


FIGURE 11.4 – Construction d'une densité de probabilité approchée à partir d'un échantillon de points : approches paramétrique ou non-paramétrique

Remarque 11.2.4. *Il n'est pas forcément nécessaire d'avoir une expression analytique de cette densité de probabilité. Selon les choix de méthodes présentées en section 11.3, on peut se contenter de son intervalle de définition, de ses moments ou bien de savoir en simuler des échantillons représentatifs. On va justement voir à présent comment estimer les moments d'une distribution (sans la caractériser de façon exhaustive).*

11.2.4.1 Estimation des moments d'une distribution

Si on a seulement besoin de caractériser l'échantillon via sa moyenne et sa dispersion, il n'est pas nécessaire de reconstruire toute sa distribution de probabilité. On peut utiliser les *estimateurs empiriques* suivants, pour un échantillon de variables aléatoires indépendantes et identiquement distribuées.

Définition 11.2.2

La *moyenne empirique* de N variables aléatoires X_i , avec $i \in \{1, \dots, N\}$ est donnée par :

$$\bar{X} = \frac{1}{N} \sum_{i=1}^N X_i \quad (11.2.1)$$

Remarque 11.2.5. *On rappelle que la notation $\bar{X}$ désigne la variable aléatoire, tandis que sa réalisation, c'est-à-dire la valeur numérique que l'on pourra calculer à partir des mesures effectuées,*

est notée $\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i$.

Définition 11.2.3

La *variance empirique* de la variable aléatoire X est calculée comme suit :

$$\hat{\sigma}_X^2 = \frac{1}{N-1} \sum_{i=1}^N (X_i - \bar{X})^2 \quad (11.2.2)$$

Le facteur $N-1$ permet d'obtenir un estimateur non-biaisé de la variance σ^2 .

Démonstration.

$$\begin{aligned} \hat{\sigma}_X^2 &= \frac{1}{N-1} \sum_{i=1}^N (X_i^2 - 2X_i\bar{X} + \bar{X}^2) \\ &= \frac{1}{N-1} \left(\sum_{i=1}^N X_i^2 - 2N\bar{X}\bar{X} + N\bar{X}^2 \right) \\ &= \frac{1}{N-1} \sum_{i=1}^N (X_i^2 - \bar{X}^2) \end{aligned}$$

d'où, en utilisant la linéarité de l'espérance et le fait que $\text{Var}[X] = \mathbb{E}[X^2] - \mathbb{E}[X]^2$:

$$\begin{aligned} \mathbb{E}[\hat{\sigma}_X^2] &= \frac{1}{N-1} \sum_{i=1}^N (\text{Var}[X_i] + \mathbb{E}[X_i]^2 - \text{Var}[\bar{X}] - \mathbb{E}[\bar{X}]^2) \\ &= \frac{1}{N-1} \sum_{i=1}^N (\sigma^2 + m^2 - \frac{\sigma^2}{N} - m^2) \\ &= \frac{1}{N-1} \sum_{i=1}^N \frac{\sigma^2(N-1)}{N} \\ &= \sigma^2 \end{aligned}$$

où on a noté m et σ^2 l'espérance et la variance de X (les 'vraies' valeurs). On peut noter que si on considérait m connu, alors un estimateur non-biaisé de la variance serait cette fois $\frac{1}{N} \sum_{i=1}^N (X_i - m)^2$. $\square$

Proposition 11.2.1

La variance de la moyenne empirique $\bar{X}$ peut être estimée par :

$$\hat{\sigma}_{\bar{X}}^2 = \frac{1}{N(N-1)} \sum_{i=1}^N (X_i - \bar{X})^2 \quad (11.2.3)$$

Démonstration. La variance de la moyenne empirique s'écrit en prenant en compte l'indépen-

dance des X_i :

$$\begin{aligned}\text{Var}[\bar{X}] &= \text{Var}\left[\frac{1}{N} \sum_{i=1}^N X_i\right] \\ &= \frac{1}{N^2} \left(\sum_{i=1}^N \text{Var}[X_i] \right) \\ &= \frac{1}{N^2} (N \cdot \sigma_X^2)\end{aligned}$$

Le paramètre σ_X^2 n'étant pas connu, on le remplace par son estimateur empirique de la définition 11.2.3, ce qui nous donne l'expression recherchée pour l'estimateur de la variance de $\bar{X}$. $\square$

Cette proposition nous indique le gain en précision que l'on peut avoir si on estime une grandeur via sa moyenne empirique sur des répétitions plutôt que de se contenter d'une seule mesure.

11.2.4.2 Estimation de la densité par approche non-paramétrique : histogramme et noyaux

Les *histogrammes* sont les estimations de densité les plus utilisés et sont, formellement, des estimateurs de la densité par des fonctions constantes sur des intervalles choisis par le modélisateur, qu'on appelle classes de l'histogramme. On se fixe un nombre de classes $J = J(N)$ qui peut dépendre du nombre de mesures N et on pose $h = \frac{1}{J}$. On décompose le support de X en une partition de J

classes $\bigcup_{j=1}^J C_{N,j}$ de même largeur h . Alors la fonction $\hat{f}_{N,h}(x)$ définissant l'histogramme s'écrit :

$$\hat{f}_{N,h}(x) = \frac{1}{h} \sum_{j=1}^J \hat{p}_{N,j} \mathbb{1}_{C_{N,j}}(x)$$

où

$$\hat{p}_{N,j} = \frac{1}{N} \sum_{i=1}^N \mathbb{1}_{C_{N,j}}(x_i)$$

représente la proportion d'observations tombant dans la $j^{\text{ème}}$ classe et $\mathbb{1}_{C_{N,j}}(x)$ est la fonction qui vaut 1 si $x \in C_{N,j}$ et 0 ailleurs. Le paramètre h contrôle le *niveau de 'lissage'* de l'estimation et doit être soigneusement choisi. Si le nombre de classes J est trop faible, on obtient des classes peu nombreuses et larges qui ne permettent pas d'approcher les variations éventuelles de f , c'est-à-dire que l'approximation de f aura un biais élevé : on parle de *sur-lissage*. A l'inverse, si les classes sont trop nombreuses et trop étroites (J trop élevé), on n'aura pas suffisamment d'observations par classe pour avoir une estimation fiable de leur proportion $p_{N,j}$ et la variance de ces estimations sera importante (très dépendante de l'échantillon tiré) : on voit trop de détails pour pouvoir tirer des conclusions générales, et ces détails sont sans doute plus liés à l'aléa de l'échantillonnage plutôt qu'à la réalité de la forme de f (*sous-lissage*). En estimant chacun de ces termes d'erreurs, une valeur optimale de h peut être obtenue comme résultat du compromis biais-variance de l'erreur d'approximation commise (on se réfère à Wasserman (2006) pour plus de précisions sur les méthodes de détermination de h).

La figure 11.5 présente une estimation par histogramme d'une densité simulée, avec un mélange de trois gaussiennes. Les données simulées sont les mêmes dans les deux cas et les histogrammes sont calculés pour deux J différents. Les classes peu nombreuses et larges obtenues avec $J = 4$ correspondent à un cas de *sur-lissage* (on ne voit plus de structure tant la densité est aplatie). Les classes trop nombreuses et trop étroites obtenues avec $J = 200$ correspondent à un cas de *sous-lissage* (le profil est 'bruité' et ne permet pas de retrouver la forme des trois gaussiennes).

L'histogramme possède des propriétés se rapprochant de celles d'une densité de probabilité : il est positif et son intégrale vaut 1. En revanche, il lui manque la continuité.

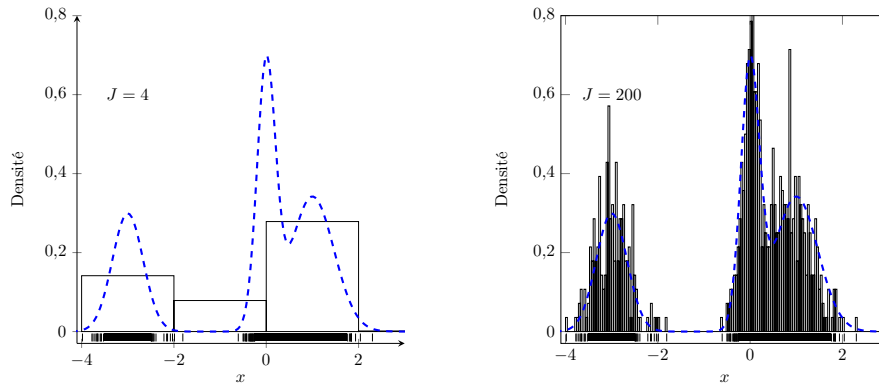


FIGURE 11.5 – Estimation par histogramme d'une densité simulée (en traits bleus), mélange de trois gaussiennes

La *méthode des noyaux* est une généralisation de l'approximation par histogramme qui pallie à ce problème. On va cette fois-ci choisir d'approximer la densité par la fonction ci-dessous, qui dépend du paramètre h :

$$\hat{f}_h(x) = \frac{1}{Nh} \sum_{i=1}^N \mathcal{K}\left(\frac{x - x_i}{h}\right)$$

où $\mathcal{K}$ vérifie des propriétés de régularité que l'on ne détaillera pas ici. On peut citer en exemple pour $\mathcal{K}$ la densité d'une fonction gaussienne standard (espérance nulle et variance unitaire). Comme dans le cas de l'histogramme, le rôle et le choix de valeur du paramètre h sont importants, comme illustré par la figure 11.6. Dans cette figure, l'estimateur à noyau de la densité f est en fait la moyenne (on divise par le nombre de courbes en cloche, i.e. 6). La variance des normales est posée à 0,5. Notons enfin que plus il y a d'observations dans le voisinage d'un point, plus sa densité est élevée.

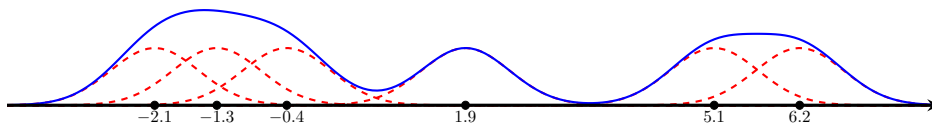


FIGURE 11.6 – Six courbes en cloche gaussiennes (rouge) et leur somme (bleu)

11.2.4.3 Estimation de la densité par approche inverse : estimation des paramètres de loi

Dans certains cas, on peut supposer que f est un élément d'un *espace paramétrique* $\{f_\theta, \theta \in \Theta \subset \mathbb{R}^p\}$, c'est-à-dire que l'on fait l'hypothèse que f se trouve dans une *famille de lois paramétriques*. Alors, estimer f se ramène à l'estimation du vecteur de paramètre θ par un certain $\hat{\theta}$ que l'on injecte dans $\hat{f}_\theta$. On pourra pour cela estimer certains moments ou certaines caractéristiques de la distribution à partir des données (en particulier ceux présentés en section 11.2.4.1) : les justifications de ces estimations, et surtout les propriétés⁶ de ces estimateurs seront vues en cours de *Statistiques*. Le choix de la famille de lois paramétriques se fait généralement après une première étude des données de l'échantillon, en utilisant notamment le graphe de l'histogramme pour identifier quel type de densité paramétrique serait le plus adapté. Cette *approche paramétrique* impose donc plus de contraintes que l'approche non-paramétrique (i.e. avec potentiellement un biais de modélisation plus élevé), mais se révèle souvent plus robuste si l'on dispose de peu de données. Les techniques de sélection de modèles vues en section 10.5.2 sont utiles pour déterminer une famille de fonctions à privilégier.

6. Sont-ce les meilleurs estimateurs que l'on peut construire pour ces paramètres ?

Exemple 11.2.1. On dispose de l'échantillon représenté en figure 11.7. L'histogramme nous laisse envisager l'hypothèse d'une loi uniforme, sachant que l'on n'a par ailleurs pas d'autre information que celle de voir que le support de l'échantillon paraît borné. On choisit les estimateurs min et max pour son intervalle et on en déduit que l'incertitude sur ce facteur est représentée par la loi uniforme $\mathcal{U}([\min_{i \in \{1, \dots, N\}} x_i, \max_{i \in \{1, \dots, N\}} x_i])$.

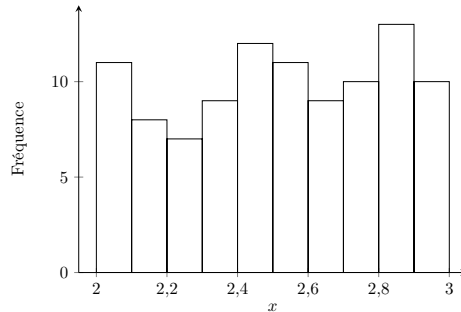


FIGURE 11.7 – Histogramme d'un échantillon

Remarque 11.2.6. On a supposé ici disposer d'un échantillon de données indépendantes et identiquement distribuées car les traitements statistiques en sont grandement simplifiés, mais il faut garder en mémoire que cette hypothèse pourrait ne pas être valide, par exemple en cas de mesures répétées dans le temps (par exemple, dans le cas d'une mesure de température : (i) si l'intervalle entre deux mesures est faible, des mesures consécutives seront corrélées et on a alors non-indépendance ; (ii) si on a une évolution temporelle, les données sont alors non-identiquement distribuées). Des méthodes statistiques plus poussées (par exemple l'analyse de séries temporelles) devraient alors être employées en toute rigueur ce qu'on ne traitera pas dans ce cours.

11.2.5 Modélisation des incertitudes de type B

La modélisation des incertitudes de type B concerne les grandeurs qui n'ont pas été observées expérimentalement, pour lesquelles on définit une densité de probabilité en s'appuyant sur :

- l'expérience ou la connaissance générale du comportement et des propriétés des matériaux et instruments utilisés ;
- les spécifications du fabricant ;
- les données des certificats d'étalonnage ;
- les valeurs de référence provenant d'ouvrages et manuels ;
- l'avis d'experts⁷, etc.

Si l'on ne dispose que d'informations partielles (par exemple une valeur moyenne et un écart-type), il existe une infinité de densités de probabilité pouvant les générer.

Théorème 11.2.1 (Principe du maximum d'entropie Jaynes (1957))

Le **principe du maximum d'entropie** spécifie que la loi de probabilité à choisir parmi un ensemble de lois possibles lors d'une modélisation d'incertitude d'un facteur X doit être celle qui maximise l'entropie de Shannon

$$S(X) = - \int_{\text{Supp}(x)} f(x) \ln(f(x)) dx = -\mathbb{E}[\ln(f(x))]$$

en tenant compte des contraintes imposées par les informations disponibles.

7. Le protocole de recueil d'avis d'experts doit comprendre un certain nombre d'étapes formelles qui ont pour but d'éviter autant que possible d'introduire des biais dans les estimations des experts : (1) sélection de l'équipe projet et de l'équipe d'experts selon des critères précis ; (2) définition rigoureuse des questions à étudier ; (3) formation ; (4) répartition des tâches ; (5) travail individuel des experts ; (5) élicitations des opinions des experts (c'est-à-dire leur formalisation en connaissances explicites et transmissibles) ; (6) analyse et agrégation des résultats ; (7) documentation.

Il s'agit de la distribution la moins informative (celle qui minimise l'information statistique), compte tenu des contraintes.

L'entropie S mesure la quantité totale d'incertitude représentée par cette distribution, après que toutes les informations pertinentes sur la situation aient été prises en compte. L'*entropie d'une distribution de probabilité* s'interprète comme l'information moyenne qu'un tirage selon cette loi pourrait apporter. On peut la considérer comme une quantification de l'intérêt d'effectuer une réalisation de l'événement. Ainsi, dans le cas discret ($S(X) = -\sum_{x \in \Omega} p(x) \ln(p(x))$), l'entropie est maximale quand toutes les possibilités sont a priori équiprobables. S'il y a N possibilités, l'entropie est alors $\ln(N)$. Inversement, si la mesure est concentrée en un point de probabilité 1, alors un tirage sous cette loi n'apporte aucune information car le résultat est connu d'avance, l'entropie est nulle.

Remarque 11.2.7. Différentes définitions existent, employant notamment $\log_2$ ou $\ln$. Lorsque les logarithmes de ces formules sont pris en base 2 l'information est mesurée en bits. Lorsque la base est e , l'unité est le nat.

La figure 11.8 montre la variation de l'entropie avec la forme de la distribution (voir Pernot (2018)).

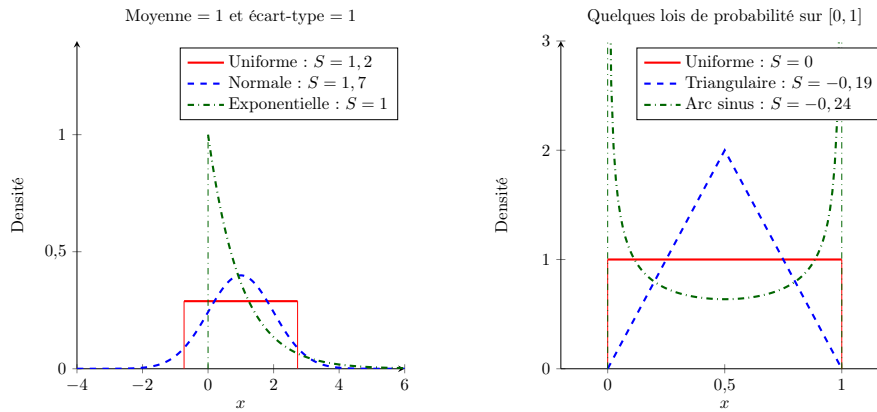


FIGURE 11.8 – Variation de l'entropie avec la forme de la distribution

Exemple 11.2.2. Pour une variable aléatoire réelle X (on se place en 1D) dont on suppose que l'on peut définir l'entropie ($f \ln f$ est Lebesgue-intégrable, pour les densités continues) :

- Si on ne dispose que d'un support borné de X , sans indications sur une valeur préférée, la distribution uniforme $\mathcal{U}([a, b])$ est d'entropie maximale parmi toutes les distributions continues sur $[a, b]$;
- Si on dispose des bornes du support de X et d'une valeur préférée au centre de l'intervalle, on choisira une distribution triangulaire ;
- Si on dispose d'une valeur moyenne μ et d'une incertitude connue sous la forme d'un écart-type absolu σ ou relatif $\sigma/\bar{x}$, alors $\mathcal{N}(\mu, \sigma)$ est la distribution d'entropie maximale sur l'intervalle $]-\infty, +\infty[$.

Démonstration. Cas loi uniforme. Montrons que parmi toutes les lois de probabilités μ dont la densité f est à support dans $[a, b]$, l'entropie est maximisée par la loi uniforme $\mathcal{U}([a, b])$.

$$\begin{aligned} S(\mathcal{U}([a, b])) - S(\mu) &= - \int_{[a, b]} \frac{1}{b-a} \ln \left(\frac{1}{b-a} \right) dx + \int_{[a, b]} f(x) \ln(f(x)) dx \\ &= \ln(b-a) \int_a^b f(x) dx + \int_a^b f(x) \ln(f(x)) dx \end{aligned}$$

Comme f est une densité donc d'intégrale 1 (et passage à Riemann), on peut écrire :

$$S(\mathcal{U}([a, b])) - S(\mu) = - \int_a^b f(x) \ln \left(\frac{1}{(b-a)f(x)} \right) dx$$

En utilisant $\ln x \leq x - 1$:

$$\begin{aligned} S(\mathcal{U}([a, b])) - S(\mu) &\geq - \int_a^b f(x) \left(\frac{1}{(b-a)f(x)} - 1 \right) dx \\ &\geq - \int_a^b \frac{1}{b-a} dx + \int_a^b f(x) dx \\ &\geq 0 \end{aligned}$$

Cas loi gaussienne. Montrons que parmi toutes les lois de probabilités μ de variance imposée σ^2 , l'entropie est maximisée par les lois gaussiennes $\mathcal{N}(m, \sigma^2)$. Par équivalence par translation, on peut supposer $m = 0$ sans perte de généralité. En notant p la densité gaussienne et q la densité de μ , on remarque tout d'abord que

$$\int q(x) \ln(p(x)) dx = \int p(x) \ln(p(x)) dx = \text{constante}$$

En effet, pour toute densité q vérifiant la condition (et donc valable également pour p) :

$$\int q(x) \ln(p(x)) dx = \int_{\mathbb{R}} q(x) \left(-\frac{1}{2} \ln(2\pi\sigma^2) - \frac{x^2}{2\sigma^2} \right) dx$$

On reconnaît dans le premier terme l'intégrale de la fonction de densité, qui vaut 1, et dans le second la variance de μ , ce qui donne :

$$\begin{aligned} - \int q(x) \ln(p(x)) dx &= \frac{1}{2} \ln(2\pi\sigma^2) + \frac{\sigma^2}{2\sigma^2} \\ &= \ln(\sigma\sqrt{2\pi e}) \end{aligned}$$

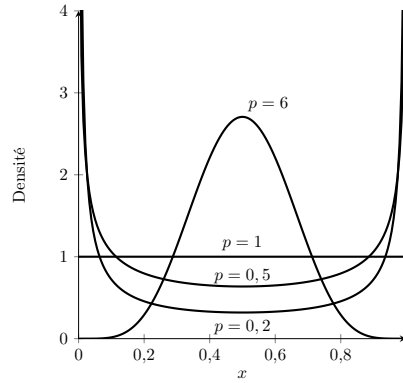
On peut noter au passage que l'on a ainsi calculé l'entropie d'une gaussienne. On écrit ensuite :

$$\begin{aligned} S(\mu) - S(\mathcal{N}(0, \sigma)) &= - \int q(x) \ln(q(x)) dx + \int p(x) \ln(p(x)) dx \\ &= - \int q(x) \ln(q(x)) dx + \int q(x) \ln(p(x)) dx \\ &\quad - \int q(x) \ln(p(x)) dx + \int p(x) \ln(p(x)) dx \\ &= \int q(x) \ln \left(\frac{p(x)}{q(x)} \right) dx \\ &\leq \int q(x) \left(\frac{p(x)}{q(x)} - 1 \right) dx, \text{ car } \ln(x) \leq x - 1 \\ &\leq 0, \text{ car } p(x) \text{ et } q(x) \text{ sont des densités d'intégrales } 1 \end{aligned}$$

□

Remarque 11.2.8 (Sur la différence entre variance et entropie). *L'exemple de la loi Beta $\mathbb{B}(p, p)$ permet d'illustrer cette différence en langage courant (figure 11.9 Maurin (1999)) :*

- Lorsque p tend vers zéro la probabilité se concentre autour des deux bornes 0 et 1, ce qui entraîne que la variance de la distribution augmente (et qu'elle croît vers $1/4$) ; alors que dans le même temps le 'champ offert' en pratique à la variabilité ou à l'errance du hasard se resserre au voisinage de ces deux bornes, ce que l'entropie traduit en décroissant en $-\frac{1}{p}$.
- A l'opposé lorsque p tend vers l'infini le champ offert au hasard se réduit également et l'entropie décroît (en $-\frac{1}{2} \ln p$), mais il est connexe autour de l'espérance ce qui a pour effet de réduire aussi la dispersion et donc la variance ; ici la différence de comportement entre variance et entropie disparaît.

FIGURE 11.9 – Densités de lois Beta $B(p, p)$

Exemple 11.2.3. On peut illustrer la différence d'approche entre les méthodes de type A et B sur un exemple traitant de la température d'un local, pour laquelle on souhaite une valeur moyenne et une incertitude-type (écart-type).

Méthode de type A. On enregistre la température pendant un an puis on traite les données :

$$\text{moyenne définie par } \bar{T} = \frac{1}{N} \sum_{i=1}^N T_i$$

$$\text{écart-type (incertitude-type) défini(e) par } u_T = \sqrt{\frac{1}{N-1} \sum_{i=1}^N (T_i - \bar{T})^2}$$

Méthode de type B. On a seulement connaissance des températures extrêmes du local et on se donne une forme de distribution nous paraissant pertinente (par exemple une distribution uniforme) :

$$\text{moyenne définie par } \bar{T} = \frac{T_{\max} + T_{\min}}{2}$$

$$\text{écart-type (incertitude-type) défini(e) par } u_T = \frac{T_{\max} - T_{\min}}{\sqrt{12}}$$

Remarque 11.2.9 (Variables dépendantes). L'hypothèse d'indépendance entre les différentes sources d'incertitude permet de les traiter de manière univariée (une par une). Si ce n'est pas le cas, c'est-à-dire que le modèle prend comme entrées des variables que l'on soupçonne d'être dépendantes (par exemple un même type de paramètre pris à différents endroits ou différentes dates), alors les mêmes méthodes sont applicables, mais doivent permettre d'estimer également ces dépendances, ce qui rend leur mise en oeuvre plus délicate. Ces dépendances pourront ensuite être prises en compte dans l'application des méthodes de propagation des incertitudes que nous allons présenter maintenant.

11.3 Méthodes de propagation des incertitudes

L'analyse d'incertitude consiste idéalement en la détermination de la loi de la sortie Y selon les lois choisies pour les facteurs X_i . Lorsque ce n'est pas possible pour des raisons de complexité du modèle ou de coût en temps de calcul, on se restreint à estimer certaines *quantités d'intérêt* (q.i.) comme par exemple :

1. un *intervalle* dans lequel la sortie se situe de manière certaine, sous réserve qu'un tel intervalle existe ;
2. les *premiers moments* (espérance, variance) de la distribution de sortie : on est dans une approche probabiliste dans laquelle l'incertitude est caractérisée par la variance ;

3. certains *quantiles* y_α , où $\mathbb{P}[Y < y_\alpha] = \alpha$ pour α donné : cette quantité d'intérêt est notamment utilisée en analyse de risque puisqu'elle concerne la probabilité qu'une sortie dépasse un seuil donné ;
4. l'*incertitude élargie*, qui est définie comme la demi-amplitude d'un intervalle dont on puisse s'attendre à ce qu'il comprenne une fraction élevée de la distribution des valeurs qui pourraient être attribuées raisonnablement à la sortie.

Le choix des méthodes de propagation d'incertitude dépend de la quantité d'intérêt considérée en objectif (ou pour la prise de décision) ainsi que des propriétés du modèle (régularité, temps de simulation, etc.). Nous décrivons ici trois méthodes.

11.3.1 Méthode des intervalles ou théorie des erreurs maximales ou approche min-max

Cette méthode est l'approche traditionnelle, que l'on appelle également *calcul d'erreur*.

La quantité d'intérêt considérée est un intervalle⁸ dans lequel la sortie se situe de manière certaine : $[\min(Y), \max(Y)]$. La modélisation de l'incertitude des facteurs ne nécessite pas de loi de probabilité, mais seulement leur intervalle d'existence $[x_i^{\min}, x_i^{\max}]$, pour $i \in \{1, \dots, k\}$. Il s'agit donc d'une méthode déterministe. On se ramène à résoudre un problème d'optimisation de la forme :

$$\max_{x_i \in [x_i^{\min}, x_i^{\max}]} f(x_1, x_2, \dots, x_k, \mathbf{d}), \text{ pour } i \in \{1, \dots, k\}$$

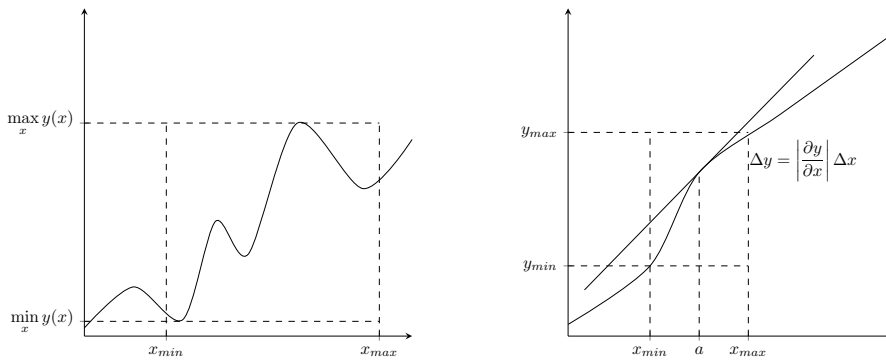


FIGURE 11.10 – Illustration de la méthode des intervalles

Ce problème d'optimisation se résout généralement par des méthodes numériques plus ou moins complexes (voir section 10.4.4). On peut avoir des traitements analytiques dans certains cas simples (figure 11.10). Ainsi, la résolution du problème d'optimisation peut être complexe si la fonction f présente de multiples optima locaux (figure 11.10, gauche). Si, en revanche, elle est monotone, les optima sont atteints aux bornes des intervalles de définition des facteurs (figure 11.10, droite).

Théorème 11.3.1

Soit $f : \mathbb{R}^k \mapsto \mathbb{R}$ monotone par rapport à chacune de ses variables et différentiable.

$$\begin{aligned} \Delta y &= |f'(a)|^T \Delta x + o(\|\Delta x\|) \\ &= \sum_{i=1}^k \left| \frac{\partial f}{\partial x_i}(a_i) \right| \Delta x_i + o(\|\Delta x\|) \end{aligned}$$

où $f'(a)$ désigne la *différentielle* de f en a (son *gradient*, plus spécifiquement) et $o(\|\Delta x\|)$ regroupe les termes d'ordres supérieurs.

8. Les bornes de l'intervalle peuvent potentiellement être infinies.

Démonstration. La monotonie de f impose que son optimum soit atteint aux limites de son espace de définition, c'est-à-dire en $b = (a_1 \pm \Delta x_1, \dots, a_k \pm \Delta x_k)$. Alors, en effectuant un développement de Taylor de f au voisinage de a et en supposant les termes $\|\Delta x\|$ assez petits, ce qui est toujours espéré pour une caractérisation d'incertitudes, on peut approximer :

$$f(b) = f(a) + \Delta y \approx f(a) + \sum_{i=1}^k \frac{\partial f}{\partial x_i}(a_i)(\pm \Delta x_i)$$

Les composantes pour lesquelles $\frac{\partial f}{\partial x_i}(a_i) < 0$ correspondent à celles où le maximum est atteint en $a_i - \Delta x_i$ (fonction décroissante selon cette direction) et réciproquement. On peut donc finalement écrire : $\Delta y = \sum_{i=1}^k \left| \frac{\partial f}{\partial x_i}(a_i) \right| (\Delta x_i)$. $\square$

Exemple 11.3.1. Pour une somme ou différence : $y = x_1 + x_2 - x_3$, les incertitudes s'additionnent :

$$\Delta y = \Delta x_1 + \Delta x_2 + \Delta x_3$$

Pour un produit ou quotient, il est pratique de considérer les incertitudes relatives : $y = \frac{x_1 \cdot x_2}{x_3}$ donne

$$\Delta y = \frac{x_2}{x_3} \Delta x_1 + \frac{x_1}{x_3} \Delta x_2 + \frac{x_1 x_2}{x_3^2} \Delta x_3 \text{ et } \frac{\Delta y}{y} = \frac{\Delta x_1}{x_1} + \frac{\Delta x_2}{x_2} + \frac{\Delta x_3}{x_3}$$

Cette approche est très conservative. Elle considère en effet le maximum d'erreur qui puisse être fait, qui est, dans la réalité, très peu susceptible de se produire, sauf en cas de fortes corrélations entre les facteurs incertains. S'ils sont au contraire indépendants (et de variance finie), le *théorème de la limite centrale* (voir théorème E.2 en annexe E) nous indique typiquement que leur somme converge vers une loi normale, donc la probabilité de tirer des valeurs extrêmes devient très faible (figure 11.11). Comme illustré dans cette figure, dès que l'on combine quelques variables aléatoires, une tendance centrale se dégage et les valeurs extrêmes deviennent non représentatives.

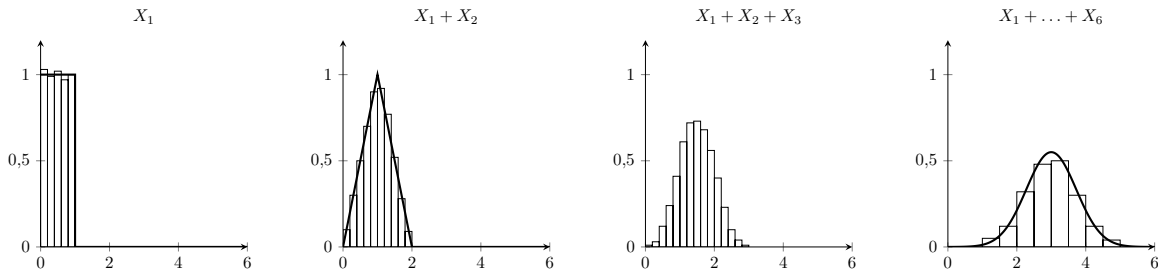


FIGURE 11.11 – Illustration du théorème de la limite centrale

La méthode de la section suivante 11.3.2 adopte une vision plus probabiliste de la propagation d'incertitudes.

11.3.2 Approximation de la variance par méthode de Taylor : méthode dite de la combinaison des variances

La quantité d'intérêt considérée ici est la *variance* ou l'*écart-type*, parfois également appelée dans ce contexte *incertitude-type*.

Théorème 11.3.2 (*Méthode de la combinaison des variances*)

Pour un modèle linéaire en ses facteurs incertains ou, en approximation, pour un modèle pouvant être approximé par un modèle linéaire sur le domaine d'incertitude considéré, avec

l'incertitude-type u_Y définie par :

$$u_Y^2 = \sum_{i=1}^k \left(\frac{\partial f}{\partial x_i} \right)_{\mathbb{E}[X]}^2 \text{Var}[X_i] + \sum_{i \neq j} \left(\frac{\partial f}{\partial x_i} \right)_{\mathbb{E}[X]} \left(\frac{\partial f}{\partial x_j} \right)_{\mathbb{E}[X]} \text{Cov}[X_i, X_j] \quad (11.3.1)$$

et l'expression suivante est vérifiée :

$$\mathbb{E}[Y] = f(\mathbb{E}[X_1], \dots, \mathbb{E}[X_k]) \quad (11.3.2)$$

Cette formule est fondée sur un développement de Taylor de f tronqué au premier ordre (linéaire). Sa validité dépend donc de la linéarité du problème dans le domaine de variation considéré pour les X_i . En particulier, même si le modèle est non-linéaire, on la considère généralement en pratique comme une approximation acceptable si les incertitudes relatives sont inférieures à 10 % et si le modèle est monotone dans la région d'intérêt. Il faut également que les densités de probabilité des X_i soient symétriques pour que ce choix de la quantité d'intérêt soit pertinent.

Démonstration. Si l'on considère un modèle linéaire selon ses paramètres incertains, ou bien plus généralement le développement en série de Taylor au premier ordre d'un modèle quelconque autour d'un point $\mathbf{x}_0$, on peut écrire :

$$f(X) = f(\mathbf{x}_0) + \nabla f^T \cdot (X - \mathbf{x}_0)$$

où les coefficients ∇f_i sont donnés par les dérivées premières du modèle au point $\mathbf{x}_0$:

$$\nabla f_i = \frac{\partial f}{\partial x_i}(X)|_{X=\mathbf{x}_0}$$

Alors, par linéarité de l'espérance :

$$\mathbb{E}[Y] = \mathbb{E}[f(\mathbf{x}_0)] + \nabla f^T \cdot (\mathbb{E}[X] - \mathbf{x}_0)$$

et en choisissant $\mathbf{x}_0 = \mathbb{E}[X]$, on obtient :

$$\mathbb{E}[Y] = f(\mathbb{E}[X])$$

Pour la variance :

$$\begin{aligned} \text{Var}[Y] &= \mathbb{E}[Y^2] - \mathbb{E}[Y]^2 \\ &= \mathbb{E}[(f(\mathbb{E}[X]) + \nabla f^T (X - \mathbb{E}[X]))^2] - (f(\mathbb{E}[X]))^2 \\ &= \mathbb{E}[f(\mathbb{E}[X])^2] + 2f(\mathbb{E}[X])\nabla f^T \mathbb{E}[X - \mathbb{E}[X]] \\ &\quad + \mathbb{E}[\nabla f^T (X - \mathbb{E}[X])(X - \mathbb{E}[X])^T \nabla f] - (f(\mathbb{E}[X]))^2 \\ &= 0 + \nabla f^T \mathbb{E}[(X - \mathbb{E}[X])(X - \mathbb{E}[X])^T] \nabla f \\ &= \nabla f^T \text{Cov}[X] \nabla f \end{aligned}$$

□

Le théorème de la limite centrale a une grande portée parce qu'il montre le rôle très important joué par les variances des lois de probabilité des grandeurs d'entrée par rapport au rôle des moments plus élevés de ces lois, pour la détermination de la forme de la loi convoluée résultante de Y . Il implique en outre : (i) que la loi convoluée converge vers une loi normale avec l'augmentation du nombre des grandeurs d'entrée qui contribuent à $\sigma^2(Y)$; (ii) que la convergence sera d'autant plus rapide que les valeurs des $\sigma^2(X_i)$ seront plus proches les unes des autres (ce qui équivaut en pratique à ce que chaque estimation d'entrée x_i contribue par une incertitude comparable à l'incertitude de l'estimation y du mesurande Y); (iii) que plus les lois des X_i seront proches de la normalité, moins il sera nécessaire d'en avoir un grand nombre pour obtenir une loi normale pour Y .

Les approches des intervalles et de la combinaison des variances sont liées par la proposition suivante :

Proposition 11.3.1

On a la relation de majoration : $u_Y \leq \sum_{i=1}^k \left| \frac{\partial f}{\partial x_i} \right|_{\mathbb{E}[X]} \text{Var}[X_i]$.

Démonstration. D'après l'*inégalité de Cauchy-Schwarz* sur le terme de covariance : $|\text{Cov}[X, Y]| \leq \text{Var}[X]\text{Var}[Y]$ d'où :

$$\begin{aligned} u_Y^2 &\leq \sum_{i=1}^k \left(\frac{\partial f}{\partial x_i} \right)_{\mathbb{E}[X]}^2 \text{Var}[X_i] + 2 \sum_{i>j} \left| \frac{\partial f}{\partial x_i} \right|_{\mathbb{E}[X]} \left| \frac{\partial f}{\partial x_j} \right|_{\mathbb{E}[X]} \text{Var}[X_i]\text{Var}[X_j] \\ &\leq \left(\sum_{i=1}^k \left| \frac{\partial f}{\partial x_i} \right|_{\mathbb{E}[X]} \text{Var}[X_i] \right)^2 \end{aligned}$$

L'expression suivante est ensuite obtenue :

$$u_Y \leq \sum_{i=1}^k \left| \frac{\partial f}{\partial x_i} \right|_{\mathbb{E}[X]} \text{Var}[X_i]$$

On reconnaît dans cette dernière expression le théorème 11.3.1 en assimilant les incertitudes-types aux intervalles de variation : ceux-ci apparaissent donc bien comme un majorant de l'incertitude lorsque l'on adopte une démarche probabiliste. $\square$

11.3.3 Propagation des incertitudes par méthode de Monte-Carlo

La quantité d'intérêt considérée consiste en une estimation de la distribution de probabilité de Y ou de ses résumés statistiques. C'est la plus informative puisqu'elle peut inclure toutes les quantités d'intérêt, et elle ne nécessite pas d'hypothèse particulière sur le modèle ; en contre-partie, elle est la méthode la plus coûteuse numériquement.

La méthode de Monte-Carlo pour la propagation des distributions se traduit par les étapes suivantes :

- *Etape 1* : définir les densités de probabilité pour toutes les variables incertaines (étape de modélisation des incertitudes, section 11.2) ;
- *Etape 2* : générer un échantillon représentatif de chacune des variables ou groupes de variables corrélées à l'aide de générateurs de nombres aléatoires ;
- *Etape 3* : calculer la sortie pour chaque point de l'échantillon $y_i = f(x_1^i, \dots, x_k^i)$, avec $i \in \{1, \dots, N\}$;
- *Etape 4* : analyser l'échantillon $\{y_i\}$, avec $i \in \{1, \dots, N\}$, ainsi obtenu (résumés statistiques, histogrammes, etc.). Par exemple, l'espérance de Y sera approximée par $\bar{f} = \frac{1}{N} \sum_{i=1}^N f(x_1^i, \dots, x_k^i)$

et la variance par $\frac{1}{N-1} \sum_{i=1}^N (f(x_1^i, \dots, x_k^i) - \bar{f})^2$.

Démonstration. Soit $g_{X_1, X_2, \dots, X_k} : \mathbb{R}^k \rightarrow \mathbb{R}$ la densité de probabilité jointe des facteurs d'entrée $g_{X_1, X_2, \dots, X_k}(x_1, x_2, \dots, x_k)$ et soit $g_Y : \mathbb{R} \rightarrow \mathbb{R}$ la densité de $Y = f(X_1, X_2, \dots, X_k)$. Il s'agit ici simplement d'explicitier pourquoi, pour obtenir des statistiques de Y comme son espérance et sa variance, on peut se contenter d'estimer des intégrales (multiples) sur les variables incertaines du modèle et qu'il n'est pas forcément nécessaire de calculer explicitement g_Y .

En effet, pour l'espérance de $Y = f(X_1, \dots, X_k)$, le *théorème de transfert* vu en cours de *Convergence, intégration, probabilités, équations aux dérivées partielles* donne :

$$\begin{aligned} \mathbb{E}[Y] &= \int_{\mathbb{R}} y g_Y(y) dy \\ &= \int_{\mathbb{R}^k} f(\mathbf{x}) g_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \end{aligned}$$

avec le vecteur $\mathbf{x} = (x_1, \dots, x_k)$.

On a alors de même, pour la variance :

$$\begin{aligned}\text{Var}[Y] &= \mathbb{E}[(f(\mathbf{x}) - \mathbb{E}[f(\mathbf{X})])^2] \\ &= \int (f(\mathbf{X}) - \mathbb{E}[f(\mathbf{X})])^2 g_{\mathbf{X}}(\mathbf{x}) d\mathbf{x}\end{aligned}$$

Il suffit ensuite d'exprimer les approximations de ces intégrales par la méthode de Monte-Carlo (annexe E) pour en déduire les expressions données à l'étape 4 de l'algorithme. $\square$

11.3.3.1 Génération d'échantillons

La méthode de référence pour générer un échantillon est la méthode de Monte-Carlo dans laquelle chaque tirage est réalisé indépendamment, selon la distribution de probabilité considérée. La taille de l'échantillon nécessaire (N) ne dépend pas du nombre de variables incertaines, mais des estimateurs statistiques d'intérêt que l'on pourrait souhaiter extraire de la distribution obtenue : par exemple, l'incertitude sur la valeur moyenne converge en $1/\sqrt{N}$. Si l'on ne s'intéresse pas aux queues de la distribution, quelques centaines à quelques milliers de points sont des tailles typiques d'échantillons.

D'autres algorithmes existent parmi lesquels on peut citer la *méthode d'échantillonnage stratifié par hypercube latin* (en anglais *Latin Hypercube Sampling – LHS*).

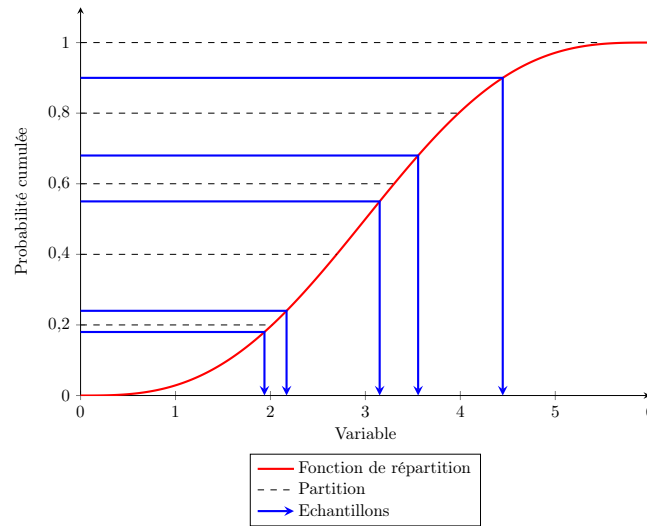


FIGURE 11.12 – Illustration de l'échantillonnage par hypercube latin

Son principe, pour le tirage indépendant de chaque variable aléatoire, est de décomposer l'intervalle $[0, 1]$ en N intervalles égaux, de calculer leurs images par la réciproque de la fonction de répartition selon laquelle on veut tirer, et de choisir une valeur dans chacun de ces intervalles (au hasard ou bien la médiane pour une certaine variante de l'algorithme), voir figure 11.12. On apparie ensuite les différents tirages de chaque variable de manière aléatoire. Cette méthode assure une meilleure couverture de l'espace de variation des facteurs incertains et certains travaux théoriques montrent qu'elle a dans certains cas une convergence plus rapide que la méthode de Monte-Carlo. Elle a en contre-partie un coût de stockage en mémoire plus important. Certaines variantes permettent également de prendre en compte les corrélations entre variables (Aistleitner et al., 2012).

Remarque 11.3.1 (Méta-modèle ou modèle surrogé). *Lorsque les temps de simulation associés à un modèle sont prohibitifs, on peut se tourner vers la construction d'un modèle surrogé : il s'agit de construire une fonction mathématique qui permette d'approcher le comportement du modèle initial à partir d'un jeu limité de simulations soigneusement sélectionnées dans l'espace d'incertitude des paramètres (problème de design expérimental). Les familles de méta-modèles les plus classiques*

incluent les polynômes, les modèles linéaires généralisés, les splines, le krigeage, l'utilisation de réseaux de neurones, des arbres de régression, etc. L'utilisation d'un modèle surrogé introduit bien évidemment une source d'incertitude supplémentaire qui doit être évaluée au regard de l'application visée.

La démarche de modélisation, puis propagation de l'incertitude peut être synthétisée par le diagramme de la figure 11.13 (Pernot, 2018). L'analyse d'incertitude mène naturellement à la question de savoir comment réduire cette incertitude sous un seuil acceptable si jamais elle était trop élevée pour permettre une bonne utilisation du modèle. C'est là qu'intervient l'analyse de sensibilité.

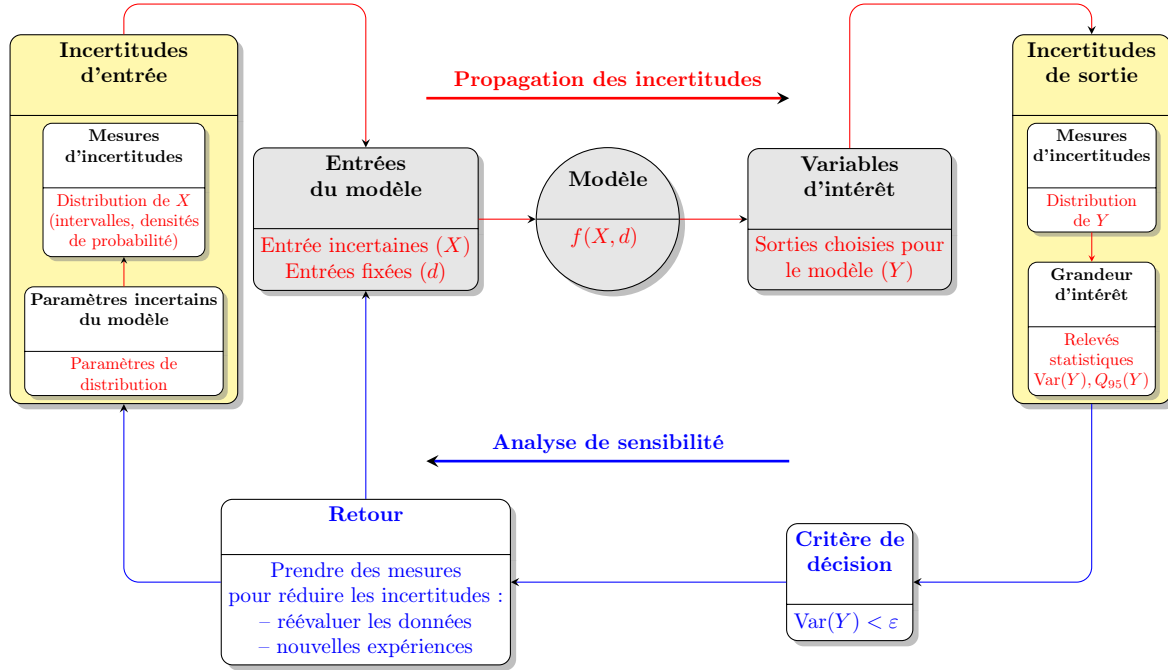


FIGURE 11.13 – Diagramme détaillé des procédures d'analyse d'incertitude et de sensibilité

11.4 Analyse de sensibilité

On rappelle les notations introduites en section 11.1 :

$$Y = f(X, \mathbf{d}) \quad (11.4.1)$$

où $Y \in \mathbb{R}$ est la variable de sortie d'intérêt, $X = (X_1, \dots, X_k) \in \mathbb{R}^{1 \times k}$ est le vecteur des facteurs sur lesquels on suppose une incertitude et $\mathbf{d} = (d_1, \dots, d_{n_d}) \in \mathbb{R}^{1 \times n_d}$ le vecteur des facteurs fixés (i.e. pris constants). On considère effectuée l'étape de modélisation des incertitudes décrite en section 11.2 pour associer une loi de probabilité à chaque facteur incertain. On souhaite savoir quelles sont les variables qui contribuent le plus à la variabilité de la réponse du modèle. Plusieurs méthodes sont détaillées par la suite.

11.4.1 Méthodes locales

Les méthodes d'analyse de sensibilité locales sont les plus répandues et les plus faciles à mettre en oeuvre. Elles s'appuient sur la dérivée (analytique ou numérique) de la fonction f définissant le modèle (11.4.1) en un point de référence donné x^* , à condition que celle-ci existe.

Définition 11.4.1

L'indice de sensibilité local se définit comme :

$$S_i^d(x^*) = \left(\frac{\partial f}{\partial x_i} \right)_{(x_1, x_2, \dots, x_k)^*}$$

Cet indice local est qualifié d'*O-A-T (One-At-a-Time)* car l'effet de la variation de chaque facteur est évalué de manière indépendante : ces indices ne peuvent donc pas fournir d'indications sur les éventuelles interactions entre facteurs.

On trouve parfois l'indice normalisé par les valeurs de référence de y et du paramètre considéré x_i^* :

Définition 11.4.2

L'indice local normalisé (aussi appelé *indice logarithmique*) est défini par l'expression :

$$S_i^l(x^*) = \frac{x_i^*}{y^*} \cdot \left(\frac{\partial f}{\partial x_i} \right)_{(x_1, x_2, \dots, x_k)^*}$$

Ces deux indices ne dépendent pas de la distribution de probabilité ou de l'intervalle de variation choisi pour le facteur X_i : ils sont donc insensibles à des variations de l'incertitude que l'on a sur chacun des facteurs. Ainsi, même si l'on acquiert de l'information sur un facteur, sa contribution à l'incertitude sur la sortie Y évaluée par ces indices ne change pas.

11.4.2 Méthode d'analyse de sensibilité globale pour un modèle linéaire ou quasi-linéaire : indices SRC

Afin de pouvoir prendre en compte une information plus globale, nous allons désormais choisir la variance comme quantité d'intérêt pour caractériser l'incertitude, et chercher à déterminer quelle part de variance de la réponse est due à la variance de chaque variable d'entrée.

11.4.2.1 Cas d'un modèle linéaire

Soit f un modèle linéaire⁹ de $\mathbb{R}^k \mapsto \mathbb{R} : Y = f(X_1, \dots, X_k) = a_0 + \sum_{i=1}^k a_i X_i$ où $\forall i, a_i \in \mathbb{R}$.

Comme les variables X_i sont supposées indépendantes, on a :

$$\text{Var}[Y] = \sum_{i=1}^k a_i^2 \text{Var}[X_i]$$

où $a_i^2 \text{Var}[X_i]$ peut être interprétée comme la part de variance due à la variable X_i . Ainsi, nous pouvons quantifier la sensibilité de Y à X_i par le rapport de la part de variance due à X_i sur la variance totale :

Définition 11.4.3

On définit les *indices SRC* ou *coefficients de régression standardisés (Standardized Regression Coefficient)* par :

$$SRC_i = \frac{a_i^2 \text{Var}[X_i]}{\text{Var}[Y]}$$

9. Pour alléger les notations, on omettra souvent les arguments fixés (i.e. $\mathbf{d}$ en (11.4.1)) pour ne garder que ceux qui sont incertains.

Proposition 11.4.1

Les indices SRC vérifient de manière immédiate (toujours pour un modèle linéaire) :

$$SRC_i \in [0, 1]$$

$$\sum_{i=1}^k SRC_i = 1$$

Remarque 11.4.1. En l'écrivant sous la forme

$$SRC_i = \frac{\text{Var}[X_i]}{\text{Var}[Y]} \cdot \left(\frac{\partial f}{\partial x_i} \right)_{(x_1, x_2, \dots, x_k)^*}^2$$

en un point de référence x^* quelconque (il n'en dépend pas), on voit qu'on retrouve l'expression d'un indice local (au carré) avec une normalisation faisant intervenir les variances respectives de X_i et Y (indice qualifié de local-hybride lorsqu'il est appliqué à un modèle non-linéaire).

Remarque 11.4.2. Cet indice de sensibilité peut également s'exprimer en fonction du coefficient de corrélation linéaire entre la réponse du modèle et les facteurs $\rho_{X_i, Y} = \frac{\text{Cov}[X_i, Y]}{\sqrt{\text{Var}[X_i]\text{Var}[Y]}}$. En effet, pour un modèle linéaire, on a la proposition E.5 donnée en annexe E :

$$\text{Cov}[X_i, Y] = a_i \text{Var}[X_i]$$

ce qui entraîne :

$$\begin{aligned} SRC_i &= \frac{a_i^2 \text{Var}[X_i]}{\text{Var}[Y]} \\ &= \frac{\text{Cov}[X_i, Y]^2}{\text{Var}[Y] \text{Var}[X_i]} \end{aligned}$$

c'est-à-dire :

$$SRC_i = \rho_{X_i, Y}^2$$

11.4.2.2 Cas d'un modèle quasi-linéaire

Lorsque le modèle n'est pas formellement linéaire, on peut tout de même estimer des indices SRC ; seulement leur interprétation n'aura de sens que si le modèle a un comportement au moins quasi-linéaire. Pour cela, à partir d'un modèle dont la forme analytique n'est pas nécessairement connue, on génère numériquement un échantillon de N valeurs pour chacun des facteurs d'entrée X_i et on forme ainsi ce que l'on appelle une *matrice d'échantillonnage* de taille $N \times k$:

$$\mathbf{M} = \begin{bmatrix} x_1(1) & x_2(1) & \dots & x_k(1) \\ x_1(2) & x_2(2) & \dots & x_k(2) \\ \vdots & \vdots & \ddots & \vdots \\ x_1(N) & x_2(N) & \dots & x_k(N) \end{bmatrix}$$

où $x_i(l)$, avec $l \in \{1, \dots, N\}$ représente la $l^{\text{ième}}$ réalisation de la variable aléatoire X_i . A chaque ligne de la matrice $\mathbf{M}$, on associe par simulation une réalisation de la réponse du modèle :

$$\begin{bmatrix} y(1) \\ y(2) \\ \vdots \\ y(N) \end{bmatrix}$$

On peut alors effectuer une régression linéaire multivariée pour estimer les coefficients du modèle :

$$\hat{Y} = \hat{a}_0 + \sum_{i=1}^k \hat{a}_i X_i \text{ et on en déduit une estimation des indices SRC :}$$

$$\widehat{SRC}_i = \frac{\hat{a}_i^2 \widehat{\text{Var}}[X_i]}{\widehat{\text{Var}}[Y]}$$

En pratique, on considère le modèle comme quasi-linéaire si le coefficient de corrélation de la régression est supérieur à 0,7 (seuil empirique et qui peut dépendre des applications). Cela s'interprète par le fait que le modèle linéaire explique alors 70% de la variance de Y : les indices SRC peuvent alors être utilisés.

11.4.3 Méthode d'analyse de sensibilité globale pour un modèle quelconque : indices de Sobol

Considérons à présent le cas d'un modèle quelconque, c'est-à-dire non forcément linéaire ou quasi-linéaire.

Pour apprécier l'importance d'une variable d'entrée X_i sur la variance de la sortie Y , on se demande de combien varierait $\text{Var}[Y]$ si jamais on pouvait réduire complètement l'incertitude sur X_i , c'est-à-dire la fixer à une valeur x_i^* .

Remarque 11.4.3. *La variance $\text{Var}[Y]$ ne diminue pas forcément quand l'incertitude sur X_i diminue.*

On s'intéresse à la quantité :

$$\text{Var}[Y] - \text{Var}[Y|X_i = x_i^*]$$

Plus cette quantité est grande, plus l'influence de X_i sur la variance de Y est importante. Comme on ne connaît pas la valeur x_i^* , on va considérer l'espérance de cette quantité pour toutes les valeurs possibles de x_i^* , c'est-à-dire :

$$\text{Var}[Y] - \mathbb{E}[\text{Var}[Y|X_i]]$$

On note que l'espérance dans cette expression est calculée sur la variable X_i tandis que la variance l'est sur toutes les autres variables $X_{j \neq i}$, noté aussi $X_{\sim i}$ ($\sim i$ signifiant *tous sauf i*). On peut donc écrire de manière plus explicite $\text{Var}[Y] - \mathbb{E}_i[\text{Var}_{\sim i}[Y|X_i]]$.

D'après le *théorème de la variance totale* E.3 en annexe E, ce terme se ré-écrit $\text{Var}[\mathbb{E}[Y|X_i]]$. En ajoutant un terme de normalisation, on aboutit finalement à la définition suivante.

Définition 11.4.4

L'indice de sensibilité de Sobol du premier ordre de Y à X_i est défini par :

$$S_i = \frac{\text{Var}[\mathbb{E}[Y|X_i]]}{\text{Var}[Y]}$$

Proposition 11.4.2

Pour un modèle linéaire, on retrouve bien $S_i = \widehat{SRC}_i$ et donc $\sum_{i=1}^k S_i = 1$.

Démonstration. En effet,

$$\begin{aligned}
 \text{Var}[\mathbb{E}[Y|X_i]] &= \text{Var} \left[\mathbb{E} \left[a_0 + \sum_{p=1}^k a_p X_p | X_i \right] \right] \\
 &= \text{Var} \left[a_0 + \sum_{p=1, p \neq i}^k a_p \mathbb{E}[X_p | X_i] + a_i \mathbb{E}[X_i | X_i] \right] \\
 &= \text{Var} \left[a_0 + \sum_{p=1, p \neq i}^k a_p \mathbb{E}[X_p] + a_i X_i \right] \\
 &= 0 + 0 + \text{Var}[a_i X_i] \\
 &= a_i^2 \text{Var}[X_i]
 \end{aligned}$$

□

Définition 11.4.5

Les modèles vérifiant $\sum_{i=1}^k S_i = 1$ sont dits *additifs*. En général on a $\sum_{i=1}^k S_i < 1$

Pour les modèles non-additifs, la part de variance manquante, c'est-à-dire non attribuable aux indices de Sobol du premier ordre, provient d'interactions entre les facteurs. On va donc définir des indices de Sobol d'ordres supérieurs à 1 :

Définition 11.4.6

L'indice de sensibilité de Sobol d'ordre 2 est défini par :

$$S_{ij} = \frac{\text{Var}[\mathbb{E}[Y|X_i, X_j]]}{\text{Var}[Y]} - S_i - S_j$$

L'indice de sensibilité de Sobol d'ordre 3 est le suivant :

$$S_{ijk} = \frac{\text{Var}[\mathbb{E}[Y|X_i, X_j, X_k]]}{\text{Var}[Y]} - S_i - S_j - S_k - S_{ij} - S_{jk} - S_{ik}$$

L'indice de sensibilité de Sobol d'ordre r , avec $r \leq k$:

$$\begin{aligned}
 S_{i_1 \dots i_r} &= \frac{\text{Var}[\mathbb{E}[Y|X_{i_1}, \dots, X_{i_r}]]}{\text{Var}[Y]} - \sum_{i \in \mathcal{I}} S_i \\
 &\quad - \sum_{\substack{1 \leq i < j \leq r \\ (i, j) \in \mathcal{I}^2}} S_{ij} - \dots - \sum_{\substack{1 \leq i_1 < \dots < i_{r-1} \leq r \\ (i_1, \dots, i_{r-1}) \in \mathcal{I}^{r-1}}} S_{i_1 \dots i_{r-1}}
 \end{aligned}$$

avec $\mathcal{I}$ l'ensemble des indices $i_1, i_2, \dots, i_r$.

On note que pour définir l'indice d'ordre n on soustrait tous les indices d'ordres inférieur ou égal à $n - 1$.

Proposition 11.4.3

La décomposition de Sobol pour la variance d'une sortie Y dans un modèle à k facteurs indépendants vérifie :

$$\sum_{i=1}^k S_i + \sum_{1 \leq i < j \leq k} S_{ij} + \dots + S_{i_1 \dots i_k} = 1$$

Démonstration. La démonstration est immédiate en constatant que les mêmes termes apparaissent avec des signes différents d'un indice d'ordre n à un indice d'ordre $n + 1$ et en utilisant que pour le terme de dernier ordre, $\text{Var}[\mathbb{E}[Y|X_1, \dots, X_k]] = \text{Var}[Y]$. $\square$

Le nombre d'indices ainsi construits, de l'ordre 1 à k , est égal à $2^k - 1$. Lorsque k est grand, l'estimation de tous les indices devient rapidement trop coûteuse en temps de calcul. On a alors recours aux indices de sensibilité totaux, qui agrègent la sensibilité de Y à une variable sous toutes ses formes, c'est-à-dire sensibilité à la variable seule et sensibilité à ses interactions avec d'autres variables.

Définition 11.4.7

L'indice de sensibilité total S_{T_i} à la variable X_i est défini comme la somme de tous les indices de sensibilité relatifs à la variable X_i :

$$S_{T_i} = \sum_{l \in \mathcal{I}(i)} S_l$$

où $\mathcal{I}(i)$ représente tous les ensembles d'indices contenant i .

Exemple 11.4.1. Pour un modèle à trois facteurs, l'indice de sensibilité totale à la variable X_1 est donné par :

$$\begin{aligned} S_{T_1} &= S_1 + S_{12} + S_{13} + S_{123} \\ &= \frac{\mathbb{E}[\text{Var}[Y|X_2, X_3]]}{\text{Var}[Y]} \end{aligned}$$

Proposition 11.4.4

L'indice de sensibilité total vérifie la relation suivante :

$$S_{T_i} = \frac{\mathbb{E}[Y|X_{\sim i}]}{\text{Var}[Y]}$$

où $X_{\sim i}$ représente l'ensemble de toutes les variables sauf X_i .

Démonstration. Pour éviter les lourdeurs d'écriture, on va considérer le cas $k = 3$ (la généralisation se conçoit facilement). L'indice de sensibilité total est calculé comme suit :

$$\begin{aligned} S_{T_1} &= S_1 + S_{12} + S_{13} + S_{123} \\ &= S_1 + S_{12} + S_{13} + \frac{\text{Var}[\mathbb{E}[Y|X_1, X_2, X_3]]}{\text{Var}[Y]} - S_1 - S_2 - S_3 - S_{12} - S_{13} - S_{23} \\ &= \frac{\text{Var}[Y]}{\text{Var}[Y]} - S_2 - S_3 - S_{23} \\ &= 1 - S_2 - S_3 - \left(\frac{\text{Var}[\mathbb{E}[Y|X_2, X_3]]}{\text{Var}[Y]} - S_2 - S_3 \right) \\ &= \frac{\mathbb{E}[\text{Var}[Y|X_2, X_3]]}{\text{Var}[Y]} \end{aligned}$$

La dernière égalité résulte de l'application du *théorème de la variance totale*. $\square$

Les indices de Sobol s'appliquent à tous les modèles sans nécessiter aucune autre hypothèse que celle de l'indépendance des facteurs. Leur inconvénient est cependant qu'ils sont numériquement coûteux à estimer dès que les modèles deviennent complexes, comme on peut le voir en annexe G : c'est pourquoi on leur préférera les indices SRCs dès qu'ils seront applicables.

Pour aller plus loin, on peut citer également des alternatives aux indices de Sobol, parmi lesquels la méthode FAST ou la méthode de Morris. On peut mentionner l'utilisation de méta-modèles ou modèles dits surrogés qui peuvent être construits pour approximer le modèle complet lorsque les temps de calcul sont rédhibitoires. Enfin, des méthodes dites d'ordre 2 permettent la prise en compte de l'incertitude sur les paramètres des distributions décrivant l'incertitude !

11.5 Conclusion de la partie III

En conclusion, on peut rappeler les étapes principales de la démarche d'analyse d'incertitude et de sensibilité :

1. Phase 'métrologique' – modélisation de l'incertitude des facteurs d'entrée, qui consiste à identifier les sources possibles d'incertitude sur les facteurs d'entrée et leur modélisation (choix d'une loi de probabilité). Cette phase requiert souvent l'aide d'experts métiers.
2. Détermination de la quantité d'intérêt visée ;
3. Propagation des incertitudes, avec une méthode dépendant de la quantité d'intérêt et des propriétés du modèle ;
4. Analyse de sensibilité, hiérarchisation des contributions.

Il faut aussi noter que les méthodes d'analyse d'incertitude et de sensibilité présentées s'appliquent principalement à des modèles à réponse dans $\mathbb{R}$. Pour des sorties multivariées ou bien, notamment, des modèles dynamiques multivariés (avec plusieurs sorties) comme ceux présentés dans la partie I par exemple, on les applique successivement pour chaque sortie ou en chaque instant d'intérêt. Il est également possible de définir des indices agrégés, consistant par exemple en une somme des indices de sensibilités pondérés par la variance de la sortie associée.

Annexe A

Transformée de Laplace de signaux usuels

Remarque préliminaire : tous les signaux considérés dans cette table sont causaux. $\delta_0(t)$ représente l'impulsion de Dirac et $\mathbb{1}(t)$ représente la fonction échelon de Heaviside.

Toutes les transformées sont définies pour $\text{Re}(p) \in]\sigma_1, +\infty[$, où σ_1 est donnée dans la dernière colonne.

$x(t)$	$X(p)$	σ_1
$\delta_0(t)$	1	$-\infty$
$\delta_0^{(n)}(t)$	p^n	$-\infty$
$\mathbb{1}(t)$	$\frac{1}{p}$	0
$\mathbb{1}(t)e^{-at}$	$\frac{1}{p+a}$	$-a$
$\mathbb{1}(t)t^n$	$\frac{n!}{p^{n+1}}$	0
$\mathbb{1}(t)t^n e^{-at}$	$\frac{n!}{(p+a)^{n+1}}$	$-a$
$\mathbb{1}(t) \sin(\omega t)$	$\frac{\omega}{p^2 + \omega^2}$	0
$\mathbb{1}(t) \cos(\omega t)$	$\frac{p}{p^2 + \omega^2}$	0
$\mathbb{1}(t) \sin(\omega t) e^{-at}$	$\frac{\omega}{(p+a)^2 + \omega^2}$	$-a$
$\mathbb{1}(t) \cos(\omega t) e^{-at}$	$\frac{p+a}{(p+a)^2 + \omega^2}$	$-a$
$x_m(t) * \sum_{k=0}^{+\infty} \delta(t - kT)$ (signal périodique de motif période $x_m(t)$)	$\frac{X_m(p)}{1 - e^{-Tp}}$	σ_{m_1}

TABLEAU A.1 – Transformées de Laplace des fonctions usuelles

Annexe B

Transformée en z de signaux usuels

Remarque préliminaire : tous les signaux considérés dans cette table sont causaux. $\delta_0(k)$ représente l'impulsion discrète et $\mathbb{1}(k)$ représente la fonction échelon de Heaviside.

Toutes les transformées sont définies pour $|z| > R_1$, où R_1 est donnée dans la dernière colonne.

$x(k)$	$X(z)$	R_1
$\delta_0(k)$	1	0
$\mathbb{1}(k)$	$\frac{z}{z-1} = \frac{1}{1-z^{-1}}$	1
$k\mathbb{1}(k)$	$\frac{z}{(z-1)^2} = \frac{z^{-1}}{(1-z^{-1})^2}$	1
$a^k \mathbb{1}(k)$	$\frac{z}{z-a} = \frac{1}{1-az^{-1}}$	$ a $
$ka^k \mathbb{1}(k)$	$\frac{za}{(z-a)^2} = \frac{az^{-1}}{(1-az^{-1})^2}$	$ a $
$\cos(k2\pi\nu) \mathbb{1}(k)$	$\frac{z(z - \cos(2\pi\nu))}{z^2 - 2z\cos(2\pi\nu) + 1} = \frac{1 - z^{-1}\cos(2\pi\nu)}{1 - 2z^{-1}\cos(2\pi\nu) + z^{-2}}$	1
$\sin(k2\pi\nu) \mathbb{1}(k)$	$\frac{z\sin(\pi\nu)}{z^2 - 2z\cos(2\pi\nu) + 1} = \frac{z^{-1}\sin(2\pi\nu)}{1 - 2z^{-1}\cos(2\pi\nu) + z^{-2}}$	1
$a^k \cos(k2\pi\nu) \mathbb{1}(k)$	$\frac{z(z - a\cos(2\pi\nu))}{z^2 - 2za\cos(2\pi\nu) + a^2} = \frac{1 - z^{-1}a\cos(2\pi\nu)}{1 - 2z^{-1}a\cos(2\pi\nu) + a^2z^{-2}}$	$ a $
$a^k \sin(k2\pi\nu) \mathbb{1}(k)$	$\frac{za\sin(\pi\nu)}{z^2 - 2za\cos(2\pi\nu) + a^2} = \frac{z^{-1}a\sin(2\pi\nu)}{1 - 2z^{-1}a\cos(2\pi\nu) + a^2z^{-2}}$	$ a $

TABLEAU B.1 – Transformées en z des fonctions usuelles

Annexe C

Termes techniques en anglais

amortissement	damping
automate à états finis	finite state machine
fonction de transfert	transfer function
machines à états	state machines
matrice de transition	transition matrix
réponse impulsionnelle	impulse response
réponse indicielle	step response
représentation d'état	state-space representation
stabilité entrée bornée sortie bornée (EBSB)	Bounded Input Bounded Output stability (BIBO)
système à déphasage minimal	minimal phase system
système à retard	time delay system
système de systèmes	system of systems
système linéaire invariant dans le temps	linear time invariant system
système monovariable	SISO (Single-Input Single-Output)
système multi-agents	multi-agent system
système multivariable	MIMO (Mingle-Input Mingle-Output)
erreur quadratique moyenne	Mean Squared Error
erreur quadratique moyenne de prédiction	Mean Squared Error of Prediction
échantillonnage stratifié par hypercube latin	Latin Hypercube Sampling

Annexe D

Convolution

Définition et conditions d'existence

Cas continu

Le produit de convolution de deux signaux continus $x(t)$ et $y(t)$ est défini par la fonction $c(t)$:

$$c(t) = \int_{-\infty}^{+\infty} x(\theta)y(t-\theta)d\theta \quad (\text{D.1})$$

Ce produit de convolution est souvent noté $*$: $c(t) = x * y(t)$. L'existence de cette convolution est liée aux propriétés de convergence de l'intégrale de support infini de l'équation (D.1). La convolution est en particulier définie pour les cas suivants :

- x est absolument sommable $\left(\int_{-\infty}^{+\infty} |x(t)| dt < +\infty\right)$ et $|y|$ est borné.
- x et y sont localement sommables et nuls en dehors d'un intervalle fini.
- x et y sont de carré sommable $\left(\int_{-\infty}^{+\infty} |x(t)|^2 dt < +\infty, \int_{-\infty}^{+\infty} |y(t)|^2 dt < +\infty\right)$. En effet, d'après l'inégalité de Schwarz (produit scalaire dans l'espace des fonctions) :

$$|c(t)|^2 = \left| \int_{-\infty}^{+\infty} x(\theta)y(t-\theta)d\theta \right|^2 \leq \int_{-\infty}^{+\infty} x^2(\theta)d\theta \cdot \int_{-\infty}^{+\infty} y^2(t-\theta)d\theta \quad (\text{D.2})$$

Cas discret

De manière similaire, la convolution de deux signaux discrets $x(k)$ et $y(k)$ est définie par le signal discret $c(k)$:

$$c(k) = \sum_{n=-\infty}^{n=+\infty} x(n)y(k-n) = x * y(k) \quad (\text{D.3})$$

L'existence de cette convolution est liée aux propriétés de convergence de la somme infinie de l'équation (D.3). La convolution est en particulier définie pour les cas suivants :

- x est sommable $\left(\sum_{n=-\infty}^{n=+\infty} x(n) < +\infty\right)$ et $|y|$ est borné.
- x et y sont localement sommables et nuls en dehors d'un intervalle fini.
- x et y sont de carré sommable $\left(\sum_{n=-\infty}^{n=+\infty} x^2(n) < +\infty, \sum_{n=-\infty}^{n=+\infty} y^2(n) < +\infty\right)$.

Support du produit de convolution

Cas continu

Si x a pour support l'intervalle $[a_x, b_x]$ et y a pour support l'intervalle $[a_y, b_y]$, alors l'équation (D.1) se réécrit par exemple :

$$c(t) = \int_{\theta=a_x}^{\theta=b_x} x(\theta)y(t-\theta)d\theta \quad (\text{D.4})$$

Cette fonction est non nulle s'il existe t tel que :

$$a_y \leq t - \theta \leq b_y \quad (\text{D.5})$$

En utilisant le fait que $a_x \leq \theta \leq b_x$, le support du produit de convolution c est donc l'intervalle $[a_x + a_y, b_x + b_y]$.

Cas discret

De la même façon que pour le cas continu, si x est une séquence de longueur N et y une séquence de longueur M , alors le produit de convolution de ces deux signaux discrets est de longueur $N + M - 1$. Si,

— $x(k)$ est non nul pour $k \in \{k_0, \dots, k_0 + N - 1\}$,

— $y(k)$ est non nul pour $k \in \{k_1, \dots, k_1 + M - 1\}$,

alors $c(k) = x * y(k)$ est non nul pour $k \in \{k_0 + k_1, \dots, k_0 + k_1 + N + M - 2\}$.

Exemple D.1 (Convolution). *Soit à convoluer deux signaux continus :*

$$\begin{aligned} x(t) &= 1 \text{ si } t \in [0, 1] \\ &= 0 \text{ sinon} \end{aligned} \quad (\text{D.6})$$

$$\begin{aligned} y(t) &= e^{-t} \text{ si } t \geq 0 \\ &= 0 \text{ si } t < 0 \end{aligned} \quad (\text{D.7})$$

Le signal y est sommable, x est borné, la convolution est donc définie et :

$$c(t) = \int_0^1 x(\theta)y(t-\theta)d\theta \quad (\text{D.8})$$

Le signal $y(t-\theta)$ est nul si $t-\theta < 0$. On obtient donc :

$$\begin{aligned} c(t) &= 0 \text{ si } t = 0 \\ c(t) &= \int_0^t e^{-(t-\theta)}d\theta = (e^t - 1)e^t = 1 - e^{-t} \text{ si } 0 \leq t \leq 1 \\ c(t) &= \int_0^t e^{-(t-\theta)}d\theta = (e - 1)e^{-t} \text{ si } t \geq 1 \end{aligned} \quad (\text{D.9})$$

Commutativité, associativité et dérivation

S'il existe, le produit de convolution est commutatif :

$$x * y = y * x$$

Sous réserve des conditions d'existence, le produit de convolution est associatif :

$$(x * y) * z = y * (x * z)$$

Dans le cas de signaux continus, la dérivation du produit de convolution de deux signaux s'exprime en fonction de la dérivée des signaux :

$$\frac{dx * y(t)}{dt} = \frac{dx(t)}{dt} * y(t) = \frac{dy(t)}{dt} * x(t)$$

Convolution de signaux périodiques

Les équations (D.1) et (D.1) ne sont pas définies lorsque les signaux x et y sont périodiques. Dans ce cas, il est possible de définir tout de même la convolution des deux signaux en utilisant les « motifs périodes ». Soient deux signaux continus x et y (resp. discrets) de période T (resp. N). Les motifs périodes sont notés x_T et y_T (resp. x_N et y_N). Les produits de convolution sont définis par :

$$\begin{aligned}
 c(t) &= x * y_T(t) = \int_0^T y_T(\theta) x(t - \theta) d\theta \\
 &= y * x_T(t) = \int_0^T x_T(\theta) y(t - \theta) d\theta \\
 c(k) &= x * y_N(k) = \sum_{n=0}^{N-1} y_N(n) x(k - n) \\
 &= y * x_N(k) = \sum_{n=0}^{N-1} x_N(n) y(k - n)
 \end{aligned} \tag{D.10}$$

Le produit de convolution est lui-même périodique de période T (resp. N). Son motif période est dit obtenu par convolution circulaire et noté $c_T(t) = x_T \otimes y_T(t)$ (resp. $c_N(k) = x_N \otimes y_N(k)$).

Interprétation graphique

La convolution de deux signaux peut s'interpréter de manière graphique. Pour calculer la valeur du produit de convolution en t_0 de deux signaux continus, il suffit :

- de tracer la fonction $y(\theta - t_0)$. Elle s'obtient par symétrie par rapport à l'axe des ordonnées puis décalage de l'origine en t_0 .
- de multiplier $x(\theta)$ par cette fonction $y(\theta - t_0)$.
- puis de calculer la valeur algébrique de l'aire placée sous cette courbe.

La figure D.1 illustre cette construction.

La même interprétation peut être donnée dans le cas de signaux discrets : pour calculer la convolution de $x(k)$ et $y(k)$ en k_0 , il suffit de :

- de tracer la suite $y(k_0 - n)$,
- de multiplier termes à termes $x(k)$ par cette suite $y(k_0 - n)$,
- d'additionner les produits obtenus.

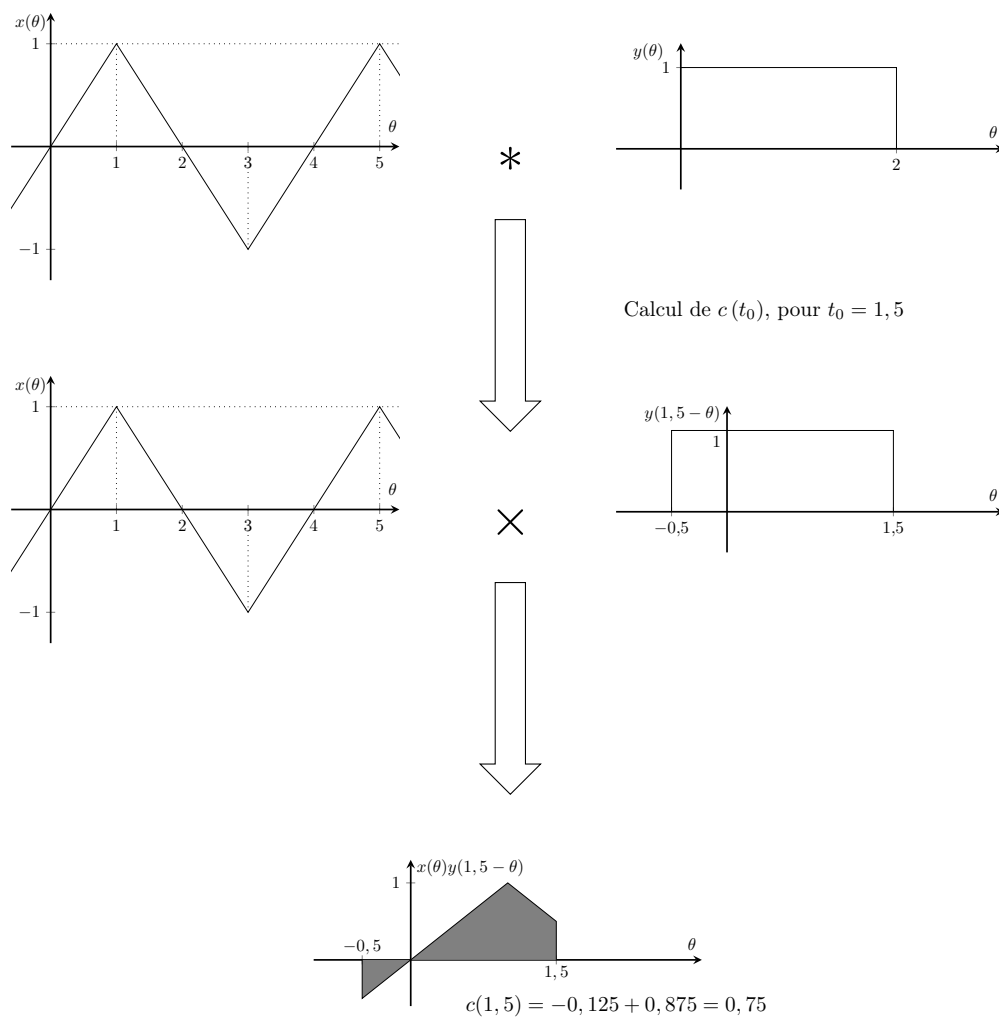


FIGURE D.1 – Calcul graphique du produit de convolution

Annexe E

Rappels de probabilités

Notions de base sur les variables aléatoires

Définition E.1

Soit un *espace de probabilité* $(\Omega, \mathcal{F}, \mathbb{P})$, où Ω est un ensemble, $\mathcal{F}$ est une tribu dans Ω et $\mathbb{P}$ une mesure de probabilité sur $(\Omega, \mathcal{F})$. Une *variable aléatoire* (v.a.) X est une application mesurable de $(\Omega, \mathcal{F})$ dans un ensemble mesurable $(E, \mathcal{E})$. La loi image est la mesure de probabilité $\mathbb{P}_X$ sur $(E, \mathcal{E})$ telle que :

$$\forall A \in \mathcal{E}, \mathbb{P}_X[A] = \mathbb{P}[X^{-1}(A)] = \mathbb{P}[\omega \in \mathcal{F} / X(\omega) \in A].$$

On note souvent abusivement $\mathbb{P}[X \in A]$.

Définition E.2 (*Variable aléatoire discrète*)

Dans le cas où E est un ensemble dénombrable et $\mathcal{E}$ est l'ensemble des parties de E , alors la loi de X s'écrit :

$$\mathbb{P}_X = \sum_{x \in E} p_x \delta_x,$$

où $p_x = \mathbb{P}[X = x]$ et δ_x est la mesure de Dirac en x .

Définition E.3 (*Variable aléatoire continue ou variable aléatoire à densité*)

Si X est à valeurs dans $(\mathbb{R}^d, \mathcal{B}(\mathbb{R}^d))$ et $\mathbb{P}_X$ absolument continue par rapport à la mesure de Lebesgue λ , alors il existe une fonction borélienne $p : \mathbb{R}^d \rightarrow \mathbb{R}^+$, appelée *densité*, telle que :

$$\mathbb{P}_X[A] = \int_A p(x) dx$$

En dimension $d = 1$, on a pour tout $a \leq b$:

$$\mathbb{P}[a \leq X \leq b] = \int_a^b p(x) dx.$$

Définition E.4

La *fonction de répartition* de X est la fonction $\mathbb{P} : \mathbb{R} \rightarrow [0, 1]$ définie par :

$$\forall t \in \mathbb{R}, F_X(t) = \mathbb{P}[X \leq t].$$

Proposition E.1

La fonction de répartition F_X est croissante, continue à droite, tend vers 0 en $-\infty$ et vers 1 en $+\infty$.

Proposition E.2

Si la fonction de répartition F_X est strictement croissante, le quantile d'ordre $\alpha \in]0, 1[$ est le nombre x_α défini par :

$$\mathbb{P}[X \leq x_\alpha] = \alpha.$$

Quelques distributions de probabilité usuelles

Fonctions de masse de quelques v.a. discrètes :

— **Bernoulli** (θ) :

$$p(k; \theta) = \theta^k + (1 - \theta)^{1-k}, \quad \forall k \in \{0, 1\},$$

— **Poisson** ($\lambda > 0$) :

$$p(n) = \frac{\lambda^n}{n!} e^{-\lambda}, \quad n \in \mathbb{N},$$

espérance λ , variance λ .

Densités de probabilité de quelques v.a. continues :

— **Uniforme** (a, b), avec $a > b$:

$$f(x) = \frac{1}{b-a} \mathbb{I}(a \leq x \leq b),$$

espérance $\frac{a+b}{2}$, variance $\frac{(b-a)^2}{12}$;

— **Exponentielle** ($\lambda > 0$) :

$$f(x) = \lambda e^{-\lambda x} \mathbb{I}(0 \leq x),$$

espérance $\frac{1}{\lambda}$, variance $\frac{1}{\lambda^2}$;

— **Gaussienne** (μ, σ^2) :

$$f(x) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} := \mathcal{N}(\mu, \sigma^2),$$

espérance μ , variance σ^2 ;

— **Triangle** (a, b, c), avec $a \leq c \leq b$:

$$f(x) = \frac{2(x-a)}{(b-a)(c-a)} \mathbb{I}(a < x \leq c) + \frac{2(b-x)}{(b-a)(b-c)} \mathbb{I}(c < x \leq b),$$

espérance $\frac{a+b+c}{3}$, variance $\frac{a^2 + b^2 + c^2 - ab - ab - bc}{18}$;

— **Beta** ($\alpha > 0, \beta > 0$) :

$$f(x) = \frac{x^{\alpha-1}(1-x)^{\beta-1}}{\int_0^1 u^{\alpha-1}(1-u)^{\beta-1} du} \mathbb{I}(0 \leq x \leq 1),$$

espérance $\frac{\alpha}{\alpha+\beta}$, variance $\frac{\alpha\beta}{(\alpha+\beta)^2(\alpha+\beta+1)}$.

Moments d'une variable aléatoire

Soit X une variable aléatoire réelle, i.e. à valeurs dans $\mathbb{R}$, continue ou discrète.

Définition E.5

L'*espérance mathématique* de X , lorsqu'elle existe, est notée par :

$$\mathbb{E}[X] = \int X(\omega) \mathbb{P}(\mathrm{d}\omega).$$

Cette définition se décline en :

- v.a. discrètes : $\mathbb{E}[X] = \sum_{i=1}^n x_i \mathbb{P}[X = x_i]$,
- v.a. à densité f_X (continues) : $\mathbb{E}[X] = \int_{-\infty}^{+\infty} x f_X(x) \mathrm{d}x$.

La propriété suivante est énoncée pour des v.a. à densité. Elle se généralise sans difficulté aux v.a. discrètes.

Proposition E.3

Pour toute fonction h mesurable, l'expression suivante est vérifiée (si l'intégrale converge) :

$$\mathbb{E}[h(X)] = \int_{-\infty}^{+\infty} h(x) f_X(x) \mathrm{d}x.$$

Définition E.6

Les *moments d'ordre k* d'une variable aléatoire $X \in L^k$ sont définis par :

$$\mathbb{E}[X^k] = \int x^k f(x) \mathrm{d}x.$$

Pour $k = 1$, on obtient l'*espérance* (ou *valeur moyenne*) de X . La *variance*, qui mesure la dispersion autour de cette valeur moyenne, s'obtient à l'aide des deux premiers moments :

$$\mathrm{Var}[X] = \mathbb{E}[(X - \mathbb{E}[X])^2] = \mathbb{E}[X^2] - \mathbb{E}[X]^2. \quad (\text{E.1})$$

Covariance

Etant données deux variables aléatoires X et Y de L^2 , leur covariance permet de quantifier leur corrélation.

Définition E.7

La *covariance* de deux variables aléatoires X et Y de carrés intégrables est définie par :

$$\mathrm{Cov}[X, Y] = \mathbb{E}[(X - \mathbb{E}[X])(Y - \mathbb{E}[Y])].$$

Proposition E.4

Pour deux variables aléatoires de carrés intégrables, les expressions suivantes sont vérifiées :

$$\mathrm{Cov}[X, Y] = \mathbb{E}[XY] - \mathbb{E}[X]\mathbb{E}[Y],$$

$$\mathrm{Cov}[X, X] = \mathrm{Var}[X].$$

Remarque E.1. Si X et Y sont indépendantes, alors $\text{Cov}[X, Y] = 0$ et la réciproque est fausse en général.

Proposition E.5

Pour le cas d'un modèle linéaire $f : \mathbb{R}^k \mapsto \mathbb{R}$, avec $Y = f(X_1, \dots, X_k) = a_0 + \sum_{i=1}^k a_i X_i$, où $\forall i, a_i \in \mathbb{R}$ et les variables aléatoires X_i sont indépendants, la propriété suivante est vérifiée :

$$\text{Cov}[X_i, Y] = a_i \text{Var}[X_i]$$

Démonstration. Pour i fixé dans $\{1, \dots, k\}$, avec $k \in \mathbb{N}$:

$$\begin{aligned} \text{Cov}[X_i, Y] &= \mathbb{E}[X_i Y] - \mathbb{E}[X_i] \mathbb{E}[Y] \\ &= \mathbb{E}[a_0 X_i + \sum_p a_p X_i X_p] - \mathbb{E}[X_i] \mathbb{E}[a_0 + \sum_p a_p X_p] \\ &= a_0 \mathbb{E}[X_i] + \sum_p a_p \mathbb{E}[X_i X_p] - (a_0 \mathbb{E}[X_i] + \sum_p a_p \mathbb{E}[X_p] \mathbb{E}[X_i]) \\ &= \sum_p a_p (\mathbb{E}[X_i X_p] - \mathbb{E}[X_p] \mathbb{E}[X_i]) \\ &= \sum_p a_p \text{Cov}[X_i, X_p] \end{aligned}$$

Comme les variables aléatoires X_i sont supposées indépendantes les unes des autres, les termes de la somme sont nuls pour tous les $p \neq i$ et on obtient $\text{Cov}[X_i, Y] = a_i \text{Cov}[X_i, X_i] = a_i \text{Var}[X_i]$. $\square$

Variable aléatoire 'moyenne empirique'

Définition E.8

Soit $X_1, \dots, X_n$ une suite de variables aléatoires iid (indépendantes et identiquement distribuées). Alors la *variable aléatoire 'moyenne empirique'* est définie par :

$$\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i.$$

Théorème E.1 (Loi forte des grands nombres)

Si les variables aléatoires X_i sont dans L^1 ($\mathbb{E}[|X_1|] < \infty$), alors :

$$\bar{X}_n \xrightarrow[n \rightarrow \infty]{\text{p.s.}} \mathbb{E}[X_1].$$

Si les v.a. X_i sont dans L^2 , alors $\bar{X}_n$ converge en loi vers une variable Gaussienne.

Théorème E.2 (Théorème de la limite centrale)

Si $\mathbb{E}[X_1^2] < \infty$, alors l'expression suivante est vérifiée :

$$\sqrt{n} \left(\frac{\bar{X}_n - \mathbb{E}[X_1]}{\sqrt{\text{Var}[X_1]}} \right) \xrightarrow{\mathcal{L}} \mathcal{N}(0, 1),$$

ce qui signifie que $\mathbb{P} \left[\sqrt{n} \frac{\bar{X}_n - \mathbb{E}[X_1]}{\sqrt{\text{Var}[X_1]}} \leq z \right] = \phi(z)$, où ϕ est la fonction de répartition de

$\mathcal{N}(0, 1)$.

Ce résultat permet notamment dans bon nombre de cas la construction d'intervalles de confiance pour l'estimation paramétrique.

Espérance conditionnelle

On se donne un ensemble Ω muni d'une tribu $\mathcal{F}$ et d'une mesure de probabilité $\mathbb{P}$.

Définition E.9 (*Probabilité conditionnelle par rapport à un évènement*)

Soit $B \in \mathcal{F}$ un évènement tel que $\mathbb{P}[B] > 0$. On peut définir une nouvelle probabilité sur $(\Omega, \mathcal{F})$, appelée *probabilité conditionnelle sachant B* , en posant pour tout $A \in \mathcal{F}$:

$$\mathbb{P}_B[A] = \mathbb{P}[A|B] = \frac{\mathbb{P}[A \cap B]}{\mathbb{P}[B]}.$$

Les propriétés suivantes sont vérifiées.

Proposition E.6 (*Formule des probabilités totales*)

Si $\Omega = \bigcup_i A_i$, avec A_i une partition de Ω , la probabilité d'un évènement B est donnée par :

$$\mathbb{P}[B] = \sum_i \mathbb{P}[B|A_i] \mathbb{P}[A_i].$$

Proposition E.7 (*Formule de Bayes*)

La probabilité conditionnelle de A sachant B peut s'écrire comme :

$$\mathbb{P}[A|B] = \frac{\mathbb{P}[B|A] \mathbb{P}[A]}{\mathbb{P}[B]}.$$

Définition E.10

De même, pour toute v.a. $Y \geq 0$, ou pour $Y \in L^1(\Omega, \mathcal{A}, P)$, l'*espérance conditionnelle de Y sachant B* est définie par :

$$\mathbb{E}[Y|B] = \frac{\mathbb{E}[Y \cdot \mathbb{I}_B]}{\mathbb{P}[B]},$$

où $\mathbb{I}_B$ est la fonction indicatrice de B , valant 1 si B est réalisé et 0 sinon. Cette quantité $\mathbb{E}[Y|B]$ s'interprète comme la valeur moyenne de Y quand B est réalisé.

Autrement dit, si Y est une v.a. discrète d'espérance finie, nous avons :

$$\mathbb{E}[Y|B] = \sum_i y_i \mathbb{P}[Y = y_i|B] = \sum_i y_i \frac{\mathbb{P}[\{Y = y_i\} \cap B]}{\mathbb{P}[B]}$$

et si Y est une v.a. continue d'espérance finie et de densité f , l'expression suivante est obtenue :

$$\mathbb{E}[Y|B] = \frac{1}{\mathbb{P}(B)} \int_B y f(y) dy.$$

Conditionnement par une variable aléatoire

La construction de l'espérance conditionnelle par une variable aléatoire est vue dans le cours de *Convergence, Intégration, Probabilités*. On s'attachera ici seulement à en présenter les principales propriétés et modes de calculs.

Définition E.11 (*Conditionnement par une variable aléatoire discrète*)

Soit X une v.a. à valeurs dans un espace dénombrable E et soit $Y \in L^1(\Omega, \mathcal{A}, \mathbb{P})$. Alors, l'espérance conditionnelle est donnée par :

$$\mathbb{E}[Y|X] = \phi(X), \text{ où } \phi(x_i) = \frac{\mathbb{E}[Y \cdot \mathbb{I}_{\{X=x_i\}}]}{\mathbb{P}[X = x_i]},$$

pour tout $x_i \in E$ tel que $\mathbb{P}[X = x_i] > 0$ et $\phi(x_i)$ peut être choisie de manière arbitraire lorsque $\mathbb{P}[X = x_i] = 0$.

Définition E.12 (*Conditionnement par une variable aléatoire à densité*)

Soient X et Y deux v.a. absolument continues de densité jointe $f_{X,Y}$. On rappelle que la densité marginale de X s'écrit : $f_X(x) = \int f_{X,Y}(x, y) dy$. Alors on définit la densité conditionnelle de Y par X comme suit :

$$f_{Y|X}(y|x) = \frac{f_{X,Y}(x, y)}{f_X(x)}, \text{ si } f_X(x) > 0.$$

Cette fonction peut prendre des valeurs arbitraires pour les valeurs de x telles que $f_X(x) = 0$. Alors, l'espérance conditionnelle est donnée par :

$$\mathbb{E}[Y|X] = \phi(X), \text{ où } \phi(x) = \int_{\mathbb{R}} y f_{Y|X}(y|x) dy.$$

Dans le cas discret, $\mathbb{E}[Y|X]$ est la v.a. qui prend les valeurs $\mathbb{E}[Y|x = x_i]$, avec les probabilités p_i (où $X = x_i$ avec la probabilité p_i). Cela signifie que c'est elle-même une variable aléatoire fonction de X , que l'on peut interpréter comme la moyenne de Y lorsqu'on connaît X .

Proposition E.8

Par extension des propriétés de l'espérance conditionnelle à une sous-tribu, on aura :

- Si X est une v.a., alors $\mathbb{E}[X|X] = X$,
- Si les v.a. X et Y sont indépendantes, alors $\mathbb{E}[Y|X] = \mathbb{E}[Y]$.

Définition E.13

La *variance conditionnelle* se définit à partir de l'*espérance conditionnelle* :

$$\text{Var}[Y|X] = \mathbb{E}[(Y - \mathbb{E}[Y|X])^2|X].$$

Théorème E.3 (*Théorème de l'espérance et de la variance totale*)

Soit un couple (X, Y) de variables aléatoires sur un même espace de probabilité, tel que X soit intégrable et Y de variance finie. L'espérance de Y s'écrit comme :

$$\mathbb{E}[Y] = \mathbb{E}[\mathbb{E}[Y|X]]$$

et la variance de Y s'écrit comme :

$$\text{Var}[Y] = \text{Var}[\mathbb{E}[Y|X]] + \mathbb{E}[\text{Var}[Y|X]]$$

Remarque E.2. Dans ce théorème, il faut simplement prendre garde aux variables sur lesquelles porte l'intégration : ainsi lorsqu'on écrit $\text{Var}[\mathbb{E}[Y|X]]$, l'espérance porte sur Y (puisque X est "fixé") et la variance porte sur X .

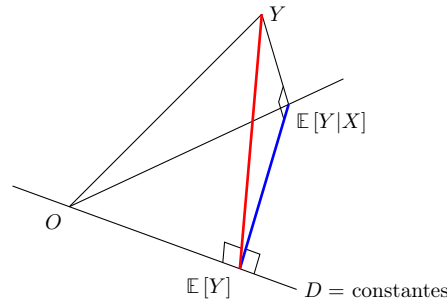


FIGURE E.1 – Illustration graphique du théorème de l'espérance totale

Exemple E.1. Soient X, Y de densité jointe : $f_{X,Y}(x, y) = 2e^{-(x+y)} \mathbb{1}_{0 \leq x \leq y}$. Calculer $\mathbb{E}[Y|X]$.

Solution : La loi marginale de X est de densité :

$$f_X(x) = \int_{\mathbb{R}} f_{X,Y}(x, y) dy = 2e^{-2x} \mathbb{1}_{[0, +\infty[}(x)$$

La densité conditionnelle de Y sachant $X = x$ s'écrit alors :

$$f_{Y|X}(y|x) = \begin{cases} \frac{f_{X,Y}(x, y)}{f_X(x)} & \text{si } f(x) > 0 \\ 0 & \text{sinon} \end{cases} \\ = e^{-(y-x)} \mathbb{1}_{y \geq x}(y)$$

D'où l'expression suivante est obtenue :

$$\mathbb{E}[Y|X = x] = \int_{\mathbb{R}} y f(y|x) dy = \int_x^{+\infty} y e^{-(y-x)} dy = x + 1$$

L'espérance conditionnelle recherchée est donnée par $\mathbb{E}[Y|X] = X + 1$.

Annexe F

Rappels de calcul différentiel

On se place dans $\mathbb{R}^n$ et la notation $\|\cdot\|$ désigne la norme euclidienne.

Définition F.1 (*Gradient*)

Une fonction $f : \mathbb{R}^n \mapsto \mathbb{R}$ est différentiable en a s'il existe une application linéaire notée *dérivée* $f'(a)$ ou *gradient* $\nabla f(a)^T$, telle que pour tout $h \in \mathbb{R}^n$, on ait :

$$f(a+h) = f(a) + \nabla f(a)^T \cdot h + \|h\| \cdot \epsilon(h)$$

où $\epsilon(\cdot)$ est une fonction continue en 0 vérifiant $\lim_{h \rightarrow 0} \epsilon(h) = 0$.

De plus, $f'(a)h = \nabla f(a)^T h$ est appelée *dérivée directionnelle* de f au point a selon la direction h .

Définition F.2 (*Convexité*)

Une fonction f définie sur un ensemble convexe $\mathbb{K}$ est convexe si elle vérifie :

$$\forall (x, y) \in \mathbb{K}^2, \forall \lambda \in [0, 1], f(\lambda x + (1-\lambda)y) \leq \lambda f(x) + (1-\lambda)f(y).$$

Pour un problème sans contraintes ou d'optimisation locale, le résultat suivant est vérifié.

Théorème F.1

Soit une fonction $f : \mathbb{R}^n \mapsto \mathbb{R}$ différentiable et $\hat{x}$ vérifiant (*globalement* ou *localement*) la relation $f(\hat{x}) \leq f(x)$, alors on a nécessairement :

$$\nabla f(\hat{x}) = 0.$$

Si de plus, la fonction f est convexe, alors la condition devient *suffisante*.

Annexe G

Algorithme d'estimation des indices de Sobol

Un algorithme (tiré de la thèse Pons-Salort (2013)) d'estimation des indices de Sobol par la méthode de Monte-Carlo est donné comme exemple ci-après.

Un hypercube latin est la généralisation du carré latin à un nombre arbitraire de dimensions ($p > 2$), chaque échantillon étant le seul dans chaque hyperplan aligné sur un axe qui le contient. Ici, nous ne nous sommes intéressés qu'aux indices de sensibilité du premier ordre et d'ordre total. Dans ce qui suit, nous détaillons comment nous avons calculé ces indices, en nous basant sur l'algorithme :

1. Créez deux matrices d'échantillons d'entrée générées avec le LHS : $\mathbf{M}_1$, typiquement appelée la matrice "échantillon", et $\mathbf{M}_2$, appelée la matrice "ré-échantillonnée". Elles comportent m lignes et p colonnes, où m est la taille de l'échantillon. Ces matrices sont de la forme :

$$\mathbf{M}_1 = \begin{bmatrix} x_1^{(1)} & x_2^{(1)} & \cdots & x_p^{(1)} \\ x_1^{(2)} & x_2^{(2)} & \cdots & x_p^{(2)} \\ \vdots & \vdots & & \vdots \\ x_1^{(m)} & x_2^{(m)} & \cdots & x_p^{(m)} \end{bmatrix} \text{ et } \mathbf{M}_2 = \begin{bmatrix} x_1^{(1')} & x_2^{(1')} & \cdots & x_p^{(1')} \\ x_1^{(2')} & x_2^{(2')} & \cdots & x_p^{(2')} \\ \vdots & \vdots & & \vdots \\ x_1^{(m')} & x_2^{(m')} & \cdots & x_p^{(m')} \end{bmatrix}$$

2. Pour chaque facteur d'entrée X_i , créez une matrice $\mathbf{N}_i$ utilisant $\mathbf{M}_1$ et $\mathbf{M}_2$ en ré-échantillonnant tous les facteurs sauf X_i :

$$\mathbf{N}_i = \begin{bmatrix} x_1^{(1')} & x_2^{(1')} & \cdots & x_i^{(1)} & \cdots & x_p^{(1')} \\ x_1^{(2')} & x_2^{(2')} & \cdots & x_i^{(2)} & \cdots & x_p^{(2')} \\ \vdots & \vdots & & \vdots & & \vdots \\ x_1^{(m')} & x_2^{(m')} & \cdots & x_i^{(m)} & \cdots & x_p^{(m')} \end{bmatrix}$$

3. Exécutez le modèle pour chaque échantillon d'entrée dans $\mathbf{M}_1$, $\mathbf{M}_2$, $\mathbf{N}_1$, ..., $\mathbf{N}_p$. Ainsi, le modèle est exécuté $m(p+2)$ fois.
4. Calculez les estimées de la moyenne et de la variance inconditionnelles de la sortie, $\hat{\mathbb{E}}(Y)$ et $\widehat{\text{Var}}(Y)$, en utilisant les valeurs de Y calculées avec l'échantillon $\mathbf{M}_1$:

$$\hat{\mathbb{E}}(Y) = \frac{1}{m} \sum_{r=1}^m f(x_1^{(r)}, x_2^{(r)}, \dots, x_p^{(r)})$$

$$\widehat{\text{Var}}(Y) = \frac{1}{m-1} \sum_{r=1}^m f^2(x_1^{(r)}, x_2^{(r)}, \dots, x_p^{(r)}) - \hat{\mathbb{E}}^2(Y)$$

5. Pour chaque facteur d'entrée X_i , calculez son indice de sensibilité d'ordre un :

$$S_i = \frac{\text{Var}(\mathbb{E}(Y|X_i))}{\text{Var}(Y)} = \frac{U_i - \mathbb{E}^2(Y)}{\text{Var}(Y)}$$

où $U_i = \mathbb{E}(\mathbb{E}^2(Y|X_i))$ est estimée à partir du produit des valeurs de f calculées à l'aide de la matrice échantillon fois les valeurs de f calculées à l'aide de $\mathbf{N}_i$ (c'est-à-dire en faisant varier tous les paramètres sauf X_i)

$$\hat{U}_i = \frac{1}{m-1} \sum_{r=1}^m f\left(x_1^{(r)}, x_2^{(r)}, \dots, x_p^{(r)}\right) f\left(x_1^{(r')}, x_2^{(r')}, \dots, x_{i-1}^{(r')}, x_i^{(r)}, x_{i+1}^{(r')}, \dots, x_p^{(r')}\right).$$

Index

A

alphabet fermé par préfixe	90
amortissement	59
analyse	
d'incertitude	147
harmonique	129
indicielle	124
analyse de sensibilité	
globale	149
approche	
directe	152
inverse	152
non-paramétrique	152
paramétrique	156
automate	
à états finis	81
co-accessible	83
hybride	111
produit	87
temporisé fini	105
automates bisimilaires	88

B

biais	141
d'un estimateur	133
bisimulation d'automate	88
blocage	
fatal	89
vivant	89

C

calcul d'erreur	161
coefficient	
de détermination	142
de régression standardisé	167
coefficient de régression standardisé	167
concaténation de langage	90
constante de temps	69
covariance	185
critère	
d'information d'Akaike	143
de Jury	51
de Routh	50

D

décomposition de Sobol	170
déphasage minimal	68
deadlock	<i>voir</i> blocage fatal
densité de probabilité	
d'une loi de Rayleigh	119
de la loi d'Erlang	119
diagramme de Bode	59
dispersion statistique	150
distribution exponentielle	117

E

écart-type	160
échelon unité ou échelon de Heaviside	17
entrée	7
de commande	7
de perturbation	7
entropie	
d'une distribution de probabilité ...	158
de Shannon	157
erreur	148
accidentelle	150
aléatoire	150
de mesure	131, 150
de modèle	131
quadratique moyenne	141
de prédiction	144
intégrée de prédiction	144
systématique	150
espérance	
conditionnelle	187
d'une variable aléatoire	185
mathématique	185
espace paramétrique	156
estimateur	132
empirique	153
estimation paramétrique	131
état	
accessible	83
co-accessible	83
marqué	82
évaluation	
graphique	140
exactitude de mesure	150

F

facteur de forme 119

Factors

Fixing 149

Mapping 149

Prioritisation 149

fermeture

de Kleene 90

par préfixe 90

fonction

coût 137

d'observation 7

de répartition 183

de transfert 48

formule

de Bayes 187

de Sylvester 31

des probabilités totales 187

fréquence réduite 17

G

gain 59

gradient 191

H

histogramme 155

I

identification 123

impulsion

de Dirac 17

discrète 17

incertitude 148

élargie 161

de mesure 151

de type A 152

de type B 157

incertitude-type 151, 160, 162, 163

indice

de sensibilité

de Sobol d'ordre 1 169

de Sobol d'ordre 2 170

de Sobol d'ordre 3 170

de Sobol d'ordre r 170

total 171

de sensibilité local 167

local normalisé 167

logarithmique 167

SRC 167

intensité 119

du processus de Poisson 117

interpolation de Sylvester 32

L

langage 89, 90

marqué 90

reconnu par l'automate 90

livelock voir blocage vivant

M

méta-modèle 165

marquage 91

matrice

d'évolution 23

d'entrée 23

d'observation 23

d'échantillonnage 168

de transition 29

de transmission directe 23

hessienne 138

jacobienne 28

mesurage 149

mesurande 150

méthode

bayésienne 131

d'analyse de sensibilité

locale 148

d'analyse de sensibilité locale 166

d'échantillonnage stratifié par

hypercube latin 165

de Newton 138

de Strejč 126

des '5M' 150

des modes 31

des moindres carrés 131, 132

non-linéaires 132

des noyaux 156

du développement en série 30

du gradient 138

du maximum de vraisemblance 131

fréquentiste 131

modélisation 6

modèle 6

linéaire 142

polynomial 142

additif 170

candidat 130

de comportement 131

de connaissance 7, 131

sémantique 8

surrogé 165

téléonomique 136

moment d'une variable aléatoire 185

moyenne empirique 153

N

non-identifiabilité 134

P

parallélisation des automates

non temporisés 87

temporisés avec gardes 109

paramétrique 152

paramètre 7

- point d'équilibre 24
 - asymptotiquement stable 24, 28
 - exponentiellement stable 24, 28
 - stable au sens de Lyapunov 24
- prédiction 140
- principe du maximum d'entropie 157
- problème d'optimisation 132
- problème d'optimisation
 - continue 137
 - global 137
 - local 137
- processus
 - à incréments indépendants 117
 - et stationnaires 117
 - de Poisson 117
 - densité de probabilité 116
 - moyenne 117
 - variance 118
 - sans mémoire 118
- produit synchrone des automates non
 - temporisés 86
- protocole leave-one-out 145
- pulsation de brisure 59
- Q**
- quantile 161
- quantité d'intérêt 160
- R**
- répétabilité des résultats de mesurage .. 150
- réponse
 - en fréquence 71
 - impulsionnelle 47
 - indicielle 47, 69, 70
- réseau
 - borné 96
 - de Petri
 - autonome 91
 - discret autonome marqué 92
 - p-temporisé 113
 - synchronisé 102
 - t-temporisé 113
 - persistant 97
 - pseudo-vivant 96
 - sauf 96
 - vivant 95
- représentation fonctionnelle 8
- reproductibilité des résultats de mesurage
 - 150
- retard pur 59, 68
- S**
- séquence binaire pseudo-aléatoire 135
- séquence d'évènements 77
- séquence de franchissement 93
- score d'un évènement 105
- simulation 6
 - d'automate 88
- sortie 7
- sous-lissage 155
- structure d'horloge 104
- sur-lissage 155
- système 5
 - à événements discrets 80
 - à événements discrets
 - non temporisé 81
 - à temps discret sous forme d'état ... 21
 - causal 16
 - continu 15
 - de systèmes 6
 - discret 15
 - dynamique 5, 9, 16
 - invariant dans le temps 16, 22
 - linéaire 15
 - monovariable 15
 - multi-agents 5
 - multivariable 15
 - non linéaire 15
 - passe-tout 66
 - permanent *voir* système invariant dans
 - le temps
 - stable
 - asymptotiquement et
 - exponentiellement 33, 35
 - entrée bornée sortie bornée 25, 49
 - stationnaire *voir* système invariant dans
 - le temps
 - statique 16
- T**
- taux du processus de Poisson . *voir* intensité
 - du processus de Poisson
- temps
 - de réponse 69
 - du premier maximum 70
- trajectoire stable 25
- transformée
 - de Fourier 37
 - de Laplace 39
 - en z 43
- transition
 - franchie 92
 - franchissable 102
 - puits 93
 - réceptive à l'évènement 102
 - source 93
 - validée 92
- transmittance 48
- trappe 97
- V**
- valeur vraie 150
- validation croisée 144

variable		
aléatoire	183	
'moyenne empirique'	186	
à densité	183	
continue	183	
discrète	183	
variance		
conditionnelle	188	
Cutting	149	
d'une variable aléatoire	185	
empirique	154	
verrou	97	

Bibliographie

- Christoph Aistleitner, Markus Hofer, and Robert Tichy. A central limit theorem for latin hypercube sampling with dependence and application to exotic basket option pricing. *International Journal of Theoretical and Applied Finance*, 15 :1–17, 2012.
- Jacques Antoine, Jean-Luc Collette, Gilles Duc, Hervé Guéguen, and Daniel Poulton. *Signaux et systèmes*. Polycopié Supélec. Supélec, 2014.
- BIPM, IEC, IFCC, ISO, IUPAC, IUPAP, and OIML. *Vocabulaire international des termes généraux et fondamentaux de métrologie (VIM)*. 2 edition, 1993.
- BIPM, IEC, IFCC, ILAC, ISO, IUPAC, IUPAP, and OIML. *Evaluation of measurement data - Guide to the expression of uncertainty in measurement (GUM)*, volume 100 of *rapport technique*. Joint Committee for Guides in Metrology, JCGM, 1 edition, Septembre 2008.
- George E.P. Box. *Robustness in the strategy of scientific model building*. in Launer, R. L. ; Wilkinson, G. N., *Robustness in Statistics*, Academic Press, pp. 201–236, 1979.
- G.W. Brams. *Réseaux de Petri : théorie et pratique*. Editions Masson, 1983.
- Christos G. Cassandras and Stéphane Lafortune. *Introduction to discrete event systems*. Springer, 2nd edition, 2008.
- Pierre Chlique. *Grafcet et réseaux de Petri*. Polycopié Supélec. Supélec, 2000.
- René David and Hassane Lotfi Alla. *Du Grafcet aux réseaux de Petri*. Hermes, 1989.
- John Joachim d’Azzo and Constantine Dino Houpis. *Linear control system analysis and design. Conventional and modern*. McGraw-Hill Higher Education, 3rd edition, 1988.
- Philippe de Larminat, Yves Thomas, and Yves Thomas. *Automatique des systèmes linéaires*. Flammarion Sciences, 1977.
- Emmanuel Godoy. *Régulation industrielle*. 2014.
- Clifford M. Hurvich and Chih-Ling Tsai. Model selection for extended quasi-likelihood models in small samples. *Biometrics*, 51 :1077–1084, 1995.
- E.T. Jaynes. Information theory and statistical mechanics. *Physical Review*, 106 :620–630, 1957.
- Henry George Liddell and Robert Scott. *An intermediate Greek-English lexicon*. Clarendon Press, 1889.
- Michel Maurin. Jeux de variance et d’entropie. *XXXI Journées de Statistique, Grenoble, France, 17-21 Mai*, 1999.
- Jurgen Mayer, Khaled Khairy, and Jonathon Howard. Drawing an elephant with four complex parameters. *American Journal of Physics*, 78(6) :648–649, March 2010.
- Tadao Murata. Petri nets : Properties, analysis and applications. *Proceedings of the IEEE*, 77(4) : 541–580, 1989.
- John P. Norton. *An introduction to identification*. Academic Press., 1986.

- Alan V. Oppenheim, Alan S. Willsky, and Syed Hamid Nawab. Signals and systems. *Prentice Hall*, 1997.
- André Pacaud. *Signaux et systèmes linéaires*. Technosup, 2001.
- Pascal Pernot. *Initiation à la propagation des incertitudes*. Formations doctorales. Laboratoire de Chimie Physique, UMR8000, CNRS/Univ. Paris-Sud, 2018.
- Margarita Pons-Salort. *Modélisation mathématique des interactions multi-hôtes et multi-pathogènes en épidémiologie des maladies infectieuses : conséquences sur la persistance, l'émergence et le contrôle de pathogènes*. Thèse de doctorat Biostatistiques/Biomathématiques Paris 6, 2013.
- Andrea Saltelli, Stefano Tarantola, and Francesca Campolongo. Sensitivity analysis as an ingredient of modeling. *Statistical Science*, 15(4) :377–395, March 2000.
- Cristina Stoica Maniu, Guillaume Sandou, and Véronique Letort-Le Chevalier. Feedback on innovative pedagogy for teaching systems modeling. In *12th IFAC Symposium on Advances in Control Education, Invited Demonstration Session*, pages 1–6, Philadelphia, PA, USA, July 7-9 2019.
- Yves Thomas. *Signaux et systèmes linéaires*. Editions Masson, 1992.
- Yves Thomas. *Signaux et systèmes linéaires – Exercices corrigés*. Editions Masson, 1993.
- Guy Vidal-Naquet and Annie Choquet-Geniet. *Réseaux de Petri et systèmes parallèles*. Armand Colin, 1992.
- Larry Wasserman. *All of nonparametric statistics*. Springer. Springer Texts in Statistics, New York, 2006.
- Janan Zaytoon. *Systèmes dynamiques hybrides*. Traité IC2 – Série Systèmes automatisés. Hermes, 2001.