

Université de Liège
Faculté des Sciences Appliquées
Département d'Électricité, Électronique, et Informatique

Modélisation et analyse des systèmes

Notes de cours

Année académique 2014–2015

Table des matières

Table des matières	3
Préface	7
1 Signaux et systèmes	9
1.1 Exemples de signaux et systèmes	9
1.1.1 Télécommunications	9
1.1.2 Enregistrement et traitement de données	12
1.1.3 Systèmes contrôlés	14
1.2 Signaux	15
1.2.1 Signaux comme fonctions	15
1.2.2 Échantillonnage et quantification	16
1.2.3 Domaine de signaux	18
1.2.4 Image de signaux	20
1.2.5 Espaces vectoriels de signaux	21
1.3 Systèmes	22
1.3.1 Systèmes comme fonctions de fonctions	22
1.3.2 Mémoire d'un système	23
1.3.3 Équations aux différences et équations différentielles	27
1.3.4 Causalité	28
1.3.5 Interconnexions de systèmes	28
2 Le modèle d'état	31
2.1 Structure générale d'un modèle d'état	31
2.2 Automates finis	32
2.2.1 Définition et représentation	32
2.2.2 Étude des automates finis	33
2.3 Modèles d'état en temps-discret	33
2.4 Modèles d'état en temps-continu	36
2.5 Interconnexions	41
2.6 Extensions	41
2.6.1 Modèles d'état dépendant du temps	42
2.6.2 Modèles d'état non déterministes	42

2.6.3	Modèles d'état hybrides	42
3	Systèmes linéaires invariants : représentation convolutionnelle	43
3.1	Linéarité et invariance	43
3.2	La convolution pour les systèmes discrets	45
3.3	L'impulsion de Dirac	47
3.4	Convolution de signaux continus	50
3.5	Causalité, mémoire, et temps de réponse des systèmes LTI . .	51
3.6	Conditions de repos initial	54
4	Modèles d'état linéaires invariants	57
4.1	Le concept d'état	58
4.2	Modèle d'état d'un système différentiel ou aux différences . .	61
4.2.1	Cas continu	61
4.2.2	Cas discret	65
4.3	Solution des équations d'état	67
4.3.1	Modèles d'état discrets	67
4.3.2	Modèles d'état continus	68
4.3.3	Calcul de l'exponentielle matricielle	69
4.4	Transformations d'état	71
4.5	Construction de modèles linéaires invariants	72
4.5.1	Séparation espace-temps	72
4.5.2	Localisation temporelle	73
4.5.3	Linéarisation	74
4.6	Représentation interne ou externe?	76
5	Décomposition fréquentielle des signaux	79
5.1	Signaux périodiques	80
5.1.1	Signaux périodiques continus	80
5.1.2	Signaux périodiques discrets	80
5.2	Série de Fourier en temps discret	81
5.2.1	Bases de signaux dans l'espace $\ell_2[0..N)$	81
5.2.2	Bases orthogonales	82
5.2.3	Série de Fourier d'un signal périodique en temps-discret	83
5.2.4	La base harmonique est spectrale pour les opérateurs linéaires invariants	84
5.3	Série de Fourier en temps continu	86
5.3.1	La base harmonique de $L_2[0, T)$	86
5.3.2	Série de Fourier d'un signal périodique en temps-continu	87
5.3.3	Convergence des séries de Fourier	88

5.3.4	La base harmonique est spectrale pour les opérateurs linéaires invariants	90
6	Transformées de signaux discrets et continus	93
6.1	Introduction	93
6.2	Exponentielles complexes et systèmes LTI	94
6.3	Transformée de Laplace, transformée en z , et transformée de Fourier	96
6.4	Transformées inverses et interprétation physique	97
6.5	Région de convergence (ROC)	99
6.6	Propriétés élémentaires	102
6.6.1	Linéarité	102
6.6.2	Décalage temporel et fréquentiel	102
6.7	La dualité convolution – multiplication	103
6.8	Différentiation et intégration	104
7	Analyse de systèmes LTI par les transformées de Laplace et en z	107
7.1	Fonctions de transfert	107
7.2	Bloc-diagrammes et algèbre de fonctions de transfert	110
7.3	Résolution de systèmes différentiels ou aux différences initialement au repos	111
7.3.1	Transformées inverses et décomposition en fractions simples : Systèmes continus	114
7.3.2	Transformées inverses et décomposition en fractions simples : Systèmes discrets	116
8	Réponses transitoire et permanente d'un système LTI	119
8.1	Transformée unilatérale	120
8.1.1	Transformée de Laplace	120
8.1.2	Transformée en z	121
8.2	Résolution de systèmes différentiels ou aux différences avec des conditions initiales non nulles	122
8.2.1	Réponse libre et réponse forcée	122
8.2.2	Les conditions initiales vues comme des impulsions de Dirac	122
8.2.3	Modes propres d'un système différentiel	124
8.3	Stabilité d'un système LTI	124
8.3.1	Critère de Routh-Hurwitz	126
8.4	Réponse transitoire et réponse permanente	128
9	Réponse fréquentielle d'un système LTI	129
9.1	Réponse fréquentielle d'un système LTI	129

9.2	Diagrammes de Bode	131
9.3	Bande passante et constante de temps	133
9.4	Réponses d'un système continu du premier ordre	134
9.5	Réponse d'un système du deuxième ordre	136
9.6	Fonctions de transfert rationnelles	139
9.6.1	Effet d'un pôle et d'un zéro sur la réponse fréquentielle	140
9.6.2	Filtres de Butterworth	140
9.7	Réponses d'un système discret du premier ordre	141
9.8	Exemple d'analyse fréquentielle et temporelle d'un modèle LTI	142
10	Des séries de Fourier aux transformées de Fourier	145
10.1	Signaux en temps continu	145
10.2	De la série de Fourier DD à la transformée de Fourier DC . . .	148
10.3	Transformée de Fourier de signaux périodiques continus . . .	150
10.4	Transformée de Fourier de signaux périodiques discrets . . .	151
10.5	Propriétés élémentaires des séries et transformées de Fourier	152
10.5.1	Transformées de signaux réfléchis et changement d'échelle	152
10.5.2	Décalage temporel et fréquentiel	153
10.5.3	Signaux conjugués	153
10.5.4	Relation de Parseval	154
10.5.5	Dualité convolution-multiplication	154
10.5.6	Intégration-Différentiation	155
10.5.7	Dualité des transformées de Fourier	156
11	Applications élémentaires des transformées de Fourier	157
11.1	Fenêtrages temporel et fréquentiel	157
11.1.1	Transformée d'un rectangle	157
11.1.2	Troncature d'un signal par fenêtrage rectangulaire . . .	159
11.1.3	Filtrage passe-bande d'un signal par fenêtrage rectangulaire	160
11.2	Échantillonnage	163
11.3	Sous-échantillonnage et repliement de spectre	166
11.4	Interpolation	167
11.5	Synthèse d'un filtre à réponse impulsionnelle finie	168
A	Espaces vectoriels et applications linéaires	173
A.1	Espace vectoriel	173
A.2	Dimension d'un espace vectoriel	174
A.3	Base d'un espace vectoriel	175
A.4	Application linéaire	175

Préface

Les notes qui suivent constituent un support écrit pour le cours de Modélisation et Analyse des Systèmes SYST002. Elles se basent principalement sur les ouvrages de référence suivants :

- *Signals and Systems* (2nd edition), A. V. Oppenheim & A. S. Willsky, Prentice-Hall 1997.
- *Modern Signals and Systems*, Kwakernaak & Sivan, Prentice-Hall, 1991.
- *Linear Systems and Signals*, B. P. Lathi, Berkeley-Cambridge Press, 1992.
- *The structure and Interpretation of Signals and Systems*, E. A. Lee & P. Varaiya, Addison-Wesley, 2003.

Les dernières révisions apportées au syllabus datent d'Aout 2014.

Chapitre 1

Signaux et systèmes

Un signal est un vecteur d'information. Un système opère sur un signal et en modifie le contenu sémantique. La théorie des systèmes fournit un cadre mathématique pour l'analyse de n'importe quel processus *vu comme système*, c'est-à-dire restreint à sa fonctionnalité de véhicule d'information. Elle constitue un maillon important entre les mathématiques sur lesquelles elle repose (équations différentielles et algèbre linéaire dans le cadre de ce cours) et les technologies de l'information auxquelles elle fournit une assise théorique (télécommunications, traitement de signal, théorie du contrôle).

Dans le cadre de ce cours, restreint aux systèmes dits *linéaires* et *invariants*, la théorie des systèmes repose sur des outils mathématiques élémentaires. Sa difficulté réside davantage dans le degré d'abstraction et de modélisation qu'elle implique. Que faut-il conserver d'un processus donné pour le décrire en tant que système ? Quelles en sont les descriptions mathématiques *efficaces*, c'est-à-dire qui conduisent à une analyse pertinente du processus considéré et qui permettent d'orienter des choix technologiques ?

La théorie des systèmes linéaires et invariants est un exemple remarquable de compromis entre le degré de généralité auquel elle prétend et le degré d'efficacité qu'elle a démontré dans le développement des technologies de l'information. Avant d'en définir les contours de manière rigoureuse, nous tenterons dans ce premier chapitre de dégager quelques aspects essentiels des signaux et systèmes à partir d'applications concrètes.

1.1 Exemples de signaux et systèmes

1.1.1 Télécommunications

Le réseau global de télécommunications que nous connaissons aujourd'hui est sans doute la grande révolution technologique de notre temps. Il fournit aussi une pléthore d'illustrations pour la théorie des systèmes, à commencer par la simple communication téléphonique à longue distance. La Figure 1.1 illustre de la manière la plus élémentaire le système « communica-

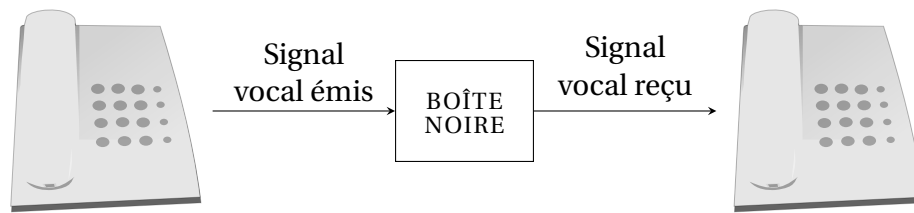


FIGURE 1.1 – Système « communication téléphonique » : le signal vocal reçu (signal de sortie) doit être le plus proche possible du signal vocal émis (signal d'entrée).

tion téléphonique », constitué d'une boîte noire dont le rôle est de transmettre un signal vocal de la manière la plus fidèle possible. Ce qui caractérise le système téléphone, ce n'est pas la technologie qui est utilisée pour transmettre le signal (par exemple transmission par satellite ou par fibre optique), mais la manière dont le téléphone *opère* sur le signal émis, appelé entrée du système, pour produire le signal reçu, appelé sortie du système. On pourra par exemple caractériser le système « communication téléphonique » par sa *réponse fréquentielle* (Chapitre 9) parce qu'un signal grave (ou *basse fréquence*) ne sera pas transmis de la même manière qu'un signal aigu (ou *haute fréquence*).

Si l'on désire comprendre les limitations du système « communication téléphonique » ou améliorer ses capacités de transmission, il faut en affiner la description. Le signal vocal émis est une onde acoustique, convertie en signal électrique par le poste de téléphone. Ce signal électrique est ensuite transmis sur une paire torsadée de fils de cuivres à une centrale téléphonique. Il est alors numérisé, c'est-à-dire transformé en une séquence de bits, généralement combiné avec d'autres signaux numériques, et modulé, c'est-à-dire monté sur une porteuse haute fréquence, avant d'être transmis sur un canal à haut débit (fibre optique, câble coaxial, ou satellite). A l'arrivée, il subit les opérations inverses : il est démodulé, reconverti en signal électrique, et transmis sur une nouvelle paire torsadée au poste récepteur avant d'être reconverti en onde acoustique. Ces étapes successives de transformation du signal original sont représentées à la Figure 1.2 sous la forme d'un *bloc-diagramme*. La boîte noire « communication téléphonique » de la Figure 1.1 a été « éclatée » en une cascade de nouvelles boîtes noires. Chacune de ces boîtes noires est un « système », opérant sur un signal d'entrée pour produire un signal de sortie.

L'analyse systémique du système « communication téléphonique » portera non pas sur la technologie des différents composants (micro, haut-parleur, canal de transmission, modulateur, convertisseur analogique-numérique,...) mais sur la caractérisation de chacun des composants du point de vue de

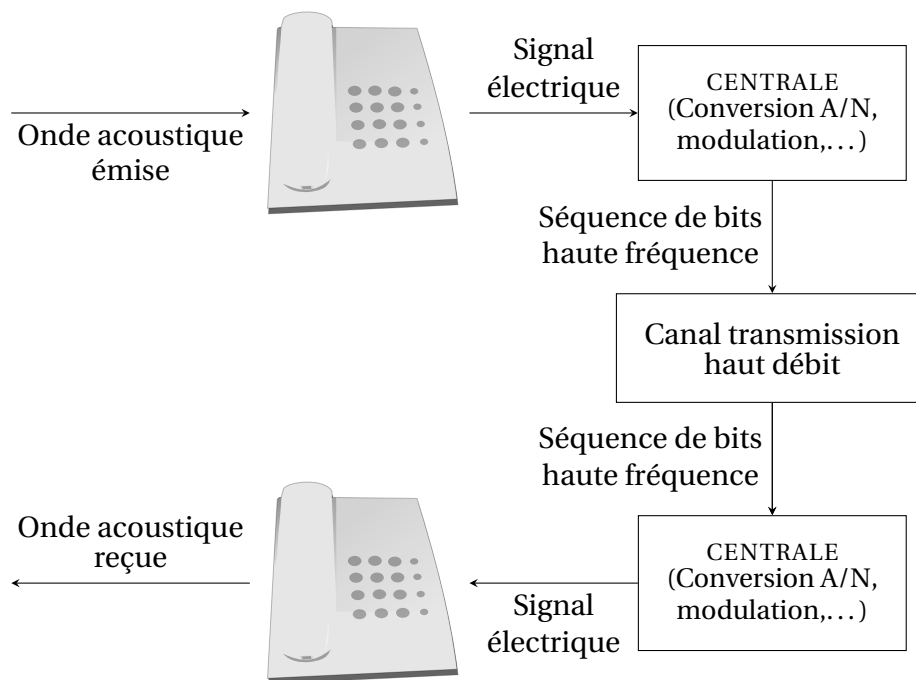


FIGURE 1.2 – Bloc-diagramme du système « communication téléphonique ».

leur participation à la dégradation finale du signal transmis : par exemple la réponse fréquentielle du micro et du haut-parleur, la fréquence d'échantillonnage lors de la conversion analogique-numérique, le retard cumulé entre l'instant d'émission et l'instant de réception.

La description du système « communication téléphonique » de la Figure 1.2 est évidemment très simplifiée mais elle est *générique*. Elle peut être étendue à d'autres types de communication qui font un usage intensif de traitement de signal (internet, télévision numérique, appareils GSM,...). Bien qu'ils se distinguent par des défis technologiques distincts, les blocs de base de ces processus ne sont pas très différents des sous-systèmes décrits dans notre exemple et font appel au même type de caractérisation systémique.

Quelle que soit la nature des signaux transmis (signal audio, vidéo, image, données informatiques), les technologies modernes de communication reposent sur leur numérisation et sur les techniques qui permettent leur transmission sur de longues distances et à très haut débit (modulation, démodulation, compression, décompression). La théorie des systèmes et signaux fournit les bases mathématiques de ces diverses opérations.

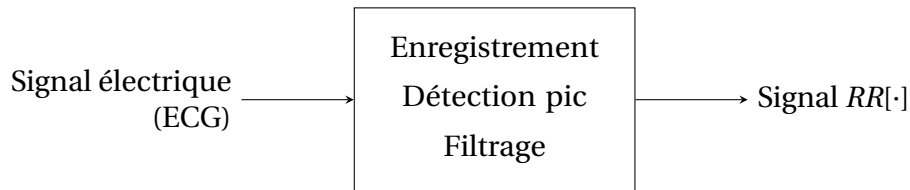
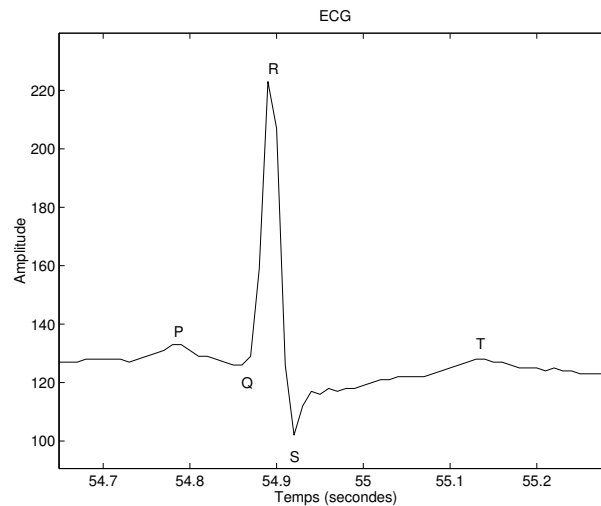
FIGURE 1.3 – Extraction du signal RR à partir d'un électrocardiogramme.

FIGURE 1.4 – Electrocardiogramme (ECG) d'un battement cardiaque.

1.1.2 Enregistrement et traitement de données

Les possibilités sans cesse croissantes de stocker un très grand nombre de données à très faible coût ont rendu presque routinière dans la majorité des disciplines scientifiques la collecte systématique au cours du temps de signaux les plus divers pouvant servir d'indicateurs économiques, sécuritaires, écologiques, ou autres. Ces *séries temporelles* sont ensuite analysées par un expert humain ou logiciel pour en extraire une information pertinente. On cherche notamment à extraire celle-ci en dépit des fluctuations aléatoires qui affectent généralement le processus étudié.

La Figure 1.3 illustre par exemple l'enregistrement du rythme cardiaque sous la forme d'un signal appelé *signal RR*. Deux électrodes mesurent l'activité électrique qui accompagne chaque contraction du muscle cardiaque. Ce signal électrique a typiquement la forme représentée à la Figure 1.4. Le pic dominant, baptisé pic R, est facile à détecter même si le signal est bruité.

L'intervalle de temps qui sépare deux pics R successifs donne une image fidèle du rythme cardiaque. Le signal $RR[.]$ est construit en numérotant les battements de cœur $k = 1, 2, \dots$ au cours de l'enregistrement et en associant au battement k l'intervalle de temps $RR[k]$ entre le battement $k - 1$ et le battement k . Un cœur parfaitement régulier dont les battements se succéderaient à intervalles constants donnerait lieu à un signal RR constant. Ce sont les

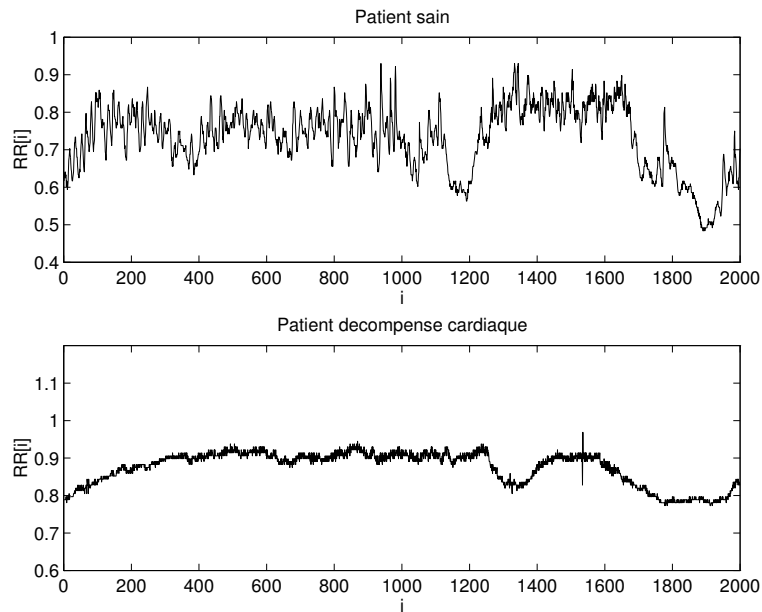


FIGURE 1.5 – Signal RR d'un patient sain et d'un patient malade. (2000 battements, soit une quinzaine de minutes d'enregistrement).

variations du signal RR autour de cette période constante qui caractérisent la variabilité cardiaque, dont l'analyse peut aider le cardiologue à diagnostiquer une maladie cardiaque, comme illustré à la Figure 1.5.

L'enregistrement du signal RR à partir du signal électrique généré au niveau du muscle cardiaque s'accompagne d'une série d'opérations que l'on peut préciser, comme dans le cas du système « communication téléphonique », en éclatant le système de la Figure 1.3 en un bloc-diagramme plus détaillé. Au minimum, l'appareil enregistreur doit *échantillonner* le signal électrique mesuré (typiquement à une fréquence de 128 Hz, c'est-à-dire 128 mesures par seconde) et *détecter* les pics R successifs du signal échantillonné (par exemple en décidant que toutes les valeurs supérieures à une valeur de seuil donnée sont interprétées comme faisant partie d'un pic R). Dans tous les enregistreurs commercialisés, ces fonctions de base sont accompagnées d'une opération de *filtrage* pour éliminer les artefacts du signal construit. Dans l'exemple considéré, on peut vouloir chercher à éliminer du signal RR des artefacts de nature physiologique (par exemple des extrasystoles) ou des artefacts liés à l'enregistrement (mouvements du patient, ...).

Toutes les applications modernes de collecte de données comprennent une opération d'échantillonnage et une opération de filtrage. Un nombre croissant d'opérations de traitement de base, telles que la détection du pic R dans l'électrocardiogramme, sont intégrées à l'enregistreur et effectuées en temps-réel. La théorie des signaux et systèmes fournit la base mathématique pour l'analyse et la conception de ces opérations.

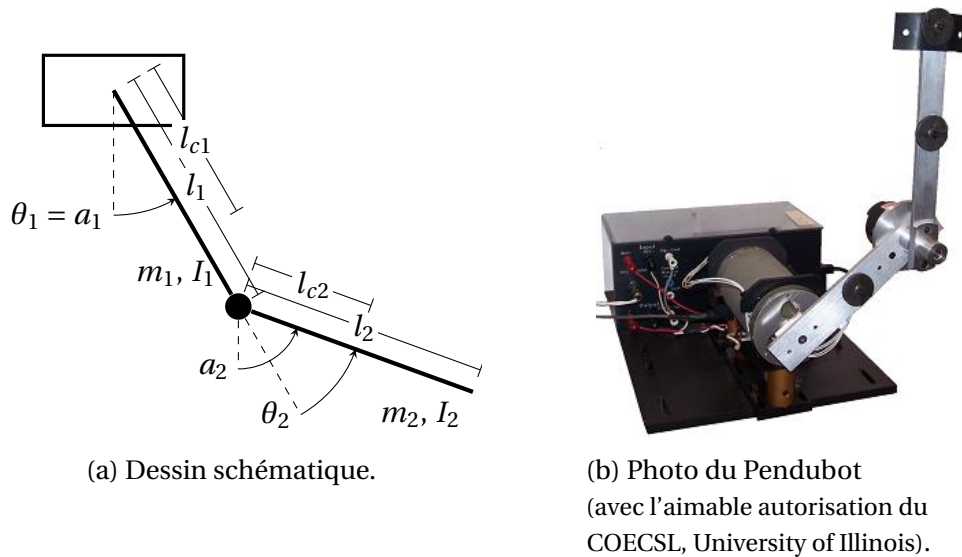


FIGURE 1.6 – Le *pendubot*, un double bras de robot commandé uniquement à l'épaule.

1.1.3 Systèmes contrôlés

Un signal peut être enregistré pour analyser l'évolution d'un processus au cours du temps mais il peut également être utilisé en temps réel pour contrôler un processus par *rétroaction* (ou *feedback*). Les exemples de systèmes contrôlés envahissent également la technologie actuelle, tant dans le domaine grand public (positionnement de la tête de lecture et stabilisation du CD, contrôle de l'injection dans les moteurs, système de freinage ABS) que dans la robotique de pointe (microchirurgie, hélicoptères autonomes, ...).

Le *pendubot*, illustré à la Figure 1.6, est un dispositif de laboratoire à finalité pédagogique. Il peut être modélisé comme un double pendule plan, c'est-à-dire un système mécanique possédant deux degrés de liberté en rotation (épaule et coude). Un couple moteur peut être appliqué à l'épaule au moyen d'un petit moteur électrique. Stabiliser le pendubot dans la configuration instable où les deux bras sont alignés vers le haut à partir de la configuration stable où les deux bras sont alignés vers le bas est un exemple d'acrobatie difficile à réaliser manuellement mais relativement simple à mettre en œuvre avec un dispositif disponible en laboratoire.

Le bloc-diagramme de la Figure 1.7 est le bloc-diagramme typique d'un système contrôlé : on y distingue deux sous-systèmes, interconnectés non pas en série comme dans nos exemples précédents, mais selon une structure bouclée. Le sous-système PENDUBOT représente le processus à contrôler. Le couple moteur $u(\cdot)$ appliqué à l'épaule du pendubot est le signal d'entrée (ou de commande) ; il influence la trajectoire des deux bras du robot. Un capteur est fixé sur chacune des articulations de manière à mesurer la position

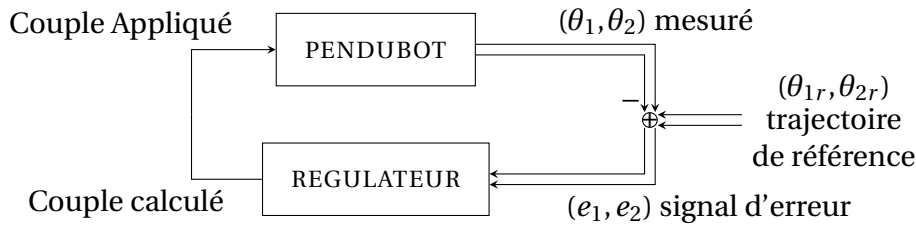


FIGURE 1.7 – Bloc-diagramme illustrant le contrôle du pendubot.

angulaire θ_1 et θ_2 de chacun des bras. Le vecteur $(\theta_1(\cdot), \theta_2(\cdot))$ est le signal de sortie (ou de mesure) du pendubot, qui peut être utilisé pour calculer la commande par rétroaction. Le sous-système REGULATEUR représente le système de contrôle du pendubot. Son entrée est un signal d'erreur entre la sortie de *référence* (une trajectoire particulière pour θ_1 et θ_2 qui réalise l'acrobatie désirée) et la sortie « mesurée ». Sur base de cette erreur, l'algorithme de contrôle calcule le couple moteur à appliquer pour réduire l'erreur de régulation.

Si le signal (θ_1, θ_2) désigne l'évolution temporelle des angles de chacun des bras du pendubot et si le signal u représente le couple mécanique appliqué à l'épaule, le sous-système REGULATEUR reliant ces signaux devrait, comme dans le cas des exemples précédents, être éclaté en un bloc-diagramme plus détaillé; on y trouverait en série le système CAPTEUR ANGULAIRE, le système ALGORITHME DE CONTRÔLE, et le système ACTIONNEUR DE COMMANDE (un moteur électrique dans le cas présent). La sortie du capteur de mesure angulaire est un signal échantillonné dans le temps et également quantifié, c'est-à-dire prenant ses valeurs dans un ensemble discret (si l'angle est codé au moyen de 8 bits, le capteur aura une résolution de $360/256$ degrés); la sortie du système ALGORITHME DE CONTRÔLE est un signal de commande digitalisé qui sera converti en tension électrique et finalement en couple mécanique par le système ACTIONNEUR. Faut-il tenir compte de toutes ces opérations de transformation dans le calcul d'une loi de commande? Comment faut-il modéliser le système pendubot? Peut-on négliger la dynamique du moteur électrique par rapport à celle du pendubot? Encore une fois, ces questions relèvent essentiellement de la théorie des signaux et systèmes.

1.2 Signaux

1.2.1 Signaux comme fonctions

Chaque système rencontré dans les exemples qui précèdent transforme un signal d'entrée en signal de sortie. Ces signaux sont de nature extrêmement diverse. Leur représentation unifiée exige de les appréhender comme *fonctions*, chaque fonction étant caractérisée par son *domaine*, son *image*, et la manière dont le domaine est appliqué sur l'image.

Un signal audio comme le signal vocal émis dans le système « communication téléphonique » est une onde acoustique qui évolue continûment au cours du temps. Son domaine est un continuum à une dimension, mathématiquement représenté par la droite réelle $\mathbb{R} = (-\infty, +\infty)$, et son image est un intervalle de pressions $(P_{\min}, P_{\max})$ exprimées dans une unité particulière. On peut formaliser cette caractérisation « fonctionnelle » de la manière classique :

$$\text{audio} : \mathbb{R} \rightarrow (P_{\min}, P_{\max}) : t \mapsto \text{audio}(t).$$

Le signal $\text{audio}(\cdot)$ est un exemple de signal *en temps-continu*. Il dépend d'une seule variable indépendante, qui peut prendre un continuum de valeurs, et qui est ordonnée, c'est-à-dire qu'on peut lui associer une notion de passé et de futur. Dans la théorie des systèmes, tous les signaux dépendant d'une seule variable *réelle* sont appelés signaux *en temps-continu*, même si la variable indépendante ne désigne pas nécessairement le temps. Les signaux qui décrivent l'évolution temporelle d'une variable physique (courant ou tension dans un système électrique, position ou vitesse dans un système mécanique) sont tous de cette nature.

Une fois son domaine et son image définis, le signal $\text{audio}(\cdot)$ peut être caractérisé sous la forme d'une fonction mathématique connue ou d'un graphe. Un outil fondamental de la théorie sera la décomposition d'un signal quelconque en signaux élémentaires, en particulier les signaux harmoniques. Par exemple, un diapason parfaitement réglé à 440 Hz produit un signal audio que l'on peut idéaliser par la fonction harmonique

$$\text{diap} : \mathbb{R} \rightarrow (P_{\min}, P_{\max}) : t \mapsto \text{diap}(t) = A \sin(2\pi 440 t).$$

Le signal $\text{diap}(\cdot)$ est caractérisé par son amplitude A et sa fréquence $f = 440$, exprimée en $\text{Hz} = 1/\text{s}$, ou sa pulsation $\omega = 2\pi 440$, exprimée en rad/s . Un résultat essentiel de la théorie des signaux consistera à décomposer un signal quelconque comme combinaison de signaux harmoniques. Les signaux harmoniques joueront donc un rôle très important dans ce cours.

1.2.2 Échantillonnage et quantification

La Figure 1.8a donne une représentation graphique du signal $\text{diap}(\cdot)$ générée par les commandes Matlab suivantes :

```
t = (0 :100) / 10^4 ;
A = 1 ;
diap = A * sin(2 * pi * 440 * t) ;
```

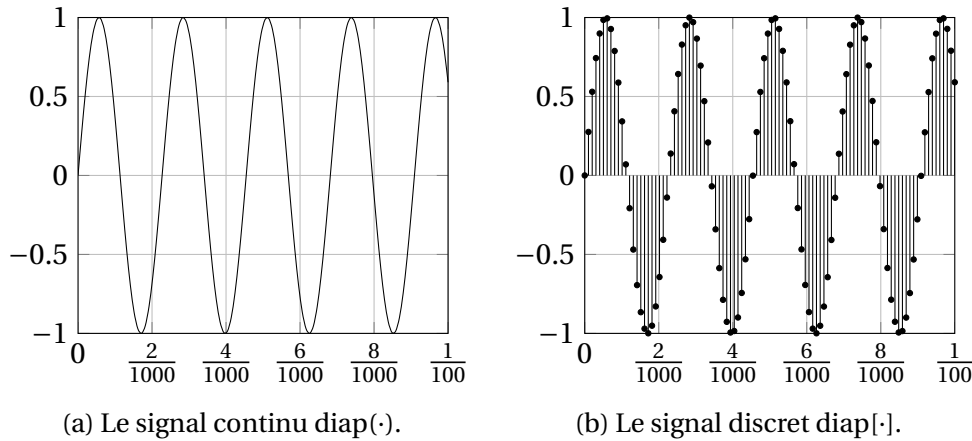


FIGURE 1.8 – Représentation graphique du signal diap.

Le signal `diap` traité par Matlab est un signal *discrétisé* ou échantillonné. Son domaine n'est plus un continuum mais un ensemble de valeurs discrètes. C'est un exemple de signal *en temps-discret*, représenté par la fonction

$$\text{diap} : \mathbb{Z} \rightarrow (P_{\min}, P_{\max}) : n \mapsto \text{diap}[n] = A \sin(2\pi 440nT), \quad T = 10^{-4}.$$

Les signaux `diap(·)` et `diap[·]` ne diffèrent que par leur domaine mais cette distinction est capitale dans la théorie des signaux et systèmes. C'est pourquoi on utilisera la notation $(\cdot)$ lorsque l'argument du signal est continu et $[\cdot]$ lorsque l'argument est discret. Graphiquement, on considérera le signal de la Figure 1.8a comme un signal continu. La représentation graphique d'un signal discret sera explicitée selon l'illustration de la Figure 1.8b. (Ces types de représentation peuvent respectivement être générés en Matlab par les commandes «`plot(t, diap)`» et «`stem(t, diap)`».)

L'expression `diap[0.5]` n'a pas de sens parce que la valeur 0.5 ne fait pas partie du domaine du signal discret. Dans la théorie des systèmes, tous les signaux dépendant d'une seule variable *entière* sont appelés signaux en *temps-discret*, même si la variable indépendante ne désigne pas le temps. Le signal *RR* discuté dans la section précédente est un exemple de signal discret dont la variable indépendante n'est pas le temps mais un numéro de battement.

Le signal `diap` n'est pas seulement échantillonné. Il est aussi *quantifié* parce que chaque valeur du vecteur `diap` est codée au moyen d'un nombre fini de bits. Tout comme son domaine, son image est donc un ensemble discret plutôt qu'un continuum. Tous les signaux codés sur un support digital sont de facto quantifiés. Les effets de quantification sont souvent négligés dans une première analyse, mais ils peuvent être déterminants. Par exemple, un filtre stable peut devenir instable lorsque ses coefficients sont implémentés en précision finie.

Les effets d'échantillonnage et de quantification de signaux continus sont omniprésents dans le traitement actuel (numérique) des signaux. Ils introduisent dans l'analyse une approximation qui dépend de la fréquence d'échantillonnage et du niveau de quantification. Si le régulateur du pendu-bot est modélisé comme un système produisant un signal en temps-continu u à partir du signal en temps-continu (θ_1, θ_2) , il faudra s'assurer que les effets d'échantillonnage et de quantification des différents éléments sont négligeables. Dans le cas contraire – par exemple si le capteur angulaire a une résolution insuffisante –, il faudra prendre ces effets en compte dans la modélisation du régulateur.

1.2.3 Domaine de signaux

Les exemples de signaux qui précèdent et plus généralement les signaux étudiés dans ce cours sont *unidimensionnels*, c'est-à-dire qu'ils ne dépendent que d'une seule variable indépendante, appelée par convention le temps. Si la variable indépendante prend ses valeurs dans $\mathbb{R}$, le signal est un signal en temps-continu. Si la variable indépendante prend ses valeurs dans $\mathbb{Z}$, le signal est un signal en temps-discret. Il ne s'agit pas d'une restriction fondamentale sur le plan des outils développés, mais ce cours ne traitera pas *directement* de signaux « images » ou « vidéo ». Le domaine naturel d'une image est un espace à deux dimensions, par exemple les coordonnées cartésiennes d'un point donné de l'image. Un signal vidéo est une image qui change au cours du temps. Son domaine naturel sera tridimensionnel (deux dimensions spatiales et une dimension temporelle).

Un *signal* sera donc décrit par une fonction $x : X \rightarrow Y$ avec un domaine réel ou entier :

- signal temps-continu $x = x(\cdot)$ avec $X \subset \mathbb{R}$,
- signal temps-discret $x = x[\cdot]$ avec $X \subset \mathbb{Z}$.

Il importe de ne pas confondre le signal x (une fonction de X dans Y) et la valeur du signal $x(t)$ à un instant donné t (un élément de Y). Lorsque l'on souhaite souligner cette distinction, on adoptera la notation $x(\cdot)$ pour un signal x en temps-continu et $x[\cdot]$ pour un signal x en temps-discret. Néanmoins, quand le domaine est clair du contexte on omettra la quantification universelle sur le domaine; e.g., abrégeant « $\forall t \in X : x(t) = e^{j\omega t}$ » à « $x(t) = e^{j\omega t}$ ».

On utilisera les deux différentes représentation graphique de la Figure 1.9 pour différencier les signaux continus et discrets.

Lorsqu'un signal est défini sur un intervalle plutôt que sur la droite réelle $\mathbb{R}$ ou l'ensemble des entiers $\mathbb{Z}$, il conviendra néanmoins pour le traitement mathématique d'étendre la définition du signal à l'ensemble $\mathbb{R}$ ou $\mathbb{Z}$. Un signal défini sur un intervalle pourra être prolongé par zéro ou répété périodiquement en dehors de son intervalle de définition.

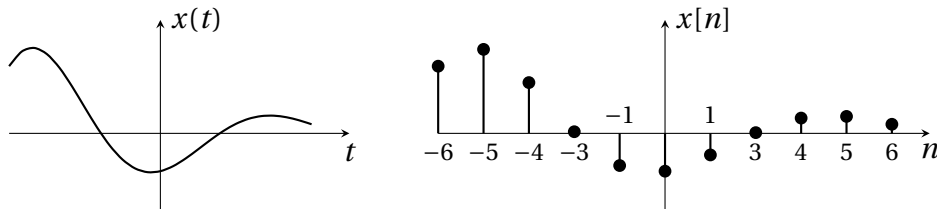
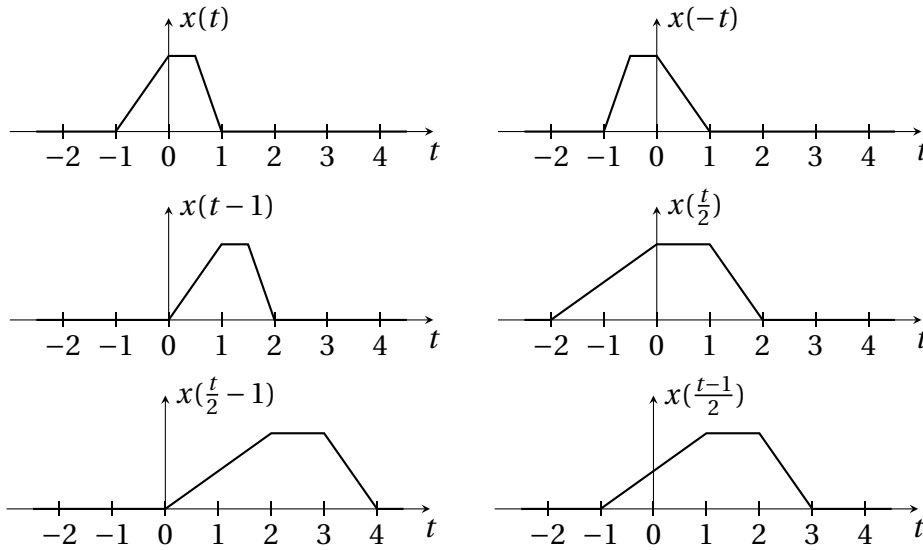
FIGURE 1.9 – Représentation graphique d'un signal continu $x(\cdot)$ et discret $x[\cdot]$.

FIGURE 1.10 – Opérations de base sur la variable indépendante d'un signal en temps-continu.

La restriction à des signaux dont le domaine est $\mathbb{R}$ ou $\mathbb{Z}$ permet de définir des opérations élémentaires utiles pour le traitement de signal :

- le *décalage temporel* ou *time-shift* qui décale le signal d'une quantité fixe sur l'axe temporel : $x(t) \mapsto \Delta_{-t_0} x(t) = x(t - t_0)$ en continu ou $x[n] \mapsto \sigma^{n_0} x[n] = x[n - n_0]$ en discret.
- la *réflexion* ou *time-reversal* qui produit un signal réfléchi par rapport à $t = 0$ ou $n = 0$: $x(t) \mapsto x(-t)$ ou $x[n] \mapsto x[-n]$.
- le *changement d'échelle de temps* ou *time-scaling* qui contracte ($\alpha > 1$) ou dilate ($\alpha < 1$) le signal selon l'axe temporel : $x(t) \mapsto x(\alpha t)$.

Ces opérations de base interviennent continuellement dans le traitement des signaux. Il convient de bien se les représenter graphiquement (voir Figure 1.10). Si le domaine des signaux est un espace vectoriel, les opérations d'addition (décalage temporel) et de multiplication par un scalaire (réflexion et changement d'échelle de temps) ne modifient pas le domaine du signal. C'est le cas des signaux définis sur $\mathbb{R}$ (pour le champ de scalaires $\mathbb{R}$). En revanche $\mathbb{Z}$ n'est pas un corps et toutes les valeurs de α ne sont par conséquent pas permises.

Un signal est *pair* s'il est invariant pour l'opération de réflexion : $x(t) = x(-t)$ ou $x[n] = x[-n]$. Il est *impair* si la réflexion produit un changement de signe : $x(t) = -x(-t)$ ou $x[n] = -x[-n]$. Tout signal peut être décomposé en la somme d'un signal pair et impair :

$$x(t) = \frac{1}{2}(x(t) + x(-t)) + \frac{1}{2}(x(t) - x(-t))$$

Un signal continu $x(\cdot)$ est *périodique* s'il existe un nombre $T > 0$ tel que $x(t + T) = x(t)$: le signal est alors invariant lorsqu'on lui applique un décalage temporel T . De même en discret, un signal $x[\cdot]$ est périodique s'il existe un entier $N > 0$ tel que $x[n + N] = x[n]$. La plus petite période T ou N qui laisse le signal invariant est appelée *période fondamentale* du signal.

1.2.4 Image de signaux

L'espace image Y d'un signal $x : X \rightarrow Y$ est l'ensemble des valeurs que peut prendre le signal.

La structure mathématique que possède l'ensemble Y conditionne les opérations de traitement de signal que l'on peut effectuer. Y peut être un simple ensemble de symboles sans structure particulière. Par exemple, Une molécule d'ADN peut être codée comme une longue séquence de symboles prenant leurs valeurs dans un alphabet réduit à quatre lettres. Cette séquence peut être traitée comme un signal en temps-discret prenant ses valeurs dans l'ensemble $Y = \{A, G, T, C\}$. Un signal binaire est un signal prenant ses valeurs dans l'ensemble $Y = \{0, 1\}$. De tels signaux sont bien surs fréquemment rencontrés dans les applications mais le manque de structure mathématique de leur espace image limite le traitement mathématique qu'on peut leur apporter.

Les signaux étudiés dans ce cours prennent leurs valeurs dans un espace vectoriel (réel ou complexe), le plus souvent $\mathbb{R}^N$ (sur le champ de scalaires $\mathbb{R}$) ou $\mathbb{C}^N$ (sur le champ de scalaires $\mathbb{C}$). Si $N > 1$, la valeur du signal à un instant donné est un vecteur à N composantes. On dénote alors la i -ème composante du vecteur par $x_i(t)$ ou $x_i[n]$, i.e.

$$x(t) = \begin{bmatrix} x_1(t) \\ x_2(t) \\ \dots \\ x_N(t) \end{bmatrix} \quad \text{ou} \quad x[n] = \begin{bmatrix} x_1[n] \\ x_2[n] \\ \dots \\ x_N[n] \end{bmatrix}.$$

1.2.5 Espaces vectoriels de signaux

Le signal x peut lui-même être considéré comme élément d'un espace de signaux $\mathcal{X}$. La structure de cet espace détermine les opérations que l'on peut effectuer sur le signal. Les espaces de signaux considérés dans ce cours sont des espaces vectoriels : une combinaison linéaire de deux signaux d'un même espace donne encore un signal du même espace :

$$x \text{ et } y \in \mathcal{X} \Rightarrow \alpha x + \beta y \in \mathcal{X}.$$

La structure d'espace vectoriel de l'espace de signaux est héritée de la structure d'espace vectoriel de l'espace image des signaux : la combinaison linéaire de signaux est définie à chaque instant t par la combinaison linéaire $\alpha x(t) + \beta y(t)$ dans l'espace vectoriel Y . Cette propriété est fondamentale pour la théorie des systèmes linéaires.

Un espace de signaux peut être muni d'une norme. La norme d'un signal physique est une mesure de son énergie. La notion de norme permet de parler de distance entre deux signaux (la distance étant définie comme la norme de la différence des deux signaux). Elle jouera un rôle important dans l'analyse de convergence des séries de Fourier (Chapitre 5).

Pour la définition des normes, on utilise la notation suivante : $|c|$ désigne la valeur absolue de c si $c \in \mathbb{R}$, le module $\sqrt{\Re(c)^2 + \Im(c)^2}$ si $c \in \mathbb{C}$, la norme euclidienne $\sqrt{\sum_{i=1}^n c_i^2}$ si c est un vecteur dans $\mathbb{R}^n$, et la norme complexe $\sqrt{\sum_{i=1}^n |c_i|^2}$ si c est un vecteur dans $\mathbb{C}^n$.

Pour un signal en temps-continu x défini sur l'intervalle avec bornes $t_1 \leq t_2$, on considèrera trois normes différentes :

$$\|x\|_1 = \int_{t_1}^{t_2} |x(t)| \, dt$$

ou

$$\|x\|_2 = \sqrt{\int_{t_1}^{t_2} |x(t)|^2 \, dt}$$

ou

$$\|x\|_\infty = \sup_{t \in (t_1, t_2)} |x(t)|.$$

De même, pour un signal discret x défini sur l'intervalle $[k_1..k_2]$ avec bornes entières $k_1 \leq k_2$, on considèrera les trois normes suivantes :

$$\|x\|_1 = \sum_{k=k_1}^{k_2} |x[k]|$$

ou

$$\|x\|_2 = \sqrt{\sum_{k=k_1}^{k_2} |x[k]|^2}$$

ou

$$\|x\|_\infty = \sup_{k \in [k_1..k_2]} |x[k]|.$$

Si le signal est défini sur un intervalle infini, i.e., $\mathbb{R}$ ou $\mathbb{Z}$, on fait tendre les limites de l'intégrale ou de la somme vers l'infini. En continu, on obtient par exemple

$$\|x\|_2 = \sqrt{\int_{-\infty}^{+\infty} |x(t)|^2 dt} \quad (1.2.1)$$

et, en discret,

$$\|x\|_2 = \sqrt{\sum_{k \in \mathbb{Z}} |x[k]|^2}. \quad (1.2.2)$$

L'ensemble des signaux de norme finie (pour une norme donnée) est un espace de signaux normé. Pour les trois normes définies sur des signaux en temps-continu, on les note respectivement $L_1(t_1, t_2)$, $L_2(t_1, t_2)$, $L_\infty(t_1, t_2)$. Pour les trois normes définies sur des signaux en temps-discret, on les note respectivement $\ell_1[k_1..k_2]$, $\ell_2[k_1..k_2]$, $\ell_\infty[k_1..k_2]$.

La norme $\|\cdot\|_2$ a un statut spécial parce qu'elle dérive d'un produit scalaire : en définissant dans $L_2(t_1, t_2)$

$$\langle x, y \rangle = \int_{t_1}^{t_2} x(t) \overline{y(t)} dt$$

ou dans $\ell_2[k_1..k_2]$

$$\langle x, y \rangle = \sum_{n=k_1}^{k_2} x[n] \overline{y[n]}$$

on peut vérifier aisément que $\langle \cdot, \cdot \rangle$ définit un produit scalaire et que $\|x\|_2^2 = \langle x, x \rangle$. La notion de produit scalaire permet de parler d'angle entre deux signaux. En particulier, deux signaux sont dits orthogonaux lorsque leur produit scalaire est nul.

1.3 Systèmes

1.3.1 Systèmes comme fonctions de fonctions

Un système établit une relation entre un signal d'entrée et un signal de sortie. Si le signal d'entrée appartient à un ensemble $\mathcal{X}$ (domaine) et que le signal de sortie appartient à un ensemble $\mathcal{Y}$ (image), le système peut être décrit comme une fonction $S : \mathcal{X} \rightarrow \mathcal{Y}$. La différence par rapport aux signaux décrits précédemment est que les éléments de $\mathcal{X}$ et $\mathcal{Y}$ sont à présent des

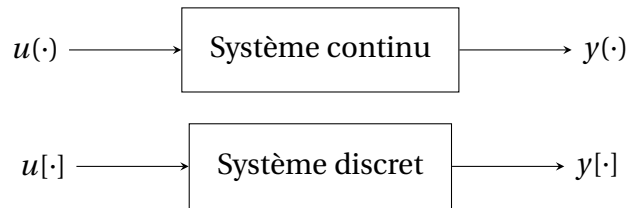


FIGURE 1.11 – Système entrée-sortie (continu et discret).

signaux, donc eux-mêmes des fonctions. Dans ce sens, un système est une fonction de fonctions, ou un *opérateur* d'un espace de signaux vers un autre espace de signaux.

S'il est important pour un système de spécifier son domaine et son image, comme dans le cas des signaux, décrire la manière dont le domaine est appliqué sur l'image est en général beaucoup plus compliqué. En effet, le domaine et l'image des signaux traités dans ce cours sont le plus souvent unidimensionnels, de sorte que le signal est aisément représenté par une fonction mathématique connue ou par un graphe. En revanche, le domaine et l'image d'un système seront le plus souvent de dimension infinie, rendant ce type de représentation impossible. La description mathématique adéquate d'un système est en fait un des objectifs importants de ce cours. Elle n'est pas univoque, et est souvent motivée par la nature du système considéré. Il importera donc également d'établir un lien entre les différentes manières de décrire un système.

Le plus souvent, la théorie des systèmes consiste à analyser comment un *signal d'entrée* u est transformé par le système considéré en un *signal de sortie* y , d'où la représentation schématique *entrée-sortie* de la Figure 1.11. Un système désigne dans ce cas une *loi* entre signaux d'entrée et signaux de sortie. Cette loi est déterminée explicitement par toutes les paires (u, y) d'entrées-sorties qui sont *solutions* du système, c'est-à-dire qui obéissent à la loi.

1.3.2 Mémoire d'un système

Un système est *statique* (ou sans mémoire) lorsque la valeur du signal de sortie à un instant donné ne dépend que de la valeur du signal d'entrée au même instant. Dans le cas contraire, le système est *dynamique* (ou à mémoire). Ainsi, la relation $y[n] = (u[n])^2$ définit un système statique, tandis que la relation $y[n] = 2u[n - 1]$ définit un système dynamique. Physiquement, un système dynamique est souvent associé à des composants pouvant emmagasiner de l'énergie (une capacité, un ressort).

Exemple 1.1. Les filtres numériques sont omniprésents dans les applications décrites précédemment. Ils constituent le prototype de systèmes spéci-

fiés par une équation aux différences. L'exemple le plus simple est le filtre à moyenne mobile spécifié par l'équation

$$y[n] = \frac{1}{2}u[n-1] + \frac{1}{2}u[n]. \quad (1.3.1)$$

La valeur du signal de sortie à l'instant n est la moyenne arithmétique entre la valeur du signal d'entrée au même instant et la valeur du signal d'entrée à l'instant précédent. Le filtre opère ainsi un « lissage » du signal d'entrée. Nous verrons dans ce cours d'autres filtres capables d'effectuer un lissage plus efficace de données bruitées. Par exemple, nous étudierons le filtre de lissage exponentiel spécifié par l'équation aux différences

$$y[n] = ay[n-1] + (1-a)u[n], \quad 0 < a < 1. \quad (1.3.2)$$

La Figure 1.12 illustre l'effet du filtre à moyenne mobile et l'effet du filtre exponentiel pour deux valeurs du paramètre a . Notons que l'implémentation (par exemple dans Matlab) d'une équation aux différences est une tâche extrêmement simple. Pourtant, l'équation (1.3.2) est une spécification implicite du système : il faut résoudre l'équation aux différences pour connaître la relation explicite entre le signal d'entrée et le signal de sortie.

Pour déterminer l'évolution de y à partir de l'instant initial n_0 , il faut spécifier une valeur initiale pour $y[n_0]$. Toutes les valeurs futures sont alors déterminées par substitution successive. On obtient facilement la formule explicite pour $n \geq 0$:

$$y[n+n_0] = a^n y[n_0] + (1-a) \sum_{k=0}^{n-1} a^k u[n+n_0-k].$$

La valeur du signal de sortie à un instant donné $n \geq n_0$ est donc une somme pondérée de toutes les valeurs passées du signal d'entrée depuis l'instant initial. Le signal de pondération est le signal en temps-discret $x[n] = a^n$ solution de l'équation homogène $x[n+1] = ax[n]$. Si $|a| < 1$, le poids alloué aux entrées passées diminue quand on remonte dans le temps. La constante a est parfois appelée *facteur d'oubli* car elle détermine à quelle rapidité on décide d'oublier les entrées passées. La solution générale d'une équation aux différences linéaire à coefficients constants sera rappelée au Chapitre 4. Elle présentera des caractéristiques analogues de mémoire. <

Exemple 1.2. Considérons le circuit RC de la Figure 1.13 où la tension source v_s est le signal d'entrée et où la tension v_C est le signal de sortie.

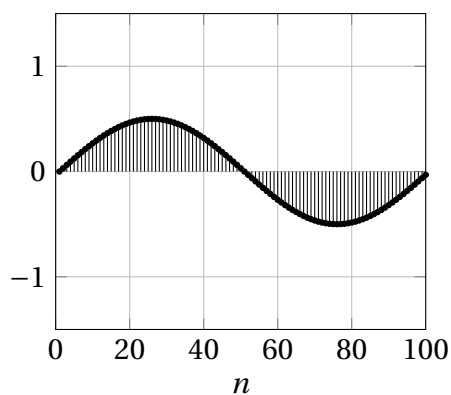
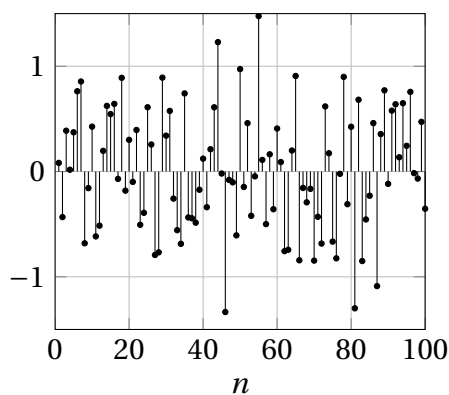
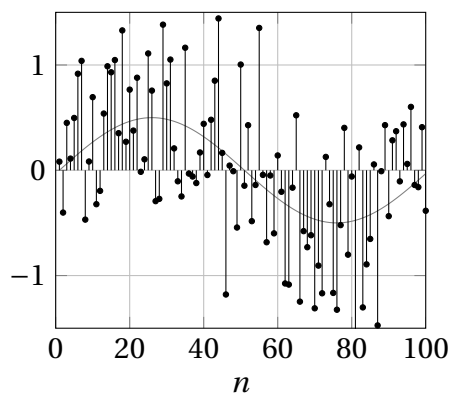
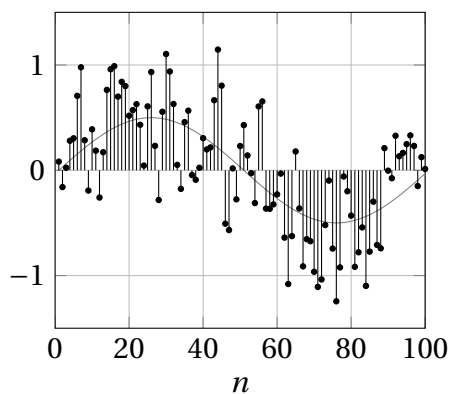
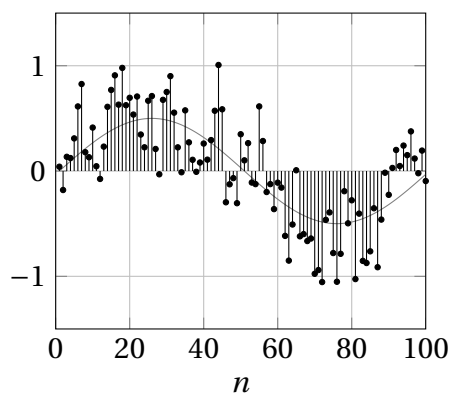
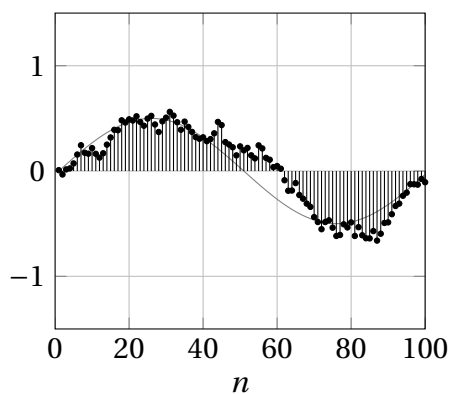
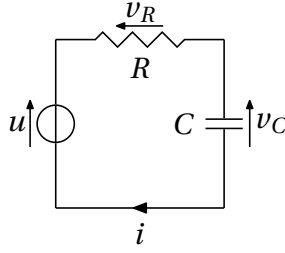
(a) $\frac{1}{2} \sin[n \frac{\pi}{50}]$ (b) $e[n]$ (c) $u[n] = \frac{1}{2} \sin[n \frac{\pi}{50}] + e[n]$ (d) $y[n]$ (moyenne mobile)(e) $y[n]$ (exponentiel, $a = \frac{1}{2}$)(f) $y[n]$ (exponentiel, $a = \frac{9}{10}$)

FIGURE 1.12 – Lissage d'un signal bruité ((a), (b), (c)). Illustré pour le filtre (1.3.1) en (d) et pour le filtre (1.3.2) pour deux valeurs du paramètre d'oubli a en (e) et (f).

FIGURE 1.13 – Un circuit RC avec une tension source $u = v_s$.

L'équation de ce circuit est obtenue en utilisant la loi d'Ohm

$$Ri = v_s - v_C$$

et l'équation constitutive d'une capacité

$$i = C \dot{v}_C.$$

En combinant ces deux équations, on obtient l'équation différentielle

$$\dot{v}_C + \frac{1}{RC} v_C = \frac{1}{RC} v_s \quad (1.3.3)$$

qui décrit le système ou la loi reliant le signal d'entrée v_s au signal de sortie v_C . Le caractère dynamique du système est dû à la capacité qui peut emmagasiner de l'énergie. Lorsqu'une tension source est appliquée au circuit à un instant donné, la capacité se charge à travers la résistance R , ce qui a pour conséquence que la tension v_C ne suit pas instantanément la tension source. Lorsque la valeur de la capacité tend vers 0, l'effet de mémoire disparaît. À la limite, le circuit tend vers un circuit ouvert, qui correspond au système statique $v_C = v_s$.

La théorie des équations différentielles permet de résoudre l'équation (1.3.3) à partir d'un temps initial t_0 pour $t \geq 0$:

$$v_C(t_0 + t) = e^{-\frac{t}{RC}} v_C(t_0) + \frac{1}{RC} \int_{t_0}^{t_0+t} e^{-\frac{(t_0+t-\tau)}{RC}} v_s(\tau) d\tau. \quad (1.3.4)$$

La solution fait apparaître l'effet de mémoire du circuit RC : le signal de sortie $v_C(t_0 + t)$ dépend de toute l'histoire du signal d'entrée v_s depuis l'instant initial t_0 jusqu'à l'instant $t_0 + t$. L'histoire passée du signal d'entrée est pondérée par le signal exponentiel $x(t) = e^{-at}$, $a = \frac{1}{RC}$, solution de l'équation homogène $\dot{x} + ax = 0$. La solution générale d'une équation différentielle linéaire à coefficients constants sera rappelée au Chapitre 4. Elle présentera des caractéristiques analogues de mémoire. $\triangleleft$

1.3.3 Équations aux différences et équations différentielles

L'exemple du filtre exponentiel (1.3.2) est un exemple d'équation *aux différences*. Beaucoup de systèmes dynamiques sont spécifiés de cette manière. Une équation aux différences entrée-sortie générale est de la forme

$$F(y, \sigma y, \dots, \sigma^{N-1} y, u, \sigma u, \dots, \sigma^{M-1} u) = 0. \quad (1.3.5)$$

Elle spécifie la loi du système sous la forme d'une relation implicite entre N valeurs successives du signal de sortie et M valeurs successives du signal d'entrée.

Ce cours se concentrera sur l'étude de système définis par des équations aux différences *linéaires à coefficients constants* : on appellera *système aux différences d'ordre N* une équation aux différences de la forme

$$\sum_{k=0}^N a_k y[n-k] = \sum_{k=0}^N b_k u[n-k], \quad (1.3.6)$$

où $a_N \neq 0$ ou $b_N \neq 0$ et $a_0 \neq 0$ ou $b_0 \neq 0$. Cette équation établit une loi entre N valeurs consécutives des signaux de sortie et d'entrée. Le filtre exponentiel est donc un système aux différences d'ordre un.

Le circuit RC décrit par l'équation (1.3.3) est un exemple d'*équation différentielle entrée-sortie*. L'équation du système exprime une relation entre le signal d'entrée, le signal de sortie, et un certain nombre de leurs dérivées. C'est le mode de description par excellence de tous les processus décrits par les lois de la physique et reliant entre eux des signaux en temps-continu, en particulier l'évolution temporelle de variables physiques. De manière plus générale, une équation différentielle entrée-sortie entre deux signaux en temps-continu sera de la forme

$$F(y, \dot{y}, \dots, y^{(N-1)}, u, \dot{u}, \dots, u^{(M-1)}) = 0 \quad (1.3.7)$$

La notation $y^{(i)}$ désigne la i -ème dérivée du signal y ; pour les dérivées première et seconde, on utilise aussi la notation $y' = \dot{y}$ et $y'' = \ddot{y}$.

On notera que la spécification d'un système par une relation différentielle telle l'équation (1.3.7) est implicite. Elle permet de vérifier si un couple entrée-sortie (u, y) est compatible avec la loi du système (c'est-à-dire vérifie l'équation) mais elle ne permet pas de calculer la sortie y à partir du signal d'entrée u . Pour cela, il faut *résoudre* l'équation différentielle. Un des objectifs de la théorie des systèmes est de décrire les propriétés d'un système *sans le résoudre*, c'est-à-dire sans calculer explicitement toutes les paires (u, y) qui satisfont l'équation du système.

Les équations différentielles étudiées dans ce cours sont des équations

linéaires à coefficients constants. On appellera *système différentiel d'ordre N* une équation différentielle de la forme

$$\sum_{k=0}^N a_k \frac{d^k y(t)}{dt^k} = \sum_{k=0}^N b_k \frac{d^k u(t)}{dt^k}, \quad (1.3.8)$$

où $a_N \neq 0$ ou $b_N \neq 0$ et $a_0 \neq 0$ ou $b_0 \neq 0$ et u et y sont des signaux scalaires. Le système est dans ce cas spécifié par une loi qui fait intervenir l'entrée u et la sortie y ainsi qu'un certain nombre de leurs dérivées. Le circuit RC est donc un système différentiel d'ordre un.

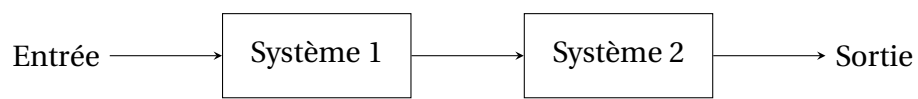
1.3.4 Causalité

Un système est *causal* lorsque le signal de sortie à l'instant t ne dépend que des valeurs passées du signal d'entrée $u(\tau)$, $\tau \leq t$. Un système statique est donc causal mais le système $y[n] = u[n+1]$ ne l'est pas car la valeur du signal de sortie à l'instant n dépend de la valeur du signal d'entrée à un instant ultérieur. Un système physique pour lequel la variable indépendante représente le temps est nécessairement causal. Néanmoins, les systèmes non causaux peuvent être utiles dans la théorie des systèmes et des signaux. En distinguant le futur du passé, la propriété de causalité établit une distinction entre les signaux mathématiques et les signaux physiques.

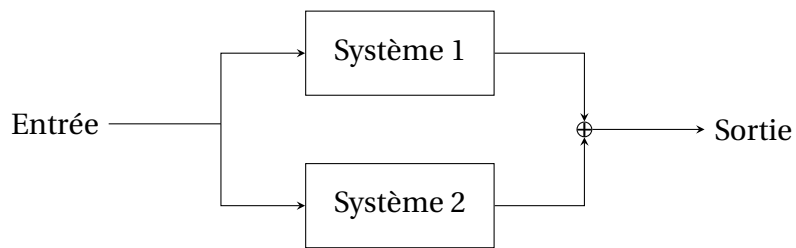
1.3.5 Interconnexions de systèmes

La notion d'interconnexion est très importante dans la théorie des systèmes. La plupart des systèmes analysés en pratique sont des interconnexions de sous-systèmes, c'est-à-dire qu'ils peuvent être décomposés en blocs de base reliés entre eux en identifiant leurs entrées ou sorties respectives. Dans certains ouvrages, la notion d'interconnexion a valeur de définition pour un système : un système *est* une interconnexion d'éléments. Par exemple, un circuit électrique est une interconnexion de résistances, de capacités, et d'inductances. La vision du système comme interconnexion de sous-systèmes souligne le fait qu'on ne s'intéresse pas tant aux éléments individuels (la description fine d'une résistance électrique ne fait pas partie de la théorie des systèmes) qu'à l'interaction qui résulte de leur interconnexion.

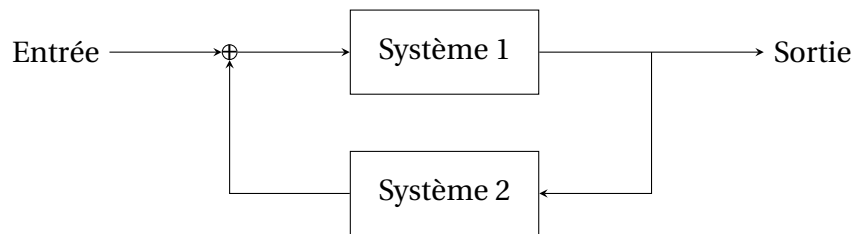
Tous les systèmes étudiés dans ce cours sont constitués de blocs de base (une résistance, un intégrateur, ...) interconnectés de trois manières différentes : en série, en parallèle, ou en rétroaction (feedback). Pour décrire les interconnexions qui définissent un système, on utilise fréquemment des *bloc-diagrammes*. Un sous-système est représenté par un bloc (boîte noire) et deux flèches pour indiquer son entrée et sa sortie. Les différentes interconnexions sont représentées à la Figure 1.14. Le + dénote l'addition de deux signaux.



(a) Série



(b) Parallèle



(c) Rétroaction

FIGURE 1.14 – Interconnexions.

Chapitre 2

Le modèle d'état

Le modèle d'état introduit dans ce chapitre fournit un moyen très efficace de spécifier la loi d'un système entre des signaux d'entrée et des signaux de sortie. Il est basé sur l'utilisation d'un signal auxiliaire appelé signal d'*état* ou vecteur d'*état* lorsqu'il a plusieurs composantes. La loi du système est spécifiée par une loi de mise à jour du vecteur d'état. La notion d'état constitue un concept important de la théorie des systèmes. Elle sera étudiée en détail dans le Chapitre 4 dans le cadre des systèmes linéaires et invariants.

Dans le chapitre présent, nous définissons la structure d'un modèle d'état et nous montrons au moyen de quelques exemples que le modèle d'état se prête bien tant à la modélisation des systèmes en temps-discret rencontrés en informatique que celle des systèmes en temps-continu rencontrés en physique.

2.1 Structure générale d'un modèle d'état

Un modèle d'état Σ (discret) est défini par trois ensembles et deux fonctions.

1. U espace d'entrée
2. Y espace de sortie
3. X espace d'état
4. $f : X \times U \rightarrow X$ fonction de mise à jour
5. $h : X \times U \rightarrow Y$ fonction de sortie

Pour chaque condition initiale $x_0 \in X$, le modèle d'état Σ définit un système (Σ, x_0) , c'est-à-dire une relation univoque entre des signaux d'entrée $u : \mathbb{N} \rightarrow U$ et des signaux de sortie $y : \mathbb{N} \rightarrow Y$.

La condition initiale x_0 et le signal d'entrée u déterminent l'évolution de l'état x pour $n \in \mathbb{N}$ grâce à l'équation de mise à jour :

$$x[n+1] = f(x[n], u[n]). \quad (2.1.1)$$

Le signal de sortie est déterminé pour $n \in \mathbb{N}$ par la fonction de sortie :

$$y[n] = h(x[n], u[n]). \quad (2.1.2)$$

Le signal x est appelé *solution* du modèle d'état pour la condition initiale x_0 et le signal d'entrée u . Il définit une *trajectoire* dans l'espace d'état, c'est-à-dire une succession d'états dans le temps. La paire (u, y) est appelée *solution* du système. Elle définit un *comportement* permis par le modèle d'état Σ . Le signal x peut avoir plusieurs composantes, auquel cas l'équation (2.1.1) est une équation vectorielle. Le nombre de composantes du vecteur x est appelé *dimension* du modèle d'état.

Exemple 2.1. Le filtre de *lissage exponentiel* avec paramètre $0 < a < 1$ introduit au chapitre précédent et défini par la « loi »

$$y[n] = ay[n-1] + (1-a)u[n] \quad (2.1.3)$$

peut être décrit par le modèle d'état suivant :

$$X = Y = U = \mathbb{R}, \quad f(\xi, v) = a\xi + (1-a)v, \quad h(\xi, v) = a\xi + (1-a)v.$$

On vérifie effectivement que la loi de mise à jour

$$x[n+1] = ax[n] + (1-a)u[n]$$

et l'équation de sortie

$$y[n] = ax[n] + (1-a)u[n]$$

impliquent la relation $x[n] = y[n-1]$ et par conséquent l'équation aux différences

$$y[n] = ay[n-1] + (1-a)u[n].$$

On étudiera dans le Chapitre 4 comment associer un modèle d'état à des équations aux différences plus générales. On peut dès à présent observer le rôle de mémoire joué par l'état du système dans le filtre exponentiel : l'état retient la dernière valeur de la sortie et cette connaissance du passé suffit à déterminer la valeur de la sortie à l'instant présent. $\triangleleft$

2.2 Automates finis

2.2.1 Définition et représentation

Lorsque l'espace d'état est un ensemble fini (c'est-à-dire lorsque chaque composante du vecteur d'état peut prendre un nombre fini de valeurs distinctes), le modèle d'état est appelé *automate fini* ou *machine d'état finie*.

L'étude des automates finis exploite le caractère fini de l'espace d'état, qui permet (en principe) une énumération de toutes les trajectoires possibles (on parle aussi d'*exploration* de l'espace d'état). Ce formalisme est particulièrement utile lorsque la cardinalité de l'ensemble X n'est pas trop élevée. Il permet de considérer des ensembles (U, Y, X) sans aucune structure particulière (on parle souvent d'alphabet pour désigner un ensemble de symboles sans relation particulière entre les éléments).

Les automates finis sont souvent représentés par un diagramme de transition, comme illustré à la Figure 2.1. Le diagramme de transition constitue une représentation graphique intuitive d'un automate fini.

2.2.2 Étude des automates finis

L'étude des automates finis n'est pas poursuivie dans le cadre de ce cours. Elle est abordée dans le cadre de cours d'informatique. Un automate fini correspond souvent à la modélisation d'un programme informatique.

On peut par exemple s'intéresser aux relations d'équivalence entre deux automates qui ont le même espace d'entrée et de sortie. Un automate est capable de *simuler* un autre automate s'il peut reproduire toutes les paires entrée-sortie de l'automate simulé. Deux automates sont en relation de *bi-simulation* s'ils produisent les mêmes paires entrée-sortie, c'est-à-dire s'ils engendrent les mêmes systèmes. La bi-simulation est une relation d'équivalence qui n'implique pas que les modèles d'états (lois de mise à jour et fonction de sortie) sont identiques.

2.3 Modèles d'état en temps-discret

Lorsqu'un modèle d'état admet un nombre d'état infinis (par exemple lorsqu'une variable d'état prend ses valeurs dans un continuum), on parle de machine d'état *infinie*. L'étude de ces modèles requiert des outils d'analyse distincts des outils d'analyse des automates finis. On ne peut par exemple plus directement recourir à des représentations par diagramme de transition. Cette généralisation s'accompagne généralement d'une restriction sur les ensembles de signaux considérés. En particulier, on rencontre couramment des modèles d'état infinis dont les espaces de signaux U , Y , et X , sont des espaces vectoriels. Par exemple, on étudie des modèles d'état dans lesquels toutes les variables sont des variables réelles (plutôt que les symboles d'un alphabet). Cette structure supplémentaire permet d'étudier des propriétés qualitatives comme la stabilité d'une solution ou la convergence vers une solution particulière malgré le fait que l'on travaille dans un espace d'état infini. Dans ces modèles, la mise à jour est le plus souvent une mise à jour temporelle, et on parle de *modèle d'état en temps-discret*.

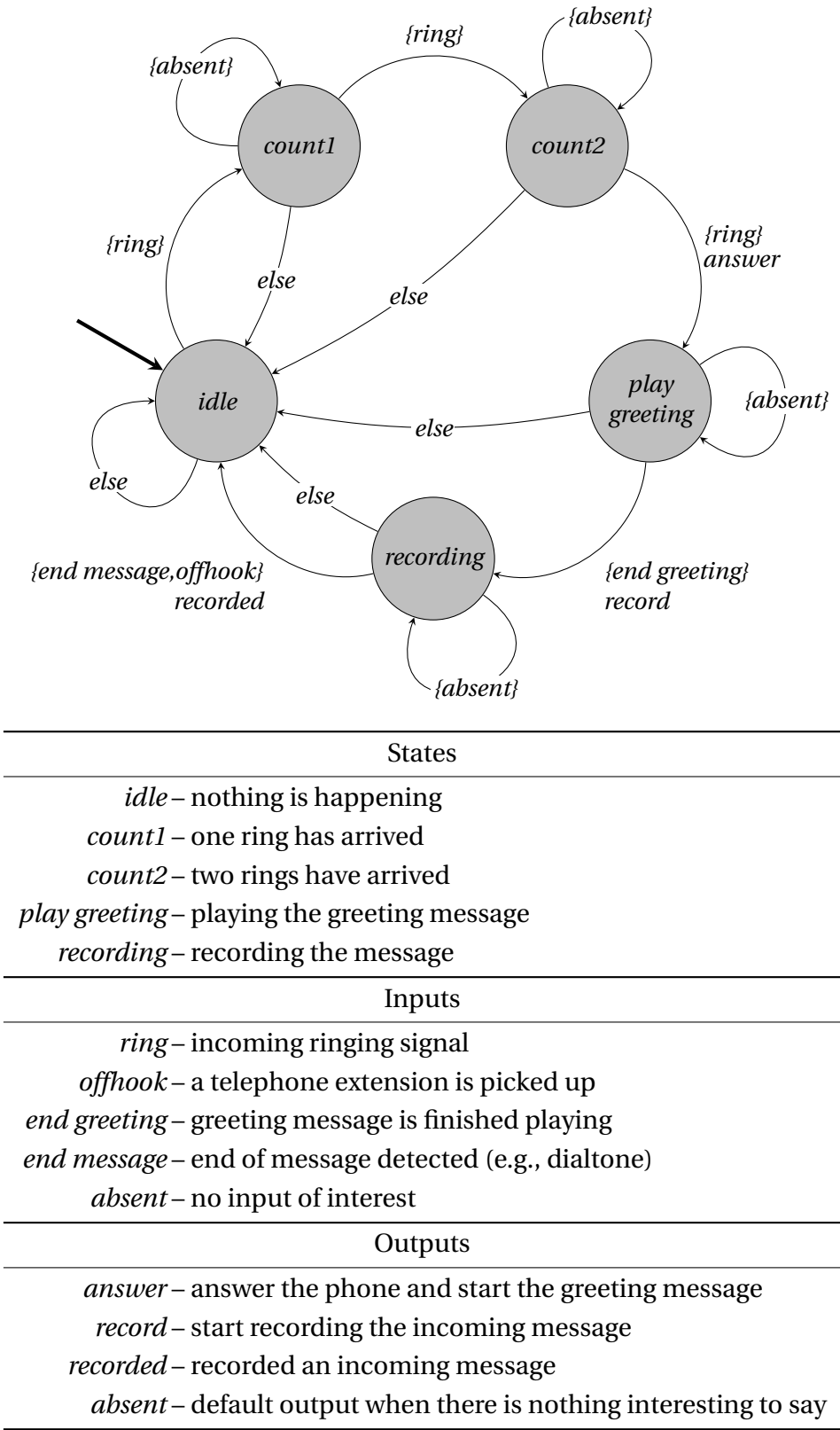


FIGURE 2.1 – Diagramme de transition représentant un automate fini associé à l'opération d'un répondeur (D'après Lee & Varaiya, p. 91).

L'étude des modèles d'état en temps-discret est largement motivée par la discipline voisine des *systèmes dynamiques* qui étudie les propriétés qualitatives de l'équation aux différences

$$x[k+1] = f(x[k]). \quad (2.3.1)$$

L'équation (2.3.1) décrit la loi de mise à jour d'un système *fermé*, c'est-à-dire une loi de mise à jour non affectée par un signal d'entrée externe. Dans ce sens, on peut considérer que les modèles d'états en temps-discret sont des systèmes dynamiques *ouverts*. Réciproquement, tout signal d'entrée constant dans un modèle d'état donne lieu à une trajectoire gouvernée par la loi d'évolution d'un système dynamique fermé (ou autonome).

Les systèmes dynamiques interviennent dans des domaines divers de la modélisation mathématique et peuvent donner lieu à des phénomènes très riches.

Exemple 2.2. Leonardo de Pisa (mieux connu sous le surnom de Fibonacci) décrit dans son traité d'arithmétique de 1202 l'évolution d'une population de lapins selon la loi suivante : une paire de lapins (mâle et femelle) engendre une nouvelle paire de lapins après deux saisons de reproduction. On suppose que les lapins ne meurent jamais et qu'ils peuvent se reproduire un nombre infini de fois. Si $N[n]$ représente le nombre de paires de lapins lors de la saison n , la loi d'évolution obéit l'équation aux différences $N[n+1] = N[n] + N[n-1]$. La solution de la récurrence donne lieu à la célèbre suite des nombres de Fibonacci :

$$1, 1, 2, 3, 5, 8, 13, 21, \dots$$

Le rapport $N[n+1]/N[n]$ converge asymptotiquement vers le nombre d'or $(\sqrt{5}-1)/2$. La suite de Fibonacci est la solution d'un système dynamique linéaire : en définissant les deux variables d'état $x_1[n] = N[n-1]$ et $x_2[n] = N[n-2]$, on obtient la loi de mise à jour

$$\begin{aligned} x_1[n+1] &= x_1[n] + x_2[n] \\ x_2[n+1] &= x_1[n] \end{aligned} \quad (2.3.2)$$

et la suite des nombres de Fibonacci est la solution x_1 obtenue pour la condition initiale $x_1[0] = 1$, $x_2[0] = 0$. Le système dynamique (2.3.2) est linéaire parce que la fonction de mise à jour f est une fonction linéaire de l'état. La solution générale des systèmes dynamiques linéaires sera étudiée au Chapitre 4. ◀

L'exemple 2.1 constitue un autre exemple de modèle d'état (infini) en temps-discret, issu d'une équation aux différences linéaire. Pour l'analyse des modèles linéaires étudiés dans les chapitres ultérieurs de ce cours, la structure d'espace vectoriel des ensembles U , Y , et X , joue un rôle fondamental.

Exemple 2.3. L'équation logistique

$$x[n+1] = \alpha x[n](1 - x[n]), \quad (2.3.3)$$

où $1 \leq \alpha \leq 4$ est un paramètre constant est un système dynamique dont l'étude a donné lieu à des livres entiers tant son comportement est riche. Une condition initiale dans l'intervalle $[0, 1]$ donne lieu à une solution dont toutes les itérées sont comprises dans le même intervalle. Le comportement limite des solutions est très dépendant du paramètre. Pour la valeur $\alpha = 4$, le système est chaotique : la plupart des solutions se promènent indéfiniment dans l'intervalle $[0, 1]$ sans jamais revenir à leur point de départ et deux conditions initiales arbitrairement proches donnent lieu à des solutions qui ont des comportements totalement différents après quelques itérations, alors même que le comportement est déterministe, c'est-à-dire que la condition initiale détermine univoquement toute la trajectoire. Ce comportement complexe résulte de la non linéarité de la fonction de mise à jour. Il n'existe pas dans les modèles linéaires étudiés dans ce cours. $\triangleleft$

2.4 Modèles d'état en temps-continu

Le modèle d'état a été décrit dans les sections précédentes dans un contexte discret : les valeurs successives des signaux correspondent à l'évolution d'une variable au cours d'une succession d'instantanés (ou événements) discrets. Le formalisme discret est bien adapté au concept de *mise à jour*. Le concept de modèle d'état peut néanmoins être très facilement adapté à des signaux en temps-continu. La définition de modèle d'état est inchangée moyennant deux modifications : d'une part, le domaine des signaux u , y , et x est le continuum $\mathbb{R}^+$ plutôt que l'ensemble discret $\mathbb{N}$; d'autre part, la loi de mise à jour pour $t \geq 0$ est remplacée par une équation différentielle

$$\dot{x}(t) = f(x(t), u(t)). \quad (2.4.1)$$

Tout comme dans le modèle d'état discret, un signal d'entrée $u : \mathbb{R}^+ \rightarrow U$ et une condition initiale x_0 déterminent le signal $x : \mathbb{R}^+ \rightarrow X$ solution de l'équation différentielle (2.4.1). Le signal de sortie y est déterminé pour $t \geq 0$ par l'équation de sortie

$$y(t) = h(x(t), u(t)). \quad (2.4.2)$$

L'existence d'une solution de l'équation différentielle (2.4.1) conditionne toutefois le caractère bien posé du modèle d'état. Il en découle des conditions plus restrictives sur l'espace d'état et sur la fonction de mise à jour que dans le cas discret. En général, on supposera que l'espace d'état est un ensemble

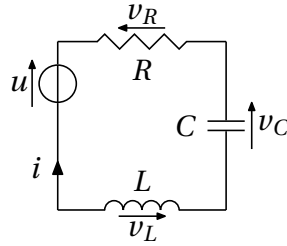


FIGURE 2.2 – Un circuit RLC.

ouvert de l'espace vectoriel $\mathbb{R}^N$ et on imposera des conditions de régularité sur la fonction f . Dans la suite du cours, on se restreindra rapidement à des fonctions f linéaires auquel cas on verra que l'existence et l'unicité des solutions ne pose pas de problème particulier.

Exemple 2.4. La loi de Newton $F = ma$ peut être interprétée comme un modèle d'état en temps-continu. La force F est le signal d'entrée u . L'état est constitué de la position q et de la vitesse $\dot{q}$ du corps considéré : $x = (q, \dot{q})$ prend ses valeurs dans $\mathbb{R}^2$. La loi de mise à jour est l'équation différentielle (2.4.1) avec $f(x(t), u(t)) = (\dot{q}(t), \frac{1}{m}F(t))$. $\triangleleft$

Le modèle d'état en temps-continu constitue le langage privilégié des modèles de la physique depuis Newton. Il est également largement utilisé dans la modélisation mathématique des phénomènes dynamiques rencontrés en biologie, en neurosciences, ou en économie. L'étude des modèles d'état dans une forme générale peut néanmoins s'avérer très complexe. Dans la suite du cours, on se limitera à l'étude des modèles linéaires pour lesquels il existe une théorie assez complète.

Exemple 2.5 (Circuit RLC). Le circuit RLC de la Figure 2.2 a comme signal d'entrée la tension u produite par le générateur.

Nous allons considérer deux choix différents pour la sortie :

- (i) la sortie est le courant i qui circule dans le circuit
- (ii) la sortie est la tension v_L au travers de l'inductance

Les équations qui décrivent le système sont obtenues en combinant la loi de Kirchhoff

$$u = v_R + v_C + v_L$$

et les équations constitutives des éléments du circuit :

$$v_R = Ri, \quad v_C = \frac{1}{C}q, \quad v_L = Li'$$

où q représente la charge de la capacité. On obtient par substitution

$$u = Ri + \frac{1}{C}q + Li'. \quad (2.4.3)$$

Pour obtenir un système différentiel, on peut dériver les deux membres de l'équation pour éliminer la variable $q = i'$, ce qui donne

$$u' = Ri' + \frac{1}{C}i + Li'' \quad (2.4.4)$$

Si la sortie y est le courant i , on obtient directement le système différentiel du deuxième ordre

$$\frac{1}{LC}y + \frac{R}{L}y' + y'' = \frac{1}{L}u'.$$

Si la sortie y est la tension $v_L = Li'$, il faut dériver une fois de plus l'équation (2.4.4) pour obtenir un nouveau système différentiel du deuxième ordre

$$\frac{1}{LC}y + \frac{R}{L}y' + y'' = u''.$$

On remarquera que les deux systèmes considérés se mettent très facilement sous forme d'état : en définissant le vecteur $x = (q, i) = [q \ i]^T$, on obtient l'équation

$$\begin{aligned} \dot{x}_1 &= x_2 \\ \dot{x}_2 &= -\frac{R}{L}x_2 - \frac{1}{LC}x_1 + \frac{1}{L}u \end{aligned}$$

Si la sortie est le courant, on a l'équation $y = x_2$; si la sortie est la tension, on obtient $y = -Rx_2 - \frac{1}{C}x_1 + u$. On a donc dans les deux cas une représentation d'état de la forme

$$\begin{aligned} \dot{x} &= Ax + Bu \\ y &= Cx + Du \end{aligned} \quad (2.4.5)$$

avec

$$A = \begin{bmatrix} 0 & 1 \\ -\frac{1}{LC} & -\frac{R}{L} \end{bmatrix}, \quad B = \begin{bmatrix} 0 \\ \frac{1}{L} \end{bmatrix}.$$

Pour l'équation de sortie, on obtient

$$C_i = \begin{bmatrix} 0 & 1 \end{bmatrix}, \quad D_i = 0$$

dans le cas $y = i$, ou

$$C_{v_L} = \begin{bmatrix} -\frac{1}{C} & -R \end{bmatrix}, \quad D_{v_L} = 1$$

dans le cas $y = v_L$.

L'exemple du circuit RLC peut bien sûr se généraliser à des circuits électriques linéaires plus complexes. <

Exemple 2.6 (Bras de robot). La Figure 2.3 représente un modèle de bras de robot très idéalisé : une masse ponctuelle m au bout d'une barre rigide et sans masse de longueur l , en rotation autour du pivot O . L'entrée du système est un couple externe v appliqué au bras, tandis que la sortie est la position

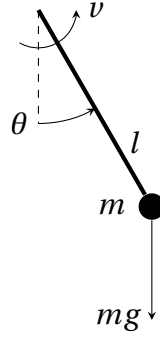


FIGURE 2.3 – Schéma simplifié d'un bras de robot.

angulaire θ de la barre par rapport à l'axe vertical. La masse m est soumise à la force de gravité, laquelle produit un couple $mg l \sin \theta$ sur la barre.

En appliquant la loi de Newton, on obtient l'équation du mouvement

$$J\ddot{\theta} = v - mgl \sin(\theta)$$

où $J = ml^2$ dénote le moment d'inertie du bras. Par substitution, on obtient l'équation différentielle

$$\ddot{\theta} + \frac{g}{l} \sin(\theta) = \frac{1}{ml^2} v. \quad (2.4.6)$$

Pour obtenir un modèle d'état, on définit l'état $x = (\theta, \dot{\theta}) = \begin{bmatrix} \theta & \dot{\theta} \end{bmatrix}^T$ et on réécrit l'équation scalaire (2.4.6) sous la forme d'une équation vectorielle (du premier ordre)

$$\begin{bmatrix} \dot{x}_1 \\ \dot{x}_2 \end{bmatrix} = \begin{bmatrix} x_2 \\ -\frac{g}{l} \sin x_1 + \frac{1}{ml^2} v \end{bmatrix}$$

qui définit un modèle d'état du type $\dot{x}(t) = f(x(t), v(t))$.

L'équation (2.4.6) est une équation non linéaire en raison du terme $\sin \theta$. A proprement parler, l'étude du bras de robot sort donc du cadre de l'étude des systèmes linéaires. On peut cependant obtenir une *approximation* linéaire du modèle (2.4.6) si l'on se restreint à l'étude de petits déplacements autour d'une solution particulière de l'équation. Par exemple, considérons la position d'équilibre du bras de robot correspondant à une entrée constante $v(t) \equiv v_0$. Si $|v_0| \leq mgl$, on vérifie que la paire (θ_0, v_0) est une solution particulière de l'équation (2.4.6) lorsque $mgl \sin \theta_0 = v_0$. Supposons que l'on s'intéresse à la loi du système pour des solutions proches de cette solution particulière. En utilisant le développement de Taylor

$$\sin(\theta) = \sin(\theta_0) + \cos(\theta_0)(\theta - \theta_0) + \mathcal{O}(\|\theta - \theta_0\|^2)$$

on voit que la variation de $\sin(\theta)$ lorsque θ varie légèrement autour de θ_0 est

donnée au premier ordre par le terme linéaire $\cos(\theta_0)(\theta - \theta_0)$. En définissant les variables d'écart $y = \theta - \theta_0$ et $u = v - v_0$, et en remplaçant la fonction non linéaire de l'équation (2.4.6) par son développement au premier ordre, on obtient l'équation

$$\ddot{y} + \frac{g}{l} \cos(\theta_0) y = \frac{1}{ml^2} u \quad (2.4.7)$$

L'équation (2.4.7) est appelée *équation variationnelle* du système (2.4.6) autour de la solution particulière (v_0, θ_0) . C'est un système différentiel linéaire à coefficients constants auquel on peut appliquer les méthodes d'analyse développées dans ce cours. Ce système différentiel est une bonne approximation du bras de robot pour décrire les solutions qui ne s'écartent pas trop de la solution d'équilibre considérée. $\triangleleft$

Exemple 2.7 (Discrétisation). Les systèmes aux différences proviennent souvent de la *discrétisation* de systèmes différentiels, par exemple lorsque l'on désire simuler le comportement du système sur un ordinateur. A titre illustratif, considérons le circuit RC de l'Exemple 1.2 sur la page 24 modélisé par le système différentiel

$$RC \dot{y} + y = u.$$

Pour une simulation numérique, il faut échantillonner les signaux, c'est-à-dire considérer l'équation aux instants discrets $t = kT$, $k \in \mathbb{Z}$, où T est la période d'échantillonnage. L'équation devient donc

$$RC \dot{y}(kT) + y(kT) = u(kT).$$

Pour obtenir une équation aux différences liant les signaux échantillonnés $u(kT)$ et $y(kT)$, il est nécessaire d'approximer la dérivée. Une approximation au premier ordre (approximation d'Euler) donne

$$\dot{y}(kT) = \frac{y(kT + T) - y(kT)}{T} + \mathcal{O}(T^2).$$

Par substitution dans l'équation différentielle, on obtient le système

$$\frac{RC}{T} (y((k+1)T) - y(kT)) + y(kT) = u(kT).$$

Ce système aux différences du premier ordre est une bonne approximation du circuit RC lorsque la période d'échantillonnage est suffisamment petite. En définissant l'état $x[n] = y(nT)$, on obtient directement l'équation de mise à jour

$$x[n+1] = \left(1 - \frac{T}{RC}\right)x[n] + \frac{T}{RC}u[n]$$

et l'équation de sortie $y = x$. $\triangleleft$

2.5 Interconnexions

La notion d'interconnexion est centrale dans la construction de modèles d'état tout comme dans la construction de systèmes (voir Section 1.3.5). Un modèle d'état complexe est souvent décrit comme l'interconnexion d'un ensemble de modèles d'état plus simples.

L'interconnexion *feedforward* ou *série* ou *cascade* de deux modèles d'état consiste à utiliser la sortie du premier modèle comme entrée du second. Le modèle $\Sigma_1(U_1, X_1, Y_1, f_1, h_1)$ et le modèle $\Sigma_2(U_2, X_2, Y_2, f_2, h_2)$ ainsi que la loi d'interconnexion $u_2 = y_1$ conduisent à un nouveau modèle d'état $\Sigma(U, X, Y, f, h)$ défini par $U = U_1$, $Y = Y_2$, $X = X_1 \times X_2$,

$$f : X \times U \rightarrow X : ((x_1, x_2), u_1) \mapsto (f_1(x_1, u_1), f_2(x_2, h_1(x_1, u_1))),$$

et $h = h_2$. Le modèle série est bien posé sous la seule condition que $Y_1 \subset U_2$.

L'interconnexion *feedback* ou *en rétroaction* de deux modèles consiste à utiliser la sortie du premier modèle comme entrée du second et vice-versa. La loi d'interconnexion est cette fois $u_2 = y_1$, $u_1 = y_2$. Il s'agit d'une loi d'interconnexion importante mais plus délicate que l'interconnexion série. En particulier, l'interconnexion peut avoir un caractère mal posé en raison de la définition implicite de certains signaux. Par exemple, si $U_1 = Y_1 = U_2 = Y_2 = \mathbb{R}$, l'interconnexion feedback n'admet pas de solution pour le modèle défini par les fonctions de sortie $h_1 = u_1$, $h_2 = u_2 - 1$ parce que la loi d'interconnexion conduit aux équations $u_2 = u_1$, $u_1 = u_2 - 1$. On peut éviter le caractère mal posé de l'interconnexion feedback de deux modèles d'état en imposant qu'une des deux fonctions de sortie au moins ne dépende pas de l'entrée.

Des lois d'interconnexion plus générales sont obtenues en combinant ces deux lois simples ou en les appliquant à une partie des signaux seulement. On trouvera fréquemment des exemples dans lesquels une partie des signaux d'entrée/sortie définissent l'interaction avec le monde extérieur et une autre partie des signaux d'entrée/sortie proviennent d'interconnexions avec d'autres systèmes.

2.6 Extensions

Le modèle d'état est un modèle général approprié à la description d'un grand nombre de systèmes dynamiques rencontrés dans le monde physique ou dans le monde informatique. Un certain nombre d'extensions méritent d'être mentionnées même si elles ne seront pas considérées dans le cadre de ce cours.

2.6.1 Modèles d'état dépendant du temps

Dans notre définition de modèle d'état, la fonction de mise à jour f et la fonction de sortie h ne dépendent pas du temps. Il en résulte un modèle d'état dit *invariant par rapport au temps*. La définition de modèle d'état peut être aisément adaptée à une fonction de mise à jour et une fonction de sortie qui dépendent du temps. Dans ce cas, la loi de mise à jour au temps t_0 peut être différente de la loi de mise à jour au temps t_1 . Dans ce cas, le temps initial t_0 ne peut plus être fixé arbitrairement à $t = 0$. Le temps initial t_0 fait partie de la condition initiale et intervient également dans la définition du domaine des signaux d'entrée et de sortie. L'étude des modèles dépendant du temps n'est pas poursuivie dans le cadre de ce cours (voir Section 3.1).

2.6.2 Modèles d'état non déterministes

Les modèles d'état que nous avons considéré sont *déterministes* : la condition initiale du modèle détermine entièrement le comportement du système. Autrement dit, pour chaque signal d'entrée, une seule trajectoire de sortie est possible une fois que la condition initiale est fixée. Un modèle non déterministe autorise une certaine indétermination au-delà du choix de la condition initiale. Par exemple, la loi de mise à jour pourrait être de type probabiliste, auquel cas on remplacerait la fonction de mise à jour par une distribution de probabilité. Un modèle d'état non déterministe peut par exemple décrire de manière plus grossière mais plus économe un modèle déterministe comportant un très grand nombre de variables.

2.6.3 Modèles d'état hybrides

Les automates finis sont principalement rencontrés en informatique, comme mode de description d'un programme. Les modèles d'état en temps continu sont principalement rencontrés en physique, comme mode de description du monde naturel qui nous entoure. Les modèles d'état hybrides considèrent l'interconnexion d'automates finis et de modèles en temps-continu. Ils sont très étudiés à l'heure actuelle car ils permettent de modéliser les processus continus commandés par des automates finis.

Chapitre 3

Systèmes linéaires invariants : représentation convolutionnelle

Deux modes de représentation prévalent pour l'analyse des systèmes : la représentation externe ou *entrée-sortie*, détaillée dans ce chapitre, et la représentation interne ou *d'état*, détaillée dans le chapitre suivant.

Dans la représentation entrée-sortie, le système est vu comme un opérateur S qui associe à chaque signal d'entrée u une et une seule sortie y :

$$y = S(u).$$

Moyennant les hypothèses fondamentales de *linéarité* et d'*invariance* sur l'opérateur, nous montrons que le système est entièrement caractérisé par sa *réponse impulsionnelle*, c'est-à-dire la réponse obtenue lorsque l'entrée est une impulsion δ . Cette propriété est une conséquence directe du principe de superposition et de la décomposition de n'importe quel signal en une combinaison d'impulsions décalées.

Dans la dernière section du chapitre, nous calculons l'opérateur linéaire temps-invariant (LTI) associé à un système différentiel ou aux différences du premier ordre, en imposant la condition dite de *repos initial*.

3.1 Linéarité et invariance

L'invariance et la linéarité sont deux propriétés cruciales pour la théorie des systèmes. Il n'y pas de commune mesure entre les résultats connus pour les systèmes linéaires invariants et ceux connus pour les systèmes variant dans le temps ou/et non linéaires. Nous allons donc définir ces deux propriétés et, dans la suite du cours, nous restreindre à l'étude des systèmes linéaires et invariants.

Un système est dit *invariant (dans le temps)* lorsque la loi qu'il établit entre entrées et sorties ne change pas au cours du temps. Dans un système invariant, la réponse à une entrée donnée sera la même aujourd'hui ou de-

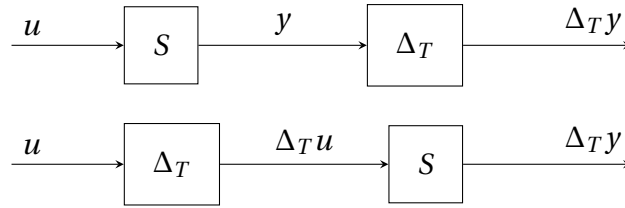


FIGURE 3.1 – Invariance de l'opérateur S : Il commute avec l'opérateur de décalage Δ_T .

main. Mathématiquement, l'invariance s'exprime par la propriété de commutation entre l'opérateur définissant le système et l'opérateur de décalage (Figure 3.1) : Si (u, y) est une solution du système, alors la paire « décalée » $(\Delta_T u, \Delta_T y) = (u(\cdot + T), y(\cdot + T))$ est aussi une solution du système, pour n'importe quelle constante T . Le système $y(t) = (u(t))^2$ est invariant; le système $y(t) = u(t) - 2t$ ne l'est pas.

Un système est *linéaire* lorsqu'il vérifie le principe de superposition : si l'entrée est la somme pondérée de deux signaux, la sortie est la même somme pondérée des deux réponses correspondantes. Mathématiquement, si (u_1, y_1) et (u_2, y_2) sont deux solutions du système, toute combinaison $(au_1 + bu_2, ay_1 + by_2)$ est également une solution, quelles que soient les scalaires a et b .

La linéarité a des conséquences très importantes pour l'analyse : si l'on connaît la réponse du système pour certaines entrées simples, alors il suffit d'exprimer une entrée plus compliquée comme une combinaison des entrées simples pour pouvoir calculer la sortie correspondante. Ce principe est le fondement des outils d'analyse développés dans les chapitres suivants.

Exemple 3.1 (Compte d'épargne). Supposons que le solde d'un compte d'épargne au jour n est représenté par $y[n]$, et que $u[n]$ représente le montant déposé au jour n . Si l'intérêt est calculé et crédité une fois par jour à un taux de α pour cent, le solde du compte au jour $n + 1$ est donné par

$$y[n + 1] = (1 + 0.01\alpha)y[n] + u[n]$$

C'est un système aux différences du premier ordre. Le système est linéaire et invariant. On notera que l'invariance du système repose sur l'hypothèse d'un taux constant et que la linéarité du système suppose que le taux d'intérêt est identique suivant que le solde du compte est positif ou négatif. $\triangleleft$

Une conséquence directe de la linéarité est la propriété d'*homogénéité* : en choisissant $b = 0$ dans la propriété de superposition, on obtient que la paire (au_1, ay_1) est solution du système pour n'importe quelle valeur de a . En

particulier, une entrée nulle produit nécessairement une sortie nulle (choisir $a = 0$ dans la propriété d'homogénéité).

La non-linéarité d'un système est souvent mise en évidence en montrant que le système n'est pas homogène. Il importe cependant de ne pas confondre les deux notions. Pour qu'un système homogène soit linéaire, il doit en outre posséder la propriété d'*additivité* : en choisissant $a = b = 1$ dans la propriété superposition, on obtient que $y_1 + y_2$ est la réponse obtenue pour l'entrée $u_1 + u_2$.

Exercice 3.2. Vérifier que le système $y[n] = \Re(u[n])$ est additif mais pas homogène. Trouver un exemple de système homogène mais pas linéaire. <

3.2 La convolution pour les systèmes discrets

Un des signaux discrets les plus simples est l'*impulsion unité* $\delta[\cdot]$ définie par

$$\delta[n] = \begin{cases} 0, & n \neq 0, \\ 1, & n = 0. \end{cases} \quad (3.2.1)$$

Pour extraire la valeur d'un signal u à l'instant $n = 0$, il suffit de le multiplier par l'impulsion $\delta[\cdot]$, i.e.

$$u[\cdot]\delta[\cdot] = u[0]\delta[\cdot]. \quad (3.2.2)$$

Le résultat de cette opération est un signal qui vaut $u[0]$ à l'instant $n = 0$ et est nul partout ailleurs. De même, on peut extraire la valeur du signal à n'importe quel autre instant $n = k$ en le multipliant par l'impulsion décalée $\delta[\cdot - k]$. Par exemple, la somme

$$u[\cdot](\delta[\cdot] + \delta[\cdot - k])$$

donne un signal qui vaut $u[0]$ à l'instant $n = 0$, $u[k]$ à l'instant $n = k$, et qui est nul partout ailleurs. Pour reconstruire le signal $u[\cdot]$ tout entier, on peut donc utiliser la somme infinie

$$u[\cdot] = \sum_{k \in \mathbb{Z}} u[k]\delta[\cdot - k] \quad (3.2.3)$$

qui exprime le fait que n'importe quel signal peut être décomposé en une somme (infinie) d'impulsions décalées. Dans cette décomposition, le poids donné à l'impulsion $\delta[\cdot - k]$ est la valeur du signal à l'instant k .

Exemple 3.3. Le signal *échelon-unité* est défini par $\mathbb{I}[n] = 1$ pour $n \geq 0$ et $\mathbb{I}[n] = 0$ pour $n < 0$. Sa représentation en somme d'impulsions donne donc

$$\mathbb{I}[n] = \sum_{k=0}^{\infty} \delta[n - k]. \quad (3.2.4)$$

<

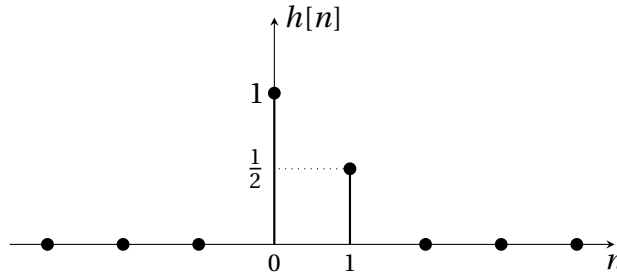


FIGURE 3.2 – La réponse impulsionnelle du système $y[n] = u[n] + \frac{1}{2}u[n-1]$.

Cherchons maintenant à calculer la réponse y d'un système LTI à une entrée quelconque u . Nous allons voir que la réponse y peut être calculée à partir de la *réponse impulsionnelle du système*, i.e. sa réponse à l'entrée $\delta[\cdot]$. La réponse impulsionnelle est notée par h . Si le système est invariant, sa réponse à l'entrée décalée $\delta[\cdot - k]$ sera par définition la réponse impulsionnelle décalée $h[\cdot - k]$. Si le système est également linéaire, la réponse à une combinaison d'entrées sera, en vertu du principe de superposition, la même combinaison des sorties correspondantes. De la représentation (3.2.3), on déduit directement que la réponse à l'entrée u est donnée par

$$y[\cdot] = \sum_{k \in \mathbb{Z}} u[k] h[\cdot - k]. \quad (3.2.5)$$

La sommation dans le membre de droite de (3.2.5) est appelée *convolution* des signaux u et h . Cette opération est notée

$$y = u * h. \quad (3.2.6)$$

La relation (3.2.6) montre qu'un système LTI est entièrement caractérisé par sa réponse impulsionnelle h .

Le produit de convolution (3.2.5) est une représentation fondamentale des systèmes LTI. Il exprime que la réponse du système à un instant n donné, par exemple $y[n]$, dépend, en principe, de toutes les valeurs passées et futures du signal d'entrée, chaque valeur $u[k]$ étant pondérée par un facteur $h[n - k]$ qui traduit l'effet plus ou moins important de « mémoire » du système. En pratique, la réponse impulsionnelle d'un système aura typiquement une fenêtre temporelle limitée (traduisant la mémoire limitée du système), c'est-à-dire que le signal $h[n]$ décroîtra rapidement vers zéro. Dans ce cas, la valeur de sortie $y[n]$ ne sera affectée que par les valeurs « voisines » de l'entrée autour de $u[n]$. La Figure 3.2 représente un cas simplifié à l'extrême où la réponse impulsionnelle décroît vers zéro en deux instants, ce qui correspond au système très simple $y[n] = u[n] + \frac{1}{2}u[n-1]$.

Pour des systèmes plus compliqués, il est très utile de recourir à une visualisation graphique de l'opération de convolution : pour calculer $y[1]$, il faut visualiser les graphes des signaux u et $h[1 - \cdot]$. Le signal $h[1 - \cdot]$ est obtenu en réfléchissant la réponse impulsionnelle $h \mapsto h[-\cdot]$ et en la décalant d'une unité vers la droite $h[-\cdot] \mapsto h[-\cdot + 1]$. Ensuite on multiplie les deux graphes et on somme toutes les valeurs ainsi obtenues. Pour obtenir $y[2]$, il suffit de décaler $h[1 - \cdot]$ une nouvelle fois vers la droite ($h[1 - \cdot] \mapsto h[2 - \cdot]$) et de recommencer l'opération. On obtient ainsi les valeurs successives de la sortie en faisant glisser vers la droite la réponse impulsionnelle réfléchie. Cette procédure est illustrée à la Figure 3.3 sur un exemple. On déduit immédiatement de la figure les valeurs de la réponse $y[0] = 0.5$, $y[1] = y[2] = 2.5$, $y[3] = 2$; pour toutes les autres valeurs de n , la réponse $y[n]$ est nulle parce que les graphes des signaux u et $h[n - \cdot]$ n'ont pas de support commun. Ceci exprime simplement le fait que dans les systèmes rencontrés en pratique, une excitation limitée dans le temps donne lieu à une réponse limitée dans le temps. Cette propriété sera formalisée plus loin.

3.3 L'impulsion de Dirac

Pour reproduire la représentation convolutionnelle d'un système LTI discret dans le cas continu, il nous faut introduire l'analogue continu $\delta(\cdot)$ du signal impulsionnel $\delta[\cdot]$. L'analogie avec les équations (3.2.2) et (3.2.4) suggère de définir un signal $\delta(\cdot)$ qui satisfait

$$u(\cdot)\delta(\cdot) = u(0)\delta(\cdot) \quad (3.3.1)$$

pour n'importe quel signal u (continu en $t = 0$), et

$$\mathbb{I}(t) = \int_0^\infty \delta(t - \tau) d\tau \quad (3.3.2)$$

où la sommation dans le cas discret a naturellement été remplacée par une intégrale. En effectuant le changement de variable $s = t - \tau$, on peut réécrire (3.3.2) comme

$$\mathbb{I}(t) = \int_{-\infty}^t \delta(s) ds. \quad (3.3.3)$$

La relation (3.3.3) signifie que l'échelon unité $\mathbb{I}(\cdot)$ est une primitive du signal $\delta(\cdot)$, ou encore que *le signal $\delta(\cdot)$ est la dérivée de la fonction $\mathbb{I}(\cdot)$* . Une telle définition pose un problème mathématique puisque la fonction $\mathbb{I}(\cdot)$, qui est discontinue à l'origine, n'est pas dérivable.

Le signal $\delta(\cdot)$ qui satisfait (3.3.3) est appelé *impulsion de Dirac*. Ce n'est pas une fonction au sens usuel, et pour cette raison, la justification rigoureuse d'un calcul différentiel faisant usage d'impulsions de Dirac nécessite une

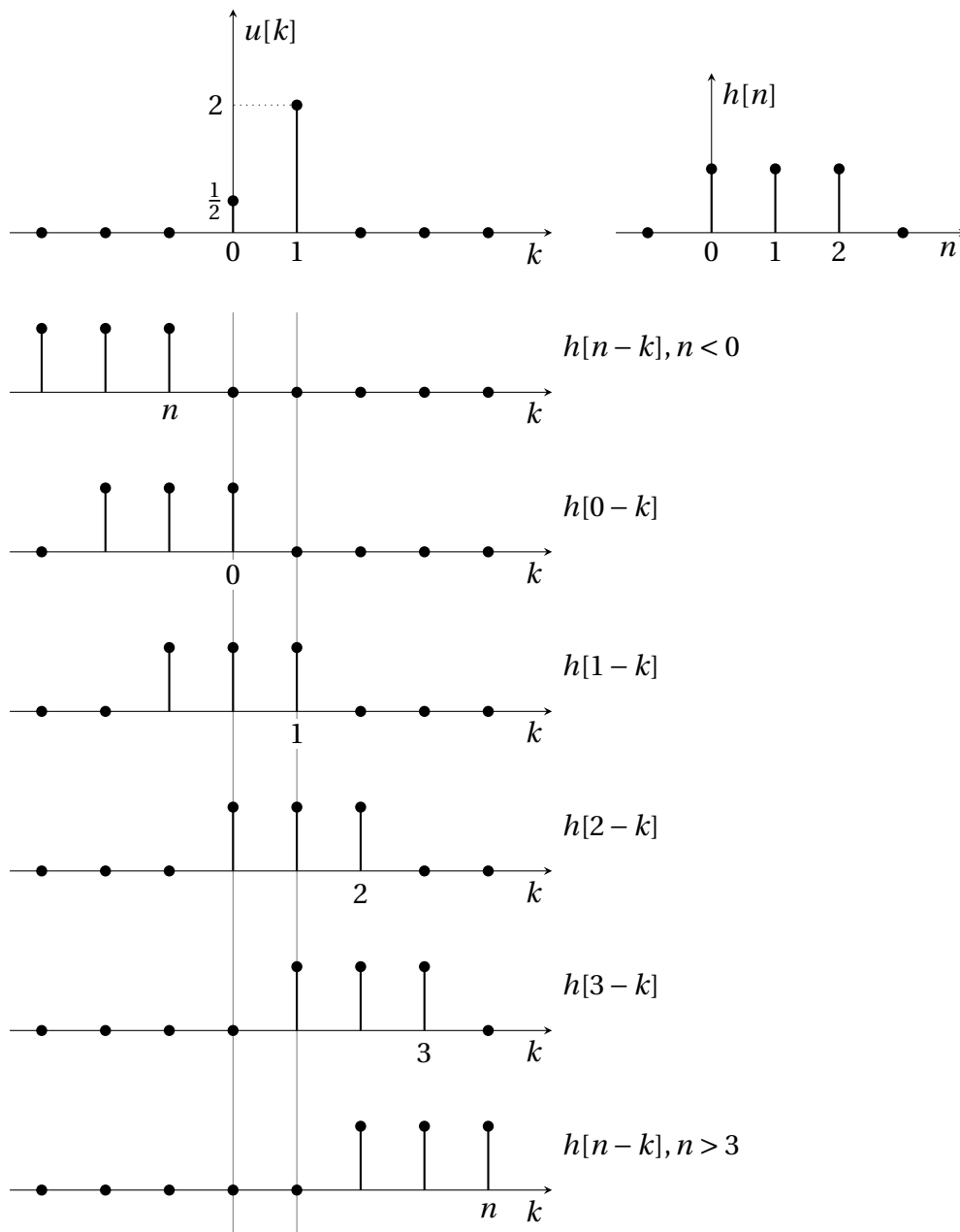


FIGURE 3.3 – Représentation graphique du produit de convolution pour calculer la sortie d'un système LTI.

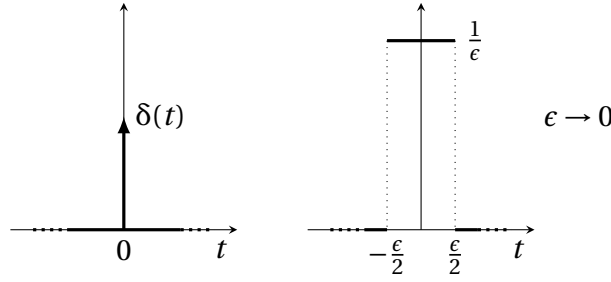


FIGURE 3.4 – Une impulsion de Dirac et son approximation.

théorie à part entière (la théorie des distributions développée par L. Schwartz dans les années 1945–1950). Néanmoins, on peut en faire un usage correct sans aucune difficulté particulière (les physiciens en firent bon usage longtemps avant sa justification mathématique). Comme illustré à la Figure 3.4, il suffit de concevoir l'impulsion $\delta(\cdot)$ comme la limite pour $\epsilon \rightarrow 0$ d'une fonction rectangulaire d'amplitude $1/\epsilon$ sur une fenêtre de largeur ϵ .

Lorsque ϵ tend vers zéro, le support de l'impulsion rétrécit mais son intégrale reste constante. La limite pour ϵ tendant vers zéro n'existe pas au sens usuel des fonctions mais donne les propriétés

$$\forall t \neq 0 : \delta(t) = 0, \quad \int_{-\infty}^{\infty} \delta(t) dt = 1. \quad (3.3.4)$$

La définition (3.3.4) est équivalente (au sens des distributions) aux définitions (3.3.1) ou (3.3.3).

Si l'on garde à l'esprit que les impulsions et a fortiori leurs dérivées ne sont pas des fonctions mais que leur usage dans le calcul différentiel peut être justifié rigoureusement au sens des distributions, on peut manipuler les impulsions *comme s'il s'agissait de fonctions* en utilisant leur définition et les règles usuelles du calcul. On peut additionner des impulsions, les multiplier par une fonction (en utilisant la propriété (3.3.1), qui est valide pour autant que la fonction u soit continue à l'origine), et leur appliquer les transformations temporelles (décalage, réflexion, changement d'échelle de temps).

Pour donner un sens aux « dérivées » de $\delta(\cdot)$, on utilise la règle de dérivation en chaîne :

$$\int_{-\infty}^{\infty} \dot{\delta}(t) u(t) dt = \delta(\cdot) u(\cdot) \Big|_{-\infty}^{\infty} - \int_{-\infty}^{\infty} \delta(t) \dot{u}(t) dt = -\dot{u}(0)$$

et la relation

$$\int_{-\infty}^{\infty} \dot{\delta}(t) u(t) dt = -\dot{u}(0)$$

a valeur de définition pour $\dot{\delta}(\cdot)$. En particulier, on vérifie facilement que

$$u(\cdot)\dot{\delta}(\cdot) = u(0)\dot{\delta}(\cdot) - \dot{u}(0)\delta(\cdot).$$

3.4 Convolution de signaux continus

Ayant caractérisé l'analogue continu $\delta(\cdot)$ de l'impulsion discrète $\delta[\cdot]$, nous pouvons à présent poursuivre la représentation de convolution d'un système LTI continu en analogie complète avec le cas discret. L'analogue de l'équation (3.2.3) donne

$$u(\cdot) = \int_{-\infty}^{+\infty} u(\tau)\delta(\cdot - \tau) d\tau \quad (3.4.1)$$

et permet donc de représenter n'importe quel signal comme une combinaison infinie d'impulsions (en concevant l'intégrale comme la limite d'une somme). Si on désigne par h la réponse impulsionnelle d'un système en temps continu LTI, on obtient directement que sa réponse à une entrée quelconque est donnée par

$$y(t) = \int_{-\infty}^{+\infty} u(\tau)h(t - \tau) d\tau. \quad (3.4.2)$$

Comme dans le cas discret, on note cette opération de convolution par

$$y = u * h.$$

L'interprétation du produit de convolution et sa visualisation graphique sont rigoureusement analogues au cas discret. A titre d'exemple, on peut calculer graphiquement en s'aidant de la Figure 3.5 les valeurs de la sortie

$$y(t) = \begin{cases} 0, & t < 0, \\ \frac{1}{2}t^2, & 0 < t < T, \\ Tt - \frac{1}{2}T^2, & T < t < 2T, \\ -\frac{1}{2}t^2 + Tt + \frac{3}{2}T^2, & 2T < t < 3T, \\ 0, & 3T < t, \end{cases}$$

correspondant à la convolution des signaux u et h représentés.

Exercice 3.4. Montrer que l'opération de convolution est commutative, distributive, et associative. En déduire que si deux systèmes ont des réponses impulsionnelles respectives h_1 et h_2 , leur interconnexion parallèle a comme réponse impulsionnelle la somme $h_1 + h_2$ et leur interconnexion série a comme réponse impulsionnelle le produit $h_1 * h_2$. $\triangleleft$

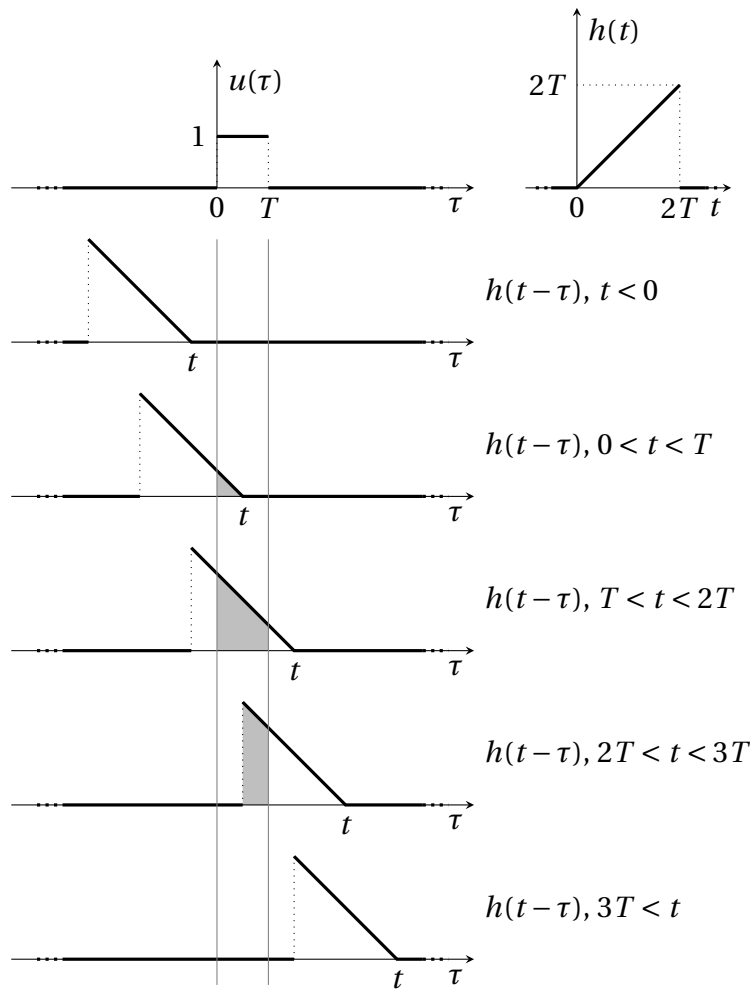


FIGURE 3.5 – Représentation graphique du produit de convolution pour calculer la sortie d'un système continu LTI.

3.5 Causalité, mémoire, et temps de réponse des systèmes LTI

La réponse impulsionnelle d'un système LTI est sa « carte d'identité » dans le domaine temporel et certaines propriétés du système se lisent très facilement sur les caractéristiques de la réponse impulsionnelle.

Par exemple, la propriété de causalité est directement lisible sur la réponse impulsionnelle. Si la sortie $y[\cdot]$ ne peut pas dépendre des entrées futures du système, il est nécessaire et suffisant qu'un poids nul leur soit affecté dans la représentation convolutionnelle (3.2.5). Ceci implique

$$h[k] = 0, \quad k < 0,$$

ce qui exprime simplement la causalité du système pour l'entrée particulière $\delta[\cdot]$.

Nous avons déjà mentionné plus haut le lien entre la mémoire d'un système LTI et la largeur de la fenêtre de sa réponse impulsionnelle. Pour un système statique, la sortie $y[n]$ à l'instant n ne peut dépendre que de l'entrée $u[n]$ à l'instant n , ce qui impose $h[n] = 0$ pour $n \neq 0$. La convolution se réduit dans ce cas à

$$y = Ku$$

qui est la forme générale d'un système LTI statique. On dit dans ce cas que le système est juste un *gain* statique; la facteur de gain K peut dépendre du temps.

Lorsque l'on veut modéliser les effets dynamiques (la mémoire) d'un système, la première caractéristique qualitative que l'on cherche à évaluer est la *constante de temps* du système. Cette constante traduit le fait observé dans la plupart des systèmes physiques que le système a un certain *temps de réponse*: contrairement à un système statique, une variation instantanée du signal d'entrée ne donne pas lieu à une variation instantanée du signal de sortie.

Une conséquence immédiate de la représentation convolutionnelle est que le temps de réponse d'un système LTI est lié à la largeur (durée) de sa réponse impulsionnelle $h(\cdot)$. En effet, si la réponse impulsionnelle a une durée T_h , alors la réponse à une entrée de durée T_u sera la convolution de ces deux signaux et aura donc une durée $T_u + T_h$. La constante de temps mesure donc la rapidité du système: un système avec une constante de temps très petite réagit presque instantanément, tandis qu'un système avec une constante de temps très grande ne sera pas capable de « suivre » une entrée qui varie rapidement.

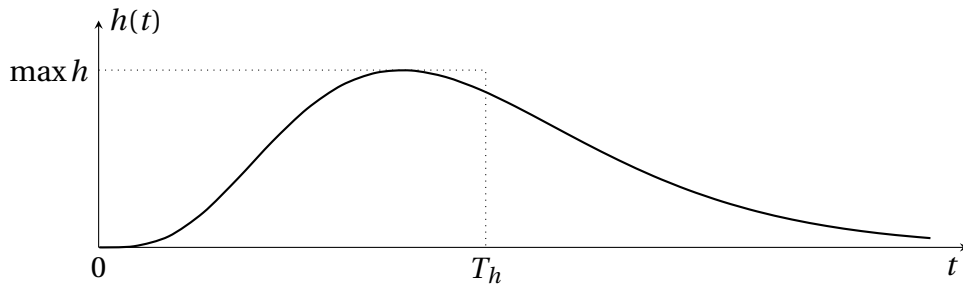
En fait, la plupart des systèmes que nous allons rencontrer sont représentés par une équation différentielle ou aux différences. Nous verrons que la réponse impulsionnelle de ces systèmes a typiquement une allure exponentielle décroissante, c'est-à-dire qu'elle a une durée infinie. Pour définir la constante de temps de tels systèmes, on adopte par exemple la convention

$$T_h = \frac{\int_{-\infty}^{\infty} h(t) dt}{\max h}$$

où $\max h$ désigne le maximum atteint par la réponse impulsionnelle. Suivant cette définition, le rectangle de hauteur $\max h$ et de largeur T_h a la même surface que l'intégrale de la réponse impulsionnelle (cf. Figure 3.6).

Pour une réponse impulsionnelle exponentielle

$$h(t) = \mathbb{I}(t) Ae^{-\lambda t}$$

FIGURE 3.6 – Le temps de réponse T_h d'un système continu LTI.

la constante de temps vaut

$$T_h = \frac{1}{A} \int_0^{\infty} A e^{-\lambda t} dt = \frac{1}{\lambda}.$$

Dans toutes les applications de la théorie des systèmes (commande, filtrage, communications), la constante de temps joue un rôle très important. En commande, un critère de performance important est le *temps de montée* du système : le temps nécessaire à la sortie pour atteindre la valeur de consigne lorsque le système est soumis à un échelon. Si la réponse impulsionnelle est un rectangle de hauteur unitaire et de largeur T_h , sa convolution avec un échelon donne

$$y(t) = \int_0^{\min\{t, T_h\}} d\tau = \min\{t, T_h\}.$$

Dans ce cas simple, il faut un temps T_h pour que la sortie atteigne la valeur constante $y = T_h$. Le temps de montée est donc égal à la constante de temps.

En filtrage, la constante de temps est inversement proportionnelle à la fréquence de coupure : si une sinusoïde de haute fréquence est appliquée à l'entrée d'un système ayant une grande constante de temps, la sortie n'est pas capable de suivre l'entrée : le système agit comme un filtre passe-bas, capable de « suivre » les basses fréquences mais éliminant les hautes fréquences.

Exercice 3.5. Soit un système de réponse impulsionnelle $h(t) = e^{-t}\mathbb{I}(t)$. Calculer la réponse à l'entrée $u(t) = \sin \omega t$. Montrer que la sortie suit bien les basses fréquences ($\omega \ll 1$) et ne suit pas les hautes fréquences ($\omega \gg 1$). ◁

En communication, on transmet des *pulses* (signaux de courte durée). Pour éviter les interférences, il faut éviter que les pulses se mélangent à la sortie du canal de transmission. Si le pulse émis est de durée T_u , le pulse de sortie est de durée $T_u + T_h$. Il faut donc espacer l'émission des pulses par des intervalles de T_h secondes, ou autrement dit, la vitesse de transmission ne peut dépasser $\frac{1}{T_h + T_u}$ pulses par seconde. La vitesse de transmission est donc limitée par la constante de temps du système de transmission.

Certaines performances caractéristiques d'un système LTI peuvent donc être évaluées à partir du graphe de sa réponse impulsionnelle. Alternative-ment – en particulier dans les applications de commande –, on utilise la réponse *indicielle* du système, i.e. la réponse obtenue avec l'entrée *échelon* $u = \mathbb{1}$. Par définition, la réponse indicielle d'un système continu prend la forme

$$s(t) = (h * \mathbb{1})(t) = \int_{-\infty}^t h(\tau) d\tau.$$

C'est donc l'intégrale de la réponse impulsionnelle. La réponse indicielle peut être déterminée expérimentalement, en soumettant le système à un signal échelon unitaire.

3.6 Conditions de repos initial

Le filtre exponentiel considéré dans l'Exemple 1.1 est un système aux différences du premier ordre décrit par l'équation

$$y[n] = ay[n-1] + (1-a)u[n], \quad 0 < a < 1 \quad (3.6.1)$$

et dont la solution est donnée par la formule explicite

$$y[n_0 + n] = a^n y[n_0] + (1-a) \sum_{k=0}^{n-1} a^k u[n_0 + n - k], \quad n \geq 0. \quad (3.6.2)$$

L'identification de la solution (3.6.2) avec un produit de convolution

$$y[n_0 + n] = \sum_{k \in \mathbb{Z}} h[k] u[n_0 + n - k]$$

suggère la réponse impulsionnelle causale

$$h[n] = (1-a)\mathbb{1}[n]a^n \quad (3.6.3)$$

qui conduit à l'expression de convolution

$$y[n_0 + n] = (1-a) \sum_{k=0}^{+\infty} a^k u[n_0 + n - k]. \quad (3.6.4)$$

Pour obtenir l'équivalence entre les expressions (3.6.2) et (3.6.4), deux hypothèses doivent être introduites :

- (i) $\forall n \leq n_0 : u[n] = 0$,
- (ii) $y[n_0] = 0$.

Les deux conditions (i) et (ii) sont appelées conditions de *repos initial*. Elles permettent d'associer au système aux différences un opérateur LTI causal

unique de réponse impulsionnelle (3.6.3). Cette propriété sera généralisée à des systèmes aux différences d'ordre quelconque dans le Chapitre 4.

On notera que si u est un signal nul pour $n \leq n_0$, l'opérateur LTI causal ainsi défini donne une solution qui correspond à la solution du système aux différences *pour la condition de repos initial* $y[n_0] = 0$. Cette restriction sur le choix de la condition initiale est une limitation de la représentation convolutionnelle.

Un traitement analogue peut être mené pour l'Exemple 1.2 du circuit RC modélisé par le système différentiel

$$\dot{y} + \frac{1}{RC}y = \frac{1}{RC}u \quad (3.6.5)$$

dont la solution à partir d'un instant initial t_0 est donnée par

$$y(t_0 + t) = e^{-\frac{t}{RC}} y(t_0) + \frac{1}{RC} \int_{t_0}^{t_0+t} e^{-\frac{(t_0+t-\tau)}{RC}} u(\tau) d\tau, \quad t \geq 0. \quad (3.6.6)$$

L'identification de la solution (3.6.6) avec un produit de convolution

$$y(t_0 + t) = \int_{-\infty}^{+\infty} h(t_0 + t - \tau) u(\tau) d\tau$$

suggère la réponse impulsionnelle causale

$$h(t) = \frac{1}{RC} \mathbb{I}(t) e^{-\frac{t}{RC}} \quad (3.6.7)$$

qui conduit à l'expression

$$y(t_0 + t) = \frac{1}{RC} \int_{-\infty}^{t_0+t} e^{-\frac{(t_0+t-\tau)}{RC}} u(\tau) d\tau. \quad (3.6.8)$$

Pour obtenir l'équivalence entre les expressions (3.6.6) et (3.6.8), deux hypothèses doivent être introduites :

- (i) $\forall t < t_0 : u(t) = 0$,
- (ii) $y(t_0) = 0$.

Les deux conditions (i) et (ii) sont appelées conditions de *repos initial*. Elles permettent d'associer au système différentiel un opérateur LTI causal unique de réponse impulsionnelle (3.6.7). Cette propriété sera généralisée à des systèmes différentiels d'ordre quelconque dans le Chapitre 4.

On notera que si u est un signal nul pour $t < t_0$, l'opérateur LTI causal ainsi défini donne une solution qui correspond à la solution du système différentiel *pour la condition de repos initial* $y(t_0) = 0$. Cette restriction sur le choix de la condition initiale est une limitation de la représentation convolutionnelle.

Chapitre 4

Modèles d'état linéaires invariants

Une classe importante de systèmes LTI est spécifiée par le modèle d'état introduit au Chapitre 2. Pour donner lieu à un système LTI, l'équation de mise à jour et l'équation de sortie du modèle d'état doivent être linéaires et invariantes, c'est-à-dire de la forme

$$\begin{aligned} \sigma^{-1}x &= Ax + Bu, & \text{ou} & & x[n+1] &= Ax[n] + Bu[n], \\ y &= Cx + Du, & & & y[n] &= Cx[n] + Du[n], \end{aligned} \quad (4.0.1)$$

pour des modèles en temps-discret et de la forme

$$\begin{aligned} \dot{x} &= Ax + Bu, & \text{ou} & & \dot{x}(t) &= Ax(t) + Bu(t), \\ y &= Cx + Du. & & & y(t) &= Cx(t) + Du(t). \end{aligned} \quad (4.0.2)$$

pour des modèles en temps-continu. Dans ces expressions, A , B , C , et D sont, en général, des matrices. La dimension n du vecteur d'état x est appelée la dimension du système. Dans cette forme générale, le signal d'entrée u peut être un vecteur de m composantes et le signal de sortie y un vecteur de p composantes, auquel cas les dimensions des matrices sont $n \times n$ pour A , $n \times m$ pour B , $p \times n$ pour C , et $p \times m$ pour D . Les espaces définissant le modèle d'état sont donc des espaces vectoriels : $X = \mathbb{R}^n$, $U = \mathbb{R}^m$, et $Y = \mathbb{R}^p$.

Ce chapitre est consacré à l'étude des modèles d'état linéaires invariants et aux systèmes LTI qui leur sont associés. Nous revenons d'abord sur la notion d'état, en faisant apparaître son rôle de paramétrisation de la mémoire du système. Nous discutons ensuite comment une représentation d'état peut être construite au départ d'un système différentiel ou aux différences, et nous montrons que ce problème est étroitement lié au problème de la *réalisation* d'un système. La solution explicite des équations d'état d'un système linéaire invariant est établie dans la section 4.3. Elle permet d'associer un et un seul système LTI causal à chaque modèle d'état en imposant une condition de repos initial. Les transformations d'état sont discutées dans la Section 4.4. La Section 4.5 établit comment un modèle d'état linéaire invariant peut être

obtenu par linéarisation d'un modèle plus général au voisinage d'une solution particulière. La dernière section du chapitre discute brièvement les mérites respectifs des deux modes de représentation de systèmes.

4.1 Le concept d'état

Nous avons vu dans le chapitre précédent que la représentation entrée-sortie des systèmes différentiels impose la condition de repos initial et est donc mal adaptée au traitement de conditions initiales non nulles. Une autre limitation réside dans la nécessité de conserver toute l'histoire passée de l'entrée pour déterminer le futur de la sortie à partir d'un instant donné, ainsi qu'exprimé par la formule de convolution d'un système causal en temps-continu

$$y(t) = \int_{-\infty}^t u(\tau) h(t - \tau) d\tau.$$

La représentation d'état des systèmes s'affranchit de ces limitations en ajoutant aux signaux d'entrée u et de sortie y un troisième signal x , appelé état du système. Le rôle de ce troisième signal est de contenir à tout instant l'information nécessaire pour pouvoir déterminer le futur de la sortie y sans connaître le passé de l'entrée u . En d'autres termes, le vecteur $x(t)$ paramétrise la mémoire que le système conserve du passé de l'entrée u sur l'intervalle de temps $(-\infty, t)$.

Physiquement, la notion d'état apparaît de manière naturelle : pour déterminer l'évolution future du courant i dans un circuit électrique, il n'est pas nécessaire de connaître la valeur de la tension v appliquée depuis l'origine des temps. On peut remplacer cette information par la valeur des charges dans les capacités et des flux dans les inductances à l'instant présent. Ces valeurs – qui déterminent l'énergie emmagasinée dans le circuit – constituent l'état du système : elles extraient de l'histoire passée du circuit l'information nécessaire pour la détermination de son futur. Cette propriété est clairement mise en évidence dans le circuit RC considéré dans l'exemple 1.2. La solution du système différentiel

$$v_c(t_0 + t) = e^{-\frac{t}{RC}} v_c(t_0) + \frac{1}{RC} \int_{t_0}^{t_0+t} e^{-\frac{t_0+t-\tau}{RC}} v_s(\tau) d\tau, \quad t \geq 0 \quad (4.1.1)$$

montre que la tension aux bornes de la capacité décrit l'état du circuit : la connaissance de $v_c(t_0)$ à un instant t_0 quelconque permet de reconstruire le futur du système à partir de l'instant t_0 sans connaître le passé du signal d'entrée v_s . Le fait que l'état du circuit RC soit décrit par une seule variable est étroitement lié au fait que le circuit RC contient un seul accumulateur d'énergie. Plus généralement, la dimension du vecteur d'état correspondra au nombre d'accumulateurs d'énergie présents dans le système.

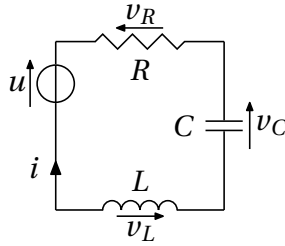


FIGURE 4.1 – Un circuit RLC.

Mettant à profit le lien entre la notion d'état et la notion d'accumulateur d'énergie, la modélisation physique conduit naturellement à une représentation d'état. Une illustration simple est fournie par l'exemple du circuit RLC représenté à la Figure 4.1. L'application des lois de Kirchhoff donne les équations

$$\begin{aligned} \dot{q} &= \frac{1}{L}\phi \\ \dot{\phi} &= -\frac{R}{L}\phi - \frac{1}{C}q + u \end{aligned} \quad (4.1.2)$$

Ces équations montrent que la variation de la charge q et du flux ϕ à l'instant t est déterminée par les valeurs de la charge, du flux, et de l'entrée *au même instant* t .

Si les valeurs de la charge et du flux sont connues à l'instant t_0 , ainsi que l'entrée $u(t)$, $t \geq t_0$, leur évolution future est univoquement déterminée pour $t \geq t_0$. Il en va de même pour la sortie qui peut être reconstituée à partir de la charge, du flux et de l'entrée. Si la sortie y est le courant i dans le circuit, on a

$$y = i = \frac{1}{L}\phi.$$

Si la sortie désigne la tension d'inductance, on a

$$y = v_L = u - \frac{R}{L}\phi - \frac{1}{C}q.$$

L'état du circuit RLC est donc un signal vectoriel $x = (q, \phi)$ qui satisfait l'équation différentielle

$$\dot{x} = \begin{bmatrix} 0 & \frac{1}{L} \\ -\frac{1}{C} & -\frac{R}{L} \end{bmatrix} x + \begin{bmatrix} 0 \\ 1 \end{bmatrix} u.$$

Cette équation est appelée *équation d'état* du système. L'*équation de sortie* est l'expression de la sortie en fonction des variables d'état et de l'entrée, par exemple

$$y = \begin{bmatrix} -\frac{1}{C} & -\frac{R}{L} \end{bmatrix} x + u$$

dans le cas où la sortie est la tension d'inductance.

L'exemple du circuit RLC est instructif à plusieurs égards :

- Il montre que dans un système physique, le vecteur d'état est naturellement associé aux accumulateurs d'énergie. L'état de chacun de ces composants à un instant donné constitue la mémoire du système, c'est-à-dire retient du passé du système ce qui est nécessaire à la détermination de son futur. Nous reviendrons sur cette notion de mémoire dans la section suivante.
- Il montre qu'une représentation d'état n'est pas unique : dans le Chapitre 2, nous avons développé un modèle d'état du circuit RLC en prenant comme variable d'état la charge q et le courant i . Deux modèles d'état peuvent donc caractériser le même système. Le passage d'une représentation à l'autre sera discuté dans la Section 4.4.

La détermination des équations d'état d'un système donné est donc un processus en trois étapes : le choix des *variables d'état* qui doit être tel que celles-ci résument l'histoire passée du système ; la dérivation de *l'équation d'état*, c'est-à-dire l'équation différentielle ou aux différences qui dicte le changement du vecteur d'état en fonction de sa valeur présente et de la valeur présente de l'entrée ; et finalement la dérivation de *l'équation de sortie*, qui exprime le signal de sortie comme une combinaison des valeurs présentes de l'état et de l'entrée.

Si un système admet une représentation d'état (4.0.2) de dimension n , cela signifie que l'influence de l'histoire passée de l'entrée – une fonction u définie sur l'intervalle $(-\infty, t_0)$ – sur le futur de la sortie peut être paramétrisée par un vecteur $x(t_0)$ de dimension n . Ceci n'est pas toujours possible, même pour un système linéaire et invariant. Par exemple, un retard pur $y(t) = u(t - T)$ est un système linéaire et invariant. Pour connaître l'évolution future de la sortie à partir de l'instant t_0 , il faut nécessairement connaître la fonction u sur l'intervalle $[t_0 - T, t_0)$. Mais l'ensemble des fonctions définies sur l'intervalle $[t_0 - T, t_0)$ ne peut pas être paramétrisé par un nombre fini de constantes (c'est un espace de dimension infinie). Ceci signifie que le système linéaire et invariant $y(t) = u(t - T)$ n'admet pas de représentation d'état de dimension finie.

Exercice 4.1. Montrer, qu'à l'inverse du cas continu, un retard pur discret $y[n] = u[n - n_0]$ admet une représentation d'état de dimension n_0 . $\triangleleft$

4.2 Modèle d'état d'un système différentiel ou aux différences

Un modèle entrée-sortie continu

$$\sum_{k=0}^N a_k \frac{d^k y(t)}{dt^k} = \sum_{k=0}^N b_k \frac{d^k u(t)}{dt^k} \quad (4.2.1)$$

admet toujours une représentation d'état si $a_N \neq 0$. Un modèle entrée-sortie discret

$$\sum_{k=0}^N a_k y[n-k] = \sum_{k=0}^N b_k u[n-k] \quad (4.2.2)$$

admet toujours une représentation d'état si $a_0 \neq 0$. Nous allons dériver une telle représentation de manière systématique et montrer que ce problème est étroitement relié à la *réalisation* d'une équation différentielle sous la forme d'un circuit analogique ou d'une équation aux différences sous la forme d'un circuit digital.

4.2.1 Cas continu

Considérons tout d'abord le système du premier ordre

$$\dot{y}(t) + a_0 y(t) = b_0 u(t) \quad (4.2.3)$$

Ce système peut être réalisé dans un circuit analogique au moyen de trois opérations de base : l'additionneur, le multiplicateur par un scalaire, et l'intégrateur. (Le choix d'un intégrateur plutôt qu'un différentiateur sera justifié au Chapitre 9 : nous verrons qu'un différentiateur est extrêmement sensible au bruit.) Comme illustré par les Figures 4.2, 4.3 et 4.4, ces trois blocs de base peuvent être réalisés physiquement au moyen d'amplificateurs, de résistances, et de capacités.

En associant un symbole à chacun de ces éléments de base, on peut construire un *bloc-diagramme* (ou *schéma bloc*) du système qui représente graphiquement l'équation (4.2.3).

En supposant que le système est initialement au repos ($y(0) = 0$), on « lit » le bloc-diagramme de la Figure 4.5 comme

$$y(t) = \int_0^t (b_0 u(\tau) - a_0 y(\tau)) d\tau, \quad t \geq 0. \quad (4.2.4)$$

C'est la représentation *entrée-sortie* du système, qui à chaque instant t utilise toute l'histoire passée de l'entrée $u(\tau)$, $0 \leq \tau < t$ pour déterminer l'évolution future de la sortie. En choisissant l'état $x = y$, on obtient l'équation

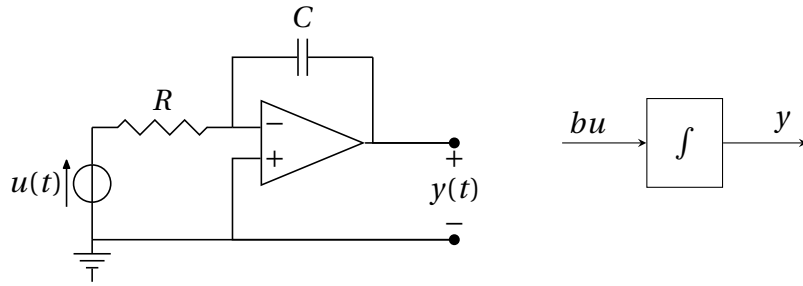


FIGURE 4.2 – Un intégrateur analogique $y(t) = b \int_{t_0}^t u(\tau) d\tau$ avec $b = -\frac{1}{RC}$ et sa représentation dans un bloc-diagramme.

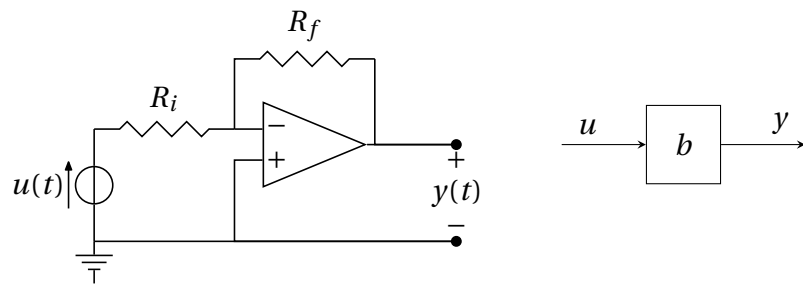


FIGURE 4.3 – Un multiplicateur analogique $y = bu$ avec $b = -\frac{R_f}{R_i}$ et sa représentation dans un bloc-diagramme.

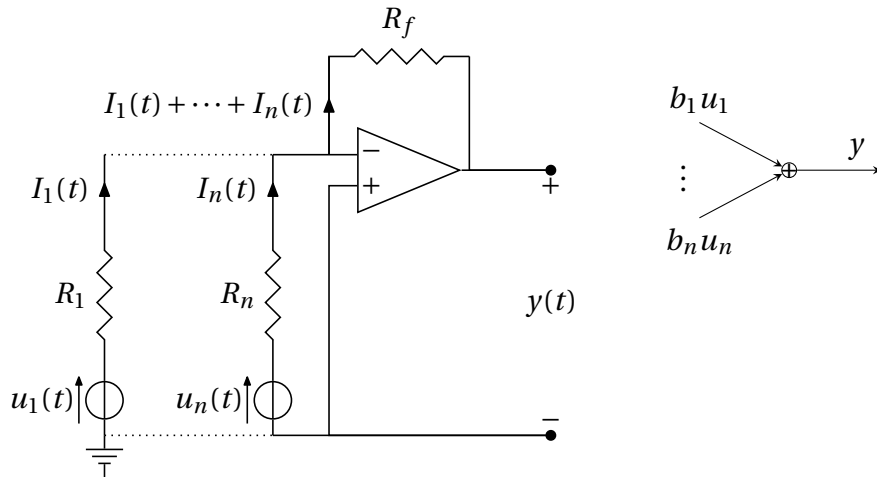


FIGURE 4.4 – Un additionneur analogique $y = \sum_{i=1}^n b_i u_i$ avec $b_i = -\frac{R_f}{R_i}$ et sa représentation dans un bloc-diagramme.

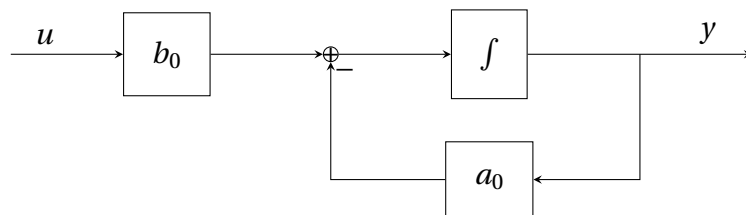


FIGURE 4.5 – Le bloc-diagramme du système (4.2.3).

d'état

$$\dot{x} = -a_0 x + b_0 u$$

et l'équation de sortie $y = x$. Cela correspond à réécrire l'équation (4.2.4) sous la forme

$$y(t) = y(t_0) + \int_{t_0}^t (b_0 u(\tau) - a_0 y(\tau)) d\tau$$

où, à l'instant t_0 , on a remplacé l'histoire passée de l'entrée u par la valeur de l'état du système.

Dans le bloc-diagramme de la Figure 4.5, le seul élément capable de stocker de l'énergie est l'intégrateur et il est donc naturel de choisir comme état du système la sortie de l'intégrateur (qui, dans le cas présent, coïncide avec la sortie du système).

On voit donc apparaître la connexion entre *la réalisation d'un système au moyen d'intégrateurs*, c'est-à-dire la description du système sous la forme d'un bloc-diagramme, et la construction d'une représentation d'état pour le système, en prenant comme choix pour le vecteur d'état toutes les sorties des intégrateurs utilisés. L'intégrateur est donc l'accumulateur d'énergie abstrait de tout modèle dynamique. Il peut représenter une capacité ou une inductance dans un circuit électrique (énergie électrique et magnétique), une position ou une vitesse dans un système mécanique (énergie potentielle et cinétique), une concentration ou une température dans un système thermodynamique (énergie chimique ou thermique).

La généralisation de notre exemple à l'équation d'ordre N

$$\sum_{k=0}^N a_k y^{(k)} = b_0 u, \quad a_N = 1, \quad (4.2.5)$$

est immédiate. En réécrivant l'équation sous la forme

$$y^{(N)} = - \sum_{k=0}^{N-1} a_k y^{(k)} + b_0 u$$

on obtient directement le bloc-diagramme de la Figure 4.6. Une représentation d'état est obtenue en choisissant comme état la sortie des N intégrateurs utilisés pour réaliser le système, ce qui correspond dans le cas présent à la sortie et ses $N - 1$ premières dérivées.

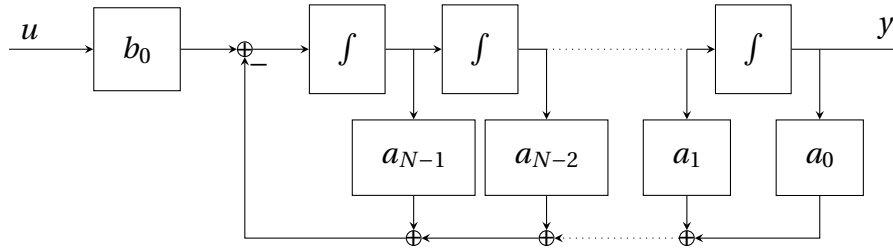


FIGURE 4.6 – Bloc-diagramme du système (4.2.5).

On obtient ainsi la représentation d'état

$$\begin{aligned}
 \dot{x}_1 &= x_2, \\
 \dot{x}_2 &= x_3, \\
 &\vdots \\
 \dot{x}_{n-1} &= x_n, \\
 \dot{x}_n &= -a_0 x_1 - \cdots - a_{N-1} x_n + b_0 u, \\
 y &= x_1.
 \end{aligned}$$

Pour traiter le cas général (4.2.1) où on pose $a_N = 1$ sans perte de généralité, on construit un signal intermédiaire v qui satisfait

$$\sum_{k=0}^N a_k v^{(k)} = u$$

et pour lequel on choisit la même réalisation que précédemment. La sortie y est déduite du signal $v(t)$ par la relation

$$y = \sum_{k=0}^N b_k v^{(k)}, \quad (4.2.6)$$

que l'on peut justifier formellement par la décomposition

$$P(D)y = Q(D)u \quad \Leftarrow \quad y = Q(D)v \text{ et } P(D)v = u.$$

La relation (4.2.6) est très facilement réalisée à partir du bloc-diagramme intermédiaire puisque les dérivées successives du signal v sont les sorties des différents intégrateurs. On obtient ainsi le bloc-diagramme complet de la Figure 4.7, réalisé au moyen de N intégrateurs.

On obtient la représentation d'état correspondante en choisissant comme état le signal v et ses $N - 1$ premières dérivées. On obtient ainsi la représenta-

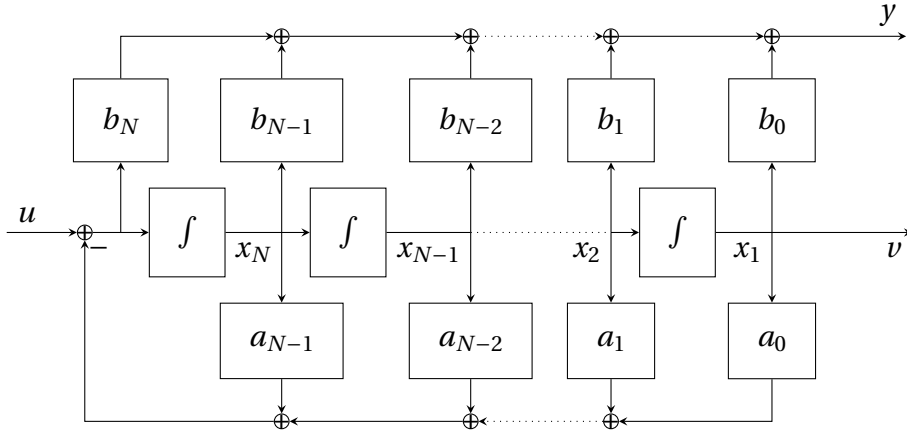


FIGURE 4.7 – Le bloc-diagramme du système différentiel (4.2.1).

tion d'état $\dot{x} = Ax + Bu$, $y = Cx + Du$ caractérisée par les matrices

$$A = \begin{bmatrix} 0 & 1 & 0 & \dots & 0 \\ \vdots & \ddots & \ddots & & 0 \\ \vdots & & \ddots & \ddots & 0 \\ 0 & \dots & 0 & 1 & \\ -a_0 & -a_1 & -a_2 & \dots & -a_{N-1} \end{bmatrix}, \quad B = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ 1 \end{bmatrix}, \quad (4.2.7)$$

$$C = \begin{bmatrix} b_0 - b_N a_0 & b_1 - b_N a_1 & \dots & b_{N-1} - b_N a_{N-1} \end{bmatrix}, \quad D = b_N.$$

On peut remarquer qu'il y a coïncidence entre l'ordre N du modèle entrée-sortie (4.2.1), le nombre d'intégrateurs nécessaires pour réaliser le système, et la dimension de l'équation d'état de la représentation.

La réalisation d'une équation aux différences (4.2.2) et l'obtention systématique d'une représentation d'état est complètement analogue au cas continu. Il suffit de remplacer l'intégrateur continu par un opérateur de décalage (retard) discret dans les blocs de base de définition du bloc-diagramme du système.

4.2.2 Cas discret

Considérons le système du premier ordre

$$y[n] + a_1 y[n-1] = b_0 u[n]. \quad (4.2.8)$$

Cette équation est représentée par le bloc-diagramme de la Figure 4.8. En choisissant comme état la sortie du décalage unitaire,

$$x[n] = y[n-1],$$

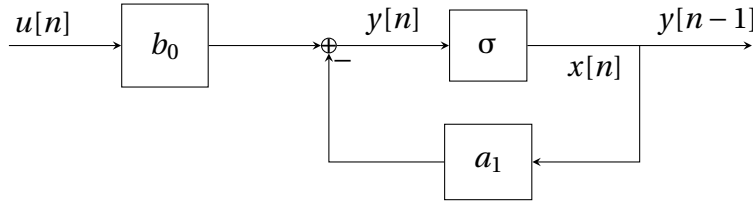


FIGURE 4.8 – Le bloc-diagramme du système (4.2.8), où σ désigne un retard pur unitaire.

on obtient l'équation d'état

$$x[n+1] = -a_1 x[n] + b_0 u[n].$$

La réalisation du cas général (4.2.2) et l'obtention systématique d'une représentation d'état est analogue au cas continu. Pour réaliser la relation entrée-sortie aux différences (4.2.2)

$$\sum_{k=0}^N a_k y[n-k] = \sum_{k=0}^N b_k u[n-k]$$

dans laquelle on pose $a_0 = 1$ sans perte de généralité, on construit un signal intermédiaire v qui satisfait

$$\sum_{k=0}^N a_k v[n-k] = u[n], \quad (4.2.9)$$

ou encore

$$v[n] = - \sum_{k=1}^N a_k v[n-k] + u[n].$$

La sortie y s'obtient alors par la relation

$$y[n] = \sum_{k=0}^N b_k v[n-k]. \quad (4.2.10)$$

Les relations (4.2.9) et (4.2.10) sont réalisées par le bloc diagramme de la Figure 4.9 dans lequel σ désigne un décalage unitaire. En choisissant comme état les sorties des décalages, c'est-à-dire

$$x[n] = (v[n-N], \dots, v[n-1]),$$

on obtient la représentation d'état $\sigma^{-1}x = Ax + Bu$, $y = Cx + Du$ caractérisée

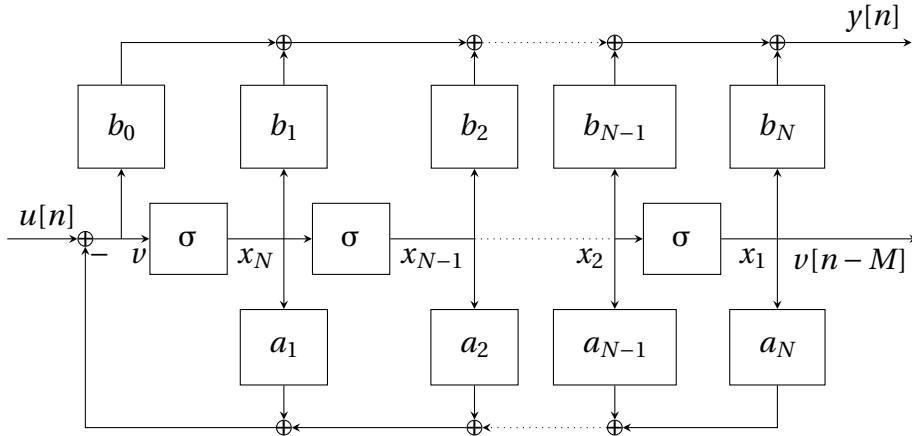


FIGURE 4.9 – Le bloc-diagramme qui réalise le système aux différences (4.2.2).

par les matrices

$$A = \begin{bmatrix} 0 & 1 & 0 & \dots & 0 \\ \vdots & \ddots & \ddots & & 0 \\ \vdots & & \ddots & \ddots & 0 \\ 0 & \dots & & 0 & 1 \\ -a_N & -a_{N-1} & -a_{N-2} & \dots & -a_1 \end{bmatrix}, \quad B = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ 1 \end{bmatrix}, \quad (4.2.11)$$

$$C = \begin{bmatrix} b_N - b_0 a_N & b_{N-1} - b_0 a_{N-1} & \dots & b_1 - b_0 a_1 \end{bmatrix}, \quad D = b_0.$$

4.3 Solution des équations d'état

Nous discutons à présent brièvement comment la solution des modèles d'état linéaires invariants peut être déterminée explicitement.

4.3.1 Modèles d'état discrets

La solution explicite de l'équation d'état

$$x[k+1] = Ax[k] + Bu[k], \quad x[0] = x_0,$$

s'obtient simplement par substitutions successives :

$$\begin{aligned} x[1] &= Ax[0] + Bu[0], \\ x[2] &= Ax[1] + Bu[1] = A(Ax[0] + Bu[0]) + Bu[1] \\ &= A^2x[0] + ABu[0] + Bu[1], \\ &\vdots \end{aligned}$$

$$x[n] = A^n x[0] + \sum_{k=0}^{n-1} A^{n-1-k} B u[k], \quad n > 0.$$

La matrice

$$A^n = \overbrace{AA \dots A}^n$$

est appelée *matrice de transition* du système homogène $x[k+1] = Ax[k]$. La solution du système homogène à deux instants différents n_1 et n_2 satisfait en effet

$$x[n_2] = A^{n_2-n_1} x[n_1].$$

La représentation entrée-sortie s'obtient en substituant la solution de l'équation d'état dans l'équation de sortie $y = Cx + Du$, ce qui donne

$$y[n] = CA^n x[0] + \sum_{k=0}^{n-1} CA^{n-1-k} B u[k] + Du[n]. \quad (4.3.1)$$

On peut en déduire la réponse impulsionnelle de l'opérateur LTI associé au modèle d'état en imposant le repos initial $x[0] = 0$ et en identifiant (4.3.1) au produit de convolution

$$y[n] = \sum_{k \in \mathbb{Z}} u[k] h[n-k]$$

pour obtenir

$$h[n] = CA^{n-1} B \mathbb{I}[n-1] + D \delta[n].$$

4.3.2 Modèles d'état continus

Intéressons-nous tout d'abord à la solution de l'équation homogène

$$\dot{x} = Ax, \quad x(0) = x_0. \quad (4.3.2)$$

Dans le cas d'une équation scalaire, $\dot{x} = ax$, $x \in \mathbb{R}$, nous avons vu que la solution est $x(t) = e^{at} x_0$. Utilisant la série de Taylor

$$e^{at} = 1 + (at) + \frac{(at)^2}{2} + \frac{(at)^3}{3!} + \dots$$

on définit l'exponentielle matricielle de la matrice At par la série

$$e^{At} = I + (At) + \frac{(At)^2}{2!} + \frac{(At)^3}{3!} + \dots$$

qui vérifie la propriété

$$\frac{d e^{At}}{dt} = \frac{d}{dt} \sum_{k=0}^{\infty} \frac{A^k t^k}{k!} = \sum_{k=1}^{\infty} k \frac{A^k t^{k-1}}{k!} = A \sum_{j=0}^{\infty} \frac{A^j t^j}{j!} = A e^{At}.$$

Il en découle que $x(t) = e^{At} x_0$ est solution de l'équation (4.3.2). Cette solution est unique car si x est une solution de (4.3.2), alors $z(t) = e^{-At} x(t)$ vérifie $\dot{z} = 0$ et donc $z(t) \equiv z(0) = x_0$.

L'exponentielle matricielle e^{At} est appelée matrice de transition de l'équation homogène $\dot{x} = Ax$. La solution à deux instants différents t_1 et t_2 satisfait en effet

$$x(t_2) = e^{A(t_2 - t_1)} x(t_1).$$

Pour calculer la solution de l'équation d'état

$$\dot{x} = Ax + Bu, \quad x(0) = x_0$$

on définit la variable auxiliaire $z(t) = e^{-At} x(t)$, qui vérifie

$$\dot{z}(t) = -A e^{-At} x(t) + e^{-At} (Ax(t) + Bu(t)) = e^{-At} Bu(t).$$

Cette équation s'intègre directement pour donner

$$z(t) = z(0) + \int_0^t e^{-As} Bu(s) ds.$$

En retournant à la variable $x(t) = e^{At} z(t)$, on obtient la solution de l'équation d'état

$$x(t) = e^{At} x_0 + \int_0^t e^{A(t-s)} Bu(s) ds. \quad (4.3.3)$$

Comme dans le cas discret, on peut en déduire la relation entrée-sortie

$$y(t) = C e^{At} x_0 + \int_0^t C e^{A(t-s)} Bu(s) ds + Du(t)$$

et la réponse impulsionnelle de l'opérateur LTI causal

$$h(t) = C e^{At} B \mathbb{I}(t) + D \delta(t). \quad (4.3.4)$$

4.3.3 Calcul de l'exponentielle matricielle

Si la formule (4.3.3) nous donne une expression explicite pour la solution de l'équation d'état, elle ne nous dit pas comment calculer l'exponentielle matricielle e^{At} . Une méthode de calcul consiste à utiliser la *forme de Jordan* de la matrice A et les propriétés suivantes de l'exponentielle matricielle.

Exercice 4.2. En appliquant la définition de l'exponentielle matricielle, établir les propriétés suivantes :

(i) Si A_1 et A_2 commutent, c'est-à-dire $A_1 A_2 = A_2 A_1$, alors

$$e^{A_1 + A_2} = e^{A_1} e^{A_2}.$$

(ii) Si A_1 et A_2 sont deux matrices carrées, alors

$$\exp \left(\begin{bmatrix} A_1 & 0 \\ 0 & A_2 \end{bmatrix} \right) = \begin{bmatrix} e^{A_1} & 0 \\ 0 & e^{A_2} \end{bmatrix}.$$

(iii) Si T est une matrice non singulière, alors

$$e^{T^{-1}AT} = T^{-1}e^A T.$$

(iv)

$$\exp \left(\begin{bmatrix} 0 & 1 & 0 & \dots & 0 \\ & \ddots & \ddots & \ddots & \vdots \\ & & \ddots & \ddots & 0 \\ & & & \ddots & 1 \\ & & & & 0 \end{bmatrix} t \right) = \begin{bmatrix} 1 & t & \frac{t^2}{2!} & \dots & \frac{t^{n-1}}{(n-1)!} \\ 0 & \ddots & \ddots & \ddots & \vdots \\ & \ddots & \ddots & \ddots & \frac{t^2}{2!} \\ & & \ddots & \ddots & t \\ & & & 0 & 1 \end{bmatrix}. \quad \triangleleft$$

On se rappellera du cours d'algèbre linéaire que toute matrice carrée peut être transformée en forme de Jordan par un changement de base, c'est-à-dire qu'il existe une matrice T non singulière telle que

$$T^{-1}AT = \begin{bmatrix} J_1 & & \\ & \ddots & \\ & & J_q \end{bmatrix}.$$

Les sous-matrices J_k sont appelées *blocs de Jordan* et sont définies comme suit. Supposons que A a q vecteurs propres indépendants $v_1, \dots, v_q$. A chaque vecteur propre correspond un bloc de Jordan de la forme

$$J_k = \begin{bmatrix} \lambda_k & 1 & & \\ 0 & \ddots & \ddots & \\ & \ddots & \ddots & 1 \\ & & 0 & \lambda_k \end{bmatrix}.$$

Le nombre de blocs de Jordan associé à une même valeur propre λ_k correspond à la multiplicité géométrique de λ_k . Dans le cas particulier où la matrice A possède n valeurs propres distinctes, la forme de Jordan est simplement la diagonale des valeurs propres.

Le calcul de l'exponentielle d'une matrice en forme de Jordan découle directement des propriétés énoncées dans l'Exercice 4.2. On a donc

$$e^{At} = Te^{Jt}T^{-1}, \quad e^{Jt} = \text{diag}(e^{J_1 t}, \dots, e^{J_q t})$$

et chaque bloc $e^{J_k t}$ est de la forme

$$e^{J_k t} = \begin{bmatrix} e^{\lambda_k t} & te^{\lambda_k t} & \frac{t^2}{2!}e^{\lambda_k t} & \dots & e^{\lambda_k t} \frac{t^{n-1}}{(n-1)!} \\ 0 & \ddots & \ddots & \ddots & \vdots \\ & \ddots & \ddots & \ddots & \frac{t^2}{2!}e^{\lambda_k t} \\ & & \ddots & \ddots & te^{\lambda_k t} \\ & & & 0 & e^{\lambda_k t} \end{bmatrix}.$$

Le calcul de l'exponentielle matricielle fait apparaître que la solution de l'équation $\dot{x} = Ax$ est toujours une combinaison linéaire de termes de la forme $\frac{t^j}{j!}e^{\lambda_k t}$ où λ_k est une valeur propre de la matrice A . En particulier, ce sont les parties réelles des différentes valeurs propres qui déterminent le caractère croissant ou décroissant des solutions.

4.4 Transformations d'état

Nous avons déjà observé dans le cas du circuit RLC qu'une représentation d'état n'est pas unique. Ceci ne doit pas surprendre si l'on conçoit l'état $x(t)$ comme un vecteur de l'espace vectoriel $\mathbb{R}^n$. De même qu'un vecteur peut s'exprimer dans différentes bases, tout changement de base donnera lieu à une représentation d'état différente. Un changement de base est une application linéaire dont la matrice est carrée, non singulière et constituée des n vecteurs de la nouvelle base exprimés dans l'ancienne base. Un changement de base caractérisé par $x = Tz$ donnera lieu à la représentation

$$\begin{aligned} \dot{z} &= T^{-1}ATz + T^{-1}Bu, \\ y &= CTz + Du. \end{aligned}$$

Si on désigne une représentation donnée par le quadruplet (A, B, C, D) , on dira donc que deux représentations d'état (A_1, B_1, C_1, D_1) et (A_2, B_2, C_2, D_2) sont équivalentes s'il existe une matrice T telle que

$$A_2 = T^{-1}A_1T, \quad B_2 = T^{-1}B_1, \quad C_2 = C_1T, \quad D_2 = D_1.$$

Exercice 4.3. Montrer que le modèle d'état du circuit RLC dérivé au Chapitre 2 est équivalent au modèle d'état dérivé au début de ce chapitre. <

Les transformations d'état (ou changements de base) sont un outil capital pour l'analyse des systèmes linéaires dans l'espace d'état. Différentes bases seront appropriées à l'analyse de différentes propriétés et la résolution de différents problèmes sera grandement facilitée par le choix d'une base adéquate. Par exemple, nous avons vu que le calcul explicite des solutions requiert la détermination de l'exponentielle matricielle e^{At} ou de la matrice A^n et que le calcul de celles-ci est particulièrement aisé lorsque la matrice A est en forme de Jordan. C'est la représentation modale du système, obtenue dans la base des vecteurs propres (généralisés) de la matrice A . En revanche, d'autres représentations d'état seront favorisées dans l'étude de propriétés comme la commandabilité ou l'observabilité.

4.5 Construction de modèles linéaires invariants

Les propriétés de *linéarité* et d'*invariance* découlent rarement de lois physiques. Le plus souvent, elles doivent être induites par le processus de modélisation. L'activité de modélisation consiste donc, pour une large part, à extraire d'une loi physique ou comportementale un modèle simplifié qui possède ces deux propriétés fortes, et à justifier l'approximation ainsi réalisée en montrant qu'elle a une valeur *locale*, c'est-à-dire qu'elle est valide dans l'échelle de temps et d'espace du phénomène étudié. L'analyse sert en partie à confirmer ou infirmer a posteriori des hypothèses introduites a priori et il y a donc très souvent un processus itératif entre les phases d'analyse et de modélisation.

Il importe donc, non pas de rejeter les modèles linéaires invariants, car ceux-ci donnent de très bons résultats lorsqu'ils sont utilisés à bon escient, mais de garder à l'esprit les approximations sous-jacentes au choix du modèle.

4.5.1 Séparation espace-temps

Une première contrainte sous-jacente aux modèles d'état introduits au Chapitre 2 réside dans le choix de modèles à paramètres *localisés* plutôt que *distribués*. Dans l'établissement des équations d'un circuit électrique, nous utilisons des lois constitutives pour les différents composants du circuit, comme par exemple la loi d'Ohm. En procédant de la sorte, nous supposons implicitement que le courant dans une résistance est le même en chaque point de cette résistance. En réalité, les signaux électriques ne se propagent pas instantanément dans le système. Ce sont des ondes électromagnétiques qui se propagent à une certaine vitesse. Ainsi, un courant électrique n'est pas seulement fonction du temps mais aussi fonction de l'espace. Ce qui justifie une hypothèse de courant constant en tout point d'une résistance, c'est que

les variations du courant dans le temps sont lentes par rapport au temps requis pour la propagation des ondes. En d'autres termes, la longueur d'onde des signaux étudiés est très grande par rapport à la dimension des composants du circuits. Cette séparation d'échelle espace/temps est à la base des modèles localisés, qui donnent lieu à des équations différentielles. Lorsque la séparation d'échelle espace/temps n'est plus valable, il faut recourir à des modèles à paramètres distribués, qui donnent lieu à des équations aux dérivées partielles. C'est par exemple le cas lorsque l'on étudie la transmission de courant électrique sur de très longues distances.

Les équations différentielles qui résultent d'un modèle à paramètres localisés sont en général couplées, non linéaires, et dépendent du temps. Elles peuvent souvent être réarrangées sous la forme de n équations du premier ordre

$$\begin{aligned}\dot{x}_1(t) &= f_1(x_1(t), x_2(t), \dots, x_n(t), u_1(t), \dots, u_m(t), t) \\ \dot{x}_2(t) &= f_2(x_1(t), x_2(t), \dots, x_n(t), u_1(t), \dots, u_m(t), t) \\ &\vdots \\ \dot{x}_n(t) &= f_n(x_1(t), x_2(t), \dots, x_n(t), u_1(t), \dots, u_m(t), t)\end{aligned}\tag{4.5.1}$$

que l'on représente sous forme vectorielle par

$$\dot{x}(t) = f(x(t), u(t), t)\tag{4.5.2}$$

avec $x = (x_1, \dots, x_n)$ et $u = (u_1, \dots, u_m)$.

4.5.2 Localisation temporelle

Un phénomène étudié a généralement un temps caractéristique. Dans une étape de modélisation, il convient donc de négliger ce qui est très rapide ou très lent par rapport au temps caractéristique du phénomène étudié.

Cette approximation temporelle est à la base des modèles *invariants*. Par exemple, les constantes R et C du circuit RC peuvent dépendre de la température et donc se modifier dans le temps. Analyser un modèle invariant revient à supposer que la variation temporelle de ces constantes est lente par rapport au temps caractéristique du circuit. Pour une équation différentielle générale (4.5.2), l'invariance est obtenue en éliminant la dépendance explicite des équations par rapport au temps, c'est-à-dire lorsque les équations peuvent se mettre sous la forme

$$\dot{x}(t) = f(x(t), u(t)).\tag{4.5.3}$$

La justification mathématique d'un modèle invariant lorsque la dépendance temporelle de l'équation (4.5.2) est suffisamment rapide ou lente fait l'objet de la théorie des perturbations régulières et de la *moyennisation* (*averaging*).

Négliger ce qui est rapide dans un système dynamique peut aussi justifier

une réduction de la dimension du système. Par exemple, nous avons vu dans l'exemple 1.2 que la variation de la tension de sortie v_c dans un circuit RC en fonction de la tension de source n'est pas instantanée à cause du chargement de la capacité. Le circuit RC définit donc un système dynamique du premier ordre entre l'entrée $u(= v_s)$ et la sortie $y(= v_c)$. Cependant, si la constante C est « petite », le transitoire devient négligeable. À la limite, quand C tend vers zéro, on obtient la relation statique $v_c = v_s$. La théorie qui justifie ce type d'approximations d'un point de vue mathématique et en quantifie la validité s'appelle la théorie des perturbations singulières.

4.5.3 Linéarisation

L'exemple du bras de robot a montré que l'obtention d'un modèle linéaire nécessite généralement de localiser le phénomène étudié autour d'une solution particulière, le modèle linéaire constituant alors une bonne approximation pour étudier les petites variations. De même, la caractéristique linéaire $y = \frac{1}{R}u$ d'une résistance est une approximation locale, c'est-à-dire qu'elle n'est valide que pour une certaine plage de courants au travers de la résistance. Lorsque la tension devient trop importante, le courant cesse de croître linéairement par exemple en raison d'un effet de saturation. Il en va de même pour la loi linéaire d'un ressort. Si la force appliquée sur le ressort devient trop importante, le ressort finit par subir des déformations permanentes et l'écartement cesse de croître linéairement.

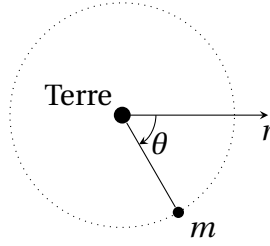
La linéarisation d'une équation différentielle générale (4.5.2) s'effectue autour d'une solution particulière (x^*, u^*) . Notons que cette solution particulière n'est pas nécessairement une solution d'équilibre, c'est-à-dire une solution caractérisée par des signaux constants. Une solution proche de la solution (x^*, u^*) peut s'écrire sous la forme $x = x^* + \bar{x}$, $u = u^* + \bar{u}$, où $\bar{x}$ et $\bar{u}$ sont de « petites » variations. Pour établir l'équation qui régit le système dans les variables $\bar{x}$ et $\bar{u}$, on développe le membre de droite de l'équation différentielle en série de Taylor (en supposant que f est différentiable), ce qui donne

$$f(x^*(t) + \bar{x}(t), u^*(t) + \bar{u}(t), t) = f(x^*(t), u^*(t), t) + A(t)\bar{x}(t) + B(t)\bar{u}(t) + \mathcal{O}(\|(\bar{x}(t), \bar{u}(t))\|^2) \quad (4.5.4)$$

où les éléments des matrices $A(t)$ et $B(t)$ sont donnés par

$$a_{ij}(t) = \frac{\partial f_i}{\partial x_j}(x^*(t), u^*(t), t) \quad \text{et} \quad b_{ij}(t) = \frac{\partial f_i}{\partial u_j}(x^*(t), u^*(t), t).$$

En réécrivant $\dot{x}(t) = f(x(t), u(t), t)$ dans les nouvelles variables et en utili-

FIGURE 4.10 – Modèle plan d'un satellite avec masse m .

sant (4.5.4), on obtient que la variable d'écart $\bar{x}$ satisfait l'équation

$$\dot{\bar{x}}(t) = A(t)\bar{x}(t) + B(t)\bar{u}(t) + \mathcal{O}(\|(\bar{x}(t), \bar{u}(t))\|^2). \quad (4.5.5)$$

Si l'on néglige les termes d'ordre supérieur, on obtient l'équation linéaire

$$\dot{z}(t) = A(t)z(t) + B(t)v(t), \quad (4.5.6)$$

qui est appelée équation linéarisée ou équation variationnelle autour de la solution (x^*, u^*) . C'est une représentation d'état linéaire qui décrit les petites variations du système (4.5.2) autour de la solution (x^*, u^*) . Si les matrices A et B ne dépendent pas du temps, le linéarisé est également invariant.

La linéarisation de l'équation de sortie $y(t) = h(x(t), u(t), t)$ obéit au même principe :

$$\begin{aligned} y^*(t) + \bar{y}(t) &= h(x^*(t) + \bar{x}(t), u^*(t) + \bar{u}(t), t) \\ &\approx h(x^*(t), u^*(t), t) + C(t)\bar{x}(t) + D(t)\bar{u}(t), \end{aligned}$$

où les éléments des matrices $C(t)$ et $D(t)$ sont donnés par

$$c_{ij}(t) = \frac{\partial h_i}{\partial x_j}(x^*(t), u^*(t), t) \quad \text{et} \quad d_{ij}(t) = \frac{\partial h_i}{\partial u_j}(x^*(t), u^*(t), t).$$

Exemple 4.4 (Modèle de satellite). On modélise le mouvement plan d'un satellite par le mouvement d'une masse ponctuelle soumise à l'action du champ gravitationnel terrestre, inversement proportionnelle au carré de la distance radiale à la terre (cf. Figure 4.10).

Si l'on suppose que le satellite peut être propulsé dans la direction radiale par une force u_1 et dans la direction tangentielle par une force u_2 , on obtient, en coordonnées polaires (r, θ) , les équations du mouvement

$$\begin{aligned} m\ddot{r} &= mr\dot{\theta}^2 - \frac{mk}{r^2} + u_1, \\ mr\ddot{\theta} &= -2m\dot{r}\dot{\theta} + u_2. \end{aligned} \quad (4.5.7)$$

Lorsque $u_1 = u_2 = 0$, les équations admettent la solution

$$r(t) = \sigma, \quad \theta(t) = \omega t, \quad \text{avec } \sigma, \omega \in \mathbb{R}^+ \text{ telles que } \sigma^3 \omega^2 = k,$$

qui correspond à une orbite circulaire. En définissant les variables d'état $x_1 = r - \sigma$, $x_2 = \dot{r}$, $x_3 = \sigma(\theta - \omega t)$, $x_4 = \sigma(\dot{\theta} - \omega)$, et en normalisant σ à 1, on obtient le système linéarisé

$$\begin{bmatrix} \dot{x}_1 \\ \dot{x}_2 \\ \dot{x}_3 \\ \dot{x}_4 \end{bmatrix} = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 3\omega^2 & 0 & 0 & 2\omega \\ 0 & 0 & 0 & 1 \\ 0 & -2\omega & 0 & 0 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{bmatrix} + \frac{1}{m} \begin{bmatrix} 0 & 0 \\ 1 & 0 \\ 0 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} u_1 \\ u_2 \end{bmatrix}. \quad (4.5.8)$$

On notera que le système linéarisé est invariant alors que la linéarisation n'a pas été effectuée autour d'une solution d'équilibre. Pour garantir l'invariance du linéarisé en toute généralité, il faut que le modèle de départ soit lui-même invariant et que la linéarisation soit effectuée autour d'un point d'équilibre. <

4.6 Représentation interne ou externe ?

Faut-il privilégier l'un ou l'autre mode de représentation des systèmes linéaires? Les avis ont divergé sur la question suivant le contexte historique et technologique mais les mérites respectifs des deux approches sont aujourd'hui unanimement reconnus.

L'approche *entrée-sortie* a émergé de la théorie des circuits. Elle privilégie donc la vision d'un système comme fonction de transfert, sorte de loi constitutive du composant décrit. Son grand atout est de permettre une analyse fréquentielle du système étudié. En revanche, elle est mal adaptée au traitement des conditions initiales et, dans une certaine mesure, à l'analyse des phénomènes transitoires.

L'approche *espace d'état* a émergé de la mécanique. Elle privilégie donc la vision d'un système comme loi d'évolution, permettant la prédiction du futur à partir de la connaissance de l'état présent du système. Elle permet d'étudier les systèmes instables et fournit une description plus complète du système étudié : le comportement entrée-sortie peut en effet être aveugle au comportement de certaines variables internes.

En dehors des circuits électriques, la modélisation des systèmes à partir des lois physiques qui régissent leur comportement débouche plus naturellement sur les modèles d'état. Mais les modèles d'état sont non robustes. Une petite variation de paramètres peut détruire des propriétés structurelles du modèle étudié telles que la commandabilité ou l'observabilité. Face à

un système complexe et incertain, on préférera souvent établir un modèle entrée-sortie sur la base d'expériences.

Les techniques graphiques d'analyse qui ont constitué l'essentiel de l'analyse classique des systèmes se sont développées dans le cadre de l'approche entrée-sortie. Le passage d'outils graphiques aux outils numériques a largement favorisé le développement de l'analyse des systèmes dans l'espace d'état.

Chapitre 5

Décomposition fréquentielle des signaux

L'étude des systèmes linéaires a montré l'intérêt d'exprimer un signal quelconque comme combinaison linéaire de signaux *de base*. Un signal en temps discret $x[\cdot]$ a été représenté tantôt comme une combinaison linéaire d'impulsions décalées (représentation temporelle) et tantôt comme une combinaison linéaire de signaux harmoniques (représentation fréquentielle). Du point de vue de l'algèbre linéaire, un même élément de l'espace de signaux considéré se voit exprimé dans deux bases différentes. Autrement dit, le passage du monde temporel au monde fréquentiel évoque un *changement de base* dans un espace vectoriel donné.

Ce chapitre établit le cadre rigoureux de ce point de vue algébrique. La série de Fourier d'un signal est l'expression de ce signal dans la base harmonique. Cette base a la double propriété d'être *orthogonale* et *spectrale* pour les opérateurs de convolution. Cette double propriété est au cœur des innombrables avantages de l'analyse fréquentielle des signaux et systèmes, dont l'existence d'algorithmes efficaces de calcul, tels la célèbre FFT (*Fast Fourier Transform*).

Contrairement aux chapitres précédents, ce chapitre traite des signaux définis sur un intervalle *fini*. L'expression de ces signaux dans la base harmonique donne lieu aux séries de Fourier. Par signal défini sur un intervalle fini, on entend un signal entièrement déterminé par les valeurs qu'il prend sur un intervalle donné : l'ensemble $[0..N) = \{0, 1, \dots, N-1\}$ dans le cas discret et l'intervalle $[0, T)$ dans le cas continu. Un tel signal désignera selon l'usage un signal qui n'a pas d'existence en dehors de cet intervalle ou un signal *périodique* défini sur l'entiereté de l'ensemble $\mathbb{R}$ ou $\mathbb{Z}$ mais entièrement caractérisé par sa définition sur une période. Dans ce sens, la théorie des séries de Fourier s'applique indifféremment aux signaux de durée finie ou aux signaux périodiques. En revanche, elle ne s'applique pas aux signaux apériodiques de durée infinie. Le chapitre suivant étend la théorie à de tels signaux, en montrant que les séries de Fourier conduisent alors aux transformées de Fourier vues au Chapitre 6.

5.1 Signaux périodiques

5.1.1 Signaux périodiques continus

Un signal périodique en temps-continu est un signal $x(\cdot)$ défini sur la droite réelle et qui vérifie la propriété

$$\forall t \in \mathbb{R} : x(t + T) = x(t)$$

pour un nombre $T > 0$ appelé *période* du signal. Si T est le plus petit nombre qui vérifie la propriété, T est appelé période *fondamentale* du signal.

Le signal sinusoïdal $x(t) = \sin(\omega_0 t)$ est un signal périodique de période $T = \frac{2\pi}{\omega_0}$. Le nombre ω_0 est appelé *fréquence* du signal *exprimée en radians par seconde* (rad/s) (ou pulsation) et le nombre $f = \frac{\omega_0}{2\pi} = \frac{1}{T}$ est appelé *fréquence* du signal exprimée en cycles par seconde (1/s). Un cycle par seconde est aussi appelé un Hertz (Hz).

Les signaux physiques ondulatoires peuvent avoir des fréquences très différentes. Les sons sont audibles par l'oreille humaine dans la bande de fréquence 20 Hz–20 kHz (kiloHertz). Les ondes électromagnétiques s'étendent dans la bande de fréquences 1 Hz to 10^{25} Hz (rayons cosmiques). Les ondes de lumière visible se situent autour de 10^{15} Hz.

Le traitement mathématique des signaux sinusoïdaux est grandement facilité par l'utilisation de l'exponentielle complexe. Tous les nombres complexes de module unité admettent la représentation polaire $e^{j\theta}$ pour un certain $\theta \in [0, 2\pi)$ et satisfont l'égalité

$$e^{j\theta} = \cos \theta + j \sin \theta.$$

Par conséquent le signal $x(t) = \sin(\omega_0 t) = \sin(2\pi f_0 t)$ de fréquence f_0 est la partie imaginaire du signal complexe $e^{j\omega_0 t}$. Ce signal peut être identifié à un vecteur plan de norme unité tournant dans le sens anti-horlogique à la vitesse angulaire ω_0 .

5.1.2 Signaux périodiques discrets

Un signal périodique en temps-discret est un signal $x[\cdot]$ défini sur l'ensemble des entiers $\mathbb{Z}$ et qui vérifie la propriété

$$\forall n \in \mathbb{Z} : x[n + N] = x[n]$$

pour un nombre $N \in \mathbb{N}_0$ appelé *période* du signal. Si N est le plus petit nombre qui vérifie la propriété, N est appelé période *fondamentale* du signal.

Contrairement aux signaux continus, le signal sinusoïdal $x[n] = \sin(\omega_0 n)$

n'est pas nécessairement un signal périodique : la propriété

$$x[n + N] = \sin(\omega_0(n + N)) = \sin(\omega_0 n) = x[n]$$

ne peut être vérifiée que si le nombre $\omega_0 N$ est un multiple de 2π , c'est-à-dire si il existe un entier k tel que $\omega_0 N = 2\pi k$, c'est-à-dire

$$f_0 = \frac{\omega_0}{2\pi} = \frac{k}{N}.$$

En d'autres termes, la fréquence d'un signal discret doit être rationnelle pour que le signal soit périodique. L'unité de fréquence en discret est un nombre de cycles par échantillon.

Tout comme en temps-continu, il est utile pour le traitement mathématique des signaux de concevoir le signal sinusoïdal $x[n] = \sin(\omega_0 n)$ comme la partie imaginaire du signal complexe $e^{j\omega_0 n}$. Ce signal peut être identifié à un vecteur plan de norme unité pivotant dans le sens anti-horlogique d'un incrément ω_0 à chaque pas de temps. Il est utile de noter qu'un incrément de $2k\pi$ pour k entier revient à ne pas bouger sur le cercle unité. Par conséquent, la fréquence d'un signal en temps-discret exprimée en radians par pas de temps est définie *modulo* 2π . Ceci explique pourquoi on se limitera à un intervalle de fréquences de longueur 2π dans l'étude fréquentielle d'un signal discret. Les hautes fréquences en discret correspondent aux valeurs de ω_0 proches de $\pm\pi$ radians par pas de temps, tandis que les basses fréquences en discret correspondent, comme en continu, aux valeurs de ω proches de 0 radians par pas de temps. Pour la valeur limite de π (correspondant à $\pm\infty$ en continu), on a le signal $e^{j\pi n}$ qui a une fréquence maximale dans le sens où elle change de signe à chaque pas de temps.

5.2 Série de Fourier en temps discret

5.2.1 Bases de signaux dans l'espace $\ell_2[0..N)$

Un signal périodique N -périodique en temps-discret est défini sur l'ensemble des entiers $\mathbb{Z}$. Néanmoins, il est entièrement caractérisé par les N valeurs qu'il prend sur l'intervalle $[0..N)$. Dans ce sens, il peut être identifié à un signal *fini* de l'espace vectoriel

$$\ell_2[0..N) = \{x \in ([0..N) \rightarrow \mathbb{C}) : \|x\|_2 < +\infty\}.$$

L'espace $\ell_2[0..N)$ est un espace vectoriel de dimension finie N . Il peut être mis en correspondance avec l'espace euclidien $\mathbb{C}^N$: les N valeurs successives du signal $x[\cdot]$ constituent les N composantes d'un vecteur de $\mathbb{C}^N$.

Une famille de vecteurs $v_1, v_2, \dots, v_N$ d'un espace vectoriel X constitue

une *base* de l'espace si à tout élément $x \in X$ correspond une suite unique de scalaires $\alpha_1, \alpha_2, \dots, \alpha_N$ telle que

$$x = \sum_{i=1}^N \alpha_i v_i.$$

Les scalaires α_i sont appelés coordonnées du vecteur x dans la base $v_1, \dots, v_N$.

Pour constituer une base, il est nécessaire et suffisant que la famille $v_1, \dots, v_N$ soit *linéairement indépendante* (aucun vecteur de la base ne peut s'exprimer comme combinaison linéaire d'autres vecteurs de la base) et *génératrice* (tout vecteur de l'espace peut être exprimé comme combinaison linéaire des vecteurs v_i). La base canonique de $\mathbb{C}^N$ est constituée des N vecteurs

$$e_1 = \begin{bmatrix} 1 & 0 & \dots & 0 \end{bmatrix}^T, \quad e_2 = \begin{bmatrix} 0 & 1 & \dots & 0 \end{bmatrix}^T, \quad \dots, \quad e_N = \begin{bmatrix} 0 & 0 & \dots & 1 \end{bmatrix}^T.$$

De même, la base canonique de l'espace $\ell_2[0..N)$ est constituée des N signaux de base

$$\delta[n - (i - 1)], \quad n \in [0..N).$$

La dimension d'un espace vectoriel est par définition le nombre de vecteurs d'une base de l'espace.

5.2.2 Bases orthogonales

Lorsque l'espace vectoriel X est muni d'un produit scalaire $\langle \cdot, \cdot \rangle$, deux éléments de l'espace sont dits *orthogonaux* lorsque leur produit scalaire est nul. Une base constituée de vecteurs orthogonaux deux à deux est appelée base orthogonale. Elle est dite orthonormée si tous les vecteurs de base v_i vérifient en outre $\langle v_i, v_i \rangle = 1$ (c'est-à-dire sont de norme unitaire pour la norme induite $\|\cdot\| = \sqrt{\langle \cdot, \cdot \rangle}$). Les bases canoniques de $\mathbb{C}^N$ et $\ell_2[0..N)$ sont orthonormées pour le produit scalaire usuel.

Les coordonnées d'un vecteur x dans une base orthogonale $(v_1, \dots, v_N)$ sont données par la formule

$$\hat{x}_i = \frac{\langle x, v_i \rangle}{\langle v_i, v_i \rangle} \quad (5.2.1)$$

obtenue à partir de l'expression $x = \sum_{k=1}^N \hat{x}_k v_k$ en effectuant le produit scalaire des deux membres avec le vecteur v_i .

Base harmonique. La base harmonique de $\ell_2[0..N)$ est définie par les N vecteurs

$$v_k[n] = \frac{1}{N} e^{jk \frac{2\pi}{N} n}, \quad k = 0, \dots, N-1. \quad (5.2.2)$$

L'orthogonalité de la famille de vecteurs (5.2.2) résulte de

$$\langle v_k, v_l \rangle = \sum_{n=0}^{N-1} v_k[n] \overline{v_l[n]} = \frac{1}{N^2} \sum_{n=0}^{N-1} e^{j(k-l)\frac{2\pi}{N}n} = \begin{cases} \frac{1}{N}, & k = l, \\ 0, & k \neq l. \end{cases}$$

(On obtiendrait une base orthonormée en modifiant le facteur de pondération $1/N$ en $1/\sqrt{N}$ dans (5.2.2).)

Une famille orthogonale est nécessairement linéairement indépendante. La famille (5.2.2) constitue donc une famille de N vecteurs indépendants dans un espace de dimension N , c'est-à-dire une base.

Les coordonnées $\hat{x}_k$ d'un signal de $\ell_2[0..N)$ dans la base harmonique sont appelés *coefficients de Fourier* du signal x . La formule (5.2.1) donne

$$\hat{x}[k] = \hat{x}_k = \frac{\langle x, v_k \rangle}{\langle v_k, v_k \rangle} = \sum_{n=0}^{N-1} x[n] e^{-jk\frac{2\pi}{N}n}. \quad (5.2.3)$$

La représentation temporelle d'un signal x par N valeurs consécutives $x[0]$, $x[1]$, ..., $x[N-1]$ correspond donc aux coordonnées du signal dans la base canonique $(e_1, \dots, e_N)$, alors que sa représentation fréquentielle par ses N coefficients de Fourier $\hat{x}_0, \hat{x}_1, \dots, \hat{x}_{N-1}$ correspond aux coordonnées du signal dans la base harmonique. Le passage de la base temporelle à la base harmonique constitue une première transformée de Fourier, la transformée *discrète-discrète* :

$$x \xleftrightarrow{\mathcal{F}_{DD}} \hat{x} \quad \text{ou} \quad [0..N) \ni n \mapsto x[n] \xleftrightarrow{\mathcal{F}_{DD}} [0..N) \ni k \mapsto \hat{x}[k]. \quad (5.2.4)$$

Du point de vue de l'algèbre linéaire, cette transformée est un changement de base.

5.2.3 Série de Fourier d'un signal périodique en temps-discret

L'implication pratique du changement de base discuté dans la section précédente est la suivante : tout signal N -périodique peut être exprimé comme une superposition de N signaux harmoniques, c'est-à-dire tout signal N -périodique x , donc avec domaine $\mathbb{Z}$, admet la décomposition fréquentielle

$$x[n] = \frac{1}{N} \sum_{k=0}^{N-1} \hat{x}_k e^{jk\omega_0 n}, \quad \omega_0 = \frac{2\pi}{N}. \quad (5.2.5)$$

Les N coefficients (complexes) de Fourier $\hat{x}_k$ du signal sont donnés par l'expression

$$\hat{x}_k = \sum_{n=0}^{N-1} x[n] e^{-jk\omega_0 n}. \quad (5.2.6)$$

Si x est un signal réel, son expression est donnée par la partie réelle de

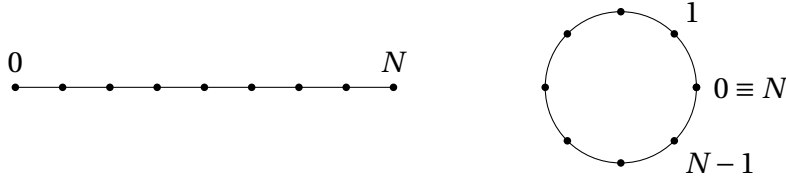


FIGURE 5.1 – Illustration du décalage cyclique.

l'expression (5.2.5). En regroupant les termes adéquats, on obtient l'expression

$$x[n] = A_0 + \frac{1}{N} \sum_{k=1}^K A_k \cos(k\omega_0 n + \phi_k)$$

où K désigne la partie entière de $N/2$. Les coefficients A_k et ϕ_k peuvent être déduits des coefficients $\hat{x}_k$.

5.2.4 La base harmonique est spectrale pour les opérateurs linéaires invariants

Considérons un opérateur H défini sur l'espace vectoriel $\ell_2[0..N)$. On suppose que H est une application *linéaire* et *invariante*. Dans le contexte des signaux *finis*, la propriété d'invariance s'exprime par rapport à un décalage *modulo* N . Si on représente l'axe temporel fini $[0..N)$ par N angles équidistants de ω_0 sur le cercle, on peut considérer le décalage modulo N comme un décalage *cyclique* sur le cercle (Figure 5.1).

Puisque l'espace vectoriel $\ell_2[0..N)$ est de dimension finie, l'image $y = H(u)$ d'un signal quelconque u exprimée dans une base particulière peut être obtenue par une expression matricielle. Par exemple, dans la base canonique, on aura l'expression matricielle

$$y_e = H_e u_e \quad (5.2.7)$$

où y_e désigne le vecteur comprenant les N valeurs successives du signal $y = H(u)$ et où H_e désigne la matrice de l'opérateur dans la base canonique. La première colonne de la matrice H_e contient les valeurs du signal image du premier vecteur de base, c'est-à-dire $H(e_1)$. La deuxième colonne contient les valeurs du signal image du deuxième vecteur de base, c'est-à-dire $H(e_2)$. Par la propriété d'invariance, le vecteur $H(e_2)$ est un décalage cyclique de $H(e_1)$ car e_2 est un décalage cyclique de e_1 . Il en va de même pour toutes les autres colonnes de la matrice H_e . La matrice est par conséquent une matrice *circulante*, entièrement déterminée par sa première colonne. L'expression (5.2.7)

prend donc la forme

$$\begin{bmatrix} y[0] \\ y[1] \\ \vdots \\ y[N-1] \end{bmatrix} = \begin{bmatrix} h[0] & h[N-1] & \dots & h[1] \\ h[1] & h[0] & \dots & h[2] \\ \vdots & \vdots & \ddots & \vdots \\ h[N-1] & h[N-2] & \dots & h[0] \end{bmatrix} \begin{bmatrix} u[0] \\ u[1] \\ \vdots \\ u[N-1] \end{bmatrix}. \quad (5.2.8)$$

Elle revient à définir le signal y par la formule

$$y[n] = \sum_{m=0}^{N-1} u[m] h[(n-m) \bmod N]$$

qui correspond à une opération de *convolution cyclique* $y = u \odot h$. La convolution cyclique peut être considérée comme l'adaptation de l'opération de convolution discrète définie au Chapitre 3 à l'espace de dimension finie $\ell_2[0..N)$.

Le produit de la matrice H_e par un vecteur v_k de la base harmonique fait apparaître une propriété fondamentale : En notant $W = e^{j\omega_0}$, on a l'abréviation $v_k[n] = \frac{1}{N} W^{nk}$ et

$$\begin{bmatrix} h[0] & h[N-1] & \dots & h[1] \\ h[1] & h[0] & \dots & h[2] \\ \vdots & \vdots & \ddots & \vdots \\ h[N-1] & h[N-2] & \dots & h[0] \end{bmatrix} \begin{bmatrix} W^{0k} \\ W^{1k} \\ \vdots \\ W^{(N-1)k} \end{bmatrix} = \left(\sum_{i=0}^{N-1} h[i] W^{-ik} \right) \begin{bmatrix} W^{0k} \\ W^{1k} \\ \vdots \\ W^{(N-1)k} \end{bmatrix}.$$

L'égalité matricielle se réécrit

$$h \odot v_k = \hat{h}_k v_k$$

et montre la propriété suivante : *Les n vecteurs propres de la matrice H_e sont les n vecteurs de la base harmonique et les n valeurs propres correspondantes sont les coefficients de Fourier $\hat{h}_k$.*

La matrice d'un opérateur dans la base des vecteurs propres de cet opérateur est la diagonale des valeurs propres. Dans la base harmonique, le produit matriciel (5.2.8) prend donc la forme

$$\begin{bmatrix} \hat{y}_0 \\ \hat{y}_1 \\ \vdots \\ \hat{y}_{N-1} \end{bmatrix} = \begin{bmatrix} \hat{h}_0 & 0 & \dots & 0 \\ 0 & \hat{h}_1 & \dots & 0 \\ \vdots & \vdots & \dots & \vdots \\ 0 & 0 & \dots & \hat{h}_{N-1} \end{bmatrix} \begin{bmatrix} \hat{u}_0 \\ \hat{u}_1 \\ \vdots \\ \hat{u}_{N-1} \end{bmatrix}. \quad (5.2.9)$$

Une base constituée de vecteurs propres est dite *spectrale*. La base harmonique est donc non seulement orthogonale mais aussi spectrale pour

les opérateurs linéaires invariants. Cette propriété a des conséquences fondamentales pour l'étude des systèmes linéaires invariants. L'opération de convolution dans la base canonique consiste en la multiplication d'une matrice $N \times N$ par un vecteur, soit $\mathcal{O}(N^2)$ opérations. Dans la base harmonique, la matrice est diagonale et le coût de la convolution est $\mathcal{O}(N)$.

En outre, il existe un algorithme rapide (Fast Fourier Transform) pour effectuer le calcul la transformée de Fourier discrète en $\mathcal{O}(N \log N)$ opérations. Le passage par la transformée de Fourier permet donc de réduire considérablement le coût numérique d'une opération de convolution.

5.3 Série de Fourier en temps continu

5.3.1 La base harmonique de $L_2[0, T)$

Un signal périodique T -périodique en temps-continu est défini sur l'ensemble de la droite réelle $\mathbb{R}$. Néanmoins, il est entièrement caractérisé par les valeurs qu'il prend sur l'intervalle $[0, T)$. Dans ce sens, il peut être identifié à un signal *fini* de l'espace vectoriel

$$L_2[0, T) = \{x \in ([0, T) \rightarrow \mathbb{C}) : \|x\|_2 < +\infty\}.$$

A l'inverse de son analogue discret, l'espace $L_2[0, T)$ est un espace vectoriel de dimension infinie. La théorie de Fourier peut néanmoins être développée de manière similaire.

Le choix de l'espace $L_2[0, T)$ est motivé par le fait que l'espace est muni d'une norme et que sa norme dérive du produit scalaire

$$\langle f, g \rangle = \int_0^T f(t) \overline{g(t)} dt.$$

On peut donc parler d'orthogonalité de deux signaux de $L_2[0, T)$ exactement comme dans l'espace $\ell_2[0..N)$ étudié dans la section précédente. La différence essentielle entre l'espace $\ell_2[0..N)$ et l'espace $L_2[0, T)$ est que nous passons d'un espace de dimension finie (N) à un espace de dimension *infinie* : il est en effet impossible d'engendrer toutes les fonctions de $L_2[0, T)$ à partir d'un nombre fini de signaux de base.

Dans un espace (normé) X de dimension infinie, une famille infinie dénombrable de vecteurs $\{\nu_k\}_{k \in \mathbb{Z}}$ est une base de l'espace si pour chaque $x \in X$, il existe une unique suite $\{\alpha_k\}_{k \in \mathbb{Z}}$ de scalaires telle que

$$x = \sum_{k \in \mathbb{Z}} \alpha_k \nu_k. \quad (5.3.1)$$

L'égalité (5.3.1) est définie au sens de la norme de X , c'est-à-dire

$$\lim_{N \rightarrow \infty} \left\| x - \sum_{k=-N}^N \alpha_k v_k \right\| = 0.$$

Les notions de base orthogonale et spectrale sont identiques en dimension finie et infinie. La base harmonique de $L_2[0, T)$ est définie par la famille infinie dénombrable de signaux

$$v_k(t) = \frac{1}{T} e^{jk \frac{2\pi}{T} t}, \quad k \in \mathbb{Z}. \quad (5.3.2)$$

L'orthogonalité des signaux v_k se vérifie aisément. En revanche, le fait que la famille (5.3.2) constitue une base de l'espace $L_2[0, T)$ est un résultat d'analyse fonctionnelle qui sort du cadre de ce cours. Ce résultat atteste qu'un signal x quelconque de $L_2[0, T)$ peut être décomposé de manière unique dans la base harmonique :

$$x(t) \sim \frac{1}{T} \sum_{k \in \mathbb{Z}} \hat{x}_k e^{jk \frac{2\pi}{T} t}.$$

Les coefficients $\hat{x}_k$ sont appelés coefficients de Fourier du signal x et valent

$$\hat{x}[k] = \hat{x}_k = \frac{\langle x, v_k \rangle}{\langle v_k, v_k \rangle} = \int_0^T x(\tau) e^{-jk \frac{2\pi}{T} \tau} d\tau. \quad (5.3.3)$$

L'expression du signal dans la base harmonique de l'Équation (5.3.2) constitue une deuxième transformée de Fourier, la transformée *continue-discrète*

$$x \xleftrightarrow{\mathcal{F}_{CD}} \hat{x} \quad \text{ou} \quad [0, T) \ni t \mapsto x(t) \xleftrightarrow{\mathcal{F}_{CD}} \mathbb{Z} \ni k \mapsto \hat{x}[k]. \quad (5.3.4)$$

5.3.2 Série de Fourier d'un signal périodique en temps-continu

L'implication pratique du changement de base discuté dans la section précédente est la suivante : tout signal T -périodique peut être exprimé comme une série de signaux harmoniques, c'est-à-dire tout signal T -périodique x admet la décomposition fréquentielle

$$x(t) \sim \frac{1}{T} \sum_{k \in \mathbb{Z}} \hat{x}_k e^{jk\omega_0 t}, \quad \omega_0 = \frac{2\pi}{T}. \quad (5.3.5)$$

Les coefficients (complexes) de Fourier $\hat{x}_k$ du signal sont donnés par l'expression

$$\hat{x}_k = \int_0^T x(\tau) e^{-jk\omega_0 \tau} d\tau. \quad (5.3.6)$$

Si x est un signal réel, sa décomposition fréquentielle est donnée par la partie réelle de l'expression (5.3.5). En regroupant les termes adéquats, on

obtient l'expression

$$x(t) \sim A_0 + \sum_{k=1}^{+\infty} A_k \cos(k\omega_0 t) + \sum_{k=1}^{+\infty} B_k \sin(k\omega_0 t).$$

Les coefficients A_k et B_k sont calculés par les formules suivantes :

$$A_k = \frac{2}{T} \int_0^T x(\tau) \cos(k\omega_0 \tau) d\tau, \quad B_k = \frac{2}{T} \int_0^T x(\tau) \sin(k\omega_0 \tau) d\tau.$$

On notera que tous les coefficients A_k sont nuls dans le cas d'un signal impair et que tous les coefficients B_k sont nuls dans le cas d'un signal pair.

5.3.3 Convergence des séries de Fourier

L'espace $L_2[0, T)$ contient des fonctions très peu régulières ayant une infinité de discontinuités. L'idée que des fonctions peu régulières puissent être représentées par une somme (infinie) de fonctions aussi lisses que les fonctions harmoniques fut avancée (sans démonstration) par Fourier en 1807 à la suite de mathématiciens du 18ème siècle tels que Euler (1748) et Bernoulli (1753). D'éminents mathématiciens comme Lagrange s'opposèrent fermement à une telle hérésie. Il convient en effet de s'interroger sur l'erreur commise lorsqu'un signal $x(t)$ est approximé par un développement en série de Fourier tronqué à l'ordre N :

$$x_N(t) = \frac{1}{T} \sum_{k=-N}^N \hat{x}_k e^{jk \frac{2\pi}{T} t}. \quad (5.3.7)$$

De même, le comportement asymptotique de cette erreur lorsque $N \rightarrow \infty$ doit être étudié.

Convergence en norme. Le fait que la base harmonique constitue une base de $L_2[0, T)$ (un résultat d'analyse fonctionnelle que nous n'avons pas démontré) garantit la convergence de la série pour tout signal x dans l'espace $L_2[0, T)$. Cette convergence est une convergence en norme, c'est-à-dire que l'erreur d'approximation $e_N = x - x_N$ satisfait

$$\lim_{N \rightarrow \infty} \|e_N\| = \lim_{N \rightarrow \infty} \sqrt{\int_0^T |e_N(t)|^2 dt} = 0.$$

L'orthogonalité de la base harmonique garantit en outre que x_N constitue la *meilleure approximation* de x dans le sous-espace $L_2^N[0, T)$ de $L_2[0, T)$ engendré par les vecteurs v_k , $k \in [-N..N]$. En effet x_N est la projection orthogonale

de x dans $L_2^N[0, T)$ et satisfait la relation

$$\|x_N - x\|_2 = \min_{y \in L_2^N[0, T)} \|y - x\|_2.$$

Cette propriété de meilleure approximation n'est en général pas vérifiée pour une base non orthogonale.

Convergence ponctuelle. La convergence en norme n'implique pas la convergence ponctuelle, définie par la propriété

$$\forall t \in [0, T) : \lim_{N \rightarrow \infty} x_N(t) = x(t).$$

La convergence ponctuelle peut néanmoins être garantie sous des hypothèses assez faibles, connues sous le nom de *conditions de Dirichlet* (lequel mit fin à la polémique entre Fourier et Lagrange en 1829) :

- le signal x est absolument intégrable sur $[0, T) : \int_0^T |x(t)| dt < \infty$;
- le signal x est à *variation bornée* : ceci signifie qu'il existe une constante α telle que $\sum_{i=1}^M |x(t_i) - x(t_{i-1})| \leq \alpha$ pour tout $M \in \mathbb{N}$ et pour toute séquence $(t_0, \dots, t_M)$ telle que $0 \leq t_0 \leq \dots \leq t_M < T$. Par exemple, si x possède un nombre fini d'extréma sur l'intervalle $[0, T)$, alors il est à variation bornée;
- le signal x a un nombre fini de discontinuités.

Sous ces conditions, on a

$$\frac{1}{T} \sum_{k \in \mathbb{Z}} \hat{x}_k e^{jk \frac{2\pi}{T} t} = \frac{x(t^+) + x(t^-)}{2}.$$

En particulier $\frac{1}{T} \sum_{k \in \mathbb{Z}} \hat{x}_k e^{jk \frac{2\pi}{T} t} = x(t)$ en tout t où x est continu.

Convergence uniforme. Si la convergence ponctuelle des séries de Fourier est garantie sous des hypothèses raisonnables, il importe de garder à l'esprit qu'elle n'implique pas la convergence uniforme, définie par la propriété

$$\lim_{N \rightarrow \infty} \sup_{t \in [0, T)} |x_N(t) - x(t)| = 0.$$

L'absence de convergence uniforme est connue en théorie des signaux sous le nom de *phénomène de Gibbs* et est illustrée par la série de Fourier d'un signal rectangulaire

$$x(t) = \begin{cases} 1, & |t| < T_1, \\ 0, & T_1 < |t| < \frac{T}{2}. \end{cases} \quad (5.3.8)$$

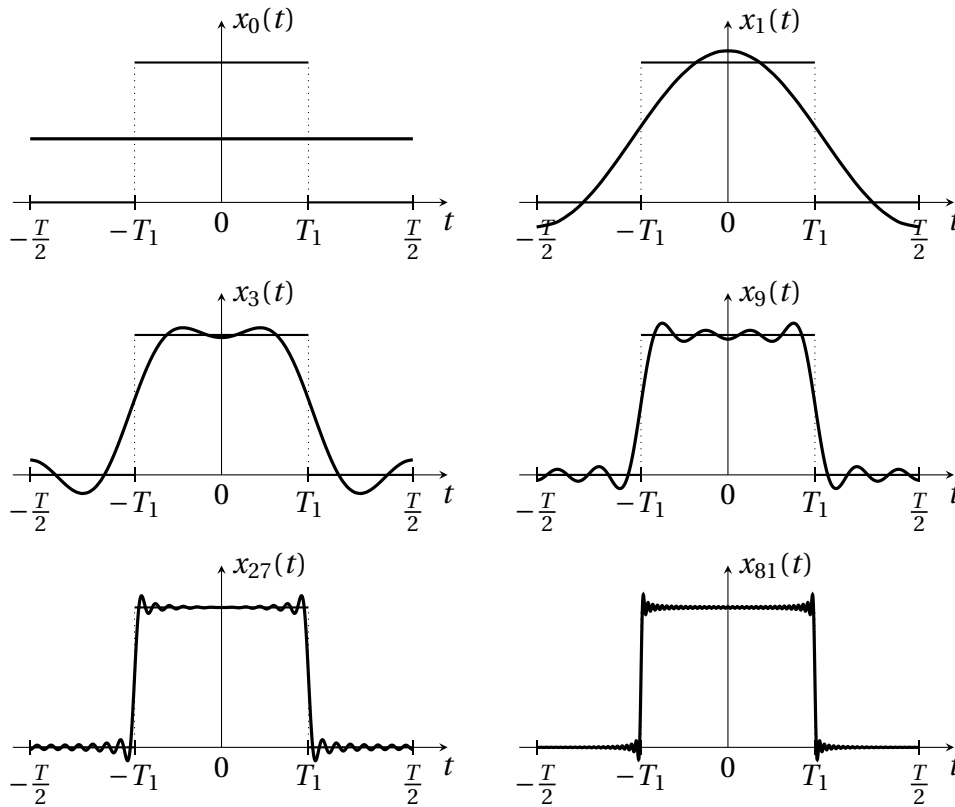


FIGURE 5.2 – Illustration du phénomène de Gibbs.

(Pour la commodité du calcul, l'intervalle de définition du signal est ici $[-\frac{T}{2}, \frac{T}{2})$.) Le calcul des coefficients de Fourier donne

$$\hat{x}_0 = \int_{-T_1}^{T_1} dt = 2T_1, \quad \hat{x}_k = \int_{-T_1}^{T_1} e^{-jk\frac{2\pi}{T}t} dt = T \frac{\sin(k\frac{2\pi}{T}T_1)}{k\pi}, \quad k \neq 0.$$

La Figure 5.2 illustre le signal x_N pour différentes valeurs de N . Le maximum de x_N se déplace vers la discontinuité lorsque N augmente, mais il reste environ 10% supérieur à la valeur maximale du signal x , quelle que soit la valeur de N .

5.3.4 La base harmonique est spectrale pour les opérateurs linéaires invariants

Tout comme en temps discret, la base harmonique possède une propriété remarquable par rapport aux opérateurs linéaires et invariants définis sur l'espace $L_2[0, T)$. L'invariance est ici exprimée en remplaçant l'opérateur de décalage usuel par un opérateur de décalage modulo T .

L'image des opérateurs linéaires et invariants est, comme dans le cas discret, définie par une opération de convolution cyclique de deux signaux

de $L_2[0, T)$:

$$y(t) = (u \odot h)(t) = \int_0^T h((t - \tau) \bmod T) u(\tau) d\tau. \quad (5.3.9)$$

On ne sera guère surpris de découvrir que, tout comme dans le cas discret, les vecteurs de la base harmonique sont des vecteurs propres pour l'opérateur LTI de réponse impulsionnelle h et que les valeurs propres correspondantes sont les coefficients de Fourier $\hat{h}_k$:

$$h \odot v_k(t) = \frac{1}{T} \int_0^T h(\tau) e^{jk \frac{2\pi}{T} ((t-\tau) \bmod T)} d\tau = \hat{h}_k \frac{1}{T} e^{jk \frac{2\pi}{T} t}. \quad (5.3.10)$$

La base harmonique diagonalise donc les opérateurs de convolution cyclique, conduisant à une nouvelle expression de la dualité convolution-multiplication (dans $L_2[0, T)$) :

$$y = h \odot u \xleftrightarrow{\mathcal{F}_{CD}} \hat{y} = \hat{h} \hat{u}. \quad (5.3.11)$$

Chapitre 6

Transformées de signaux discrets et continus

6.1 Introduction

Les outils de *transformation* de la théorie des systèmes (transformée de Laplace, transformée en z , transformée de Fourier,...) sont parmi les plus importants pour l'analyse.

L'idée de transformation d'un problème comme reformulation ou pré-traitement en vue d'une résolution simplifiée est au cœur même des mathématiques appliquées. Un exemple de transformée bien connu en arithmétique est la transformée logarithmique qui permet de remplacer l'opération de multiplication par l'opération (plus simple) d'addition : l'opération $y = u_1 u_2$ est décomposée en trois étapes : une transformée logarithmique

$$U_1 = \log u_1, \quad U_2 = \log u_2,$$

une opération d'addition

$$Y = U_1 + U_2,$$

qui remplace la multiplication, et une transformée logarithmique inverse

$$y = \log^{-1}(Y).$$

Grâce au pré-traitement (transformée) et post-traitement (transformée inverse) du problème, l'opération de multiplication est ainsi simplifiée en une opération d'addition. La simplification de la partie *calculatoire* du problème est équilibrée par le coût de la transformée. Mais si l'on utilise par exemple des tables pour celle-ci, il s'agit d'un gain net pour l'opérateur qui additionne au lieu de multiplier.

Il en va de même pour les transformées de signaux abordées dans ce chapitre. Elles simplifieront drastiquement la partie « calculatoire » de l'analyse des systèmes, conduisant par exemple à une résolution purement algébrique

des équations différentielles ou aux différences vues au chapitre précédent. Comme simplification la plus célèbre associée aux transformées de signaux, on peut citer le remplacement de l'opération de convolution entre deux signaux $y = u * h$ vue au Chapitre 3 par une opération de multiplication de leurs transformées : $Y = UH$.

L'attrait des transformées de signaux n'est pas seulement calculatoire. Il réside également dans l'interprétation physique des transformées comme pont entre l'analyse *temporelle* et l'analyse *fréquentielle* des systèmes et signaux. L'analyse temporelle du Chapitre 3 nous a conduit à concevoir un signal comme une combinaison d'impulsions δ se succédant dans le temps. Un système est alors étudié en analysant son effet sur une impulsion, c'est-à-dire sa réponse impulsionnelle h . Les transformées nous amèneront à concevoir un signal comme une combinaison d'exponentielles complexes (et en particulier une combinaison de sinusoides de différentes fréquences dans le cas de la transformée de Fourier). Un système sera alors étudié en analysant son effet sur une exponentielle complexe particulière, c'est-à-dire sa *fonction de transfert*.

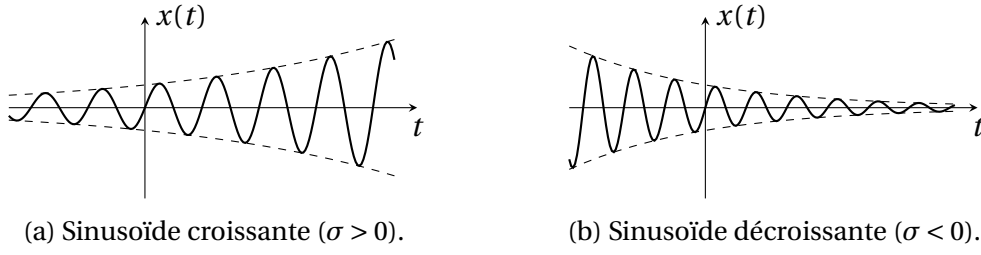
On verra dans ce chapitre que la fonction de transfert d'un système est la transformée de sa réponse impulsionnelle. C'est dans ce sens que les outils mathématiques de transformée introduisent un pont rigoureux entre l'analyse temporelle et l'analyse fréquentielle.

6.2 Exponentielles complexes et systèmes LTI

Nous avons vu dans le chapitre précédent qu'un opérateur linéaire invariant défini sur l'espace des signaux discrets finis $\ell_2[0..N)$ admet comme vecteurs propres les N signaux harmoniques $e^{jk\omega_0 n}$, $\omega_0 = \frac{2\pi}{N}$, $k = 0, \dots, N-1$. De manière analogue, un opérateur linéaire invariant défini sur l'espace des signaux continus finis $L_2[0, T)$ admet comme vecteurs propres les signaux harmoniques $e^{jk\omega_0 t}$, $\omega_0 = \frac{2\pi}{T}$, $k \in \mathbb{Z}$.

La définition des transformées repose sur la généralisation de cette propriété aux systèmes LTI dont les signaux d'entrée et de sortie ont un domaine non borné. Les vecteurs propres de ces opérateurs sont des signaux exponentiels.

En temps-continu, l'exponentielle complexe e^{st} généralise le signal $e^{jk\omega_0 t}$ à un nombre complexe $s = \sigma + j\omega$ quelconque du plan complexe. La factorisation $e^{st} = e^{\sigma t} e^{j\omega t}$ montre que la partie réelle du scalaire s détermine l'allure croissante ou décroissante du module de e^{st} . Si $\sigma > 0$, le module de e^{st} croît exponentiellement. Si $\sigma < 0$, le module décroît exponentiellement. La partie imaginaire du scalaire s détermine la fréquence d'oscillation du signal e^{st} . Ces comportements sont illustrés dans la Figure 6.1.

FIGURE 6.1 – Sinusoïde $x(t) = \Im(e^{st}) = e^{\sigma t} \sin(\omega t)$.

Un signal d'entrée $u(t) = e^{st}$ appliqué à un système LTI de réponse impulsionnelle h donne par définition la sortie

$$y(t) = h * u(t) = \int_{-\infty}^{+\infty} h(\tau) e^{s(t-\tau)} d\tau = e^{st} \int_{-\infty}^{+\infty} h(\tau) e^{-s\tau} d\tau,$$

en extrayant e^{st} de l'intégrale. En supposant que l'intégrale converge, on obtient bien un nombre complexe qui ne dépend que de la variable s et qui vaut

$$H(s) = \int_{-\infty}^{+\infty} h(\tau) e^{-s\tau} d\tau. \quad (6.2.1)$$

Le signal e^{st} est donc un vecteur propre du système LTI et la valeur propre correspondante est $H(s)$.

En temps-discret, l'exponentielle complexe z^n généralise le signal $e^{jk\omega_0 n}$ à un scalaire complexe quelconque $z = \rho e^{j\omega}$. La factorisation $z^n = \rho^n e^{j\omega n}$ montre que le module de z détermine l'allure croissante ou décroissante du module du signal z^n . Si $\rho > 1$, le module de z^n croît exponentiellement. Si $\rho < 1$, le module de z^n décroît exponentiellement.

On peut procéder de la même manière dans le cas discret : le signal d'entrée $u[n] = z^n$ appliqué au système LTI de réponse impulsionnelle h donne la sortie

$$y[n] = h * u[n] = \sum_{k \in \mathbb{Z}} h[k] z^{n-k} = z^n \sum_{k \in \mathbb{Z}} h[k] z^{-k}.$$

En supposant que la somme converge, c'est un nombre complexe qui ne dépend que de la variable complexe z et vaut

$$H(z) = \sum_{k \in \mathbb{Z}} h[k] z^{-k}. \quad (6.2.2)$$

La fonction H définie par (6.2.1) dans le cas continu ou par (6.2.2) dans le cas discret s'appelle *fonction de transfert* du système LTI de réponse impulsionnelle h . Elle caractérise le système d'une manière analogue à celle décrite au Chapitre 3 puisque la réponse à n'importe quelle entrée exprimée comme combinaison d'exponentielles complexes $u(t) = \sum_{i=1}^n a_i e^{s_i t}$ ou $u[n] = \sum_{i=1}^n a_i z_i^n$ est directement obtenue comme la même combinaison

pondérée par la fonction de transfert :

$$y(t) = \sum_{i=1}^n a_i H(s_i) e^{s_i t} \quad \text{ou} \quad y[n] = \sum_{i=1}^n a_i H(z_i) z_i^n.$$

6.3 Transformée de Laplace, transformée en z , et transformée de Fourier

La *transformée de Laplace* (bilatérale) d'un signal continu $x(t)$ est définie par

$$X(s) = \int_{-\infty}^{\infty} x(t) e^{-st} dt. \quad (6.3.1)$$

Son analogue discret pour un signal $x[n]$ est la *transformée en z*

$$X(z) = \sum_{k \in \mathbb{Z}} x[k] z^{-k}. \quad (6.3.2)$$

On utilise les notations $X = \mathcal{L}(x)$ et

$$x \xleftrightarrow{\mathcal{L}} X \quad \text{ou} \quad \mathbb{R} \ni t \rightarrow x(t) \xleftrightarrow{\mathcal{L}} \mathbb{C} \ni s \rightarrow X(s).$$

Similairement, $X = \mathcal{Z}(x)$ et

$$x \xleftrightarrow{\mathcal{Z}} X \quad \text{ou} \quad \mathbb{Z} \ni n \rightarrow x[n] \xleftrightarrow{\mathcal{Z}} \mathbb{C} \ni z \rightarrow X(z).$$

En comparant la définition (6.3.1) et l'expression (6.2.1), on vérifie que la fonction de transfert d'un système LTI est la transformée de Laplace (ou la transformée en z) de sa réponse impulsionnelle.

La variable indépendante des transformées (6.3.1) et (6.3.2) n'est plus le temps mais une variable complexe : $s = \sigma + j\omega$ pour la transformée de Laplace, et $z = \rho e^{j\omega}$ pour la transformée en z . On comprendra par la suite pourquoi on décompose s en partie réelle $\sigma = \Re(s)$ et imaginaire $\omega = \Im(s)$ et z en module $\rho = |z|$ et phase $\omega = \angle z$.

La *transformée de Fourier (continue-continue)* d'un signal $x(\cdot)$ est la valeur de la transformée de Laplace sur l'axe imaginaire, c'est-à-dire $s = j\omega$:

$$X(j\omega) = X(s) \Big|_{s=j\omega} = \int_{-\infty}^{\infty} x(t) e^{-j\omega t} dt. \quad (6.3.3)$$

La *transformée de Fourier (discrete-continue)* d'un signal $x[\cdot]$ est la valeur de la transformée en z sur le cercle unité $|z| = 1$, c'est-à-dire $z = e^{j\omega}$:

$$X(e^{j\omega}) = X(z) \Big|_{z=e^{j\omega}} = \sum_{k \in \mathbb{Z}} x[k] e^{-j\omega k}. \quad (6.3.4)$$

On obtient donc le lien suivant entre les différentes transformées : la transformée de Laplace le long de l'axe $\Re(s) = \sigma$ est la transformée de Fourier du signal $x(\cdot)$ multiplié par l'exponentielle réelle $e^{-\sigma t}$. Similairement, la transformée en z le long du cercle $|z| = \rho > 0$ est la transformée de Fourier (discrète) du signal $x[\cdot]$ multiplié par l'exponentielle réelle ρ^{-n} .

La transformée de Fourier est très largement utilisée pour l'analyse fréquentielle des signaux et on dispose d'algorithmes très performants pour le calcul numérique de cette transformée. Dans le reste de ce chapitre, on se concentrera sur les propriétés mathématiques des transformées plus générales : Laplace (6.3.1) et en z (6.3.2).

6.4 Transformées inverses et interprétation physique

A chacune des transformées définies dans la section précédente, on peut associer une transformée *inverse* qui permet de passer du signal transformé X au signal original x . La *transformée de Laplace inverse* est définie par

$$x(t) = \frac{1}{2\pi j} \int_{\sigma-j\infty}^{\sigma+j\infty} X(s) e^{st} ds \quad (6.4.1)$$

que l'on peut aussi réécrire comme

$$x(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} X(\sigma + j\omega) e^{(\sigma + j\omega)t} d\omega. \quad (6.4.2)$$

La *transformée en z inverse* est définie par

$$x[n] = \frac{1}{2\pi j} \oint X(z) z^{n-1} dz \quad (6.4.3)$$

que l'on peut aussi réécrire comme

$$x[n] = \frac{1}{2\pi} \int_{2\pi} X(\rho e^{j\omega}) (\rho e^{j\omega})^n d\omega. \quad (6.4.4)$$

Les formules (6.4.1) et (6.4.3) sont des formules d'intégration dans le plan complexe : le long de l'axe $\Re(s) = \sigma$ pour un certain σ constant dans le cas continu, et le long du cercle $|z| = \rho$ pour un certain ρ constant dans le cas discret. Nous ne les utiliserons jamais comme telles pour le calcul de transformées inverses et nous reportons leur justification mathématique à un chapitre ultérieur (Sections 10.1 et 10.2). Nous accepterons donc pour le moment sans justification que les formules (6.4.1) et (6.4.3) caractérisent les

transformées inverses $\mathcal{L}^{-1}$ et $\mathcal{F}^{-1}$ de sorte que

$$x(\cdot) \xleftrightarrow{\mathcal{L}} X \xleftrightarrow{\mathcal{L}^{-1}} x(\cdot) \quad \text{et} \quad x[\cdot] \xleftrightarrow{\mathcal{F}} X \xleftrightarrow{\mathcal{F}^{-1}} x[\cdot].$$

Les formules de transformées inverses sont importantes pour l'intuition physique contenue dans les transformées comme passage du domaine temporel au domaine fréquentiel. La relation (6.4.2) exprime le signal $x(\cdot)$ comme une combinaison (infinie) de sinusoides exponentiellement croissantes ou décroissantes. La partie réelle de l'exponentielle ($e^{\sigma t}$) est fixée mais la fréquence de la sinusoïde varie de $-\infty$ à $+\infty$. Le poids d'une fréquence particulière ω^* dans cette décomposition spectrale est donné par $X(\sigma + j\omega^*)$. La transformée de Fourier est un cas particulier où le signal $x(\cdot)$ est décomposé en une somme d'exponentielles imaginaires pures $e^{j\omega t}$. Cette représentation *fréquentielle* du signal $x(\cdot)$ est à comparer avec sa représentation *temporelle*

$$x(t) = \int_{-\infty}^{\infty} x(\tau) \delta(t - \tau) d\tau,$$

qui exprime le signal comme une combinaison (infinie) d'impulsions décalées $\delta(t - \tau)$. Le poids d'une impulsion à un instant donné τ^* dans cette décomposition temporelle est donné par $x(\tau^*)$.

Exercice 6.1. En se basant sur la formule inverse (6.4.2), calculer la transformée de Fourier d'une exponentielle imaginaire $e^{j\nu t}$ de fréquence ν . $\triangleleft$

La discussion est encore une fois analogue dans le cas discret. La formule (6.4.4) exprime le signal $x[\cdot]$ comme une combinaison (infinie) de sinusoides exponentiellement croissantes ou décroissantes. La partie réelle de l'exponentielle (ρ^n) est fixée mais la fréquence de la sinusoïde varie de 0 à 2π . Le poids d'une fréquence particulière ω^* dans cette décomposition spectrale est donné par $X(\rho e^{j\omega^*})$. La transformée de Fourier est un cas particulier où $x[\cdot]$ est décomposé en une somme de sinusoides pures $e^{j\omega n}$. Cette représentation *fréquentielle* du signal $x[\cdot]$ est à comparer avec sa représentation *temporelle*

$$x[n] = \sum_{k \in \mathbb{Z}} x[k] \delta[n - k],$$

qui exprime le signal comme une combinaison (infinie) d'impulsions décalées $\delta[n - k]$. Le poids d'une impulsion à un instant donné k^* dans cette décomposition temporelle est donné par $x[k^*]$.

On peut se demander la raison d'être de la partie réelle de l'exponentielle dans une décomposition fréquentielle du signal. En d'autres termes, pourquoi ne pas se limiter aux transformées de Fourier ($\sigma = 0$ ou $\rho = 0$) ? La raison est que nous n'avons pas encore discuté les conditions d'existence des intégrales infinies qui définissent les transformées. Il apparaîtra dans la section suivante

qu'un signal peut ne pas avoir de transformée de Fourier – l'intégrale (6.3.3) ou la somme (6.3.4) diverge – mais que la même intégrale (ou somme) devient convergente si le signal est multiplié par une exponentielle décroissante $e^{-\sigma t}$ (ou ρ^{-n}). Dans ce sens, la partie réelle de l'exponentielle est parfois appelé *facteur de convergence* de l'intégrale (ou somme).

6.5 Région de convergence (ROC)

Calculons d'abord la transformée de Laplace de l'exponentielle décroissante $x_1(t) = \mathbb{I}(t)e^{-at}$, $a > 0$. Par définition, on a

$$X_1(s) = \int_{-\infty}^{\infty} \mathbb{I}(t)e^{-(a+s)t} dt = \int_0^{\infty} e^{-(a+s)t} dt = -\frac{1}{a+s} e^{-(a+s)t} \Big|_0^{\infty}.$$

La transformée est donc définie si $\Re(a+s) > 0$ et elle vaut dans ce cas

$$X_1(s) = \frac{1}{s+a}, \quad \Re(s) > -a. \quad (6.5.1)$$

Si on recommence le même calcul pour le signal $x_2(t) = -\mathbb{I}(-t)e^{-at}$, on a

$$X_2(s) = \int_{-\infty}^{\infty} -\mathbb{I}(-t)e^{-(a+s)t} dt = -\int_{-\infty}^0 e^{-(a+s)t} dt = \frac{1}{a+s} e^{-(a+s)t} \Big|_{-\infty}^0$$

et donc

$$X_2(s) = \frac{1}{s+a}, \quad \Re(s) < -a. \quad (6.5.2)$$

En comparant les expressions (6.5.1) et (6.5.2), on constate que l'on a obtenu une même expression algébrique pour la transformée de Laplace de deux signaux différents. Ce qui différencie les transformées X_1 et X_2 , c'est le domaine du plan complexe sur lequel les transformées existent, c'est-à-dire le domaine du plan complexe sur lequel les intégrales existent. Il est donc important de spécifier ce domaine – appelé *région de convergence* (ROC) – en plus de l'expression algébrique de la transformée.

On peut répéter le même type de calcul dans le cas discret : la transformée en z du signal $\mathbb{I}[n]a^n$ est

$$X(z) = \sum_{k=0}^{\infty} (az^{-1})^k.$$

Cette somme converge si $|az^{-1}| < 1$, ce qui donne la région de convergence $|z| > |a|$ et la transformée

$$X(z) = \frac{z}{z-a}, \quad |z| > |a|.$$

Le signal $-\mathbb{I}[-n-1]a^n$ donnerait la même transformée mais une région de convergence $|z| < |a|$.

Exercice 6.2. Vérifier la table de transformées suivantes en appliquant la définition :

$$\begin{aligned}\mathbb{I}(\cdot) &\xleftrightarrow{\mathcal{L}} s \mapsto \frac{1}{s}, \Re(s) > 0, \\ \delta(\cdot) &\xleftrightarrow{\mathcal{L}} 1, \mathbb{C}, \\ t \mapsto \delta(t - t_0) &\xleftrightarrow{\mathcal{L}} s \mapsto e^{-st_0}, s \in \mathbb{C}.\end{aligned}$$

De manière analogue, dans le cas discret :

$$\begin{aligned}\mathbb{I}[\cdot] &\xleftrightarrow{\mathcal{Z}} z \mapsto \frac{z}{z-1}, |z| > 1, \\ \delta[\cdot] &\xleftrightarrow{\mathcal{Z}} 1, \mathbb{C}, \\ n \mapsto \delta[n - n_0] &\xleftrightarrow{\mathcal{Z}} z \mapsto z^{-n_0}, \begin{cases} |z| > 0, & n_0 > 0, \\ z \in \mathbb{C}, & \text{sinon.} \end{cases} \quad \triangleleft\end{aligned}$$

La région de convergence de la transformée de Laplace est l'ensemble des valeurs complexes de la variable s pour lesquelles l'intégrale

$$\int_{-\infty}^{\infty} x(t) e^{-st} dt$$

existe. En utilisant $|\int_{-\infty}^{\infty} x(t) e^{-st} dt| \leq \int_{-\infty}^{\infty} |x(t)| |e^{-st}| dt = \int_{-\infty}^{\infty} |x(t)| |e^{-\sigma t}| dt$, une condition suffisante d'existence de la transformée est obtenue sous la forme

$$\int_{-\infty}^{\infty} |x(t)| e^{-\sigma t} dt < \infty \quad (6.5.3)$$

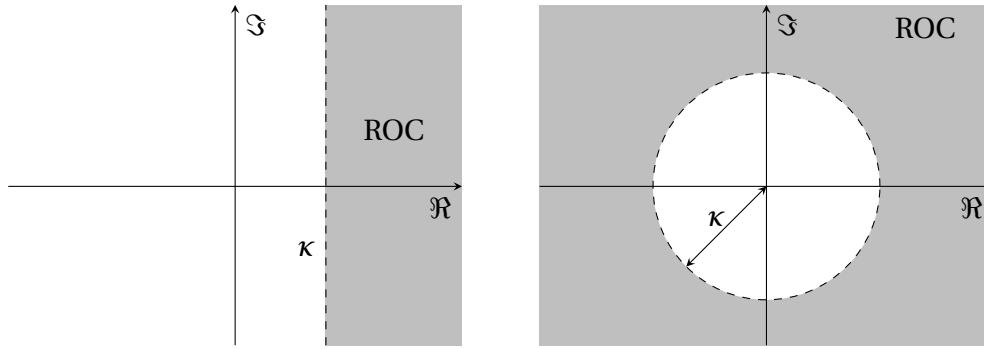
En désignant par (T_1, T_2) l'intervalle contenant le support du signal $x(\cdot)$, c'est-à-dire un intervalle de la droite réelle en dehors duquel le signal $x(\cdot)$ est identiquement nul, la condition (6.5.3) devient

$$\int_{T_1}^{T_2} |x(t)| e^{-\sigma t} dt < \infty \quad (6.5.4)$$

et on en déduit les conséquences suivantes :

- Un signal de *support fini* ($-\infty < T_1 \leq T_2 < +\infty$) et borné a une transformée de Laplace définie dans le plan complexe tout entier.
- Un signal de *support fini à gauche* ($-\infty < T_1$) et borné par une exponentielle (il existe deux constantes $\kappa \in \mathbb{R}$ et $C > 0$ telles que $|x(t)| \leq Ce^{\kappa t}$) a une transformée de Laplace définie dans le demi-plan $\{s \in \mathbb{C} : \Re(s) > \kappa\}$.
- Un signal de *support fini à droite* ($T_2 < \infty$) et borné par une exponentielle (il existe deux constantes $\kappa \in \mathbb{R}$ et $C > 0$ telles que $|x(t)| \leq Ce^{\kappa t}$) a une transformée de Laplace définie dans le demi-plan $\{s \in \mathbb{C} : \Re(s) < \kappa\}$.

En pratique, la majorité des signaux rencontrés dans la théorie des sys-



(a) ROC d'un signal continue ayant un support fini à gauche et étant borné par une exponentielle du type $Ce^{\kappa t}$.

(b) ROC d'un signal discret ayant un support fini à gauche et étant borné par une exponentielle du type $C\kappa^n$.

FIGURE 6.2 – Régions de convergence.

tèmes ont un support fini à gauche et sont bornés par une exponentielle. Pour de tels signaux, la région de convergence est au moins un demi-plan, comme illustré à la Figure 6.2a.

La discussion est analogue pour la transformée en z . La somme

$$\sum_{k \in \mathbb{Z}} x[k] z^{-k}$$

converge sous la condition

$$\sum_{k \in \mathbb{Z}} |x[k]| \rho^{-k} < \infty \quad (6.5.5)$$

et en désignant par $[k_1..k_2]$ le support du signal $x[\cdot]$, on en déduit les conséquences suivantes :

- Un signal de *support fini* ($-\infty < k_1 \leq k_2 < +\infty$) a une transformée en z définie dans le plan complexe tout entier.
- Un signal de *support fini à gauche* ($-\infty < k_1$) et borné par une exponentielle (il existe deux constantes $\kappa > 0$ et $C > 0$ telles que $|x[n]| \leq C\kappa^n$) a une transformée en z définie dans la région $\{s \in \mathbb{C} : |z| > \kappa\}$.
- Un signal de *support fini à droite* ($k_2 < \infty$) et borné par une exponentielle (il existe deux constantes $\kappa > 0$ et $C > 0$ telles que $|x[n]| \leq C\kappa^n$) a une transformée en z définie dans le disque $\{z \in \mathbb{C} : |z| < \kappa\}$.

En pratique, la majorité des signaux rencontrés dans la théorie des systèmes ont un support fini à gauche et sont bornés par une exponentielle. Pour de tels signaux, la région de convergence est illustrée à la Figure 6.2b.

6.6 Propriétés élémentaires

6.6.1 Linéarité

Une propriété immédiate mais cruciale des transformées est leur *linéarité*. Si x_1 a une transformée X_1 et une région de convergence R_1 , et si x_2 a une transformée X_2 et une région de convergence R_2 , alors le signal $x = ax_1 + bx_2$ a une transformée $X = aX_1 + bX_2$ et une région de convergence R qui contient l'intersection $R_1 \cap R_2$. S'il n'y pas d'intersection commune, la transformée de x n'est pas définie. Si $R_1 \cap R_2 \neq \emptyset$, la région de convergence R peut être plus grande que l'intersection $R_1 \cap R_2$. Par exemple le signal $x - x = 0$ a une région de convergence $R = \mathbb{C}$ quelle que soit la région de convergence de x .

Exemple 6.3. La transformée de Laplace du signal $x(t) = e^{-b|t|}$, $b \in \mathbb{R}$, est calculée comme suit : on a

$$x(t) = e^{-bt}\mathbb{I}(t) + e^{bt}\mathbb{I}(-t).$$

A l'inverse du signal x , chacun des deux signaux dans la somme est à support fini à droite ou à gauche et est borné par une exponentielle. On a donc

$$e^{-bt}\mathbb{I}(t) \xleftrightarrow{\mathcal{L}} \frac{1}{s+b}, \Re(s) > -b$$

et

$$e^{bt}\mathbb{I}(-t) \xleftrightarrow{\mathcal{L}} \frac{-1}{s-b}, \Re(s) < b.$$

Si $b < 0$, les deux régions de convergence n'ont pas d'intersection et le signal x n'a pas de transformée de Laplace. Si $b > 0$, on obtient par linéarité

$$x \xleftrightarrow{\mathcal{L}} \frac{1}{s+b} - \frac{1}{s-b} = \frac{-2b}{s^2 - b^2}, -b < \Re(s) < b.$$

La région de convergence est dans ce cas une bande verticale dans le plan complexe. $\triangleleft$

6.6.2 Décalage temporel et fréquentiel

Considérez un signal $x(\cdot)$ avec transformée de Laplace X dans une région de convergence R . Un décalage temporel correspond à la multiplication de la transformée par une exponentielle complexe :

$$t \mapsto x(t - t_0) \xleftrightarrow{\mathcal{L}} s \mapsto e^{-st_0} X(s), R.$$

Réciproquement, l'opération de décalage dans le domaine de Laplace correspond à une multiplication du signal par une exponentielle :

$$t \mapsto e^{s_0 t} x(t) \xleftrightarrow{\mathcal{L}} s \mapsto X(s - s_0), \quad R - \Re(s_0).$$

Le domaine de convergence est décalé vers la droite de $\Re(s_0)$.

Un cas particulier important est la multiplication du signal $x(\cdot)$ par une exponentielle imaginaire (une opération très fréquente dans les applications de télécommunications) :

$$t \mapsto e^{j\omega t} x(t) \xleftrightarrow{\mathcal{L}} s \mapsto X(s - j\omega), \quad R.$$

Les conclusions sont analogues dans le cas discret :

$$n \mapsto x[n - n_0] \xleftrightarrow{\mathcal{Z}} z \mapsto z^{-n_0} X(z), \quad \begin{cases} R \setminus \{0\}, & n_0 > 0, \\ R, & n_0 < 0, \end{cases} \quad (6.6.1)$$

et

$$n \mapsto z_0^n x[n] \xleftrightarrow{\mathcal{Z}} z \mapsto X(z/z_0), \quad |z_0| > R. \quad (6.6.2)$$

Lors d'un décalage temporel vers la droite, le point $z = 0$ doit être retiré du domaine de convergence de X . Lors d'un décalage temporel vers la gauche, le point à l'infini doit être « retiré » du domaine de convergence de X . Lors d'une multiplication par l'exponentielle z_0^n , le domaine de convergence de X est dilaté par $|z_0|$.

Comme dans le cas continu, le cas particulier de la multiplication de $x[\cdot]$ par une exponentielle imaginaire est important ; La transformée en z subit une rotation d'angle dans le plan z :

$$n \mapsto e^{j\omega_0 n} x[n] \xleftrightarrow{\mathcal{Z}} z \mapsto X(e^{-j\omega_0} z), \quad R.$$

6.7 La dualité convolution – multiplication

Nous avons vu dans l'introduction de ce chapitre que pour un système LTI de réponse impulsionnelle h , la réponse à une exponentielle complexe e^{st} est la même exponentielle complexe multipliée par H , la transformée de Laplace de h :

$$y(t) \Big|_{u=e^{st}} = H(s)e^{st}.$$

Exprimons à présent une entrée quelconque u comme une combinaison d'exponentielles complexes. Nous avons vu dans la Section 6.4 qu'une telle décomposition est obtenue par la transformée de Laplace inverse : si u admet

une transformée de Laplace U et que l'axe $\Re(s) = \sigma$ appartient à la région de convergence, on a

$$u(t) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} U(\sigma + j\omega) e^{(\sigma + j\omega)t} d\omega.$$

Puisque pour chaque exponentielle e^{st} , la réponse du système s'exprime par le produit $H(s)e^{st}$, on obtient directement en appliquant le principe de superposition que la réponse à une entrée u vaut

$$y(t) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} H(\sigma + j\omega) U(\sigma + j\omega) e^{(\sigma + j\omega)t} d\omega. \quad (6.7.1)$$

Mais si y admet une transformée de Laplace Y , on a également

$$y(t) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} Y(\sigma + j\omega) e^{(\sigma + j\omega)t} d\omega. \quad (6.7.2)$$

En identifiant les expressions (6.7.1) et (6.7.2) pour chaque valeur de σ , on obtient

$$Y(s) = H(s)U(s)$$

Mais y est la réponse du système à l'entrée u , c'est-à-dire $y = h * u$. On peut donc conclure que l'opération de convolution dans le domaine temporel est remplacée par une opération de multiplication dans le domaine fréquentiel :

$$t \mapsto y(t) = x_1 * x_2(t) \xleftrightarrow{\mathcal{L}} s \mapsto Y(s) = X_1(s)X_2(s). \quad (6.7.3)$$

La propriété (6.7.3) est capitale pour l'analyse des systèmes et explique en grande partie le succès des transformées comme outil d'analyse. Nous y reviendrons souvent dans les chapitres suivants.

On a rigoureusement la même propriété en discret :

$$n \mapsto y[n] = x_1 * x_2[n] \xleftrightarrow{\mathcal{Z}} z \mapsto Y(z) = X_1(z)X_2(z). \quad (6.7.4)$$

6.8 Différentiation et intégration

L'opération de différentiation dans le domaine temporel est également remplacée par une opération particulièrement simple dans le domaine fréquentiel : en dérivant (par rapport à t) les deux membres de l'expression

$$x(t) = \frac{1}{2\pi j} \int_{\sigma - j\infty}^{\sigma + j\infty} X(s) e^{st} ds,$$

on obtient

$$\dot{x}(t) = \frac{1}{2\pi j} \int_{\sigma-j\infty}^{\sigma+j\infty} sX(s)e^{st} ds,$$

ce qui signifie que x' est la transformée inverse de $sX(s)$. On conclut

$$\dot{x} \xleftrightarrow{\mathcal{L}} s \mapsto sX(s). \quad (6.8.1)$$

De manière réciproque, l'intégration temporelle d'un signal $x(t)$ devient une simple multiplication par $\frac{1}{s}$:

$$t \mapsto \int_{-\infty}^t x(\tau) d\tau \xleftrightarrow{\mathcal{L}} s \mapsto \frac{1}{s} X(s). \quad (6.8.2)$$

Nous pouvons interpréter la propriété (6.8.2) de la manière suivante : l'intégrateur est un système LTI dont la réponse impulsionnelle est l'échelon $\mathbb{I}(\cdot)$.

On a donc

$$\int_{-\infty}^t x(\tau) d\tau = x * \mathbb{I}(t).$$

La transformée de Laplace de l'échelon est $\frac{1}{s}$ et la propriété de convolution donne donc directement

$$x * \mathbb{I} \xleftrightarrow{\mathcal{L}} s \mapsto \frac{1}{s} X(s).$$

L'analogue discret d'un intégrateur est un retard ou décalage $\sigma x[n] = x[n-1]$. La réponse impulsionnelle d'un retard est $\sigma \delta[n] = \delta[n-1]$. On a donc

$$\sigma x = x * \sigma \delta.$$

La transformée en z de $\delta[n-1]$ est z^{-1} et la propriété de convolution donne

$$x * \sigma \delta \xleftrightarrow{\mathcal{Z}} z \mapsto z^{-1} X(z).$$

Chapitre 7

Analyse de systèmes LTI par les transformées de Laplace et en z

7.1 Fonctions de transfert

La *fonction de transfert* d'un système LTI est la transformée de sa réponse impulsionnelle : dans le cas continu, la transformée de Laplace

$$H(s) = \int_{-\infty}^{+\infty} h(\tau) e^{-s\tau} d\tau \quad (7.1.1)$$

et, dans le cas discret, la transformée en z

$$H(z) = \sum_{k \in \mathbb{Z}} h[k] z^{-k}. \quad (7.1.2)$$

Par la propriété de convolution des transformées, la réponse d'un système LTI continu est caractérisée par

$$Y = HU$$

où Y et U sont les transformées de Laplace ou en z de la sortie y et de l'entrée u .

En pratique, l'étude des systèmes LTI se fait essentiellement via l'analyse de la fonction de transfert du système. Nous allons voir en effet que le remplacement de l'opération de convolution par une opération de multiplication simplifie drastiquement les calculs.

Le calcul d'une fonction de transfert est particulièrement simple pour un système LTI décrit par une équation différentielle ou aux différences (cfr. Chapitre 3). Si la relation entrée-sortie est définie par une équation aux différences

$$\sum_{k=0}^M a_k y[n-k] = \sum_{k=0}^M b_k u[n-k]$$

l'application de la transformée en z aux deux membres donne

$$\sum_{k=0}^M a_k z^{-k} Y(z) = \sum_{k=0}^M b_k z^{-k} U(z).$$

On en déduit la fonction de transfert

$$H(z) = \frac{Y(z)}{U(z)} = \frac{\sum_{k=0}^M b_k z^{-k}}{\sum_{k=0}^M a_k z^{-k}} = \frac{N(z)}{D(z)}. \quad (7.1.3)$$

De même, si la relation entrée-sortie est définie par une équation différentielle

$$\sum_{k=0}^M a_k \frac{d^k y(t)}{dt^k} = \sum_{k=0}^M b_k \frac{d^k u(t)}{dt^k}, \quad (7.1.4)$$

l'application de la transformée de Laplace aux deux membres donne

$$\sum_{k=0}^M a_k s^k Y(s) = \sum_{k=0}^M b_k s^k U(s).$$

On en déduit la fonction de transfert

$$H(s) = \frac{Y(s)}{U(s)} = \frac{\sum_{k=0}^M b_k s^k}{\sum_{k=0}^M a_k s^k} = \frac{N(s)}{D(s)}. \quad (7.1.5)$$

La fonction de transfert d'un système LTI différentiel est donc une fonction *rationnelle*, c'est-à-dire le quotient de deux polynômes en la variable s ou z . Dans le cas discret, on exprime souvent $H(z)$ comme le quotient de deux polynômes en la variable z^{-1} plutôt que z car les coefficients des polynômes sont alors directement identifiés aux coefficients de l'équation aux différences.

Les *racines* des polynômes N et D qui définissent le numérateur et dénominateur de la fonction de transfert H jouent un rôle très important dans l'analyse. Les racines du numérateur N sont appelées les *zéros* du système : ce sont les valeurs complexes pour lesquelles la fonction de transfert s'annule. Les racines du dénominateur D sont appelées les *pôles* du système : ce sont les valeurs complexes pour lesquelles la fonction de transfert devient infinie.

Le calcul d'une fonction de transfert peut aussi être réalisé à partir d'un modèle d'état

$$\begin{aligned} \dot{x} &= Ax + Bu, \\ y &= Cx + Du. \end{aligned} \quad (7.1.6)$$

La transformée de Laplace des deux membres de l'équation donne

$$\begin{aligned} sX(s) &= AX(s) + BU(s), \\ Y(s) &= CX(s) + DU(s). \end{aligned}$$

L'élimination de $X(s)$ dans l'expression de la sortie fournit

$$Y(s) = C(sI - A)^{-1}BU(s) + DU(s) \quad (7.1.7)$$

et on en déduit que la fonction de transfert du système (7.1.6) est

$$H(s) = \frac{Y(s)}{U(s)} = C(sI - A)^{-1}B + D. \quad (7.1.8)$$

En comparant (7.1.8) à l'expression de la réponse impulsionnelle calculée en (4.3.4), on a

$$e^{At}\mathbb{I}(t) \xleftrightarrow{\mathcal{L}} (sI - A)^{-1},$$

qui est l'analogie matriciel de la relation $e^{at}\mathbb{I}(t) \xleftrightarrow{\mathcal{L}} \frac{1}{s-a}$.

Exemple 7.1. La fonction de transfert associée à l'équation différentielle du premier ordre

$$\dot{y} + 3y = u \quad (7.1.9)$$

est donnée par

$$\frac{1}{s+3}. \quad (7.1.10)$$

Nous avons vu dans la Section 6.5 que l'on peut associer deux réponses impulsionnelles différentes à la fonction de transfert (7.1.10) : si la région de convergence est $\Re(s) > -3$, on obtient la réponse impulsionnelle

$$h(t) = e^{-3t}\mathbb{I}(t). \quad (7.1.11)$$

Par contre, si la région de convergence est $\Re(s) < -3$, on obtient la réponse impulsionnelle

$$h(t) = -e^{-3t}\mathbb{I}(-t). \quad (7.1.12)$$

L'expression (7.1.11) est la solution de l'équation différentielle (7.1.9) pour l'entrée $u = \delta$, la condition initiale $y(0^-) = 0$, et pour $t \geq 0$. C'est la réponse impulsionnelle du système LTI *causal* associé à l'équation différentielle (7.1.9). Par contre, l'expression (7.1.12) est la solution de l'équation différentielle (7.1.9) pour l'entrée $u = \delta$, la condition « initiale » $y(0^+) = 0$, et pour $t \leq 0$. C'est la réponse impulsionnelle du système LTI *anti-causal* associé à l'équation différentielle (7.1.9). $\triangleleft$

L'exemple précédent montre qu'il y a une relation étroite entre la propriété de causalité d'un système LTI et la région de convergence de sa fonction de transfert. Pour obtenir la réponse impulsionnelle du système LTI causal associé à l'équation différentielle (7.1.4), il faut prendre la transformée de Laplace inverse de la fonction de transfert (7.1.3) dont la région de convergence est le demi-plan à droite de l'axe vertical fixé par la partie réelle du pôle

le plus à droite. Cette propriété est en accord avec notre discussion dans la Section 6.5 à propos de la région de convergence d'un signal dont le support est fini à gauche, ce qui, par définition, est le cas de la réponse impulsionnelle d'un système causal.

7.2 Bloc-diagrammes et algèbre de fonctions de transfert

Nous avons vu au Chapitre 1 qu'un système est très souvent décrit comme une interconnexion de sous-systèmes élémentaires. La reconstitution de la fonction de transfert du système global à partir des fonctions de transfert des différents sous-systèmes est particulièrement aisée et en fait purement algébrique.

Soient deux systèmes LTI ayant pour réponse impulsionnelle h_1 et h_2 et pour fonction de transfert H_1 et H_2 . La réponse impulsionnelle de leur interconnexion parallèle (Exercice 3.4) est la somme $h_1 + h_2$. Puisque les transformées sont linéaires, la fonction de transfert de l'interconnexion parallèle est également la somme

$$H_{\text{parallèle}} = H_1 + H_2 \quad (7.2.1)$$

Considérons maintenant leur interconnexion série : sa réponse impulsionnelle est la convolution $h_1 * h_2$. En vertu de la dualité convolution-multiplication, la fonction de transfert de l'interconnexion série est donc le simple produit

$$H_{\text{série}} = H_1 H_2 \quad (7.2.2)$$

Le troisième type d'interconnexion de base est l'interconnexion feedback. le calcul de sa réponse impulsionnelle n'est pas aisé. En revanche, la fonction de transfert entre u et y est directement déduite de la Figure 7.1 :

$$Y = H_1 E, \quad E = U - Z, \quad Z = H_2 Y \quad (7.2.3)$$

d'où on tire la relation

$$\frac{Y}{U} = H_{\text{feedback}} = \frac{H_1}{1 + H_1 H_2} \quad (7.2.4)$$

En combinant les formules élémentaires (7.2.1), (7.2.2), et (7.2.4), il est donc possible de déterminer la fonction de transfert de n'importe quel bloc-diagramme constitué d'interconnexions séries, parallèles, ou feedback.

Inversement, on peut utiliser ces formules élémentaires pour construire le bloc-diagramme d'un système à partir de sa fonction de transfert. Le bloc-diagramme ainsi obtenu dépend bien sûr de la manière utilisée pour dé-

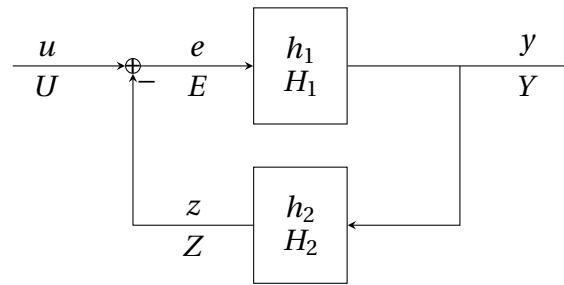


FIGURE 7.1 – Bloc-diagramme d'une interconnexion feedback.

composer la fonction de transfert en fonctions élémentaires. Enfin, nous avons vu au Chapitre 4 qu'une représentation d'état correspond à chaque bloc-diagramme. Nous pouvons donc utiliser différents blocs-diagramme pour obtenir différentes représentations d'état du système.

Exemple 7.2. Soit un système LTI causal du second ordre défini par la fonction de transfert

$$H(s) = \frac{1}{(s+1)(s+2)} = \frac{1}{s^2 + 3s + 2}. \quad (7.2.5)$$

C'est la fonction de transfert du système différentiel

$$\ddot{y} + 3\dot{y} + 2y = u. \quad (7.2.6)$$

Au Chapitre 4, nous avons associé à un tel système le bloc-diagramme de la Figure 7.2a, appelé forme directe. Alternativement, on peut concevoir H comme la mise en série de deux systèmes du premier ordre :

$$H(s) = \left(\frac{1}{s+1} \right) \left(\frac{1}{s+2} \right).$$

Le bloc-diagramme correspondant (Figure 7.2b) est appelé forme cascade. Enfin, un développement de H en fractions simples exprime le système comme la mise en parallèle de deux systèmes du premier ordre :

$$H(s) = \frac{1}{s+1} - \frac{1}{s+2}.$$

Le bloc-diagramme correspondant (Figure 7.2c) est appelé forme parallèle. <

7.3 Résolution de systèmes différentiels ou aux différences initialement au repos

Les transformées de Laplace et en z fournissent un outil mathématique pour la résolution de systèmes différentiels ou aux différences initialement

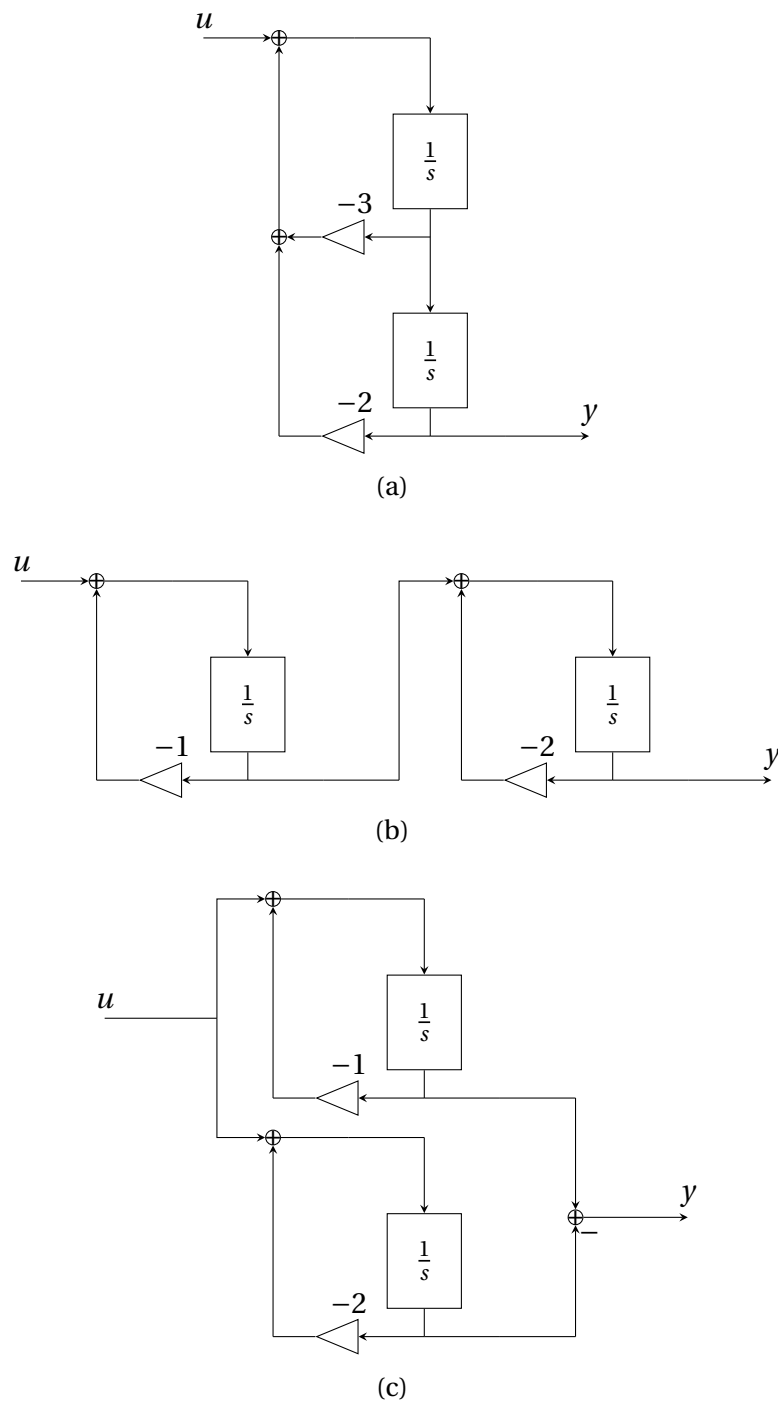


FIGURE 7.2 – Blocs-diagrammes de la fonction de transfert (7.2.5).

au repos. Nous illustrons cet outil sur un exemple. Soit le système différentiel décrit par l'équation

$$\ddot{y} + 3\dot{y} + 2y = a\mathbb{I}. \quad (7.3.1)$$

En supposant une condition initiale nulle

$$y(0) = \dot{y}(0) = 0,$$

le signal $y(t)$ est la sortie d'un système LTI causal. La région de convergence de sa fonction de transfert est un demi-plan ouvert à droite. Il en va de même pour la transformée (bilatérale) de Laplace $U(s) = \frac{a}{s}$ du signal d'entrée $u(t) = a\mathbb{I}(t)$. Dans la région de convergence commune des deux signaux, on obtient la transformée de l'équation (7.3.1)

$$s^2 Y(s) + 3sY(s) + 2 = \frac{a}{s}, \quad (7.3.2)$$

ou encore

$$Y(s) = H(s)U(s) = \left(\frac{1}{s^2 + 3s + 2} \right) \frac{a}{s}, \quad \Re(s) > 0. \quad (7.3.3)$$

La décomposition de $Y(s)$ en fractions simples donne

$$Y(s) = \frac{a}{2s} - \frac{a}{s+1} + \frac{a}{2(s+2)}, \quad \Re(s) > 0$$

et on en déduit la solution de l'équation différentielle

$$y(t) = a \left(\frac{1}{2} - e^{-t} + \frac{1}{2} e^{-2t} \right) \mathbb{I}(t).$$

La procédure ainsi décrite est valable pour calculer la solution de n'importe quelle équation différentielle linéaire à coefficients constants pour un signal d'entrée nul pour $t < 0$ et pour une condition initiale nulle. La propriété de causalité garantit que les transformées de Laplace de u et de y ont une région de convergence commune (un demi-plan ouvert à droite si u est borné par une exponentielle) et permet donc la résolution d'une équation algébrique dans le domaine de Laplace. La décomposition en fractions simples permet une inversion simple de Y pour reconstruire le signal de sortie y .

Comme illustré dans les sections précédentes, le calcul des fonctions de transfert permet de déterminer la transformée Y du signal de sortie y par des opérations algébriques. La détermination du signal y demande alors de déterminer la transformée inverse de Y . Lorsque Y est une fraction rationnelle, une méthode aisée de calcul consiste à développer Y en *fractions simples*.

7.3.1 Transformées inverses et décomposition en fractions simples : Systèmes continus

Soit l'expression suivante du signal de sortie d'un système LTI causal en temps-continu dans le domaine de Laplace :

$$Y(s) = \frac{b_M s^M + b_{M-1} s^{M-1} + \dots + b_1 s + b_0}{s^N + a_{N-1} s^{N-1} + \dots + a_1 s + a_0}. \quad (7.3.4)$$

Si $M - N = Q \geq 0$, cette dernière peut encore s'écrire, après une simple division de polynômes, sous la forme

$$Y(s) = c_Q s^Q + c_{Q-1} s^{Q-1} + \dots + c_1 s + c_0 + \frac{d_L s^L + d_{L-1} s^{L-1} + \dots + d_1 s + d_0}{s^N + a_{N-1} s^{N-1} + \dots + a_1 s + a_0}, \quad (7.3.5)$$

où $L < N$. La région de convergence du polynôme de coefficients c_i est le plan complexe tout entier. De plus, nous savons que $\delta \xleftrightarrow{\mathcal{L}} 1$, $\text{ROC} = \mathbb{C}$, et que la multiplication par s de la transformée correspond à la dérivée dans le domaine temporel. Ainsi,

$$\delta^{(m)} \xleftrightarrow{\mathcal{L}} s \mapsto s^m. \quad (7.3.6)$$

La fraction rationnelle de (7.3.5), notée $G(s)$, peut être traitée efficacement en la décomposant en une somme de fractions simples de la forme $\frac{1}{(s-p)^m}$.

La transformée inverse de chaque terme est alors obtenue grâce à la formule suivante :

$$t \mapsto \frac{t^{m-1} e^{pt}}{(m-1)!} \mathbb{I}(t) \xleftrightarrow{\mathcal{L}} s \mapsto \frac{1}{(s-p)^m}, \quad \Re(s) > \Re(p). \quad (7.3.7)$$

La décomposition en fractions simples s'effectue de la manière suivante.

Dans une première étape, on factorise le dénominateur de $G(s)$ afin de faire apparaître ses pôles :

$$G(s) = \frac{d_L s^L + d_{L-1} s^{L-1} + \dots + d_1 s + d_0}{(s-p_1)^{m_1} (s-p_2)^{m_2} \dots (s-p_r)^{m_r}}. \quad (7.3.8)$$

Selon cette notation, $G(s)$ a r pôles p_i , $1 \leq i \leq r$, de multiplicité respective m_i . On a évidemment

$$\sum_{i=1}^r m_i = N.$$

Ensuite, le théorème de décomposition en série de Laurent des fonctions complexes analytiques dans $\mathbb{C} \setminus \{p_i\}$ nous assure que (7.3.8) peut s'écrire de

manière unique sous la forme d'une somme de fractions simples

$$G(s) = \sum_{k=1}^{m_1} \frac{A_{1k}}{(s-p_1)^k} + \sum_{k=1}^{m_2} \frac{A_{2k}}{(s-p_2)^k} + \cdots + \sum_{k=1}^{m_r} \frac{A_{rk}}{(s-p_r)^k}, \quad (7.3.9)$$

où les A_{ik} sont N constantes à déterminer.

Cette dernière étape peut s'effectuer par différentes méthodes. La plus simple à justifier, mais aussi la plus longue d'utilisation, consiste à réécrire (7.3.9) sous la forme (7.3.8) en additionnant les fractions et à identifier les coefficients des puissances de s dans le numérateur ; il reste alors à résoudre un système de N équations à N inconnues. Une autre technique bien plus rapide consiste à utiliser directement la formule suivante :

$$A_{ik} = \frac{1}{(m_i - k)!} \left[\frac{d^{m_i-k}}{ds^{m_i-k}} ((s-p_i)^{m_i} G(s)) \right]_{s=p_i}. \quad (7.3.10)$$

La formule (7.3.10) pour l'obtention des coefficients de Laurent d'une fraction rationnelle ressemble au théorème des résidus de l'analyse des fonctions complexes, dont elle consiste en quelque sorte une généralisation : si p_i est un pôle simple de $G(s)$, la formule (7.3.10) se réduit au seul coefficient

$$A_{i1} = (s-p_i)G(s) \Big|_{s=p_i},$$

qui est le résidu du pôle p_i .

La validité de (7.3.10) s'illustre à l'aide d'un exemple.

Exemple 7.3. Soit la fraction rationnelle

$$G(s) = \frac{b_2 s^2 + b_1 s + b_0}{(s-p_1)^2 (s-p_2)}. \quad (7.3.11)$$

Nous écrivons sa décomposition en fractions simples

$$G(s) = \frac{A_{11}}{s-p_1} + \frac{A_{12}}{(s-p_1)^2} + \frac{A_{21}}{s-p_2}. \quad (7.3.12)$$

En additionnant les 3 fractions de (7.3.12) et en égalant le numérateur obtenu à celui de (7.3.11), on obtient

$$b_2 s^2 + b_1 s + b_0 = A_{11}(s-p_1)(s-p_2) + A_{12}(s-p_2) + A_{21}(s-p_1)^2.$$

Notons que les trois termes sont bien nécessaires afin d'obtenir un système bien posé de 3 équations à 3 inconnues en égalant les coefficients des termes en s^2 , en s^1 et en s^0 .

Afin de vérifier la formule (7.3.10), multiplions (7.3.12) par $(s-p_1)^2$ pour

obtenir

$$(s - p_1)^2 G(s) = A_{11}(s - p_1) + A_{12} + \frac{A_{21}(s - p_1)^2}{s - p_2}, \quad (7.3.13)$$

ce qui donne bien $[(s - p_1)^2 G(s)]|_{s=p_1} = A_{12}$. On conçoit aisément comment $[(s - p_2)G(s)]|_{s=p_2} = A_{21}$ de manière similaire. Enfin, dérivant (7.3.13) par rapport à s , on supprime la constante A_{12} tandis que le terme en A_{21} sera toujours multiplié par $(s - p_1)$:

$$\frac{d}{ds}(s - p_1)^2 G(s) = A_{11} + A_{21} \left(\frac{2(s - p_1)}{s - p_2} - \frac{(s - p_1)^2}{(s - p_2)^2} \right). \quad (7.3.14)$$

Finalement, l'évaluation de (7.3.14) en $s = p_1$ isole donc la valeur de A_{11} .

On peut ainsi imaginer comment (7.3.10) se justifie pour des ordres de dérivées plus élevés. $\triangleleft$

La région de convergence de chaque terme (7.3.6) ou (7.3.7) de la décomposition (7.3.9) de Y contient la région de convergence de Y . En effet, cette dernière est $\{s \in \mathbb{C} : \Re(s) > \Re(p_{\max})\}$ où $p_{\max}$ est le pôle de Y de partie réelle maximale. Dès lors, grâce à la propriété $\text{ROC}_{X_1+X_2} \supseteq \text{ROC}_{X_1} \cap \text{ROC}_{X_2}$, nous pouvons prendre la transformée inverse terme à terme dans (7.3.9), en utilisant les formules (7.3.6) et (7.3.7). Ceci fournit le signal y recherché.

7.3.2 Transformées inverses et décomposition en fractions simples : Systèmes discrets

Le raisonnement est analogue au cas continu. Cependant, nous préférons partir de l'expression suivante de la sortie dans le domaine en z :

$$Y(z) = \frac{b_M z^{-M} + b_{M-1} z^{-(M-1)} + \dots + b_1 z^{-1} + b_0}{a_N z^{-N} + a_{N-1} z^{-(N-1)} + \dots + a_1 z^{-1} + 1}. \quad (7.3.15)$$

(On vérifie facilement qu'une fraction rationnelle en z^{-1} est en fait aussi une fraction rationnelle en z .)

En effet, d'une part les équations aux différences mènent directement à des fonctions de transfert qui sont des fractions rationnelles en z^{-1} et d'autre part les formules de transformée inverse s'expriment également sous cette forme. Pour un système causal, on utilisera les transformées suivantes :

$$n \mapsto \frac{(n + m - 1)!}{n!(m - 1)!} p^n \mathbb{I}[n] \xleftrightarrow{\mathcal{Z}} z \mapsto \frac{1}{(1 - pz^{-1})^m}, \quad |z| > |p|. \quad (7.3.16)$$

A nouveau, lorsque $M > N$, une division polynomiale permet d'exprimer $Y(z)$ comme la somme d'une fraction rationnelle propre (degré du numérateur inférieur au degré du dénominateur) et d'un polynôme (en z^{-1}). La

transformée inverse de ce dernier s'obtient tout simplement en se rappelant la propriété

$$n \mapsto \delta[n-k] \xleftrightarrow{\mathcal{Z}} z \mapsto z^{-k}, z \neq 0. \quad (7.3.17)$$

La fraction rationnelle peut être décomposée en fractions simples comme dans le cas continu : la seule différence est le remplacement de la variable s par la variable z^{-1} . Le choix de la valeur 1 pour le terme indépendant du dénominateur permet une factorisation de ce dernier sous la forme d'un produit $(1 - p_1 z^{-1})^{m_1} (1 - p_2 z^{-1})^{m_2} \dots (1 - p_r z^{-1})^{m_r}$ qui fournira directement des termes du type (7.3.16).

On obtient ainsi

$$G(z^{-1}) = \sum_{i=1}^r \sum_{k=1}^{m_i} \frac{B_{ik}}{(1 - p_i z^{-1})^k},$$

où les coefficients B_{ik} peuvent être obtenus analytiquement soit comme précédemment par identification des coefficients d'égales puissances de z^{-1} , soit à l'aide d'une adaptation de la formule (7.3.10) à la factorisation présente du dénominateur, ce qui donne

$$B_{ik} = \frac{1}{(m_i - k)! (-p_i)^{m_i - k}} \left[\frac{d^{m_i - k}}{(dz^{-1})^{m_i - k}} \left((1 - p_i z^{-1})^{m_i} G(z^{-1}) \right) \right]_{z=p_i}.$$

Les régions de convergence des termes de type (7.3.16) et (7.3.17) ainsi obtenus contiennent à nouveau la région de convergence de $Y(z)$, si bien que le signal causal $y[n]$ associé est obtenu par application directe, terme à terme, des transformées inverses (7.3.16) et (7.3.17).

Chapitre 8

Réponses transitoire et permanente d'un système LTI

La Figure 8.1 illustre la propriété fondamentale sur laquelle repose l'analyse d'un système par sa fonction de transfert : un signal exponentiel en entrée d'un système LTI donne le même signal exponentiel en sortie, multiplié par la fonction de transfert du système. Cette propriété constitue toutefois une abstraction mathématique car le signal exponentiel n'est pas limité dans le temps. Une expérience physique impose d'appliquer un signal d'entrée à partir d'un instant initial, choisi en $t = 0$ par convention. De plus, dans bon nombre d'expériences, la condition initiale du système n'est pas nulle. L'objectif de ce chapitre est d'évaluer ce que devient la propriété illustrée à la Figure 8.1 dans cette situation plus réaliste.

La transformée unilatérale, couverte dans la Section 8.1, permet d'étudier l'effet de conditions initiales non nulles dans un système LTI causal. La propriété de *stabilité*, introduite dans la Section 8.3, permet de décomposer la sortie en une partie *transitoire* et une partie *permanente*. Cette décomposition permet de retrouver la propriété de transfert des signaux exponentiels à *un transitoire près*.

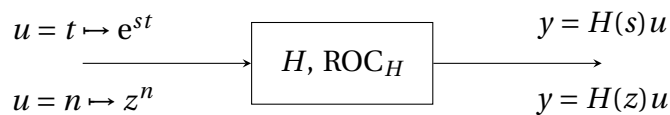


FIGURE 8.1 – La propriété de transfert d'un système LTI (valable dans le domaine de convergence de H).

8.1 Transformée unilatérale

8.1.1 Transformée de Laplace

La transformée de Laplace *unilatérale* est définie par

$$\mathcal{X}(s) = \int_{0^-}^{\infty} x(t) e^{-st} dt, \quad (8.1.1)$$

où la limite inférieure d'intégration 0^- signifie que l'on inclut dans l'intervalle d'intégration l'effet d'une impulsion ou de toute autre singularité concentrée en $t = 0$.[†] Par opposition, la transformée de Laplace définie au Chapitre 6 est appelée *bilatérale*.

La transformée de Laplace unilatérale d'un signal x est la transformée bilatérale du signal $x\mathbb{I}$. La région de convergence de la transformée unilatérale est donc toujours un demi-plan ouvert à droite. Les transformées unilatérale et bilatérale coïncident pour tous les signaux identiquement nuls pour $t < 0$. Il en va de même pour toutes leurs propriétés. En particulier, la fonction de transfert d'un système LTI causal est indifféremment la transformée unilatérale ou bilatérale de la réponse impulsionnelle.

La différence fondamentale entre les transformées unilatérales et bilatérales réside dans la propriété de différentiation. Soit x un signal et $\mathcal{X}$ sa transformée de Laplace unilatérale. En intégrant par parties, on obtient que la transformée unilatérale de sa dérivée vaut

$$\int_{0^-}^{\infty} \dot{x}(t) e^{-st} dt = x(t) e^{-st} \Big|_{0^-}^{\infty} + s \int_{0^-}^{\infty} x(t) e^{-st} dt = s\mathcal{X}(s) - x(0^-). \quad (8.1.2)$$

On constate donc l'apparition d'un nouveau terme $-x(0^-)$ par rapport à la propriété de différentiation obtenue dans le cas de la transformée bilatérale. Il en va de même pour les dérivées d'ordre supérieur. En dérivant une nouvelle fois, on obtient que la transformée unilatérale de Laplace de $\ddot{x}$ vaut

$$s^2 \mathcal{X}(s) - sx(0^-) - \dot{x}(0^-),$$

où $\dot{x}(0^-)$ désigne la dérivée de x évaluée en $t = 0^-$.

[†]. Par convention, on suppose qu'une impulsion de Dirac a lieu dans l'« intervalle » $(0^-, 0^+)$. Par exemple, si x est une impulsion de Dirac, sa transformée unilatérale de Laplace est égale à sa transformée bilatérale, parce que l'intervalle d'intégration inclut l'impulsion. Par contre, si l'intégration (8.1.1) débutait en $t = 0^+$, l'intégrale serait nulle dans le cas d'une impulsion.

Théorèmes de la valeur initiale et finale. La relation (8.1.2) est utile pour déterminer la valeur « initiale »

$$x(0^+) = \lim_{t \downarrow 0} x(t)$$

et la valeur « finale »

$$x(+\infty) = \lim_{t \rightarrow +\infty} x(t)$$

d'un signal x à partir de sa transformée unilatérale $\mathcal{X}$. Le théorème de la valeur initiale s'exprime comme

$$x(0^+) = \lim_{s \rightarrow \infty} s\mathcal{X}(s). \quad (8.1.3)$$

Inversement, si la région de convergence de $s \mapsto s\mathcal{X}(s)$ contient l'axe imaginaire, alors le théorème de la valeur finale s'exprime comme

$$x(+\infty) = \lim_{s \rightarrow 0} s\mathcal{X}(s). \quad (8.1.4)$$

Les deux théorèmes se démontrent à partir de la relation (8.1.2).

Exercice 8.1. Démontrer le théorème de la valeur finale en prenant la limite pour $s \rightarrow 0$ de l'équation (8.1.2). $\triangleleft$

8.1.2 Transformée en z

Les considérations sur la transformée de Laplace unilatérale sont analogues pour la transformée en z unilatérale, définie par

$$\mathcal{X}(z) = \sum_{k=0}^{\infty} x[k]z^{-k}. \quad (8.1.5)$$

Comme dans le cas continu, la transformée unilatérale d'un signal x est la transformée bilatérale du même signal multiplié par un échelon-unité.

La différence majeure avec la transformée bilatérale réside dans la propriété de décalage : si le signal x a pour transformée unilatérale $\mathcal{X}$, la transformée unilatérale du signal décalé $y[n] = x[n-1]$ s'obtient de la manière suivante :

$$\mathcal{Y}(z) = \sum_{k=0}^{\infty} x[k-1]z^{-k} = x[-1] + z^{-1} \sum_{k=0}^{\infty} x[k]z^{-k} = x[-1] + z^{-1}\mathcal{X}(z). \quad (8.1.6)$$

On constate l'apparition du nouveau terme $x[-1]$ par rapport à la propriété de décalage de la transformée bilatérale.

8.2 Résolution de systèmes différentiels ou aux différences avec des conditions initiales non nulles

La transformée unilatérale permet d'étendre la méthode de résolution décrite dans la Section 7.3 à des systèmes avec conditions initiales non-nulles. A titre d'illustration, nous reprenons l'exemple traité dans la Section 7.3 :

$$\ddot{y} + 3\dot{y} + 2y = a\mathbb{I} \quad (8.2.1)$$

avec cette fois une condition initiale arbitraire $y(0^-) = y_0$, $\dot{y}(0^-) = \dot{y}_0$. En appliquant la transformée unilatérale aux deux membres de l'équation (7.3.1), on obtient

$$s^2\mathcal{Y}(s) - y_0s - \dot{y}_0 + 3s\mathcal{Y}(s) - 3y_0 + 2\mathcal{Y}(s) = \frac{a}{s} \quad (8.2.2)$$

ou encore

$$\mathcal{Y}(s) = \frac{y_0(s+3)}{s^2+3s+2} + \frac{\dot{y}_0}{s^2+3s+2} + \frac{a}{s(s^2+3s+2)}. \quad (8.2.3)$$

Pour chaque valeur particulière des constantes a , y_0 , et $\dot{y}_0$, on peut développer l'expression (8.2.3) en fractions simples pour obtenir la solution $y(t)$. Par exemple, si $a = 2$, $y_0 = 3$, et $\dot{y}_0' = -5$, on a

$$\mathcal{Y}(s) = \frac{1}{s} - \frac{1}{s+1} + \frac{3}{(s+2)}$$

et la transformée inverse donne, pour $t > 0$,

$$y(t) = 1 - e^{-t} + 3e^{-2t}.$$

8.2.1 Réponse libre et réponse forcée

En comparant les équations (7.3.3) et (8.2.3), on voit que la solution d'une équation différentielle linéaire quelconque peut être exprimée comme la somme de deux réponses distinctes : la réponse *forcée*, qui est la réponse du système LTI causal initialement au repos associé à l'équation différentielle et soumis à l'entrée u , et la réponse *libre*, due à l'effet des conditions initiales en l'absence d'excitation extérieure. La réponse forcée est la réponse obtenue quand la condition initiale est nulle (en anglais, on parle aussi de *zero-state response*), tandis que la réponse libre est la réponse obtenue quand l'entrée est identiquement nulle (en anglais, on parle aussi de *zero-input response*).

8.2.2 Les conditions initiales vues comme des impulsions de Dirac

Le fait que la réponse totale du système soit la somme de la réponse libre et de la réponse forcée est une conséquence du principe de superposition : une combinaison linéaire de deux entrées distinctes donne en sortie la même

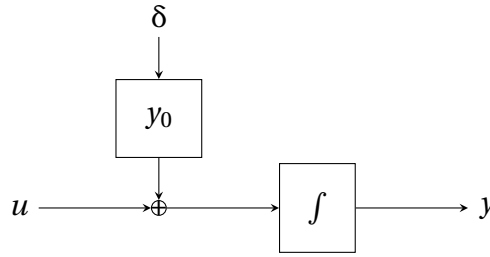


FIGURE 8.2 – Le bloc-diagramme d'un intégrateur avec une condition initiale non nulle.

combinaison linéaire des sorties distinctes. Il convient à ce titre de représenter l'effet des conditions initiales comme l'action d'une entrée particulière sur un système causal initialement au repos. Pour un intégrateur, cette opération est très aisée : le système différentiel

$$\dot{y} = u, \quad y(0^-) = y_0$$

est « équivalent » au système différentiel

$$\dot{y} = u + y_0 \delta, \quad y(0^-) = 0$$

dans le sens où la solution $y(t)$ des deux systèmes est identique pour $t > 0$ (on vérifie que la transformée unilatérale de Laplace donne dans les deux cas $s\mathcal{Y}(s) = \mathcal{U}(s) + y_0$). Le bloc-diagramme d'un intégrateur à condition initiale non nulle est illustré à la Figure 8.2. L'effet de la condition initiale est modélisé par l'action d'une impulsion superposée à l'action de l'entrée. D'une manière plus générale, on peut donc construire le bloc-diagramme d'une équation différentielle quelconque, en modélisant l'effet des conditions initiales comme des impulsions de Dirac à l'entrée de chaque intégrateur.

Les trois termes de la réponse (8.2.3) peuvent maintenant être interprétés comme les effets de trois entrées distinctes sur trois systèmes LTI causaux distincts :

$$\mathcal{Y}(s) = H_1(s)y_0 + H_0(s)y'_0 + H(s)\frac{a}{s} \quad (8.2.4)$$

Les différentes fonctions de transfert H , H_0 , et H_1 auront le même dénominateur et donc les mêmes pôles. De plus, l'ordre du numérateur de H_i est égal à l'ordre de la dérivée de la condition initiale $y_0^{(i)}$. Pour l'Exemple (7.3.1), on a donc

$$\mathcal{Y}(s) = H(s)(s+3)y_0 + H(s)y'_0 + H(s)\frac{a}{s} \quad (8.2.5)$$

avec $H(s) = \frac{1}{s^2+3s+2}$ la fonction de transfert associée au système différentiel (7.3.1).

8.2.3 Modes propres d'un système différentiel

Les pôles de la fonction de transfert $H(s)$ associée à un système différentiel sont parfois appelés modes propres du système (parce qu'ils ne dépendent pas de l'entrée appliquée). A chaque pôle p_i correspond une exponentielle $e^{p_i t}$ et la réponse du système fait intervenir chacune de ces exponentielles. L'expression (8.2.5) fait clairement apparaître que les modes propres influencent aussi bien la réponse libre du système que sa réponse forcée.

8.3 Stabilité d'un système LTI

Un système est *stable* si toutes les paires entrée-sortie (u, y) solutions du système satisfont une borne du type

$$\sup_{t \in \mathbb{R}} |y(t)| \leq K \sup_{t \in \mathbb{R}} |u(t)|$$

où la constante K est parfois appelée gain du système. Ce type de stabilité est parfois désigné dans la littérature sous le nom de *stabilité BIBO* (*bounded input bounded output*) puisque toute entrée bornée produira une sortie bornée. Un système est *instable* s'il n'est pas stable, c'est-à-dire si il existe une entrée bornée qui conduit à une sortie non bornée.

Dans toutes les applications d'ingénierie, la stabilité du système est une spécification de base. Un système physique est constamment soumis à des perturbations et l'amplification potentielle de ces perturbations par un système instable conduit inévitablement à un système incapable de réaliser la tâche à laquelle il est destiné (transmission, filtrage, régulation) ou opérant en dehors de la région linéaire des « petits signaux » (effets de saturation, destruction du système, ...).

Nous allons caractériser la stabilité d'un système LTI par une propriété de sa réponse impulsionnelle : dans le cas continu, on a

$$y(t) = h * u(t) = \int_{-\infty}^{+\infty} h(\tau) u(t - \tau) d\tau$$

et donc la majoration

$$|y(t)| \leq \int_{-\infty}^{+\infty} |h(\tau)| |u(t - \tau)| d\tau.$$

Si l'entrée $u(t)$ est bornée, i.e. $|u(t)| \leq C < \infty$ pour tout t , alors

$$|y(t)| \leq C \int_{-\infty}^{+\infty} |h(\tau)| d\tau.$$

Une condition suffisante pour la stabilité BIBO d'un système LTI est donc que sa réponse impulsionnelle soit absolument intégrable. Mathématiquement, cela signifie que le signal h appartient à l'espace $L_1(\mathbb{R})$, c'est-à-dire

$$\|h\|_1 = \int_{-\infty}^{+\infty} |h(\tau)| d\tau = K < \infty. \quad (8.3.1)$$

Dans ce cas, on aura toujours

$$\sup_{t \in \mathbb{R}} |y(t)| \leq K \sup_{t \in \mathbb{R}} |u(t)|$$

et donc n'importe quelle entrée bornée donnera une sortie bornée. Dans le cas discret, on obtient de la même manière comme condition suffisante que le signal h appartienne à l'espace $\ell_1(\mathbb{Z})$, c'est-à-dire

$$\sum_{k \in \mathbb{Z}} |h[k]| = K < \infty. \quad (8.3.2)$$

On démontre facilement que les conditions (8.3.1) et (8.3.2) sont non seulement suffisantes mais également nécessaires pour la stabilité BIBO : si elles ne sont pas satisfaites, on peut toujours construire une entrée bornée qui fera diverger la sortie.

Il est intéressant d'exprimer le critère de stabilité BIBO dans le domaine fréquentiel : l'intégrabilité de la réponse impulsionnelle implique l'existence de la transformée de Laplace $H(s)$ sur l'axe imaginaire :

$$H(j\omega) = \int_{-\infty}^{\infty} e^{-j\omega\tau} h(\tau) d\tau$$

et

$$\left| \int_{-\infty}^{\infty} e^{-j\omega\tau} h(\tau) d\tau \right| \leq \int_{-\infty}^{\infty} |h(\tau)| d\tau = K < \infty.$$

Une condition suffisante de stabilité BIBO est donc que l'axe imaginaire soit inclus dans la région de convergence de la fonction de transfert $H(s)$. Cette condition est également nécessaire : si $H(j\omega_p)$ n'est pas défini pour un certain ω_p , cela signifie que l'entrée bornée $u(t) = e^{j\omega_p t}$ produit une sortie divergente. On conclut donc que les trois propositions suivantes sont équivalentes :

- un système LTI continu (resp. discret) est BIBO stable ;
- sa réponse impulsionnelle est absolument intégrable (resp. sommable) ;
- l'axe imaginaire (resp. le cercle unité) est inclus dans la région de convergence de sa fonction de transfert.

Exercice 8.2. Vérifier que le système LTI continu défini par un retard pur $y(t) = u(t - 1)$ est BIBO stable. $\triangleleft$

Examinons à présent le critère de stabilité BIBO dans le cas particulier

d'un système LTI *causal* décrit par un système différentiel ou aux différences. Dans ce cas, la fonction de transfert H est rationnelle et sa région de convergence est le demi-plan borné à gauche par la partie réelle du pôle le plus à droite (respectivement, dans le cas discret, le complément du disque centré à l'origine et de rayon correspondant au maximum des modules des différents pôles). En utilisant le critère développé ci-dessus, on obtient pour le cas discret :

Le système LTI causal discret associé à un système aux différences linéaire à coefficients constants est BIBO stable si et seulement si tous les pôles de sa fonction de transfert sont de module strictement inférieur à un.

Similairement pour le cas continue :

Le système LTI causal continu associé à un système différentiel linéaire à coefficients constants est BIBO stable si et seulement si tous les pôles de sa fonction de transfert sont à partie réelle strictement négative et le degré du numérateur n'est pas supérieur à celui du dénominateur.

Si le degré du numérateur est supérieur à celui de dénominateur, on a des termes en s^k , avec $k > 0$, dans la fonction de transfert. Ces termes correspondent avec des dérivées dans le domaine temporelle; la dérivation est une opération instable.

Exemple 8.3. Considérons le signal borné $\sin(t^2)$. Alors $\frac{d \sin(t^2)}{dt} = 2t \cos(t^2)$, ce qui n'est pas borné. $\triangleleft$

8.3.1 Critère de Routh-Hurwitz

La stabilité d'un système différentiel LTI causal de fonction de transfert

$$H(s) = \frac{q_M s^M + q_{M-1} s^{M-1} + \dots + q_1 s + q_0}{p_N s^N + p_{N-1} s^{N-1} + \dots + p_1 s + p_0}$$

est donc caractérisée par une condition simple : tous les pôles du système doivent être à partie réelle strictement négative. On peut donc tester la stabilité d'un système différentiel en calculant ses pôles, c'est-à-dire les N racines du polynôme

$$p_N s^N + p_{N-1} s^{N-1} + \dots + p_1 s + p_0 = 0, \quad p_N > 0. \quad (8.3.3)$$

À l'heure où les ordinateurs n'existaient pas, il était particulièrement utile de pouvoir déterminer la stabilité d'un système *sans* recourir au calcul explicite

de ses pôles. Maxwell fut le premier à formuler ce problème, sans pouvoir y apporter une réponse satisfaisante. Le problème fit date dans la littérature des mathématiques appliquées avant d'être élégamment résolu par Routh et par Hurwitz.

Pour énoncer le critère de Routh–Hurwitz, considérons la manière suivante de générer une nouvelle suite de nombres réels à partir de deux suites données :

$$\begin{aligned} \text{suite 1 : } & a_1 \quad a_2 \quad a_3 \quad \dots, \\ \text{suite 2 : } & b_1 \quad b_2 \quad b_3 \quad \dots, \\ \Rightarrow & c_1 \quad c_2 \quad c_3 \quad \dots, \end{aligned} \quad (8.3.4)$$

avec $c_k = b_1 a_{k+1} - a_1 b_{k+1}$. (On notera que c_k est simplement le déterminant de la matrice $\begin{bmatrix} a_1 & a_{k+1} \\ b_1 & b_{k+1} \end{bmatrix}$ changé de signe.)

Construisons à présent un tableau dont les deux premières rangées sont les coefficients des parties paires et impaires du polynôme (8.3.3) :

$$\begin{aligned} \text{rangée 1 : } & p_0 \quad p_2 \quad p_4 \quad \dots, \\ \text{rangée 2 : } & p_1 \quad p_3 \quad p_5 \quad \dots, \end{aligned}$$

où $p_i = 0$ pour $i > N$. La troisième rangée du tableau de Routh–Hurwitz est construite à partir des deux premières en appliquant la procédure (8.3.4) aux rangées 1 et 2; la quatrième rangée, en appliquant (8.3.4) aux rangées 2 et 3, et ainsi de suite jusqu'à obtenir une rangée complète de zéros (on vérifie que la dernière rangée non nulle est la $N + 1$ -ième rangée). Si on note r_i le premier élément de la rangée $i + 1$, on a $r_0 = p_0$, $r_1 = p_1$, $r_2 = p_1 p_2 - p_0 p_3$, etc. Le critère de Routh–Hurwitz s'énonce alors comme suit :

Les n racines du polynôme (8.3.3) sont à partie réelle strictement négative si et seulement si les coefficients $r_0, r_1, \dots, r_N$ du tableau de Routh–Hurwitz sont strictement positifs.

On retiendra une condition nécessaire de stabilité impliquée par ce test et très simple à vérifier : tous les coefficients du polynôme (8.3.3) doivent être strictement positifs : $p_i > 0$ pour $i = 0, 1, 2, \dots, N - 1$. Cette condition est également suffisante pour un polynôme d'ordre 2, mais pas pour des polynômes d'ordre plus élevé. Par exemple, pour un polynôme du troisième degré, on obtient la condition

$$p_2 > 0, \quad p_1 > \frac{p_0}{p_2}, \quad p_0 > 0$$

Remarque. Il existe un critère « de type Routh » qui permet de vérifier si toutes les racines d'un polynôme sont de module strictement inférieur à 1, et ainsi de tester la stabilité d'un système LTI causal décrit par une équation aux différences. Sa formulation est moins directe et sort du cadre de ce cours. <

8.4 Réponse transitoire et réponse permanente

Si un système différentiel LTI est stable, alors le signal harmonique $u(t) = e^{j\omega t}$ en entrée donne la sortie (bornée) $y = H(j\omega)u$. Nous allons à présent étudier comment cette propriété se modifie dans le cas où l'entrée $u(t) = e^{j\omega t}\mathbb{I}(t)$ est un signal harmonique *initialisé* en $t = 0$. On suppose dans le développement que $H(s)$ est une fraction rationnelle.

Par définition, la transformée de Laplace du signal de sortie vaut dans ce cas

$$\mathcal{Y}(s) = H(s)\mathcal{U}(s) = \frac{N(s)}{D(s)} \frac{1}{s - j\omega}, \quad (8.4.1)$$

que l'on peut exprimer sous la forme

$$\mathcal{Y}(s) = \frac{\gamma(s)}{D(s)} + \frac{H(j\omega)}{s - j\omega}, \quad (8.4.2)$$

où le polynôme γ a un degré inférieur au degré de N . Puisque toutes les racines de D sont à partie réelle strictement négative, la transformée inverse donne

$$y(t) = y_{\text{trans}}(t) + H(j\omega)e^{j\omega t}\mathbb{I}(t), \quad t \geq 0 \quad (8.4.3)$$

avec la propriété $\lim_{t \rightarrow \infty} y_{\text{trans}}(t) = 0$.

La réponse (8.4.3) concerne la réponse *forcée* du système. Si la condition initiale du système est non nulle, il faut y ajouter la réponse *libre* du système. Si le système est stable, cette dernière est également une somme d'exponentielles décroissantes et tend donc asymptotiquement vers zéro. L'effet des conditions initiales sur la réponse du système est donc également un effet *transitoire*.

La partie $H(j\omega)e^{j\omega t}\mathbb{I}(t)$ de la réponse (8.4.3) est appelée *réponse permanente* du système à l'entrée $e^{j\omega t}\mathbb{I}(t)$ car elle constitue le seul terme qui ne tend pas asymptotiquement vers zéro dans la réponse du système. Nous avons rappelé au début du chapitre la propriété de transfert $y = H(j\omega)u$ valable pour tout système LTI stable soumis à l'entrée harmonique $u(t) = e^{j\omega t}$. Le développement qui précède montre que cette propriété de transfert s'étend à *un transitoire près* aux signaux d'entrée $u(t) = \mathbb{I}(t)e^{j\omega t}$ et en présence d'une condition initiale non nulle. L'importance de l'effet transitoire dépend de la position des pôles du système dans le plan complexe. Dans bon nombre d'applications, on néglige cet effet transitoire et on étudie les performances du système sur base de la réponse permanente.

Chapitre 9

Réponse fréquentielle d'un système LTI

La réponse impulsionnelle h d'un système LTI caractérise sa réponse temporelle par la représentation de convolution. La fonction de transfert H d'un système LTI caractérise sa réponse fréquentielle car sa valeur $H(j\omega)$ sur l'axe imaginaire (ou $H(e^{j\omega})$ sur le cercle unité) exprime l'effet du système sur un signal harmonique de fréquence arbitraire. Dans ce chapitre, nous détaillons la représentation graphique usuelle de la réponse fréquentielle $H(j\omega)$ et $H(e^{j\omega})$ au moyen de diagrammes de Bode.

Le diagramme de Bode d'un système est une information graphique aussi importante pour l'analyse fréquentielle du système que ne l'est le graphe de la réponse impulsionnelle ou indicielle pour l'analyse temporelle. Ces graphes temporel et fréquentiel constituent l'information la plus utilisée dans les applications de filtrage, de télécommunications, ou d'asservissement. Ils sont redondants sur le plan mathématique (la réponse fréquentielle est la transformée de Fourier de la réponse impulsionnelle) mais en pratique l'analyse ou la synthèse se font en parallèle dans le domaine temporel et fréquentiel car certaines propriétés sont immédiatement lues sur le graphe de la réponse impulsionnelle ou indicielle (par exemple le temps de réponse) tandis que d'autres sont immédiatement lues sur le graphe de la réponse fréquentielle (par exemple la bande passante).

L'objectif de ce chapitre est de montrer la complémentarité de l'approche temporelle et de l'approche fréquentielle dans l'analyse des systèmes. Leur utilisation parallèle fait apparaître des compromis fondamentaux qui sont au cœur même de l'ingénierie des systèmes. On détaillera en particulier les graphes de systèmes du premier et deuxième ordre qui constituent les blocs de base pour la plupart des applications.

9.1 Réponse fréquentielle d'un système LTI

L'analyse d'un système LTI au moyen de la transformée de Laplace repose sur la propriété observée au début du Chapitre 6 : la réponse à une expo-

entielle complexe est la même exponentielle multipliée par la fonction de transfert :

$$u(t) = e^{st} \Rightarrow y(t) = H(s)e^{st} \quad (9.1.1)$$

En utilisant la relation $2\cos\omega t = e^{j\omega t} + e^{-j\omega t}$, on obtient en particulier la réponse à un signal sinusoïdal :

$$u(t) = \cos\omega t \Rightarrow y(t) = \frac{1}{2}(H(j\omega)e^{j\omega t} + H(-j\omega)e^{-j\omega t}) = \Re(H(j\omega)e^{j\omega t}),$$

où on a utilisé $\overline{H(j\omega)} = H(-j\omega)$, ce qui suit du fait que la réponse impulsionnelle h est réelle. La représentation polaire

$$H(j\omega) = |H(j\omega)|e^{j\angle H(j\omega)}$$

donne donc

$$u(t) = \cos\omega t \Rightarrow y(t) = |H(j\omega)|\cos(\omega t + \angle H(j\omega)). \quad (9.1.2)$$

La propriété (9.1.2) est une reformulation de la propriété (9.1.1) pour des signaux sinusoïdaux : lorsqu'un signal sinusoïdal de fréquence ω passe au travers d'un système LTI, son amplitude est multipliée par $|H(j\omega)|$ et sa phase est décalée d'un angle $\angle H(j\omega)$. Nous avons vu dans le chapitre précédent que cette propriété s'étend, à un transitoire près, aux signaux initialisés $u(t) = \mathbb{I}(t)\cos(\omega t)$.

Le nombre positif $|H(j\omega)|$ est appelé le *gain* du système à la fréquence ω . Le graphe de $|H(j\omega)|$ en fonction de ω est appelé la réponse en amplitude du système. L'angle $\angle H(j\omega)$ est appelé la *phase* du système à la fréquence ω . Le graphe de $\angle H(j\omega)$ en fonction de ω est appelé la réponse en phase du système. Les deux graphes ensemble constituent la *réponse fréquentielle* du système.

Exemple 9.1. Calculons la réponse fréquentielle d'un différentiateur pur $y = \dot{u}$. Sa fonction de transfert est $H(s) = s$ et donc

$$H(j\omega) = j\omega = \omega e^{j\frac{\pi}{2}}.$$

On en déduit

$$|H(j\omega)| = \omega, \quad \angle H(j\omega) = \frac{\pi}{2}.$$

On reconnaît cette propriété en examinant la relation $\frac{d\cos(\omega t)}{dt} = -\omega \sin(\omega t)$.

L'amplitude d'un différentiateur pur croît donc linéairement avec la fréquence. Ceci signifie qu'un différentiateur amplifie fortement les hautes fréquences. C'est un phénomène indésirable en pratique puisque dans les applications, le rapport bruit/signal est plus élevé dans les hautes fréquences. C'est la raison pour laquelle un système n'est pas physiquement réalisable

au moyen de différentiateurs. Par contre, on peut vérifier que l'amplitude d'un intégrateur pur est $\frac{1}{\omega}$, c'est-à-dire que l'intégrateur atténue les hautes fréquences. $\triangleleft$

L'exemple ci-dessus illustre que le diagramme d'*amplitude* d'un système révèle ses propriétés de filtrage. Si le gain du système est élevé dans une certaine plage de fréquences $[\omega_1, \omega_2]$ et faible en dehors, le système assure une bonne transmission des signaux dans la gamme de fréquences considérée tandis qu'il atténue les signaux qui ont une fréquence inférieure à ω_1 ou supérieure à ω_2 . Ces systèmes sont appelés *passe-bande* car ils éliminent le contenu fréquentiel des signaux transmis en dehors d'une certaine bande de fréquences. De même, le différentiateur pur a une caractéristique *passe-haut* car il amplifie davantage les hautes fréquences et l'intégrateur a une caractéristique *passe-bas* car il amplifie davantage les basses fréquences.

La *phase* d'un système est également une caractéristique très importante. Qualitativement, une variation de phase est souvent associée à un décalage (avance ou retard) temporel du signal. Un retard pur $y(t) = u(t - t_0)$ a comme fonction de transfert $H(s) = e^{-st_0}$, ce qui implique un gain unité et une phase linéaire :

$$|H(j\omega)| = 1, \quad \angle H(j\omega) = -\omega t_0.$$

La *pente* de la courbe de phase correspond donc au retard appliqué au signal temporel. Pour un système plus général, la courbe de phase est une fonction non linéaire de la fréquence. Dans ce cas, la tangente à la courbe en un point donné ω_0 approxime le retard appliqué aux signaux temporels de fréquence proche de ω_0 . Cet effet de retard (ou d'avance) différent sur chaque fréquence peut causer une distorsion importante du signal temporel. Dans les applications de régulation, un retard – et plus généralement des variations importantes dans la phase du système – est souvent associé à une détérioration de la stabilité et de la robustesse. Certaines propriétés de dissipativité (l'impossibilité physique pour un système de créer de l'énergie de manière interne) sont, pour un système LTI, entièrement caractérisées par des propriétés de phase du système, sans aucune restriction sur l'amplitude.

9.2 Diagrammes de Bode

Pour représenter le diagramme d'amplitude d'un système, on utilise le plus souvent une échelle logarithmique, de sorte que la relation multiplicative $|Y(j\omega)| = |H(j\omega)||U(j\omega)|$ devient additive comme dans le cas de la phase :

$$\log|Y(j\omega)| = \log|H(j\omega)| + \log|U(j\omega)|$$

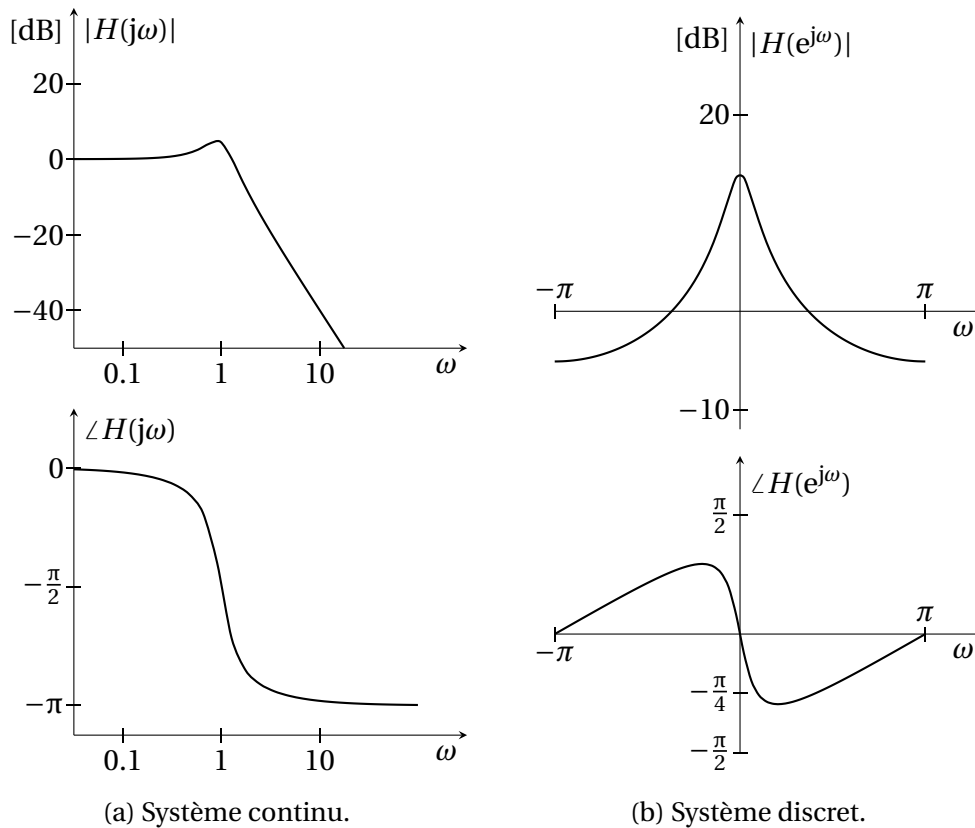


FIGURE 9.1 – Exemples de diagrammes de Bode.

et

$$\angle Y(j\omega) = \angle H(j\omega) + \angle U(j\omega)$$

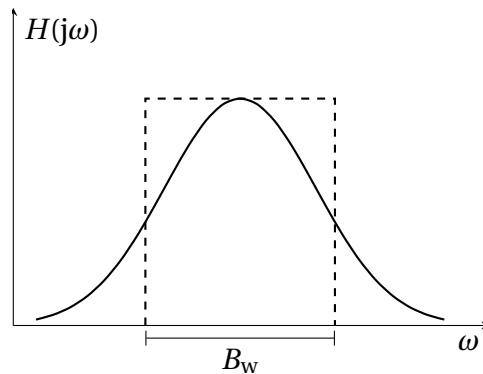
Avec de telles relations additives, il devient extrêmement aisé de construire graphiquement la réponse fréquentielle d'une cascade de systèmes à partir de leurs réponses fréquentielles individuelles.

L'échelle logarithmique typiquement utilisée est en unités de $20\log_{10}$, que l'on appelle *décibels* (dB). Ainsi, 0 dB correspond à un gain unité, 6 dB correspond à peu près à un gain double, 20 dB correspond à un gain de 10, et -20 dB correspond à une atténuation de 0.1.

Pour les systèmes continus, on utilise aussi couramment une échelle logarithmique pour les fréquences. Les graphes en décibels pour l'amplitude et en radians pour la phase en fonction de $\log_{10} \omega$ constituent les très célèbres *diagrammes de Bode*. Un diagramme de Bode typique est représenté à la Figure 9.1a.

On notera que le graphe n'est représenté que pour $\omega > 0$. Ceci résulte du fait que si la réponse impulsionnelle h est une fonction réelle, l'amplitude de sa transformée de Fourier est paire et la phase de sa transformée de Fourier est impaire (cf. Section 10.5.3).

Pour les systèmes discrets, on n'utilise pas une échelle logarithmique pour

FIGURE 9.2 – Caractérisation de la bande passante B_w d'un système.

les fréquences puisque l'intervalle $[-\pi, \pi]$ caractérise entièrement l'axe des fréquences. Un diagramme de Bode typique en discret est représenté à la Figure 9.1b.

On reviendra dans les sections suivantes sur l'utilité d'un diagramme de Bode et sa facilité de construction, en particulier pour des fonctions de transfert rationnelles de systèmes continus. Remarquons que la construction d'un diagramme de Bode peut se faire expérimentalement, sans connaître la fonction de transfert du système : en appliquant à l'entrée du système une sinusoïde de fréquence ω_0 et en la comparant à la sortie mesurée (après disparition des transitoires), on obtient facilement l'amplitude $|H(j\omega_0)|$ et la phase $\angle H(j\omega_0)$ du système à cette fréquence particulière. En répétant l'expérience pour différentes fréquences, on obtient différents points du diagramme de Bode.

9.3 Bande passante et constante de temps

De même que nous avons défini la constante de temps d'un système comme caractéristique qualitative de base de sa réponse impulsionnelle, on peut définir la *bande passante* d'un système comme caractéristique qualitative de base de sa réponse fréquentielle. Cette caractéristique traduit le fait que le système atténue les fréquences pour lesquelles $|H(j\omega)|$ est petit et laisse « passer » les fréquences ω pour lesquelles $|H(j\omega)|$ est grand. En accord avec notre définition de la constante de temps (Section 3.5), on peut par exemple définir la bande passante comme la largeur d'un rectangle dont l'aire est égale à l'intégrale de $|H(j\omega)|$ (cf. Figure 9.2) :

$$B_w = \frac{\int_{-\infty}^{\infty} |H(j\omega)| d\omega}{\max H}.$$

Il est très instructif de comparer la définition de bande passante à la

définition de constante de temps (Section 3.5) :

$$T_h = \frac{\int_{-\infty}^{\infty} h(t) dt}{\max h}.$$

Pour simplifier, considérons un système dont la réponse fréquentielle est réelle et positive : $|H(j\omega)| = H(j\omega)$. On a alors

$$B_w = \frac{2\pi h(0)}{\max H}, \quad T_h = \frac{H(0)}{\max h}.$$

Si en outre, $H(0) = \max H$ et $h(0) = \max h$, on en déduit la relation

$$B_w T_h = 2\pi \quad (9.3.1)$$

Cette relation qualitative exprime un compromis fondamental dans toute application de synthèse de système : il n'est pas possible d'assigner indépendamment la bande passante et le temps de réponse. Par exemple, si l'on veut que le système puisse réagir rapidement, il doit avoir une constante de temps petite, ce qui, par la relation (9.3.1), implique nécessairement une bande passante élevée. Dans les applications de régulation, c'est un compromis fondamental de robustesse : on veut que le système puisse suivre « rapidement » le signal de consigne, tout en limitant la bande passante du système pour réduire la sensibilité au bruit.

9.4 Réponses d'un système continu du premier ordre

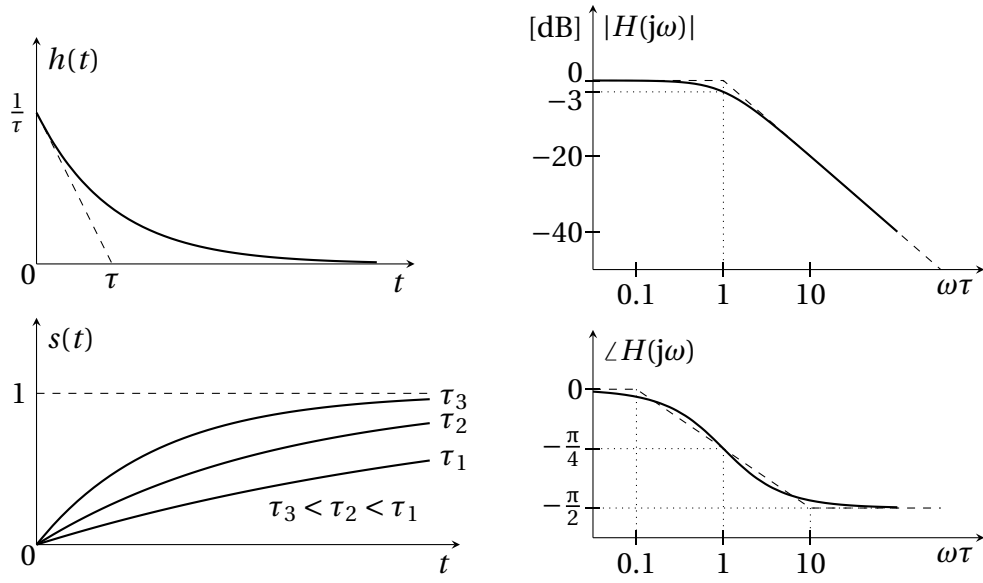
Un système du premier ordre est le premier modèle utilisé pour capturer les caractéristiques dynamiques les plus grossières d'un système : un temps de réponse et une bande passante. Dans les applications simples de régulation de procédés industriels, un tel modèle peut être suffisant et la simple prise en compte du temps de réponse dans la synthèse du régulateur représente un gain considérable par rapport à un modèle statique.

Un modèle du premier ordre est décrit par le système différentiel

$$\tau \dot{y} + y = u. \quad (9.4.1)$$

La constante positive τ est directement identifiée à la constante de temps du système (cfr. Section 3.5). La réponse fréquentielle est

$$H(j\omega) = \frac{1}{j\omega\tau + 1} \quad (9.4.2)$$



(a) Réponses impulsionnelle et indicielle.

(b) Diagramme de Bode.

FIGURE 9.3 – Caractéristiques graphiques d'un système continu du premier ordre.

et la réponse impulsionnelle est

$$h(t) = \frac{1}{\tau} e^{-\frac{t}{\tau}} \mathbb{I}(t).$$

La réponse impulsionnelle et la réponse indicielle $s(t) = (1 - e^{-\frac{t}{\tau}}) \mathbb{I}(t)$ sont représentées à la Figure 9.3a. Lorsque la constante de temps tend vers zéro, la réponse indicielle tend vers celle d'un système statique. On peut également noter que la réponse indicielle d'un système du premier ordre n'a pas de dépassement (et donc pas d'oscillations). La Figure 9.3b représente le diagramme de Bode d'un système du premier ordre. Le graphe est très facilement esquissé à partir de son approximation asymptotique. Pour l'amplitude, on déduit de l'équation (9.4.2) la relation

$$20 \log_{10} |H(j\omega)| = -10 \log_{10} ((\omega\tau)^2 + 1).$$

On peut observer que pour $\omega\tau \ll 1$, l'amplitude (logarithmique) vaut presque zéro, tandis que pour $\omega\tau \gg 1$, la relation est presque linéaire en la variable $\log_{10}(\omega)$:

$$20 \log_{10} |H(j\omega)| \approx 0, \quad \omega\tau \ll 1$$

et

$$20 \log_{10} |H(j\omega)| \approx -20 \log_{10}(\omega) - 20 \log_{10}(\tau), \quad \omega\tau \gg 1.$$

On en conclut que pour un système du premier ordre, les asymptotes hautes

et basses fréquences de l'amplitude logarithmique sont des lignes droites : l'abscisse 0 dB pour les basses fréquences, et une diminution de 20 dB par décade pour les hautes fréquences. L'intersection des deux asymptotes a lieu à la fréquence $\omega = \frac{1}{\tau}$ et il est très facile d'esquisser la courbe d'amplitude à partir de cette approximation. On peut observer la caractéristique passe-bas du système du premier ordre : le système est passant jusqu'à la fréquence $\omega = \frac{1}{\tau}$ et atténue les fréquences supérieures. Néanmoins, la transition est très douce et peut être insuffisante pour une application de filtrage.

La courbe de phase peut également être facilement esquissée :

$$\angle H(j\omega) = -\tan^{-1}(\omega\tau) \approx \begin{cases} 0, & \omega\tau \leq 0.1, \\ -\frac{\pi}{4}[\log_{10}(\omega\tau) + 1], & 0.1 \leq \omega\tau \leq 10, \\ -\frac{\pi}{2}, & \omega\tau \geq 10. \end{cases} \quad (9.4.3)$$

Cette approximation est linéaire sur l'intervalle $[0.1/\tau, 10/\tau]$. Dans cet intervalle, la phase du système est approximée par la tangente à la courbe au point $\omega = \frac{1}{\tau}$ et équivaut à l'effet d'un retard pur de $-\frac{\pi}{4}$. On retiendra qu'un système du premier ordre provoque un *retard de phase* qui n'excède pas 90 degrés.

9.5 Réponse d'un système du deuxième ordre

Les modèles du deuxième ordre sont typiques de systèmes soumis à une force de rappel, comme un ressort ou un circuit RLC, qui peut causer des oscillations dans la réponse. L'équation différentielle du deuxième ordre est mise sous la forme

$$\ddot{y} + 2\zeta\omega_n\dot{y} + \omega_n^2 y = \omega_n^2 u. \quad (9.5.1)$$

Par exemple, une masse ponctuelle m soumise à l'effet d'un ressort linéaire de raideur k , un frottement linéaire visqueux b , et une excitation extérieure f donne le système du deuxième ordre

$$m\ddot{y} + b\dot{y} + ky = f$$

et on identifie

$$\omega_n = \sqrt{\frac{k}{m}}, \quad \zeta = \frac{b}{2\sqrt{km}}, \quad u = \frac{1}{k}f.$$

La réponse fréquentielle du système (9.5.1) est

$$H(j\omega) = \frac{\omega_n^2}{(j\omega)^2 + 2\zeta\omega_n(j\omega) + \omega_n^2}. \quad (9.5.2)$$

Les deux pôles du système sont

$$c_{1,2} = -\zeta\omega_n \pm \omega_n\sqrt{\zeta^2 - 1}.$$

En utilisant un développement en fractions simples de la fonction de transfert et en appliquant la transformée inverse, on obtient la réponse impulsionnelle

$$h(t) = \frac{\omega_n}{2\sqrt{\zeta^2 - 1}} (e^{c_1 t} - e^{c_2 t}) \mathbb{I}(t), \quad \zeta \neq 1, \quad (9.5.3)$$

ou, dans le cas de deux pôles confondus,

$$h(t) = \omega_n^2 t e^{-\omega_n t} \mathbb{I}(t), \quad \zeta = 1.$$

On peut observer que h est une fonction de la variable $\omega_n t$. De même, la réponse fréquentielle peut se réécrire comme

$$H(j\omega) = \frac{1}{(j\frac{\omega}{\omega_n})^2 + 2\zeta j(\frac{\omega}{\omega_n}) + 1}.$$

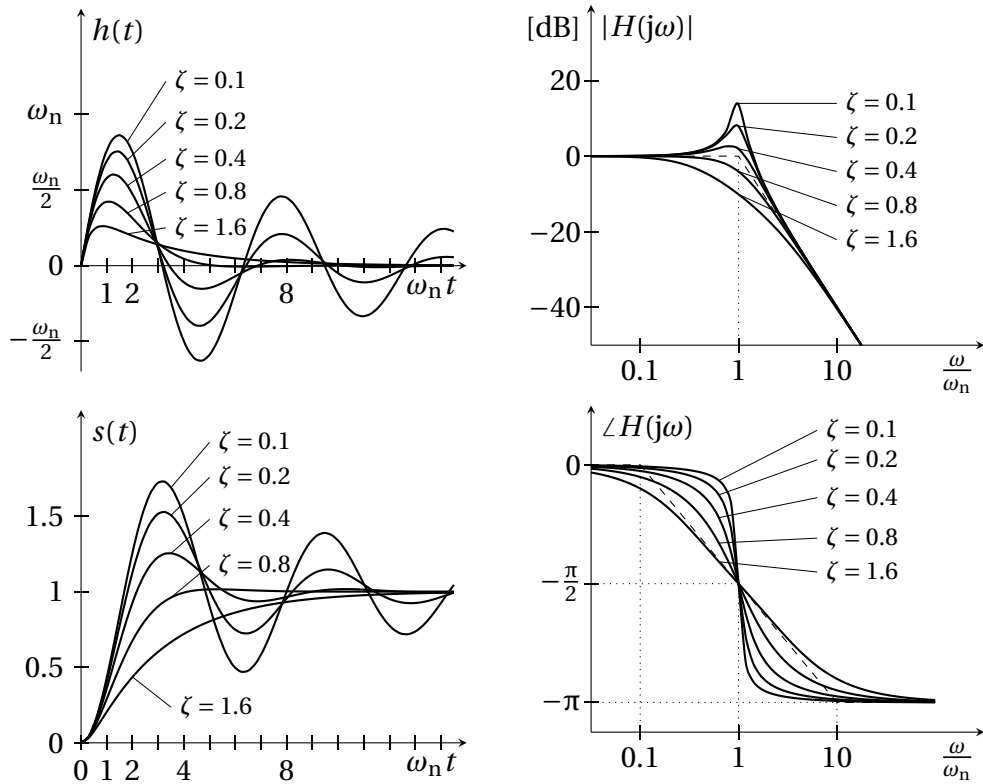
Une variation du paramètre ω_n , appelé *fréquence propre* du système, revient donc à un changement d'échelle temporel et fréquentiel. Si on définit un nouveau « temps » $\tau = \omega_n t$, l'équation différentielle elle-même peut être normalisée, c'est-à-dire rendue indépendante de la fréquence propre :

$$\frac{d^2 y(\tau)}{d\tau^2} + 2\zeta \frac{dy(\tau)}{d\tau} + y(\tau) = u(\tau). \quad (9.5.4)$$

Le paramètre ζ est appelé *facteur d'amortissement* du système. Cette terminologie s'explique en réécrivant la réponse impulsionnelle (9.5.3) comme

$$\begin{aligned} h(t) &= \frac{\omega_n e^{-\zeta\omega_n t}}{2j\sqrt{1-\zeta^2}} \left(\exp(j(\omega_n\sqrt{1-\zeta^2})t) - \exp(-j(\omega_n\sqrt{1-\zeta^2})t) \right) \mathbb{I}(t) \\ &= \frac{\omega_n e^{-\zeta\omega_n t}}{\sqrt{1-\zeta^2}} \sin(\omega_n\sqrt{1-\zeta^2}t) \mathbb{I}(t), \quad 0 < \zeta < 1 \end{aligned} \quad (9.5.5)$$

Ainsi, pour $0 < \zeta < 1$, la réponse impulsionnelle d'un système du deuxième ordre est une sinusoïde amortie. Pour $\zeta > 1$, les deux pôles sont réels et la réponse impulsionnelle est la somme de deux exponentielles décroissantes. La Figure 9.4a illustre l'allure des réponses impulsionnelle et indicielle pour différentes valeurs du paramètre ζ . Pour $\zeta < 1$, on observe le caractère oscillatoire de la réponse indicielle et son *dépassement* (elle atteint des valeurs supérieures à sa valeur finale durant son transitoire). On observe également le compromis entre le temps de réponse et le temps nécessaire pour atteindre sa valeur de régime (*settling time*).



(a) Réponses impulsionnelle et indicielle. (b) Diagrammes d'amplitude et de phase.

FIGURE 9.4 – Caractéristiques graphiques d'un système continu du deuxième ordre pour différentes valeurs de ζ .

La Figure 9.4b représente le diagramme de Bode de la réponse fréquentielle (9.5.2) pour différentes valeurs de ζ . Comme pour le système du premier ordre, on a des asymptotes linéaires pour le diagramme d'amplitude. On déduit de l'expression

$$20 \log_{10} |H(j\omega)| = -10 \log_{10} \left(\left(1 - \left(\frac{\omega}{\omega_n} \right)^2 \right)^2 + 4\zeta^2 \left(\frac{\omega}{\omega_n} \right)^2 \right) \quad (9.5.6)$$

les asymptotes

$$20 \log_{10} |H(j\omega)| \approx \begin{cases} 0, & \omega \ll \omega_n, \\ -40 \log_{10} \omega + 40 \log_{10} \omega_n, & \omega \gg \omega_n. \end{cases} \quad (9.5.7)$$

L'asymptote des basses fréquences est donc la ligne des 0 dB, tandis que l'asymptote des hautes fréquences a une pente de -40 dB par décade. Les deux asymptotes s'intersectent au point $\omega = \omega_n$.

Une approximation linéaire par morceaux est également obtenue pour la

phase : on déduit de l'expression

$$\angle H(j\omega) = -\tan^{-1} \left(\frac{2\zeta(\omega/\omega_n)}{1 - (\omega/\omega_n)^2} \right) \quad (9.5.8)$$

l'approximation

$$\angle H(j\omega) \approx \begin{cases} 0, & \omega \leq 0.1\omega_n, \\ -\frac{\pi}{2} (\log_{10}(\omega/\omega_n) + 1), & 0.1\omega_n \leq \omega \leq 10\omega_n, \\ -\pi, & \omega \geq 10\omega_n. \end{cases} \quad (9.5.9)$$

Les approximations ne dépendent pas de ζ et il est donc important de corriger ces approximations afin de mieux approcher le graphe réel, surtout lorsque l'amortissement ζ est faible. En particulier, il est important dans les applications de bien estimer le pic observé dans le diagramme d'amplitude pour les faibles valeurs de ζ . Pour $\zeta < \sqrt{2}/2 \approx 0.7$, on calcule facilement que la fonction $|H(j\omega)|$ atteint un maximum à la fréquence

$$\omega_{\max} = \omega_n \sqrt{1 - 2\zeta^2},$$

et que le maximum vaut

$$|H(j\omega_{\max})| = \frac{1}{2\zeta \sqrt{1 - \zeta^2}}$$

Pour $\zeta > 0.707$, la décroissance de la fonction $|H(j\omega)|$ est monotone et il n'y a donc plus de pic.

9.6 Fonctions de transfert rationnelles

Les diagrammes de Bode des systèmes du premier et deuxième ordre constituent les blocs de base pour la construction du diagramme de Bode de n'importe quelle réponse fréquentielle $H(j\omega) = N(j\omega)/D(j\omega)$. En factorisant les numérateur et dénominateur en produits de facteurs du premier et deuxième ordre, il suffit d'additionner les diagrammes (logarithmiques) de chaque terme pris séparément. Chaque terme du numérateur est donc de la forme

$$H(j\omega) = 1 + j\omega\tau \quad \text{ou} \quad H(j\omega) = 1 + 2\zeta \left(\frac{j\omega}{\omega_n} \right) + \left(\frac{j\omega}{\omega_n} \right)^2.$$

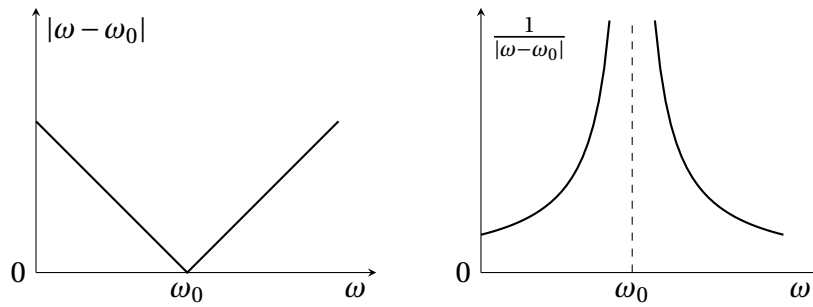


FIGURE 9.5 – Diagrammes d’amplitude pour la fonction de transfert $s - j\omega_0$ et la fonction de transfert $\frac{1}{s - j\omega_0}$.

On déduit les diagrammes de ces numérateurs à partir de leurs correspondants (9.4.2) et (9.5.2) en utilisant les relations

$$20\log_{10}|H(j\omega)| = -20\log_{10}\left|\frac{1}{H(j\omega)}\right| \quad \text{et} \quad \angle(H(j\omega)) = -\angle\frac{1}{H(j\omega)}.$$

9.6.1 Effet d’un pôle et d’un zéro sur la réponse fréquentielle

La synthèse d’un filtre au moyen d’une fonction de transfert rationnelle consiste à « placer » les pôles et les zéros de la fonction de transfert de manière à obtenir une réponse fréquentielle passante/bloquante aux fréquences souhaitées. Mathématiquement, un zéro $s = j\omega_0$ sur l’axe imaginaire annule la fréquence ω_0 et un pôle $s = j\omega_0$ sur l’axe imaginaire correspond à une amplification infinie de cette même fréquence (cf. Figure 9.5).

En pratique, la contrainte de stabilité du filtre impose de placer les pôles dans le demi-plan de gauche, et la sélectivité fréquentielle du pôle diminue lorsque l’on s’éloigne de l’axe imaginaire. La réalisation du filtre impose également de placer les pôles complexes par paire conjuguée (pour que les coefficients de la fonction de transfert soient réels), et d’assurer un nombre de pôles égal ou supérieur au nombre de zéros.

9.6.2 Filtres de Butterworth

Un filtre passe-bas a un gain maximum à la fréquence nulle, ce qui suggère de placer un pôle sur l’axe réel : c’est le système du premier ordre décrit dans la Section 9.4, dont l’amplitude décroît de manière monotone. Pour améliorer la sélectivité du filtre sur un intervalle $[0, \omega_c]$, il faut rajouter des pôles ayant une partie imaginaire comprise entre 0 et ω_c , et d’autant plus près de l’axe imaginaire que l’on veut augmenter leur sélectivité. On peut montrer qu’une caractéristique d’amplitude optimale – i.e. se rapprochant le plus de la caractéristique idéale passe-bas – pour un nombre de pôles N

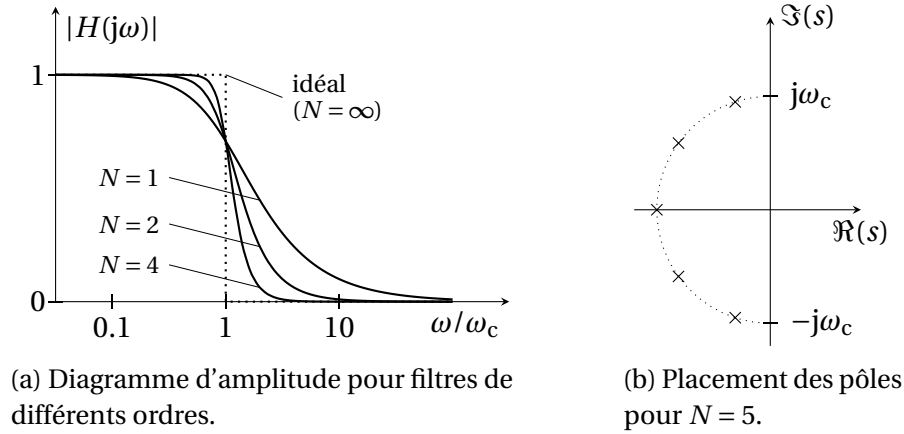


FIGURE 9.6 – Filtres de Butterworth.

est obtenue en plaçant les pôles à distance égale sur un demi-cercle centré à l'origine et de rayon ω_c . On obtient ainsi une famille de filtres appelés filtres de Butterworth. La Figure 9.6 illustre le diagramme d'amplitude obtenu pour différentes valeurs de N et le placement des pôles pour un filtre de Butterworth d'ordre 5.

Pour la synthèse d'un filtre passe-bande centré sur la fréquence ω_0 , on applique le même principe mais les pôles sont placés sur un demi-cercle centré au point $s = j\omega_0$ plutôt qu'à l'origine.

9.7 Réponses d'un système discret du premier ordre

Un système *discret* du premier ordre est mis sous la forme

$$y[n] - ay[n-1] = u[n]$$

où $|a| < 1$. La réponse fréquentielle de ce système est

$$H(e^{j\omega}) = \frac{1}{1 - ae^{-j\omega}}$$

et sa réponse impulsionnelle est

$$h[n] = a^n \mathbb{I}[n].$$

La réponse indicielle

$$s[n] = h * \mathbb{I}[n] = \frac{1 - a^{n+1}}{1 - a} \mathbb{I}[n]$$

est représentée à la Figure 9.7 pour différentes valeurs de a . La valeur absolue de a joue un rôle similaire à la constante de temps dans le cas continu. Une

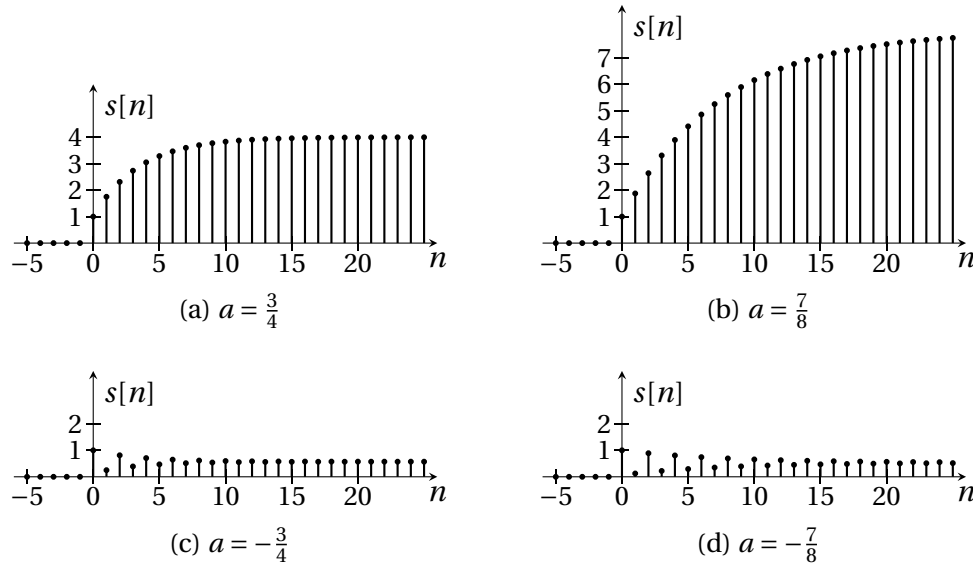


FIGURE 9.7 – Réponses indicielles d'un système discret du premier ordre pour différentes valeurs de a .

différence notoire avec le cas continu est que la réponse d'un système discret du premier ordre peut être oscillatoire (lorsque $a < 0$).

La courbe d'amplitude et la courbe de phase sont données par les relations

$$|H(e^{j\omega})| = \frac{1}{(1 + a^2 - 2a \cos \omega)^{1/2}}$$

et

$$\angle H(e^{j\omega}) = -\tan^{-1} \frac{a \sin \omega}{1 - a \cos \omega}.$$

Il n'existe pas de règles simples pour la construction de ces courbes comme dans le cas continu. À titre indicatif, la Figure 9.8 donne les courbes d'amplitude et de phase pour différentes valeurs de a . On peut observer que pour $a > 0$, la caractéristique est *passé-bas*, tandis que pour $a < 0$, la caractéristique est *passé-haut*. En outre, la courbe d'amplitude est relativement plate pour $|a|$ petit, tandis que les pics sont accentués pour $|a|$ proche de un.

9.8 Exemple d'analyse fréquentielle et temporelle d'un modèle LTI

Les diagrammes temporels (réponses impulsionnelle et indicielle) et fréquentiels (diagrammes de Bode) d'un système constituent les informations de base pour l'analyse et la synthèse de systèmes dans divers domaines d'application. Dans chaque discipline (théorie du filtrage, traitement du signal, asservissement,...), une expertise spécifique est requise pour traduire les spécifications temporelles et fréquentielles propres à l'application consi-

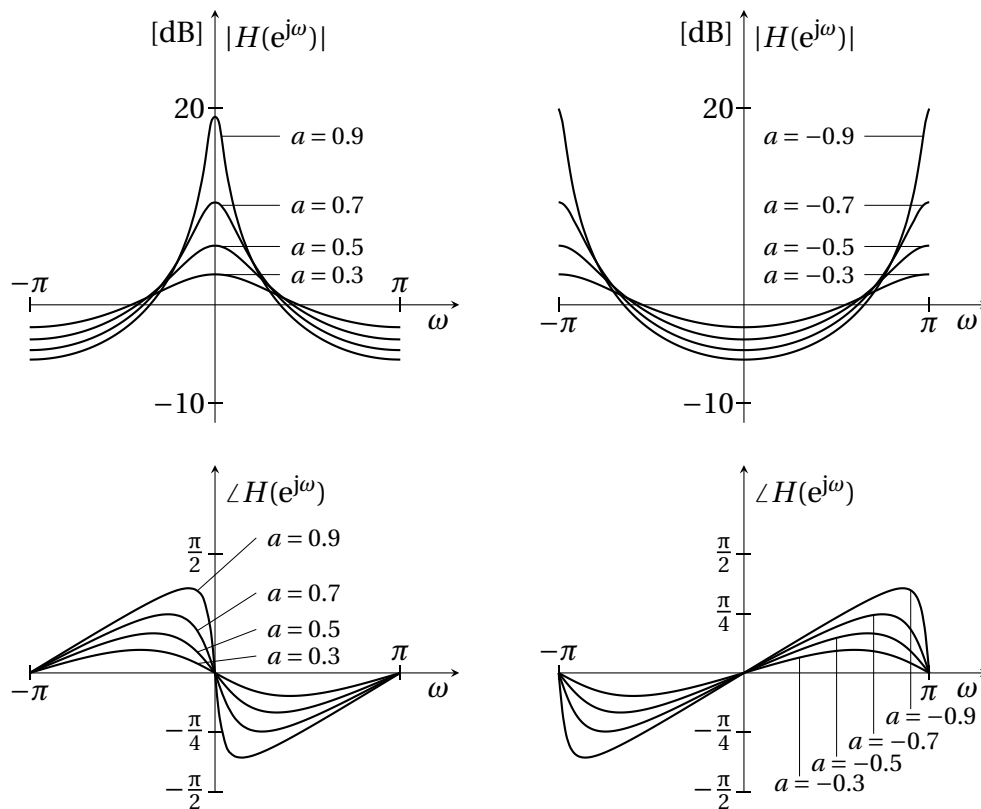


FIGURE 9.8 – Amplitude et phase de la réponse fréquentielle d'un système du premier ordre discret.

dérée en spécifications « graphiques » sur l'allure des diagrammes précités. Néanmoins, on retrouve dans chaque discipline les mêmes compromis fondamentaux entre temps de réponse et bande passante, temps de montée et dépassement,...

Pour un exemple illustratif nous référons à

Section 6.7.1, « Analysis of an Automobile Suspension System »,
Signals and Systems, 2nd edition, A. Oppenheim and A. Willsky,
 Prentice-Hall, 1997, pp. 473-482.

Cet exemple est simplifié à l'extrême mais illustre l'utilisation des diagrammes temporels et fréquentiels dans des applications spécifiques.

Chapitre 10

Des séries de Fourier aux transformées de Fourier

Ce chapitre montre que les transformées de Fourier définies au Chapitre 6 pour des signaux apériodiques correspondent à la limite des séries de Fourier vues au Chapitre 5 pour des signaux périodiques lorsque la période de ces derniers tend vers l'infini. On obtient ainsi quatre types de transformées de Fourier distinctes (deux pour les signaux périodiques, deux pour les signaux apériodiques). Le chapitre met en évidence la similarité des propriétés associées à ces quatre types de transformées et les propriétés de dualité (ou symétrie) qui les relient.

10.1 Signaux en temps continu

Les séries de Fourier permettent de représenter un signal défini sur un intervalle fini ou son extension périodique comme une combinaison linéaire de signaux harmoniques. Une des contributions fondamentales de Fourier fut de généraliser cette idée aux signaux *apériodiques*, un signal apériodique étant conçu comme un signal périodique de période infinie.

Pour illustrer le raisonnement de Fourier, nous allons chercher à calculer la « série de Fourier » du signal apériodique x représenté en trait plein à la Figure 10.1 et défini par

$$x(t) = \begin{cases} 1, & |t| < T_1, \\ 0, & |t| > T_1. \end{cases} \quad (10.1.1)$$

Le signal (10.1.1) est un signal apériodique défini sur l'entièreté de la droite réelle. Il diffère du signal (5.3.8) défini sur l'intervalle *fini* $(-\frac{T}{2}, \frac{T}{2})$ ainsi que de l'extension périodique $\tilde{x}$ de ce dernier. En revanche, le signal (10.1.1) peut être conçu comme la limite du signal $\tilde{x}$ pour $T \rightarrow \infty$.

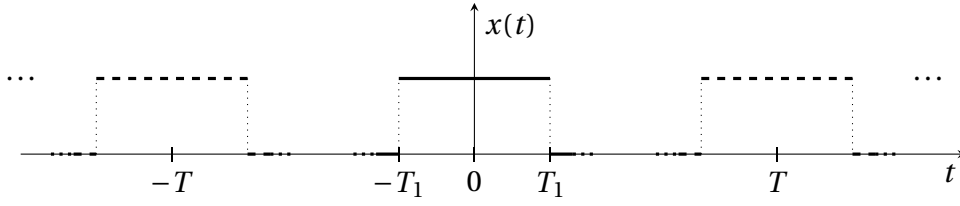


FIGURE 10.1 – Le signal apériodique x (en trait plein) et son extension T -périodique $\tilde{x}$ (en trait plein et en tirets).

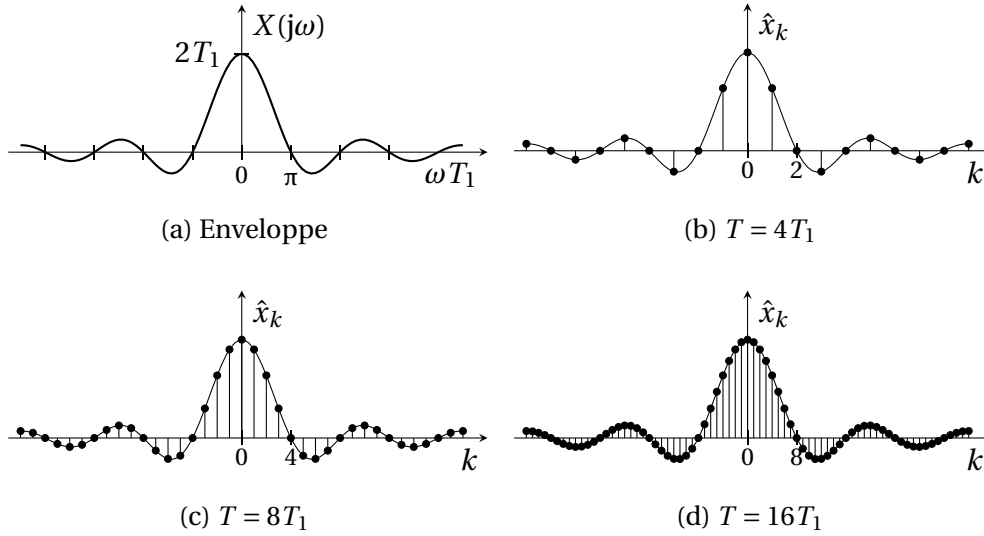


FIGURE 10.2 – Coefficients de Fourier de $\tilde{x}$ et leur enveloppe pour différentes valeurs de T .

Dans le Chapitre 5, nous avons déterminé la série de Fourier du signal $\tilde{x}$:

$$\tilde{x}(t) = \frac{1}{T} \sum_{k \in \mathbb{Z}} \hat{x}_k e^{jk \frac{2\pi}{T} t}, \quad (10.1.2)$$

$$\hat{x}_k = 2T_1 \operatorname{sinc}\left(k \frac{2\pi}{T} T_1\right) = \begin{cases} 2T_1, & k = 0, \\ 2 \frac{\sin(k \frac{2\pi}{T} T_1)}{k \frac{2\pi}{T}}, & k \neq 0. \end{cases} \quad (10.1.3)$$

L'équation (10.1.3) peut se réécrire sous la forme

$$\hat{x}_k = 2T_1 \operatorname{sinc}(\omega T_1) \Big|_{\omega=k\omega_0}, \quad \omega_0 = \frac{2\pi}{T}, \quad (10.1.4)$$

où le coefficient $\hat{x}_k$ est interprété comme la valeur d'une fonction enveloppe $2T_1 \operatorname{sinc}(\omega T_1)$ à la fréquence $\omega = k\omega_0$.

La fonction enveloppe ne dépend pas de la période T . Lorsque la période T augmente, ou, de manière équivalente, lorsque la fréquence fondamentale $\omega_0 = \frac{2\pi}{T}$ décroît, l'enveloppe est « échantillonnée » à une fréquence de plus en plus élevée, comme illustré à la Figure 10.2.

Lorsque $T \rightarrow \infty$, l'ensemble des coefficients de Fourier tendent vers la fonction enveloppe elle-même, qui n'est autre que la transformée de Fourier de x :

$$X(j\omega) = 2T_1 \operatorname{sinc}(\omega T_1). \quad (10.1.5)$$

Le raisonnement tenu ci-dessus s'étend à n'importe quel signal x de support $[-T_1, T_1]$.

L'extension périodique $\tilde{x}$, construite de manière à avoir $\tilde{x}(t) = x(t)$ sur la période $[-T, T]$ satisfait

$$\tilde{x}(t) = \frac{1}{T} \sum_{k \in \mathbb{Z}} \hat{x}_k e^{jk\omega_0 t} \quad (10.1.6)$$

$$\hat{x}_k = \int_{-T/2}^{T/2} \tilde{x}(t) e^{-jk\omega_0 t} dt \quad (10.1.7)$$

où $\omega_0 = \frac{2\pi}{T}$. Puisque $\tilde{x}(t) = x(t)$ sur une période et que $x(t) = 0$ en dehors de l'intervalle $[-\frac{T}{2}, \frac{T}{2}]$, l'équation (10.1.7) peut se réécrire comme

$$\hat{x}_k = \int_{-T/2}^{T/2} x(t) e^{-jk\omega_0 t} dt = \int_{-\infty}^{\infty} x(t) e^{-jk\omega_0 t} dt. \quad (10.1.8)$$

En définissant la fonction enveloppe X par

$$X(j\omega) = \int_{-\infty}^{\infty} x(t) e^{-j\omega t} dt \quad (10.1.9)$$

on obtient la relation

$$\hat{x}(k) = X(jk\omega_0).$$

Pour l'équation (10.1.6), on a donc

$$\tilde{x}(t) = \frac{1}{T} \sum_{k \in \mathbb{Z}} X(jk\omega_0) e^{jk\omega_0 t} = \frac{1}{2\pi} \sum_{k \in \mathbb{Z}} X(jk\omega_0) e^{jk\omega_0 t} \omega_0. \quad (10.1.10)$$

Lorsque $T \rightarrow \infty$, $\tilde{x}(t)$ tend vers $x(t)$ et (10.1.10) doit tendre vers une représentation du signal x . Chaque terme de la somme est l'aire d'un rectangle de largeur ω_0 et de hauteur $X(jk\omega_0) e^{jk\omega_0 t}$. Lorsque $\omega_0 \rightarrow 0$, le produit $X(jk\omega_0) e^{jk\omega_0 t} \omega_0$ tend vers $X(j\omega) e^{j\omega t} d\omega$ et la somme tend vers une intégrale (cf. Figure 10.3). A la limite, l'équation (10.1.10) devient

$$x(t) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} X(j\omega) e^{j\omega t} d\omega. \quad (10.1.11)$$

Les équations (10.1.9) et (10.1.11) correspondent bien aux définitions de la transformée de Fourier et de la transformée de Fourier inverse vues au Chapitre 6.

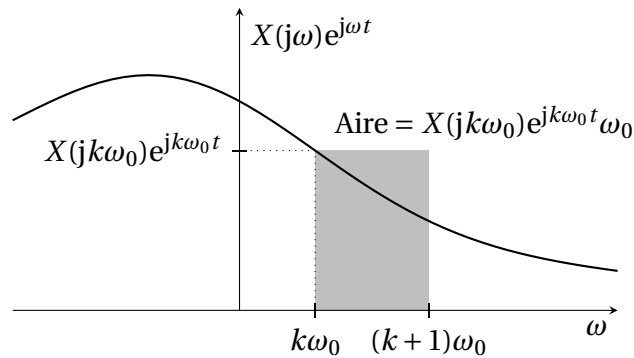


FIGURE 10.3 – Interprétation graphique de (10.1.11).

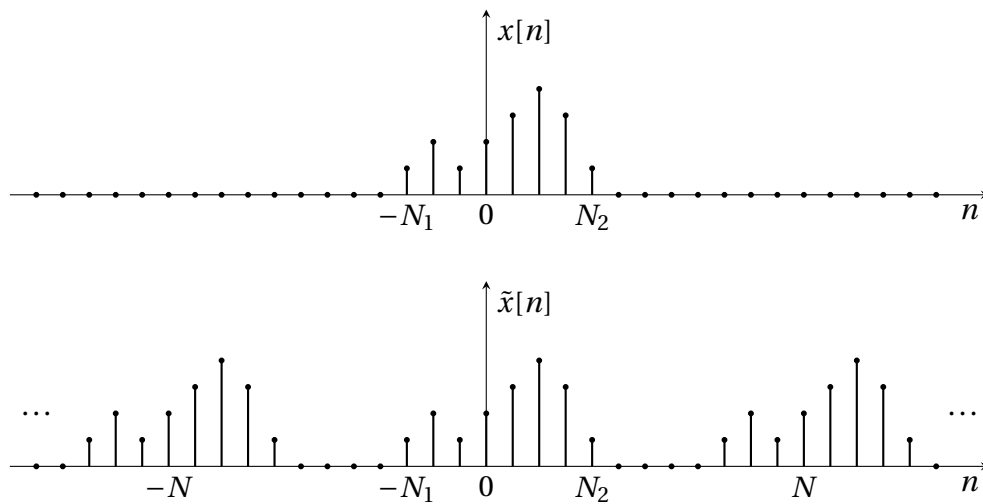


FIGURE 10.4 – Signal à support borné et son prolongement périodique.

Alors qu'un signal périodique continu avait une représentation de Fourier discrète (coefficients de Fourier dans $\mathbb{Z}$), un signal apériodique continu a une représentation de Fourier continue

$$X(j\omega), \quad \omega \in \mathbb{R},$$

d'où la terminologie de *transformée de Fourier continue-continue (CC)*.

10.2 De la série de Fourier DD à la transformée de Fourier DC

La représentation de Fourier d'un signal apériodique discret $x[n]$ s'obtient, comme dans le cas continu, en concevant le signal apériodique comme la limite (pour $N \rightarrow \infty$) d'une suite de signaux périodiques de période N .

Soit un signal x de support $[-N_1..N_2]$ et un prolongement périodique $\tilde{x}$ de période $N > N_1 + N_2$, comme illustré à la Figure 10.4.

La série de Fourier de $\tilde{x}$ vaut

$$\tilde{x}[n] = \frac{1}{N} \sum_{k=0}^{N-1} \hat{x}_k e^{jk \frac{2\pi}{N} n}, \quad (10.2.1)$$

$$\hat{x}_k = \sum_{n=-N_1}^{N_2} \tilde{x}[n] e^{-jk \frac{2\pi}{N} n}. \quad (10.2.2)$$

En remplaçant $\tilde{x}[n]$ par $x[n]$ on obtient

$$\hat{x}_k = \sum_{n=-N_1}^{N_2} x[n] e^{-jk \frac{2\pi}{N} n} = \sum_{n \in \mathbb{Z}} x[n] e^{-jk \frac{2\pi}{N} n}. \quad (10.2.3)$$

En définissant la fonction enveloppe

$$X(e^{jk\omega}) = \sum_{n \in \mathbb{Z}} x[n] e^{-jk\omega n} \quad (10.2.4)$$

on obtient les relations

$$\hat{x}_k = X(e^{jk\omega_0}), \quad \omega_0 = \frac{2\pi}{N}, \quad (10.2.5)$$

$$\tilde{x}[n] = \frac{1}{N} \sum_{k=0}^{N-1} X(e^{jk\omega_0}) e^{jk\omega_0 n} \quad (10.2.6)$$

$$= \frac{1}{2\pi} \sum_{k=0}^{N-1} X(e^{jk\omega_0}) e^{jk\omega_0 n} \omega_0. \quad (10.2.7)$$

De manière analogue au cas continu, $\tilde{x}[n]$ tend vers $x[n]$ lorsque $N \rightarrow \infty$. La somme (10.2.7) devient une intégration, la largeur de l'intervalle d'intégration valant $\omega_0 N = 2\pi$ pour toute valeur de N . On obtient donc les relations

$$x[n] = \frac{1}{2\pi} \int_{2\pi} X(e^{j\omega}) e^{j\omega n} d\omega,$$

$$X(e^{j\omega}) = \sum_{n \in \mathbb{Z}} x[n] e^{-j\omega n},$$

qui correspondent aux définitions de transformée et de transformée inverse données au Chapitre 6.

Alors que la série de Fourier d'un signal périodique discret de période N était discrète (N coefficients de Fourier), la transformée de Fourier d'un signal apériodique discret est cette fois un signal continu (périodique), d'où la terminologie de *transformée de Fourier discrète-continue (DC)*.

10.3 Transformée de Fourier de signaux périodiques continus

Dans les deux sections précédentes, nous avons considéré des signaux à support compact. Bien que ce ne soit pas établi rigoureusement dans le cadre de ce cours, le processus limite utilisé pour la définition de la transformée de Fourier et de son inverse peut être étendu aux signaux de $L_2(\mathbb{R})$. Par contre, la transformée d'un signal $x \notin L_2(\mathbb{R})$, comme par exemple un signal périodique non nul, n'est en général pas défini au sens des fonctions usuelles. L'utilisation d'impulsions de Dirac permet néanmoins une extension naturelle des transformées de Fourier aux signaux périodiques.

Pour un signal harmonique $x(t) = e^{j\omega_0 t}$, la définition de $X(j\omega)$ est suggérée par la formule de transformée inverse

$$x(t) = e^{j\omega_0 t} = \frac{1}{2\pi} \int_{-\infty}^{\infty} X(j\omega) e^{j\omega t} d\omega. \quad (10.3.1)$$

L'identification des deux membres donne

$$X(j\omega) = 2\pi\delta(\omega - \omega_0). \quad (10.3.2)$$

La transformée de Fourier du signal harmonique $e^{j\omega_0 t}$ est donc une impulsion de Dirac à la fréquence ω_0 . Ce résultat n'est pas surprenant dans la mesure où il exprime que toute l'énergie du signal $x(t) = e^{j\omega_0 t}$ est concentrée à la seule fréquence ω_0 .

Pour calculer la transformée de Fourier d'un signal périodique quelconque, il suffit de connaître sa série de Fourier

$$x(t) = \frac{1}{T} \sum_{k \in \mathbb{Z}} \hat{x}_k e^{jk\omega_0 t}, \quad (10.3.3)$$

d'où l'on déduit

$$X(j\omega) = \frac{2\pi}{T} \sum_{k \in \mathbb{Z}} \hat{x}_k \delta(\omega - k\omega_0). \quad (10.3.4)$$

La transformée de Fourier d'un signal périodique quelconque est donc un train d'impulsions aux fréquences harmoniques $k\omega_0$, $k \in \mathbb{Z}$, et d'amplitudes proportionnelles aux coefficients de Fourier $\hat{x}_k$.

Un cas particulier très utile dans la théorie de l'échantillonnage concerne le train d'impulsions périodique

$$x(t) = \sum_{k \in \mathbb{Z}} \delta(t - kT), \quad (10.3.5)$$

dont les coefficients de Fourier valent

$$\hat{x}_k = \int_{-T/2}^{T/2} \delta(t) e^{-jk\omega_0 t} dt = 1. \quad (10.3.6)$$

La transformée de Fourier donne dans ce cas

$$X(j\omega) = \frac{2\pi}{T} \sum_{k \in \mathbb{Z}} \delta(\omega - k\omega_0). \quad (10.3.7)$$

On en conclut que la transformée de Fourier d'un train périodique d'impulsions de période T dans le domaine temporel est un train d'impulsions de période $2\pi/T$ dans le domaine fréquentiel.

10.4 Transformée de Fourier de signaux périodiques discrets

L'analogie avec le cas continu suggère que la transformée de Fourier du signal harmonique $x(n) = e^{j\omega_0 n}$ est une impulsion à la fréquence ω_0 . Le signal $X(e^{j\omega})$ doit toutefois être périodique de période 2π , ce qui suggère de répéter l'impulsion ω_0 aux fréquences $\omega_0 + k2\pi$, $k \in \mathbb{Z}$, pour obtenir

$$n \mapsto e^{j\omega_0 n} \xleftrightarrow{\mathcal{F}_{\text{CD}}} \omega \mapsto \sum_{k \in \mathbb{Z}} 2\pi \delta(\omega - \omega_0 - k2\pi). \quad (10.4.1)$$

La formule de transformée inverse

$$\frac{1}{2\pi} \int_{2\pi} X(e^{j\omega}) e^{j\omega n} d\omega = \int_0^{2\pi} \sum_{k \in \mathbb{Z}} \delta(\omega - \omega_0 - k2\pi) e^{j\omega n} d\omega = e^{j\omega_0 n} \quad (10.4.2)$$

confirme le résultat.

Pour un signal périodique x , on a la série de Fourier

$$x[n] = \frac{1}{N} \sum_{k=0}^{N-1} \hat{x}_k e^{jk \frac{2\pi}{N} n} \quad (10.4.3)$$

et la transformée de Fourier

$$X(e^{j\omega}) = \frac{2\pi}{N} \sum_{k \in \mathbb{Z}} \hat{x}_k \delta(\omega - k \frac{2\pi}{N}). \quad (10.4.4)$$

Dans le cas particulier d'un train d'impulsions de période N ,

$$x[n] = \sum_{k \in \mathbb{Z}} \delta[n - kN] \quad (10.4.5)$$

les coefficients de Fourier valent

$$\hat{x}_k = \sum_{n=0}^{N-1} x[n] e^{-jk \frac{2\pi}{N} n} = 1 \quad (10.4.6)$$

et la transformée de Fourier vaut

$$X(e^{j\omega}) = \frac{2\pi}{N} \sum_{k \in \mathbb{Z}} \delta(\omega - k \frac{2\pi}{N}). \quad (10.4.7)$$

C'est l'analogie du résultat obtenu pour un train d'impulsions en continu.

10.5 Propriétés élémentaires des séries et transformées de Fourier

Si les propriétés principales des transformées de Fourier ont été vues au Chapitre 6, il importe d'avoir une vue unifiée de ces propriétés pour les quatre transformées vues dans ce chapitre.

- $x[n]$ (N -périodique discret) $\xleftrightarrow{\mathcal{F}_{DD}}$ $\hat{x}[k]$ (N -périodique discret)
- $x(t)$ (T -périodique continu) $\xleftrightarrow{\mathcal{F}_{CD}}$ $\hat{x}[k]$ (apériodique discret)
- $x(t)$ (apériodique continu) $\xleftrightarrow{\mathcal{F}_{CC}}$ $X(j\omega)$ (apériodique continu)
- $x[n]$ (apériodique discret) $\xleftrightarrow{\mathcal{F}_{DC}}$ $X(e^{j\omega})$ (2π -périodique continu)

On utilisera tout au long de cette section la notation $\hat{x}[k]$ pour le coefficient de Fourier $\hat{x}_k$. Outre la simplification notationnelle, on verra l'intérêt de concevoir l'ensemble des coefficients de Fourier comme un signal discret. L'ensemble des propriétés reprises ci-dessous se démontre de manière élémentaire à partir des définitions de transformées et de transformées inverses.

10.5.1 Transformées de signaux réfléchis et changement d'échelle

$$\begin{aligned} x[-n] &\xleftrightarrow{\mathcal{F}_{DD}} \hat{x}[-k] \\ x(-t) &\xleftrightarrow{\mathcal{F}_{CD}} \hat{x}[-k] \\ x(-t) &\xleftrightarrow{\mathcal{F}_{CC}} X(-j\omega) \\ x[-n] &\xleftrightarrow{\mathcal{F}_{DC}} X(e^{-j\omega}) \end{aligned}$$

On notera en particulier qu'un signal pair (resp. impair) a une transformée paire (resp. impaire).

Plus généralement, on a pour chaque $\alpha \in \mathbb{R}$ que

$$x(\alpha t) \xleftrightarrow{\mathcal{F}_{CC}} \frac{1}{|\alpha|} X\left(-j\frac{\omega}{\alpha}\right).$$

10.5.2 Décalage temporel et fréquentiel

Un décalage temporel correspond à une multiplication fréquentielle par un signal harmonique, c'est-à-dire un décalage de phase :

$$\begin{aligned} x[n - n_0] &\xleftrightarrow{\mathcal{F}_{\text{DD}}} e^{-jk\frac{2\pi}{N}n_0} \hat{x}[k] \\ x(t - t_0) &\xleftrightarrow{\mathcal{F}_{\text{CD}}} e^{-jk\frac{2\pi}{T}t_0} \hat{x}[k] \\ x(t - t_0) &\xleftrightarrow{\mathcal{F}_{\text{CC}}} e^{-j\omega t_0} X(j\omega) \\ x[n - n_0] &\xleftrightarrow{\mathcal{F}_{\text{DC}}} e^{-j\omega n_0} X(e^{j\omega}) \end{aligned}$$

Réciproquement, un décalage fréquentiel correspond à une multiplication temporelle par un signal harmonique :

$$\begin{aligned} e^{jM\frac{2\pi}{N}n} x[n] &\xleftrightarrow{\mathcal{F}_{\text{DD}}} \hat{x}[k - M] \\ e^{jM\frac{2\pi}{T}t} x(t) &\xleftrightarrow{\mathcal{F}_{\text{CD}}} \hat{x}[k - M] \\ e^{j\omega_0 t} x(t) &\xleftrightarrow{\mathcal{F}_{\text{CC}}} X(j(\omega - \omega_0)) \\ e^{j\omega_0 n} x[n] &\xleftrightarrow{\mathcal{F}_{\text{DC}}} X(e^{j(\omega - \omega_0)}) \end{aligned}$$

10.5.3 Signaux conjugués

$$\begin{aligned} \bar{x}[n] &\xleftrightarrow{\mathcal{F}_{\text{DD}}} \bar{\hat{x}}[-k] \\ \bar{x}(t) &\xleftrightarrow{\mathcal{F}_{\text{CD}}} \bar{\hat{x}}[-k] \\ \bar{x}(t) &\xleftrightarrow{\mathcal{F}_{\text{CC}}} \bar{X}(-j\omega) \\ \bar{x}[n] &\xleftrightarrow{\mathcal{F}_{\text{DC}}} \bar{X}(e^{-j\omega}) \end{aligned}$$

En combinant cette propriété avec la propriété de réflexion (Section 10.5.1), on obtient les conclusions suivantes pour un signal $x(\cdot)$ réel : la transformée d'un signal réel pair est un signal réel pair, tandis que la transformée d'un signal réel impair est un signal imaginaire impair. En particulier, la décomposition d'un signal réel en parties paire et impaire correspond à la décomposition de sa transformée en parties réelle et imaginaire.

$$x = x_p + x_i \xleftrightarrow{\mathcal{F}} X = \Re(X) + j\Im(X) \quad (x \text{ réel}).$$

10.5.4 Relation de Parseval

Le signal périodique discret $x[n]$ peut s'exprimer dans la base temporelle ($\delta_k[n] = \delta[n - k]$) ou harmonique ($v_k[n] = \frac{1}{N} e^{jk\frac{2\pi}{N}n}$) :

$$x = \sum_{k=0}^{N-1} x[k] \delta_k = \sum_{k=0}^{N-1} \hat{x}[k] v_k.$$

Comme les deux bases sont orthogonales, il vient

$$\|x\|_2^2 = \sum_{k=0}^{N-1} |x[k]|^2 \|\delta_k\|_2^2 = \sum_{k=0}^{N-1} |\hat{x}[k]|^2 \|v_k\|_2^2,$$

d'où l'on déduit la relation de Parseval

$$\sum_{k=0}^{N-1} |x[k]|^2 = \frac{1}{N} \sum_{k=0}^{N-1} |\hat{x}[k]|^2. \quad (10.5.1)$$

La relation de Parseval est analogue pour les autres transformées :

$$\begin{aligned} \int_T |x(t)|^2 dt &= \frac{1}{T} \sum_{k \in \mathbb{Z}} |\hat{x}[k]|^2, \\ \int_{-\infty}^{\infty} |x(t)|^2 dt &= \frac{1}{2\pi} \int_{-\infty}^{\infty} |X(j\omega)|^2 d\omega, \\ \sum_{k \in \mathbb{Z}} |x[k]|^2 &= \frac{1}{2\pi} \int_{2\pi} |X(e^{j\omega})|^2 d\omega. \end{aligned}$$

10.5.5 Dualité convolution-multiplication

$$\begin{array}{ll} x \odot y[n] \xleftrightarrow{\mathcal{F}_{DD}} \hat{x}[k] \hat{y}[k] & x[n] y[n] \xleftrightarrow{\mathcal{F}_{DD}} \frac{1}{N} \hat{x} \odot \hat{y}[k] \\ x \odot y(t) \xleftrightarrow{\mathcal{F}_{CD}} \hat{x}[k] \hat{y}[k] & x(t) y(t) \xleftrightarrow{\mathcal{F}_{CD}} \frac{1}{T} \hat{x} * \hat{y}[k] \\ x * y(t) \xleftrightarrow{\mathcal{F}_{CC}} X(j\omega) Y(j\omega) & x(t) y(t) \xleftrightarrow{\mathcal{F}_{CC}} \frac{1}{2\pi} X * Y(j\omega) \\ x * y[n] \xleftrightarrow{\mathcal{F}_{DC}} X(e^{j\omega}) Y(e^{j\omega}) & x[n] y[n] \xleftrightarrow{\mathcal{F}_{DC}} \frac{1}{2\pi} X \odot Y(e^{j\omega}) \end{array}$$

La dualité convolution-multiplication est évidemment une propriété centrale de la théorie des signaux et systèmes. Elle a été mise en évidence dans les Chapitres 6 et 5. Elle fait également apparaître la symétrie qui existe entre les différentes transformées.

10.5.6 Intégration-Différentiation

La différentiation dans le domaine temporel donne

$$\begin{aligned} x[n] - x[n-1] &\xleftrightarrow{\mathcal{F}_{DD}} (1 - e^{-jk\omega_0})\hat{x}[k], \\ \dot{x} &\xleftrightarrow{\mathcal{F}_{CD}} jk\omega_0\hat{x}[k], \\ \dot{x} &\xleftrightarrow{\mathcal{F}_{CC}} j\omega X(j\omega), \\ x[n] - x[n-1] &\xleftrightarrow{\mathcal{F}_{DC}} (1 - e^{-j\omega})X(e^{j\omega}). \end{aligned}$$

L'intégration dans le domaine temporel donne

$$\begin{aligned} \sum_{k=-\infty}^n x[k] &\xleftrightarrow{\mathcal{F}_{DD}} \begin{cases} \frac{1}{1-e^{-jk\omega_0}}\hat{x}[k], & k \neq lN, \\ 0, & k = lN, \end{cases} \quad (\text{si } \hat{x}[0] = 0), \\ \int_{-\infty}^t x(\tau) d\tau &\xleftrightarrow{\mathcal{F}_{CD}} \begin{cases} \frac{1}{jk\omega_0}\hat{x}[k], & k \neq 0, \\ 0, & k = 0, \end{cases} \quad (\text{si } \hat{x}[0] = 0), \\ \int_{-\infty}^t x(\tau) d\tau &\xleftrightarrow{\mathcal{F}_{CC}} \begin{cases} \frac{1}{j\omega}X(j\omega), & \omega \neq 0, \\ \pi X(0)\delta(\omega), & \omega = 0, \end{cases} \\ \sum_{k=-\infty}^n x[k] &\xleftrightarrow{\mathcal{F}_{DC}} \begin{cases} \frac{1}{1-e^{-j\omega}}X(e^{j\omega}), & \omega \neq 0, \\ \pi X(e^{j0})\sum_{k \in \mathbb{Z}} \delta(\omega - k2\pi), & \omega = 0. \end{cases} \end{aligned}$$

La transformée de Fourier du signal $\int_{-\infty}^t x(\tau) d\tau$ fait apparaître un nouveau terme par rapport à la transformée de Laplace. Ceci ne doit pas surprendre puisque

$$\int_{-\infty}^t x(\tau) d\tau = \mathbb{1} * x(t)$$

et que la transformée de l'échelon ne comprend pas l'axe imaginaire dans sa région de convergence ($H(s) = \frac{1}{s}, \Re(s) > 0$).

Il nous faut donc établir la propriété

$$\mathbb{1}(t) \xleftrightarrow{\mathcal{F}_{CC}} \frac{1}{j\omega} + \pi\delta(\omega). \quad (10.5.2)$$

La décomposition en partie paire et impaire de l'échelon donne $\mathbb{1} = \frac{1}{2} + \frac{1}{2}\text{sign}$ et l'identification avec les parties réelle et imaginaire de (10.5.2) donne

$$\frac{1}{2} \xleftrightarrow{\mathcal{F}_{CC}} \pi\delta \quad \text{et} \quad \frac{1}{2}\text{sign} \xleftrightarrow{\mathcal{F}_{CC}} \frac{1}{j\omega}. \quad (10.5.3)$$

La deuxième relation s'obtient en prenant la transformée de Fourier des deux membres dans l'égalité $\delta = \frac{1}{2}\text{sign}'$.

10.5.7 Dualité des transformées de Fourier

Les propriétés qui précèdent ont fait apparaître à maintes reprises une symétrie entre les différentes transformées. Une manifestation supplémentaire de cette dualité réside dans le fait que la double application de la transformée de Fourier d'un signal continu apériodique ou discret périodique redonne le signal initial réfléchi (à un facteur multiplicatif près, comme dans la relation de Parseval).

Ainsi, pour un signal continu apériodique, on a

$$x(t) \xleftrightarrow{\mathcal{F}_{CC}} X(j\omega) \xleftrightarrow{\mathcal{F}_{CC}} 2\pi x(-t).$$

Pour un signal discret N -périodique, on a de la même manière

$$x[n] \xleftrightarrow{\mathcal{F}_{DD}} \hat{x}[k] \xleftrightarrow{\mathcal{F}_{DD}} Nx[-n].$$

Les deux dernières transformées présentent une dualité « croisée » : la transformée d'un signal discret apériodique donne un signal continu 2π -périodique et la transformée de ce dernier redonne le signal discret apériodique initial (réfléchi) :

$$x[n] \xleftrightarrow{\mathcal{F}_{DC}} X(j\omega) \xleftrightarrow{\mathcal{F}_{CD}} 2\pi x[-n].$$

Le k -ième coefficient de Fourier du signal périodique $X(e^{j\omega})$ donne en effet

$$\begin{aligned} \widehat{X(e^{j\omega})}[n] &= \int_{2\pi} X(e^{j\omega}) e^{-jn\omega} d\omega \\ &= \int_{2\pi} \sum_{k \in \mathbb{Z}} x[k] e^{j(-k-n)\omega} d\omega \\ &= \sum_{k \in \mathbb{Z}} x[k] \int_{2\pi} e^{-j(k+n)\omega} d\omega = 2\pi x[-n]. \end{aligned}$$

Chapitre 11

Applications élémentaires des transformées de Fourier

La majorité des signaux traités aujourd'hui dans les applications d'ingénierie subissent un traitement numérique. L'objectif premier de ce traitement numérique variera d'une application à l'autre mais certaines opérations de base y seront invariablement associées : l'échantillonnage des signaux continus (conversion analogique-numérique), l'interpolation des signaux discrets (conversion numérique-analogique), et le fenêtrage des signaux de longueur infinie (c'est-à-dire leur approximation par un signal de longueur finie).

Ces opérations génèrent des erreurs d'approximation dans le traitement de signal escompté. L'analyse de ces erreurs d'approximation au moyen des transformées de Fourier constitue une belle illustration des propriétés développées dans le chapitre précédent.

Nous clôturons ce chapitre par un exemple élémentaire de traitement de signal intégrant ces diverses opérations.

11.1 Fenêtrages temporel et fréquentiel

11.1.1 Transformée d'un rectangle

Plusieurs opérations de traitement de signal correspondent à la multiplication d'un signal par une fenêtre rectangulaire dans le domaine temporel ou fréquentiel.

Le signal

$$x_1(t) = \begin{cases} 1, & |t| \leq T_1, \\ 0, & |t| > T_1, \end{cases} \quad (11.1.1)$$

a pour transformée de Fourier

$$X_1(j\omega) = \int_{-T_1}^{T_1} e^{-j\omega t} dt = 2T_1 \operatorname{sinc}(\omega T_1),$$

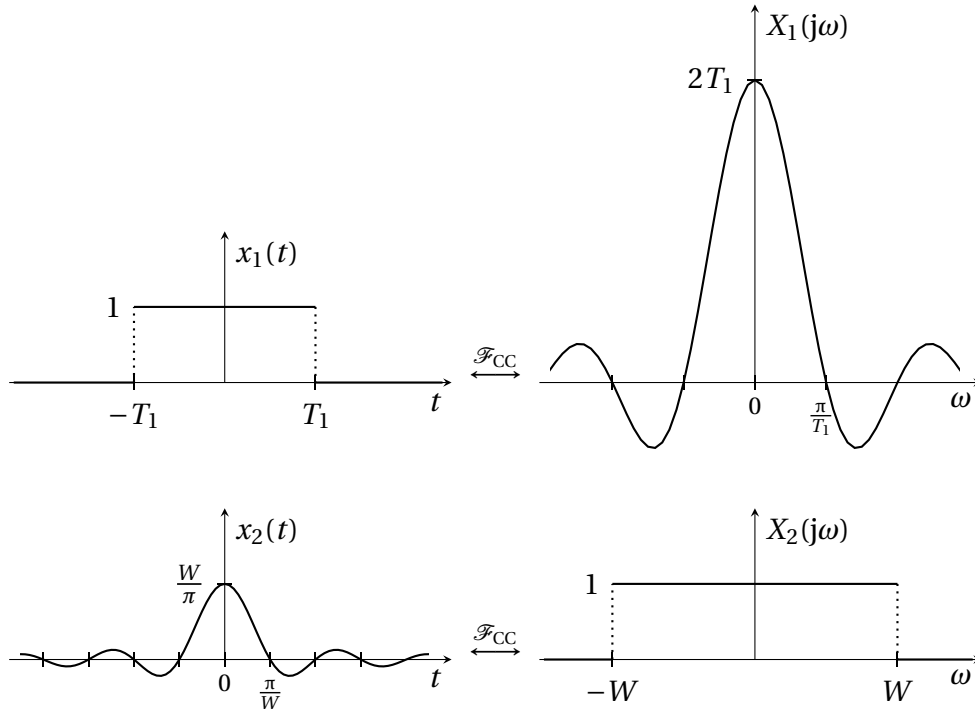


FIGURE 11.1 – La transformée d’une fenêtre rectangulaire est la fonction sinc et vice versa.

où

$$\text{sinc}(\theta) = \begin{cases} 1, & \theta = 0, \\ \frac{\sin(\theta)}{\theta}, & \theta \neq 0. \end{cases}$$

Réciproquement, le signal

$$X_2(j\omega) = \begin{cases} 1, & |\omega| \leq W, \\ 0, & |\omega| > W, \end{cases} \quad (11.1.2)$$

a pour transformée de Fourier inverse

$$x_2(t) = \frac{1}{2\pi} \int_{-W}^W e^{j\omega t} d\omega = \frac{W}{\pi} \text{sinc}(Wt).$$

Cette dualité est illustrée sur la Figure 11.1.

La multiplication d’un signal par une fenêtre rectangulaire dans le domaine temporel ou fréquentiel correspond à une convolution du signal avec la fonction sinc dans le domaine transformée. Cette simple propriété trouve plusieurs applications en traitement de signal.

Notons enfin que si la fenêtre n’est pas centrée à l’origine mais décalée de D , la transformée est la fonction sinc multipliée par le nombre com-

plexe $e^{\pm jD}$. Le décalage de la fenêtre affecte donc uniquement la phase du signal transformé.

11.1.2 Troncature d'un signal par fenêtrage rectangulaire

Tout signal stocké dans un ordinateur est de support fini. La troncature d'un signal de support infini correspond à une multiplication du signal par une fenêtre rectangulaire unitaire de support fini.

Un signal fenêtré est une approximation plus ou moins fidèle du signal original. Il est instructif d'analyser l'effet de cette approximation dans le domaine transformé, c'est-à-dire dans le domaine fréquentiel lorsque le fenêtrage est temporel et dans le domaine temporel lorsque le fenêtrage est fréquentiel. Le fenêtrage étant une multiplication par un signal rectangulaire, son effet dans le domaine transformé est une convolution par la fonction sinc. Lorsque la longueur de la fenêtre devient infinie, la fonction sinc tend vers une impulsion, et sa convolution avec le signal transformé laisse le signal inchangé. Par contre, pour une longueur de fenêtre finie, la convolution avec la fonction sinc provoque un double effet : une perte de résolution et une dispersion des phénomènes localisés (phénomène de fuite ou *leakage*). La perte de résolution provient du moyennage local causé par la largeur non nulle du lobe central de la fonction sinc. La convolution avec le lobe central équivaut à un filtrage passe-bas. Le phénomène de fuite est quant à lui causé par les lobes latéraux de la fonction sinc. Si le signal non fenêtré présente un pic localisé, la convolution de ce pic avec les lobes latéraux en répercutera l'effet sur l'ensemble de l'axe des temps (dans le cas d'un fenêtrage fréquentiel) ou des fréquences (dans le cas d'un fenêtrage temporel). Lorsque la longueur de fenêtrage diminue, la largeur du lobe centrale et l'aire des lobes latéraux augmente, ce qui amplifie les effets de perte de résolution et de dispersion.

Une manifestation bien connue des deux effets indésirables associés au fenêtrage est le phénomène de Gibbs vu au Chapitre 5 et dont l'illustration est reprise à la Figure 11.2. La disparité entre le signal rectangulaire et sa reconstruction au moyen d'un nombre fini de ses coefficients de Fourier équivaut à l'effet d'un fenêtrage fréquentiel dans le domaine temporel. La perte de résolution se manifeste dans la pente finie du signal approximé aux points de discontinuité du signal original. L'effet de dispersion des phénomènes localisés se manifeste dans les petites oscillations du signal approximé au voisinage des points de discontinuité.

Les effets indésirables du fenêtrage peuvent être atténués en utilisant des fenêtres non rectangulaires. Le choix d'une fenêtre particulière résulte d'un compromis entre la minimisation des effets d'approximation indésirables et la complexité de la fenêtre. La Figure 11.3 illustre quelques fenêtres couram-

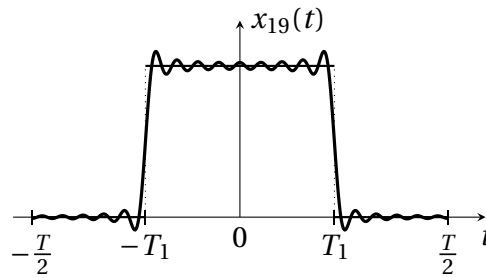


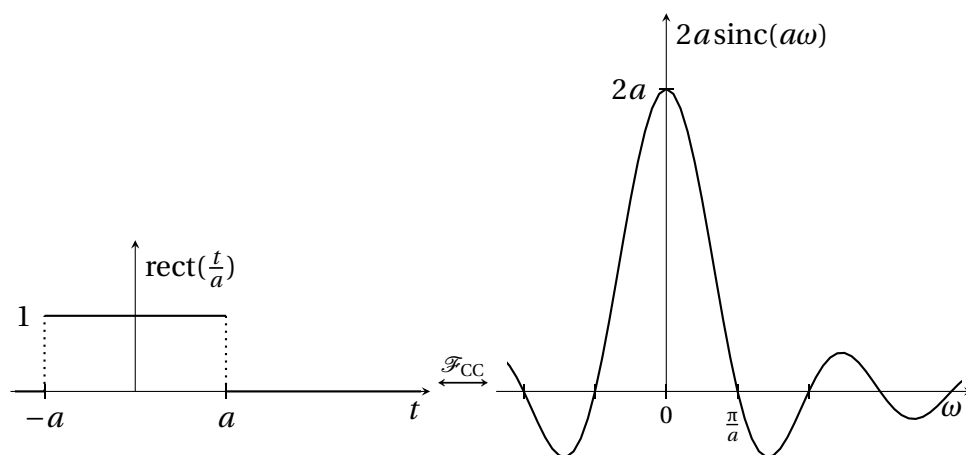
FIGURE 11.2 – Effets du fenêtrage fréquentiel dans le domaine temporel : perte de résolution et dispersion des phénomènes localisés dans le signal approximé (Phénomène de Gibbs).

ment utilisées. Le signal transformé montre que l'effet de dispersion observé dans le cas d'une fenêtre rectangulaire peut être pratiquement éliminé en recourant à des fenêtres continues, au prix d'une perte de résolution légèrement accrue. Cet exemple illustre l'intérêt de répartir les effets d'approximation sur le signal et sur sa transformée.

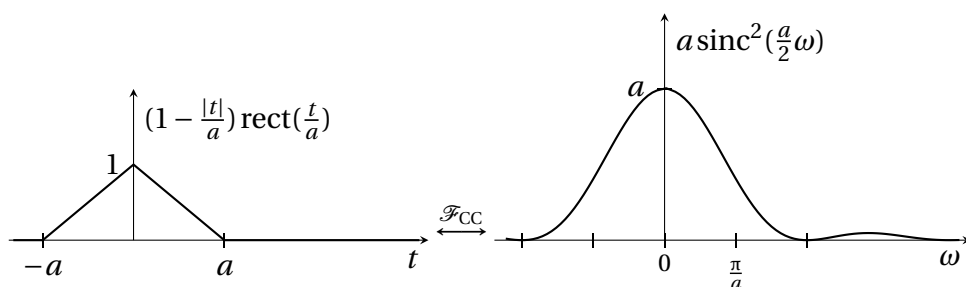
11.1.3 Filtrage passe-bande d'un signal par fenêtrage rectangulaire

Une des applications les plus courantes de la théorie du filtrage consiste à synthétiser un système LTI qui est « passant » dans une certaine plage de fréquences et « bloquant » aux autres fréquences. Dans un enregistrement vocal par exemple, en supposant que le bruit est essentiellement concentré dans les hautes fréquences, un filtre *passe-bas* permet d'enregistrer le signal vocal sans enregistrer le bruit. Dans la transmission d'information en *modulation d'amplitude* (AM), différents signaux sont transmis sur un même canal en assignant à chaque signal une certaine bande de fréquence. Pour séparer les différents signaux transmis, le récepteur doit être muni d'un filtre passe-bande. Les fréquences de coupure qui délimitent la bande passante d'un filtre sont ses caractéristiques de base. Nous allons voir cependant que la synthèse pratique d'un filtre s'expose à différents compromis qui requièrent la combinaison d'une analyse temporelle et fréquentielle.

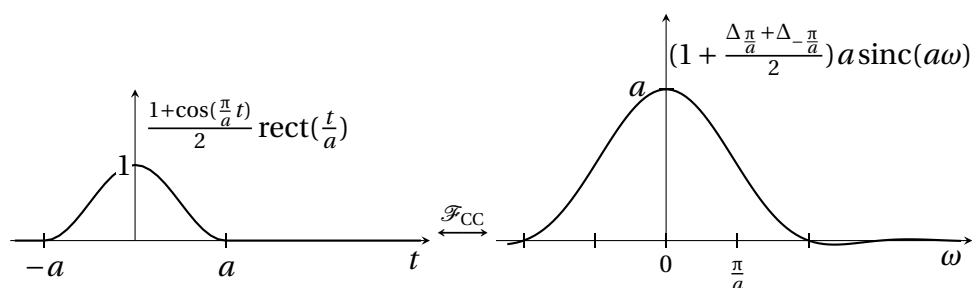
Considérons l'exemple d'un filtre passe-bas avec fréquence de coupure ω_c . Idéalement, on souhaite que le filtre laisse passer (sans distorsion en amplitude et en phase, c'est-à-dire $|H(j\omega)| = 1$, $\angle H(j\omega) = 0$) tous les signaux dans la gamme de fréquence $[-\omega_c, \omega_c]$ et bloque les fréquences restantes (c'est-à-dire $|H(j\omega)| = 0$). La réponse fréquentielle d'un tel filtre est une fenêtre rectangulaire de largeur $2\omega_c$. Sa réponse impulsionnelle est une fonction sinc. Le filtrage « idéal » consiste donc à convoluer le signal à filtrer avec la fonction sinc. Quand $\omega_c \rightarrow \infty$, la réponse impulsionnelle tend vers une impulsion, ce qui est en accord avec la transformée de Fourier (inverse) d'un signal constant.



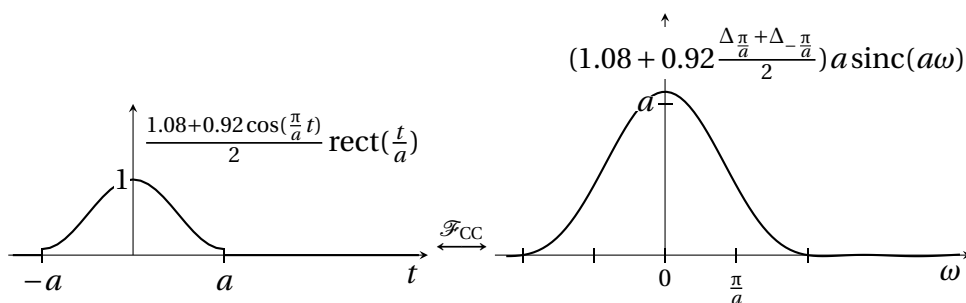
(a) Fenêtre rectangulaire



(b) Fenêtre triangulaire



(c) Fenêtre de Hann



(d) Fenêtre de Hamming

FIGURE 11.3 – Différentes fenêtres couramment utilisées pour tronquer un signal.

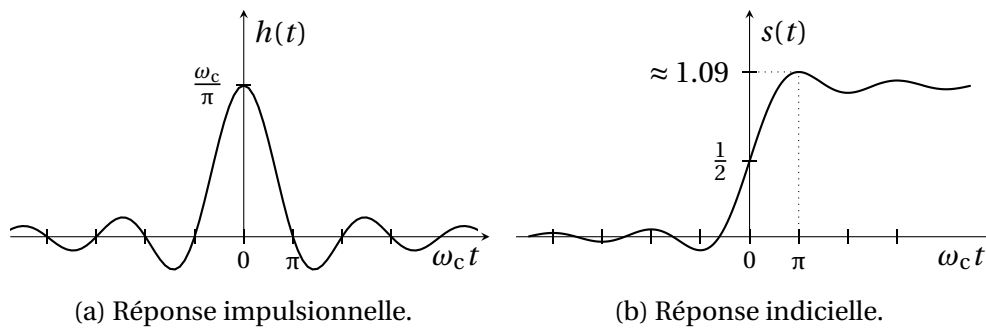
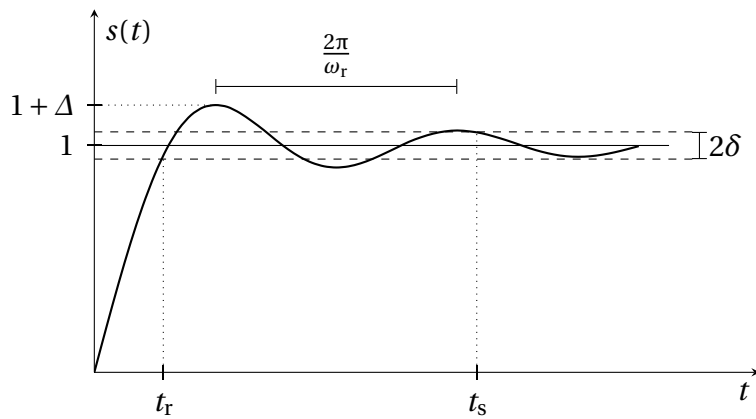


FIGURE 11.4 – Caractéristiques temporelles d'un filtre passe-bas « idéal ».

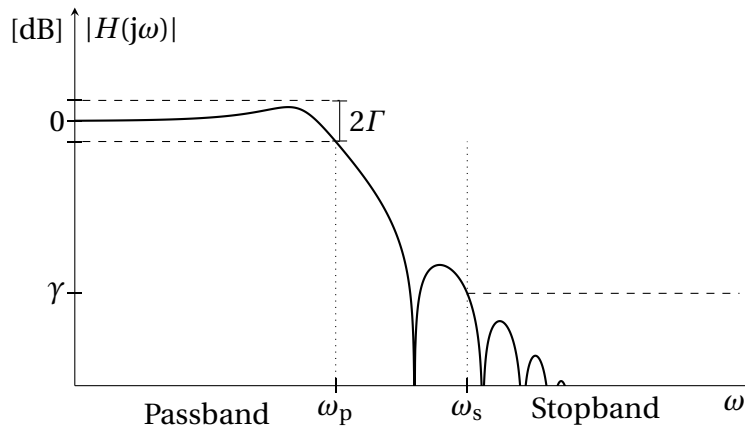
La réalisation du filtre idéal se heurte à plusieurs obstacles. Tout d'abord, le filtre idéal n'est pas causal, et sa réponse impulsionnelle a un support infini (cf. Figure 11.4a). L'implémentation sous forme d'un filtre causal à support fini (Finite Impulse Response) sera brièvement discutée dans la dernière section du chapitre (Section 11.5) mais elle introduit un retard.

En outre, les performances temporelles du filtre, par exemple évaluées sur la réponse indicielle, ne sont pas satisfaisantes. La réponse indicielle du filtre idéal est représentée à la Figure 11.4b. En général, une telle réponse est jugée insatisfaisante en raison de son caractère oscillatoire et de son dépassement trop élevé (elle atteint des valeurs de 10% supérieures à la valeur de régime). On peut également observer que le temps de montée est inversement proportionnel à la fréquence de coupure.

On peut aisément imaginer que les effets indésirables observés sur la réponse indicielle du filtre idéal sont dus à des spécifications trop « dures » dans le domaine fréquentiel. De manière analogue à l'application de fenêtrage discutée plus haut, il conviendra d'équilibrer les spécifications temporelles et fréquentielles. Les spécifications fréquentielles et temporelles typiques d'un filtre passe-bas sont représentées à la Figure 11.5b et 11.5a. En fréquence, on définit une région de transition entre la caractéristique passante et la caractéristique bloquante du filtre. On admet aussi une certaine tolérance sur le caractère strictement passant ($|H| = 1$) ou bloquant ($|H| = 0$) du filtre. Dans le domaine temporel, on spécifie un certain temps de montée (tout en sachant qu'il ne peut pas être choisi indépendamment de la bande passante), ainsi qu'une tolérance maximale sur le dépassement et sur le temps nécessaire pour que la réponse se stabilise (avec une tolérance δ) sur sa valeur finale. Aux spécifications temporelles et fréquentielles ainsi décrites s'ajoutent bien sûr des considérations pratiques d'implémentation et de coût. Les filtres les plus couramment utilisés en pratique sont les filtres *Butterworth* qui permettent de réaliser un bon compromis entre les spécifications temporelles et fréquentielles au moyen d'une équation différentielle ou aux différences d'ordre peu élevé (voir Section 9.6).



(a) Spécifications temporelles.



(b) Spécifications fréquentielles.

FIGURE 11.5 – Spécifications d'un filtre passe-bas.

11.2 Échantillonnage

Chaque fois qu'un signal continu $x(\cdot)$ est traité numériquement, il doit d'abord être *échantillonné*. Le signal échantillonné est un signal discret $x[\cdot]$ défini par $x[n] = x(nT)$, où la constante T est la période d'échantillonnage (cf. Figure 11.6).

L'intuition suggère que l'échantillonnage d'un signal est généralement associé à une perte d'information, d'autant plus importante que la période

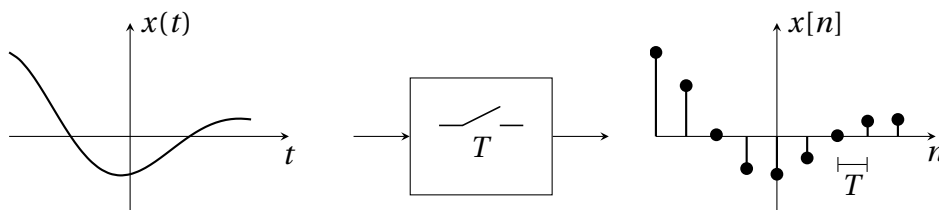


FIGURE 11.6 – Échantillonnage d'un signal en temps continu.

d'échantillonnage est élevée, puisque les valeurs du signal continu entre deux instants d'échantillonnage sont perdues dans le processus. Une infinité de signaux différents en temps continu peuvent interpoler le signal discret $x[\cdot]$. Nous allons cependant voir que sous certaines conditions, un signal en temps continu peut être parfaitement reconstruit à partir du signal échantillonné. C'est l'objet du célèbre théorème d'échantillonnage (souvent attribué à Shannon).

L'échantillonnage est une opération hybride qui associe un signal en temps discret $x[\cdot]$ à un signal en temps continu $x(\cdot)$. Une opération équivalente très utile pour l'analyse consiste à multiplier le signal $x(\cdot)$ par un *train d'impulsions*

$$p(t) = \sum_{n \in \mathbb{Z}} \delta(t - nT).$$

Le signal $x_p(\cdot) = x(\cdot)p(\cdot)$ est un signal en temps continu qui contient la même information que le signal échantillonné $x[\cdot]$. On peut en effet écrire

$$x_p(t) = \sum_{n \in \mathbb{Z}} x(nT)\delta(t - nT) = \sum_{n \in \mathbb{Z}} x[n]\delta(t - nT).$$

L'effet de l'échantillonnage dans le domaine fréquentiel peut être analysé en étudiant la transformée X_p du signal x_p . La propriété de multiplication-convolution donne

$$x_p(t) = x(t)p(t) \xleftrightarrow{\mathcal{F}_{CC}} \frac{1}{2\pi} X * P(j\omega).$$

Par ailleurs, la transformée du train d'impulsions p est un nouveau train d'impulsions (10.3.7)

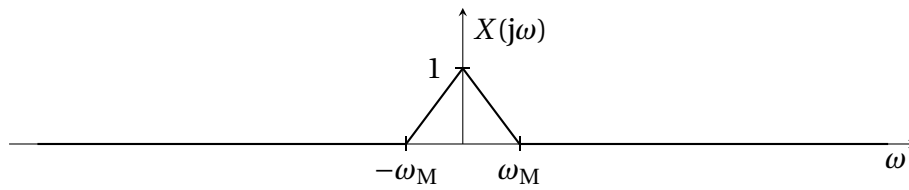
$$P(j\omega) = \frac{2\pi}{T} \sum_{k \in \mathbb{Z}} \delta(\omega - k\omega_s), \quad \omega_s = \frac{2\pi}{T}$$

Puisque la convolution avec une impulsion produit un simple décalage, on obtient

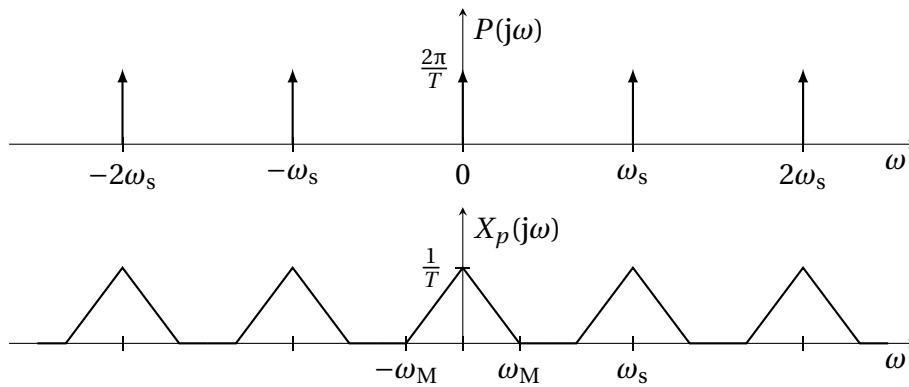
$$X_p(j\omega) = \frac{1}{T} \sum_{k \in \mathbb{Z}} X(j(\omega - k\omega_s)). \quad (11.2.1)$$

Le signal X_p est donc une fonction périodique obtenue par la superposition de copies décalées du signal $\frac{1}{T}X$.

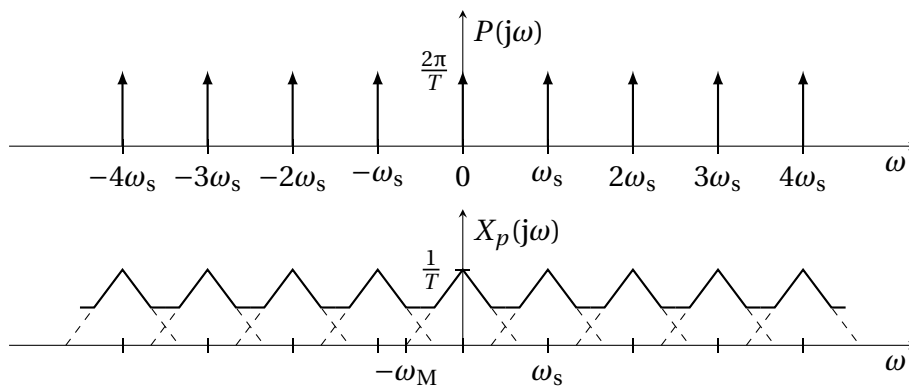
La Figure 11.7 illustre que l'échantillonnage d'un signal $x(\cdot)$ de largeur de bande ω_M peut conduire à deux situations différentes suivant la fréquence d'échantillonnage ω_s . Si $\omega_s > 2\omega_M$, le signal X est simplement recopié aux multiples entiers de la fréquence d'échantillonnage. Il n'y a pas recouvrement des copies du signal. Dans ce cas, il n'y a pas de perte d'information car le contenu fréquentiel du signal $x(\cdot)$ peut être extrait de celui du signal x_p : il suffit de filtrer x_p à l'aide d'un filtre passe-bas idéal de gain T et de fréquence



(a) Le signal original.



(b) Dans ce premier cas, le signal original peut être reconstitué par filtrage.

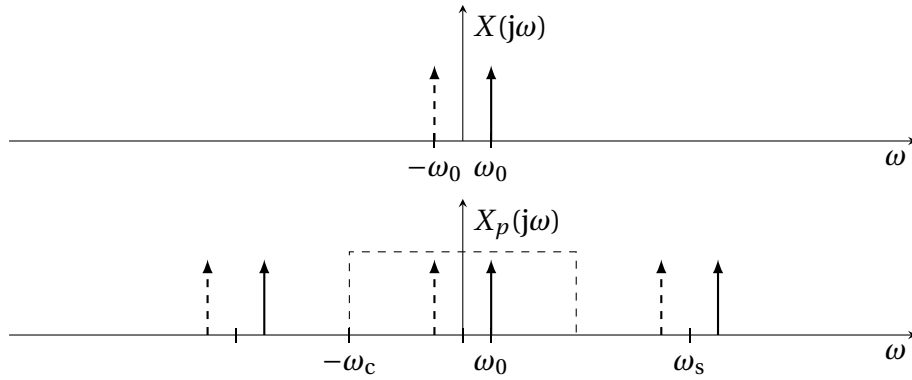


(c) Dans ce second cas, il y a recouvrement des spectres.

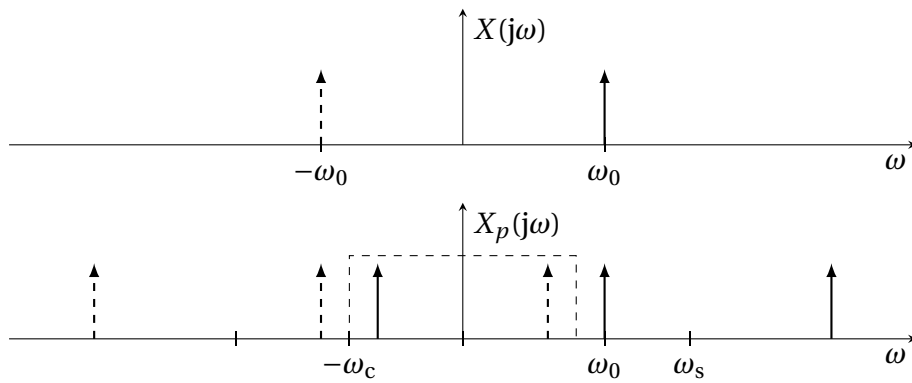
FIGURE 11.7 – Effet de l'échantillonnage dans le domaine fréquentiel.

de coupure comprise entre ω_M et $\omega_s - \omega_M$. Par contre, si $\omega_s \leq 2\omega_M$, les copies de X se recouvrent partiellement et il y a perte d'information : deux signaux différents en temps continu pourront dans ce cas donner lieu à un même signal x_p .

Le résultat que nous venons d'établir est un résultat de base connu sous le nom de théorème d'échantillonnage ou théorème de Shannon : un signal de largeur de bande limitée ω_M peut être échantillonné sans perte d'information si la fréquence d'échantillonnage ω_s est supérieure à deux fois la largeur de bande ω_M .



(a) Dans ce premier cas $\omega_s = 8\omega_0$ et le signal reconstruit est identique.



(b) Dans ce second cas $\omega_s = \frac{8}{5}\omega_0$ et il y a sous-échantillonnage, le signal reconstruit est $x_r(t) = \cos((\omega_s - \omega_0)t)$.

FIGURE 11.8 – Effet de l'échantillonnage sur le signal $x(t) = \cos(\omega_0 t)$.

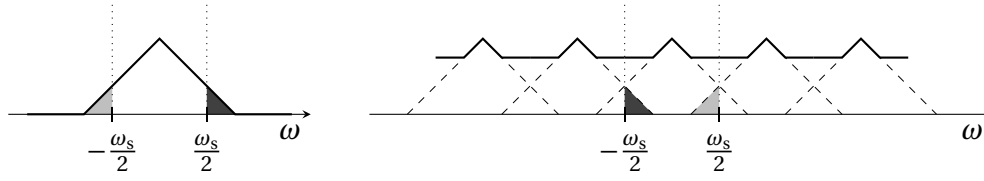
11.3 Sous-échantillonnage et repliement de spectre

Lorsque la condition de Shannon $\omega_s > 2\omega_M$ n'est pas respectée, on parle de sous-échantillonnage. Le sous-échantillonnage produit des effets indésirables connus sous le nom de repliement de spectre (*aliasing*). Ces effets peuvent être examinés sur un simple signal sinusoïdal $x(t) = \cos(\omega_0 t)$, cf. Figure 11.8.

Le spectre de $x(\cdot)$ est composé de deux impulsions aux fréquences ω_0 et $-\omega_0$. Le spectre du signal échantillonné x_p contient en outre des impulsions aux fréquences $\omega_s \pm \omega_0, 2\omega_s \pm \omega_0, \dots$. En supposant que le signal reconstruit x_r est obtenu par filtrage idéal avec un filtre passe-bas de fréquence de coupure $\omega_c = \omega_s/2$, on voit que $x_r(\cdot) = x(\cdot)$ pour $\omega_0 < \omega_s/2$, conformément au théorème de Shannon. Par contre, pour $\omega_s/2 < \omega_0 < \omega_s$, on obtient

$$x_r(t) = \cos((\omega_s - \omega_0)t).$$

La fréquence du signal reconstruit décroît à présent avec ω_0 pour finalement



(a) Spectre du signal initial. (b) Spectre du signal échantillonné à la fréquence ω_s .

FIGURE 11.9 – Phénomène de repliement du spectre.

s'annuler en $\omega_s = \omega_0$. À cette fréquence, la période d'échantillonnage est identique à la période du signal $x(\cdot)$ et le signal échantillonné est identique à celui d'un signal constant.

Les effets associés au sous-échantillonnage sont connus sous le phénomène de repliement de spectre car le spectre de x_r est le spectre de $x(\cdot)$ « replié » sur la bande de fréquences $[-\frac{\omega_s}{2}, \frac{\omega_s}{2}]$ comme illustré sur la Figure 11.9.

Le sous-échantillonnage constitue le principe de fonctionnement du stroboscope. Le flash périodique du stroboscope correspond à un échantillonnage par train d'impulsions. Le sous-échantillonnage permet d'observer des phénomènes haute fréquence en « repliant » leur spectre dans un intervalle de fréquence admissible pour l'œil.

11.4 Interpolation

La reconstruction d'un signal en temps continu $x_i(\cdot)$ à partir d'un signal discret $x[\cdot]$ est un processus d'*interpolation*. Opération inverse de l'échantillonnage, l'interpolation est une nouvelle opération hybride qui associe cette fois un signal en temps continu à un signal en temps discret.

Le signal x_p associé au signal discret $x[\cdot]$ nous sera à nouveau très utile pour l'analyse. Nous allons en effet générer différents processus d'interpolation par simple filtrage (ou convolution) du signal x_p :

$$x_i(t) = x_p * h(t) \xleftrightarrow{\mathcal{F}_{CC}} X_i(j\omega) = X_p(j\omega)H(j\omega).$$

Le signal interpolé aura dès lors la forme

$$x_i(t) = \sum_{k \in \mathbb{Z}} x[k]h(t - kT). \quad (11.4.1)$$

La fonction d'interpolation h constitue un choix pour l'utilisateur. Un exemple d'interpolation a été donné dans la section précédente : nous avons vu que sous l'hypothèse de Shannon, le signal $x(\cdot)$ était parfaitement reconstitué au moyen d'un filtre idéal passe-bas. La fonction H est dans ce cas un rectangle de hauteur T (Figure 11.10) et la réponse impulsionnelle du filtre

idéal vaut

$$h(t) = \frac{\omega_c T}{\pi} \text{sinc}(\omega_c t).$$

On obtient donc la formule d'interpolation

$$x_i(t) = \frac{\omega_c T}{\pi} \sum_{n \in \mathbb{Z}} x(nT) \text{sinc}(\omega_c(t - nT)). \quad (11.4.2)$$

Cette formule d'interpolation est exacte (c'est-à-dire $x_i(\cdot) = x(\cdot)$) si la condition de Shannon est satisfaite et si la fréquence de coupure ω_c . L'interpolation exacte est illustrée dans la Figure 11.11.

En pratique, on utilise des formules d'interpolation plus simples que la formule (11.4.2). L'interpolation la plus simple est le *bloqueur d'ordre zéro* : le signal x_i est obtenu en maintenant le signal à une valeur constante entre deux instants d'échantillonnage :

$$x_i(t) = x(nT), \quad nT \leq t < (n+1)T$$

L'interpolation par un bloqueur d'ordre zéro correspond à un filtrage par un filtre de réponse impulsionnelle

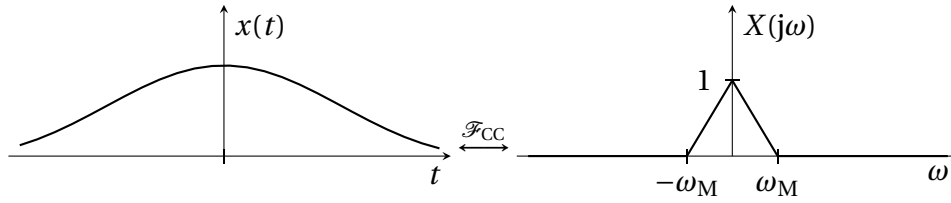
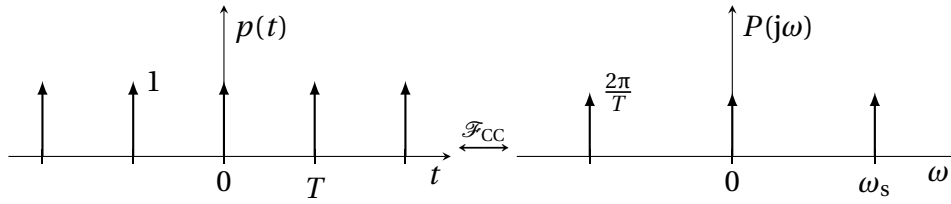
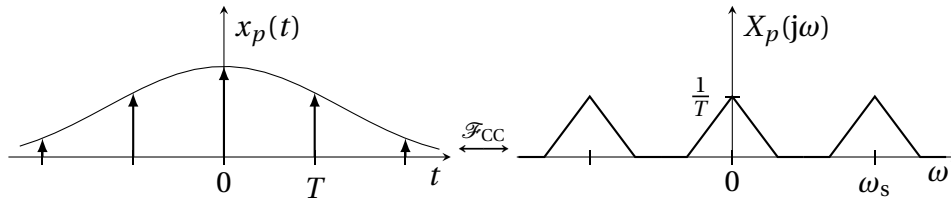
$$h(t) = \begin{cases} 1, & t \in [0, T], \\ 0, & t \notin [0, T]. \end{cases}$$

La Figure 11.12 illustre que la réponse fréquentielle de ce filtre constitue une approximation grossière du filtre passe-bas idéal qui permettrait une reconstruction parfaite du signal.

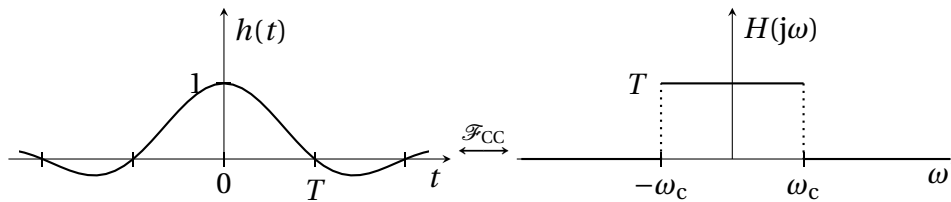
Tout comme dans les applications de fenêtrage temporel et fréquentiel discutées précédemment, le choix d'un filtre d'interpolation résulte encore une fois d'un compromis entre la complexité de sa réponse impulsionnelle h et sa capacité à approximer de manière satisfaisante la réponse fréquentielle du filtre passe-bas idéal. Le bloqueur d'ordre zéro et le filtre idéal constituent les deux extrêmes de ce compromis. Un choix intermédiaire consiste en une interpolation linéaire du signal discret (bloqueur d'ordre un). La Figure 11.12 illustre la qualité intermédiaire du filtre qui en résulte (les lobes latéraux sont fortement réduits).

11.5 Synthèse d'un filtre à réponse impulsionnelle finie

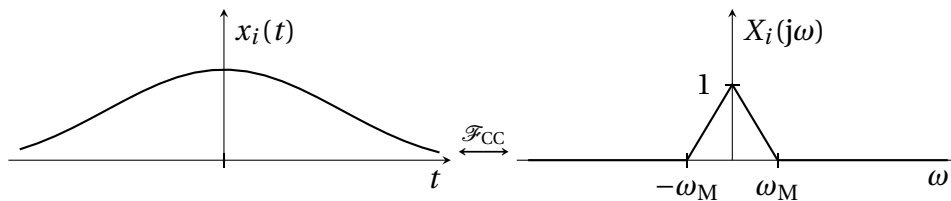
Le filtrage numérique de signaux continus comporte trois étapes : l'échantillonnage du signal d'entrée $u(\cdot)$, la génération du signal de sortie $y[\cdot]$ à partir du signal d'entrée $u[\cdot]$ par un filtre, et l'interpolation du signal $y[\cdot]$ pour générer un signal de sortie continu $y(\cdot)$.

(a) Signal continu $x(\cdot)$.(b) Train d'impulsions $p(t) = \sum_{n \in \mathbb{Z}} \delta(t - nT)$.

(c) Signal « pulsé » $x_p(t) = x(t)p(t)$ contenant la même information que le signal échantillonné $x(nT)$. Sa transformée de Fourier $X_p = X * P$ est une superposition de copies décalées du signal $\frac{1}{T}X$. On suppose que la condition de Shannon $\omega_s > 2\omega_M$ est respectée et qu'il n'y a donc pas de recouvrement.

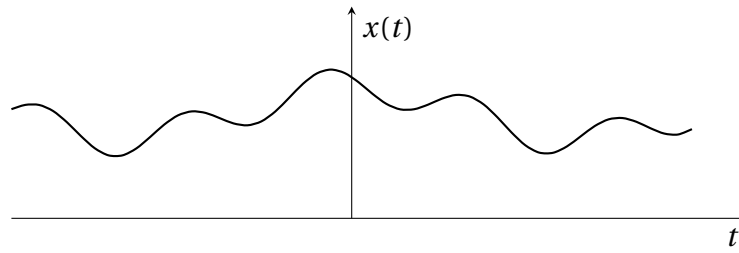


(d) Filtre idéal passe-bas $h(t) = \frac{\omega_c T}{\pi} \text{sinc}(\omega_c t)$ avec $\omega_M < \omega_c < \omega_s - \omega_M$.

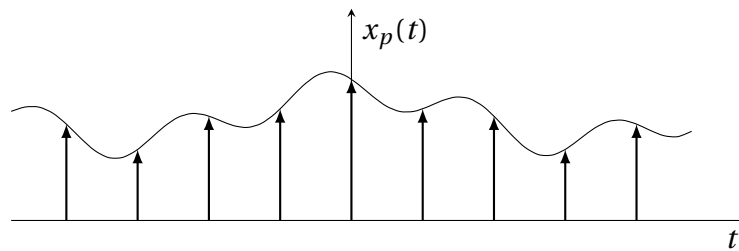
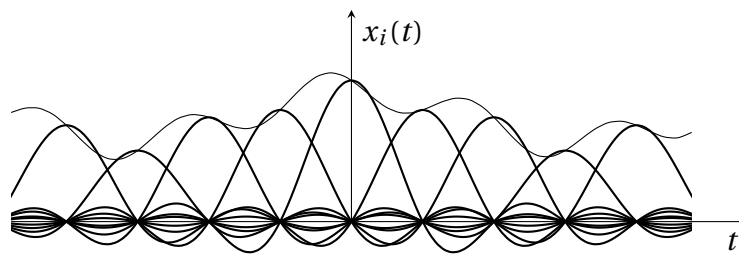


(e) Résultat du passage de x_p dans le filtre idéal passe-bas. Vu que la condition de Shannon est respectée, on observe que $X_i = X_p H = X(j\omega)$ et donc $x_i(t) = x_p * h(t) = x(t)$.

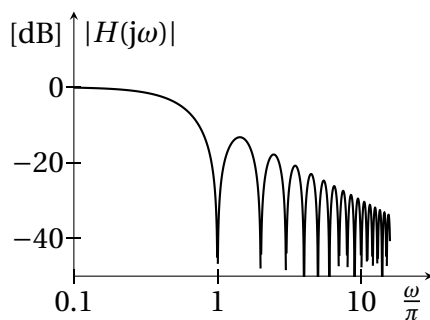
FIGURE 11.10 – Illustration des processus d'échantillonnage et d'interpolation. La colonne de gauche contient des signaux temporels dont le contenu fréquentiel est affiché en vis-à-vis dans la colonne de droite.



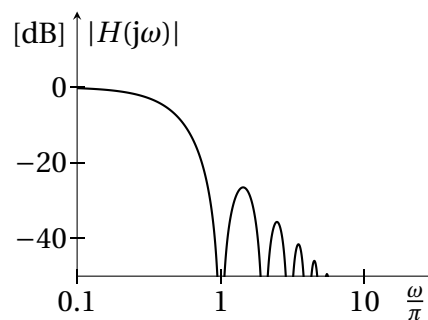
(a) Signal de départ.

(b) Signal δ -échantillonné.

(c) Interpolation du signal.

FIGURE 11.11 – Interpolation exacte d'un signal à bande limitée au moyen d'un filtre passe-bas idéal avec $\omega_c = \omega_s/2$.

(a) Bloqueur d'ordre zéro.



(b) Bloqueur d'ordre un.

FIGURE 11.12 – Amplitude de la réponse fréquentielle d'un bloqueur d'ordre zéro (interpolation constante par morceaux) et d'un bloqueur d'ordre un (interpolation linéaire).

Un filtre à *réponse impulsionnelle finie* (FIR pour *Finite Impulse Response*) est spécifié sous la forme explicite

$$y[n] = h[0]u[n] + h[1]u[n-1] + \dots + h[M]u[n-M],$$

où la séquence finie $h[\cdot]$ est la réponse impulsionnelle du filtre.

La réalisation d'un tel filtre sous la forme d'un bloc-diagramme est directe. Par opposition, un filtre à *réponse impulsionnelle infinie* (IIR pour *Infinite Impulse Response*) est spécifié sous la forme d'une équation aux différences

$$\sum_{k=1}^N a_k y[n+k] = \sum_{m=1}^M b_m u[n+m],$$

ou, de manière équivalente, par sa fonction de transfert $H(z) = \frac{b_M z^M + \dots + b_0}{a_N z^N + \dots + a_0}$. La réalisation d'un tel filtre sous forme d'un bloc-diagramme a été discutée au Chapitre 4. Le choix d'un filtre FIR ou d'un filtre IIR dépend du contexte et de l'utilisation. L'implémentation d'un filtre IIR est souvent plus compacte mais il est plus sensible aux erreurs numériques vu son implémentation en boucle fermée. Le filtre FIR offre l'avantage d'être BIBO stable quelle que soit la précision des coefficients du filtre.

Une méthode simple pour la conception de filtres FIR est basée sur le fenêtrage d'une réponse impulsionnelle infinie : Partant d'une réponse fréquentielle désirée, on calcule par transformée inverse la réponse impulsionnelle correspondante, et on tronque cette dernière au moyen d'une fenêtre. En général, la réponse impulsionnelle obtenue est celle d'un filtre non causal. On décale alors le signal $h[\cdot]$ vers la droite pour rendre le filtre causal. (Ceci introduit bien sûr un retard dans le filtrage, ce qui constitue une limitation des filtres FIR).

En guise d'illustration, on décrira brièvement la synthèse d'un filtre FIR passe-bas échantillonné à 180 kHz et présentant une fréquence de coupure de 60 kHz. La réponse fréquentielle du filtre idéal passe-bas vaut donc

$$H(j\omega) = \begin{cases} 1, & |\omega| < 120\pi (= \omega_c), \\ 0, & |\omega| \geq 120\pi. \end{cases} \quad (11.5.1)$$

La réponse impulsionnelle du filtre parfait vaut $h(t) = \frac{\omega_c}{\pi} \text{sinc}(\omega_c t)$. C'est la réponse impulsionnelle infinie d'un filtre non causal.

Il faut donc fenêtrer la réponse impulsionnelle pour la rendre finie. L'utilisation d'une fenêtre de Hann ou Hamming élimine pratiquement les oscillations de la réponse approximée. Une fenêtre de longueur MT , où T est la période d'échantillonnage, donnera un lobe principal de largeur $\frac{2}{MT}$, ce qui déterminera la résolution de la réponse fréquentielle approximée. Une

tolérance de 5% (soit 3 kHz) sur la fréquence de coupure du filtre conduit donc à choisir une fenêtre de $M = 2 \frac{180}{3} = 120$ échantillons.

Pour rendre le filtre causal, la réponse impulsionnelle fenêtrée doit être décalée de 60 échantillons, ce qui introduit un retard de $60/180 = 0.33$ ms. L'implémentation du filtre requiert 119 décalages unitaires et 120 coefficients. Un filtre (IIR) de Butterworth donnerait un résultat équivalent avec 4 décalages unitaires et 10 coefficients.

Annexe A

Espaces vectoriels et applications linéaires

A.1 Espace vectoriel

Un *espace vectoriel* $\mathcal{E}$ sur $\mathbb{C}$ (le *corps* de l'espace vectoriel) est composé d'un ensemble d'éléments, appelés *vecteurs*, auxquels sont associées deux opérations :

- l'addition vectorielle $+: \mathcal{E} \times \mathcal{E} \rightarrow \mathcal{E} : v, w \mapsto v + w$
- la multiplication par un scalaire $\cdot: \mathbb{C} \times \mathcal{E} \rightarrow \mathcal{E} : \alpha, v \mapsto \alpha \cdot v = \alpha v$.

Pour tous $u, v, w \in \mathcal{E}$ et $\alpha, \beta \in \mathbb{C}$, ces dernières doivent satisfaire les propriétés suivantes.

1. Commutativité de l'addition : $v + w = w + v$.
2. Associativité de l'addition : $(u + v) + w = u + (v + w)$.
3. Existence d'un élément identité pour l'addition : il existe un élément $0 \in \mathcal{E}$ tel que $0 + v = v$.
4. Existence de l'inverse pour l'addition : pour tous $v \in \mathcal{E}$, il existe $-v \in \mathcal{E}$ tel que $v + (-v) = 0$.
5. Associativité de la multiplication par un scalaire : $\alpha \cdot (\beta \cdot v) = (\alpha\beta) \cdot v$.
6. Identité pour la multiplication par un scalaire : $1 \cdot v = v$.
7. Double distributivité : $(\alpha + \beta) \cdot v = \alpha \cdot v + \beta \cdot v$ et $\alpha \cdot (v + w) = \alpha \cdot v + \alpha \cdot w$.

On a une définition similaire pour un espace vectoriel sur $\mathbb{R}$, mais pas sur $\mathbb{Z}$ qui ne satisfait pas les conditions pour former un corps.

Exemples

L'ensemble $\mathbb{C}^n$, formé de tous les ensembles ordonnés $x = (x_1, x_2, \dots, x_n)$ où $x_1, x_2, \dots, x_n \in \mathbb{C}$, constitue un espace vectoriel sur $\mathbb{C}$ si l'on définit l'addition composante par composante et la multiplication par un scalaire par $\alpha \cdot x = (\alpha x_1, \alpha x_2, \dots, \alpha x_n)$. $\mathbb{C}^n$ est également un espace vectoriel sur $\mathbb{R}$. $\mathbb{R}^n$ est un espace vectoriel sur $\mathbb{R}$ mais pas sur $\mathbb{C}$.

De même, l'ensemble des matrices $m \times n$ d'éléments complexes ou réels forme un espace vectoriel sur $\mathbb{C}$ ou $\mathbb{R}$ respectivement lorsque l'on considère les règles habituelles d'addition et de multiplication par un scalaire.

On se convainc aisément que l'ensemble des matrices $n \times n$ symétriques forme un espace vectoriel, tout comme les matrices $n \times n$ antisymétriques.

Comme les éléments de ces espaces vectoriels sont contenus dans l'ensemble $\mathbb{R}^{n \times n}$ des matrices $n \times n$ en considérant le même corps et les mêmes définitions des opérations, on dit qu'ils forment des *sous-espaces vectoriels* de $\mathbb{R}^{n \times n}$. On vérifie aisément que l'union de ces deux ensembles ne forme pas un espace vectoriel, contrairement à leur intersection (qui ne contient que l'élément identité la matrice nulle) et leur somme (définie comme l'ensemble des vecteurs résultant de l'addition de vecteurs des deux sous-espaces vectoriels). Au vu de la propriété de décomposition d'une matrice en partie paire et impaire, on voit que ce dernier espace vectoriel est en fait identique à $\mathbb{R}^{n \times n}$.

L'ensemble $C(\mathbb{C}, \mathbb{C})$ des fonctions continues de $\mathbb{C}$ dans $\mathbb{C}$ forme un espace vectoriel lorsque l'on définit l'addition de deux fonctions et la multiplication par un scalaire « point par point »; i.e., on définit $(f_1 + f_2)(z) = f_1(z) + f_2(z)$ et $(\alpha f_1)(z) = \alpha f_1(z)$ pour tous $z \in \mathbb{C}$ et $f_1, f_2 \in C(\mathbb{C}, \mathbb{C})$.

Avec les mêmes définitions des opérations, l'ensemble $C^n(\mathbb{C}, \mathbb{C})$ des fonctions n fois continument dérivables forme un espace vectoriel.

De même, les espaces fonctionnels L_1, L_2, L_∞ et $\ell_1, \ell_2, \ell_\infty$ définis page 22 constituent des espaces vectoriels.

A.2 Dimension d'un espace vectoriel

Les vecteurs $v_1, v_2, \dots, v_n$ d'un espace vectoriel $\mathcal{E}$ sont *linéairement indépendants* lorsque $\alpha_1 \cdot v_1 + \alpha_2 \cdot v_2 + \dots + \alpha_n \cdot v_n = 0$ implique nécessairement $\alpha_1 = \alpha_2 = \dots = \alpha_n = 0$.

Dès lors, la *dimension* d de $\mathcal{E}$ est définie comme étant le nombre maximal possible de vecteurs de $\mathcal{E}$ linéairement indépendants.

Exemples

La dimension de $\mathbb{R}^n$ est n . La dimension de $\mathbb{C}^n$ sur $\mathbb{C}$ est également n mais la dimension de $\mathbb{C}^n$ sur $\mathbb{R}$ est $2n$.

Encore, la dimension d_S de l'ensemble des matrices $n \times n$ symétriques est $n(n+1)/2$; la dimension d_A des matrices $n \times n$ antisymétriques est $n(n-1)/2$. La dimension de l'espace résultant de leur somme vaut $n^2 = d_S + d_A$. Cette dernière égalité est valable pour la somme de deux sous-espaces vectoriels quelconques si et seulement si ils contiennent pour seul élément commun l'identité; on dit alors qu'il s'agit de la *somme directe* de ces deux sous-espaces vectoriels.

Les espaces vectoriels de type $C^n(\mathbb{C}, \mathbb{C})$ sont de dimension infinie, en ce sens que pour tout ensemble fini de fonctions n fois continument dérivables, il sera toujours possible de trouver une fonction supplémentaire tel que

l'ensemble reste linéairement indépendant. Il en est de même des espaces vectoriels L_1, L_2, L_∞ ou $\ell_1, \ell_2, \ell_\infty$.

A.3 Base d'un espace vectoriel

Un ensemble de vecteurs linéairement indépendants forme une *base* de l'espace vectoriel $\mathcal{E}$, lorsque tout vecteur de $\mathcal{E}$ peut être écrit comme une combinaison linéaire des vecteurs de la base (la famille de vecteurs est *génératrice*). On montre que le nombre de vecteurs de base ne dépend pas de la base particulière choisie pour l'espace vectoriel $\mathcal{E}$, et correspond en fait à sa dimension. Ainsi, la dimension d de $\mathcal{E}$ peut aussi être définie comme le nombre de vecteurs formant une base de $\mathcal{E}$.

Etant donnés une base $B = (v_1, v_2, \dots, v_d)$ et un vecteur v de $\mathcal{E}$, l'unique suite d'éléments du corps $\alpha_1, \alpha_2, \dots, \alpha_d$ tels que $v = \alpha_1 v_1 + \alpha_2 v_2 + \dots + \alpha_d v_d$ sont les *coordonnées* de v dans la base B . Tout vecteur d'un espace vectoriel de dimension d peut donc être représenté, dans une base donnée B , par le vecteur des coordonnées

$$\alpha = \begin{bmatrix} \alpha_1 & \alpha_2 & \dots & \alpha_d \end{bmatrix}^T.$$

A.4 Application linéaire

Soit une application $f : \mathcal{E} \rightarrow \mathcal{F} : v \mapsto w = f(v)$, où $\mathcal{E}$ et $\mathcal{F}$ sont deux espaces vectoriels sur le corps $\mathbb{K}$. Pour être une *application linéaire*, f doit satisfaire

1. $f(u + v) = f(u) + f(v)$ et
2. $f(\alpha v) = \alpha f(v)$

pour tous $u, v \in \mathcal{E}$ et $\alpha \in \mathbb{K}$.

L'*image* de f , notée $\text{im}(f)$, est l'ensemble des $f(v)$ pour $v \in \mathcal{E}$.

Le *noyau* de f , noté $\ker(f)$, est l'ensemble des $v \in \mathcal{E}$ tels que $f(v) = 0$.

Lorsque les espaces vectoriels $\mathcal{E}$ et $\mathcal{F}$ sont munis de bases, respectivement $B_E = (v_1, v_2, \dots, v_{d_E})$ et $B_F = (w_1, w_2, \dots, w_{d_F})$, on peut représenter f de la manière suivante.

Soit un vecteur quelconque $v \in \mathcal{E}$ de composantes $(\alpha_1, \alpha_2, \dots, \alpha_{d_E})$ dans la base B_E . Par les propriétés de l'application linéaire, on a

$$\begin{aligned} f(v) &= f(\alpha_1 v_1 + \alpha_2 v_2 + \dots + \alpha_{d_E} v_{d_E}) \\ &= \alpha_1 f(v_1) + \alpha_2 f(v_2) + \dots + \alpha_{d_E} f(v_{d_E}). \end{aligned}$$

L'application linéaire est donc définie par l'image des vecteurs de base. De plus, les vecteurs $f(v_i)$ peuvent être exprimés dans la base B_F :

$$\begin{aligned} f(v_1) &= a_{11}w_1 + a_{21}w_2 + \cdots + a_{d_F1}w_{d_F} \\ f(v_2) &= a_{12}w_1 + a_{22}w_2 + \cdots + a_{d_F2}w_{d_F} \\ &\vdots \\ f(v_{d_E}) &= a_{1d_E}w_1 + a_{2d_E}w_2 + \cdots + a_{d_Fd_E}w_{d_F}. \end{aligned}$$

Ainsi, les a_{ij} (éléments du corps) définissent complètement l'application linéaire. Les coordonnées de $w = f(v)$ dans la base B_F , i.e., les β_i tels que

$$f(v) = \beta_1w_1 + \beta_2w_2 + \cdots + \beta_{d_F}w_{d_F},$$

valent donc

$$\beta_i = \sum_{j=1}^{d_E} a_{ij}\alpha_j$$

pour $i = 1, 2, \dots, d_F$.

Finalement, l'application linéaire est entièrement caractérisée par la correspondance qu'elle établit entre les coordonnées de départ α_j et les coordonnées transformées β_i . Cela peut s'exprimer par la relation matricielle

$$\begin{bmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_{d_F} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1d_E} \\ a_{21} & a_{22} & \cdots & a_{2d_E} \\ \vdots & \vdots & & \vdots \\ a_{d_F1} & a_{d_F2} & \cdots & a_{d_Fd_E} \end{bmatrix} \begin{bmatrix} \alpha_1 \\ \alpha_2 \\ \vdots \\ \alpha_{d_E} \end{bmatrix}. \quad (\text{A.4.1})$$

Toute application linéaire de $\mathcal{E}$ dans $\mathcal{F}$ peut donc être représentée par une matrice A de dimension $d_F \times d_E$. Chaque colonne de la matrice A représente les coordonnées dans la base B_F de l'image d'un vecteur de la base B_E . On peut noter la représentation matricielle (A.4.1) sous la forme compacte

$$\beta = A\alpha.$$