

Modélisation des systèmes asservis



Robot NAO (ALDEBARAN Robotics)

Objectifs

Les systèmes asservis sont des systèmes automatisés qui assurent le pilotage d'une grandeur physique d'un processus en la mesurant et en adaptant la commande pour suivre une consigne. La conception d'un système asservi et en particulier l'élaboration de sa loi de commande nécessite une modélisation mathématique du système. L'objectif de ce cours est de poser les bases de la modélisation des systèmes asservis (et plus généralement des systèmes dynamiques) et l'évaluation de leurs performances.

Table des matières

1	Introduction	3
1.1	Commande des systèmes automatisés	3
1.2	Structure d'un système asservi	6
1.3	Performances de la commande d'un système	8
2	Modélisation du comportement dynamique d'un système	12
2.1	Objectif de la modélisation	13
2.2	Hypothèses de modélisation	13
2.3	Systèmes linéaires continus invariants	17
3	Résolution par la transformée de Laplace	19
3.1	Définition et intérêts	19
3.2	Propriétés élémentaires	21
3.3	Transformée de fonctions usuelles	25
3.4	Fonction de transfert	28
3.5	Propriétés asymptotiques	30
3.6	Transformée de Laplace inverse	31
4	Représentation par schéma-blocs	37
4.1	Schéma-blocs	37
4.2	Chaîne fermée : système asservi	39
4.3	Manipulation de schéma-blocs	40



1 Introduction

Un système est une agrégation d'éléments inter-connectés constitué naturellement ou artificiellement pour accomplir une tâche pré-définie. Son état est affecté par une ou plusieurs variables : les **entrées** du système. Le résultat de l'action des entrées est la **réponse** du système qui peut être caractérisée par le comportement d'une ou plusieurs variables de sorties. La notion de système en automatique est indissociable de celle de signal. Ils sont généralement représentés par un **schéma fonctionnel** consistant en un bloc, c'est-à-dire un rectangle, sur lequel les signaux d'entrée sont représentés par des flèches entrantes sur la gauche. L'action des entrées produit de manière causale des effets mesurés par des signaux de sortie représentés par des flèches sortantes sur la droite (figure 1).

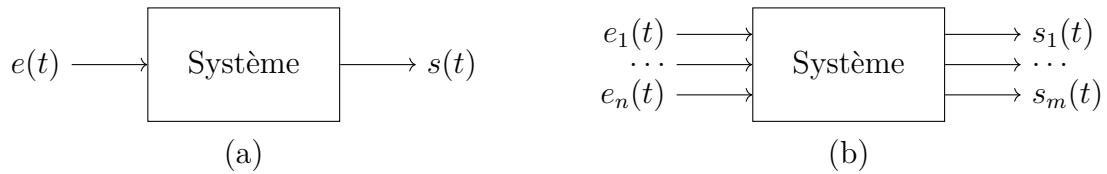


FIGURE 1 – Schéma fonctionnel d'un système mono-variable (a) et multi-variables (b).

Chaque système peut être caractérisé par un nombre fini de variables qui peuvent être de natures différentes. Parmi les entrées affectant un système, on peut distinguer les **commandes** qui ont pour but d'exercer des actions entraînant le fonctionnement souhaité du système et les **perturbations** qui troublent le fonctionnement désiré (figure 2).

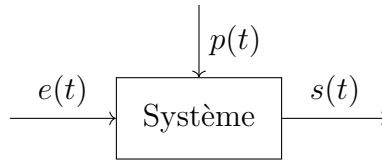


FIGURE 2 – Commande $e(t)$ et perturbation $p(t)$ d'un système.

En l'absence de perturbation, il est imaginable de créer un signal de commande d'un système produisant le signal de sortie souhaité. Par contre, en présence de perturbations, il est fort probable que les signaux de sortie ne correspondent pas à ceux souhaités. L'introduction d'un système de commande devient alors nécessaire.

1.1 Commande des systèmes automatisés

L'automatique est la discipline scientifique qui traite de la caractérisation, de la conception et de la réalisation du système de commande des systèmes automatisés.

Définition 1.1 (Système automatisé)

Un système automatisé est un système réalisant des opérations et pour lequel l'intervention humaine est limitée à la programmation du système et à son réglage préalable.

Définition 1.2 (Système de commande)

Un système de commande est un système dont la vocation est de piloter un processus.

Le but d'un système de commande est d'exercer des actions entraînant une amélioration du comportement et donc des performances d'un système ou d'un processus. L'ensemble des méthodes permettant l'analyse du comportement d'un système, la synthèse d'un système de commande satisfaisant les exigences de performances spécifiées dans un cahier des charges définit la **théorie de la commande**. La théorie de la commande est une branche de la théorie des systèmes, par nature pluri-disciplinaire et développée à partir de solides fondements mathématiques, qui a pour objectif très concret de développer des **correcteurs/régulateurs** pouvant être mis en œuvre sur des systèmes techniques réels dans les domaines de l'aéronautique, le spatial, l'industrie chimique, l'automobile, etc.

Un système de commande est toujours associé à un processus (la partie opérative) qu'il doit commander. Quelle que soit sa nature, il est toujours possible de classer les différentes structures de commande en deux catégories :

- les structures de commande en boucle ouverte ou directe (figure 3a) ;
- les structures de commande en boucle fermée ou à rétroaction (figure 3b).

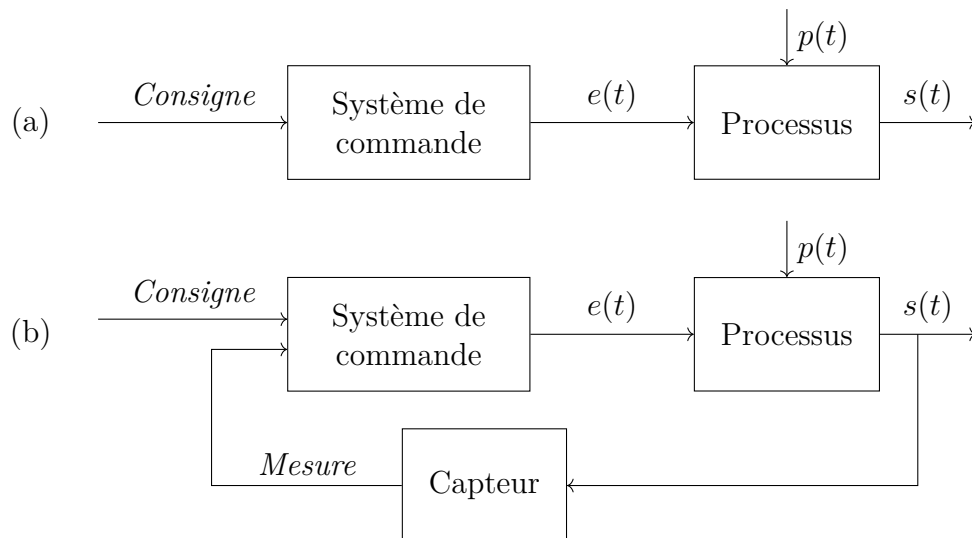
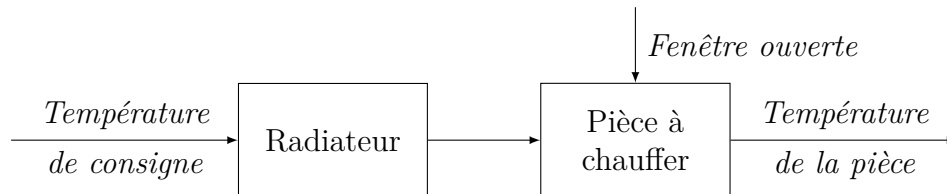


FIGURE 3 – Commande en boucle ouverte (a) et en boucle fermée (b) d'un processus.

Un système de commande en boucle ouverte nécessite la connaissance parfaite du processus et des perturbations pour que la réponse observée corresponde à la consigne. En effet, sans dispositif de vérification de la réponse du système, il est impossible de corriger ou modifier la commande pour l'adapter à d'éventuels aléas ou imperfections de prévision du comportement du système car en aucun cas les signaux de commande ne dépendent des signaux de sortie. Si ces structures de commande présentent l'avantage d'une grande simplicité de mise en œuvre (et donc de coût réduit), elles sont très sensibles aux perturbations et doivent donc être proscrites pour les applications nécessitant des niveaux de performances élevés ou associées à des perturbations difficilement prévisibles.

Exemple 1.1 (Chauffage d'une pièce)

Considérons le chauffage par un radiateur d'une pièce dont on souhaite maintenir la température à 20°C (température de consigne). Pour une température extérieure donnée, il est probable qu'un réglage du radiateur permette d'avoir une température de 20°C dans la pièce.

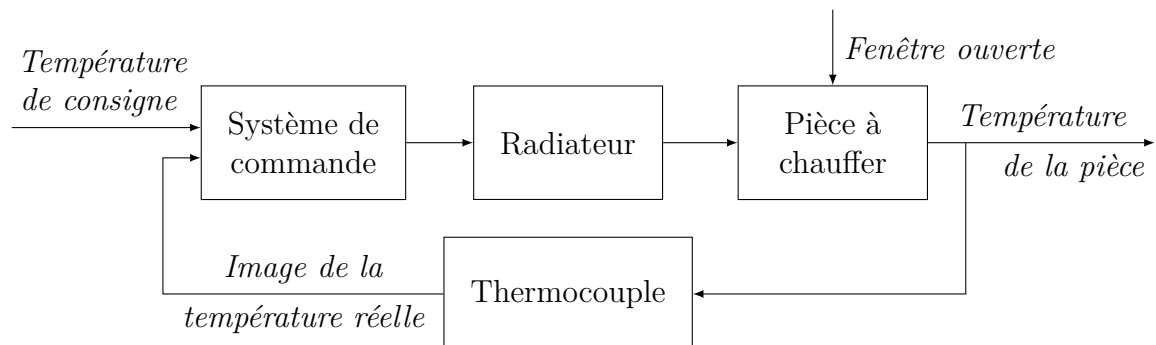


Seulement, si une fenêtre donnant sur l'extérieur, à une température plus faible, est ouverte (la perturbation), il est fort probable que la température de la pièce va diminuer et ne respectera plus la température de consigne. Pour recouvrer les 20°C souhaités, il sera alors nécessaire d'augmenter artificiellement la température de consigne pour augmenter l'apport de chaleur.

Dans le cas où une image de la réponse du système est renvoyée en temps réel au système de commande par un (ou plusieurs) capteur, il devient possible d'adapter la commande pour que la réponse du système corresponde au mieux à la consigne. On parle alors de structure de commande à rétroaction (*feedback* en anglais). Contrairement à un système de commande en boucle ouverte, un système de commande en boucle fermée « contrôle » en temps réel l'effet de son action. L'intérêt de ce contrôle est de pouvoir respecter une consigne tout en tenant compte de perturbations inconnues (ou mal connues) ou de variations imprévisibles du comportement. Cette structure de commande permet ainsi d'améliorer les performances dynamiques du système (rapidité, précision, insensibilité aux perturbations ou stabilité).

Exemple 1.2 (Chauffage d'une pièce)

Pour maintenir effectivement la température d'une pièce à 20°C, il est nécessaire de disposer d'un capteur de température appelé un thermocouple et d'agir sur la commande du radiateur. Pour que cette action soit automatisée, il est nécessaire de relier le capteur à un système de commande afin qu'il compare la valeur réelle de la température à la consigne et détermine la commande du radiateur en fonction.



Avec un tel système de commande, si une fenêtre donnant sur l'extérieur est ouverte, il est fort probable que la température de la pièce diminue momentanément : c'est l'effet de la perturbation. Seulement, avec une commande adaptée (plus élevée), l'apport de chaleur par les radiateurs augmentera afin que la pièce recouvre automatiquement les 20°C souhaités.

Si une structure de commande avec rétroaction permet d'améliorer les performances d'un processus, il est toutefois important de remarquer que cette structure de commande ne présente pas que des avantages : elle nécessite l'emploi de capteurs qui augmentent le coût d'une installation et les réglages de stabilité et de précision de ces systèmes sont plus délicats et nécessitent de mettre en œuvre des outils plus complexes.

1.2 Structure d'un système asservi

Les systèmes asservis forment la plus grande partie des systèmes automatisés que nous allons étudier cette année. Leur objectif principal est de maîtriser l'évolution d'une ou plusieurs grandeurs physiques (température, pression, vitesse, pH, ...) d'un processus en fonction de contraintes de fonctionnement fixées par un cahier des charges fonctionnel et ceci dans un environnement perturbé.

Définition 1.3 (Système asservi)

Un système asservi est l'association d'un processus, de capteurs et d'un système de commande permettant, à partir des consignes, de piloter une ou plusieurs grandeurs physiques du processus.

Les systèmes asservis sont caractérisés par la présence de **capteurs** permettant d'obtenir une image des grandeurs physiques que l'on souhaite piloter appelées les **grandeurs asservies**. La présence de capteurs est indispensable car ils permettent de contrôler l'état du processus dont le système de commande, par l'intermédiaire du régulateur, tient compte pour élaborer (corriger) sa commande. Les systèmes asservis sont généralement décomposés en une partie opérative (le processus), une partie d'observation (les capteurs) et une partie commande (figure 4).

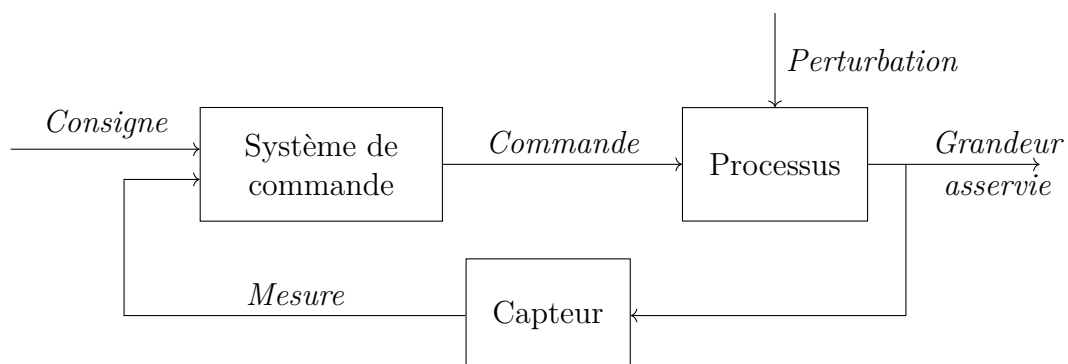


FIGURE 4 – Structure d'un système asservi.

Cette structure fait apparaître une chaîne directe comprenant une partie de la chaîne d'information associée au système de commande et la chaîne d'énergie associée au processus et une boucle de rétroaction comprenant un capteur. Ce capteur mesure en permanence l'évolution de la sortie et en retourne une image à la partie commande. Pour pouvoir être comparées, la consigne et la mesure doivent être de même nature ce qui nécessite la plupart du temps de mettre en forme la consigne avec un adaptateur de consigne. Ensuite, en fonction de l'écart fourni par le comparateur, une nouvelle commande va être calculée par le régulateur et envoyée au processus après amplification (figure 5).

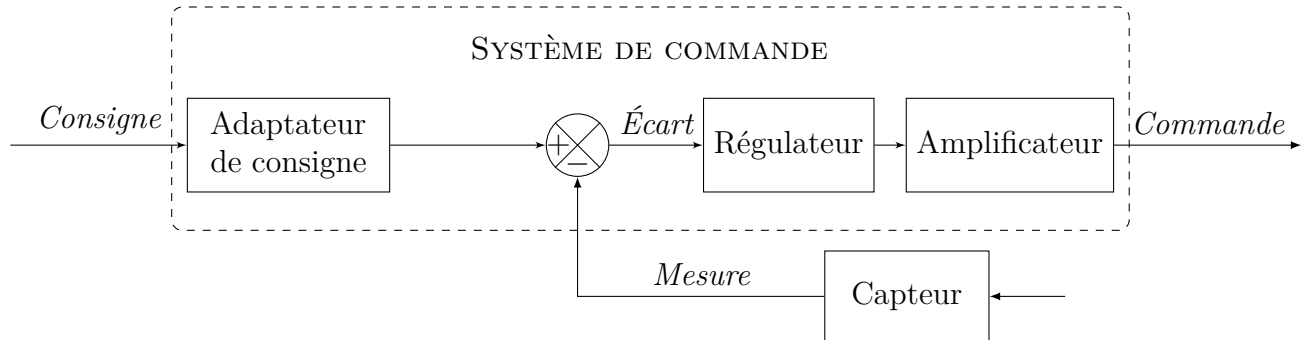


FIGURE 5 – Structure générique d'un système de commande.

Le système de commande et le processus sont reliés par l'information de commande reçue par le pré-actionneur. Ce composant, à l'interface des chaînes d'information et d'énergie, module l'apport en énergie de l'actionneur à partir de l'information de commande. L'actionneur qui lui est relié convertit un type d'énergie (électrique, hydraulique, etc.) en un autre (généralement en énergie mécanique). L'énergie est ensuite transmise, *via* les transmetteurs de puissance, à l'effecteur sur lequel s'appliquent les perturbations (figure 6).

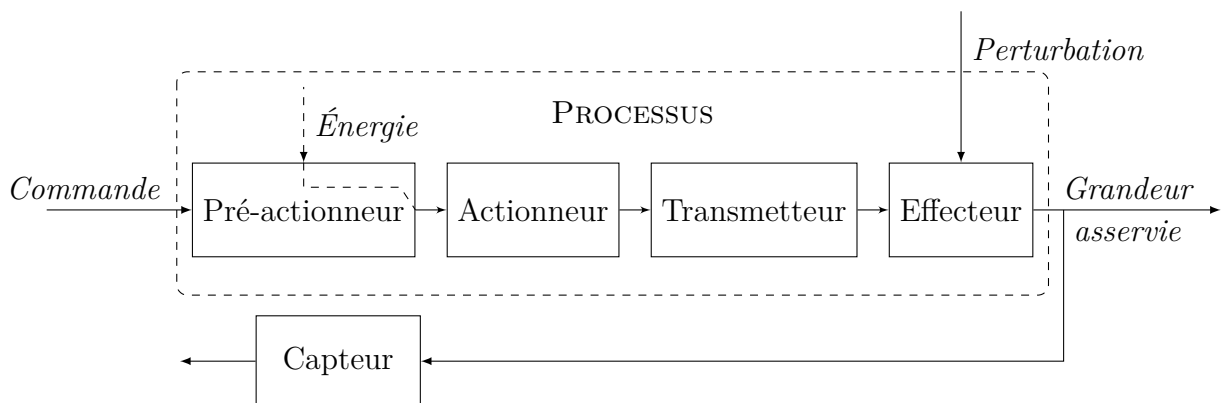


FIGURE 6 – Structure générique d'un processus.

En sortie du schéma fonctionnel d'un processus asservi, on retrouve toujours un point de prélèvement auquel est rattaché un capteur. Il permet de « prélever » sur le système la valeur de la grandeur physique asservie.

Les différents constituants permettant d'asservir un processus sont donc :

Capteur : il « prélève » sur le système la valeur de la grandeur régulée ou asservie (information physique) et la transforme en un signal « compréhensible » par le régulateur. La précision, la rapidité et l'amplitude de mesure sont les caractéristiques importantes d'un capteur.

Adapateur de consigne : il permet de mettre en forme la consigne pour qu'elle soit de même nature et donc comparable à la mesure fournie par le capteur.

Comparateur : il est chargé de comparer une image de la consigne et de la grandeur asservie (ici notée *mesure*). Le comparateur renvoie l'écart entre ces deux informations.

Régulateur : il est chargé d'élaborer la loi de commande et son évolution à partir de l'écart calculé.

Amplificateur : il amplifie le signal de commande pour passer d'un faible niveau d'énergie (information) à un niveau plus élevé.

On considère deux types de systèmes asservis :

Régulation : on appelle régulation un système asservi qui doit maintenir constante la sortie conformément à la consigne (constante) indépendamment des perturbations (régulation de température d'un four, régulateur de vitesse automobile) ;

Asservissement : on appelle asservissement un système asservi dont la sortie doit suivre le plus fidèlement possible la consigne quelle que soit son évolution (suivi de trajectoire d'un robot, profil de vitesse). Suivant la nature de la grandeur à asservir, on parle d'asservissement de position, de vitesse, etc.

1.3 Performances de la commande d'un système

Avant toute chose, précisons que le comportement des systèmes asservis repose sur un principe de causalité, c'est-à-dire que si un phénomène (nommé cause) produit un autre phénomène (nommé effet), alors l'effet ne peut précéder la cause. Ce qui se traduit en automatique par :

Définition 1.4 (Système causal)

Un système est dit causal si sa réponse à un instant donné ne dépend que des entrées précédant cet instant.

Le principe de causalité est très souvent invoqué en automatique pour rendre une théorie mathématique physiquement admissible. Suivant ce principe de causalité, le comportement d'un système asservi est généralement évalué suivant différents critères de performance :

- **stabilité** : la réponse du système converge-t-elle pour une entrée constante ?
- **précision** : le système asservi atteint-il la valeur de consigne ?
- **rapidité** : le combien de temps faut-il au système pour se stabiliser ?
- **dépassements** : la grandeur de sortie a-t-elle tendance à dépasser la valeur à convergence et osciller autour avant de se stabiliser ?
- **sensibilité aux perturbations** : les perturbations agissant sur le système modifient-elles la valeur à convergence de la grandeur asservie ?

Sachant que la réponse d'un système dépend évidemment des signaux auxquels il est soumis, ces critères de performances sont généralement évalués à partir de la réponse à une entrée constante nommée échelon.

Définition 1.5 (Stabilité)

Un système est stable si pour toute entrée bornée, la sortie est bornée.

La stabilité traduit la propriété de convergence temporelle vers un état d'équilibre. La réponse à un échelon d'un système stable est caractérisée par une sortie bornée qui converge de façon asymptotique vers une valeur finale (figure 7a). Dans le cas où sa réponse diverge, le système sera dit instable (figure 7b).

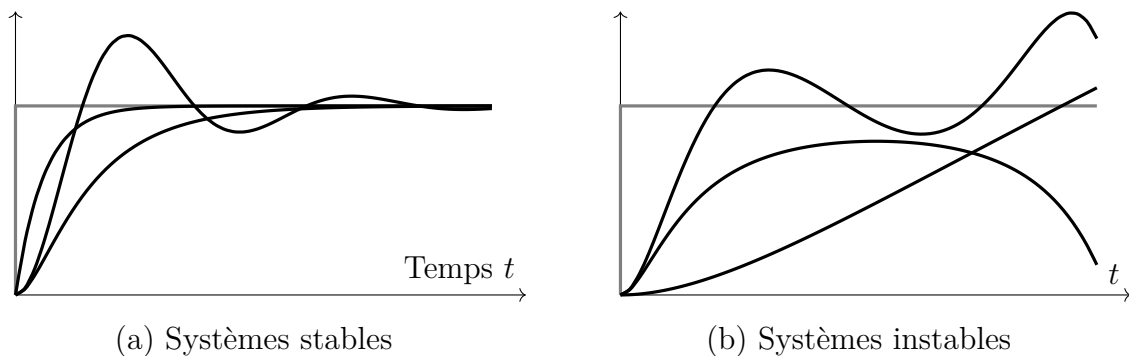


FIGURE 7 – Réponses de systèmes stables (a) et instables (b).

Définition 1.6 (Rapidité)

Un système est rapide s'il converge en un temps court au regard de son contexte d'utilisation.

La rapidité est associée au temps de réaction du système soumis à une variation brusque de la grandeur d'entrée (temps de réponse). La valeur à convergence (ou finale) étant le plus souvent atteinte de manière asymptotique, le critère d'évaluation de la rapidité d'un système est le temps de réponse à n %. Le critère le plus utilisé en ingénierie est le temps de réponse à 5 %.

Définition 1.7 (Temps de réponse à 5 %)

Le temps de réponse à 5 % d'un système, noté $t_{5\%}$, correspond au temps mis par la réponse de ce système soumis à un échelon pour entrer définitivement dans une bande à ± 5 % de sa valeur finale.

Pour déterminer le temps de réponse à 5 % d'un système soumis à un échelon, il faut :

- trouver la valeur asymptotique (finale) de la réponse du système (à ne pas confondre avec la consigne) ;
- tracer deux droites parallèles à ± 5 % de l'asymptote ;
- chercher le dernier point d'entrée de la réponse dans la bande (point à partir duquel la courbe ne sort plus de la bande) dont l'abscisse correspond au $t_{5\%}$ recherché.

Cette construction est illustrée sur la figure 8.

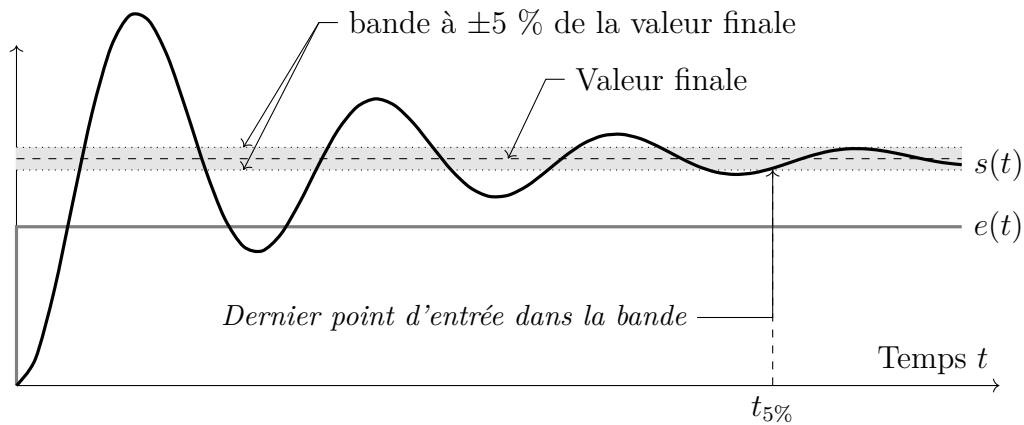


FIGURE 8 – Caractérisation de la rapidité d'un système par son temps de réponse à 5 %.

Définition 1.8 (Précision)

La précision qualifie l'aptitude d'un système à respecter la consigne.

La précision peut être soit absolue, soit relative mais est toujours associée à une quantification de l'erreur avec laquelle la sortie réalisée suit la loi de sortie désirée.

Définition 1.9 (Erreur)

L'erreur $\varepsilon(t)$ est la différence entre la consigne $e(t)$ et la sortie $s(t)$:

$$\varepsilon(t) = e(t) - s(t)$$

Elle n'est définie que si la consigne et la sortie sont de même nature.

Définition 1.10 (Erreur statique)

L'erreur statique ε_S est la limite au temps long de l'erreur associée à la réponse d'un système soumis à un échelon :

$$\varepsilon_S = \lim_{t \rightarrow +\infty} \varepsilon(t)$$

Un système est dit **précis** si l'erreur statique est nulle.

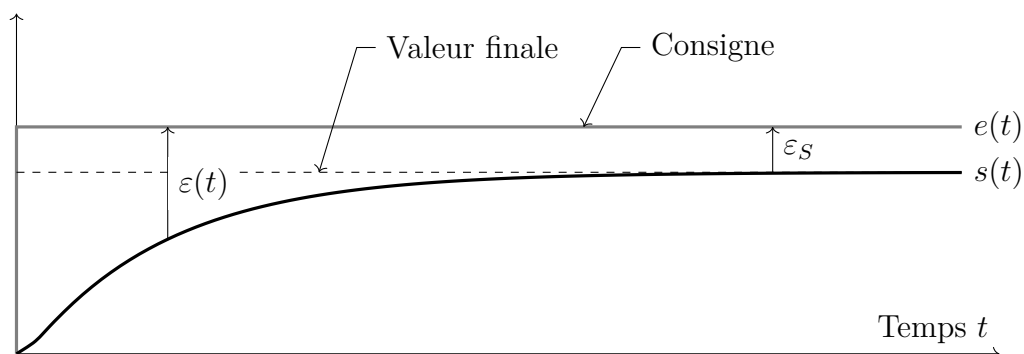


FIGURE 9 – Caractérisation de la précision d'un système.

Définition 1.11 (Dépassements)

Un système stable présente des dépassements si sa réponse oscille de façon transitoire autour de la valeur finale avant de l'atteindre.

Les dépassements sont définis par rapport à la valeur finale de la réponse d'un système à un échelon. Ils peuvent être inférieurs ou supérieurs à cette valeur et sont numérotés par ordre croissant d'apparition. La convergence de la réponse vers une valeur finale implique que l'amplitude des dépassements diminue progressivement : on parle d'**amortissement** du régime transitoire de la réponse d'un système (figure 10).

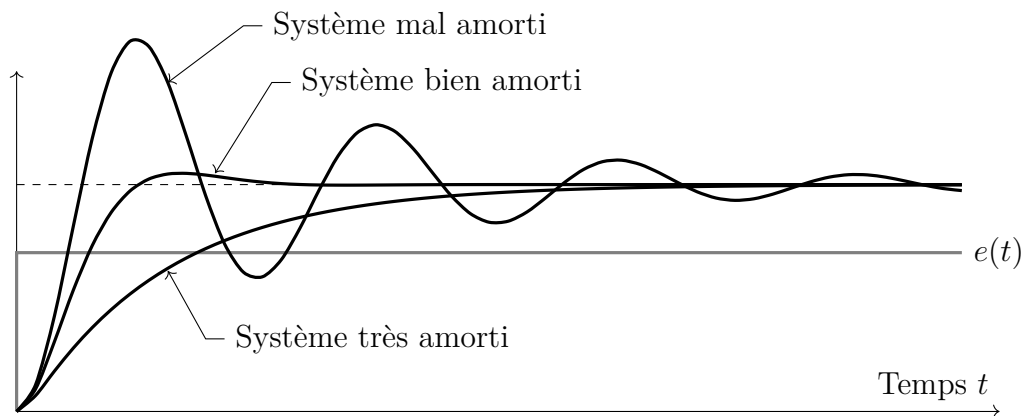


FIGURE 10 – Caractérisation des dépassements de la réponse d'un système.

Suivant l'application étudiée, un cahier des charges peut exiger – souvent pour des questions de sécurité – une amplitude maximale des dépassements voire l'absence de dépassement, par exemple pour éviter toute collision ou débordement d'un vase lors de son remplissage. Comme l'amplitude des dépassements est le plus souvent proportionnelle à la consigne, la valeur maximale est souvent spécifiée en pourcentage de la valeur finale. Le premier dépassement ayant le plus souvent l'amplitude la plus importante, c'est le plus pénalisant et donc celui qui est comparé au cahier des charges.

Définition 1.12 (Sensibilité aux perturbations)

Un système est sensible aux perturbations s'il ne converge pas vers les mêmes valeurs finales selon qu'une perturbation s'applique ou non.

Un système insensible aux perturbations peut voir sa grandeur de sortie évoluer de façon transitoire après l'apparition d'une perturbation mais, une fois stabilisé, tendra vers la même valeur finale qu'avant la perturbation (figure 11b). À l'opposé, un système sensible aux perturbations atteindra, après une période transitoire, une valeur finale différente de celle qu'il aurait atteinte sans perturbation (figure 11a).

La sensibilité aux perturbations est étroitement liée à la notion de **robustesse** d'un système. Un système insensible aux perturbations sera d'autant plus robuste que les effets (transitoires) d'une perturbation sur son comportement sont faibles.

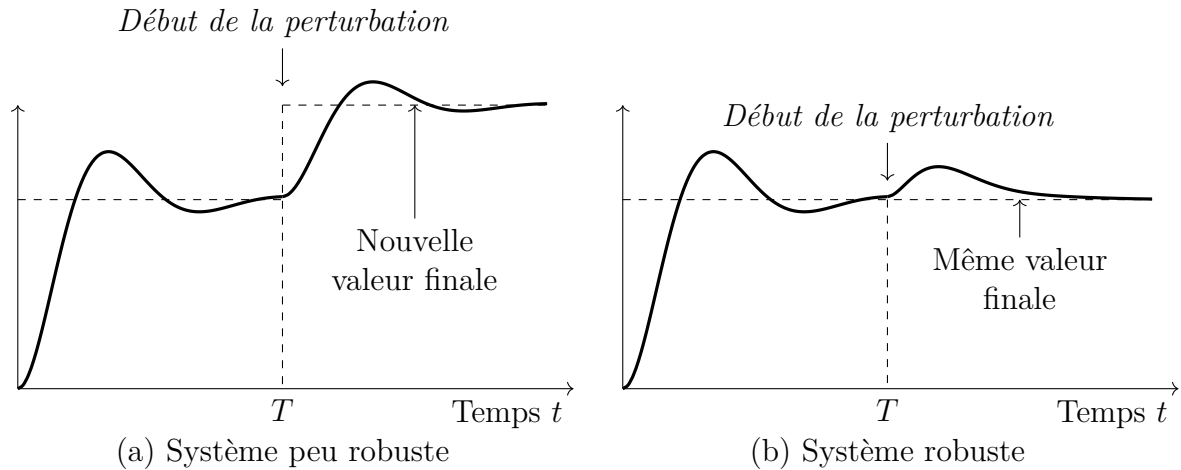


FIGURE 11 – Influence d’une perturbation à l’instant T sur la réponse d’un système sensible (a) et insensible (b) aux perturbations.

Remarque 1.1 (Performances des systèmes asservis)

Dans l’étude des systèmes asservis, on cherchera à rendre compatible les quatre notions de performance que sont la stabilité, précision, la rapidité et la robustesse.

2 Modélisation du comportement dynamique d’un système

Pour garantir les spécifications imposées par le cahier des charges d’un asservissement (stabilité, rapidité, précision, robustesse), on ne peut pas choisir et dimensionner le système de commande au hasard. L’obtention des meilleures performances nécessite de tenir compte des propriétés et paramètres du système dynamique à asservir. L’ensemble des propriétés qui gouvernent le comportement du système est avantageusement condensé dans le modèle mathématique du système.

Définition 2.1 (Modèle mathématique d’un système)

Le modèle mathématique d’un système dynamique correspond à l’ensemble des lois mathématiques régissant la causalité entre ses entrées et ses sorties.

Parmi les modèles mathématiques de systèmes dynamiques, on peut distinguer :

- les modèles de connaissance, construits de façon analytique à partir des lois de la physique régissant le comportement des constituants d’un système ;
- les modèles de comportement, construits par identification à partir de résultats expérimentaux obtenus avec les systèmes réels.

Dans ce qui suit, nous ne présentons que les outils permettant d’établir les modèles de connaissances des systèmes dynamiques. Les démarches associées à l’identification seront présentées dans le chapitre suivant.

2.1 Objectif de la modélisation

Le processus de développement d'un modèle mathématique à partir d'une réalité physique est appelé modélisation. C'est une étape préliminaire de l'**analyse** d'un système devenue essentielle pour la synthèse des systèmes de commande. L'objectif de la modélisation du comportement dynamique d'un système est de le décrire avec un ensemble d'équations permettant de simuler et donc de prédire son comportement. Pour cela, il est nécessaire de :

- proposer une modélisation des variables d'entrée (consignes, perturbations, etc.) ;
- modéliser le système global ;
- simuler le comportement du système (à la main ou par ordinateur) ;
- analyser la réponse de la simulation du comportement du système, c'est-à-dire analyser l'évolution de la sortie du modèle du système.

Pour que les prédictions du modèle soient fiables, il est indispensable qu'il ne soit pas trop simpliste – au risque de ne pas représenter assez bien la réalité – mais quand même suffisamment simple pour ne pas rendre inutilement complexes les étapes d'analyse des propriétés du système. Pour établir un bon compromis entre la précision d'un modèle et sa complexité, il est nécessaire de formuler un certain nombre d'hypothèses de modélisation.

2.2 Hypothèses de modélisation

Le modèle de connaissance d'un système est obtenu en écrivant les lois de la physique régissant le comportement de ses constituants. Cette étape résulte en l'écriture d'équations différentielles et algébriques. Un certain nombre d'hypothèses de travail sont formulées, définissant la classe des modèles utilisés. Parmi toutes les classes de modèles, on se restreindra à la classe la plus simple et la plus couramment utilisée : les **systèmes linéaires continus à temps invariant** (SLCI).

2.2.1 Systèmes continus

Définition 2.2 (Système continu)

Un système est dit continu si les fonctions définissant ses entrées et sorties sont continues, c'est-à-dire définies pour tout instant.

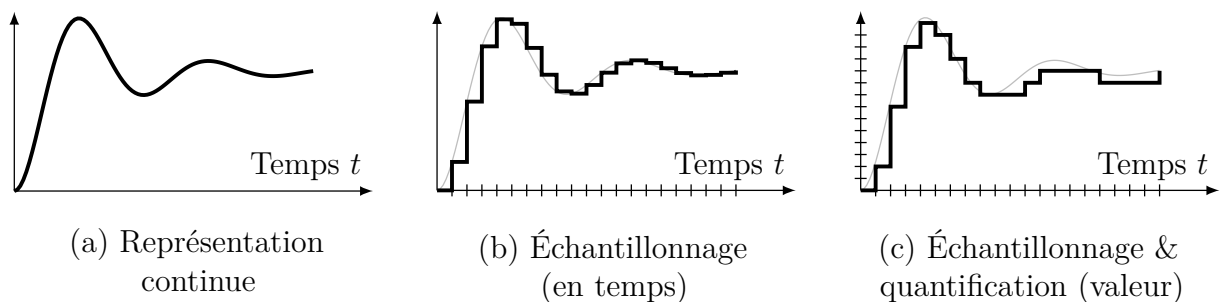


FIGURE 12 – Représentation numérique des données.

Les signaux d'un système continu sont forcément de nature analogique. Cependant, les systèmes de commande modernes manipulent des données numériques, c'est-à-dire échantillonnées (valeurs prélevées toutes les T secondes) et quantifiées (valeurs numériques décomposées en N pas). Toutefois, si la période d'échantillonnage est très inférieure au temps de réponse du système et que le nombre de pas de quantification est suffisamment grand, alors les données numériques pourront être assimilées à celles d'un système continu.

2.2.2 Systèmes linéaires

Les modèles de connaissance des systèmes sont établis à partir d'équations reliant les fonctions continues d'entrée et de sortie et leurs dérivées appelées équations différentielles.

Définition 2.3 (Système linéaire)

Un système linéaire est un système continu décrit par des équations différentielles linéaires.

Un système linéaire vérifie les principes de proportionnalité et de superposition.

Définition 2.4 (Principe de proportionnalité)

Un système vérifie le principe de proportionnalité (ou d'homogénéité) si à un multiple d'une entrée quelconque correspond le même multiple de la sortie correspondante.

Si $s(t)$ est la réponse à l'entrée $e(t)$ alors $\lambda s(t)$ est la réponse à l'entrée $\lambda e(t)$ (figure 13).

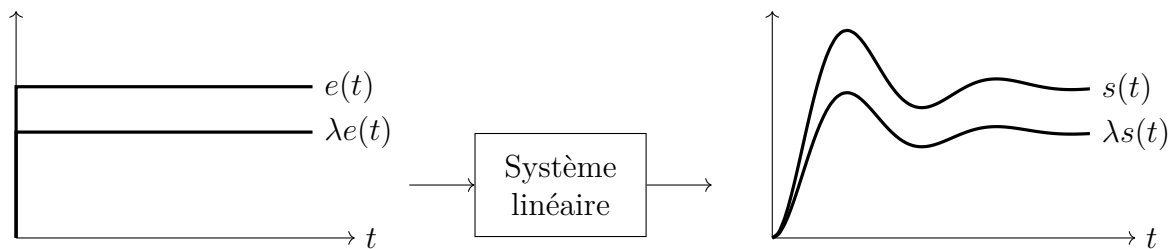


FIGURE 13 – Principe de proportionnalité.

Définition 2.5 (Principe de superposition)

Un système vérifie le principe de superposition si à la somme de deux entrées quelconques correspond la somme des deux sorties correspondantes.

Si $s_1(t)$ et $s_2(t)$ sont les réponses respectives des entrées $e_1(t)$ et $e_2(t)$, alors $s_1(t) + s_2(t)$ est la réponse à l'entrée $e_1(t) + e_2(t)$ (figure 14).

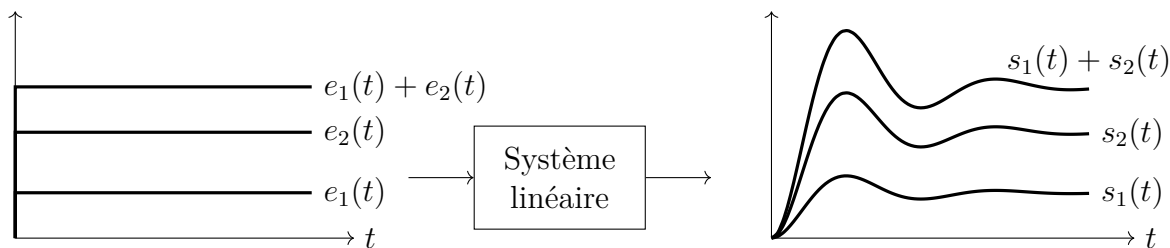


FIGURE 14 – Principe de superposition.

La généralisation de ces deux principes est que la réponse $s(t)$ d'un système linéaire soumis à une combinaison linéaire des entrées $e_1(t)$ et $e_2(t)$ est la combinaison linéaire des réponses à chacune d'elle, respectivement notées $s_1(t)$ et $s_2(t)$. Ce qui se traduit par :

$$\forall \alpha_1, \alpha_2 \in \mathbb{R}, \quad \alpha_1 e_1(t) + \alpha_2 e_2(t) \longrightarrow \boxed{\text{Système linéaire}} \longrightarrow \alpha_1 s_1(t) + \alpha_2 s_2(t)$$

Remarque 2.1 (Linéarité du comportement, pas de la réponse !)

Il est important de noter que la linéarité dont il est question ici est celle de la relation de comportement du système et non celle de l'évolution temporelle de sa sortie qui est très souvent non-linéaire. Cette linéarité peut être appréhendée facilement à partir de la caractéristique du système $s = f(e)$ obtenue en relevant les valeurs asymptotiques de la réponse du système pour des échelons de différentes amplitudes. Dans le cas de systèmes linéaires, cette caractéristique est une droite, de pente K appelée le gain du système.

Les caractéristiques des systèmes réels ne sont que très rarement linéaires sur la totalité de leur domaine d'utilisation et peuvent présenter des non-linéarités :

- de courbure (saturation d'un amplificateur) ;
- de saturation (circuit électronique, distributeur, commande en butée mécanique,...) ;
- de seuil (frottement sec) ;
- d'hystérésis (jeux mécaniques).

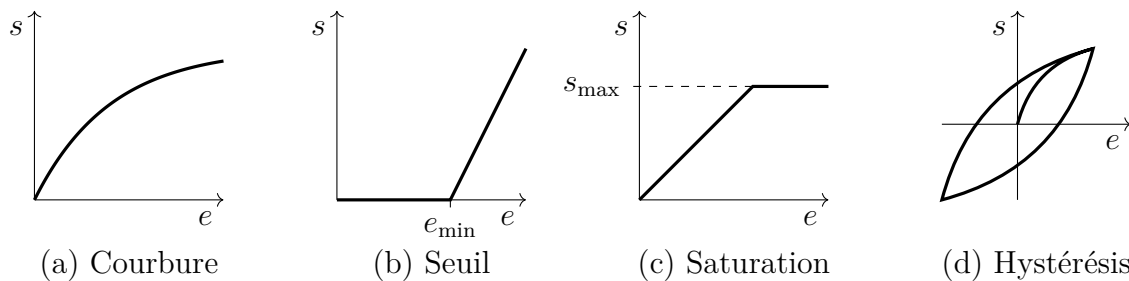


FIGURE 15 – Exemples de non-linéarités.

Dans la pratique pour étudier ces systèmes, on est amené à limiter la zone d'étude du fonctionnement à une région où le comportement est « à peu près » linéaire. Si ce n'est pas possible, il est nécessaire de linéariser la caractéristique du système au voisinage d'un point de fonctionnement. Il s'agit le plus souvent d'une approximation linéaire tangente au point de fonctionnement étudié où la courbe de la caractéristique est remplacée par une droite : sa tangente (figure 16). Le comportement du système est alors dit linéarisé autour du point de fonctionnement étudié. En pratique, il conviendra de toujours de préciser le domaine de validité de cette approximation.

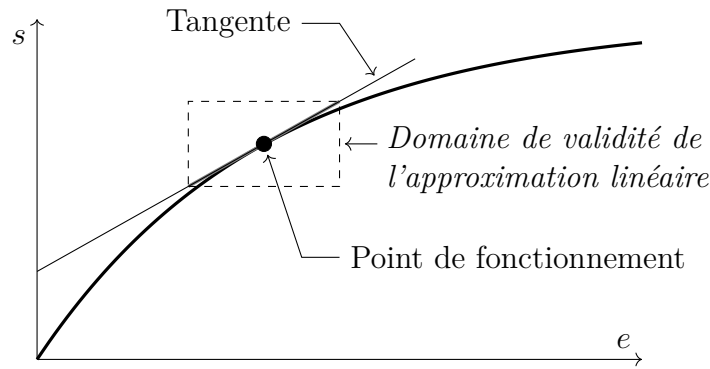


FIGURE 16 – Linéarisation autour d'un point de fonctionnement.

2.2.3 Systèmes invariants

Parmi les hypothèses de modélisation, on ajoutera que les systèmes étudiés possèdent toujours les mêmes caractéristiques et propriétés physiques (masse, dimensions, etc.) que l'on considérera donc comme des constantes indépendantes du temps.

Définition 2.6 (Système invariant)

Un système est dit invariant si ses caractéristiques et propriétés régissant son comportement sont invariables dans le temps.

Si une même entrée se produit à deux instants distincts t_1 et $t_1 + \tau$ alors les deux sorties temporelles $s_1(t)$ et $s_2(t)$ seront identiques et simplement décalées de τ (figure 17).

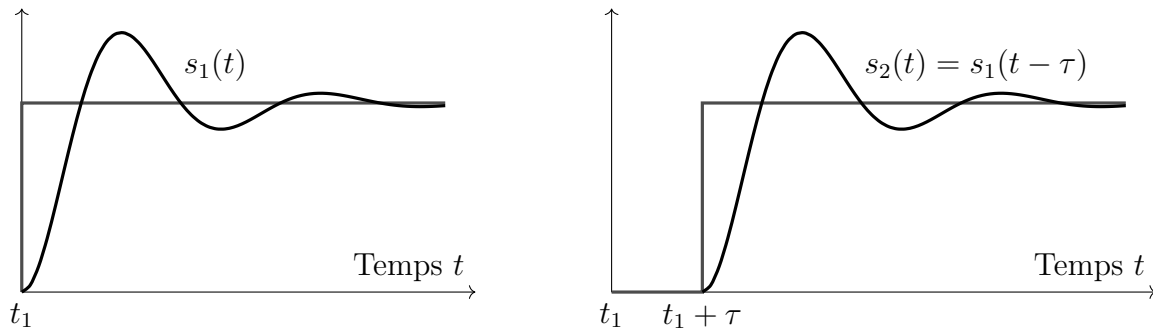


FIGURE 17 – Invariance dans le temps du comportement d'un système.

2.2.4 Systèmes mono-variables

Enfin, et pour des raisons de simplicité, on se restreindra à l'étude des systèmes mono-variables ou SISO (*Single Input Single Output*).

Définition 2.7 (Système mono-variable)

Un système mono-variable ne possède qu'une seule entrée et une seule sortie.

L'étude des systèmes multi-variables ou MIMO (*Multiple Input Multiple Output*) sera abordée en écoles d'ingénieurs.

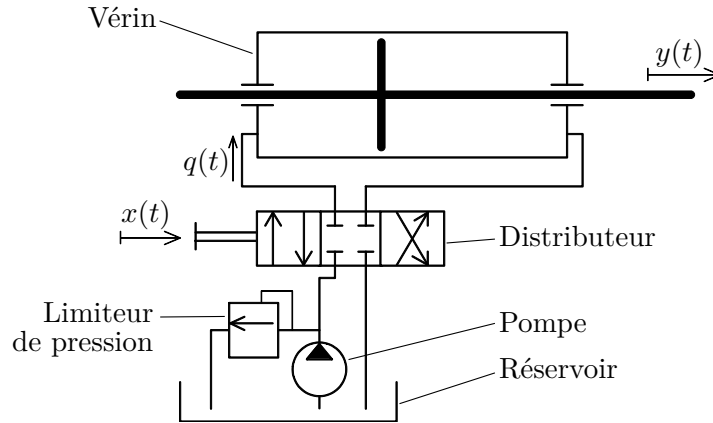
2.3 Systèmes linéaires continus invariants

Pour modéliser le comportement dynamique des systèmes, on se limitera donc aux systèmes linéaires continus et invariants mono-variables. Si H représente le modèle mathématique d'un système – qui relie la sortie $s(t)$ et l'entrée $e(t)$ –, alors :

- un système est linéaire : $H(\alpha_1 e_1 + \alpha_2 e_2) = \alpha_1 H(e_1) + \alpha_2 H(e_2)$
- un système est continu : la fonction H est une fonction continue du temps t ;
- un système est invariant : la réponse du système est indépendante du temps, c'est-à-dire si $e(t)$ induit $s(t)$ alors $e(t + \tau)$ induit $s(t + \tau)$.

Il est à noter que la propriété de linéarité permet de prendre en compte facilement une perturbation : le signal de sortie sera la somme du signal issu de l'entrée $e(t)$ et du signal issu de la perturbation $p(t)$.

Exemple 2.1 (Commande de vérin hydraulique)



Lorsque la commande $x(t)$ du distributeur est non nulle, le débit de fluide $q(t)$ en entrée et en sortie du vérin se traduit par un déplacement $y(t)$ de la tige du vérin. Déterminons la relation $y = H(x)$ sachant que :

- la caractéristique du distributeur (pré-actionneur) induit une relation de proportionnalité entre l'entrée $x(t)$ et la sortie du distributeur, qui correspond au débit de fluide $q(t)$ entrant dans le vérin (actionneur) :

$$q(t) = K x(t)$$

- la section utile du vérin, notée S , permet de relier le déplacement de la tige du vérin $y(t)$ avec le débit dans le vérin selon

$$q(t) = \frac{d(S \cdot y(t))}{dt}$$

En combinant ces deux relations, on obtient immédiatement l'expression de la vitesse du déplacement de la tige du vérin :

$$\dot{y}(t) = \frac{dy}{dt}(t) = \frac{1}{S} q(t) = \frac{K}{S} x(t)$$

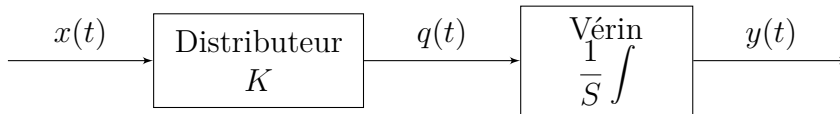
Pour un signal causal ($x(t) = 0$ pour $t < 0$) et une condition initiale nulle ($y(0) = 0$), on peut finalement écrire :

$$y(t) = \frac{K}{S} \int_0^t x(\tau) d\tau$$

Le modèle mathématique associé à l'ensemble distributeur vérin peut alors se mettre **formellement** sous la forme :

$$H = \frac{K}{S} \int$$

On peut vérifier les propriétés mathématiques de H . Sous forme de blocs, on peut décomposer ce système de la manière suivante :



Description par des équations différentielles

Sous les hypothèses de continuité, de linéarité et d'invariance dans le temps, la relation de comportement dynamique d'un système peut se mettre sous la forme d'une équation différentielle linéaire à coefficients constants de la forme :

$$a_0 e(t) + a_1 \frac{de}{dt}(t) + \dots + a_m \frac{d^m e}{dt^m}(t) = b_0 s(t) + b_1 \frac{ds}{dt}(t) + \dots + b_n \frac{d^n s}{dt^n}(t)$$

où $a_0, a_1, \dots, a_m$ et $b_0, b_1, \dots, b_n$ sont des constantes réelles. Pour que le système soit physiquement réalisable, il doit vérifier le principe de causalité qui implique que l'on aura toujours $n \geq m$, avec n l'ordre du système.

Remarque 2.2 (Notation des dérivées)

On note :

$$\dot{y}(t) = \frac{dy}{dt}(t) \quad \text{dérivée 1^{re} de } y \text{ par rapport au temps}$$

$$\ddot{y}(t) = \frac{d^2 y}{dt^2}(t) \quad \text{dérivée 2^{de} de } y \text{ par rapport au temps}$$

et plus généralement $\frac{d^n y}{dt^n}(t)$ la dérivée n -ième de la fonction y par rapport au temps.

À partir de cette représentation il est possible de déterminer l'évolution temporelle de la réponse d'un système en résolvant l'équation différentielle. Les hypothèses de modélisation des systèmes (linéaires, continus et invariants) permettent d'envisager une résolution analytique de ces équations, par exemple à partir de la transformée de Laplace.

3 Résolution par la transformée de Laplace

3.1 Définition et intérêts

La modélisation des systèmes dynamiques conduit toujours à l'expression d'une ou plusieurs équations intégro-différentielles. Une équation intégro-différentielle est une équation dans laquelle la fonction à rechercher figure avec ses intégrales et dérivées. De façon « classique », la résolution d'une équation intégro-différentielle suit une démarche en deux temps permettant d'assembler les composantes transitoires et permanentes (figure 18).

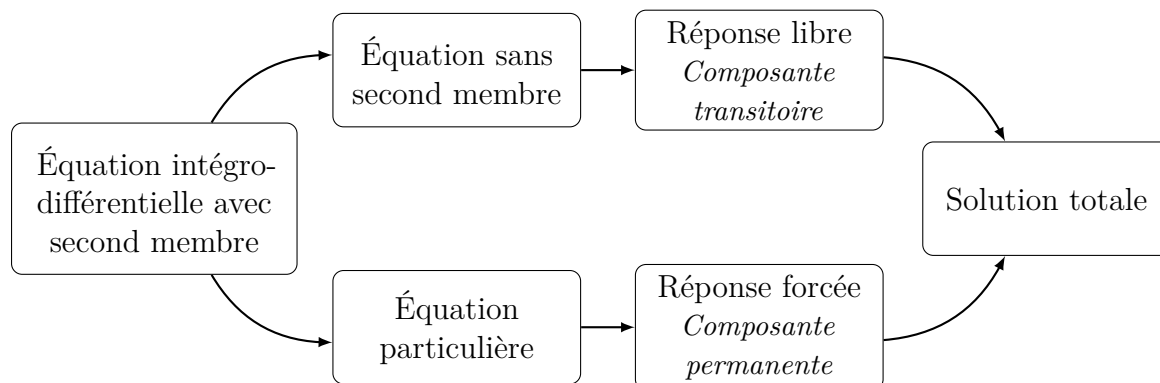


FIGURE 18 – Méthode de résolution « classique » d'une équation intégro-différentielle.

Cette façon de résoudre les équations devient vite complexe. La transformation de Laplace permet de résoudre simplement ces **équations différentielles** linéaires à coefficients constants en les transformant en **équations algébriques**.

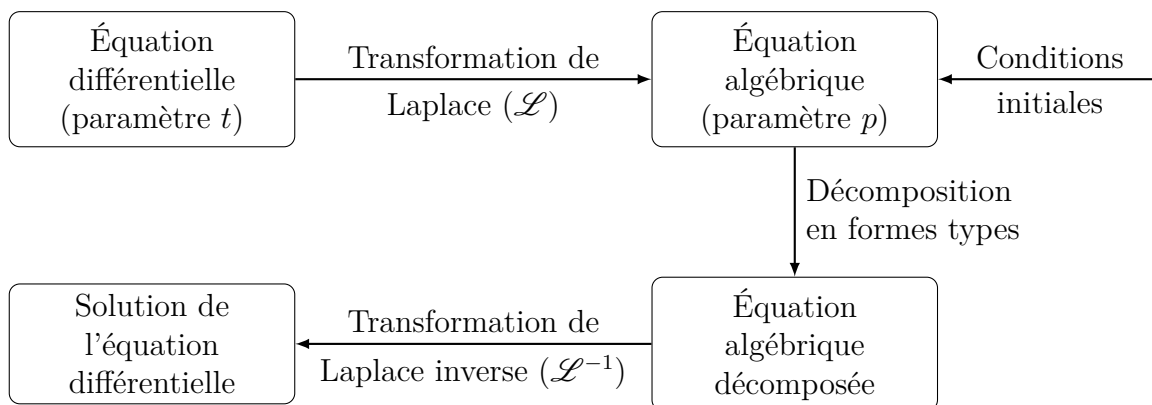


FIGURE 19 – Technique de résolution utilisée par les automaticiens.

L'équation différentielle de la variable réelle t est transformée en une équation algébrique de la variable complexe p par une **transformation dite de Laplace** qui tient compte des conditions initiales. Après mise sous forme « type » de l'équation algébrique, on obtient la solution temporelle de l'équation différentielle par **transformée inverse**.

Définition 3.1 (Transformée de Laplace)

Soit f une fonction causale continue par morceaux sur $\mathbb{R}^+$ telle que $f(t) = 0$ pour tout $t < 0$. On appelle transformée de Laplace de f , la fonction $F = \mathcal{L}[f]$ de la variable complexe p définie par la relation intégrale :

$$F(p) = \lim_{\tau \rightarrow +\infty} \int_0^\tau e^{-pt} f(t) dt = \int_0^{+\infty} e^{-pt} f(t) dt$$

pour toutes les valeurs de $p \in \mathbb{C}$ pour laquelle cette intégrale impropre converge.

La transformée de Laplace permet de passer du domaine temporel au domaine dit de Laplace. $f(t)$ est appelée *originale* et sa transformée $F(p)$ est appelée *image*. La variable complexe p est la variable conjuguée de t . Elle est parfois notée $p = \sigma + j\omega$ avec j la variable complexe. Comme t représente le temps, la dimension de p est une fréquence (l'inverse d'un temps).

Remarque 3.1 (Notation des fonctions images)

Il est d'usage de noter une fonction temporelle avec une lettre minuscule et sa transformée de Laplace avec une lettre majuscule. Dans le cas où la notation d'une fonction temporelle est déjà une lettre majuscule, on conserve la même lettre. Par exemple :

$$\mathcal{L}[f(t)] = F(p) , \quad \mathcal{L}[C(t)] = C(p) , \quad \mathcal{L}[\theta(t)] = \Theta(p) \quad \text{ou} \quad \mathcal{L}[\omega(t)] = \Omega(p)$$

La transformée de Laplace ainsi définie ne s'applique qu'aux fonctions causales. Pour rendre causale une fonction temporelle $f(t)$, on la combine avec la fonction de Heaviside (ou échelon unité) notée $u(t)$ et définie par :

$$u(t) = \begin{cases} 0 & \text{si } t < 0 \\ 1 & \text{sinon} \end{cases}$$

Ainsi, si $f(t)$ est une fonction quelconque, définie sur $\mathbb{R}$, la fonction $f(t)u(t)$ est son image causale, nulle pour tout $t < 0$ (figure 20).

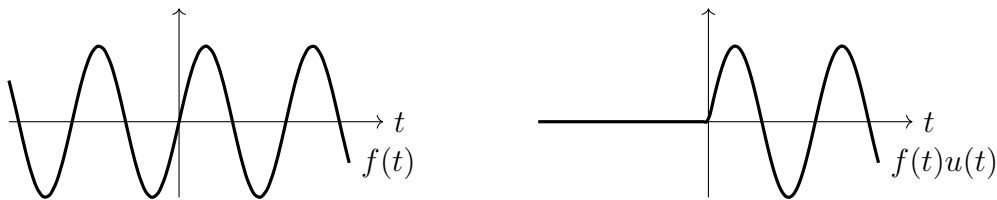


FIGURE 20 – D'une fonction quelconque à son image causale.

On veillera à toujours prendre l'image causale d'une fonction avant de calculer sa transformée de Laplace.

3.2 Propriétés élémentaires

Propriété 3.1 (Unicité)

À une fonction f correspond une seule transformée de Laplace $F = \mathcal{L}[f]$.

Démonstration. Pour évaluer l'unicité de la transformée de Laplace, considérons qu'il existe deux transformées de Laplace de $f(t)$ notées $F_1(p)$ et $F_2(p)$ et évaluons leur différence :

$$\begin{aligned} F_1(p) - F_2(p) &= \int_0^{+\infty} e^{-pt} f(t) dt - \int_0^{+\infty} e^{-pt} f(t) dt \\ &= \int_0^{+\infty} e^{-pt} (f(t) - f(t)) dt = \int_0^{+\infty} 0 dt = 0 \end{aligned}$$

Comme la différence entre ces deux transformées de Laplace d'une même fonction est nulle, alors la transformée de Laplace est unique. $\square$

Théorème 3.1 (Linéarité)

La transformée de Laplace d'une combinaison linéaire de deux fonctions f et g (dont la transformée de Laplace existe) est la combinaison linéaire des transformées de Laplace de ces deux fonctions.

$$\forall \alpha, \beta \in \mathbb{R}, \quad \mathcal{L}[\alpha f(t) + \beta g(t)] = \alpha F(p) + \beta G(p)$$

Démonstration. Par définition de la transformée de Laplace :

$$\mathcal{L}[\alpha f(t) + \beta g(t)] = \int_0^{+\infty} e^{-pt} (\alpha f(t) + \beta g(t)) dt$$

or comme α et β sont deux réels indépendants du temps, alors :

$$\mathcal{L}[\alpha f(t) + \beta g(t)] = \alpha \underbrace{\int_0^{+\infty} e^{-pt} f(t) dt}_{F(p)} + \beta \underbrace{\int_0^{+\infty} e^{-pt} g(t) dt}_{G(p)}$$

$\square$

Théorème 3.2 (Dérivation)

Si f est une fonction dérivable dont la transformée de Laplace existe alors la transformée de Laplace de sa dérivée est définie par :

$$\mathcal{L}\left[\frac{df}{dt}(t)\right] = pF(p) - f(0)$$

Démonstration. Par définition de la transformée de Laplace :

$$\mathcal{L} \left[\frac{df}{dt}(t) \right] = \int_0^{+\infty} \underbrace{e^{-pt}}_u \underbrace{\left(\frac{df}{dt}(t) \right)}_{v'} dt$$

où l'on reconnaît immédiatement une forme d'intégration par parties du type

$$\int u v' dt = [u v] - \int u' v dt$$

conduisant à

$$\mathcal{L} \left[\frac{df}{dt}(t) \right] = [e^{-pt} f(t)]_0^{+\infty} - \int_0^{+\infty} \frac{de^{-pt}}{dt} f(t) dt$$

Pour le premier terme, en notant que

$$\lim_{t \rightarrow \infty} e^{-pt} = 0$$

on obtient :

$$[e^{-pt} f(t)]_0^{+\infty} = \underbrace{\lim_{t \rightarrow \infty} e^{-pt} f(t)}_{\rightarrow 0} - f(0) = -f(0)$$

et pour le second terme, il faut simplement dériver la fonction exponentielle pour obtenir :

$$- \int_0^{+\infty} \frac{de^{-pt}}{dt} f(t) dt = \int_0^{+\infty} p e^{-pt} f(t) dt = p \int_0^{+\infty} e^{-pt} f(t) dt = p \mathcal{L} [f(t)]$$

car la variable p ne dépend pas de t . Soit finalement :

$$\mathcal{L} \left[\frac{df}{dt}(t) \right] = p \mathcal{L} [f(t)] - f(0)$$

□

Théorème 3.3 (Dérivation d'ordre 2)

Si f est une fonction deux fois dérivable et dont la transformée de Laplace existe alors la transformée de Laplace de sa dérivée seconde est définie par :

$$\mathcal{L} \left[\frac{d^2 f}{dt^2}(t) \right] = p^2 F(p) - pf(0) - \dot{f}(0)$$

Démonstration. Par une démarche identique à celle utilisée pour la dérivée première, on en déduit :

$$\begin{aligned} \mathcal{L} \left[\frac{d^2 f}{dt^2}(t) \right] &= \int_0^{+\infty} \underbrace{e^{-pt}}_u \underbrace{\left(\frac{d^2 f}{dt^2}(t) \right)}_{v'} dt \\ &= \left[e^{-pt} \frac{df}{dt}(t) \right]_0^{+\infty} + \int_0^{+\infty} p e^{-pt} \frac{df}{dt}(t) dt \\ &= \underbrace{\lim_{t \rightarrow \infty} e^{-pt} \frac{df}{dt}(t)}_{\rightarrow 0} - \dot{f}(0) + p \underbrace{\int_0^{+\infty} e^{-pt} \frac{df}{dt}(t) dt}_{p \mathcal{L}[f(t)] - f(0)} \\ &= p^2 \mathcal{L} [f(t)] - pf(0) - \dot{f}(0) \end{aligned}$$

□

Théorème 3.4 (Dérivation d'ordre n)

Si f est une fonction n fois dérivable et dont la transformée de Laplace existe alors la transformée de Laplace de sa n -ième dérivée est définie par :

$$\mathcal{L} \left[\frac{d^n f}{dt^n}(t) \right] = p^n F(p) - p^{n-1} f(0) - p^{n-2} \dot{f}(0) - \dots - \frac{d^{n-1} f}{dt^{n-1}}(0)$$

Démonstration. Par récurrence, on en déduit immédiatement :

$$\begin{aligned} \mathcal{L} \left[\frac{d^n f}{dt^n}(t) \right] &= \left[e^{-pt} \frac{d^{n-1} f}{dt^{n-1}}(t) \right]_0^{+\infty} + p \int_0^{+\infty} e^{-pt} \frac{d^{n-1} f}{dt^{n-1}}(t) dt \\ &= -\frac{d^{n-1} f}{dt^{n-1}}(0) + p \left(\left[e^{-pt} \frac{d^{n-2} f}{dt^{n-2}}(t) \right]_0^{+\infty} + p \int_0^{+\infty} e^{-pt} \frac{d^{n-2} f}{dt^{n-2}}(t) dt \right) \\ &= -\frac{d^{n-1} f}{dt^{n-1}}(0) - p \frac{d^{n-2} f}{dt^{n-2}}(0) + p^2 (\dots) \\ &= -\frac{d^{n-1} f}{dt^{n-1}}(0) - p \frac{d^{n-2} f}{dt^{n-2}}(0) - \dots - p^{n-1} f(0) + p^n \underbrace{\int_0^{+\infty} e^{-pt} f(t) dt}_{F(p)} \quad \square \end{aligned}$$

Remarque 3.2 (Dérivation et conditions de Heaviside)

Si le système est dans des conditions de Heaviside, c'est-à-dire si $f(0) = 0$, $\dot{f}(0) = 0$ et toutes les dérivées en $t = 0$ sont nulles, alors :

$$\mathcal{L} \left[\frac{df}{dt}(t) \right] = p F(p) \quad \mathcal{L} \left[\frac{d^2 f}{dt^2}(t) \right] = p^2 F(p) \quad \dots \quad \mathcal{L} \left[\frac{d^n f}{dt^n}(t) \right] = p^n F(p)$$

Ce que l'on peut résumer par « dériver dans le domaine temporel revient à multiplier par p dans le domaine de Laplace ».

Propriété 3.2 (Intégration d'une fonction causale)

La transformée de Laplace d'une fonction g définie par :

$$g(t) = \int_0^t f(\tau) d\tau$$

et de valeur initiale nulle est égale à :

$$\mathcal{L} \left[\int_0^t f(\tau) d\tau \right] = \frac{F(p)}{p}$$

Démonstration.

$$\begin{aligned} \frac{dg}{dt}(t) = f(t) &\Leftrightarrow \mathcal{L} \left[\frac{dg}{dt}(t) \right] = \mathcal{L} [f(t)] \\ p G(p) - \underbrace{g(0)}_{=0} &= F(p) \Leftrightarrow \mathcal{L} [g(t)] = \frac{F(p)}{p} \quad \square \end{aligned}$$

Théorème 3.5 (Retard)

Soient f et g deux fonctions temporelles d'allure identique mais décalées dans le temps d'une valeur τ . Si g est définie par :

$$\begin{cases} g(t) = 0 & \text{pour } 0 \leq t < \tau \\ g(t) = f(t - \tau) & \text{pour } t \geq \tau \end{cases}$$

alors

$$\mathcal{L} [f(t - \tau)] = e^{-p\tau} \mathcal{L} [f(t)] = e^{-p\tau} F(p)$$

Démonstration.

$$\mathcal{L} [g(t)] = \mathcal{L} [f(t - \tau)] = \int_0^{+\infty} e^{-pt} f(t - \tau) dt$$

En procédant au changement de variable $t = \xi + \tau$, avec $\xi \in [-\tau, +\infty[$ et tel que $dt = d\xi$ car τ est une constante, il vient :

$$\begin{aligned} \mathcal{L} [g(t)] &= \mathcal{L} [f(\xi)] = \int_{-\tau}^{+\infty} e^{-p(\xi+\tau)} f(\xi) d\xi \\ &= e^{-p\tau} \int_{-\tau}^{+\infty} e^{-p\xi} f(\xi) d\xi \\ &= e^{-p\tau} \left(\underbrace{\int_{-\tau}^0 e^{-p\xi} f(\xi) d\xi}_{=0} + \underbrace{\int_0^{+\infty} e^{-p\xi} f(\xi) d\xi}_{F(p)} \right) \end{aligned}$$

car f est une fonction causale, nulle sur l'intervalle $[-\tau, 0[$. □

Théorème 3.6 (Modulation)

Soient $f(t)$ une fonction causale et e^{-at} une fonction de modulation avec $a \in \mathbb{C}$. La transformée de Laplace de la fonction modulée $f(t) e^{-at}$ est définie par :

$$\mathcal{L} [f(t) e^{-at}] = F(p + a)$$

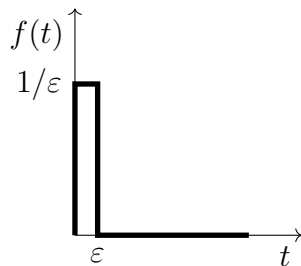
Démonstration. Par définition de la transformée de Laplace, on a directement :

$$\mathcal{L} [f(t) e^{-at}] = \int_0^{+\infty} e^{-pt} f(t) e^{-at} dt = \underbrace{\int_0^{+\infty} e^{-(p+a)t} f(t) dt}_{F(p+a)} \quad \square$$

Dans le cas où $a \in \mathbb{R}^+$, le théorème de modulation est dit d'amortissement. Il permet de déterminer la transformée de Laplace de signaux amortis.

3.3 Transformée de Laplace de fonctions usuelles

3.3.1 Impulsion



L'impulsion unité ou distribution de Dirac est définie par la limite de la fonction

$$f(t) = \begin{cases} \frac{1}{\varepsilon} & \text{si } t \in [0, \varepsilon[\\ 0 & \text{sinon} \end{cases}.$$

quand ε tend vers 0

FIGURE 21 – Impulsion ou Dirac.

$$\delta(t) = \lim_{\varepsilon \rightarrow 0} f(t)$$

Théorème 3.7

La transformée de Laplace de l'impulsion unité est égale à

$$\mathcal{L}[\delta(t)] = 1$$

Démonstration. Commençons par écrire la transformée de Laplace :

$$\mathcal{L}[f(t)] = \int_0^{+\infty} e^{-pt} f(t) dt$$

et restreindre le domaine d'intégration à celui sur lequel la fonction f est non nulle

$$\mathcal{L}[f(t)] = \int_0^\varepsilon e^{-pt} \frac{1}{\varepsilon} dt = \frac{1}{\varepsilon} \left[\frac{-e^{-pt}}{p} \right]_0^\varepsilon = \frac{1 - e^{-p\varepsilon}}{p\varepsilon}$$

puis exploiter la définition de l'impulsion par passage à la limite

$$\mathcal{L}[\delta(t)] = \lim_{\varepsilon \rightarrow 0} \mathcal{L}[f(t)] = \lim_{\varepsilon \rightarrow 0} \left(\frac{1 - e^{-p\varepsilon}}{p\varepsilon} \right)$$

Se rappelant que la fonction exponentielle est définie par la décomposition :

$$e^x = \sum_{i=0}^{\infty} \frac{x^i}{i!} = 1 + \frac{x}{1!} + \frac{x^2}{2!} + \dots$$

il vient pour la fonction étudiée au voisinage de 0

$$\lim_{\varepsilon \rightarrow 0} \left(\frac{1 - e^{-p\varepsilon}}{p\varepsilon} \right) = \lim_{\varepsilon \rightarrow 0} \frac{1}{p\varepsilon} \left(1 - 1 + p\varepsilon + \frac{(p\varepsilon)^2}{2} + \dots \right) = 1$$

□

3.3.2 Échelon

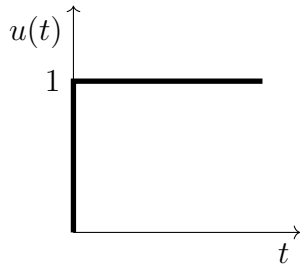


FIGURE 22 – Échelon.

L'échelon unitaire ou fonction de Heaviside est définie par la fonction

$$u(t) = \begin{cases} 1 & \text{si } t \geq 0 \\ 0 & \text{sinon} \end{cases}$$

généralement notée u .

Théorème 3.8

La transformée de Laplace de l'échelon unitaire est égale à

$$\mathcal{L}[u(t)] = \frac{1}{p}$$

Démonstration. Il suffit d'écrire la transformée de Laplace en se rappelant que la fonction u vaut 1 pour tout $t \geq 0$

$$\mathcal{L}[u(t)] = \int_0^{+\infty} e^{-pt} \underbrace{u(t)}_1 dt = \left[\frac{-e^{-pt}}{p} \right]_0^{\infty} = \frac{1}{p}$$

□

Exemple 3.1 (Commande de vérin hydraulique)

Reprenons l'exemple de la commande de vérin hydraulique. Nous avons trouvé l'équation régissant le déplacement de la tige du vérin :

$$\dot{y}(t) = \frac{K}{S} x(t)$$

La transformée de Laplace de cette équation donne :

$$p Y(p) = \frac{K}{S} X(p)$$

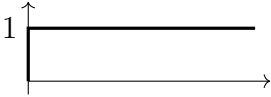
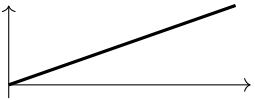
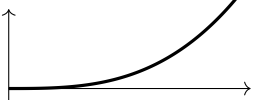
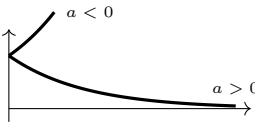
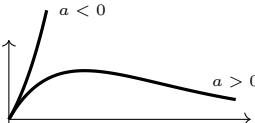
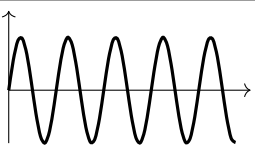
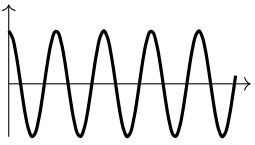
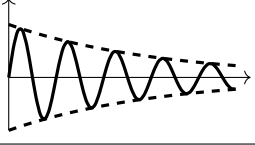
FONCTION	ALLURE	$f(t)$	$F(p) \equiv \mathcal{L}[f(t)]$
Impulsion		$\delta(t)$	1
Échelon		$u(t)$	$\frac{1}{p}$
Rampe		$a t u(t)$	$\frac{a}{p^2}$
Puissance		$t^n u(t)$	$\frac{n!}{p^{n+1}}$
Exponentielle		$e^{-at} u(t)$	$\frac{1}{p + a}$
		$t e^{-at} u(t)$	$\frac{1}{(p + a)^2}$
Sinus		$\sin(\omega t) u(t)$	$\frac{\omega}{p^2 + \omega^2}$
Cosinus		$\cos(\omega t) u(t)$	$\frac{p}{p^2 + \omega^2}$
Sinus amorti		$\sin(\omega t) e^{-at} u(t)$	$\frac{\omega}{(p + a)^2 + \omega^2}$
Cosinus amorti		$\cos(\omega t) e^{-at} u(t)$	$\frac{p + a}{(p + a)^2 + \omega^2}$

TABLE 1 – Transformées de Laplace des fonctions élémentaires.

3.4 Fonction de transfert

On détermine la fonction de transfert d'un système à partir des équations différentielles régissant son comportement. La notion de fonction de transfert est étroitement liée au produit de convolution.

Définition 3.2 (Produit de convolution)

On définit la convolution de deux fonctions f et g , notée $f * g$, par l'opération suivante :

$$(f * g)(t) = \int_0^t f(\tau) g(t - \tau) d\tau$$

Le produit de convolution joue un rôle important dans l'analyse du comportement dynamique des systèmes. Si $h(t)$ est une fonction caractéristique d'un système soumis à une entrée $e(t)$ alors la relation exprime que l'on accumule dans le temps t tout ce que le système possède comme information pendant le laps de temps $t - \tau$. La réponse du système est alors donnée par le produit de convolution $(e * h)(t)$ (figure 23).

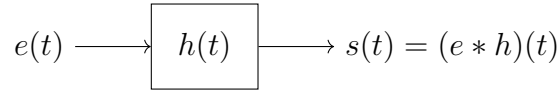


FIGURE 23 – Réponse d'un système définie par convolution.

Théorème 3.9 (Convolution)

La transformée de Laplace du produit de convolution de deux fonctions f et g (dont les transformées existent) est définie par :

$$\mathcal{L}[(f * g)(t)] = F(p) G(p)$$

Démonstration. Par définition de la transformée de Laplace, on a :

$$\mathcal{L}[(f * g)(t)] = \int_0^{+\infty} e^{-pt} \left(\int_0^t f(\tau) g(t - \tau) d\tau \right) dt$$

L'inversion de l'ordre d'intégration donne :

$$\mathcal{L}[(f * g)(t)] = \int_0^{+\infty} \left(\int_\tau^{+\infty} e^{-pt} f(\tau) g(t - \tau) dt \right) d\tau$$

En procédant au changement de variable $t = \xi + \tau$, avec $\xi \in [0, +\infty[$, il vient :

$$\mathcal{L}[(f * g)(t)] = \int_0^{+\infty} f(\tau) \left(\int_0^{+\infty} e^{-p(\xi+\tau)} f(\tau) g(\xi) d\xi \right) d\tau$$

que l'on peut évidemment dissocier sous la forme :

$$\mathcal{L}[(f * g)(t)] = \underbrace{\int_0^{+\infty} e^{-p\tau} f(\tau) d\tau}_{F(p)} \times \underbrace{\int_0^{+\infty} e^{-p\xi} g(\xi) d\xi}_{G(p)}$$

□

En exploitant la propriété de la transformée de Laplace du produit de convolution, il vient immédiatement que si un système de caractéristique $h(t)$ est soumis à une entrée $e(t)$, alors la transformée de Laplace de sa réponse $s(t) = (e * h)(t)$ sera :

$$S(p) = E(p)H(p)$$

La fonction $H(p)$ est appelée la fonction de transfert du système.

Définition 3.3 (Fonction de transfert)

On appelle fonction de transfert ou transmittance le rapport entre la transformée de Laplace du signal d'entrée $e(t)$ et celle du signal de sortie $s(t)$:

$$H(p) = \frac{S(p)}{E(p)}$$

Une fonction de transfert est une fraction rationnelle de la variable complexe p qui peut généralement se mettre sous la forme :

$$H(p) = \frac{a_0 + a_1 p + a_2 p^2 + \cdots + a_m p^m}{b_0 + b_1 p + b_2 p^2 + \cdots + b_n p^n} = \frac{N(p)}{D(p)}$$

On appelle **pôles** de $H(p)$ les racines qui annulent le dénominateur $D(p)$ et **zéros** les racines du numérateur $N(p)$.

Définition 3.4 (Forme canonique)

On appelle forme canonique d'une fonction de transfert la représentation qui consiste à imposer une constante unitaire aux polynômes

$$H(p) = \frac{K}{p^\alpha} \frac{1 + a_{1'} p + a_{2'} p^2 + \cdots + a_{m'} p^{m'}}{1 + b_{1'} p + b_{2'} p^2 + \cdots + b_{r'} p^{r'}}$$

de façon à faire apparaître le gain K et la classe α du système.

À chaque fonction de transfert on peut associer un schéma fonctionnel dans le domaine de Laplace appelé schéma-blocs (figure 24).

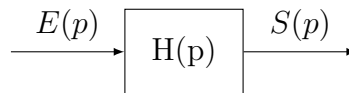


FIGURE 24 – Schéma-blocs d'un système.

D'après cette représentation, il apparaît clairement que la fonction de transfert $H(p)$ représente le comportement du système indépendamment du signal d'entrée. Dès lors, comme la transformée de Laplace d'une impulsion vaut 1, on peut interpréter la fonction de transfert comme la réponse du système à une impulsion de Dirac $e(t) = \delta(t)$.

3.5 Propriétés asymptotiques

Dans de nombreuses situations pratiques, il peut être nécessaire de déterminer les valeurs asymptotiques d'une réponse temporelle à partir de la seule connaissance de sa transformée de Laplace. C'est ce que permettent les deux théorèmes, dits des valeurs initiales et finales.

Théorème 3.10 (Valeur initiale)

Si une fonction f est causale et que la transformée de Laplace de sa dérivée existe alors sa valeur initiale peut être obtenue à partir de sa transformée de Laplace selon :

$$\lim_{t \rightarrow 0} f(t) = \lim_{p \rightarrow \infty} p F(p)$$

Démonstration. En partant de la définition de la transformée de Laplace de la dérivée de la fonction f

$$\mathcal{L} [\dot{f}(t)] = \int_0^{+\infty} e^{-pt} \dot{f}(t) dt = pF(p) - f(0)$$

et en passant à la limite quand p tend vers l'infini, il vient :

$$\lim_{p \rightarrow \infty} \int_0^{+\infty} e^{-pt} \dot{f}(t) dt = \lim_{p \rightarrow \infty} pF(p) - f(0)$$

Comme cette intégrale converge uniformément en p , on peut passer à la limite sous l'intégrale pour obtenir :

$$\int_0^{+\infty} \underbrace{\left(\lim_{p \rightarrow \infty} e^{-pt} \right)}_{\rightarrow 0} \dot{f}(t) dt = \lim_{p \rightarrow \infty} pF(p) - f(0)$$

conduisant finalement à :

$$f(0) = \lim_{t \rightarrow 0} f(t) = \lim_{p \rightarrow \infty} pF(p)$$

□

Théorème 3.11 (Valeur finale)

Si une fonction f admet une limite finie pour $t \rightarrow +\infty$ et que la transformée de Laplace de sa dérivée existe alors sa valeur finale peut être obtenue à partir de sa transformée de Laplace selon :

$$\lim_{t \rightarrow \infty} f(t) = \lim_{p \rightarrow 0} p F(p)$$

Démonstration. L'existence d'une limite $f(\infty) = \lim_{t \rightarrow \infty} f(t)$ pour $f(t)$ implique que cette fonction est bornée et que l'intégrale

$$\int_0^{+\infty} \dot{f}(t) dt$$

existe puisqu'elle est égale à $f(\infty) - f(0)$. Par ailleurs, on a toujours

$$\mathcal{L} [\dot{f}(t)] = \int_0^{+\infty} e^{-pt} \dot{f}(t) dt = pF(p) - f(0)$$

Là encore, comme cette intégrale converge uniformément en p , on peut passer à la limite sous l'intégrale pour obtenir :

$$\lim_{p \rightarrow 0} \int_0^{+\infty} e^{-pt} \dot{f}(t) dt = \int_0^{+\infty} \underbrace{\left(\lim_{p \rightarrow 0} e^{-pt} \right)}_{\rightarrow 1} \dot{f}(t) dt = \int_0^{+\infty} \dot{f}(t) dt = f(\infty) - f(0)$$

d'où finalement :

$$\lim_{p \rightarrow 0} (p F(p) - f(0)) = f(\infty) - f(0)$$

□

Remarque 3.3 (Conditions d'utilisation)

Deux remarques importantes sur les conditions d'application de ces deux théorèmes :

- le théorème de la valeur initiale ne s'applique que si la fonction $f(t)$ est causale, c'est-à-dire que si le degré du numérateur de $F(p)$ est inférieur au degré de son dénominateur ;
- le théorème de la valeur finale ne s'applique que si le système est stable, ce qui se traduit par le fait que les pôles de $pF(p)$ doivent être à parties réelles négatives, sinon $f(t)$ n'a pas de limite finie à temps infini.

3.6 Transformée de Laplace inverse

Comme son nom l'indique, la transformée de Laplace inverse permet de revenir dans le domaine temporel. Son application est indispensable pour déterminer l'expression de la réponse temporelle $s(t)$ d'un système dont on connaît la transformée de Laplace $S(p)$ sous la forme d'une fraction rationnelle :

$$S(p) = \frac{a_0 + a_1 p + a_2 p^2 + \dots + a_m p^m}{b_0 + b_1 p + b_2 p^2 + \dots + b_n p^n} = \frac{N(p)}{D(p)}$$

La méthode la plus commode pour déterminer cette expression temporelle consiste à faire une **décomposition en éléments simples** de la fraction rationnelle de façon à obtenir une somme algébrique de fractions simples dont on connaît la transformée de Laplace inverse et à revenir dans le domaine temporel.

L'étape de décomposition en éléments simples est sans conteste la plus technique. Une fois réalisée, la transformation inverse de chacun des éléments se fait par identification avec les fonctions élémentaires données dans la table 1. Enfin, l'expression de la réponse temporelle recherchée est simplement obtenue par assemblage des différentes contributions.

Exemple 3.2 (Commande de vérin hydraulique)

Reprenons l'exemple de la commande de vérin hydraulique pour lequel nous avons l'expression du déplacement de la tige du vérin dans le domaine de Laplace :

$$Y(p) = \frac{K}{S p} X(p)$$

Si on impose une entrée (commande du distributeur) sous forme d'échelon unitaire $x(t) = u(t)$ tel que $X(p) = 1/p$, la fonction traduisant le déplacement de la tige sera alors :

$$Y(p) = \frac{K}{S p^2}$$

En regardant dans la table 1 des transformées de Laplace, on remarque que cette forme est déjà simple et correspond à la fonction rampe avec $a = K/S$; d'où :

$$\forall t \geq 0, \quad y(t) \stackrel{\text{déf.}}{=} \mathcal{L}^{-1}[Y(p)] = \frac{K}{S} t$$

3.6.1 Méthode de décomposition en éléments simples

Étape 0 : Causalité ?

On se limitera volontairement aux fractions rationnelles $\frac{N(p)}{D(p)}$ vérifiant le principe de causalité des systèmes physiques, c'est-à-dire telles que :

$$\deg N(p) \leq \deg D(p)$$

Si ce n'est pas le cas, on effectuera la division de $N(p)$ par $D(p)$ tel que le reste vérifie alors la condition. On notera que dans le cas où $\deg N(p) = \deg D(p)$ la décomposition en éléments simples admettra une constante.

Étape 1 : Factorisation du dénominateur

Si ce n'est déjà le cas, on mettra le dénominateur sous une forme factorisée

$$D(p) = \prod_{i=1}^{\deg D(p)} (p - p_i)$$

en recherchant ses pôles (racines) réels ou complexes conjugués – qu'ils soient simples ou multiples (puissance entière naturelle) – de façon à avoir :

$$D(p) = \underbrace{(p - a_1)}_{\text{réel simple}} \dots \underbrace{(p - a_2)^K}_{\text{réel multiple}} \dots \underbrace{(p - z_1)(p - \bar{z}_1)}_{\text{complexe conjugué simple}} \dots \underbrace{(p - z_2)^K(p - \bar{z}_2)^K}_{\text{complexe conjugué multiple}}$$

Étape 2 : Forme de la décomposition

Pour tous les pôles, on recherche la forme de la décomposée associée.

► Pôle réel

Pour chaque pôle réel a , de multiplicité $K \geq 1$, on cherchera une décomposition sous la forme :

$$\sum_{\beta=1}^K \frac{A_\beta}{(p - a)^\beta} = \frac{A_1}{p - a} + \frac{A_2}{(p - a)^2} + \dots + \frac{A_K}{(p - a)^K}$$

nécessitant la détermination des K coefficients réels $A_1, A_2, \dots, A_K$.

► **Couple de pôles complexes conjugués**

Pour chaque couple de pôles complexes conjugués $z = \alpha + j\beta$ et $\bar{z} = \alpha - j\beta$, de multiplicité $K \geq 1$, on exploitera la forme polynômiale

$$(p - z)(p - \bar{z}) = p^2 + b p + c$$

avec $b = -2\alpha$ et $c = \alpha^2 + \beta^2$ et on cherchera une décomposition sous la forme :

$$\sum_{\beta=1}^K \frac{B_{\beta} p + C_{\beta}}{(p^2 + b p + c)^{\beta}} = \frac{B_1 p + C_1}{(p^2 + b p + c)} + \frac{B_2 p + C_2}{(p^2 + b p + c)^2} + \dots + \frac{B_K p + C_K}{(p^2 + b p + c)^K}$$

nécessitant la détermination des $2K$ coefficients réels $B_1, C_1, B_2, C_2, \dots, B_K$ et C_K .

Toutes les contributions sont ensuite sommées – en prenant soin de bien distinguer les coefficients à déterminer – pour recomposer la fraction rationnelle de départ.

Étape 3 : Détermination des coefficients

Pour déterminer tous les coefficients de la fraction rationnelle

$$\frac{N(p)}{D(p)} = \sum_{\beta=1}^K \frac{A_{\beta}}{(p - a)^{\beta}} + \sum_{\beta=1}^K \frac{B_{\beta} p + C_{\beta}}{(p^2 + b p + c)^{\beta}} + \dots$$

il est nécessaire de l'identifier à la fraction d'origine en passant par des calculs de limites ou en utilisant des valeurs particulières. Seules les deux règles suivantes sont toujours valables :

- le coefficient de degré le plus élevé A_K associé à un pôle réel a , de multiplicité $K \geq 1$, est égal à :

$$A_K = \lim_{p \rightarrow a} (p - a)^K \frac{N(p)}{D(p)}$$

- les coefficients de degré le plus élevés B_K et C_K associés à un couple de pôles complexes conjugués $z = \alpha + j\beta$ et $\bar{z} = \alpha - j\beta$, de multiplicité $K \geq 1$, sont respectivement égaux à :

$$\begin{cases} B_K = \frac{1}{\beta} \Im \left(\lim_{p \rightarrow \alpha + j\beta} (p^2 + b p + c)^K \frac{N(p)}{D(p)} \right) \\ C_K = \Re \left(\lim_{p \rightarrow \alpha + j\beta} (p^2 + b p + c)^K \frac{N(p)}{D(p)} \right) - \alpha B_K \end{cases}$$

où $\Re$ et $\Im$ sont respectivement les fonctions partie réelle et partie imaginaire.

Exemple 3.3 (Résolution d'une équation différentielle)

Résoudre l'équation différentielle suivante en utilisant la transformée de Laplace

$$y''(t) + 6y'(t) + 5y(t) = 3 \quad \text{avec } y(0) = 1 \text{ et } y'(0) = 2$$

Commençons par écrire les transformées de Laplace :

$$\begin{aligned}\mathcal{L}[y(t)] &= Y(p) \\ \mathcal{L}[y'(t)] &= p Y(p) - 1 \\ \mathcal{L}[y''(t)] &= p^2 Y(p) - p - 2 \\ \mathcal{L}[3u(t)] &= \frac{3}{p}\end{aligned}$$

où nous avons pris soin de convertir le second membre en fonction causale avec un échelon implicite. L'équation différentielle s'écrit dans le domaine de Laplace :

$$(p^2 + 6p + 5) Y(p) = \frac{3}{p} + p + 8$$

que l'on peut factoriser sous la forme :

$$Y(p) = \frac{p^2 + 8p + 3}{p(p^2 + 6p + 5)}$$

On cherche maintenant à décomposer en éléments simples cette expression pour déterminer celle de $y(t)$. Pour identifier tous les pôles, calculons le discriminant associé à l'équation $p^2 + 6p + 5 = 0$,

$$\Delta = 6^2 - 4 \times 5 = 16 > 0$$

admettant deux racines réelles :

$$p_{1/2} = \frac{-6 \pm \sqrt{16}}{2} \Rightarrow p_1 = -1 \text{ et } p_2 = -5$$

telles que le dénominateur s'écrive sous forme factorisée :

$$D(p) = p(p+1)(p+5)$$

et que la fonction $Y(p)$ admette la forme de décomposition :

$$Y(p) = \frac{A}{p} + \frac{B}{p+1} + \frac{C}{p+5}$$

On exploitant la première des deux règles données, on peut déterminer tous les coefficients selon :

$$\begin{aligned}A &= \lim_{p \rightarrow 0} p \left(\frac{p^2 + 8p + 3}{p(p+1)(p+5)} \right) = \frac{3}{5} \\ B &= \lim_{p \rightarrow -1} (p+1) \left(\frac{p^2 + 8p + 3}{p(p+1)(p+5)} \right) = \frac{1 - 8 + 3}{-1 \times 4} = 1 \\ C &= \lim_{p \rightarrow -5} (p+5) \left(\frac{p^2 + 8p + 3}{p(p+1)(p+5)} \right) = \frac{25 - 40 + 3}{-5 \times -4} = \frac{-3}{5}\end{aligned}$$

d'où :

$$Y(p) = \frac{3}{5p} + \frac{1}{p+1} - \frac{3}{5(p+5)}$$

Par identification dans le tableau des transformées, il vient l'expression temporelle :

$$y(t) = \frac{3}{5} + e^{-t} - \frac{3}{5} e^{-5t}$$

Exemple 3.4 (Résolution d'une équation différentielle)

Résoudre l'équation différentielle suivante en utilisant la transformée de Laplace

$$y''(t) + y'(t) + y(t) = \sin(t) \quad \text{avec } y(0) = 1 \text{ et } y'(0) = -1$$

Commençons par écrire les transformées de Laplace :

$$\begin{aligned}\mathcal{L}[y(t)] &= Y(p) \\ \mathcal{L}[y'(t)] &= p Y(p) - 1 \\ \mathcal{L}[y''(t)] &= p^2 Y(p) - p + 1 \\ \mathcal{L}[\sin(t)] &= \frac{1}{p^2 + 1}\end{aligned}$$

L'équation différentielle s'écrit dans le domaine de Laplace :

$$(p^2 + p + 1) Y(p) - p = \frac{1}{p^2 + 1}$$

que l'on peut factoriser sous la forme :

$$Y(p) = \frac{p^3 + p + 1}{(p^2 + p + 1)(p^2 + 1)}$$

Les deux polynômes du dénominateur ont des racines complexes conjuguées. La décomposition de $Y(p)$ prend alors la forme :

$$Y(p) = \frac{Ap + B}{p^2 + p + 1} + \frac{Cp + D}{p^2 + 1}$$

Par identification, il vient immédiatement le système d'équations :

$$\begin{cases} 1 &= A + C \\ 0 &= B + C + D \\ 1 &= A + C + D \\ 1 &= B + D \end{cases} \quad \Rightarrow \quad \begin{cases} A &= 2 \\ B &= 1 \\ C &= -1 \\ D &= 0 \end{cases}$$

d'où

$$Y(p) = \frac{2p + 1}{p^2 + p + 1} - \frac{p}{p^2 + 1}$$

L'identification du second terme est immédiate :

$$\mathcal{L}^{-1} \left[\frac{p}{p^2 + 1} \right] = \cos(t)$$

Pour le premier terme, bien qu'il soit déjà sous forme simple, il n'est pas trivial d'identifier directement la fonction. Par déduction, on est amené à rechercher pour le dénominateur une expression de la forme $(p + a)^2 + \omega^2$ qui, avec $a = 1/2$ et $\omega = \sqrt{3}/2$, conduit à :

$$\frac{2p + 1}{p^2 + p + 1} = \frac{2 \left(p + \frac{1}{2} \right)}{\left(p + \frac{1}{2} \right)^2 + \left(\frac{\sqrt{3}}{2} \right)^2} = \mathcal{L} \left[2 \cos \left(\frac{\sqrt{3}}{2} t \right) e^{-\frac{t}{2}} \right]$$

Finalement, on obtient l'expression temporelle :

$$y(t) = 2 \cos \left(\frac{\sqrt{3}}{2} t \right) e^{-\frac{t}{2}} - \cos(t)$$

Exercice 3.1

Décomposer en éléments simples la fraction rationnelle suivante :

$$H(p) = \frac{(p + 3)^2}{(p^2 + 3p + 2)(p - 1)^2(2p^2 - 3p + 2)^2}$$

4 Représentation par schéma-blocs

4.1 Schéma-blocs

La représentation par schéma-blocs permet de représenter de manière graphique un système linéaire. Chaque bloc du schéma caractérise une des fonctions du système (un des constituants) et on associe à chaque bloc la fonction de transfert du constituant qu'il représente. Les arcs qui relient les blocs portent les informations des signaux d'entrée et de sortie et l'allure globale du schéma nous renseigne sur la structure du système (boucle ouverte, boucle fermée).

Dans ce formalisme de schéma-blocs, on trouve (figure 25) :

- le bloc qui représente un élément ou un groupe d'éléments du système ;
- la flèche qui représente une grandeur physique entrant ou sortant d'un élément ;
- le comparateur qui soustrait des grandeurs de même nature physique ;
- le sommateur qui ajoute des grandeurs de même nature physique ;
- le branchement qui représente un prélèvement d'information analogue à un capteur parfait.

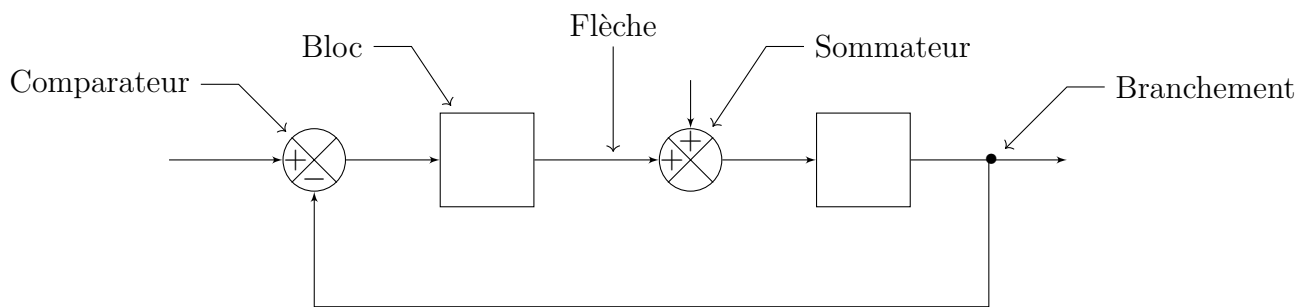


FIGURE 25 – Structure et composition d'un schéma-blocs.

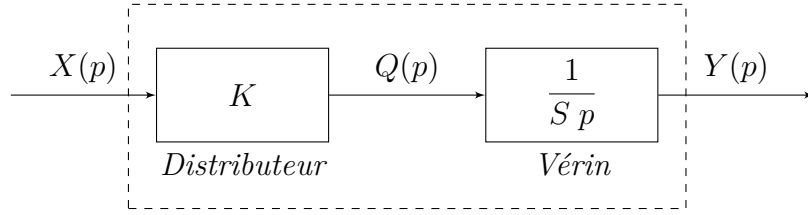
Un système complexe se compose en réalité d'une combinaison d'un certain nombre de sous-systèmes plus simples. Un schéma-blocs global fait donc apparaître un ensemble de systèmes élémentaires auxquels on peut associer individuellement une transmittance à partir des équations différentielles régissant leur comportement. Compte tenu des liaisons entre les différents blocs, il est souvent nécessaire de calculer la transmittance globale d'un système entre deux signaux.

Exemple 4.1 (Commande de vérin hydraulique)

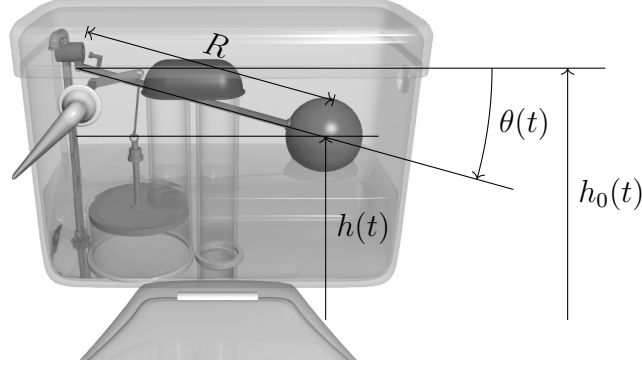
On peut écrire la transmittance du système global {distributeur + vérin}

$$\begin{array}{ccc} \xrightarrow{X(p)} & \boxed{H(p)} & \xrightarrow{Y(p)} \end{array} \qquad H(p) = \frac{Y(p)}{X(p)} = \frac{K}{S p}$$

ou décomposer le système structurellement :



Exemple 4.2 (Commande de chasse d'eau)



On considère une chasse d'eau reliée au réseau d'eau courante de débit maximal constant q_v . Quand le clapet est fermé, le débit de fuite $q_f(t)$ est faible. Au fur et à mesure que le niveau d'eau monte, le débit d'arrivée d'eau $q_c(t)$ diminue car le flotteur ferme la vanne progressivement. On supposera cette variation comme linéaire. À chaque instant, la position du flotteur varie en fonction de la hauteur d'eau $h(t)$ ce qui fait tourner le levier associé à la vanne. Le levier a une longueur utile de R et on note $\theta(t) \in [-\pi/4; 0]$ l'angle du levier par rapport à l'horizontale. Cet angle est associé à la différence de hauteur d'eau $\Delta h(t) = h_0(t) - h(t)$ où $h_0(t)$ est la hauteur de consigne. Enfin, on note S la section du réservoir supposée constante sur toute la hauteur.

Modélisation du système

Sur le domaine de définition de θ , un comportement linéaire de l'ouverture de la vanne se traduit par :

$$q_c(t) = q_v \frac{\theta(t)}{-\pi/4} \quad \xrightarrow{\mathcal{L}} \quad Q_c(p) = \frac{-4q_v}{\pi} \Theta(p)$$

Le débit utile pour remplir le réservoir est la différence entre le débit venant de la vanne $q_c(t)$ et le débit de fuite du réservoir $q_f(t)$. Ce débit utile est associé à une variation de hauteur d'eau dans le réservoir :

$$S\dot{h}(t) = q_c(t) - q_f(t) \quad \xrightarrow{\mathcal{L}} \quad H(p) = \frac{1}{Sp} (Q_c(p) - Q_f(p))$$

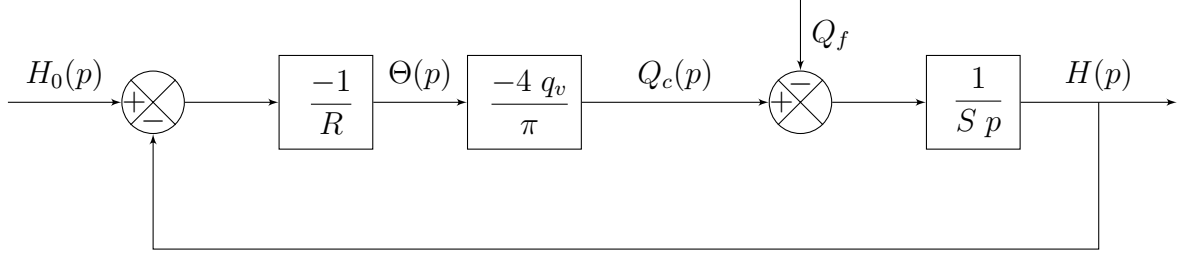
Enfin, l'angle du levier par rapport à l'horizontale est lié à la hauteur d'eau dans le réservoir par la relation :

$$\Delta h(t) = h_0(t) - h(t) = R \sin(-\theta(t))$$

Cette relation est non linéaire. Pour se retrouver dans le cadre des SLCI, on va linéariser cette expression autour du point de fonctionnement $\theta \approx 0$ qui permet de faire l'approximation linéaire :

$$\Delta h(t) = h_0(t) - h(t) \approx -R\theta(t) \quad \xrightarrow{\mathcal{L}} \quad \Theta(p) = \frac{-1}{R} (H_0(p) - H(p))$$

On peut synthétiser ces trois relations par le schéma-blocs :



4.2 Chaîne fermée : système asservi

On appelle **asservissement** un système de commande qui possède une amplification et/ou une correction de puissance (dans la chaîne directe) et une boucle de retour munie d'un capteur. Les schéma-blocs des systèmes asservis ont de façon générale la même structure qui contient un comparateur, une chaîne d'action, une jonction et une chaîne de retour (figure 26).

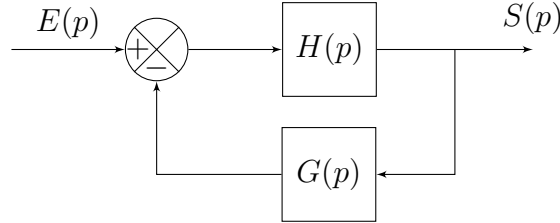


FIGURE 26 – Schéma-blocs minimal d'un système asservi.

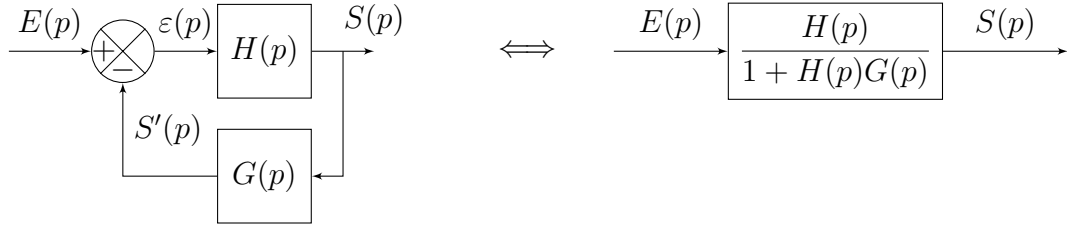
On note $H(p)$ la fonction de transfert de la chaîne directe et $G(p)$ la fonction de transfert de la chaîne de retour. On peut facilement calculer la fonction de transfert en boucle fermée (FTBF) du système reliant l'entrée et la sortie en parcourant les signaux depuis l'entrée. Par lecture directe, il vient :

$$\begin{aligned} [E(p) - G(p) S(p)] H(p) &= S(p) \\ E(p) H(p) &= S(p) [1 + H(p) G(p)] \end{aligned}$$

d'où finalement :

$$\text{FTBF}(p) = \frac{S(p)}{E(p)} = \frac{H(p)}{1 + H(p) G(p)}$$

Ce que l'on peut traduire graphiquement par :



Il est utile dans l'analyse des systèmes d'introduire la notion de fonction de transfert en boucle ouverte (FTBO). Elle est définie comme :

$$\text{FTBO}(p) = \frac{S'(p)}{\varepsilon(p)} = H(p) G(p)$$

On peut alors écrire :

$$\text{FTBF}(p) = \frac{\text{chaîne directe}(p)}{1 + \text{FTBO}(p)}$$

avec chaîne directe(p) = $H(p)$.

Remarque 4.1 (Cas d'un retour unitaire)

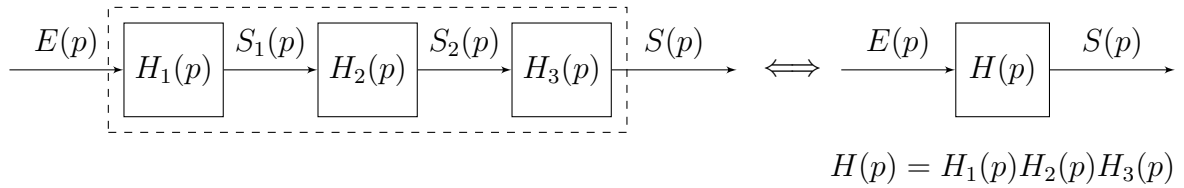
Dans le cas particulier d'un retour unitaire, c'est-à-dire quand $G(p) = 1$, on a chaîne directe(p) = FTBO(p).

4.3 Manipulation de schéma-blocs

4.3.1 Blocs en série

Propriété 4.1 (Association série)

La fonction de transfert équivalente à l'association de blocs en série est égale au produit des fonctions de transfert des blocs.



Démonstration. Par définition de la fonction de transfert de chaque bloc, on a :

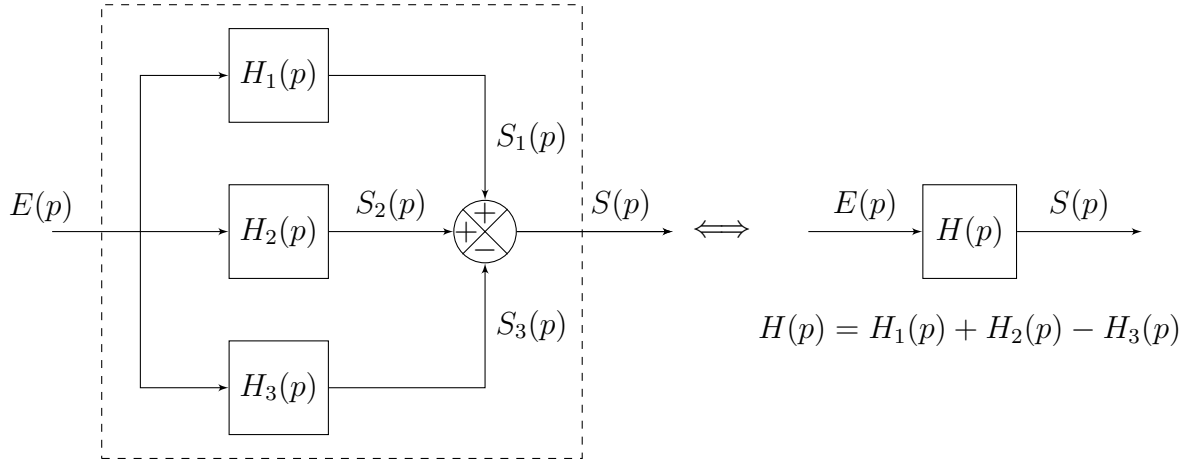
$$\begin{cases} S_1(p) &= H_1(p)E(p) \\ S_2(p) &= H_2(p)S_1(p) \\ S(p) &= H_3(p)S_2(p) \end{cases} \quad \text{tel que :} \quad S(p) = \underbrace{H_3(p)H_2(p)H_1(p)}_{H(p)} E(p)$$

□

4.3.2 Blocs en parallèle

Propriété 4.2 (Association parallèle)

La fonction de transfert équivalente à l'association de blocs en parallèle est égale à la somme des fonctions de transfert des blocs.

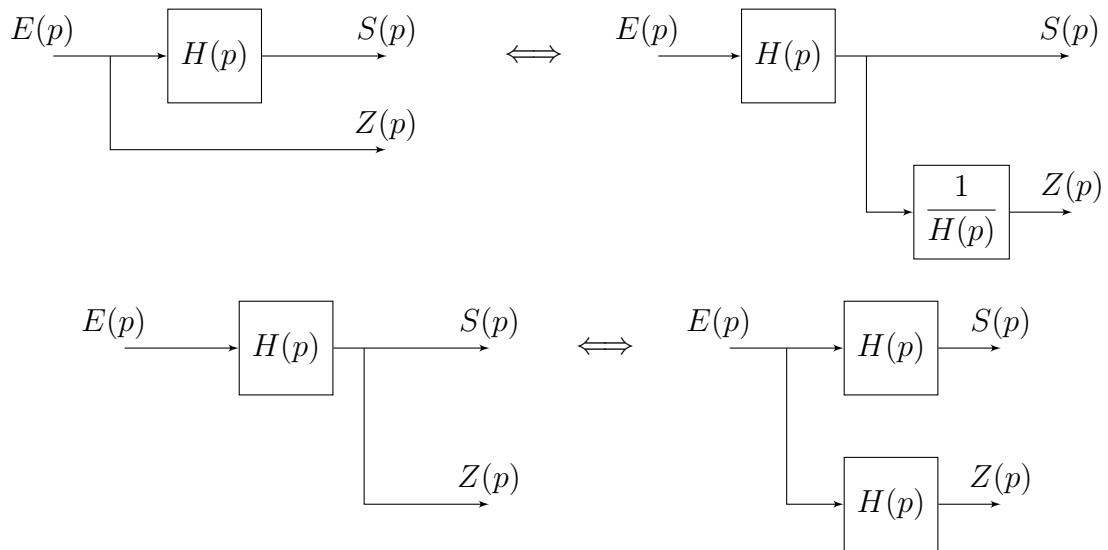


Démonstration. Par définition de la fonction de transfert de chaque bloc, on a :

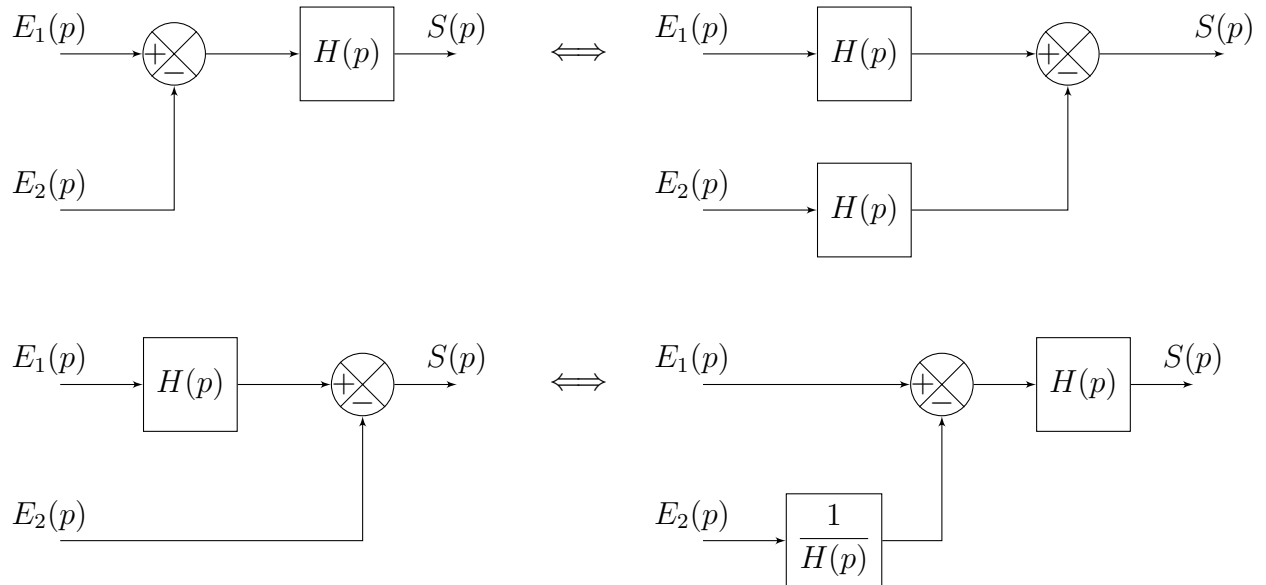
$$\begin{cases} S_1(p) = H_1(p)E(p) \\ S_2(p) = H_2(p)E(p) \\ S_3(p) = -H_3(p)E(p) \end{cases} \quad \text{tel que :} \quad \begin{aligned} S(p) &= S_1(p) + S_2(p) + S_3(p) \\ &= \underbrace{(H_1(p) + H_2(p) - H_3(p))}_{H(p)} E(p) \end{aligned}$$

□

4.3.3 Déplacement d'un point de prélèvement



4.3.4 Déplacement d'un comparateur



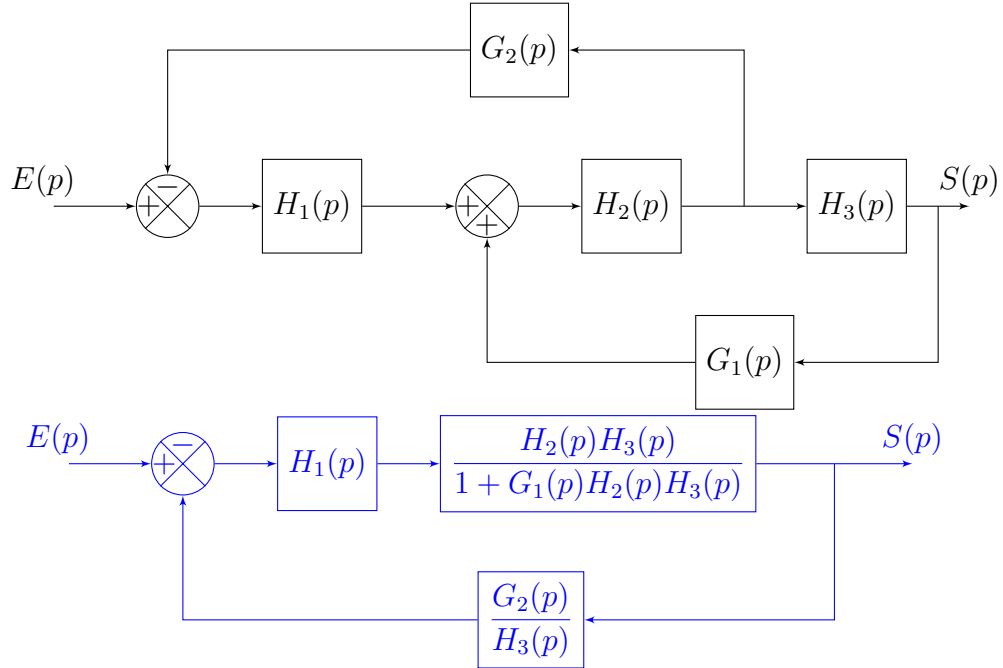
4.3.5 Simplification d'un schéma fonctionnel

Il s'agit de trouver un schéma-blocs équivalent avec une seule boucle ou de trouver la transmittance équivalente. Deux cas peuvent se présenter :

Boucles concentriques – Il faut remplacer la boucle la plus interne par un schéma rectiligne et recommencer l'opération jusqu'à ce qu'il n'y ait plus qu'une boucle.

Boucles imbriquées – Il faut se ramener au cas précédent en déplaçant les points de prélèvement ou les comparateurs qui provoquent les enchevêtrements.

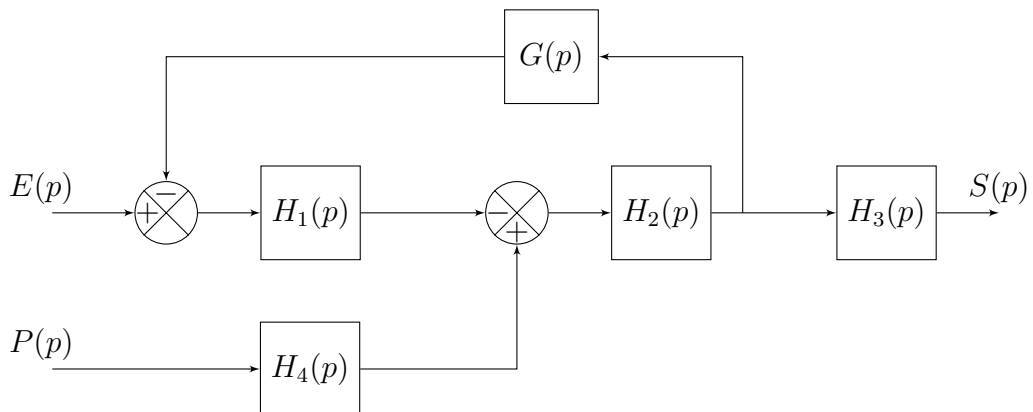
Exemple 4.3 (Simplification de schéma-blocs)

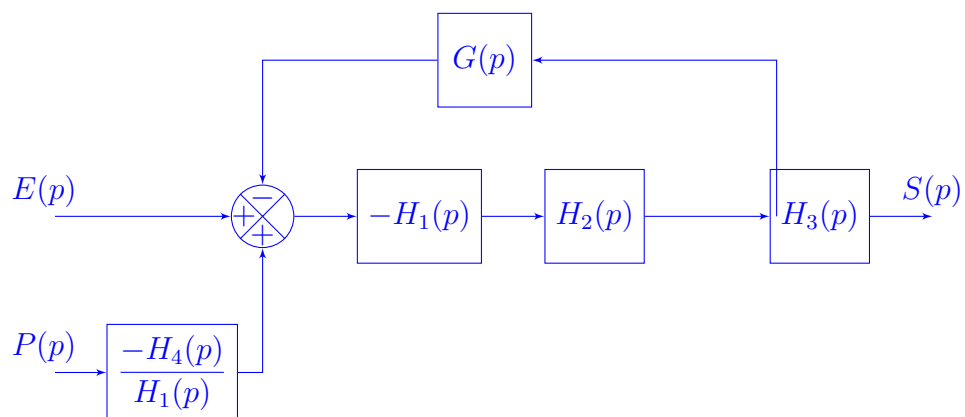


4.3.6 Cas des systèmes multi-variables

Dans le cas de problèmes avec plusieurs entrées, on utilise le théorème de superposition. On annule successivement toutes les entrées sauf une. On calcule les transmittances partielles. La fonction de transfert du système sera la somme des transmittances partielles.

Exemple 4.4 (Simplification de schéma-blocs)





* *

*