



Modélisation de systèmes à paramètres répartis par développements d'impédances dans l'espace de Laplace.

Fabien Soulier

► To cite this version:

Fabien Soulier. Modélisation de systèmes à paramètres répartis par développements d'impédances dans l'espace de Laplace. : Application aux protections numériques et à la localisation de défauts sur les lignes triphasées. Sciences de l'ingénieur [physics]. Univ. Poitiers, 2003. Français. NNT: . tel-01784852

HAL Id: tel-01784852

<https://theses.hal.science/tel-01784852>

Submitted on 3 May 2018

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

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

THÈSE

pour l'obtention du grade de
DOCTEUR DE L'UNIVERSITÉ DE POITIERS

(Faculté des sciences fondamentales et appliquées)
(Diplôme national — Arrêté du 25 avril 2002)

École doctorale : sciences pour l'ingénieur
Secteur de recherche : génie électrique

Présentée par

Fabien SOULIER

Ingénieur ENSIEG, agrégé de physique appliquée

MODÉLISATION DE SYSTÈMES À PARAMÈTRES RÉPARTIS
PAR DÉVELOPPEMENTS D'IMPÉDANCES DANS L'ESPACE DE
LAPLACE

*Application aux protections numériques et à la localisation de
défauts sur les lignes triphasées*

Directeur de Thèse

Patrick LAGONOTTE

Soutenue le 4 décembre 2003

devant la commission d'examen

— JURY —

<i>Rapporteurs</i>	Philippe AURIOL	Professeur à l'École centrale de Lyon
	Youssef TOURÉ	Professeur à l'université d'Orléans
<i>Examineurs</i>	François COSTA	Professeur à l'IUFM de Créteil
	Patrick LAGONOTTE	Maître de conférences à l'université de Poitiers
	Jean-Bernard SAULNIER	Professeur à l'ENSMA
	Jean-Claude TRIGEASSOU	Professeur à l'université de Poitiers
<i>Invités</i>	Khalid ALEM	Ingénieur Alstom
	Volker LEITLOFF	Docteur, ingénieur RTE

À Marion
À mes parents

Remerciements

Les recherches qui sont l'objet de ce mémoire se sont déroulées au sein de l'équipe analyse système et optimisation du laboratoire d'études thermiques de Poitiers et de l'équipe recherche et développement produits de la société *Alstom* de Lattes.

Pour m'avoir accueilli pendant ces trois années, je remercie M. Daniel Petit, directeur du laboratoire d'études thermiques et M. Khalid Alem, directeur de la recherche et du développement de la ligne produits d'*Alstom*.

Je tiens tout particulièrement à remercier mon directeur de thèse, M. Patrick Lagonotte, maître de conférences à l'université de Poitiers, pour m'avoir fait confiance en me proposant ce sujet ambitieux et d'une richesse exceptionnelle, pour son implication exemplaire dans l'encadrement dans mon travail, ses nombreuses idées, ses propositions pertinentes, mais avant tout pour sa sympathie et son amitié.

Je voudrais remercier M. Philippe Auriol, professeur à l'École centrale de Lyon, pour avoir accepté d'être rapporteur de mon travail de recherche. Et je voudrais, pour cette même raison, remercier M. Youssoufi Touré, professeur à l'université d'Orléans, et aussi parce qu'il a su, par la clarté de son exposé à l'école d'automatique de Lille sur le contrôle des systèmes à paramètres répartis, démystifier cette « étrange théorie » du *semigroupe*.

Pour avoir accepté de faire partie de mon jury de thèse, je remercie M. François Costa, professeur à l'IUFM de Créteil, M. Jean-Bernard Saulnier, professeur à l'École nationale supérieure de mécanique et d'aérotechnique, M. Jean-Claude Trigeassou, professeur à l'université de Poitiers, ainsi que M. Khalid Alem, ingénieur Alstom T&D et M. Volker Leitloff, docteur et ingénieur RTE, les représentants du contexte industriel dans lequel ces travaux se sont déroulés.

Merci aussi à M. Claude Brezinski qui a trouvé assez d'intérêt à nos travaux pour m'accueillir pour un séminaire au laboratoire d'analyse et d'optimisation de Lille. Et encore plus particulièrement à sa collaboratrice, Jeannette Van Iseghem, qui n'a pas ménagé ses efforts pour m'expliquer la théorie des polynômes orthogonaux et les subtilités de leurs relations avec les approximants de Padé. Merci aussi à M. Denis Matignon pour l'intérêt qu'il a montré pour mes recherches à l'université Notre-Dame et pour ses réponses, parfois passionnées mais toujours argumentées.

Je remercie M. Éric Legrand, qui est à l'initiative de ce travail et qui s'occupe toujours de lignes aériennes, mais dans un autre domaine, M. Damien Tholomier qui en a assuré un temps l'encadrement industriel et en a toujours suivi la progression, M. Frédéric Bel, responsable de la recherche et du développement produits à Lattes, et M. Thierry Bardou

qui a eu l'éprouvante tâche d'encadrer ce travail.

Et je remercie également toute l'équipe de recherche et développement de la ligne produit d'Alstom Lattes, Nathalie, qui m'a précédé à Grenoble, Karen, Laurence, Patricia, Jutta, les Christopes (tous les trois), Patrick, Fred, François l'artiste, René pour les discussions toujours intéressantes et la relecture, l'encyclopédique M. Carel, Jean-Philippe, spécialiste de *Matlab* et des modèles ARMA, et les beaux-frères Fabrice et Denis qui m'ont initié aux protections numériques.

Pour mon séjour poitevin, je remercie les (ex-)doctorants thermiciens : Jean-François qui a défriché le thème de la présente thèse pendant son DEA, Marc, Béa, le Manu du Poitou, les cantalous Yo et Jack, Seb, Maria, Stéphane pour son Coin, Nicolas et Cédric pour les échanges d'idées constructifs.

Notations et symboles

Symboles et paramètres physiques

c	capacité linéique, pour un modèle à constantes réparties, ou selon le contexte, $c = \frac{1}{\sqrt{\epsilon\mu}}$: vitesse des ondes électromagnétiques dans le milieu,
C	capacité globale d'un modèle : $C = cX$ en coordonnées cartésiennes,
Δ	coefficient d'influence électrostatique,
ϵ	permittivité diélectrique absolue du milieu,
g	conductance linéique, pour un modèle à constantes réparties,
G	conductance globale d'un modèle : $G = gX$ en coordonnées cartésiennes,
i	courant,
l	inductance linéique, pour un modèle à constantes réparties,
L	inductance globale d'un modèle : $L = lX$ en coordonnées cartésiennes,
μ	perméabilité magnétique absolue du milieu,
Φ	flux d'un vecteur, ou flux de chaleur,
q	charge électrique, ou quantité de chaleur,
r	résistance linéique, pour un modèle à constantes réparties,
R	résistance globale d'un modèle : $R = rX$ en coordonnées cartésiennes,
t	variable temporelle,
v	tension,
x	variable d'espace dans la direction de propagation ou de diffusion, en coordonnées cartésiennes,
X	dimension totale du système étudié en coordonnées cartésiennes,
y	admittance transversale linéique, pour un modèle à constantes réparties,
Y	admittance transversale globale d'un modèle : $Y = yX$ en coordonnées cartésiennes,
z	impédance longitudinale linéique, pour un modèle à constantes réparties,
Z	impédance longitudinale globale d'un modèle : $Z = zX$ en coordonnées cartésiennes,
$\mathcal{Z}$	impédance d'ordre infini.

Notations mathématiques

- $\square$ transformation de Laplace,
- ∇ opérateur *nabla* de dérivation vectorielle,
- $\mathbb{C}$ ensemble des nombres complexes,
- Δ Laplacien, opérateur de dérivation d'ordre deux,
- e base des logarithmes népériens,
- i nombre complexe vérifiant $i^2 = -1$,
- o, O notations de Landau pour les développements limités,
- $\mathbf{m}$ représente une matrice,
- $\mathbb{N}$ ensemble des entiers naturels,
- ν fréquence dans la transformée de Fourier,
- p variable complexe de la transformation de Laplace,
- $\mathbb{R}$ ensemble des nombres réels,
- T transposition d'une matrice,
- $\mathbf{v}$ représente un vecteur,
- z variable symbolique de la transformée en Z ,
- $\mathbb{Z}$ ensemble des entiers relatifs.

Table des matières

Introduction	1
1 La protection des lignes de transmission	3
1.1 Protections des réseaux électriques	3
1.1.1 Divers types de réseaux	3
1.1.2 Utilité de la protection	4
1.1.3 Défauts d'isolement	5
1.1.4 Qualités des protections	5
1.1.5 Technologies des protections	6
1.2 La protection de distance	6
1.2.1 Principe de la protection de distance	6
1.2.2 Exemple de protection numérique	7
1.2.3 Modèle de ligne	9
1.2.4 Algorithme de distance	10
1.2.5 Algorithme directionnel	11
1.3 Amélioration des performances	11
2 Modélisation des lignes électriques	15
2.1 L'évolution de la modélisation des lignes électriques	15
2.1.1 Du télégraphe au téléphone	15
2.1.2 L'équation des télégraphistes	21
2.1.3 Significations physiques des paramètres de la ligne	22
2.2 Résolution dans l'espace de Laplace	24
2.2.1 Transformation de Laplace	26
2.2.2 Modèles équivalents	27
2.2.3 Modèles à constantes localisées	29
2.3 Les systèmes à paramètres répartis	33
2.3.1 Équations aux dérivées partielles	33
2.3.2 Systèmes d'ordre infini	38
2.4 Axe de recherche développé dans ce mémoire	40

3	De l'approximation polynomiale à l'approximation rationnelle	43
3.1	L'approximation polynomiale	43
3.1.1	Développement limité	44
3.1.2	Polynômes orthogonaux	45
3.2	Les approximants de Padé	49
3.2.1	Définition des approximants	49
3.2.2	Propriétés des approximants	52
3.2.3	Calcul des approximants	54
3.2.4	Approximants de Padé et polynôme orthogonaux	56
3.3	Les Fractions Continues	58
3.3.1	Le développement d'un nombre	58
3.3.2	Le développement d'une fonction	61
3.3.3	L'équivalence entre les fractions continues et les approximants de Padé	63
3.4	Théorème de Mittag-Leffler	68
3.4.1	Développement suivant les pôles d'une fonction	69
3.4.2	Développement suivant les zéros d'une fonction	70
3.5	Méthodes d'approximation retenues	70
4	Synthèse de réseaux par développements d'impédances	71
4.1	Extension des synthèses de Foster et Cauver à l'ordre infini	71
4.1.1	Développement en éléments simples de l'impédance	72
4.1.2	Développement en éléments simples de l'admittance	74
4.1.3	Développement en fraction continue de l'impédance	74
4.1.4	Comparaison des réponses fréquentielles	77
4.2	Application à une ligne monophasée	77
4.2.1	Synthèse de Cauver	79
4.2.2	Capacités de compensation	79
4.2.3	Modèle développé	80
4.2.4	Simulation d'une réponses temporelle	83
4.3	Application à un câble avec effet de peau	83
4.3.1	Calcul de la répartition des filets de courants	84
4.3.2	Calcul de l'impédance interne	87
4.3.3	Synthèse de Cauver	88
4.3.4	Modélisation d'ordre non entier	89
4.3.5	Comparaison du modèle de Cauver et du modèle non entier	89
4.3.6	Application à un câble réel	93
4.4	Application à une ailette de refroidissement	98
4.4.1	Le système physique étudié	98
4.4.2	La modélisation nodale classique	98
4.4.3	Les synthèses de Foster et Cauver	99
4.4.4	Comparaison des modélisations	101
4.5	Application à la conduction thermique dans un cylindre	103

4.5.1	Modèle thermique d'ordre infini	104
4.5.2	Modèles réduits	106
4.5.3	Comparaison des modèles	107
4.6	Conclusion sur les développements d'impédances	111
5	Extension au cas des lignes polyphasées	113
5.1	Équation des télégraphistes matricielle	113
5.1.1	Modélisation de l'élément de ligne multifilaire	113
5.1.2	Passage en composantes modales	114
5.1.3	Résolution dans l'espace modal	115
5.1.4	Conditions aux limites	116
5.2	Développement d'impédances matricielles	116
5.2.1	Définition d'une impédances matricielles	117
5.2.2	Modèle en « Π » équivalent	118
5.2.3	Fractions continues matricielles	119
5.3	Synthèse de Cauer matricielle	121
5.3.1	Réseau d'impédances matricielles	121
5.3.2	Représentation du couplage capacitif	121
5.3.3	Représentation d'une résistance de défaut triphasée	122
5.4	Exemple de simulation d'une ligne triphasée	123
5.4.1	Méthode modale	124
5.4.2	Modèle approché de la ligne polyphasée	125
5.5	Avantages du développement d'impédances matricielles	125
6	Application à la détection de défauts	129
6.1	Modèle numérique d'une ligne en défaut	129
6.1.1	Synthèse de Cauer	129
6.1.2	Modèle numérique monophasé	131
6.1.3	Modèle numérique polyphasé	132
6.2	Méthode d'identification paramétrique	132
6.2.1	Méthode du gradient	132
6.2.2	Méthode de Gauss-Newton	133
6.2.3	Méthode de Levenberg-Marquardt	133
6.3	Mise en œuvre	133
6.3.1	Calcul du gradient	134
6.3.2	Calcul du Hessien	134
6.3.3	Calcul des fonctions de sensibilité	134
6.3.4	Simulation pour une ligne monophasée	135
7	Implémentation numérique des réseaux en échelle	143
7.1	Équivalence entre réseaux en échelle et filtres en treillis	143
7.1.1	Relations de récurrence	144
7.1.2	Filtres en treillis	144

Table des matières

7.2	Synthèse du filtre numérique	145
7.2.1	Discrétisation	145
7.2.2	Troncature	145
7.2.3	Résultats de simulation	146
7.2.4	Extension au cas polyphasé	150
7.3	Remarques sur le cas d'une ligne sans perte	150
7.3.1	Réponse impulsionnelle	150
7.3.2	Décomposition de la réponse impulsionnelle	151
7.3.3	Expression à l'aide des polynômes de Legendre	153
7.4	Conclusion sur filtrage récursif	159
Conclusion		161
A La transformation de Laplace		163
A.1	Définition et premières propriétés	163
A.2	Propriétés opérationnelles	164
A.3	Images de fonction usuelles	165
A.4	Transformation de Laplace inverse	165
A.5	Lien avec la transformée de Fourier	166
B Notes sur la résolution de l'équation des télégraphistes		167
B.1	Résolution par découplage des équations	167
B.1.1	Conditions aux limites	168
B.2	Utilisation de l'équation de Ricatti	169
B.2.1	Impédance locale	169
B.2.2	Résolution	169
B.2.3	Exemples de calcul d'impédances	170
B.3	Utilisation d'une exponentielle de matrice	171
C Polynômes orthogonaux		173
C.1	Familles de fonctions orthogonales	173
C.1.1	Produit scalaire et norme	173
C.1.2	Forme linéaire et moments	174
C.2	Familles de polynômes orthogonaux	175
C.2.1	Calcul direct	175
C.2.2	Récurrence à trois termes	176
C.2.3	Fonctions génératrices	177
C.3	Polynômes orthogonaux usuels	177
C.3.1	Polynômes de Legendre	177
C.3.2	Polynômes de Tchebycheff de première espèce	179
C.3.3	Polynômes de Tchebycheff de deuxième espèce	179
C.3.4	Polynômes de Gegenbauer	179
C.3.5	Polynômes de Jacobi	182

C.3.6	Polynômes de Hermite	182
C.3.7	Polynômes de Laguerre	182
D	Caractéristiques électriques d'un conducteur cylindrique	187
E	Représentation matricielle de réseaux élémentaires	189
E.1	Impédances et admittances monophasées	189
E.1.1	Résistance longitudinale	189
E.1.2	Conductance transversale	190
E.1.3	Inductance longitudinale	190
E.1.4	Capacité transversale	191
E.2	Impédances et admittances biphasées	191
E.2.1	Inductances longitudinales couplées	191
E.2.2	Capacités transversales couplées	192
E.3	Propriétés	194
E.3.1	Passage entre les représentations	194
E.3.2	Élément infinitésimal d'une ligne	195
E.3.3	Mise en parallèle d'impédances longitudinales	196
E.3.4	Mise en série d'impédances longitudinales	197
E.3.5	Couplage capacitif dans une impédance longitudinale	197

Table des matières

Introduction

Nous présentons dans ce mémoire la synthèse des travaux de recherche effectués en partie au laboratoire d'études thermiques de Poitiers et en partie au sein de l'usine *Alstom Transmission & Distribution* de Lattes (Hérault) dans le cadre d'une convention CIFRE.

Présentation

La sûreté d'alimentation en électricité dans les pays développés est cruciale pour l'économie et la vie courante. Les pannes majeures survenues dernièrement, le 20 août 2003 en Amérique du nord et le 28 septembre 2003 en Italie sont là pour nous le rappeler à juste titre.

Dans un réseau de transport et de distribution en exploitation, les lignes de transmission à haute tension sont soumises à différentes agressions de longue durée, comme le givre, la neige collante ou le défaut d'un isolateur, ou de très courte durée, mais beaucoup plus nombreuses, comme la foudre. Le rôle des protections contre les défauts d'isolation est d'éliminer l'élément du réseau concerné en ouvrant les organes de coupure adéquats, après avoir détecté et localisé le défaut. Il importe que cette action soit à la fois rapide et sélective. La rapidité limite les contraintes sur le matériel et les risques de pertes de synchronisme dus au déséquilibre. La sélectivité permet d'éviter la mise hors tension d'un trop grand nombre d'ouvrages.

Au niveau des algorithmes utilisés dans les protections numériques, il est absolument nécessaire de trouver le juste compromis entre simplicité, dont dépend la rapidité, et la précision dont dépend la sélectivité. C'est en gardant à l'esprit cette double contrainte qu'a été réalisé le travail de recherche présenté dans ce mémoire.

Objectif initial de l'étude

La modélisation des lignes électriques est apparemment un problème résolu depuis de nombreuses années. Cependant, des travaux initiés par M. Patrick Lagonotte au laboratoire d'études thermiques, à la fois à partir de développements des solutions analytiques ou de techniques de réduction de réseaux aboutissaient à un nouveau type de modèles sous la forme de réseaux comportant des condensateurs à capacités négatives de compensation. Sans augmenter le nombre de nœuds et la dimension des calculs matriciels, cette approche

permettait d'améliorer de manière sensible aussi bien le comportement en amplitude et en déphasage du régime harmonique à 50 Hz, que les réponses en régimes transitoires.

Du point de vue de la société *Alstom*, partenaire industriel de ce travail, la puissance des nouveaux processeurs de traitement du signal intégrés aux protections numériques permettait de concevoir des algorithmes de plus en plus complexes tout en conservant une très bonne réactivité du système de détection de défauts. L'objectif initial de cette thèse CIFRE, initiée par M. Éric Legrand, était d'intégrer dans un environnement temps réel embarqué de nouveaux algorithmes, afin d'améliorer et de sécuriser le comportement des protections.

Contributions apportées

Pour mener à bien ce projet ambitieux, nous avons donc décidé de rechercher une méthode de modélisation à la fois précise, et fournissant des modèles de lignes simples, pour ne pas nuire à la vitesse de l'algorithme de détection de défauts. Nous avons à notre disposition des résultats intéressants constituant une voie de recherche à explorer, mais sans réels fondements théoriques.

Nous avons donc mis en place dans un premier temps un formalisme mathématique rigoureux qui rattache les modèles de lignes proposés par M. Patrick Lagonotte aux fractions continues et aux approximants de Padé. Ce formalisme nous a permis d'aborder de façon unifiée la modélisation des systèmes linéaires à paramètres répartis et de généraliser les synthèses de Foster et de Cauer à une large classe de problèmes physiques.

Certaines extensions faites pour traiter des cas relevant de la modélisation thermique peuvent paraître surprenantes dans le cadre d'un mémoire traitant de la modélisation des lignes électriques. Cependant, cette thèse étant encadrée par le laboratoire d'études thermique de Poitiers, il nous a, au contraire, semblé opportun de souligner l'importance de nos travaux dans ce domaine. L'équation de diffusion de la chaleur et l'équation des télégraphistes constituant en effet les archétypes des systèmes à paramètres répartis dans le contexte de l'automatique, nous avons décidé de les traiter par une approche commune.

Bien que d'une certaine manière nous ne soyons pas allés jusqu'au bout de l'étude initialement prévue par la validation d'un prototype de protection de ligne et de localisation de défauts, cette dernière étape pourra être prise en charge au niveau recherche et développement industriel. En effet, nous montrons dans ce mémoire comment obtenir des modèles de lignes polyphasées par développement d'impédances et nous proposons également une méthode d'implémentation numérique.

Chapitre 1

La protection des lignes de transmission

1.1 Protections des réseaux électriques

Les réseaux électriques servent à assurer le transport d'énergie entre les producteurs et les utilisateurs, industriels ou domestiques. L'énergie est produite en France principalement par des centrales nucléaires (78 %), hydrauliques (12,1 %) ou thermiques (charbons, fuel, gaz pour 9,2 %)¹. Or ces centrales sont en général éloignées des centres de consommation car elle dépendent d'approvisionnement en combustible et en eau de refroidissement ou de contraintes géographiques (hydraulique).

1.1.1 Divers types de réseaux

Pour acheminer de façon efficace l'énergie électrique depuis les centres de productions jusqu'aux consommateurs finaux, on utilise plusieurs types de réseaux selon la fonction à assurer [Cor91].

Réseaux de transport et d'interconnexion Les réseaux de transport et d'interconnexion servent à collecter la puissance produite par les centrales importantes et à l'acheminer vers les zones de consommations. C'est la fonction de *transport*. Ils permettent aussi de répartir la production afin de permettre l'exploitation des ressources les plus économiques et d'optimiser ainsi les coûts d'exploitation. C'est la fonction d'*interconnexion*.

Réseaux de répartition régionale Leur fonction est d'amener l'énergie prélevée sur les réseaux de transport jusqu'aux gros utilisateurs industriels ou aux points d'injection sur les réseaux de distributions.

Réseaux de distribution publique Les réseaux de distribution publique ont pour rôle d'alimenter un grand nombre d'utilisateurs, industriels petits et moyens consommateurs et domestiques.

¹Chiffres RTE de la répartition de la production par sources d'énergie pour l'année 2002 en France.

1.1. Protections des réseaux électriques

Notons que différents domaines de tension reflètent ces diverses fonctions des réseaux électriques. Les réseaux de transports (réseau à 400 kV en France) et de répartition (225 kV, 90 kV et 63 kV en France) fonctionnent dans le domaine de la haute tension HTB (tableau 1.1). La fonction de distribution est assurée quant à elle par des réseaux fonctionnant dans les domaines allant de la basse tension (BTA et BTB) jusqu'en haute tension (HTA).

TAB. 1.1 – Domaines de tension des réseaux électriques en France.

Domaine de tension	Tension nominale en alternatif	Tension nominale en continu lissé
TBT	$U_n \leq 50 \text{ V}$	$U_n \leq 120 \text{ V}$
BTA	$50 \text{ V} < U_n \leq 50 \text{ V}$	$120 \text{ V} < U_n \leq 750 \text{ V}$
BTB	$500 \text{ V} < U_n \leq 1000 \text{ V}$	$750 \text{ V} < U_n \leq 1500 \text{ V}$
HTA	$1000 \text{ V} < U_n \leq 50000 \text{ V}$	$1500 \text{ V} < U_n \leq 75000 \text{ V}$
HTB	$50000 \text{ V} < U_n$	$75000 \text{ V} < U_n$

1.1.2 Utilité de la protection

Il y a plusieurs raisons de protéger un réseau de transport d'énergie. L'exploitant doit à la fois assurer la continuité et la qualité de service aux consommateurs, et assurer une durée de vie raisonnable au matériel d'exploitation. Pour cela, il est nécessaire d'éviter des situations anormales de fonctionnement du réseau, citons, d'après [GEC] :

- les situations anormales de fonctionnement d'ensemble du réseau, comme les insuffisances de production, les instabilités et les surtensions,
- les situations anormales de fonctionnement d'éléments du réseau, comme les surcharges des lignes et des transformateurs,
- l'apparition de court-circuits ou de défauts d'isolement sur certains éléments du réseau, en particulier les lignes ou les transformateurs.

L'exploitant se doit donc d'assurer à la fois la *sécurité* des personnes et du matériel, et la *disponibilité* de la fourniture d'énergie. Pour cela, il utilise des dispositifs de *protection* et des dispositifs *d'automatismes et de surveillance*.

Ces dispositifs servent à détecter les dysfonctionnements du réseau et à apporter une réponse appropriée, comme d'isoler une section de ligne en défaut. Il convient de remarquer que la sécurité s'oppose à la disponibilité puisque les dispositifs automatiques de protection provoquent souvent des interruptions de service [Sas94]. Remarquons que les automatismes de réenclenchement sont une solution technique permettant d'améliorer la disponibilité sans détériorer la sécurité. Ils permettent en effet d'éliminer les défauts fugitifs en limitant l'interruption de service (voir la section 1.1.3).

1.1.3 Défauts d'isolement

Les défauts de fonctionnement des éléments du réseau peuvent avoir plusieurs causes. Les lignes aériennes sont soumises aux aléas météorologiques. Les tempêtes et surtout la foudre constituent la principale origine des défauts sur ce type de liaisons. Les câbles souterrains sont soumis quant à eux aux agressions extérieures telles que les engins de terrassement qui peuvent endommager leurs gaines de protection.

Les défauts se traduisent par l'apparition de courts-circuits, parmi lesquels on distingue :

- les courts-circuits monophasés, entre une phase et la terre ou une masse,
- les courts-circuits biphasés, entre deux phases, et éventuellement la terre,
- les courts-circuits triphasés, entre les trois phases, et éventuellement la terre.

Ces courts-circuits se distinguent aussi par leur durée. Il existe des défauts *permanents*, qui résultent par exemple de l'endommagement d'un isolant, et des défauts *fugitifs* (sur les lignes aériennes), qui disparaissent à la mise hors tension de la ligne (extinction de l'arc). Dans ce dernier cas, la reprise du service est assurée par un simple réenclenchement. La fréquence d'apparition de ces défauts peut être estimée à partir des données statistiques fournies pour la France par le Réseau de transport d'électricité (RTE) [Rés02] dans le tableau 1.2.

TAB. 1.2 – Statistiques sur les défauts des réseaux aériens en France (source RTE).

Niveaux de tension	400 kV	230 kV	90 kV	63 kV
Nombre/100 km/an	3,8	12,1	14,3	29,3
Fugitifs	3,1	10,5	12,7	24,8
Monophasés	2,8	8,8	10,1	14,3
Polyphasés	0,3	1,6	2,6	4,5
Permanents	0,7	1,6	1,6	4,5
Monophasés	0,5	1,1	0,9	3
Polyphasés	0,2	0,5	0,7	1,5

1.1.4 Qualités des protections

Pour assurer une bonne continuité de service, le dispositif de protection doit être *sélectif*, c'est-à-dire qu'il doit éliminer du réseau, par l'intermédiaire du ou des disjoncteurs associés, l'élément affecté par un défaut, et seulement celui-ci.

Pour assurer une sécurité maximale, le dispositif doit être *rapide* pour limiter au mieux les effets des perturbations, de manière à éviter les détériorations du matériel et les risques d'instabilité du réseau. Enfin, il doit être *sensible* pour fonctionner même si les circonstances sont telles que les courants de défaut sont réduits (défauts résistants) [GEC].

Enfin, la *fiabilité* traduit l'aptitude d'une protection à remplir un double objectif. À savoir [Sas94] :

1.2. La protection de distance

- être sûr du déclenchement (sécurité),
- ne pas avoir de déclenchement intempestif (disponibilité).

Remarque En termes de théorie de la détection, ce compromis peut être quantifié en tant que probabilité de « non détection » et de probabilité de « fausse alarme ». La valeur maximale de la probabilité de fausse alarme étant fixée, on montre qu'on obtient la meilleure probabilité de détection avec un critère de décision basé sur le test de *Neymann-Pearson* (voir par exemple [Duv91], ch. 4 ou [Cha96], ch. 4).

1.1.5 Technologies des protections

Les premières protections électriques étaient basées sur des systèmes électromécaniques [Cor91]. Les grandeurs électriques mesurées étaient directement converties en actions mécaniques qui pouvaient déclencher des relais. Évidemment les parties mobiles de ces protections constituaient leurs faiblesses, soit par leur fragilité, soit par le risque de grippage.

Dans les années 70, les organes électromécaniques des protections ont été remplacés par des cartes électroniques assurant les mêmes fonctions. Plus fiables et plus rapides, les protections se sont aussi perfectionnées grâce à la souplesse apportée par cette nouvelle technologie.

Les protections numériques sont alors apparues dans les années 80 et ce sont elles qui dominent actuellement le marché. Cette évolution technologique s'accompagne de nombreux avantages. La puissance de calcul offerte par les microprocesseurs autorise l'utilisation d'algorithmes complexes et l'utilisation de mémoires permet d'enregistrer les signaux pour une analyse et un traitement différés. C'est grâce à l'utilisation des techniques numériques que les protections ont pu intégrer plus facilement la localisation de défauts.

1.2 La protection de distance

1.2.1 Principe de la protection de distance

Pour protéger les lignes aériennes ou les câbles souterrains, il existe des protections à *liaison pilote longitudinale*, parmi lesquelles on distingue les protections différentielles, les plus répandues, et les protections à comparaison de phases et à verrouillage directionnel presque exclusivement employées aux États-unis. Ces protections nécessitent l'échange d'informations entre les extrémités de la ligne par fils pilotes, liaisons par fibres optiques dans les câbles de garde, courants porteurs à haute fréquence, ou encore par ondes radio.

Au contraire, le principe de la protection de distance n'utilise que les mesures des courants et des tensions disponibles au point de mesure. À partir de ces grandeurs, il est nécessaire d'évaluer la distance X et les caractéristiques (résistance R_{def} , nature monophasée, biphasée ou triphasée) du défaut. Cela se fait en général en estimant l'impédance de la ligne vue de la position où se trouve la protection. La distance est alors déduite à

l'aide d'un modèle de ligne comme le résume le schéma de la figure 1.1. Les éléments clés d'une protection de distance sont donc le modèle de ligne et l'algorithme d'identification.

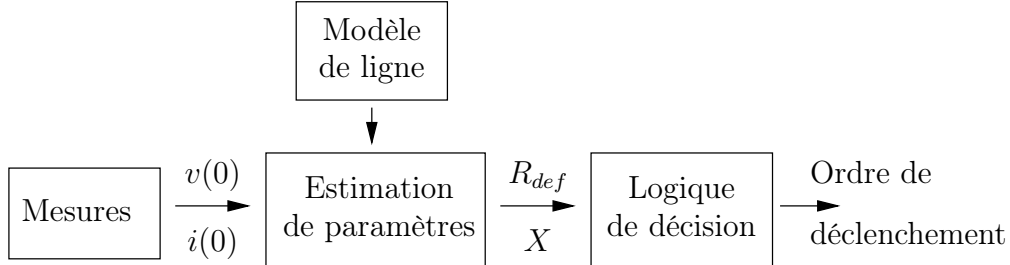


FIG. 1.1 – Principe général d'une protection de distance.

1.2.2 Exemple de protection numérique

Pour illustrer notre propos, détaillons les principales fonctions d'une protection de distance de type *MiCOM* P440 de la société Alstom (fig. 1.2). Le schéma de la figure 1.3 est une vue générale décrivant les entrées et les sorties de la protection. Nous pouvons tout d'abord remarquer, en haut et à gauche du schéma, les dispositifs de mesure des tensions et des courants de la ligne à protéger.



FIG. 1.2 – Une protection de distance de la série *MiCOM* P440 de la société Alstom.

La mesure de la tension de chacun des conducteurs se fait par l'intermédiaire d'un transformateur de tension (TT ou VT pour *voltage transformer*). Les trois phases plus le neutre correspondent aux entrées analogiques *C19* à *C22*. Les entrées optionnelles *C23* et *C24* servent à mesurer la tension sur un jeu de barres, dans le cas où le disjoncteur est ouvert. Cette mesure correspond à un test de synchronisation qui empêche la fermeture du disjoncteur en cas de déphasage trop important entre les tensions des barres et de la ligne.

Les entrées *C1* à *C9* sont reliées à des transformateurs de courant (TC ou CT pour *current transformers*, qui fournissent une image du courant de chacune des phases. Les entrées supplémentaires *C10* à *C12* sont utilisées pour la mesure directe de la composante homopolaire du courant à l'aide d'un CT spécifique. Il est possible d'utiliser deux câbles différents pour chacune des mesures de courant. Toutes ces entrées analogiques sont

1.2. La protection de distance

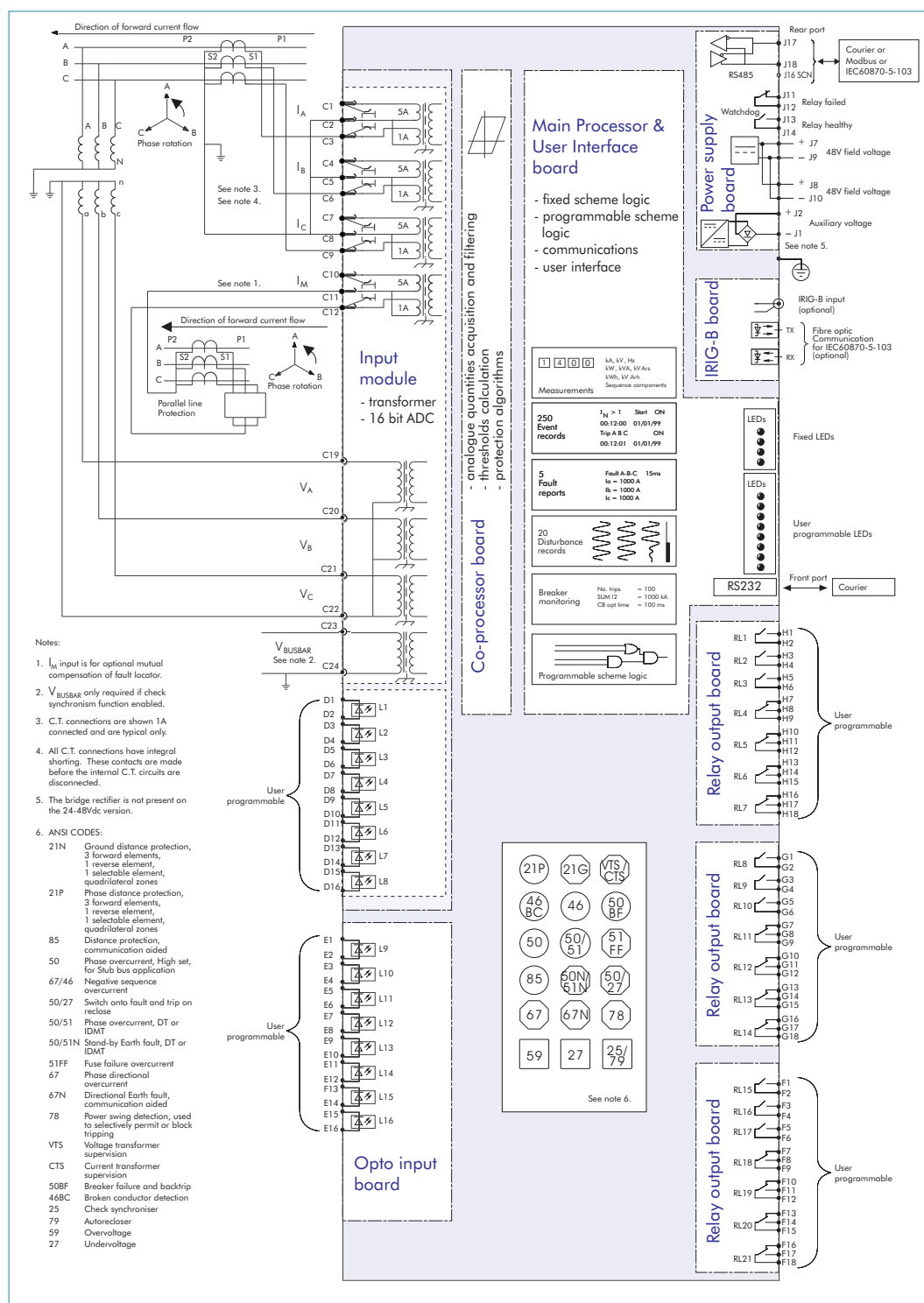


FIG. 1.3 – Schéma détaillé des entrées et sorties d'une protection de distance de type *MiCOM P440* de la société Alstom.

ensuite transformées par un convertisseur analogique-numérique, CAN pour être traitées numériquement.

Le traitement des signaux se fait au niveau de la carte coprocesseur. Il consiste en un certain nombre d'algorithmes visant à extraire les informations utiles des grandeurs mesurées. En particulier, l'algorithme de distance que nous détaillerons à la section 1.2.4 traduit les mesures des tensions et des courants « zones » d'impédances (voir fig. 1.5).

Les informations issues de la carte coprocesseur sont ensuite utilisées pour définir les actions à effectuer par l'intermédiaire de fonctions logiques fixes, ou programmées par l'opérateur. Par exemple, en fonction de la zone de détection d'un défaut, une temporisation de durée variable est utilisée pour que la protection la plus proche déclenche en premier, améliorant ainsi la selectivité. D'autre part, la partie logique gère les fonctions de réenclenchement, destinées à limiter l'interruption de service pour les défauts fugitifs. Là encore, des temporisations variables peuvent être programmées entre les réenclenchements successifs.

Des entrées numériques optiques ($D1$ à $D16$ et $E1$ à $E16$) sont aussi disponibles pour que l'utilisateur puisse adapter le comportement de la protection à sa politique de gestion du réseau. Ces entrées influent sur la *logique de décision* (cf. fig. 1.1). Parallèlement, un certain nombre d'événements et de signaux sont enregistrés pour post-traitement éventuel et analyse du défaut à l'aide de mémoires numériques.

Les actions possibles pilotées par la protection se font par l'intermédiaire de signaux de commande sur des relais qui constituent les sortie $F1-F18$, $G1-G18$ et $H1-H18$. Ces relais sont utilisés pour commander les disjoncteurs capables de couper le courant sur le conducteur où le défaut a été décelé.

Enfin, un certain nombre de connecteurs (RS232, fibre optique) sont disponibles pour assurer la communication de la protection avec un ordinateur de poste ou sa configuration et la récupération d'informations, par l'intermédiaire d'un logiciel de paramétrage (*MiCOM S1*) ou la synchronisation horaire (IRIG-B).

1.2.3 Modèle de ligne

À titre de référence, nous allons détailler le fonctionnement d'une protection de distance telle qu'on en trouve dans l'industrie. Le modèle de ligne est un modèle de ligne courte pour lequel on néglige les phénomènes capacitifs et les pertes diélectriques².

$$c = 0, \quad g = 0. \quad (1.1)$$

Nous représentons de plus le défaut par une simple résistance représentant la *résistance d'arc* R_{def} . Le circuit représentant le système total est représenté sur la figure 1.4. L'impédance d'entrée s'exprime immédiatement en fonction de la longueur de la ligne X .

$$\mathcal{Z} = (r + lp)X + R_{\text{def}} \quad (1.2)$$

²Les inductances mutuelles ne sont pas non plus prises en compte ici. Nous y reviendrons au chapitre 2 et plus en détails au chapitre 5.

1.2. La protection de distance

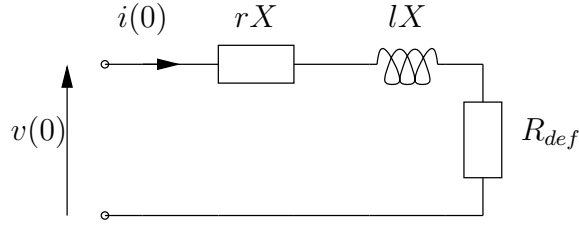


FIG. 1.4 – Modèle de ligne en défaut ($y = 0$).

1.2.4 Algorithme de distance

L'impédance d'entrée de la ligne à protéger est estimée à partir des mesures de la tension et du courant en $x = 0$. (schéma 1.1). Les protections numériques de distance PXLN ou *MiCOM*, série P440, du constructeur *Alstom* utilisent un algorithme de type moindre carrés récursifs pour estimer les caractéristiques du défaut (algorithme de Gauss-Seidel). Le vecteur de paramètres s'écrit donc

$$\theta = \begin{pmatrix} X \\ R_{\text{def}} \end{pmatrix}, \quad (1.3)$$

Connaissant les caractéristiques électriques du modèle (fig. 1.4) on calcule à partir des mesures de $i(0)$ la chute de tension normalisée le long de la ligne

$$w = ri(0) + l \frac{di(0)}{dt}. \quad (1.4)$$

Les mesures de la tension et du courant étant effectuées sur N points (échantillonnage), nous pouvons les présenter sous forme matricielle

$$\mathbf{v} = \begin{pmatrix} v_1(0) \\ v_2(0) \\ \vdots \\ v_N(0) \end{pmatrix}, \quad \mathbf{H} = \begin{pmatrix} w_1(0) & i_1(0) \\ w_2(0) & i_2(0) \\ \vdots & \vdots \\ w_N(0) & i_N(0) \end{pmatrix}. \quad (1.5)$$

L'identification des paramètres θ se ramène donc à la résolution d'un système linéaire, noté sous forme matricielle

$$\mathbf{v} = \mathbf{H}\theta. \quad (1.6)$$

Ce système d'équations est surdéterminé puisque l'on a en pratique beaucoup plus de points de mesure que de paramètres. La solution fournie par les moindres carrés est l'estimation

$$\hat{\theta}_N = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{v}. \quad (1.7)$$

Cette dernière relation s'écrit sous la forme récursive [Mun90]

$$\hat{X}_N = \frac{\sum_{n=1}^N v_n w_n - \hat{R}_{\text{def}N-1} \sum_{n=1}^N w_n i_n}{\sum_{n=1}^N w_n^2} \quad (1.8)$$

$$\hat{R}_{\text{def}N} = \frac{\sum_{n=1}^N v_n i_n - \hat{X}_{N-1} \sum_{n=1}^N w_n i_n}{\sum_{n=1}^N i_n^2} \quad (1.9)$$

On introduit enfin un *facteur d'oubli* exponentiel en remplaçant chacune des sommes de la forme $\sum_{n=1}^N a_n b_n$ par la quantité $S(a, b)_N$ définie par

$$S(a, b)_N = a_N b_N + \lambda S(a, b)_{N-1} \quad 0 < \lambda \leq 1 \quad (1.10)$$

1.2.5 Algorithme directionnel

L'algorithme présenté fournit une estimation de la distance X et de la résistance de défaut R_{def} , donc de l'impédance d'entrée

$$\hat{Z} = (r + lp)\hat{X} + \hat{R}_{\text{def}} \quad (1.11)$$

En posant $p = i\omega$, sa représentation dans le plan de Fresnel (fig. 1.5) permet de définir différentes zones associées à la distance du défaut. Sur le schéma, les zones (1), (2) et (3) correspondent à des défauts *aval*, la zone (4) correspondant quant à elle à un défaut *amont*.

En associant différentes temporisations précédant l'ordre de déclenchement des disjoncteurs à ces différentes zones, il est possible de protéger de façon redondante des portions de lignes, tout en respectant l'exigence de sélectivité (cf. section 1.1.4).

1.3 Amélioration des performances

La précision et la rapidité d'une protection dépendent donc des qualités de la modélisation choisie pour la ligne et de l'algorithme d'identification. Cela est d'autant plus vrai si l'on désire tirer profit des signaux de hautes fréquences qui peuvent résulter soit d'une injection, comme le proposait déjà J. Fallou [Fal32] en 1932, soit des fréquences propres de la ligne traduisant les réflexions de l'onde électromagnétique provoquée par l'apparition du défaut lui-même [Val00].

Les progrès constants dans le domaine des microprocesseurs et des DSP³ ne justifient plus de nos jours l'utilisation d'un modèle de ligne aussi simpliste que celui de la figure 1.4. Pour améliorer les performances des protections, il paraît intéressant de disposer de modèles de lignes ou de câbles plus précis et performants, tout en restant suffisamment simples pour être facilement mis en œuvre numériquement sous des contraintes de calculs en *temps réel*.

³Acronyme de *digital signal processor*, processeur de signaux numériques

1.3. Amélioration des performances

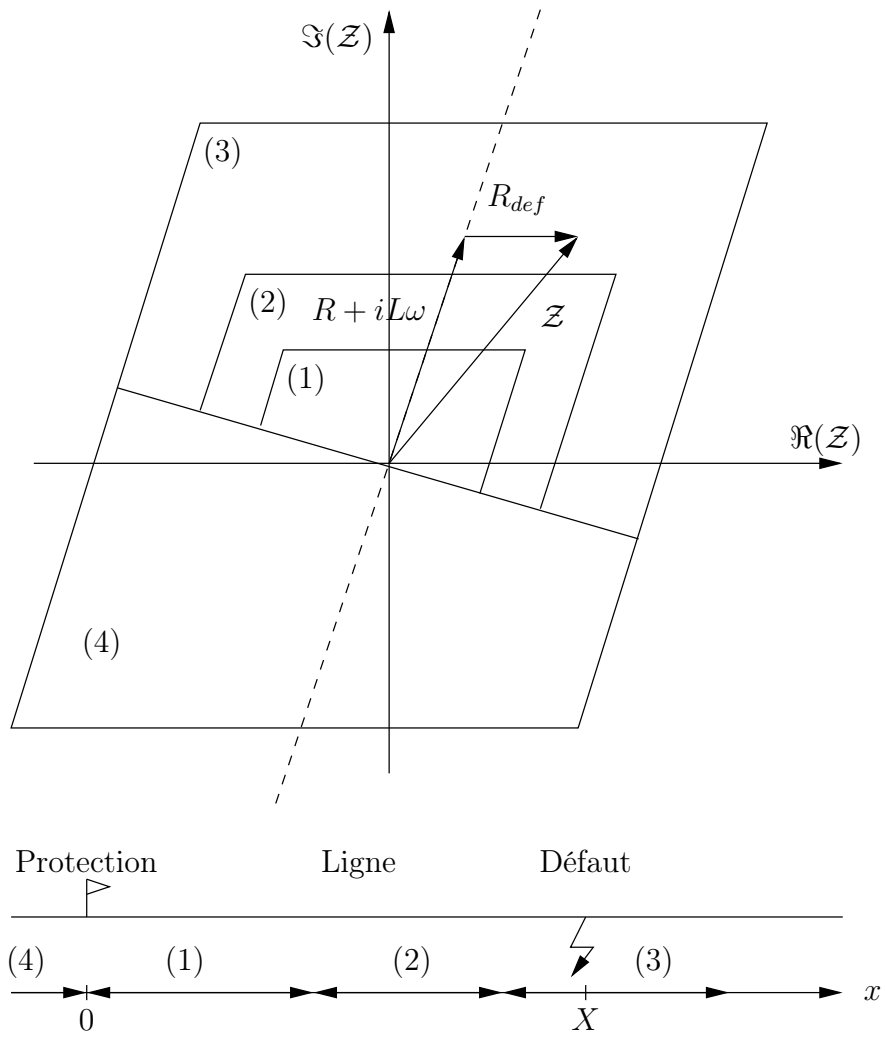


FIG. 1.5 – Zones dans le plan de Fresnel pour une protection de distance directionnelle.

Chapitre 1. La protection des lignes de transmission

Pour répondre au souhait de la société Alstom de faire évoluer ses protections numériques, il est important de rechercher une méthode de modélisation originale assurant le meilleur compromis entre simplicité et performance.

Avant de présenter le principe de modélisation choisi, il nous est apparu nécessaire de détailler dans le chapitre suivant les méthodes classiques de modélisation des lignes de transmission afin de peser les avantages et les inconvénients des différentes approches.

1.3. Amélioration des performances

Chapitre 2

Modélisation des lignes électriques

Nous désignons de nos jours par la locution *ligne de transmission* divers types de guides d'ondes électromagnétiques composés essentiellement de conducteurs filiformes. Ces conducteurs servent à véhiculer un signal électrique sur des distances relativement longues. Dès lors que ces distances ne sont plus négligeables par rapport à la longueur d'onde du signal, l'approximation des états stationnaires, appliquée souvent implicitement en électronique et en électrotechnique, n'est plus valable. En d'autres termes, deux éléments reliés par un conducteur ne sont plus équipotentiels. C'est pourquoi l'étude des lignes a nécessité une analyse théorique particulière, initiée à partir du milieu du XIX^e siècle, et qui est toujours d'actualité.

Si la téléphonie, et plus généralement la transmission d'informations sous forme d'impulsions électriques (réseaux informatiques) sont les utilisations les plus ostensibles des lignes de transmission, leur domaine d'application dépasse largement ce cadre. En ce qui concerne l'électrotechnique, par exemple, les lignes de transport d'énergie électrique des centrales de production aux zones de consommation sont un exemple de ligne de transmission. Les guides d'onde utilisés en radio-électronique, voire en hyperfréquences (radar, micro-ondes) en sont un autre exemple.

Ce vaste domaine d'utilisation est à mettre en relation avec la vaste gamme de fréquences des phénomènes étudiés, allant du continu à quelques dizaines de hertz pour les fréquences industrielles d'alimentation en énergie, jusqu'à plusieurs gigahertz pour des communications par satellites et des applications comme le système GPS.

2.1 L'évolution de la modélisation des lignes électriques

2.1.1 Du télégraphe au téléphone

Si les utilisations des lignes de transmission sont nombreuses, la première à donner lieu à un développement industriel fut le télégraphe. Les lignes télégraphiques parcouraient déjà les états d'Europe et d'Amérique bien avant qu'une approche théorique de la propagation d'un signal électrique ne fût conceptualisée. Jusqu'aux années 1850, en

2.1. L'évolution de la modélisation des lignes électriques

effet, la faible longueur des lignes (quelques dizaines de kilomètres) et la faible bande passante nécessaire à la transmission d'un signal télégraphique permettaient de considérer les câbles comme de simples contacts, tout au plus des résistances, ce qui suffisait à expliquer l'atténuation constatée avec la longueur de la ligne (fig. 2.1).

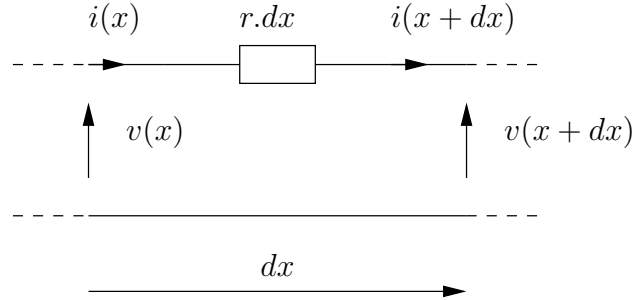


FIG. 2.1 – Segment élémentaire de ligne purement résistif ($z = r$, $y = 0$).

Dès 1851, pourtant, le premier câble sous-marin reliant Calais à Douvre permit de mesurer le temps de propagation des signaux et de mettre en évidence leur déformation. L'installation de câbles sous-marins encore plus longs, comme celui reliant Varna à Balaklava (741 km) en Mer Noire ou le premier câble transatlantique opérationnel (1866) posèrent alors le problème de la conception des lignes de transmission. Lord Kelvin (fig. 2.2) fut le premier à résoudre ce problème en proposant en 1856 un modèle de ligne qui tenait compte de la capacité répartie dans les câbles sous-marins (fig. 2.3). Dans ce type de câble, en effet, et aux faibles fréquences considérées, les effets inductifs sont négligeables au regard des effets capacitifs.

Pendant une trentaine d'années, le modèle proposé par Lord Kelvin fit office de dogme en matière de transmission de signaux. Ses faiblesses apparurent pourtant avec le développement du télégraphe multiplexe, qui utilise une plage plus large de fréquences afin de faire transiter simultanément plusieurs messages sur une même ligne. En effet, les hypothèses utilisées par le célèbre physicien, dans le cas d'un câble sous-marin et aux basses fréquences, se révélaient caduques pour des lignes aériennes. Ces lignes présentent des phénomènes d'induction importants et d'autant plus sensibles que le signal présente des hautes fréquences.

Le problème devint crucial avec l'avènement du téléphone, inventé conjointement par Graham Bell et Élisha Gray en 1876. La déformation des signaux analogiques due à la nature dispersive des lignes de l'époque limitait la transmission de la parole humaine à quelques dizaines de kilomètres. C'est alors qu'Oliver Heaviside (fig. 2.4), ancien télégraphiste et pionnier de l'électromagnétisme, se basant sur un modèle de ligne propagatif qui tient compte de l'inductance (fig. 2.5), montra que les signaux se propageaient sans déformation à la condition, dite depuis « de Heaviside », que

$$\frac{l}{c} = \frac{r}{g}, \quad (2.1)$$

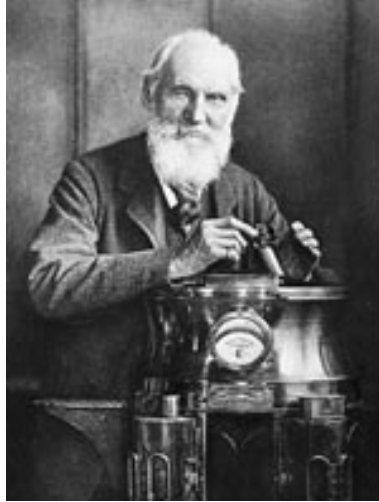


FIG. 2.2 – William Thomson, Lord Kelvin (1824-1907).
Source de l'image : université de Glasgow (<http://www.gla.ac.uk>).

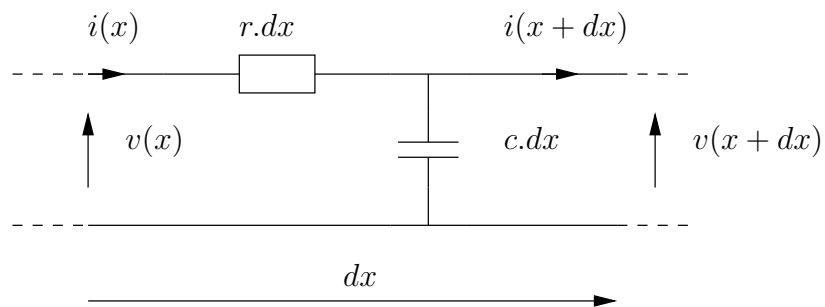


FIG. 2.3 – Segment élémentaire de ligne résistif et capacitif ($z = r$, $y = cp$).

2.1. L'évolution de la modélisation des lignes électriques

avec les notations de la figure 2.9.

La capacité linéique c étant en général bien plus faible, pour des lignes aériennes, que l'inductance linéique l , il proposa en 1887 l'emploi de bobines d'induction intercalées régulièrement le long des lignes pour éviter la distorsion du signal.



FIG. 2.4 – Oliver Heaviside (1850-1925).
Source de l'image : université de Saint Andrews
(<http://www-history.mcs.st-andrews.ac.uk>).

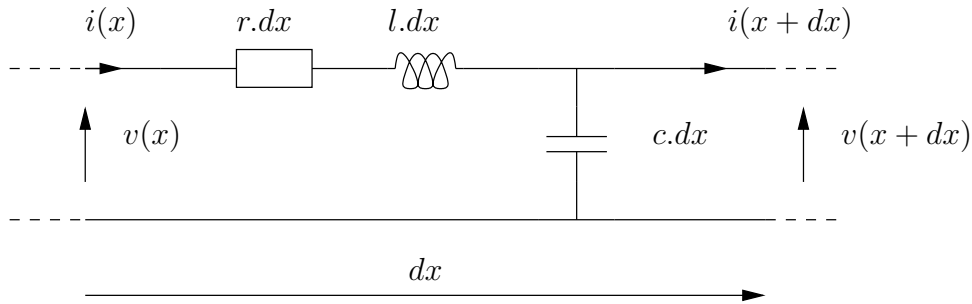


FIG. 2.5 – Segment élémentaire de ligne résistif, inductif et capacitif ($z = r + lp$, $y = cp$).

Cette proposition allait totalement à l'encontre des préceptes de l'époque qui visaient à réduire au minimum l'induction et il s'en suivit une bataille épistolaire mémorable entre Heaviside et William Preece, un fervent défenseur du modèle de Kelvin. Parallèlement, un français, Aimé Vaschy (fig. 2.6), arrivait aux mêmes conclusions qu'Heaviside en ce qui concerne l'inductance. Il écrivit dès 1886 dans les *Annales télégraphiques* qu'au contraire d'être « nuisible [...] l'effet de la self-induction de la ligne est essentiellement utile » [DC17].

L'histoire lui donna raison et, en 1902, le danois Carl Emil Krarup (1872-1909) proposa de réaliser l'augmentation de l'auto-inductance répartie (l) des lignes de transmission en



FIG. 2.6 – Aimé Vaschy (1857-1899).
Source de l'image : *Revue générale de l'électricité*.

enveloppant le conducteur de cuivre avec un matériau possédant des propriétés magnétiques, un fil de fer par exemple (fil *compound*) [Car93]. Le procédé préconisé par Heaviside fut finalement breveté en 1904 par l'américain d'origine yougoslave Michael Pupin sous une forme simplifiée, appelé pupinisation. Il consiste à placer en série et à intervalles réguliers des inductances bobinées sur une ligne téléphonique (fig. 2.8).

Pour bien saisir le contexte de ces avancées théoriques en électrotechnique, il est nécessaire de souligner la réelle difficulté expérimentale à cette fin de XIX^e siècle pour étudier et comprendre les phénomènes transitoires. Le condensateur est découvert en 1745 au travers de « la bouteille de Leyde », mais sa décharge ne traduit que l'existence d'une charge statique. Il faut attendre 1827 pour que Ohm publie son mémoire *Die galvanische Kette, mathematisch bearbeitet* comparant la conduction de l'électricité à celle de la chaleur, étudiée par Joseph Fourier (1768-1830) et publiée dans son mémoire *Théorie analytique de la chaleur* en 1822 [Fou22].

En 1830, Joseph Henry (1797-1878) découvre qu'un courant peut être induit dans un conducteur par déplacement d'un champ magnétique, c'est le principe de l'électromagnétisme qu'il ne publiera pas. En revanche, Michael Faraday (1791-1867), publiera sa découverte de l'induction électromagnétique en 1831. En 1833, Heinrich Lenz (1804-1865) découvre la loi donnant le sens des courants induits (loi de Lenz). Maxwell ne publiera ses « équations », qu'il expose dans son grand ouvrage, *Treatise on Electricity and Magnetism* qu'en 1873. Cependant le concept d'inductance et de circuit magnétique semble avoir été difficile à maîtriser. Il faudra attendre 1883 pour que John Hopkinson (1849-1898) publie sa loi qui est pour les circuits magnétiques, l'homologue de la loi d'Ohm pour les circuits électriques. Ce n'est que depuis 1893, que l'unité de résistance inductive porte le nom de Henry, alors que les unités de capacité électrique (farad) et d'intensité électrique (ampère) avaient été définies au congrès de Paris en 1881.

2.1. L'évolution de la modélisation des lignes électriques



FIG. 2.7 – Michael Idvorsky Pupin (1854-1935).

Source de l'image : institut Mihajlo Pupin (<http://www.imp.bg.ac.yu>).

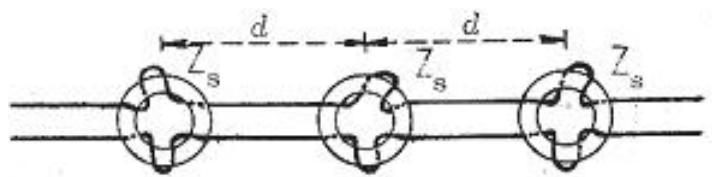


FIG. 2.8 – Ligne téléphonique *pupinisée*.

Source de l'image : *Théorie des réseaux de Kirchhoff*. [Bay54].

2.1.2 L'équation des télégraphistes

Le développement rapide des moyens de communication a donc été le moteur d'évolutions techniques majeures dans le domaine des transmissions. Elles ne sont cependant que le reflet des immenses progrès théoriques réalisés dans le domaine de l'électromagnétisme. Si Maxwell ne formule qu'en 1861 la notion d'onde électromagnétique, c'est dès 1857 que Kirchhoff établit, à partir des travaux de Lord Kelvin, l'*équation des télégraphistes*. Or cette équation suffit à rendre compte des phénomènes de propagation des signaux électriques le long d'une ligne, ainsi que le montrera Heaviside en 1876¹. Considérons en effet un modèle de ligne un peu plus général (fig. 2.9). Les *constantes linéiques réparties* qui représentent les caractéristiques électriques de la ligne sont au nombre de quatre :

- la résistance linéique r ($\Omega \cdot \text{m}^{-1}$) qui représente la chute de tension due à la dissipation d'énergie le long de la ligne (pertes Joule) ;
- l'inductance linéique l ($\text{H} \cdot \text{m}^{-1}$) qui traduit le stockage d'énergie sous la forme d'un champ magnétique autour et dans la ligne (auto-induction) ;
- le couplage capacitif (capacité linéique c en $\text{F} \cdot \text{m}^{-1}$) entre la ligne et le sol, ou entre la ligne et le conducteur de retour supposé au potentiel de référence, qui dénote l'énergie stockée sous la forme d'un champ électrique ;
- enfin, la conductance transversale linéique g ($\text{S} \cdot \text{m}^{-1}$) regroupant les pertes par conduction à travers le diélectrique (courants de fuite), et les pertes par « effet couronne » dans le cas de lignes à très haute tension [AI87, ch. 9].

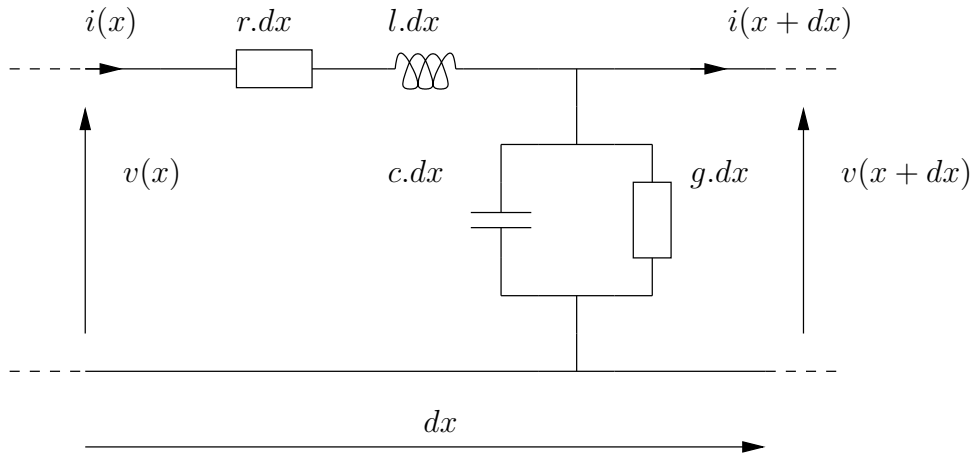


FIG. 2.9 – Segment élémentaire de ligne à quatre constantes linéiques ($z = r + lp$, $y = g + cp$).

À l'aide du schéma électrique de la figure 2.9, nous pouvons écrire les équations aux dérivées partielles suivantes qui décrivent les variations de la tension et du courant en

¹Le lecteur intéressé par l'apport de Heaviside à la théorie des lignes électriques pourra consulter l'ouvrage de L. Cohen [Coh35].

2.1. L'évolution de la modélisation des lignes électriques

fonction du temps t et de la variable d'espace x .

$$\begin{cases} \frac{\partial v}{\partial x} = -ri - l \frac{\partial i}{\partial t}, \\ \frac{\partial i}{\partial x} = -gv - c \frac{\partial v}{\partial t}. \end{cases} \quad (2.2)$$

Dérivons la première équation par rapport à la variable x , puis substituons la dérivée de i par rapport à x à l'aide de la deuxième équation afin de « découpler », c'est-à-dire de séparer, les variables v et i .

$$\begin{aligned} \frac{\partial^2 v}{\partial x^2} &= -r \frac{\partial i}{\partial x} - l \frac{\partial^2 i}{\partial x \partial t} \\ &= r \left(gv + c \frac{\partial v}{\partial t} \right) + l \frac{\partial}{\partial t} \left(gv + c \frac{\partial v}{\partial t} \right) \\ &= lc \frac{\partial^2 v}{\partial t^2} + (rc + lg) \frac{\partial v}{\partial t} + rgv \end{aligned} \quad (2.3)$$

Nous pouvons agir de même en éliminant v de la deuxième équation de 2.2.

$$\frac{\partial^2 i}{\partial x^2} = lc \frac{\partial^2 i}{\partial t^2} + (rc + lg) \frac{\partial i}{\partial t} + rgi \quad (2.4)$$

Les grandeurs v et i vérifient donc la même équation aux dérivées partielles du second degré, qualifiée de « terrible équation » par Vaschy [DC17], et connue sous le nom d'*équation des télégraphistes*.

2.1.3 Significations physiques des paramètres de la ligne

Nous avons jusqu'à présent caractérisé une ligne de transmission à l'aide des constantes linéiques r , l , c et g , brièvement définies à la section 2.1.2. Précisons maintenant la signification physique de ces paramètres et les conditions nécessaires à leurs définitions rigoureuses.

En effet, il est délicat de transposer directement les lois de Kirchhoff aux phénomènes transitoires qui ont lieu sur une ligne de transmission. Ces lois qui régissent l'*électrocinétique* dans l'*approximation des états quasi-stationnaires* atteignent leur limite de validité lorsque la dimension caractéristique d'un système électrique n'est plus négligeable par rapport à la longueur d'onde des ondes électromagnétiques qui diminue avec la fréquence des signaux. Les lignes de transport peuvent couvrir plusieurs centaines de kilomètres et les variations rapides des grandeurs électriques se traduisent par des composantes en haute fréquence.

Nous appliquerons donc l'approche proposée par M. Péliissier [Pél75] en rappelant que le, ou les conducteurs, formant une ligne de transmission sont des guides d'ondes électromagnétiques. C'est le champ électromagnétique autour des conducteurs qui véhicule l'énergie et ses caractéristique sont couplées par les équations de Maxwell. La tension v et

le courant i que nous avons définis en un point de la ligne sont les quantités macroscopiques aisément mesurables induites par la présence de ce champ.

Les ondes électromagnétiques se déplacent à la vitesse de la lumière dans le diélectrique entourant les conducteurs,

$$c = \frac{1}{\sqrt{\epsilon\mu}}, \quad (2.5)$$

en notant ϵ la permittivité et μ la perméabilité absolues du milieu. Dans l'air, qui constitue le diélectrique des lignes aériennes, la vitesse de la lumière est proche de celle dans le vide, soit $299\,792\,458 \text{ km s}^{-1}$.

Au niveau d'un conducteur, le passage d'une onde électromagnétique se traduit par la présence d'une charge linéique q . Le principe de conservation de la charge électrique s'écrit

$$\frac{\partial i}{\partial x} + \frac{\partial q}{\partial t} + i_f = 0, \quad (2.6)$$

le courant de fuite i_f étant par exemple dû au courant superficiel des isolateurs.

Une relation de type capacitif peut être établie entre la charge électrique et la tension à la condition que le champ électrique entre les conducteurs et entre les conducteurs et le sol soit peu différent du champ *électrostatique*. Cette condition est remplie pour une propagation avec des pertes faibles. Dans ce cas, le champ électrique $\mathbf{E}$ est approximativement contenu dans le plan perpendiculaire aux conducteurs (fig. 2.10).

$$E_n \gg E_t. \quad (2.7)$$

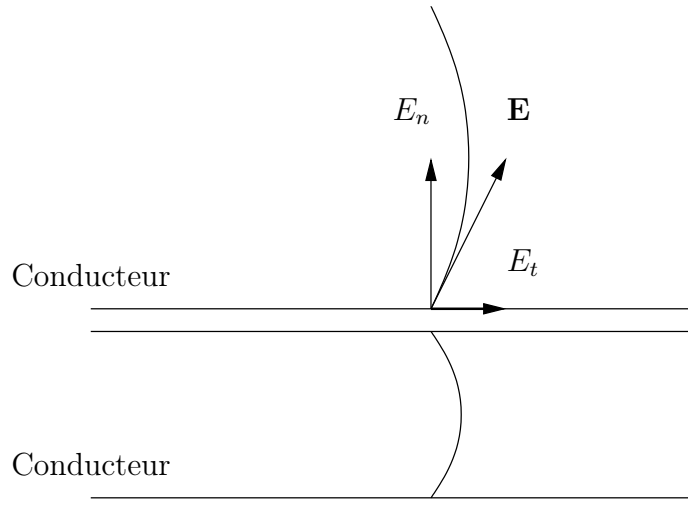


FIG. 2.10 – Ligne de champ électrique autour des conducteurs ([Pél75]).

Le potentiel du conducteur n s'exprime par une relation linéaire à l'aide des *coefficients*

2.1. L'évolution de la modélisation des lignes électriques

d'influence électrostatique Δ_{nk} et des charges q_k des N conducteurs formant la ligne,

$$v_n = \sum_{k=1}^N \Delta_{nk} q_k \quad (2.8)$$

ou, pour une ligne monophasée,

$$v = \Delta q. \quad (2.9)$$

La capacité linéique se définit alors par

$$c = \frac{1}{\Delta}. \quad (2.10)$$

Il reste à supposer que les courants de fuites sont proportionnels à la tension pour introduire la conductance transversale g ,

$$i_f = gv, \quad (2.11)$$

pour obtenir à partir de la loi de conservation de la charge (2.6) la deuxième équation du système (2.2).

La première équation du système représente quant à elle la chute de tension le long de la ligne. On peut l'attribuer à deux effets. Nous avons déjà cité la chute *ohmique* par effet Joule que reflète le paramètre r . Elle provient du fait que la conductivité des métaux (aluminium et cuivre) constituant les conducteurs n'est pas infinie.

La deuxième cause de la chute de tension est la variation du flux d'induction

$$\frac{\partial \Phi}{\partial t} \quad (2.12)$$

pour lequel on distingue le flux intérieur au conducteur Φ_i et le flux extérieur Φ_e . La propagation des ondes électromagnétiques à l'extérieur du conducteur fournit la relation

$$\Phi_e = l_e i, \quad l_e = \frac{\Delta}{c^2}. \quad (2.13)$$

L'expression du flux intérieur est plus délicate. Il provient aussi du fait que les matériaux ne sont pas infiniment conducteurs (auquel cas les courants seraient purement superficiels). Nous utiliserons en première approximation un coefficient d'induction constant,

$$\Phi_i \approx l_i i, \quad (2.14)$$

ce qui est justifié à condition que les sections des conducteurs et la fréquence des signaux soient *assez faibles*, comme nous le verrons au chapitre 4.

Le paramètre inductif de la ligne l se définit alors par

$$l = l_e + l_i. \quad (2.15)$$

2.2 Résolution dans l'espace de Laplace

L'équation des télégraphistes n'est résolue que sous la forme d'une série, à la fin du XIX^e siècle². Il faut attendre André Blondel (fig. 2.11) pour que la solution soit exprimée à l'aide de fonctions hyperboliques en 1909 [BR09, Blo20, Blo37].



FIG. 2.11 – André Blondel (1863-1938).
Source de l'image : *Revue générale de l'électricité*.

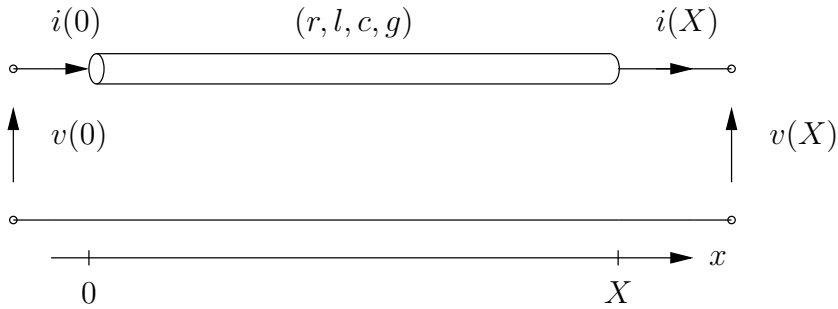


FIG. 2.12 – Modèle de ligne à paramètres répartis.

2.2.1 Transformation de Laplace

Considérons le système d'équations 2.2. Chacune de ses équations fait intervenir à la fois la tension v et le courant i . On parle alors de système d'équations *couplées*. La

²Heaviside, dans *the Electrician*, le 3 juin 1887, et Vaschy, dans les *Annales télégraphiques* de novembre-décembre 1888.

2.2. Résolution dans l'espace de Laplace

résolution d'un tel système peut se faire facilement en remplaçant l'opérateur de dérivation temporelle $\frac{\partial}{\partial t}$ par la variable symbolique $p \in \mathbb{C}$. Cette méthode fut utilisée par Heaviside sous le nom de *calcul opérationnel*³.

Par abus de notation, et pour éviter la multiplication des symboles, nous écrirons respectivement $v(p, x)$ et $i(p, x)$ les transformées de Laplace de $v(t, x)$ et de $i(t, x)$.

$$v(p, x) = \int_0^\infty e^{-pt} v(t) dt \quad i(p, x) = \int_0^\infty e^{-pt} i(t) dt \quad (2.16)$$

En supposant la tension et le courant nuls dans toute la ligne à l'instant $t = 0^+$, nous réécrivons le système formé des équations aux dérivées partielles (2.2) sous la forme d'un système d'équations différentielles :

$$\frac{dv}{dx} + (r + lp)i = 0 \quad \frac{di}{dx} + (g + cp)v = 0 \quad (2.17)$$

Pour des commodités de notations, mais aussi à des fins de généralisation, nous introduisons l'*impédance longitudinale* par unité de longueur z ($\Omega \cdot \text{m}^{-1}$) et l'*admittance transversale* par unité de longueur y ($\text{S} \cdot \text{m}^{-1}$).

$$z = r + lp \quad y = g + cp \quad (2.18)$$

Dérivons ensuite par rapport à la variable x chacune des équations précédentes, nous découplons les variables, et nous obtenons donc la transformée de Laplace des équations (2.3) et (2.4).

$$\frac{d^2 v}{dx^2} + z \frac{di}{dx} = 0, \quad \frac{d^2 i}{dx^2} + y \frac{dv}{dx} = 0. \quad (2.19)$$

Et en utilisant (2.17) et (2.18),

$$\frac{d^2 v}{dx^2} - zyv = 0, \quad \frac{d^2 i}{dx^2} - zyi = 0. \quad (2.20)$$

Nous reconnaissons la transformée de Laplace de l'équation des télégraphistes (2.4) vérifiée à la fois par v et i . C'est une équation du second degré en x . L'opérateur $\frac{\partial^2}{\partial x^2}$ est appelé *laplacien* (ici pour un problème unidimensionnel en coordonnées cartésiennes). Ce type d'équation se rencontre dans de nombreux domaines de la physique donnant lieu à des phénomènes de *diffusion* ou de *propagation*. Nous y reviendrons dans la section 2.3.

³Heaviside, en résolvant ces équations grâce au calcul opérationnel, ne se soucia pas de trouver une démonstration rigoureuse de sa méthode, ce qui lui a été reproché par les mathématiciens de l'époque. La justification du calcul opérationnel se trouve dans les transformations intégrales. Nous utiliserons en particulier la *transformation de Laplace* (A). Nommée ainsi par les mathématiciens George Boole (1815-1864) et Jules-Henri Poincaré (1854-1912), cette technique semble pourtant être en grande partie due aux travaux du hongrois Petzval (1807-1891). Mais la transformation de Laplace est une approche *analytique* du calcul opérationnel de Heaviside, remarquons qu'il se justifie aussi plus directement de façon *algébrique* notamment depuis les travaux de Mikusiński [Mik83]

La résolution de ces équations différentielles ne présente pas de difficultés particulières. Pour alléger ce chapitre tout en produisant un document complet, nous en avons reporté les détails en annexe B. Cette annexe présente par ailleurs trois méthodes de résolution différentes. La première est la résolution par découplage des équations, de loin la plus répandue, mais il est aussi possible de se ramener à une équation de Riccati pour le calcul direct de l'impédance (deuxième méthode), ou de faire une résolution globale du système en utilisant une exponentielle de matrice. Cette dernière méthode permet d'exprimer les relations entre les tensions et les courants aux extrémités de la ligne sous la forme

$$\begin{cases} v(p, x) = \cosh(\sqrt{zy}x)v(p, 0) - \sqrt{\frac{z}{y}} \sinh(\sqrt{zy}x)i(p, 0), \\ i(p, x) = -\frac{\sinh(\sqrt{zy}x)}{\sqrt{\frac{z}{y}}}v(p, 0) + \cosh(\sqrt{zy}x)i(p, 0). \end{cases} \quad (2.21)$$

Ces équations décrivent totalement le modèle de ligne choisi vu de ses extrémités (0, X). Exprimées sous forme matricielle, et en introduisant les grandeurs caractéristiques

- facteur de propagation, $\gamma = \sqrt{zy}$,
- impédance caractéristique, $Z_c = \sqrt{\frac{z}{y}}$,

elles donnent naturellement lieu à une représentation de type *quadripolaire*.

$$\begin{pmatrix} v \\ i \end{pmatrix} (p, X) = \begin{pmatrix} \cosh(\gamma X) & -Z_c \sinh(\gamma X) \\ -\frac{\sinh(\gamma X)}{Z_c} & \cosh(\gamma X) \end{pmatrix} \begin{pmatrix} v \\ i \end{pmatrix} (p, 0) \quad (2.22)$$

2.2.2 Modèles équivalents

Il est possible de définir des quadripôles équivalents à la ligne vue de ses bornes extrêmes sous forme de réseaux présentant différentes topologies. Les modèles quadripolaires les plus utilisés sont le modèle en Π , le modèle en T, et dans une moindre mesure le modèle en treillis ([Bay54], ch. 15). Ils sont représentés sur la figure 2.13. Il suffit de procéder par identification de la matrice chaîne (2.22) avec celles des modèles pour en déduire les valeurs des impédances qui les composent (voir par exemple [Fel86], ch. 11).

Par exemple, considérons le premier schéma de la figure 2.13 (modèle en Π). Les lois de Kirchhoff nous permettent d'exprimer les relations liant les courants et les tensions d'entrée et de sortie.

$$\begin{cases} v(p, X) = v(p, 0) - Z_{\Pi}(i(p, 0) - Y_{\Pi}v(0)) \\ i(p, X) = i(p, 0) - Y_{\Pi}v(p, 0) - Y_{\Pi}v(p, X) \end{cases} \quad (2.23)$$

D'où la matrice chaîne

$$\begin{pmatrix} v \\ i \end{pmatrix} (p, X) = \begin{pmatrix} 1 + Z_{\Pi}Y_{\Pi} & -Z_{\Pi} \\ -2Y_{\Pi} - Z_{\Pi}Y_{\Pi}^2 & 1 + Z_{\Pi}Y_{\Pi} \end{pmatrix} \begin{pmatrix} v \\ i \end{pmatrix} (p, 0) \quad (2.24)$$

2.2. Résolution dans l'espace de Laplace

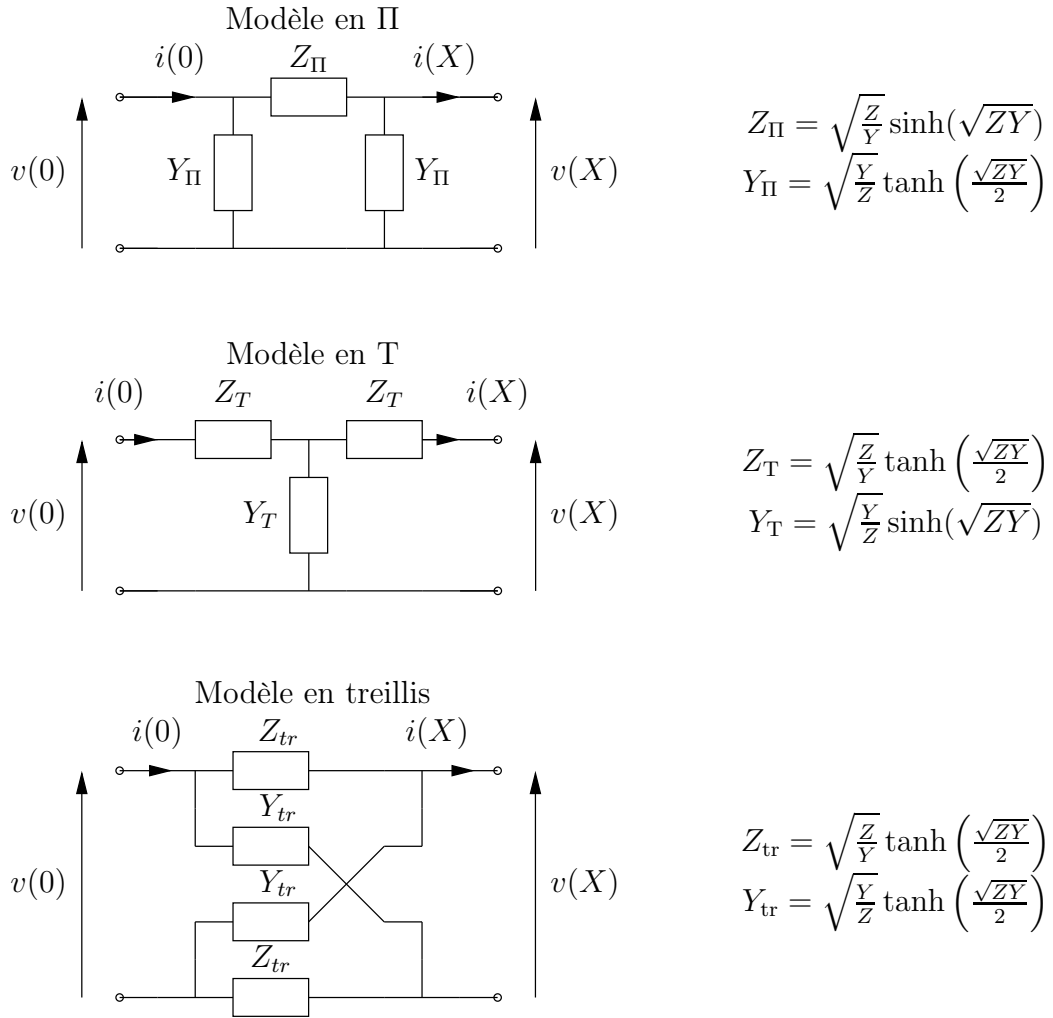


FIG. 2.13 – Modèles quadripolaires d'une ligne de transmission.

En identifiant termes à termes les éléments de la matrice aux éléments de la matrice 2.22, on détermine les valeurs de Z_{Π} et Y_{Π} (fig. 2.13).

Sous certaines hypothèses, les modèles de lignes peuvent se ramener à des dipôles (fig. 2.14). C'est en particulier le cas lorsque la ligne étudiée est soit ouverte, soit en court-circuit à son extrémité $x = X$. Le modèle se confond alors avec l'impédance d'entrée $\mathcal{Z}(p, X)$, par définition,

$$\mathcal{Z}(p, X) = \frac{v(p, 0)}{i(p, 0)}. \quad (2.25)$$

Dans la cas d'une ligne en court-circuit, $v(p, X) = 0$, la première équation de (2.21) nous permet d'écrire immédiatement :

$$\mathcal{Z}(p, X) = \sqrt{\frac{Z}{Y}} \tanh(\sqrt{ZY}). \quad (2.26)$$

De même, dans la cas d'une ligne ouverte, $i(p, X) = 0$ la deuxième équation de (2.21) donne

$$\mathcal{Z}(p, X) = \sqrt{\frac{Z}{Y}} \coth(\sqrt{ZY}). \quad (2.27)$$

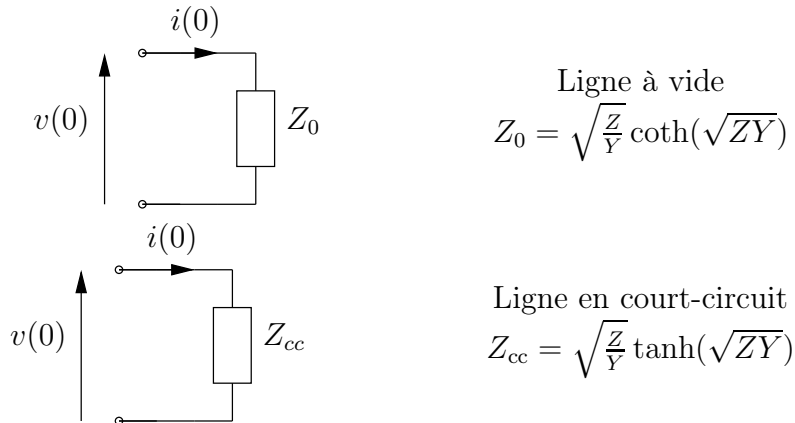


FIG. 2.14 – Modèles dipolaires d'une ligne de transmission.

Remarquons immédiatement que les expressions des impédances qui composent ces modèles, aussi bien quadripolaires que dipolaires, font toutes intervenir des fonctions *transcendantes* telles les fonctions usuelles de la trigonométrie hyperbolique. Cette propriété est caractéristique des systèmes à paramètres répartis. Nous verrons dans la section 2.3 qu'elle se traduit en général par un ordre infini, ce qui constitue la principale difficulté de leur étude.

2.2.3 Modèles à constantes localisées

Nous venons de voir que les modèles équivalents à une ligne de transmission peuvent être représentés par des impédances dont les expressions sont difficiles à manipuler. Pour

2.2. Résolution dans l'espace de Laplace

cette raison, il est souvent préférable d'utiliser des modèles de lignes approchés, plus usuels, et que l'on peut manipuler avec les outils classiques de la théorie des réseaux ou de l'automatique.

Plus précisément, en ce qui concerne l'approche « réseaux », ces modèles approchés sont à paramètres localisés, c'est-à-dire que les éléments électriques qui les composent (résistances, inductances et capacités) sont ponctuels. Ils n'ont pas de dimension géométrique (on ne s'intéresse pas aux grandeurs électriques internes) mais se comportent comme des opérateurs mathématiques : intégrateurs, dérivateurs purs et facteurs de proportionnalité⁴.

En ce qui concerne l'approche « automatique », les impédances des modèles approchés, assimilables à des *fonctions de transfert* liant la tension et le courant sont d'ordre fini, c'est à dire que ce sont des fonctions rationnelles de la variable de Laplace p .

Lignes courtes

Dans le cas où la ligne à étudier est suffisamment courte, il est habituel de procéder à une approximation du premier ordre des impédances d'ordre infini des modèles de la section 2.2.2. Pour cela on effectue un développement de Taylor (ou plutôt de Mac-Laurin, puisque c'est un développement au voisinage de l'origine) suivant la variable X . Considérons par exemple le modèle en Π de la figure 2.13.

$$Z_{\Pi} = \sqrt{\frac{z}{y}} \sinh(\sqrt{zy}X) = zX + \frac{z^2y}{6}X^3 + \frac{y^2z^3}{120}X^5 + O(X^6) \quad (2.28)$$

$$Y_{\Pi} = \sqrt{\frac{y}{z}} \tanh\left(\frac{\sqrt{zy}}{2}X\right) = \frac{y}{2}X - \frac{zy^2}{24}X^3 + \frac{z^2y^3}{240}X^5 + O(X^6) \quad (2.29)$$

En ne conservant que les termes du premier ordre,

$$Z_{\Pi} \approx zX = Z, \quad Y_{\Pi} \approx \frac{y}{2}X = \frac{Y}{2}. \quad (2.30)$$

En revenant à la définition de $z = r + lp$ et $y = g + cp$, et en introduisant les grandeurs électrique « globales » :

$$R = rX, \quad L = lX, \quad C = cX, \quad G = gX, \quad (2.31)$$

il vient

$$Z_{\Pi} \approx R + Lp, \quad Y_{\Pi} \approx \frac{G}{2} + \frac{C}{2}p. \quad (2.32)$$

Nous obtenons le modèle de ligne courte de la figure 2.15. Cette approximation n'est bien sûr valable que pour des valeurs de X suffisamment faibles. Pour une ligne aérienne fonctionnant à 50 Hz, nous pouvons retenir l'ordre de grandeur de la limite de 200 km, donnée par M. Escané ([Esc97], ch. 9).

⁴Ce formalisme est parfaitement décrit dans le premier chapitre de [BN96]

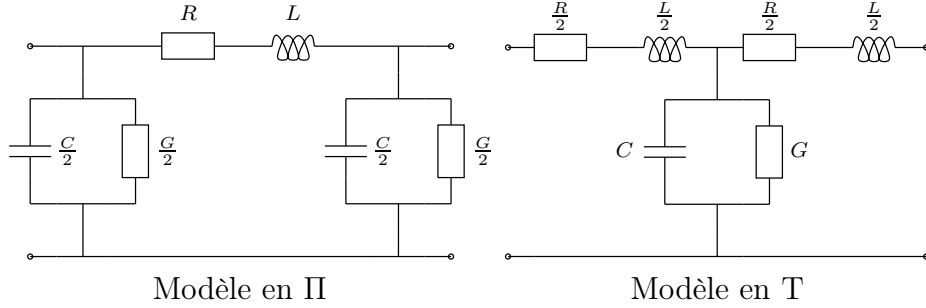


FIG. 2.15 – Modèles quadripolaires d’une ligne courte.

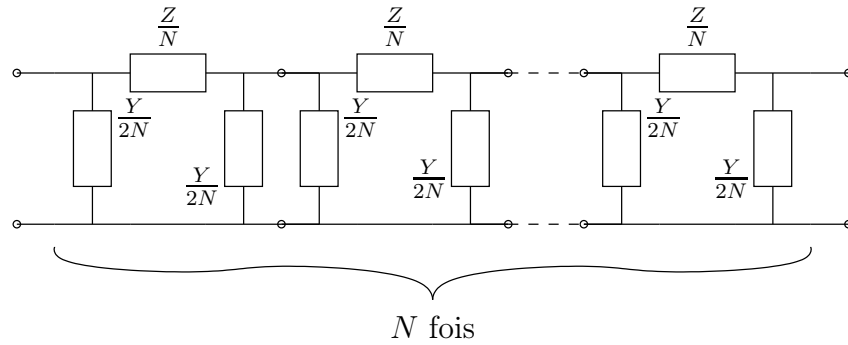
La même approximation peut être faite pour les modèles quadripolaires en T et en treillis (fig. 2.13). Le modèle en treillis étant plus rarement utilisé nous ne fournirons que le modèle en T sur la figure 2.15.

$$Z_T \approx \frac{R}{2} + \frac{L}{2}p, \quad Y_T \approx G + Cp, \quad (2.33)$$

$$Z_{tr} \approx \frac{R}{2} + \frac{L}{2}p, \quad Y_{tr} \approx \frac{G}{2} + \frac{C}{2}p. \quad (2.34)$$

Chaînage des cellules

Pour des lignes de quelques centaines de kilomètres, l’approximation des lignes courtes n’est donc plus valable. Il est alors courant de subdiviser la ligne en N portions de longueur $\frac{X}{N}$ compatible avec l’approximation du premier ordre précédente. Les bornes des modèles de lignes courtes étant accessibles, il est alors possible de *chaîner* les *cellules* constituées par les différents tronçons de ligne (fig. 2.16) [MK00].


 FIG. 2.16 – Chaînage de cellules en Π .

Nous pouvons simplifier le modèle de ligne obtenu en remplaçant les admittances en parallèle de valeur $\frac{Y}{2N}$ par leur somme. Remarquons que si nous avions procédé de même à partir des modèles de lignes courtes en T, ce sont des impédances en série de valeur $\frac{Z}{2N}$ que

2.2. Résolution dans l'espace de Laplace

nous aurions remplacées. Dans les deux cas nous obtenons pratiquement le même réseau de Kirchhoff à paramètres localisés, les différences résidant aux extrémités des réseaux⁵.

Discrétisation spatiale

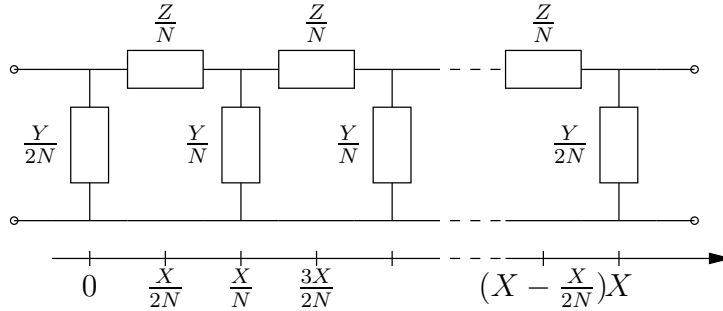


FIG. 2.17 – Chaînage de cellules en Π , après simplification.

Toutes simplifications faites, le modèle de ligne obtenu par chaînage est donc celui de la figure 2.17. Il est intéressant de le comparer au modèle à paramètres répartis décrits en 2.1.2 par les équations (2.17).

$$\begin{cases} \frac{dv}{dx} + z(p, x)i = 0, \\ \frac{di}{dx} + y(p, x)v = 0. \end{cases} \quad (2.35)$$

Dans ce cas, z et y sont constants le long de la ligne.

$$z(p, x) = r + lp = z(p), \quad (2.36)$$

$$y(p, x) = g + cp = y(p). \quad (2.37)$$

Dans le cas du modèle de la figure 2.17, et en excluant les extrémités, les lois de Kirchhoff fournissent des relations de la forme

$$\begin{cases} v \left(\left(n + \frac{1}{2} \right) \frac{X}{N} \right) - v \left(\left(n - \frac{1}{2} \right) \frac{X}{N} \right) + \frac{Z}{N} i = 0 \\ , i \left(n \frac{X}{N} \right) - i \left((n - 1) \frac{X}{N} \right) + \frac{Y}{N} v = 0. \end{cases} \quad (2.38)$$

⁵En effet, le réseau obtenu par chaînage de cellules en Π débute et se termine par une admittance $\frac{Y}{2N}$, alors que celui obtenu par chaînage de cellules en T débute et se termine par des impédances $\frac{Z}{2N}$. Notons enfin que cette différence sera d'autant plus faible que N sera grand puisque le reste du réseau est identique dans les deux cas [LBS99].

Ce qui peut s'écrire :

$$\begin{cases} \int_{(n-\frac{1}{2})\frac{X}{N}}^{(n+\frac{1}{2})\frac{X}{N}} \left(\frac{dv}{dx} + z_e(p, x)i \right) dx = 0, \\ \int_{(n-1)\frac{X}{N}}^{n\frac{X}{N}} \left(\frac{di}{dx} + y_e(p, x)v \right) dx = 0. \end{cases} \quad (2.39)$$

Nous retrouvons donc l'expression des relations locales entre courant et tension, similaire à (2.35) à condition de poser

$$z_e = z \frac{1}{N} \sum_{n=0}^{N-1} \delta \left(x - \left(n + \frac{1}{2} \right) \frac{X}{N} \right) \quad (2.40)$$

$$y_e = y \left(\frac{1}{2N} (\delta(x) + \delta(x - X)) + \frac{1}{N} \sum_{n=1}^{N-1} \delta \left(x - n \frac{X}{N} \right) \right) \quad (2.41)$$

$$\begin{cases} \frac{dv}{dx} + z_e(p, x)i = 0 \\ \frac{di}{dx} + y_e(p, x)v = 0 \end{cases} \quad (2.42)$$

Le symbole δ représentant la distribution de Dirac (au sens des distributions de Schwartz [Sch66]). Le passage du modèle à paramètres répartis (continu) au modèle à paramètres localisés (discret) se fait par multiplication par un *peigne de Diracs*. C'est le formalisme d'un *échantillonnage spatial*. Le modèle que nous avons obtenu par chaînage de modèles de lignes courtes est donc assimilable à une simple discrétisation spatiale⁶.

2.3 Les systèmes à paramètres répartis

La particularité des systèmes physiques à *paramètres répartis* est d'être décrits par des *équations aux dérivées partielles*, par opposition aux systèmes à constantes *localisées* (section 2.2.3) décrits par des *équations différentielles ordinaires*, c'est-à-dire d'une seule variable. Nous avons vu que l'équation des télégraphistes (2.3) et (2.4) faisait partie des équations aux dérivées partielles du second degré linéaires de deux variables x et t . Nous allons préciser quelques propriétés et le vocabulaire associé.

2.3.1 Équations aux dérivées partielles

La classe des équations aux dérivées partielles du second degré linéaires se rencontre dans de nombreux problèmes de la physique des phénomènes de transports. Ils regroupent principalement toutes les ondes propagatives (acoustiques, mécaniques, électromagnétiques, etc.) et les phénomènes diffusifs (conduction de la chaleur, diffusion d'espèces

⁶Cette méthode est aussi appelée méthode *nodale*, notamment dans le domaine de la modélisation thermique et elle est à rapprocher des méthodes par éléments finis.

2.3. Les systèmes à paramètres répartis

chimiques, etc.). Les mathématiciens ont donc très tôt étudié ce type d'équations et en ont dégagé une classification.

Classification

La forme la plus générale pour une équation aux dérivées partielles du second degré de deux variables indépendantes x et t peut s'exprimer par

$$A \frac{\partial^2 f}{\partial x^2} + 2B \frac{\partial^2 f}{\partial x \partial t} + C \frac{\partial^2 f}{\partial t^2} + D \frac{\partial f}{\partial x} + E \frac{\partial f}{\partial t} + Ff = G. \quad (2.43)$$

Si les coefficients A, B, C, D, E, F et G ne sont fonctions que de x et t , cette équation est de plus linéaire. Selon les valeurs de ces coefficients, on montre de plus que l'on peut se ramener par changement de variables à l'une des trois formes canoniques suivantes [Gar64] :

- équation hyperbolique si $B^2 - AC > 0$

$$\frac{\partial^2 f}{\partial \tilde{x}^2} - \frac{\partial^2 f}{\partial \tilde{t}^2} + \dots = 0, \quad (2.44)$$

- équation parabolique si $B^2 - AC = 0$

$$\frac{\partial^2 f}{\partial \tilde{x}^2} + \dots = 0, \quad (2.45)$$

- équation elliptique $B^2 - AC < 0$

$$\frac{\partial^2 f}{\partial \tilde{x}^2} + \frac{\partial^2 f}{\partial \tilde{t}^2} + \dots = 0. \quad (2.46)$$

Les points de suspension représentent les termes qui ne font intervenir que f et ses dérivées premières.

Appliquons cette classification à l'équation des télégraphistes (2.3) ou 2.4. Par identification, nous avons immédiatement

$$\begin{cases} A = 1 & B = 0 \\ C = -lc & D = 0 \\ E = -rc - lg & F = -rg \\ G = 0. \end{cases} \quad (2.47)$$

L'inductance et la capacité linéiques étant réelles positives, la quantité $B^2 - AC = lc$ est donc positive ou nulle. L'équation des télégraphistes ne peut donc prendre que deux formes canoniques, celle de l'équation hyperbolique (2.44) si l et c sont strictement positifs, ou celle de l'équation parabolique (2.45) si l ou bien c est nul⁷. Nous laisserons donc de côté pour la suite de cette étude l'équation elliptique.

⁷Remarquons que si l et c sont nuls, nous n'avons plus à faire à une équation aux dérivées partielles, mais à une équation différentielle ordinaire, et nous sortons du cadre de cette étude.

Forme réduite hyperbolique

Afin de simplifier les expressions des équations sans sacrifier à leur généralité, il peut être utile d'effectuer des changements de variables. Nous cherchons ainsi à obtenir une forme réduite de l'équation des télégraphistes.

Commençons par nous intéresser au cas de l'équation hyperbolique (2.44), encore appelée *équation des ondes*. Dans ce cas, l et c sont non nuls. Dans un premier temps, supposons qu'il en est de même pour r , c'est-à-dire

$$l > 0, \quad c > 0, \quad r > 0, \quad (2.48)$$

ce qui autorise les changements de variables dans le système (2.2)

$$\tilde{v} = \sqrt{c}v, \quad \tilde{i} = \sqrt{l}i, \quad \tilde{t} = \frac{r}{l}t, \quad \tilde{x} = r\sqrt{\frac{c}{l}}x. \quad (2.49)$$

Pour la première équation, nous obtenons la forme réduite

$$\begin{aligned} \frac{\partial \tilde{v}}{\partial \tilde{x}} &= \frac{\sqrt{l}}{r} \frac{\partial v}{\partial x} \\ &= \frac{\sqrt{l}}{r} (-ri - l \frac{\partial i}{\partial t}) \\ &= -\sqrt{l}i - \frac{l\sqrt{l}}{r} \frac{\partial i}{\partial t} \\ &= -\tilde{i} - \frac{\partial \tilde{i}}{\partial \tilde{t}}. \end{aligned} \quad (2.50)$$

Pour la deuxième équation de (2.2), le même changement de variables donne

$$\frac{\partial \tilde{i}}{\partial \tilde{x}} = -\frac{lg}{rc}\tilde{v} - \frac{\partial \tilde{v}}{\partial \tilde{t}}. \quad (2.51)$$

Il s'avère donc nécessaire d'introduire un coefficient sans dimension $\alpha = \frac{lg}{rc}$ pour garder toute la généralité des équations. Le système réduit est donc

$$\begin{cases} \frac{\partial \tilde{v}}{\partial \tilde{x}} = -\tilde{i} - \frac{\partial \tilde{i}}{\partial \tilde{t}} \\ \frac{\partial \tilde{i}}{\partial \tilde{x}} = -\alpha \tilde{v} - \frac{\partial \tilde{v}}{\partial \tilde{t}} \end{cases} \quad (2.52)$$

L'étude générale de l'équation des télégraphistes dans le cas hyperbolique pourra donc se faire en prenant pour les expressions des impédances et admittances réparties

$$z = 1 + p, \quad y = \alpha + p. \quad (2.53)$$

Les cas particuliers suivants sont à signaler.

2.3. Les systèmes à paramètres répartis

1. $\alpha = 0$, ligne sans pertes « transversales » ($g = 0$) qui correspond au schéma de la figure 2.5 ;
2. $\alpha = 1$, ligne non dispersive. L'impédance caractéristique $Z_c = 1$ ne dépend pas de la fréquence, ce qui correspond à la condition de Heaviside. C'est ce que l'on cherche à obtenir par les procédés de pupinisation ou krarupisation (voir page 16) ;

Pour illustrer l'influence du paramètre α sur la forme de la réponse fréquentielle du système, nous pouvons considérer le cas simple d'une ligne en court-circuit. Nous avons vu (figure 2.14) que dans ce cas, l'impédance se résume à

$$\mathcal{Z} = \sqrt{\frac{Z}{Y}} \tanh(\sqrt{ZY}), \quad (2.54)$$

ou encore, dans le cas d'équations réduites, et en prenant $x = 1$,

$$\mathcal{Z} = \sqrt{\frac{1+p}{\alpha+p}} \tanh(\sqrt{(1+p)(\alpha+p)}) \quad (2.55)$$

Pour tracer la réponse du système en fonction de la fréquence ν sur les figures 2.18 à 2.20, nous avons simplement posé

$$p = 2i\pi\nu, \quad (2.56)$$

ce qui permet d'obtenir la transformée de Fourier à partir de la transformée de Laplace pour les fonctions causales (annexe A).

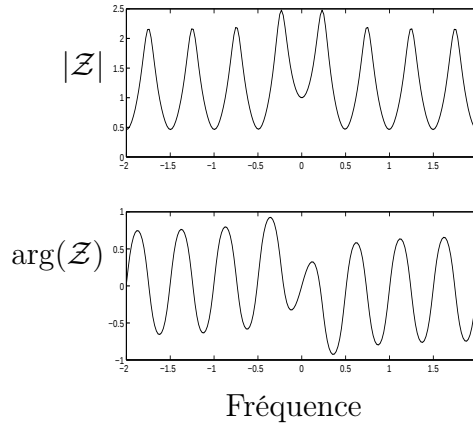


FIG. 2.18 – Réponse fréquentielle de la ligne en court-circuit ($\alpha = 0$).

Il nous faut traiter maintenant le cas où r est égal à zéro, que nous avons laissé de côté pour effectuer le changement de variables (2.49). Tout d'abord, si g est différent de zéro, nous pouvons nous ramener au cas précédent par *dualité* tension-courant. Le cas de la *ligne sans perte*, pour lequel

$$r = 0, \quad g = 0 \quad (2.57)$$

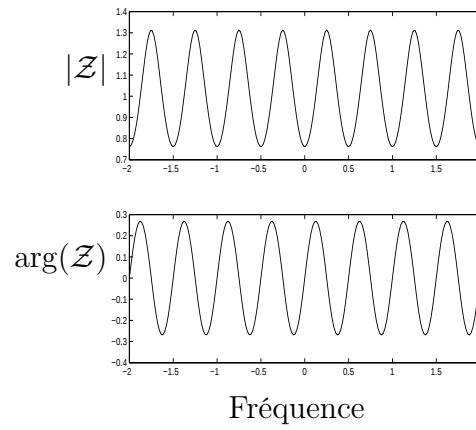


FIG. 2.19 – Réponse fréquentielle de la ligne en court-circuit ($\alpha = 1$).

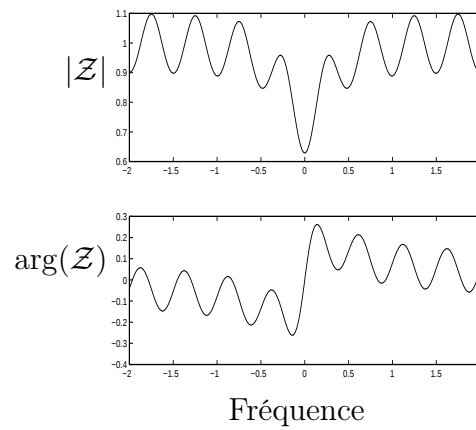


FIG. 2.20 – Réponse fréquentielle de la ligne en court-circuit ($\alpha = 2$).

2.3. Les systèmes à paramètres répartis

nécessite un autre changement de variables.

$$\tilde{v} = \sqrt{c}v, \quad \tilde{i} = \sqrt{l}i, \quad \tilde{t} = \frac{t}{\sqrt{lc}}, \quad \tilde{x} = x. \quad (2.58)$$

Appliquée au système (2.2), nous avons simplement

$$\begin{cases} \frac{\partial \tilde{v}}{\partial \tilde{x}} = -\tilde{i} \\ \frac{\partial \tilde{i}}{\partial \tilde{x}} = -\tilde{v}, \end{cases} \quad (2.59)$$

ce qui correspond aux constantes linéiques

$$z = p, \quad y = p. \quad (2.60)$$

En équations réduites, l'impédance d'une ligne sans pertes en court-circuit est donc

$$\mathcal{Z} = \tanh(px). \quad (2.61)$$

Ce qui donne une expression extrêmement simple de la transformée de Fourier, d'après (2.56), et avec $x = 1$,

$$\mathcal{Z} = \tan(2\pi\nu) \quad (2.62)$$

Cette réponse est tracée sur la figure 2.21.

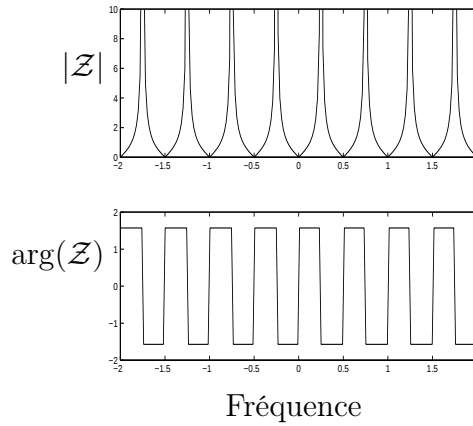


FIG. 2.21 – Réponse fréquentielle de la ligne sans pertes en court-circuit.

Le cas de la ligne sans pertes présente des particularités très intéressantes. C'est un système qui comporte des retards purs. Signalons aussi que ce cas particulier fait apparaître des liens entre la synthèse de Cauer présentée au chapitre 4 et les polynômes de Legendre (annexe C) qui seront étudiés dans le chapitre 7.

Forme réduite parabolique

L'équation parabolique (2.45) peut être vue comme une forme *dégénérée* de l'équation des télégraphistes pour laquelle soit l , soit c est égal à zéro. Posons par exemple $l = 0$. L'équation (2.3) devient

$$\frac{\partial^2 v}{\partial x^2} = rc \frac{\partial v}{\partial t} + rgv. \quad (2.63)$$

Nous verrons au chapitre 4 que cette expression décrit la température au sein d'une ailette de refroidissement, qui est un élément de dissipateur thermique (section 4.4). Nous pouvons aussi dire, que le cas encore plus simple où g est pris égal à zéro, et qui correspond au schéma de la figure 2.3, donne la température au sein d'un « mur » thermique [Mar90, MAB⁺00]. Pour cette raison, l'équation parabolique (2.45) est encore appelée *équation de la chaleur*.

2.3.2 Systèmes d'ordre infini

Dans les précédents exemples, nous avons pu remarquer que les systèmes à paramètres répartis se modélisaient dans le cas général par des impédances exprimées par des fonctions transcendentes. Le terme *impédance* est utilisé dans les domaines de l'électrotechnique et de l'électronique, mais peut très bien s'interpréter comme une *fonction de transfert* pour utiliser une approche *automatique*.

En effet si le courant i désigne la variable d'entrée du système, et v la variable de sortie, la loi d'Ohm devient une équation de transfert.

$$i(p) \rightarrow \boxed{\mathcal{Z}(p)} \rightarrow v(p) = \mathcal{Z}(p)i(p). \quad (2.64)$$

Ces fonctions de transfert peuvent avoir un nombre infini de zéros et de pôles, ce qui correspond à un *ordre infini*. L'étude des systèmes d'ordre infini ne peut pas se faire exactement de la même façon que celle des système d'ordre fini, pour laquelle les outils mathématiques sont bien connus.

Théorie du semigroupe

La théorie du semigroupe des opérateurs linéaires sur un espace de Hilbert est utilisée pour généraliser les concepts de l'automatique d'ordre fini et étendre le formalisme d'état à l'ordre infini. En effet, si l'on considère l'opérateur qui fait passer le système de l'état initial $X(0)$ à l'état $X(t)$ à l'instant t ,

$$X(t) = T(t)X(0), \quad (2.65)$$

alors on a les propriétés

$$T(0) = I, \quad I \text{ étant l'opérateur identité,} \quad (2.66)$$

$$T(t+s) = T(t)T(s), \quad s, t \geq 0. \quad (2.67)$$

2.3. Les systèmes à paramètres répartis

Et s'il existe une limite

$$AX = \lim_{t \rightarrow 0} \frac{1}{t} (T(t) - I)X, \quad (2.68)$$

alors A est appelé *générateur infinitésimal* du semigroupe T et vérifie formellement l'« équation d'état »

$$\frac{dX}{dt} = AX + Bu \quad (2.69)$$

Avec cette définition, A coïncide pour les systèmes d'ordre n avec les matrices d'ordre $n \times n$, et généralise effectivement la notion d'état. Le cas des systèmes d'ordre infini est étudié en détails dans [CZ95] et représente encore à l'heure actuelle un domaine de recherche particulièrement fertile. Citons par exemple les travaux récents de M. Touré sur le contrôle d'un échangeur de chaleur [Tou02] ou de M. Rouchon et M. Fliess sur la « commande plate » [Rou02].

Systèmes d'ordres non-entiers

Une autre voie récemment explorée pour l'étude des systèmes à paramètres distribués est la dérivation fractionnaire. La plupart de ces systèmes physiques ont en effet une réponse fréquentielle se comportant comme une puissance demi-entière sur une large bande de fréquence (processus à « mémoire longue »). Il se prête donc très bien à une modélisation par des opérateurs d'ordre fractionnaires, qui généralisent les notions de dérivation et d'intégration, [Mat02] dans [AGV02], [Ous98, LSOH99]. Cette approche a notamment été utilisée pour modéliser l'induction dans les machines tournantes [Riu01], pour les transferts thermiques [BPA02] ou bien en acoustique [Mat94].

La représentation par un système d'ordre fractionnaire offre aussi des possibilités très intéressantes de synthèse physique. Ceci a notamment été mis à profit pour la mise au point de commandes ou d'asservissements robustes [Ous83, Ous95]. Et la représentation sous forme de réseaux de Kirchhoff, par le biais d'analogies électriques, a été en particulier étudié par [Sab98].

Modèles d'ordres réduits

Il nous reste à évoquer le cas où un problème ne peut pas se traiter directement par la modélisation d'ordre infini ou non-entier. Cela arrive notamment quand interviennent des contraintes dans le domaine temporel. Par exemple lorsqu'un système n'est pas invariant dans le temps⁸

Une approche très simple consiste alors à approcher le système par un ordre fini (modèle d'ordre réduit). Cependant, de multiples stratégies d'approximation s'offrent à l'ingénieur qui désire se ramener au cadre classique des ordres finis. Nous en présenterons quelques unes dans le chapitre 3.

Parmi les approximations dans le domaine fréquentiel, nous pouvons citer les développements de Taylor que nous avons déjà rencontré à la section 2.2.3, les développements

⁸C'est en particulier le cas lors de l'apparition d'un défaut sur une ligne de transmission.

en séries de Fourier, les développements en éléments simples, les approximants de Padé et les méthodes modales.

2.4 Axe de recherche développé dans ce mémoire

Toutes ces différentes approches récemment développées montrent l'intérêt scientifique de la résolution du problème de modélisation des systèmes à paramètres répartis, dans laquelle s'inscrit celle des lignes électriques. D'un point de vue économique, l'industrie des protections et de la localisation de défauts demande aussi des modèles performants, c'est-à-dire à la fois rapide et précis, pour pouvoir développer de nouvelles applications et réduire les temps et les coûts d'intervention.

Les méthodes basées sur la théorie du semigroupe, si elles ont l'avantage d'une grande rigueur, sont délicates à mettre en œuvre dans le cadre des protections de distance. En effet, les contraintes du temps réel (imposées par la puissance limitée des DSP) nous font privilégier la simplicité du modèle et surtout l'interprétation physique des paramètres.

La modélisation non entière ne reflète en général que le comportement asymptotique de la fonction de transfert. De plus, cette approche se double souvent d'une approximation fréquentielle dans nombre de cas pratiques. On a donc au final deux approximations successives. D'un point de vue plus fondamental, les ordres non-entiers ont le défaut d'introduire des singularités essentielles dans le plan fréquentiel complexe qui n'ont pas de réelles justifications physiques⁹, au contraire des pôles, que l'on relie facilement aux modes ou aux résonances d'un système.

Ces considérations nous ont fait choisir l'utilisation de modèles d'ordres réduits. Cependant, les approximations fréquentielles classiques aboutissent en général à des modèles de type « boîte noire », dont les paramètres sont difficiles à interpréter physiquement, par exemple la méthode modale appliquée aux lignes de transmission par [WC97, WC98].

Nous montrerons cependant qu'en modifiant légèrement ces méthodes d'approximations, il est possible d'obtenir des modèles à la fois simples et performants. De plus, nous pourrions représenter ces modèles sous forme de réseaux de Kirchhoff, à l'instar des méthodes basées sur les ordres non-entiers, tout en gardant une *interprétation physique* directe de la valeurs des paramètres. Ceci représente un avantage de taille pour l'application à la protection des lignes de transmission où l'identification des paramètres est le but recherché *in fine*.

Cette approche découle de l'approfondissement des travaux de P. Lagonotte [LN00, LPC00] et vise à répondre au souhait de la société *Alstom* de s'intéresser à cette nouvelle voie.

⁹Avec toutefois une réserve pour les systèmes possédant une structure fractale comme le montre [Ous95].

2.4. Axe de recherche développé dans ce mémoire

Chapitre 3

De l'approximation polynomiale à l'approximation rationnelle

Dans le chapitre précédent, nous avons présenté succinctement les systèmes à paramètres répartis. Nous avons montré qu'ils se traduisent, pour le cas général, par des fonctions de transfert d'ordre infini dans le domaine fréquentiel. La voie que nous allons explorer consiste à trouver des approximations de ces fonctions de transfert sous la forme de fonctions d'ordre fini.

Aussi, il paraît nécessaire de présenter quelques unes des approximations exploitables. Plus précisément, nous développerons dans ce chapitre l'approximation polynomiale qui découle naturellement du développement de Taylor largement exploités pour la résolution de problèmes physiques. Puis nous soulignerons ses limites et nous proposerons des méthodes d'approximations rationnelles. En particulier, nous détaillerons le théorème de Mittag-Leffler permettant de conserver les pôles ou les zéros des fonctions approchées, puis nous présenterons les théories des *fractions continues* et des *approximants de Padé* que nous avons retenus pour la suite de cette étude.

Ces puissants outils mathématiques qui, comme nous le verrons, sont intimement liés, sont utilisés dans des domaines aussi variés que l'analyse, les probabilités, la mécanique ou la théorie des nombres [Khi64] [Niv56]. Par la suite, ces approximations purement formelles donneront lieu à une représentation sous forme de réseaux électriques. Nous obtiendrons ainsi des modèles physiques d'ordres réduits, c'est-à-dire finis, de systèmes à paramètres distribués.

3.1 L'approximation polynomiale

Le physicien qui doit trouver une approximation d'une fonction, recherche bien souvent les hypothèses plausibles qui lui permettent de qualifier une grandeur physique comme étant négligeable (ou au contraire prépondérante) devant les autres. Ceci autorise à tronquer un développement en série de la fonction suivant cette grandeur particulière.

Nous avons déjà vu un exemple de cette démarche dans l'hypothèse des *lignes courtes*

3.1. L'approximation polynomiale

de la section 2.2.3. Le développement se fait dans ce cas suivant la longueur de la ligne électrique, et on s'autorise à ne conserver que le premier terme du développement en faisant l'hypothèse que cette longueur est *suffisamment petite*.

3.1.1 Développement limité

Nous allons maintenant généraliser cette approche. Considérons la série entière formelle

$$f(x) = \sum_{n=0}^{\infty} c_n x^n. \quad (3.1)$$

Nous recherchons une approximation de $f(x)$ au voisinage de $x = 0$. En tronquant la série (3.1) à un ordre arbitraire $N \in \mathbb{N}$, nous assurons que l'erreur d'approximation est d'ordre $N + 1$.

$$f_N(x) = \sum_{n=0}^N c_n x^n \quad (3.2)$$

$$f(x) - f_N(x) = O(x^{n+1}) \quad (3.3)$$

La notation de Landau O (« grand o ») ci-dessus signifie qu'il existe un voisinage de V de 0 et un réel positif M vérifiant

$$|f(x) - f_N(x)| \leq M|x^{m+n+1}|, \quad \forall x \in V. \quad (3.4)$$

Les approximations obtenues sont alors, au sens mathématique du terme, des *équivalents* de $f(x)$ au voisinage de $x = 0$.

Exemple Considérons la fonction tangente hyperbolique $\tanh$. Son développement en série de Taylor au voisinage de zéro s'écrit

$$f(x) = \tanh(x) = x - \frac{1}{3}x^3 + \frac{2}{15}x^5 - \frac{17}{315}x^7 + \frac{62}{2835}x^9 + \dots \quad (3.5)$$

Nous pouvons donc considérer les approximations successives

$$\begin{aligned} f_0(x) &= 0, \\ f_1(x) &= f_2(x) = x, \\ f_3(x) &= f_4(x) = x - \frac{1}{3}x^3, \\ f_5(x) &= f_6(x) = x - \frac{1}{3}x^3 + \frac{2}{15}x^5, \\ f_7(x) &= f_8(x) = x - \frac{1}{3}x^3 + \frac{2}{15}x^5 - \frac{17}{315}x^7, \\ f_9(x) &= f_{10}(x) = x - \frac{1}{3}x^3 + \frac{2}{15}x^5 - \frac{17}{315}x^7 + \frac{62}{2835}x^9. \end{aligned} \quad (3.6)$$

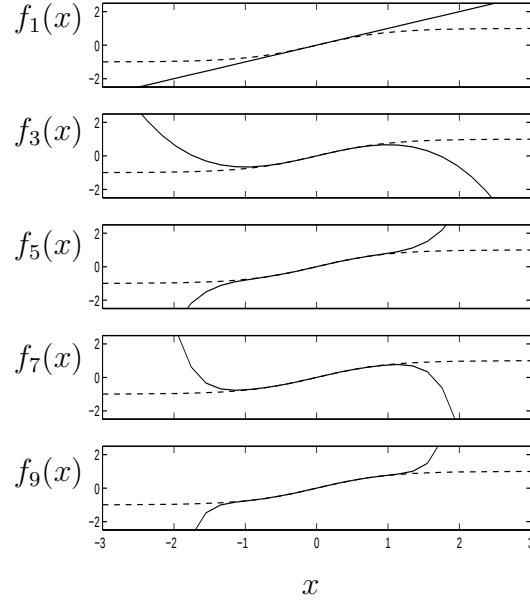


FIG. 3.1 – Approximations polynomiales de la tangente hyperbolique par développement de Taylor.

Ces approximations polynomiales sont tracées sur la figure 3.1.

D'un point de vue plus algébrique, nous pouvons associer le développement de Taylor à une décomposition sur la *base canonique* des polynômes

$$(1, x, x^2, \dots, x^n, \dots) \quad (3.7)$$

L'approximation (3.2) obtenue en tronquant à l'ordre N la série (3.1) est donc la projection de f sur le sous-espace vectoriel des polynômes de degrés inférieurs ou égaux à N . Mais le choix de la base canonique, dicté par la simplicité de son expression n'est pas forcément le plus judicieux dans le cadre de l'approximation polynomiale.

3.1.2 Polynômes orthogonaux

Le problème à résoudre est de quantifier l'erreur d'approximation induite par les termes négligés, c'est à dire les termes de degrés strictement supérieurs à N . Pour cela, il est nécessaire de pourvoir l'espace d'une norme, le plus souvent induite par un produit scalaire (voir annexe C).

$$\|\varphi\| = \sqrt{\langle \varphi, \varphi \rangle}. \quad (3.8)$$

Réduire l'erreur d'approximation au sens des moindres carrés revient donc à minimiser l'erreur quadratique

$$e_N = \|f - f_N\|^2. \quad (3.9)$$

3.1. L'approximation polynomiale

Considérons une base $(P_n)_{n \in \mathbb{N}}$ de *polynômes orthogonaux* pour le produit scalaire choisi. Si l'on suppose de plus que cette base est orthonormée, les coefficients de la décomposition de $f(x)$ sont appelés coefficients de Fourier et ils vérifient

$$f(x) = \sum_{n=0}^{\infty} c_n P_n(x), \quad c_n = \langle f, P_n \rangle. \quad (3.10)$$

On montre alors que l'approximation polynomiale

$$f_N(x) = \sum_{n=0}^N c_n P_n(x). \quad (3.11)$$

est celle qui minimise l'erreur quadratique [Wal94]. En utilisant la linéarité du produit scalaire, cette erreur est alors donnée par

$$\begin{aligned} e_N &= \left\| f - \sum_{n=0}^N c_n P_n \right\|^2 \\ &= \left\langle f - \sum_{n=0}^N c_n P_n, f - \sum_{n=0}^N c_n P_n \right\rangle \\ &= \langle f, f \rangle - 2 \sum_{n=0}^N c_n \langle f, P_n \rangle + \left\langle \sum_{n=0}^N c_n P_n, \sum_{n=0}^N c_n P_n \right\rangle \\ &= \langle f, f \rangle - 2 \sum_{n=0}^N c_n^2 + \sum_{n=0}^N c_n^2 \\ &= \langle f, f \rangle - \sum_{n=0}^N c_n^2. \end{aligned} \quad (3.12)$$

En effet, notons f'_N une autre approximation polynomiale d'ordre N de f . Comme $(P_n)_{n \leq N}$ est une base du sous-espace vectoriel des polynômes de degré inférieur ou égal à N , nous pouvons écrire

$$f'_N(x) = \sum_{n=0}^N c'_n P_n(x). \quad (3.13)$$

En utilisant l'expression (3.12) nous avons alors

$$\begin{aligned} e'_N &= \|f - f'_N\|^2 \\ &= \left\langle f - \sum_{n=0}^N c'_n P_n, f - \sum_{n=0}^N c'_n P_n \right\rangle \\ &= \langle f, f \rangle - 2 \sum_{n=0}^N c'_n \langle f, P_n \rangle + \sum_{n=0}^N c_n'^2 \\ &= \langle f, f \rangle - 2 \sum_{n=0}^N c'_n c_n + \sum_{n=0}^N c_n'^2 + \sum_{n=0}^N c_n^2 - \sum_{n=0}^N c_n^2 \\ &= \langle f, f \rangle + \sum_{n=0}^N (c'_n - c_n)^2 - \sum_{n=0}^N c_n^2 \\ &= e_N + \sum_{n=0}^N (c'_n - c_n)^2. \end{aligned} \quad (3.14)$$

Chapitre 3. De l'approximation polynomiale à l'approximation rationnelle

Comme la quantité $\sum_{n=0}^N (c'_n - c_n)^2$ est positive, nous avons nécessairement une moins bonne approximation au sens de la norme quadratique.

$$e_N \leq e'_N \quad (3.15)$$

Exemple Reprenons le cas de la tangente hyperbolique. Nous allons trouver une approximation polynomiale à l'aide de polynômes orthogonaux. Soit $\{P_n\}$ la base des polynômes de Legendre. Le calcul des coefficients de Fourier se fait à l'aide de la formule (3.10). Cependant les polynômes de Legendre n'étant pas normés, nous substituons

$$\frac{P_n}{\|P_n\|} \quad \text{à} \quad P_n. \quad (3.16)$$

Nous avons alors, avec la définition du produit scalaire correspondant aux polynômes de Legendre (voir en annexe la section C.3.1)

$$c_n = \left\langle \tanh, \frac{P_n}{\|P_n\|} \right\rangle \quad (3.17)$$

$$= \frac{1}{\|P_n\|} \int_{-1}^1 \tanh(x) P_n(x) dx \quad (3.18)$$

Nous avons donc formellement la série

$$\tanh(x) = \sum_{n=0}^{\infty} c_n \frac{P_n}{\|P_n(x)\|}, \quad (3.19)$$

ou encore, en notant

$$c'_n = \frac{c_n}{\|P_n\|} = \frac{1}{\|P_n\|^2} \int_{-1}^1 \tanh(x) P_n(x) dx = \frac{\int_{-1}^1 \tanh(x) P_n(x) dx}{\int_{-1}^1 P_n(x)^2 dx}, \quad (3.20)$$

$$\tanh(x) = \sum_{n=0}^{\infty} c'_n P_n. \quad (3.21)$$

Le calcul des c'_n nous donne alors les approximations successives (les polynômes P_n sont définis dans l'annexe C)

$$\begin{aligned} c'_0 &= 0, & c'_1 &= .8436022096 \\ c'_2 &= 0, & c'_3 &= .090725174 \\ c'_4 &= 0, & c'_5 &= .009587934 \\ c'_6 &= 0, & c'_7 &= -.000919 \\ c'_8 &= 0, & c'_9 &= .0000461. \end{aligned} \quad (3.22)$$

3.1. L'approximation polynomiale

$$\begin{aligned}
f_0(x) &= 0, \\
f_1(x) &= f_2(x) = c'_1 P_1(x), \\
f_3(x) &= f_4(x) = c'_1 P_1(x) + c'_3 P_3(x), \\
f_5(x) &= f_6(x) = c'_1 P_1(x) + c'_3 P_3(x) + c'_5 P_5(x), \\
&\vdots
\end{aligned} \tag{3.23}$$

Ces approximations polynomiales sont tracées sur la figure 3.2. On remarque que du fait de la définition du produit scalaire, l'approximation converge très rapidement sur l'intervalle $x \in [-1, 1]$.

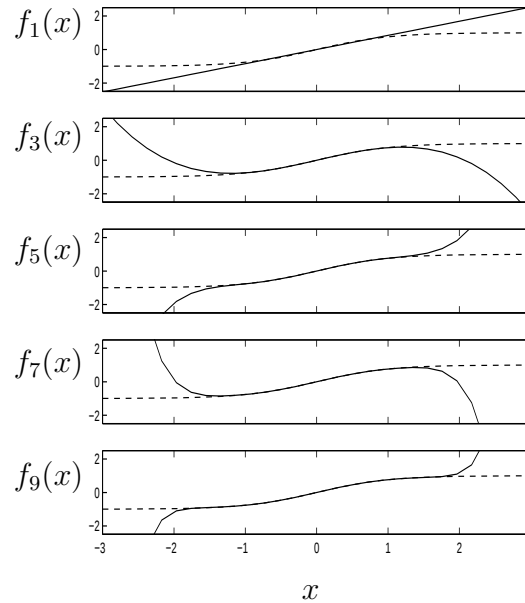


FIG. 3.2 – Approximations de la tangente hyperbolique par les polynômes orthogonaux de Legendre.

Le choix de la famille de polynômes orthogonaux (Hermite, Legendre, Tchebychev, etc.) et du produit scalaire associé se fait selon la nature du problème. Dans le cas de l'approximation d'une fonction de transfert, c'est bien souvent la facilité d'inversion de la transformation de Laplace qui dicte ce choix.

Évidemment, cette démarche peut aussi être utilisée avec d'autres familles de fonctions orthogonales. Par exemple, R. Morvan [Mor00] utilise les polynômes de Laguerre et de Kautz, ainsi que les fonctions dérivées, pour modéliser des systèmes d'ordre infini comme une ligne de transmission.

Malgré ses propriétés intéressantes de continuité et de dérivabilité, l'approximation polynomiale présente certains défauts qui en limitent les applications. Tout d'abord, quel que soit $N \geq 0$, $f_N(x)$ est non bornée quand x tend vers l'infini. Cela peut se révéler gênant pour certaines intégrations.

D'autre part, la série (3.1) ne converge pas nécessairement pour tout $x \in \mathbb{C}$. On définit alors un *rayon de convergence*, c'est-à-dire le plus grand réel positif R tel que si $|x| \leq R$, la série converge. Pour des valeurs en dehors du domaine de convergence ($|x| > R$), l'augmentation du nombre de termes retenus dans $f_N(x)$ détériorera alors la qualité de l'approximation !

Ce sont ces limitations qui nous amènent à nous intéresser à une nouvelle classe de fonctions d'approximation : les fractions rationnelles.

3.2 Les approximants de Padé

L'approximation de fonctions à l'aide de fractions rationnelles a été étudié de façon rigoureuse par le mathématicien français Henri Eugène Padé (1863-1953).



FIG. 3.3 – Henri Eugène Padé 1863-1953
Source de l'image : université de Saint Andrews
(<http://www-history.mcs.st-andrews.ac.uk>).

Padé se proposa de résoudre le problème de l'approximation sous la condition énoncée par l'équation 3.3. L'approximation polynomiale, déjà détaillée en 3.1, est une des solutions possibles, dont nous avons souligné les limitations. Mais Padé montra comment remplir cette condition avec une fonction rationnelle.

3.2.1 Définition des approximants

Voici comment Padé a défini ses fonctions d'approximation en 1892. Soit une fonction f admettant un développement en série entière au voisinage de $x = 0$. Étant donnés deux entiers $(m, n) \in \mathbb{N}^2$, nous pouvons trouver un couple de polynômes (P, Q) tels que P soit de degré inférieur ou égal à m et Q de degré inférieur ou égal à n , et vérifiant les relations

$$Q(x)f(x) - P(x) = O(x^{m+n+1}) \text{ et } Q(x) \not\equiv 0, \quad (3.24)$$

3.2. Les approximants de Padé

avec la notation de Landau définie par l'équation (3.4).

Le rapport de ces deux polynômes définit une unique fraction rationnelle [Gau75] appelée *approximant de Padé* d'ordre (n, m) et notée :

$$[n, m]_f(x) = \frac{P(x)}{Q(x)} \quad (3.25)$$

Les approximants de Padé sont habituellement présentés dans un tableau, disposés selon leurs ordres. Ce tableau est appelé *table de Padé* de la fonction $f : x \mapsto f(x)$.

$$\begin{array}{cccc} [0, 0]_f(x) & [0, 1]_f(x) & [0, 2]_f(x) & \cdots \\ [1, 0]_f(x) & [1, 1]_f(x) & [1, 2]_f(x) & \cdots \\ [2, 0]_f(x) & [2, 1]_f(x) & [2, 2]_f(x) & \cdots \\ \vdots & \vdots & \vdots & \ddots \end{array} \quad (3.26)$$

Les comportements asymptotiques des approximants de Padé présentent l'avantage de pouvoir être fixés par le choix des degrés n et m . De plus, une suite d'approximants (par exemple la suite $([n, n]_f(x))_{n=0,1,\dots}$) possède souvent¹ un rayon de convergence supérieur à la série entière générée par la même fonction.

Exemple d'approximants d'une fonction

Considérons la fonction exponentielle $f(x) = e^x$. La figure 3.4 présente les premiers approximants de Padé de cette fonction, ainsi que leurs allures au voisinage de l'origine. En utilisant le développement de Taylor de e^x , il est aisé de vérifier par exemple que l'approximant d'ordre $(1, 1)$ vérifie les relations de définition (3.24). Nous avons en effet :

$$[1, 1]_{exp}(x) = \frac{2+x}{2-x} = \frac{P(x)}{Q(x)} \quad (3.27)$$

$$e^x = 1 + x + \frac{1}{2}x^2 + O(x^3) \quad (3.28)$$

$$\begin{aligned} Q(x)f(x) - P(x) &= (2-x)(1+x+\frac{1}{2}x^2+O(x^3)) - 2-x \\ &= 2+x+O(x^3) - 2-x \\ &= O(x^3) \end{aligned} \quad (3.29)$$

Approximation d'un retard

Certains de ces approximants sont bien connus des automaticiens qui les utilisent pour modéliser les retards dans les fonctions de transferts. En effet, la transformée de Laplace d'un retard pur est donnée par :

$$f(t-T) \sqsupset e^{-Tp}F(p) \quad \text{avec} \quad f(t) \sqsupset F(p) \quad (3.30)$$

¹G. Baker [GABG70] remarque qu'il n'existe cependant pas encore de formulation générale des rayons de convergence des approximants de Padé.

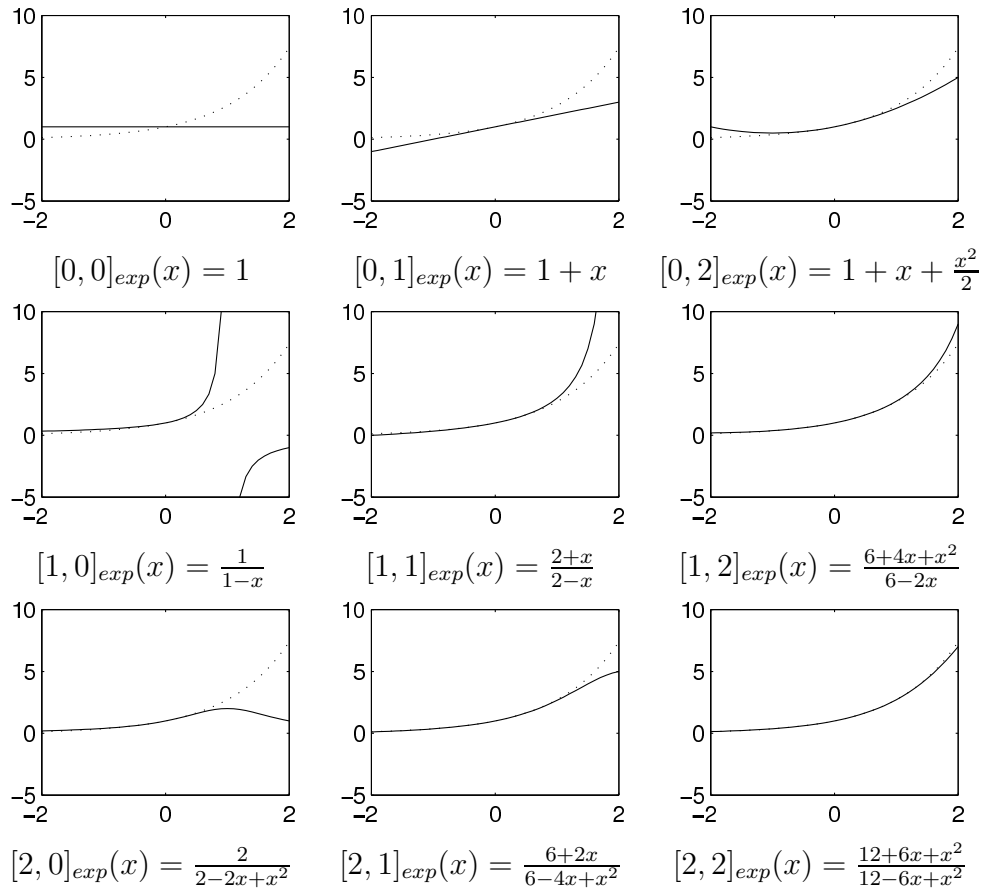


FIG. 3.4 – Les premiers approximants de Padé de la fonction exponentielle.

3.2. Les approximants de Padé

Une très bonne approximation en est donnée par les approximants diagonaux :

$$[1, 1]_{exp}(-Tp) = \frac{2 - Tp}{2 + Tp} \quad (3.31)$$

$$[2, 2]_{exp}(-Tp) = \frac{12 - 6Tp + T^2p^2}{12 + 6Tp + T^2p^2} \quad (3.32)$$

$$[3, 3]_{exp}(-Tp) = \frac{120 - 60Tp + 12T^2p^2 - T^3p^3}{120 + 60Tp + 12T^2p^2 + T^3p^3} \quad (3.33)$$

3.2.2 Propriétés des approximants

Premières propriétés

Nous pouvons d'ores et déjà présenter quelques propriétés immédiates des approximants de Padé.

1. Tout d'abord, les éléments de la première rangée de la table (3.26) présentent des dénominateurs de degré nul. Ils s'identifient donc aux sommes partielles du développement de $f(x)$ en série de Taylor (3.1).

$$\begin{aligned} [0, 0]_f(x) &= c_0, \\ [0, 1]_f(x) &= c_0 + c_1x, \\ [0, 2]_f(x) &= c_0 + c_1x + c_2x^2, \\ &\vdots \end{aligned} \quad (3.34)$$

Les approximants de Padé peuvent donc être considérés comme une extension de l'approximation polynomiale présentée en introduction et nous pouvons écrire

$$[0, N]_f(x) = f_N(x). \quad (3.35)$$

2. D'autre part, remarquons que la différence entre les degrés du numérateur et du dénominateur reste constante quand on parcourt la table (3.26) selon l'une de ses diagonales.
3. En revanche, si l'on parcourt la table de Padé de f selon une *antidiagonale* (par exemple la séquence $[0, 2]$ $[1, 1]$ $[2, 0]$) on obtient des approximants faisant intervenir un nombre constant de paramètres.
4. Enfin, les relations de définition (3.24) permettent de vérifier la relation suivante entre les approximants d'une fonction f et ceux de son inverse $\frac{1}{f}$:

$$[n, m]_f = \frac{1}{[m, n]_{\frac{1}{f}}} \quad (3.36)$$

L'ordre de contact d'un approximant

Considérons maintenant les développements en séries entières de $f(x)$ et de $[n, m]_f(x)$. Supposons qu'ils coïncident jusqu'à l'ordre r , c'est-à-dire que r est la valeur maximale pour laquelle est vérifiée la relation

$$f(x) - [n, m]_f(x) = O(x^{r+1}). \quad (3.37)$$

Nous pouvons dire que l'approximation est d'autant meilleure que r est grand et donc, nous caractériserons l'approximant par son *ordre de contact* r ([Gau75]). En ce qui concerne la fonction exponentielle, par exemple, l'ordre de contact des approximants $[n, m]_{exp}$ est strictement

$$r = n + m. \quad (3.38)$$

Dans ce cas, la table de Padé de la fonction est dite *normale*. La condition nécessaire et suffisante pour que la table de Padé soit normale s'exprime à partir des coefficients du développement en série entière de $f(x)$ (équation (3.1)) :

$$\Delta_{m,n} = \det \begin{vmatrix} c_{n-m} & c_{n-m+1} & \cdots & c_n \\ c_{n-m+1} & c_{n-m+2} & \cdots & c_{n+1} \\ \vdots & \vdots & \ddots & \vdots \\ c_n & c_{n+1} & \cdots & c_{n+m} \end{vmatrix} \neq 0 \quad \forall (m, n) \in \mathbb{N}^2 \quad (3.39)$$

Ces déterminants sont calculés avec la convention : $c_k = 0$ pour $k < 0$.

Si la table est *anormale*, c'est-à-dire s'il existe des approximants pour lesquels l'ordre de contact r est différent de $m + n$, certains approximants de f sont identiques. Ces derniers se répartissent sous forme de « blocs » carrés, de la forme $[n + i, m + j]$ avec $(i, j) \in \{0, 1, \dots, k\}^2$. Et ils ont pour ordre de contact $r = m + n + k$.

Réexaminons à la lumière de ces nouvelles définitions la fonction tangente hyperbolique qui nous a servi d'exemple précédemment. En comparant son développement à celui de l'un de ses approximants, nous concluons que cette fonction présente une table anormale.

$$\tanh(x) = x - \frac{1}{3}x^3 + \frac{2}{15}x^5 + O(x^7) \quad (3.40)$$

$$[2, 1]_{\tanh}(x) = \frac{3x}{3 + x^2} = x - \frac{1}{3}x^3 + \frac{1}{9}x^5 + O(x^7) \quad (3.41)$$

Par définition, cet approximant admet un ordre de contact $r = 4$, qui est supérieur à la somme de ces degrés $m + n = 3$. Explicitons les premiers éléments de la table pour

3.2. Les approximants de Padé

vérifier l'existence d'approximants identiques.

$\frac{0}{1}$	$\frac{x}{1} \quad \frac{x}{1}$	$\frac{3x - x^3}{3} \quad \frac{3x - x^3}{3}$	(3.42)
$\frac{0}{x}$	$\frac{x}{1} \quad \frac{x^2}{x}$	$\frac{3x - x^3}{3} \quad \frac{3x^2 - x^4}{3x}$	
$\frac{0}{x^2}$	$\frac{3x}{3 + x^2} \quad \frac{3x}{3 + x^2}$	$\frac{15x + x^3}{15 + 6x^2} \quad \frac{15x + x^3}{15 + 6x^2}$	
$\frac{0}{x^3}$	$\frac{3x}{3 + x^2} \quad \frac{3x^2}{3x + x^3}$	$\frac{15x + x^3}{15 + 6x^2} \quad \frac{15x^2 + x^4}{15x + 6x^3}$	

Des blocs d'approximants identiques ont été mis en évidence dans la table. Considérons le premier. Il s'agit de : $[0, 1]_{\tanh(x)}$, $[0, 2]_{\tanh(x)}$, $[1, 1]_{\tanh(x)}$ et $[1, 2]_{\tanh(x)}$. Ils possèdent tous le même ordre de contact $r = 2$.

- Pour le premier, cet ordre de contact se révèle être supérieur à la somme de ses degrés $m + n = 1 < r$.
- Pour les deux suivants, l'ordre de contact est exactement égal à la somme des degrés : $m + n = 2 = r$.
- Pour le dernier, en revanche, nous avons : $m + n = 3 > r$. Remarquons cependant que cet approximant vérifie bien la relation de définition (3.24). En effet, sous forme non réduite,

$$P(x) = x^2 \text{ et } Q(x) = x \quad (3.43)$$

Ce qui nous permet d'écrire :

$$\begin{aligned}
 Q(x) \tanh(x) - P(x) &= x(x + O(x^3)) - x^2 \\
 &= x^2 + O(x^4) - x^2 \\
 &= O(x^4)
 \end{aligned} \quad (3.44)$$

3.2.3 Calcul des approximants

Calcul direct

Le calcul d'un approximant peut se faire à l'aide de l'équation (3.24) et du développement (3.1) [GABG70]. On vérifie que la solution formelle suivante, exprimée comme le

rapport de deux déterminants d'ordre $n + 1$, satisfait au problème :

$$[n, m](x) = \frac{\det \begin{vmatrix} c_{m-n+1} & c_{m-n+2} & \cdots & c_{m+1} \\ \vdots & \vdots & \ddots & \vdots \\ c_m & c_{m+1} & \cdots & c_{m+n} \\ \sum_{k=n}^m c_{k-n} x^k & \sum_{k=n-1}^m c_{k-n+1} x^k & \cdots & \sum_{k=0}^m c_k x^k \end{vmatrix}}{\det \begin{vmatrix} c_{m-n+1} & c_{m-n+2} & \cdots & c_{m+1} \\ \vdots & \vdots & \ddots & \vdots \\ c_m & c_{m+1} & \cdots & c_{m+n} \\ x^n & x^{n-1} & \cdots & 1 \end{vmatrix}} \quad (3.45)$$

Pour que cette expression soit valable dans le cas général, nous adoptons les conventions suivantes :

- comme précédemment, $c_k = 0$ pour $k < 0$;
- si la valeur initiale d'un indice de sommation est strictement supérieure à sa valeur finale, la somme est remplacée par zéro.

Exemple

Nous allons détailler le calcul direct de l'approximant d'ordre $(2, 1)$ de la tangente hyperbolique. Commençons par expliciter les coefficients du développement de la fonction jusqu'à l'ordre $m + n = 3$.

$$\tanh(x) = x - \frac{1}{3}x^3 + O(x^5) \quad (3.46)$$

Les calculs des déterminants nous fournissent les polynômes :

$$P(x) = \det \begin{vmatrix} 0 & 1 & 0 \\ 1 & 0 & -\frac{1}{3} \\ 0 & 0 & x \end{vmatrix} = -x \quad (3.47)$$

$$Q(x) = \det \begin{vmatrix} 0 & 1 & 0 \\ 1 & 0 & -\frac{1}{3} \\ x^2 & x & 1 \end{vmatrix} = -1 - \frac{1}{3}x^2 \quad (3.48)$$

Nous retrouvons donc l'expression fournie précédemment pour $[2, 1]_{\tanh}(x)$.

$$\frac{P(x)}{Q(x)} = \frac{3x}{3 + x^2} \quad (3.49)$$

Algorithmes

Dans les cas ou une partie de la table de Padé doit être calculée, et non plus seulement un approximant particulier, des algorithmes itératifs sont plus adaptés que le calcul direct.

3.2. Les approximants de Padé

Le plus cité dans la littérature est l'algorithme *qd*, qui tient son nom de l'anglicisme *quotient difference*. Il consiste à construire récursivement une table de coefficients qui permet d'obtenir les approximants sous la forme d'une fraction continue. Cet algorithme pose cependant des problèmes de stabilité numérique, et ne fonctionne que pour des tables normales.

Des travaux récents concernant les polynômes orthogonaux ont permis de proposer différents nouveaux algorithmes de calcul. A. Draux et P. Van Ingelandt [DI87] proposent par exemple un logiciel en arithmétique rationnelle capable de parcourir une table de Padé anormale avec une très bonne stabilité numérique.

3.2.4 Approximants de Padé et polynôme orthogonaux

Il existe des liens étroits entre les polynômes orthogonaux et les approximants de Padé. Bien que cet aspect ne soit pas indispensable pour la compréhension des chapitres suivants, il nous est cependant apparu difficile de le passer totalement sous silence.

Reprenons l'expression (3.1) de la fonction f sous forme d'une série entière et définissons maintenant une forme linéaire à partir des coefficients c_n . Une forme linéaire c , qui associe un scalaire à tout polynôme (et par continuité à toute série entière, ou fonctions holomorphe en zéro [Ise88]) une fonction, est en effet parfaitement définie par ses moments, c'est-à-dire par la connaissance des $c(t^n)$ pour tout n^2 .

$$c(t^n) = c_n, \quad \forall n \in \mathbb{N}. \quad (3.50)$$

On montre alors facilement [BI94] que cette forme associe $f(x)$ à la fonction $g : t \rightarrow (1 - xt)^{-1}$. En effet, en utilisation la linéarité et les définitions de c et f ,

$$\begin{aligned} c(g(t)) &= c\left(\frac{1}{1 - xt}\right) \\ &= c(1 + xt + x^2t^2 + x^3t^3 + \dots) \\ &= c(1) + xc(t) + x^2c(t^2) + x^3c(t^3) + \dots \\ &= c_0 + c_1x + c_2x^2 + c_3x^3 + \dots = f(x). \end{aligned} \quad (3.51)$$

D'où la relation

$$f(x) = c(g(t)) = c\left(\frac{1}{1 - xt}\right). \quad (3.52)$$

Toujours dans le but de trouver une approximation de $f(x)$, il est alors intéressant de ne plus procéder directement, mais plutôt de chercher une approximation polynomiale de $g(t)$. Plus précisément, on définit un *polynôme d'interpolation de Hermite* de degré $k - 1$ qui vérifie la propriété

$$\frac{d^j R_k}{dt^j}(t_i) = \frac{d^j g}{dt^j}(t_i) \quad (3.53)$$

²Précisons pour éviter toute confusion que dans ce chapitre, t et x ne dénotent plus des variables de temps et d'espace.

pour t_i , point d'interpolation d'ordre k_i et $j = 0, \dots, k_i - 1$.

La continuité justifie alors qu'une approximation de $f(x)$ soit obtenue en prenant l'image du polynôme $R_k(t)$ par la forme c . On montre que cette image est une *fraction rationnelle* dont le numérateur est de degré $k - 1$, le dénominateur de degré k , et dont les pôles sont les points d'interpolation t_i . Cette approximation rationnelle est appelée approximant *de type Padé*.

Dans le cas général, l'erreur d'approximation entre cette fraction rationnelle et la fonction f est de l'ordre de x^k .

$$f(x) - c(R_k(t)) = O(x^k). \quad (3.54)$$

Néanmoins, on démontre que si les N points d'interpolation t_i sont choisis de telle façon que le polynôme

$$v_k(t) = (t - t_1)^{k_1} \dots (t - t_N)^{k_N} \quad (3.55)$$

vérifie la relation

$$c(t^i v_k(t)) = 0, \quad i = 1, \dots, k - 1, \quad (3.56)$$

l'erreur d'approximation est ramenée à

$$f(x) - c(R_k(t)) = O(x^{2k}). \quad (3.57)$$

Comme d'autre part, le degré du numérateur est égal à $k - 1$ et celui du dénominateur à k , la relation (3.57) correspond à la définition 3.24. Donc si (3.56) est vérifiée, l'approximant de type Padé $c(R_k(t))$ est donc un approximant de Padé.

$$c(R_k(t)) = [k, k - 1]_f(x) \quad (3.58)$$

Le schéma de la figure 3.5 résume ces relations.

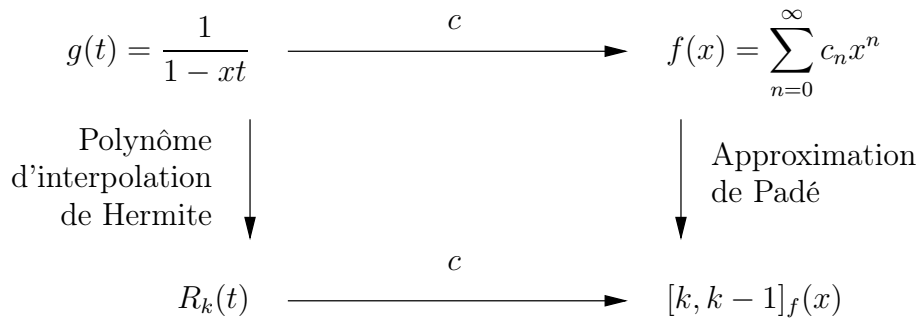


FIG. 3.5 – Relation entre approximants de Padé et polynômes orthogonaux.

Ce cas précis peut s'interpréter en remarquant que la relation (3.56) exprime que les polynômes $v_k(t)$ forment une famille *orthogonale* par rapport à la forme c (voir annexe C).

$$\langle v_n, v_m \rangle = c(v_n v_m) = 0 \quad \text{pour tout } n \neq m. \quad (3.59)$$

3.3. Les Fractions Continues

Ces liens étroits avec les polynômes orthogonaux forment la base de l'étude actuelle des propriétés algébriques des approximants de Padé. Ces propriétés, ainsi que les démonstrations qui ont été omises ici, sont détaillées dans [BI94] et [BI95]. Pour notre part, nous allons maintenant nous intéresser à une autre théorie mathématique, qui est peut-être encore plus liée aux approximants de Padé. Nous allons en effet introduire la théorie des fractions continues qui nous permettra d'utiliser de façon extrêmement intuitive les approximants de Padé pour résoudre des problèmes physiques.

3.3 Les Fractions Continues

Bien avant que ne furent élaborés les concepts d'irrationalité³ ou de transcendance⁴, les mathématiciens s'intéressèrent aux propriétés particulières de nombres tels que π , e , $\sqrt{2}$, etc. Ces nombres, bien que parfaitement définis, échappaient à une représentation « classique », ce qui rendait leur étude délicate. C'est donc tout naturellement qu'ils en cherchèrent des approximations sous forme rationnelle. Une méthode générale, connue sous le nom de *fraction continue* permet d'obtenir de telles approximations. Par la suite, cette même méthode fut justement utilisée afin de démontrer l'irrationalité de ces nombres.

3.3.1 Le développement d'un nombre

La théorie des nombres fut à l'origine de l'utilisation des fractions continues. Le développement du nombre π constitue donc un exemple de choix⁵ pour appréhender la démarche mise en œuvre dans cette méthode. C'est pourquoi nous le détaillerons dans un premier temps, puis nous verrons que l'approximation rationnelle d'une fonction d'ordre infini se déduit facilement du développement en fraction continue.

Commençons par isoler la partie entière de π . Nous pouvons écrire :

$$\pi = 3 + 0,1415926\dots = 3 + \frac{1}{\frac{1}{0,1415926\dots}} \quad (3.60)$$

L'inverse de la partie décimale étant supérieure à un, il possède une partie entière non nulle. Nous pouvons donc à nouveau isoler cette partie entière.

$$\frac{1}{0,1415926\dots} = 7,062513\dots \quad \pi = 3 + \frac{1}{7 + 0,062513\dots} \quad (3.61)$$

³Un nombre est dit rationnel s'il peut s'écrire sous la forme $\frac{a}{b}$ (fraction) où a et b sont des entiers relatifs.

⁴Un nombre est dit transcendant s'il ne peut être racine d'aucun polynôme à coefficients rationnels (dans le cas contraire, on parle de nombre algébrique). Exemple : e (la base des logarithmes népériens) et π sont des nombres transcendants, $\sqrt{2}$ et i (la racine carrée complexe de -1) sont des nombres algébriques. Nous devons la notion de nombres transcendants à Liouville (1844) et celle de nombres algébriques à Abel (1825).

⁵Cette partie est fortement inspirée du site Web de l'IREM consacré à l'histoire des mathématiques (<http://www.sciences-en-ligne.com/momo/chronomath/accueil.htm>).

Chapitre 3. De l'approximation polynomiale à l'approximation rationnelle

Cette forme de représentation nous permet déjà d'obtenir une première approximation du nombre π . Cette fraction porte le nom d'approximation d'Archimède, le célèbre mathématicien et physicien grec (287-212 av. J.-C.)

$$\pi \approx 3 + \frac{1}{7} = \frac{22}{7} = 3,142857\dots \quad (3.62)$$

Pour poursuivre le développement, il suffit à nouveau d'inverser la partie décimale du « résidu » puis d'en isoler la partie entière :

$$\frac{1}{0,062513\dots} = 15,99659\dots \quad \pi = 3 + \frac{1}{7 + \frac{1}{15 + 0,99659\dots}} \quad (3.63)$$

Nous obtenons alors une nouvelle approximation de π , dite approximation de Métius (géomètre et astronome hollandais, 1571-1635).

$$\pi \approx 3 + \frac{1}{7 + \frac{1}{15 + 1}} = \frac{355}{113} = 3,1415929\dots \quad (3.64)$$

Le développement peut se poursuivre ainsi indéfiniment. Les approximations successives, appelées *réduites*⁶, s'obtiennent en tronquant la fraction continue à un certain ordre n . Nous obtenons donc une suite de fractions rationnelles qui converge vers π . Donc, par passage à la limite :

$$\pi = 3 + \frac{1}{7 + \frac{1}{15 + \frac{1}{1 + \frac{1}{292 + \frac{1}{1 + \ddots}}}}} \quad (3.65)$$

De façon plus générale, partant d'un réel quelconque x , nous pouvons construire la suite de réduites $(r_n)_{n \in \mathbb{N}}$:

$$r_n = q_0 + \frac{1}{q_1 + \frac{1}{q_2 + \ddots + \frac{1}{q_n}}} \quad (3.66)$$

Il se conçoit assez facilement que si la suite des réduites est finie, alors x est rationnel. Inversement, Lambert (philosophe, physicien et astronome suisse, 1728-1777) montra qu'un développement infini est la preuve que x est irrationnel.

3.3. *Les Fractions Continues*



FIG. 3.6 – Johann Heinrich Lambert (1728-1777)
Source de l'image : université de Saint Andrews
(<http://www-history.mcs.st-andrews.ac.uk>).



FIG. 3.7 – Adrien-Marie Legendre (1752-1833)
Source de l'image : université de Saint Andrews
(<http://www-history.mcs.st-andrews.ac.uk>).

Nous savons par ailleurs, grâce à Legendre (mathématicien français, 1752-1834), que la convergence de la suite $(r_n)_{n \in \mathbb{N}}$ des réduites se déduit des valeurs des éléments de la suite $(q_n)_{n \in \mathbb{N}}$. Il prouva en effet que pour tout x irrationnel,

$$|x - r_n| < \frac{1}{2q_n^2} \quad (3.67)$$

3.3.2 Le développement d'une fonction

Nous allons adapter la méthode précédente au cas d'une fonction non rationnelle. Afin d'appliquer rapidement nos résultats à l'expression d'une impédance, nous choisissons de développer la fonction tangente hyperbolique qui apparaît dans l'expression de l'impédance d'une ligne en court-circuit.

$$\tanh : x \mapsto \frac{e^x - e^{-x}}{e^x + e^{-x}} \quad (3.68)$$

Cette fonction est holomorphe⁷ à l'intérieur du disque $D = \{x \in \mathbb{C} \mid |x| < \pi\}$. Elle admet donc en particulier un développement en série entière de puissances de x .

$$\tanh(x) = x - \frac{1}{3}x^3 + \frac{2}{15}x^5 + \dots \quad (3.69)$$

Le rôle joué précédemment par la partie entière du réel sera tenu par le terme du développement en x^{-1} . À l'instar du passage d'une expression décimale à une expression sous forme de fraction continue, nous pouvons écrire :

$$\tanh(x) = 0 + \frac{1}{\frac{1}{x - \frac{1}{3}x^3 + \frac{2}{15}x^5 + \dots}}} \quad (3.70)$$

L'inverse de la tangente hyperbolique est holomorphe à l'intérieur d'une couronne $D = \{x \in \mathbb{C} \mid \epsilon < |x| < 2\pi\}$. Aussi pouvons-nous le développer en série de Laurent de puissance de x (développement asymptotique) : L'origine $x = 0$ étant un pôle simple⁸ pour cette fonction, ce développement est limité au terme x^{-1} du côté des puissances négatives.

$$\frac{1}{x - \frac{1}{3}x^3 + \frac{2}{15}x^5 + \dots} = \frac{1}{x} + \frac{1}{3}x - \frac{1}{45}x^3 + \dots \quad (3.71)$$

⁶en anglais *convergent*s

⁷Une fonction f de la variable complexe est dérivable ou *holomorphe* dans un domaine D si f est holomorphe en tout point x_0 de D . C'est-à-dire que $\frac{f(x)-f(x_0)}{x-x_0}$ admet une limite lorsque $x \rightarrow x_0$

⁸Un pôle p est d'ordre $n \in \mathbb{N}^*$ pour la fonction f si $\lim_{x \rightarrow p} (x-p)^n f(x)$ existe et est non nul. Un pôle simple est un pôle d'ordre 1.

3.3. Les Fractions Continues

Nous reportons ce résultat dans l'expression précédente, nous isolons le terme en x^{-1} , puis nous développons l'inverse du « résidu » en série de Laurent :

$$\tanh(x) = 0 + \frac{1}{\frac{\frac{1}{x} + \frac{1}{\frac{1}{\frac{1}{\frac{1}{3}x - \frac{1}{45}x^3 + \dots}}}}}{\quad} \quad (3.72)$$

$$\frac{1}{\frac{1}{3}x - \frac{1}{45}x^3 + \dots} = \frac{3}{x} + \frac{1}{5}x - \frac{1}{175}x^3 + \dots \quad (3.73)$$

$$\tanh(x) = 0 + \frac{1}{\frac{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{1}{5}x - \frac{1}{175}x^3 + \dots}}}{\quad} \quad (3.74)$$

Finalement, nous pouvons réitérer les étapes précédentes pour obtenir le développement de la fonction sous forme de fraction continue. Ce résultat fut obtenu pour la première fois par Lambert et lui servit notamment à démontrer l'irrationalité du nombre π .

$$\tanh(x) = \frac{1}{\frac{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{1}{\frac{5}{x} + \frac{1}{\frac{7}{x} + \dots}}}}}{\quad} \quad (3.75)$$

De même que pour le développement de π , la suite des réduites $(r_n(x))_{n \in \mathbb{N}}$ nous donne des approximations successives de $\tanh(x)$ sous forme de fractions rationnelles.

$$r_1(x) = x \quad (3.76)$$

$$r_2(x) = \frac{1}{\frac{1}{x} + \frac{x}{3}} = \frac{3x}{x^2 + 3} \quad (3.77)$$

$$r_3(x) = \frac{1}{\frac{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{x}{5}}}{\quad}} = \frac{15x + x^3}{15 + 6x^2} \quad (3.78)$$

$$r_4(x) = \frac{1}{\frac{\frac{1}{x} + \frac{1}{\frac{\frac{3}{x} + \frac{1}{\frac{5}{x} + \frac{x}{7}}}{\quad}}}{\quad}} = \frac{105x + 10x^3}{105 + 45x^2 + x^4} \quad (3.79)$$

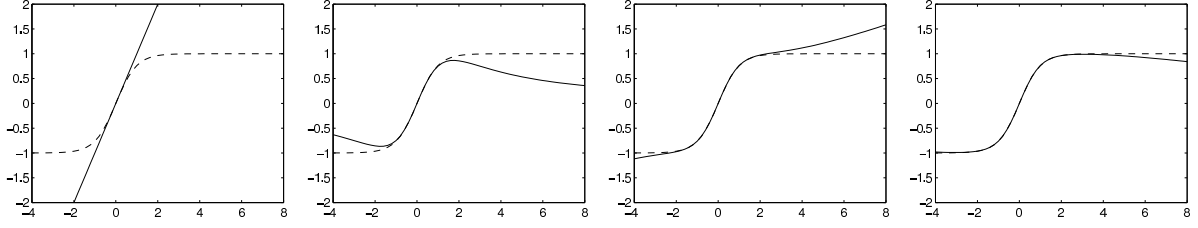


FIG. 3.8 – Les approximations rationnelles successives de la fonction tangente hyperbolique. De gauche à droite : $r_1(x)$, $r_2(x)$, $r_3(x)$ et $r_4(x)$.

Les différentes fonctions rationnelles ainsi obtenues sont représentées sur la figure 3.8. Il est particulièrement intéressant de remarquer que les approximations ainsi obtenues coïncident avec les *approximants de Padé* de la fonction tangente hyperbolique. Aussi convient-il de s'attarder quelque peu sur cette classe d'approximation rationnelle.

3.3.3 L'équivalence entre les fractions continues et les approximants de Padé

Nous venons de présenter deux façons d'obtenir des approximations rationnelles de fonctions. Ces méthodes sont radicalement différentes dans leur approche de la problématique. La première, dérivée de la théorie des nombres, utilise le développement d'une fonction sous la forme d'une fraction continue. Ce développement est ensuite tronqué à un ordre arbitraire, en fonction de la précision souhaitée. La seconde impose un critère d'optimalité (maximiser l'ordre de contact) pour calculer les coefficients de la fraction, une fois fixés les degrés du numérateur et du dénominateur.

Comme nous l'avons déjà signalé à la fin de la section 3.3 et vérifié sur l'exemple de calcul de $[2, 1]_{\tanh}$, il est remarquable que ces deux approches aboutissent à des résultats concomitants.

Les réduites de fractions continues sont des approximants de Padé

Cas général Dans le cas de fonctions présentant des tables de Padé normales, l'équivalence entre réduites de fractions continues et approximants de Padé est un résultat fort bien connu. Ainsi, soit f développable en série entière telle que les coefficients de son développement vérifient :

$$\Delta_{n,n} \neq 0 \text{ et } \Delta_{n,n+1} \neq 0 \quad \forall n \in \mathbb{N} \quad (3.80)$$

Ceci est une condition suffisante [Gau75] pour que $f(x)$ admette une représentation

3.3. Les Fractions Continues

sous forme de fraction continue *simple* :

$$f(x) = \frac{s_0}{1 - \frac{s_1 x}{1 - \frac{s_2 x}{\ddots}}} \quad (3.81)$$

Les réduites successives obtenues à partir de cette fraction continue sont alors les approximants de Padé de $f(x)$ répartis au voisinage de la diagonale de la manière suivante :

$$\begin{array}{cc} [0, 0]_f(x) & [0, 1]_f(x) \\ [1, 1]_f(x) & [1, 2]_f(x) \\ [2, 2]_f(x) & [2, 3]_f(x) \\ & \ddots \end{array} \quad (3.82)$$

Cas des fonctions impaires La suite de l'étude se limitera à l'approximation de certaines fonctions, dont la tangente hyperbolique, qui ont toutes la particularité d'être impaires,

$$f(-x) = -f(x). \quad (3.83)$$

Nous savons que les développements de telles fonctions ne comportent pas de termes d'ordre pair. Elles ont donc toutes des tables de Padé anormales. Pour s'en convaincre, il suffit de vérifier que (*cf.* équation (3.39)), soit

$$\Delta_{0,0} = c_0 = 0. \quad (3.84)$$

Nous supposons de plus que l'origine $x = 0$ est un zéro ou un pôle *simple* pour la fonction f . Le développement que nous utiliserons pour ces fonctions diffère quelque peu des développements classiques. C'est celui qui a été présenté dans la section 3.3.

$$\frac{s_0}{x} + \frac{1}{\frac{s_1}{x} + \frac{1}{\frac{s_2}{x} + \frac{1}{\frac{s_3}{x} + \ddots}}} \quad (3.85)$$

Les approximants de Padé correspondant aux réduites successives dépendent alors de la valeurs de s_0 . En effet deux cas sont à distinguer :

- si $x = 0$ est un zéro simple de la fonction f , $s_0 = 0$;
- si $x = 0$ est un pôle simple de la fonction f , $s_0 \neq 0$.

Plaçons-nous dans le premier cas et gardons à l'esprit que nous avons à faire à une table de Padé anormale : nous retrouvons des « blocs » d'approximants égaux.

$$\frac{1}{\frac{s_1}{x}} = [0, 1]_f = [0, 2]_f = [1, 1]_f = [1, 2]_f \quad (3.86)$$

$$\frac{1}{\frac{s_1}{x} + \frac{1}{\frac{s_2}{x}}} = [2, 1]_f = [2, 2]_f = [3, 1]_f = [3, 2]_f \quad (3.87)$$

$$\frac{1}{\frac{s_1}{x} + \frac{1}{\frac{s_2}{x} + \frac{1}{\frac{s_3}{x}}}} = [2, 3]_f = [2, 4]_f = [3, 3]_f = [3, 4]_f \quad (3.88)$$

$$\frac{1}{\frac{s_1}{x} + \frac{1}{\frac{s_2}{x} + \frac{1}{\frac{s_3}{x} + \frac{1}{\frac{s_4}{x}}}}} = [4, 3]_f = [4, 4]_f = [5, 3]_f = [5, 4]_f \quad (3.89)$$

$$\begin{array}{|c|c|} \hline [0, 1]_f & [0, 2]_f \\ \hline [1, 1]_f & [1, 2]_f \\ \hline [2, 1]_f & [2, 2]_f \\ \hline [3, 1]_f & [3, 2]_f \\ \hline \end{array} \quad \begin{array}{|c|c|} \hline [2, 3]_f & [2, 4]_f \\ \hline [3, 3]_f & [3, 4]_f \\ \hline [4, 3]_f & [4, 4]_f \\ \hline [5, 3]_f & [5, 4]_f \\ \hline \end{array} \quad \begin{array}{|c|c|} \hline [4, 5]_f & [4, 6]_f \\ \hline [5, 5]_f & [5, 6]_f \\ \hline \end{array} \quad (3.90)$$

Le cas où l'origine est un pôle simple ($s_0 \neq 0$) se déduit simplement du précédent. Il suffit pour cela d'utiliser le développement en fraction continue de $\frac{1}{f}$. Nous faisons ensuite appel à la propriété des approximants de Padé de l'inverse d'une fonction (équation (3.36)).

$$\frac{1}{f} = \frac{1}{\frac{s_0}{x} + \frac{1}{\frac{s_1}{x} + \frac{1}{\frac{s_2}{x} + \frac{1}{\frac{s_3}{x} + \dots}}}} \quad (3.91)$$

3.3. Les Fractions Continues

$$\frac{s_0}{x} = [0, 1]_{\frac{1}{f}} = [0, 2]_{\frac{1}{f}} = [1, 1]_{\frac{1}{f}} = [1, 2]_{\frac{1}{f}} \quad (3.92)$$

$$= [1, 0]_f = [2, 0]_f = [1, 1]_f = [2, 1]_f \quad (3.93)$$

$$\frac{s_0}{x} + \frac{1}{\frac{s_1}{x}} = [2, 1]_{\frac{1}{f}} = [2, 2]_{\frac{1}{f}} = [3, 1]_{\frac{1}{f}} = [3, 2]_{\frac{1}{f}} \quad (3.94)$$

$$= [1, 2]_f = [2, 2]_f = [1, 3]_f = [2, 3]_f \quad (3.95)$$

$$\frac{s_0}{x} + \frac{1}{\frac{s_1}{x} + \frac{1}{\frac{s_2}{x}}} = [2, 3]_{\frac{1}{f}} = [2, 4]_{\frac{1}{f}} = [3, 3]_{\frac{1}{f}} = [3, 4]_{\frac{1}{f}} \quad (3.96)$$

$$= [3, 2]_f = [4, 2]_f = [3, 3]_f = [4, 3]_f \quad (3.97)$$

Et ainsi de suite.

Ces développements particuliers présentent des propriétés de convergence remarquables et des expressions simples, ce qui en fait tout l'intérêt pour l'approximation de fonctions. Les liens entre les fonctions *usuelles* telles les fonctions trigonométriques, exponentielles, logarithmiques, de Bessel, etc. et les fractions continues sont établies par les propriétés de *confluence* des fonctions *hypergéométriques* [Wal48, JT80] et ont notamment été mis à profit en modélisation par [Sab98].

Les approximants de Padé sous forme de fractions continues

Nous venons de voir que les réduites des fractions continues sont des approximants de Padé particuliers, et plus précisément des termes proches de la diagonale de la table de Padé. Il est aisée de montrer que, réciproquement, chaque approximant de Padé d'une fonction peut se mettre sous la forme d'une réduite de fraction continue. En effet toute fraction rationnelle est développable sous la forme d'une fraction continue finie.

L'existence du développement non classique en $\frac{1}{x}$ est plus délicate à prouver. Contentons-nous de le vérifier sur un exemple, tout en signalant que si la fonction f est impaire et admet l'origine comme zéro (respectivement pôle) simple, ses approximants de Padé sont eux aussi des fonctions impaires de la variable x et admettent $x = 0$ comme zéro (respectivement pôle) simple.

À titre d'exemple nous réutilisons les approximants de la tangente hyperbolique.

$$[0, 1]_{\tanh}(x) = x = \frac{1}{\frac{1}{x}} \quad (3.98)$$

$$[0, 3]_{\tanh}(x) = \frac{3x - x^3}{3} = \frac{1}{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{1}{-\frac{1}{x}}}} \quad (3.99)$$

$$[0, 5]_{\tanh}(x) = \frac{2x^5 - 5x^3 + 15x}{15} = \frac{1}{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{1}{\frac{5}{x} + \frac{1}{-\frac{1}{12x} + \frac{1}{-\frac{6}{x}}}}}} \quad (3.100)$$

$$[2, 1]_{\tanh}(x) = \frac{3x}{3 + x^2} = \frac{1}{\frac{1}{x} + \frac{1}{\frac{3}{x}}}} \quad (3.101)$$

$$[2, 3]_{\tanh}(x) = \frac{x^3 + 15x}{15 + x^2} = \frac{1}{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{1}{\frac{5}{x}}}} \quad (3.102)$$

$$[2, 5]_{\tanh}(x) = \frac{-x^5 + 45x^3 + 630x}{630 + 255x^2} = \frac{1}{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{1}{\frac{5}{x} + \frac{1}{\frac{7}{x} + \frac{1}{-\frac{6}{x}}}}}} \quad (3.103)$$

$$[4, 1]_{\tanh}(x) = \frac{45x}{45 + 15x^2 + x^4} = \frac{1}{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{1}{\frac{5}{x} + \frac{1}{\frac{3}{x}}}}}} \quad (3.104)$$

3.4. Théorème de Mittag-Leffler

$$[4, 3]_{\tanh}(x) = \frac{10x^3 + 105x}{105 + 45x^2 + x^4} = \frac{1}{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{1}{\frac{5}{x} + \frac{1}{\frac{7}{x}}}}} \quad (3.105)$$

$$[4, 5]_{\tanh}(x) = \frac{x^5 + 105x^3 + 945x}{945 + 420x^2 + 15x^4} = \frac{1}{\frac{1}{x} + \frac{1}{\frac{3}{x} + \frac{1}{\frac{5}{x} + \frac{1}{\frac{7}{x} + \frac{1}{\frac{9}{x}}}}}} \quad (3.106)$$

Remarquons que les approximants de Padé éloignés de la diagonale sont obtenus en modifiant les derniers termes du développement « naturel » de la fonction en fraction continue. Finalement, nous avons une totale équivalence entre les approximations rationnelles issues des fractions continues et les approximants de Padé.

3.4 Théorème de Mittag-Leffler

Nous avons présenté la méthode la plus courante pour obtenir des approximations rationnelles de fonctions. Mais, à la différence d'une approximation polynomiale, ces fonctions approchées présentent en général des pôles dans le plan complexe. L'utilisateur d'une approximation rationnelle peut donc, légitimement, se poser la question de la localisation de ces pôles. Sont-ils confondus avec ceux de la fonction à approcher ?

D'autre part, peut-il assurer que les comportements de la fonction à approcher et de son approximation rationnelle sont suffisamment proches au voisinage de ces pôles ? Dans le raisonnement qui nous a conduit à construire les approximants de Padé, rien ne garantit une telle propriété. Elle est même fausse en général. Aussi allons-nous examiner la possibilité de construire un autre type d'approximation rationnelle en prenant pour contrainte l'équivalence de la fonction et de son approximation aux voisinages des pôles.

Le théorème de Mittag-Leffler peut être considéré comme une généralisation aux fonctions méromorphes⁹ de la décomposition en éléments simples des fractions rationnelles. Toute fraction rationnelle complexe est en effet décomposable en une somme d'*éléments simples de première espèce*. En particulier, si une fraction rationnelle $\frac{P}{Q}$ ne possède que

⁹Une fonction f de la variable complexe est méromorphe dans un ouvert connexe U de $\mathbb{C}$ si elle est définie et analytique dans un ouvert U' obtenu en enlevant à U un ensemble de points isolés, dont chacun est un pôle pour f . [BG83]



FIG. 3.9 – Magnus Gösta Mittag-Leffler (1846-1927)
Source de l'image : université de Saint Andrews
(<http://www-history.mcs.st-andrews.ac.uk>).

des pôles simples p_n , $1 \leq n \leq N$,

$$\frac{P(x)}{Q(x)} = A(x) + \sum_{n=1}^N \frac{a_n}{(x - p_n)} \quad (3.107)$$

où A est un polynôme et où a_n est le résidu correspondant au pôle p_n de la fraction.

Considérons le cas d'une fonction f méromorphe, qui n'a que des pôles simples p_n , que nous supposons classés par ordre croissant de leurs modules. S'il existe une suite de cercles C_n de rayons R_n dans le plan complexe, ne passant par aucun pôle de f et sur lesquels $|f(x)| < M$, M étant indépendant de n , et tels que $R_n \rightarrow \infty$ quand $n \rightarrow \infty$, alors le théorème de Mittag-Leffler fournit la décomposition en série suivante [Spi73] :

$$f(x) = f(0) + \sum_{n=1}^{\infty} a_n \left(\frac{1}{x - p_n} + \frac{1}{p_n} \right) \quad (3.108)$$

En particulier, la fonction tangente hyperbolique qui admet une infinité de pôles simples sur l'axe imaginaire remplit les conditions d'application du théorème.

3.4.1 Développement suivant les pôles d'une fonction

Les pôles de la tangente hyperbolique sont les

$$p_n = i\left(n\pi - \frac{\pi}{2}\right), \quad a_n = 1, \quad n \in \mathbb{Z}. \quad (3.109)$$

3.5. Méthodes d'approximation retenues

Ces pôles étant opposés deux à deux, après simplifications, nous obtenons une série d'éléments simples (et non plus une somme comme dans le cas d'une fraction).

$$\tanh(x) = \sum_{n=-\infty}^{\infty} \frac{1}{x - i(n\pi - \frac{\pi}{2})} \quad (3.110)$$

Cette série peut être écrite, en rapprochant les pôles conjugués :

$$\tanh(x) = \sum_{n=1}^{\infty} \frac{2x}{x^2 + (n\pi - \frac{\pi}{2})^2} \quad (3.111)$$

3.4.2 Développement suivant les zéros d'une fonction

De la même façon, l'inverse de la tangente hyperbolique peut s'écrire :

$$\coth(x) = \frac{1}{x} + \sum_{n=1}^{\infty} \frac{2x}{x^2 + (n\pi)^2} \quad (3.112)$$

Deux autres décompositions intéressantes concernent le rapport de fonctions de Bessel. Nous considérons les *fonctions de Bessel modifiées de première espèce* d'ordre n , I_n [SL00].

$$\frac{I_0(x)}{I_1(x)} = \sum_{n=1}^{\infty} \frac{2x}{x^2 + \lambda_{1,n}^2}, \quad (3.113)$$

$$\frac{I_1(x)}{I_0(x)} = \sum_{n=1}^{\infty} \frac{2x}{x^2 + \lambda_{0,n}^2}, \quad (3.114)$$

les $\lambda_{k,n} \in \mathbb{R}^+$, étant les zéros des *fonctions de Bessel de première espèce*, J_k .

3.5 Méthodes d'approximation retenues

Au cours de cette brève présentation des méthodes d'approximations, nous avons pu nous rendre compte des limites présentées par les méthodes d'approximations polynomiales pourtant largement répandues. Les approximants de Padé ont l'avantage d'une convergence bien plus rapide, spécialement pour les fonctions *usuelles* telles que celles rencontrées dans les expressions des impédances d'ordre infini présentées au chapitre 2 (fonctions de trigonométrie hyperbolique principalement).

Les approximations basées sur le théorème de Mittag-Leffler présentent une autre approche intéressante de par le contrôle des pôles (ou des zéros) de la fonction d'approximation. Nous retiendrons donc pour la suite ces deux types d'approximations rationnelles, et nous allons maintenant nous intéresser à leurs représentations physiques sous la forme de réseaux de Kirchhoff.

Chapitre 4

Synthèse de réseaux par développements d'impédances

Dans ce chapitre, nous nous proposons de donner une interprétation physique aux approximations rationnelles obtenues par les approximants de Padé et le théorème de Mittag-Leffler. Puisque nous avons représenté les systèmes d'ordre infini à l'aide d'impédances et d'admittances, la construction d'un système d'ordre réduit s'identifiera à la synthèse d'un réseau de Kirchhoff.

En particulier, le développement en fraction continue s'identifie à la *synthèse de Cauer* [Cau26], et le développement en éléments simples à la *synthèse de Foster* [Fos24b, Fos24a]. Ces méthodes ont connu un grand succès pour la synthèse de filtres. Elles permettent en effet de synthétiser une fonction de transfert donnée en électronique analogique [Fel86], ce qui a été démontré par [Bru31] pour les fonctions réelles positives.

Cependant notre application des synthèses de Foster et Cauer se distinguera par deux aspects importants. Tout d'abord, nous modélisons des impédances d'ordre infini, dans le but d'obtenir des approximations. La synthèse de filtres analogiques vise à obtenir de façon exacte une fonction rationnelle (d'ordre fini)¹.

D'autre part, les développements se feront selon une variable d'espace (x), alors qu'elles se font classiquement selon la variable fréquentielle (p). Nous verrons que ceci permet de conserver une simplicité remarquable pour l'expression des éléments constituant les réseaux.

4.1 Extension des synthèses de Foster et Cauer à l'ordre infini

Nous nous proposons d'exposer plusieurs méthodes de synthèse par des réseaux de Kirchhoff de modèles approchés de systèmes d'ordre infini. Ces méthodes relèvent d'une

¹Il semble toutefois que Cauer (fig. 4.1) ait étudié les ordres infinis dans son travail initial de thèse, malheureusement perdu aujourd'hui [CM95]

4.1. Extension des synthèses de Foster et Cauer à l'ordre infini



FIG. 4.1 – Wilhelm Cauer (1900-1945).

Source de l'image : université technique de Berlin (<http://www.tu-berlin.de>).

démarche commune, à savoir :

1. modélisation du système par une impédance d'ordre infini ;
2. développement de cette impédance (sous forme de fraction continue ou bien à l'aide du théorème de Mittag-Leffler) ;
3. modélisation du système par un réseau de Kirchhoff infini (sans approximation) ;
4. troncature du réseau infini pour obtenir un modèle d'ordre réduit (approximation rationnelle).

Chaque modèle approché sous forme d'un réseau tronqué présente une impédance qui est une approximation rationnelle particulière de l'impédance d'ordre infini. Ces différentes approximations² donneront ensuite lieu à une comparaison dans le domaine fréquentiel.

4.1.1 Développement en éléments simples de l'impédance

La première méthode utilise le théorème de Mittag-Leffler pour développer l'impédance d'ordre infini en série. Reprenons l'exemple de l'impédance d'une ligne de transmission en court-circuit (figure 2.14) qui peut s'écrire

$$Z = \sqrt{\frac{Z}{Y}} \tanh(\sqrt{ZY}) \quad (4.1)$$

en notant

$$Z = zx = R + Lp, \quad Y = yx = G + Cp. \quad (4.2)$$

²Remarquons que M. Bayard [Bay54] présente ces méthodes dans le cas particulier de la synthèse de « lignes artificielles ».

En utilisant le développement en série de la fonction tangente hyperbolique (3.111), nous pouvons écrire cette impédance sous la forme

$$\mathcal{Z} = \sqrt{\frac{z}{y}} \sum_{n=1}^{\infty} \frac{2\sqrt{zy}x}{zyx^2 + (n\pi - \frac{\pi}{2})^2} \quad (4.3)$$

Cette dernière expression se simplifiant :

$$\mathcal{Z} = \sum_{n=1}^{\infty} \frac{1}{\frac{Y}{2} + \frac{(n\pi - \frac{\pi}{2})^2}{2Z}}. \quad (4.4)$$

Cette forme est immédiatement identifiable à un réseau de Kirchhoff infini, formé de cellules en série, chaque cellule étant elle-même formée d'une admittance et d'une impédance en parallèle. En ne conservant que les $2N$ premiers termes de la série, c'est-à-dire les $2N$ pôles dominants, nous obtenons l'impédance équivalente du réseau tronqué de la figure 4.3, les valeurs des éléments étant :

$$Z_n = \frac{2Z}{(n\pi - \frac{\pi}{2})^2} \quad Y_n = \frac{Y}{2} \quad (4.5)$$

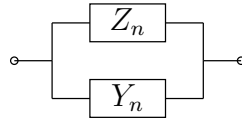


FIG. 4.2 – Cellule parallèle.

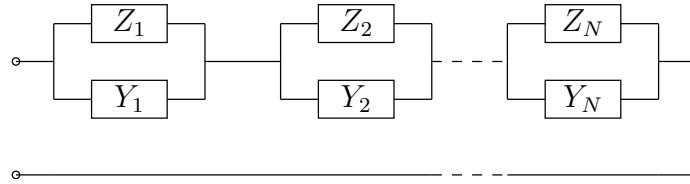


FIG. 4.3 – Modèle d'ordre réduit obtenu par développement en série d'une impédance d'ordre infini.

Chaque cellule peut être vue comme un *circuit bouchon* qui peut être le lieu d'une *résonance parallèle*. Les cellules étant mises en série, il suffit alors que l'impédance de l'une d'elle tende vers l'infini, ou prenne une valeur très grande³, pour que l'impédance globale du réseau tende elle aussi vers l'infini, ou prenne elle aussi une valeur très grande.

Dans le cas particulier où $Z = lpx$ et $Y = cpx$ (ligne sans pertes), ce développement s'identifie à la *première synthèse canonique de Foster* d'une impédance [Fel86]. Cependant,

³L'impédance ne tend vers l'infini ou vers zéro en régime harmonique que dans le cas sans pertes, avec soit $r = 0$, soit $g = 0$, voir fig. 2.21.

4.1. Extension des synthèses de Foster et Cauer à l'ordre infini

nous l'appliquons à une fonction non rationnelle. La convergence n'est donc pas garantie par un nombre fini de pôles, mais par le théorème de Mittag-Leffler ⁴.

4.1.2 Développement en éléments simples de l'admittance

Si nous développons cette fois-ci l'inverse de l'impédance à l'aide de l'équation (3.112) de la section 3.4, nous obtenons :

$$\frac{1}{\mathcal{Z}} = \sqrt{\frac{Y}{Z}} \coth(\sqrt{ZY}) = \frac{1}{Z} + \sum_{n=1}^{\infty} \frac{1}{\frac{Z}{2} + \frac{(n\pi)^2}{2Y}} \quad (4.6)$$

Cette expression peut être interprétée comme un réseau infini de cellules en parallèle, chaque cellule étant composée d'une impédance et d'une admittance mises en série. Les $2N$ premiers termes correspondent maintenant aux zéros dominants de l'impédance $\mathcal{Z}$ qui sont conservés dans le modèle d'ordre réduit de la figure 4.5. Les éléments composant le réseau étant :

$$Z_1 = Z, \quad Y_1 \equiv \infty, \quad Z_n = \frac{Z}{2}, \quad Y_n = \frac{2Y}{((n-1)\pi)^2}, \text{ pour } n \geq 2. \quad (4.7)$$

De façon duale par rapport au développement en série d'éléments simples de l'impédance, nous avons ici des cellules pouvant présenter des *résonances parallèles*. Si l'impédance de l'une d'entre elles s'annule (ou devient faible, avec la même réserve que précédemment), elle annule (ou affaiblit) alors l'impédance globale du réseau car les cellules sont disposées en parallèle.

De la même façon, remarquons que nous aboutissons à une topologie duale, semblable à la *deuxième synthèse canonique de Foster* [Fel86] d'une impédance rationnelle.

4.1.3 Développement en fraction continue de l'impédance

Reprenons maintenant le développement en fraction continue de la fonction tangente hyperbolique qui nous a servi d'exemple dans la section 3.3. À l'aide de ce résultat, nous pouvons écrire l'impédance du modèle

$$\mathcal{Z} = \sqrt{\frac{z}{y}} \cfrac{1}{\cfrac{1}{\sqrt{zyx}} + \cfrac{1}{\cfrac{3}{\sqrt{zyx}} + \cfrac{1}{\cfrac{5}{\sqrt{zyx}} + \cfrac{1}{\cfrac{7}{\sqrt{zyx}} + \ddots}}}}} \quad (4.8)$$

⁴En effet, dans le cas de la synthèse d'une fonction rationnelle, la décomposition en éléments simples aboutit à une somme finie et non à une série. La limite est donc atteinte dès que le nombre de cellules N égale le nombre de pôles de la fonction.



FIG. 4.4 – Cellule série.

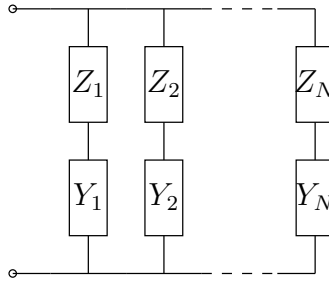


FIG. 4.5 – Modèle d'ordre réduit obtenu par développement en série d'une admittance d'ordre infini.

Remarquons que le développement se fait suivant la variable x . En utilisant les règles de calcul sur les fractions continues, nous simplifions cette expression :

$$\mathcal{Z} = \frac{1}{\sqrt{\frac{y}{z}} \frac{1}{\sqrt{zyx}} + \frac{1}{\sqrt{\frac{z}{y}} \frac{3}{\sqrt{zyx}} + \frac{1}{\sqrt{\frac{y}{z}} \frac{5}{\sqrt{zyx}} + \frac{1}{\sqrt{\frac{z}{y}} \frac{7}{\sqrt{zyx}} + \dots}}} = \frac{1}{\frac{1}{Z} + \frac{1}{\frac{3}{Y} + \frac{1}{\frac{5}{Z} + \frac{1}{\frac{7}{Y} + \dots}}}} \quad (4.9)$$

Considérons maintenant le réseau de la figure 4.7. Il est aisé de vérifier que son impédance équivalente s'exprime par la fraction continue :

$$\frac{1}{Y_1} + \frac{1}{\frac{1}{Z_1} + \frac{1}{\frac{1}{Y_2} + \dots + \frac{1}{Z_N}}} \quad (4.10)$$

4.1. Extension des synthèses de Foster et Cauer à l'ordre infini

Nous obtenons donc le modèle approché par simple identification [Fry29] :

$$Y_1 \equiv \infty, \quad Z_n = \frac{Z}{4n-3}, \quad Y_n = \frac{Y}{4n-5}, \text{ pour } n \geq 2. \quad (4.11)$$

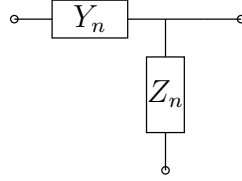


FIG. 4.6 – Cellule en échelle.

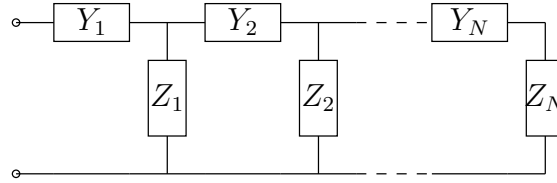


FIG. 4.7 – Modèle d'ordre réduit obtenu par développement d'une impédance d'ordre infini en fraction continue.

Les résonances et anti-résonances des cellules du réseau en échelle sont plus difficiles à appréhender. En effet, chaque nouvelle cellule vient modifier les fréquences de résonance des précédentes. Nous savons cependant qu'elles ne sont pas exactement confondues avec les zéros et les pôles de l'impédance $\mathcal{Z}$.

D'autre part, le cas particulier $Z = lpx$, $Y = cpx$ s'identifie cette fois-ci à la *première synthèse canonique de Cauer* [Fel86] appliquée à une impédance d'ordre infini.

Remarque Précisons que le développement de l'admittance du système en fraction continue aboutit à la synthèse du même réseau en échelle. Il suffit en effet de remarquer que :

$$\frac{1}{\mathcal{Z}} = \frac{1}{Z} + \frac{1}{\frac{3}{Y} + \frac{1}{\frac{5}{Z} + \frac{1}{\frac{7}{Y} + \dots}}} \quad (4.12)$$

Si nous tronquons cette fraction continue, et que nous l'identifions à l'admittance d'entrée du réseau de la figure 4.7, nous retrouvons les mêmes valeurs pour les éléments.

4.1.4 Comparaison des réponses fréquentielles

Trois modèles d'ordres réduits ont été obtenus, en se basant sur différents développements de fonctions. Ces modèles ont en commun de pouvoir être représentés par des réseaux de Kirchhoff. Ils peuvent donc être facilement comparés, à nombre équivalent d'« étages » que nous choisirons successivement égal à trois puis à cinq. La comparaison se fait en étudiant les variations du module et de la phase des impédances en régime harmonique, à la fréquence ν :

$$p = 2i\pi\nu \quad (4.13)$$

Les constantes physiques du système étant prises unitaires ($R = L = C = G = 1$), ce qui correspond au cas $\alpha = 1$ vu en 2.3.1, nous obtenons les courbes de la figure 4.8.

Nous pouvons conclure de cette première comparaison que :

- le modèle obtenu à l'aide du théorème de Mittag-Leffler appliqué à l'impédance reflète mieux les maxima de $\mathcal{Z}$ puisque ses pôles sont parfaitement *placés* dans le plan complexe ;
- celui obtenu par application du même théorème à l'admittance reflète mieux les minima de $\mathcal{Z}$ puisque ses zéros sont parfaitement placés ;
- le modèle obtenu par développement en fraction continue offre le meilleur comportement dans les basses fréquences, et la convergence la plus rapide.

Chaque méthode présente donc une convergence particulière dans le domaine fréquentiel. Cependant, il est remarquable que les modèles obtenus par développement de Foster ou de Cauer *selon la variable spatiale x* s'expriment comme des réseaux de Kirchhoff constitués d'*éléments discrets*. Ces éléments discrets s'expriment très simplement comme des fractions des constantes globales du système, soit R , L , C et G dans le cas général.

Nous avons donc montré que l'on peut représenter un système à constantes réparties (d'ordre infini) par un système à constantes localisées (réseau infini) [SL02]. C'est sur ce point que notre approche se distingue le plus de la méthode nodale (section 2.2.3) ou l'ordre infini disparaît immédiatement du fait de la discrétisation spatiale.

Nous allons maintenant étudier les modèles d'ordres réduits correspondant à des systèmes physiques réels.

4.2 Application à une ligne monophasée

Les méthodes de synthèse présentées pour le développement de l'impédance d'une ligne en court-circuit sont généralisables aux autres impédances d'ordre infini obtenues dans les modèles issus de l'équation des télégraphistes (fig. 2.13).

Nous allons commencer par appliquer la synthèse de Cauer au modèle de ligne en Π , chargé par une résistance de défaut (fig. 4.9).

4.2. Application à une ligne monophasée

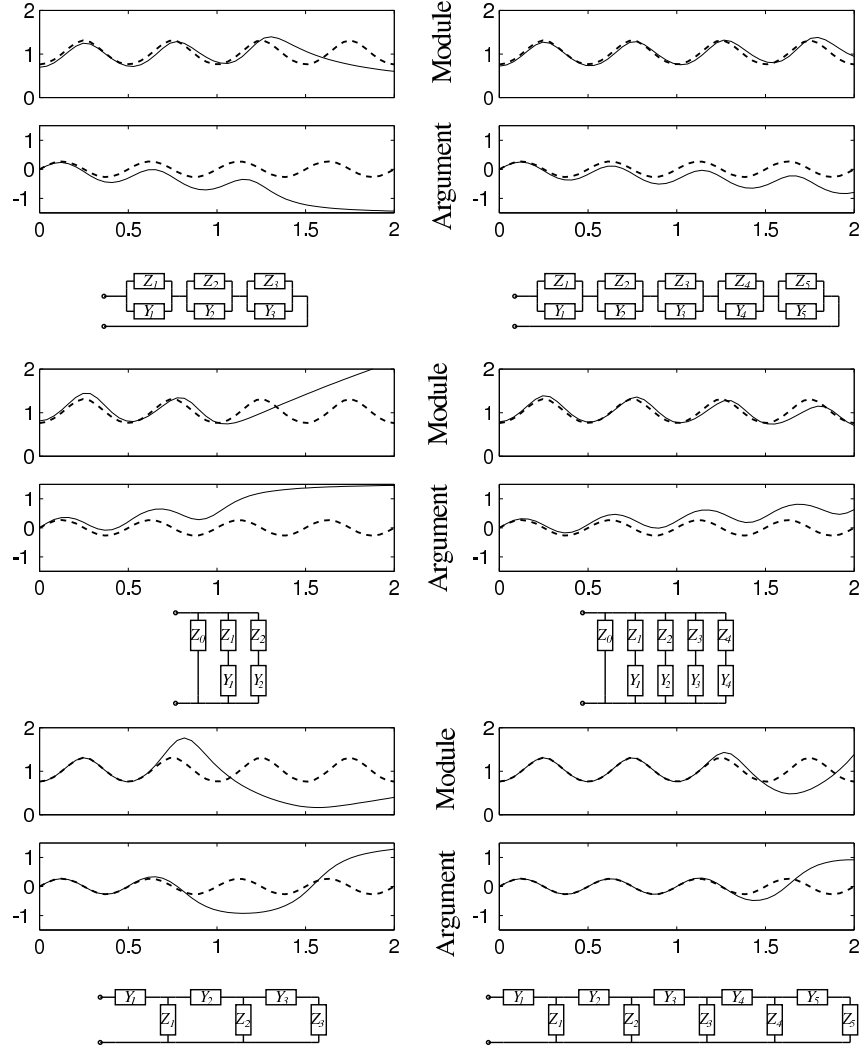


FIG. 4.8 – Réponses fréquentielles des modèles approchés obtenus par les trois méthodes présentées (cas $N = 3$ et $N = 5$) comparées à l'impédance d'ordre infini (en pointillés).

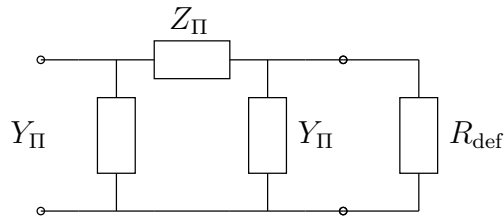


FIG. 4.9 – Modèle de ligne monophasée en défaut.

4.2.1 Synthèse de Cauer

La synthèse de Cauer des impédances d'ordre infini figurant dans 2.13 est réalisée par les développements en fractions continues suivants.

$$Z_{\Pi} = \frac{1}{\frac{1}{Z} + \frac{1}{-\frac{6}{Y} + \frac{1}{-\frac{10}{7Z} + \frac{1}{\frac{686}{11Y} + \ddots}}}}, \quad Y_{\Pi} = \frac{1}{\frac{2}{Y} + \frac{1}{\frac{6}{Z} + \frac{1}{\frac{10}{Y} + \frac{1}{\frac{14}{Z} + \ddots}}}}, \quad (4.14)$$

$$Z_T = \frac{1}{\frac{2}{Z} + \frac{1}{\frac{6}{Y} + \frac{1}{\frac{10}{Z} + \frac{1}{\frac{14}{Y} + \ddots}}}}, \quad Y_T = \frac{1}{\frac{1}{Y} + \frac{1}{-\frac{6}{Z} + \frac{1}{-\frac{10}{7Y} + \frac{1}{\frac{686}{11Z} + \ddots}}}}. \quad (4.15)$$

4.2.2 Capacités de compensation

Choisissons Le modèle en Π d'une ligne monophasée. En tronquant les fractions continues au premier ordre, nous avons les approximations

$$Z_{\Pi} \approx \frac{1}{\frac{1}{Z}} = Z, \quad Y_{\Pi} \approx \frac{1}{\frac{2}{Y}} = \frac{Y}{2}, \quad (4.16)$$

Nous retrouvons le modèle de ligne courte vu à la section 2.2.3.

Nous prenons maintenant un ordre de plus pour Z_{Π} ,

$$Z_{\Pi} \approx \frac{1}{\frac{1}{\frac{1}{Z} + \frac{1}{-\frac{6}{Y}}}}, \quad (4.17)$$

et nous supposons de plus les pertes transversales sont négligeables, $g = 0$ (cas $\alpha = 0$ dans la section 2.3.1), ce qui est souvent le cas pour la modélisation des lignes de transmission industrielles. En rappelant que

$$Z = R + Lp \quad Y = Cp \quad (4.18)$$

nous obtenons facilement le modèle de la figure 4.10

4.2. Application à une ligne monophasée

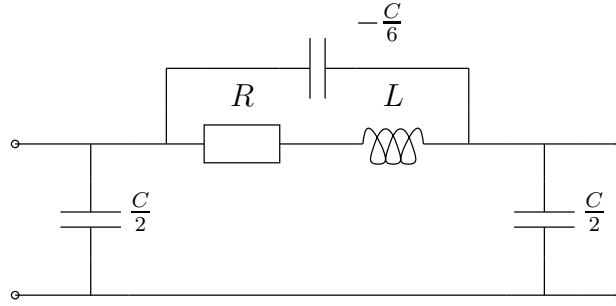


FIG. 4.10 – Modèle de ligne courte avec capacité de compensation.

Nous retrouvons le modèle précédent de ligne courte, mais auquel on a adjoint une capacité négative, ce qui en fait toute l'originalité. Ce nouveau modèle, possédant une *capacité de compensation* a été proposé par M. Lagonotte [LPC00] en remplacement des cellules chaînées (fig. 2.16) et mis en pratique notamment par [RM00].

Cette capacité négative apporte une nette amélioration de la réponse fréquentielle du modèle, comme on peut le voir sur la figure 4.12, à comparer avec 4.11. Ces figures ont été tracées en prenant les valeurs numériques d'une ligne à haute tension (tab. 4.1), chargée par une résistance de défaut R_{def} .

Paramètre	Symbole	Valeur
longueur	x	300 km
résistance linéique	r	$3 \cdot 10^{-2} \Omega \text{km}^{-1}$
inductance linéique	l	$1,05 \cdot 10^{-3} \text{Hkm}^{-1}$
conductance linéique	g	$\sim 0 \text{Skm}^{-1}$
capacité linéique	c	$11 \cdot 10^{-9} \text{Fkm}^{-1}$
résistance de charge	R_{def}	$10^3 \Omega$

TAB. 4.1 – Valeurs numériques du modèle de ligne.

4.2.3 Modèle développé

Au vu de l'amélioration apportée par l'ajout de capacités de compensation dans le modèle de la ligne courte, il est tentant de continuer la synthèse de Caueur des impédances du modèle.

Ainsi, à l'ordre six, nous obtenons la réponse fréquentielle de la figure 4.16.

Nous remarquons que la réponse fréquentielle converge bien vers la réponse du modèle d'ordre infini, sauf sur une étroite bande de fréquence, qui dépend de l'ordre d'approximation. Pour connaître l'effet de cette erreur, nous avons décidé d'effectuer une simulation dans le domaine temporel.

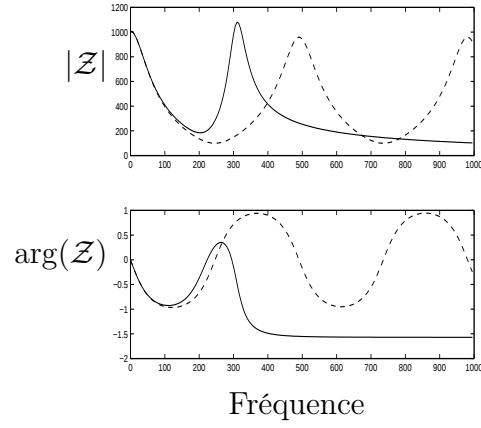


FIG. 4.11 – Réponse fréquentielle du modèle de ligne courte, comparée au modèle d'ordre infini (en pointillés).

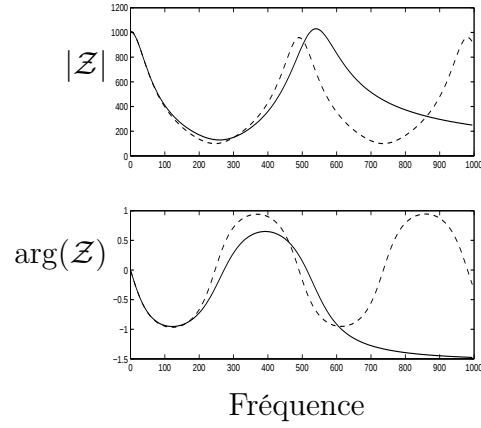


FIG. 4.12 – Réponse fréquentielle du modèle à capacité de compensation, comparée au modèle d'ordre infini (en pointillés).

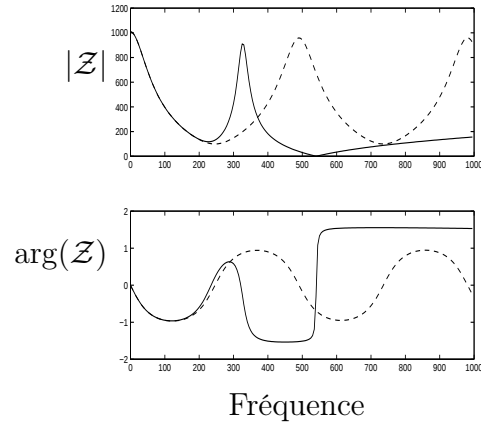


FIG. 4.13 – Réponse fréquentielle du modèle développé à l'ordre deux, comparée au modèle d'ordre infini (en pointillés).

4.2. Application à une ligne monophasée

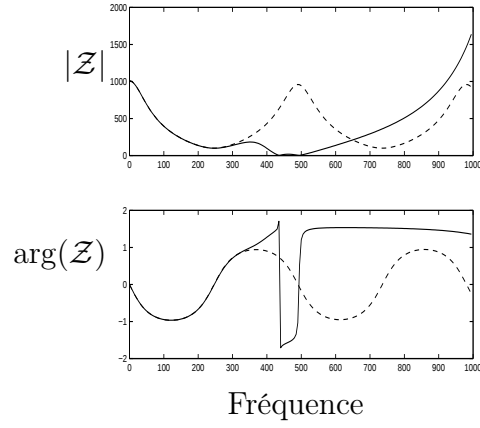


FIG. 4.14 – Réponse fréquentielle du modèle développé à l'ordre trois, comparée au modèle d'ordre infini (en pointillés).

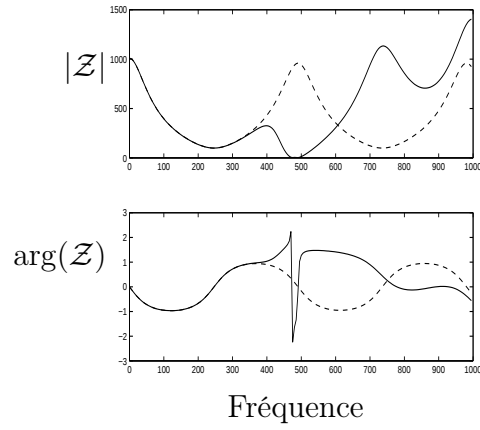


FIG. 4.15 – Réponse fréquentielle du modèle développé à l'ordre quatre, comparée au modèle d'ordre infini (en pointillés).

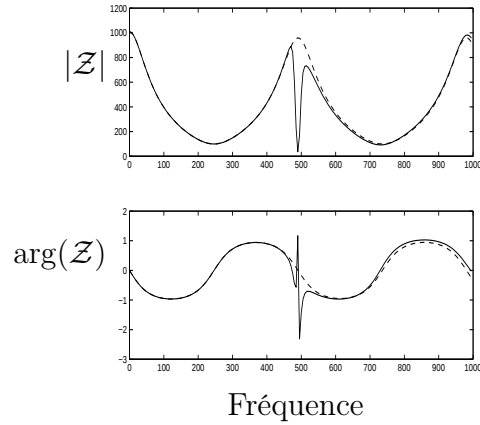


FIG. 4.16 – Réponse fréquentielle du modèle développé à l'ordre six, comparée au modèle d'ordre infini (en pointillés).

4.2.4 Simulation d'une réponses temporelle

Le modèle de ligne monophasée obtenu par la synthèse de Cauer est directement représenté sous la forme d'un réseau de Kirchhoff. Nous pouvons donc effectuer une simulation de réponse à une mise sous tension en utilisant un logiciel de simulation de réseau, tel que *Spice* ou *Esacap*⁵.

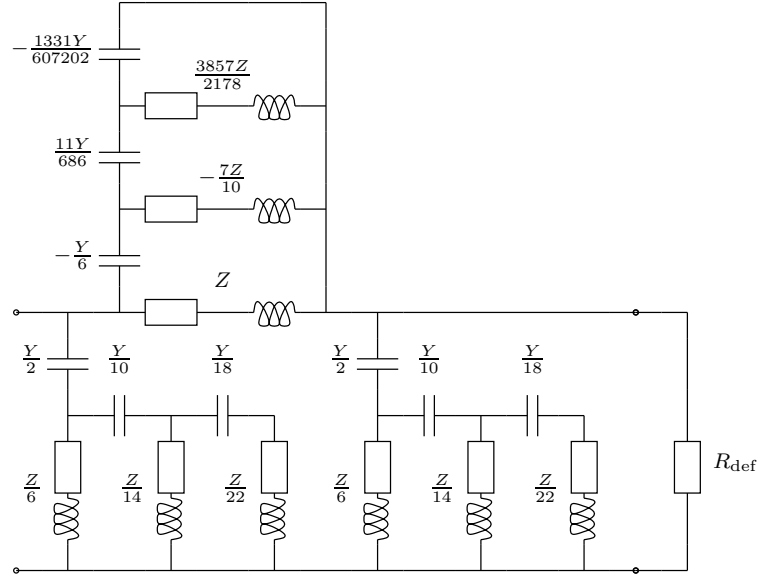


FIG. 4.17 – Modèle de ligne monophasée approché à l'ordre six sous la forme d'un réseau de Kirchhoff.

La réponse de référence est calculée par une transformée de Fourier inverse, à partir de l'impédance d'ordre infini du modèle en Π (fig. 4.9). Le calcul se fait numériquement à l'aide du logiciel *Matlab* par FFT⁶.

Nous avons utilisé *Esacap* pour calculer la réponse du modèle approché à une tension d'entrée imposée par intégration à pas constant dans le domaine temporel. Les réponses sont comparées sur la figure 4.18.

Nous voyons que même avec un ordre réduit, la pointe de courant initiale est bien reproduite, ainsi que les oscillations propres du système.

4.3 Application à un câble avec effet de peau

Jusqu'à présent, nous avons assimilé l'impédance interne des conducteurs composant une ligne à une résistance linéique r . Cette approximation n'est en fait valable que dans le domaine des très basses fréquences. Pour obtenir un modèle réaliste, nous devons tenir

⁵<http://bwrc.eecs.berkeley.edu/Classes/IcBook/SPICE/>, <http://eswww.it.dtu.dk/~el/ecs/esacap.htm>

⁶Fast Fourier Transform.

4.3. Application à un câble avec effet de peau

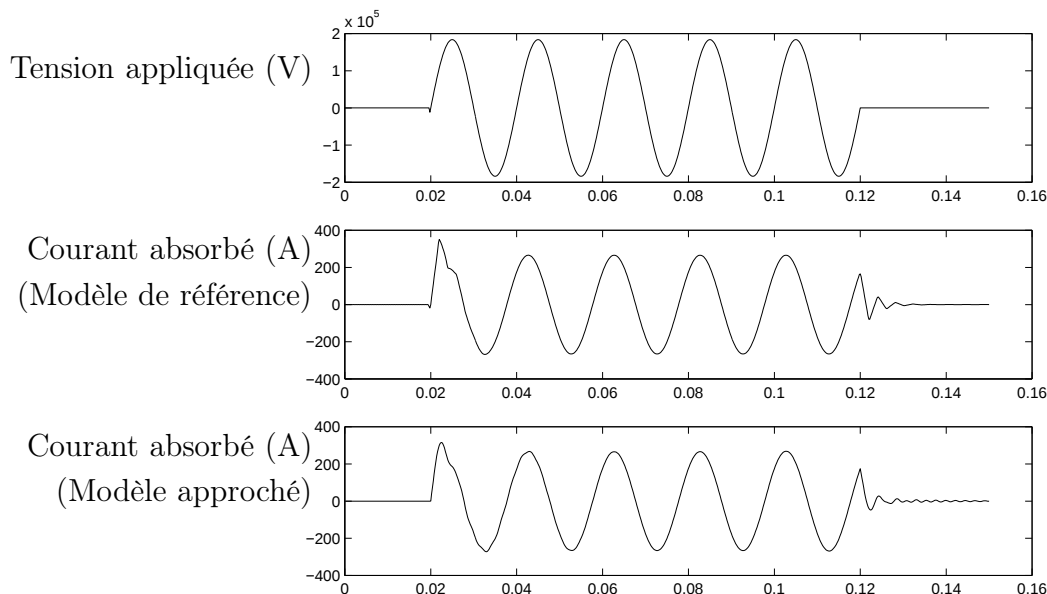


FIG. 4.18 – Simulation temporelle de la réponse du modèle de ligne à une mise sous tension.

compte de l'effet de peau. Commençons par citer la définition qu'en propose M. Fournet dans les *Techniques de l'ingénieur* [Fou93].

« L'expression *effet de peau* est relative aux phénomènes qui se produisent quand un conducteur est parcouru par un courant électrique dépendant du temps. L'expérience et la théorie montrent alors que la densité de courant n'est pas uniforme : il y a concentration des lignes de courant vers la surface extérieure du conducteur, cet aspect étant d'autant plus marqué que la vitesse temporelle de variation du courant est plus grande. »

Ce phénomène, qui tend à restreindre le courant dans une faible partie de la section des conducteurs, provoque une augmentation de la résistance apparente. Nous verrons que cette augmentation s'accompagne d'une variation de la phase du courant.

Pour modéliser cet effet, nous allons dans un premier temps effectuer une analyse de la répartition du courant dans un conducteur cylindrique, et nous en déduirons une expression de son impédance interne apparente. Dans un second temps, nous développerons des modèles sous forme de réseaux de Kirchhoff pour rendre compte des variations de cette impédance. Enfin nous proposerons un modèle de câble que nous confronterons à des mesures réelles.

4.3.1 Calcul de la répartition des filets de courants

Nous considérons ici un conducteur cylindrique et nous admettons que les *filets de courants* sont exclusivement orientés selon l'axe des x (figure 4.19). Nous savons que la densité

de courant est proportionnelle au champ électrique : $j = \sigma E$, σ étant la conductivité du matériau. Elle vérifie donc la même équation de propagation

$$\Delta j - \frac{p^2}{v^2} \left(1 + \frac{\sigma}{\epsilon p} \right) j = 0. \quad (4.19)$$

où c est la vitesse de propagation des ondes électromagnétiques, ϵ la permittivité diélectrique absolue, et p la variable complexe de Laplace.

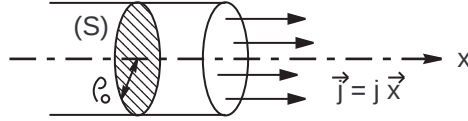


FIG. 4.19 – Filets de courants dans un conducteur cylindrique.

En utilisant l'expression du laplacien en coordonnées cylindriques, cette équation devient :

$$\frac{\partial^2 j}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial j}{\partial \rho} + \frac{1}{\rho^2} \frac{\partial^2 j}{\partial \theta^2} + \frac{\partial^2 j}{\partial x^2} - \frac{p^2}{c^2} \left(1 + \frac{\sigma}{\epsilon p} \right) j = 0. \quad (4.20)$$

Or, la permittivité électrique d'un conducteur métallique est $\epsilon \approx \epsilon_0 = \frac{1}{36\pi} 10^{-9} F.m^{-1}$ (permittivité du vide). Et d'autre part, la conductivité du cuivre est $\sigma \approx 6.10^7 S.m^{-1}$, et celle de l'aluminium, $\sigma \approx 3,5.10^7 S.m^{-1}$. La constante de temps $\frac{\epsilon}{\sigma}$ étant alors de l'ordre de $10^{-18} s$, le terme $\frac{\sigma}{\epsilon p}$ est négligeable devant 1 pour des fréquences inférieures à $10^{17} Hz$ (c'est le domaine des ultraviolets) !

D'autre part, dans le cas de la propagation sans pertes, nous avons

$$\frac{\partial^2 j}{\partial x^2} = \frac{p^2}{c^2} j. \quad (4.21)$$

Si les pertes sont faibles, nous pouvons alors supposer que ces deux termes restent du même ordre de grandeur [Pél75]. Nous les négligerons donc tous deux de la même manière, et en remarquant que $c^2 = \frac{1}{\epsilon\mu}$ nous obtenons l'équation qui détermine la répartition du courant.

$$\frac{\partial^2 j}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial j}{\partial \rho} + \frac{1}{\rho^2} \frac{\partial^2 j}{\partial \theta^2} - \sigma\mu p j = 0 \quad (4.22)$$

Au vu de la symétrie cylindrique du problème, nous pouvons rechercher des « solutions élémentaires » de la forme $j = f(\rho)g(\theta)$. En reportant dans l'équation précédente, nous obtenons

$$f''g + \frac{f'g}{\rho} + \frac{fg''}{\rho^2} - \sigma\mu p fg = 0. \quad (4.23)$$

Nous multiplions alors chacun des termes par $\frac{\rho^2}{fg}$ pour isoler les expressions respectivement fonction de ρ et de θ .

$$\left(\rho^2 \frac{f''}{f} + \rho \frac{f'}{f} \right) + \frac{g''}{g} - \sigma\mu p = 0. \quad (4.24)$$

4.3. Application à un câble avec effet de peau

À p fixé, nous en déduisons que nécessairement

$$\frac{g''}{g} = C \quad \text{C étant une constante.} \quad (4.25)$$

De plus, comme g est périodique de période 2π , nous pouvons écrire : $g(\theta) = \cos(n\theta + \phi_n)$, d'où

$$\rho^2 f'' + \rho f' - (n^2 + \sigma\mu p\rho^2) f = 0. \quad (4.26)$$

Nous reconnaissons l'équation différentielle de Bessel modifiée [Lam78]. La fonction f est donc la somme de fonctions de Bessel modifiées de première (I_n) et de seconde espèce (K_n).

$$f(\rho) = a_n I_n(\sqrt{\mu\sigma p}\rho) + b_n K_n(\sqrt{\mu\sigma p}\rho). \quad (4.27)$$

Mais les fonctions de Bessel modifiées de seconde espèce admettent une singularité à l'origine (figure 4.20). La densité de courant ne pouvant tendre vers l'infini au centre du conducteur, la seule solution physiquement acceptable est celle correspondant à $b_n = 0$.

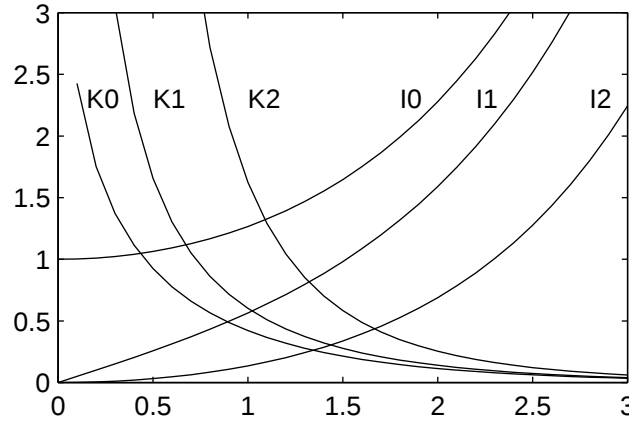


FIG. 4.20 – Fonctions de Bessel modifiées de première (I_n) et seconde espèce (K_n).

La solution générale de l'équation (4.22) s'exprime comme une combinaison linéaire de « solutions élémentaires » déterminées ci-dessus.

$$j = \sum_{n=0}^{\infty} a_n I_n(\sqrt{\mu\sigma p}\rho) \cos(n\theta + \phi_n) \quad (4.28)$$

Cette expression autorise le calcul du courant total i qui parcourt le conducteur. Par définition, il suffit de sommer j sur une section du cylindre.

$$i = \iint_{(S)} j dS = \sum_{n=0}^{\infty} 2\pi a_n \int_0^{\rho_0} \rho I_n(\sqrt{\mu\sigma p}\rho) d\rho \int_0^{2\pi} \cos(n\theta + \phi_n) d\theta. \quad (4.29)$$

La seconde intégrale s'annulant pour tout n non nul, l'expression du courant se réduit à

$$i = 2\pi a_0 \cos \phi_0 \int_0^{\rho_0} \rho I_0(\sqrt{\mu\sigma p}\rho) d\rho. \quad (4.30)$$

En utilisant un résultat connu sur l'intégration des fonctions de Bessel, à savoir que pour tout réel ν ,

$$\int x^\nu I_{\nu-1}(\alpha x) dx = \frac{x^\nu}{\alpha} I_\nu(\alpha x), \quad (4.31)$$

nous avons au final

$$i = \frac{2\pi a_0 \cos \phi_0 \rho_0}{\sqrt{\mu\sigma p}} I_1(\sqrt{\mu\sigma p}\rho_0). \quad (4.32)$$

4.3.2 Calcul de l'impédance interne

Supposons maintenant que la composante longitudinale du champ électrique E_x est uniforme à la surface du conducteur ($\rho = \rho_0^+$). Par continuité de la composante tangentielle du champ électrique à l'interface air-métal, la densité de courant $j = \sigma E_x$ à la périphérie ($\rho = \rho_0^-$) est elle aussi uniforme. Cette condition n'est remplie que si les termes de la somme (4.28) sont tous nuls à l'exception de celui correspondant à $n = 0$,

$$j = a_0 \cos \phi_0 I_0(\sqrt{\mu\sigma p}\rho). \quad (4.33)$$

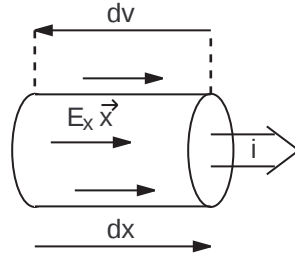


FIG. 4.21 – Segment de conducteur cylindrique de longueur dx .

D'autre part, nous définissons l'impédance interne par unité de longueur z_i à l'aide de la relation analogue à la loi d'Ohm portant sur la différence de potentiel selon la direction $\mathbf{x}$ et le courant total (figure 4.21).

$$-dv = z_i dx i. \quad (4.34)$$

Nous avons donc, par définition du potentiel,

$$z_i = -\frac{1}{i} \frac{dv}{dx} = \frac{E_x(\rho_0)}{i} = \frac{j(\rho_0)}{\sigma i}, \quad (4.35)$$

4.3. Application à un câble avec effet de peau

et d'après les équations (4.32) et (4.33),

$$z_i = \frac{1}{2\pi\rho_0} \sqrt{\frac{\mu p}{\sigma}} \frac{I_0(\sqrt{\mu\sigma p}\rho_0)}{I_1(\sqrt{\mu\sigma p}\rho_0)} \quad (4.36)$$

Cette dernière expression est due à Sergei A. Schelkunoff [Sch34] (selon [Joh97]). Elle fournit une expression de l'impédance interne d'ordre infini. Nous allons donc naturellement rechercher un modèle approché de l'impédance interne, comme nous l'avons fait pour l'impédance de la ligne.

4.3.3 Synthèse de Cauer

La méthode pour générer un modèle de Kirchhoff du conducteur à effet de peau s'applique ici. En effet, remarquons que le rapport des fonctions de Bessel $I_0(x)$ et $I_1(x)$ forme une fonction impaire admettant l'origine comme pôle d'ordre un. Son développement en fraction continue s'écrit :

$$\frac{I_0(x)}{I_1(x)} = \frac{2}{x} + \frac{1}{\frac{4}{x} + \frac{1}{\frac{6}{x} + \frac{1}{\frac{8}{x} + \ddots}}} \quad (4.37)$$

Utilisons maintenant les définitions de R , la résistance totale du conducteur en continu, et de L_i l'inductance interne données en annexe D. Notons :

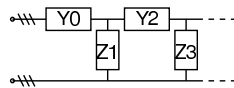
$$Y = \frac{1}{R} = \frac{\pi\rho_0^2\sigma}{X} \quad Z = L_i p = \frac{\mu X p}{4\pi} \quad (4.38)$$

Nous pouvons alors réécrire l'impédance interne totale apparente du conducteur,

$$z_i X = \sqrt{\frac{Z}{Y}} \frac{I_0(2\sqrt{ZY})}{I_1(2\sqrt{ZY})} \quad (4.39)$$

et en déduire une expression de z_i sous forme de fraction continue, puis d'impédance de réseau en échelle.

$$z_i X = \frac{1}{Y} + \frac{1}{\frac{2}{Z} + \frac{1}{\frac{3}{Y} + \frac{1}{\frac{4}{Z} + \ddots}}} \quad (4.40)$$



4.3.4 Modélisation d'ordre non entier

Reprenons l'expression (4.36) de l'impédance interne de Schelkunoff. Les comportements asymptotiques des fonctions de Bessel nous permettent d'en déduire des équivalents classiques. Aux basses fréquences, nous avons en effet ⁷ :

$$I_0(\sqrt{\mu\sigma p\rho_0}) \underset{0}{\sim} 1 \quad (4.41)$$

$$I_1(\sqrt{\mu\sigma p\rho_0}) \underset{0}{\sim} \sqrt{\mu\sigma p} \frac{\rho_0}{2} \quad (4.42)$$

$$\text{d'où } z_i \underset{0}{\sim} \frac{1}{\pi\rho^2\sigma} \quad (4.43)$$

Nous retrouvons donc pour les basses fréquences $z_i = r$, la résistance du conducteur en continu par unité de longueur. De même, pour les hautes fréquences, les comportements asymptotiques ⁸ fournissent :

$$I_0(\sqrt{\mu\sigma p\rho_0}) \underset{\infty}{\sim} I_1(\sqrt{\mu\sigma p\rho_0}) \quad (4.44)$$

$$z_i \underset{\infty}{\sim} \frac{1}{2\pi\rho_0} \sqrt{\frac{\mu p}{\sigma}} \quad (4.45)$$

Ce dernier équivalent traduit le comportement non-entier (d'ordre un demi) de l'impédance z_i aux hautes fréquences.

Les variations de l'impédance et de ses équivalents figurent sur le graphique 4.22 avec une échelle logarithmique, afin de mettre en évidence la bande de croissance à 10 dB par décades. Les valeurs numériques correspondent à un conducteur fictif en cuivre de un centimètre de rayon. Nous voyons que la fréquence de transition est de l'ordre de la centaine de hertz. Aussi serait-il illusoire de vouloir négliger l'effet de peau pour les gammes de fréquence auxquelles s'étudient les phénomènes transitoires.

4.3.5 Comparaison du modèle de Cauet et du modèle non entier

Cette modélisation de l'effet de peau par un réseau de résistances et d'inductances a notamment été utilisé en simulation par [RM00]. Les éléments du réseau choisis varient comme des puissances de deux. La comparaison des réponses fréquentielles de ce modèle (figure 4.24) et de celui obtenu à l'aide du développement en fraction continue (figure 4.23) montre que ce dernier se rapproche plus de la réponse théorique calculée.

Des modèles similaires ont aussi été obtenus sous forme de réseaux récursifs par découpage arbitraire du conducteur en éléments cylindriques [YFW82]. Ce découpage respectant un rapport constant entre les résistances successives, il est à rapprocher des approches de type « non-entiers » [Ous95].

⁷ $I_\nu(x) \underset{0}{\sim} \frac{1}{\Gamma(\nu+1)} \left(\frac{x}{2}\right)^\nu$ pour tout ν réel, $\nu \neq -1, -2, -3, \dots$

⁸ $I_\nu(x) \underset{\infty}{\sim} e^x \sqrt{\frac{2}{\pi x}}$ pour tout ν réel.

4.3. Application à un câble avec effet de peau

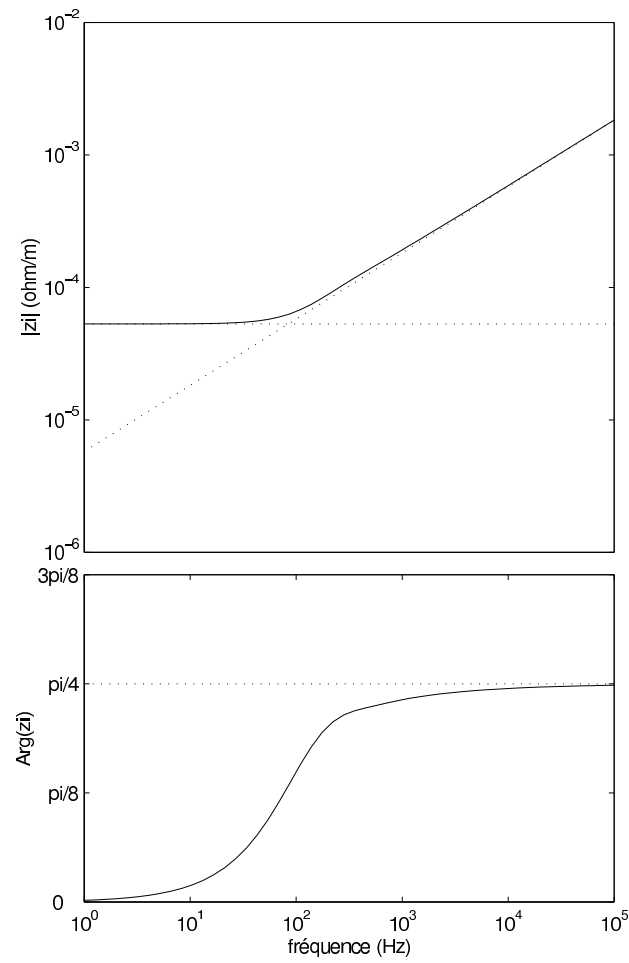


FIG. 4.22 – Variation du module et de l'argument de l'impédance interne du conducteur.

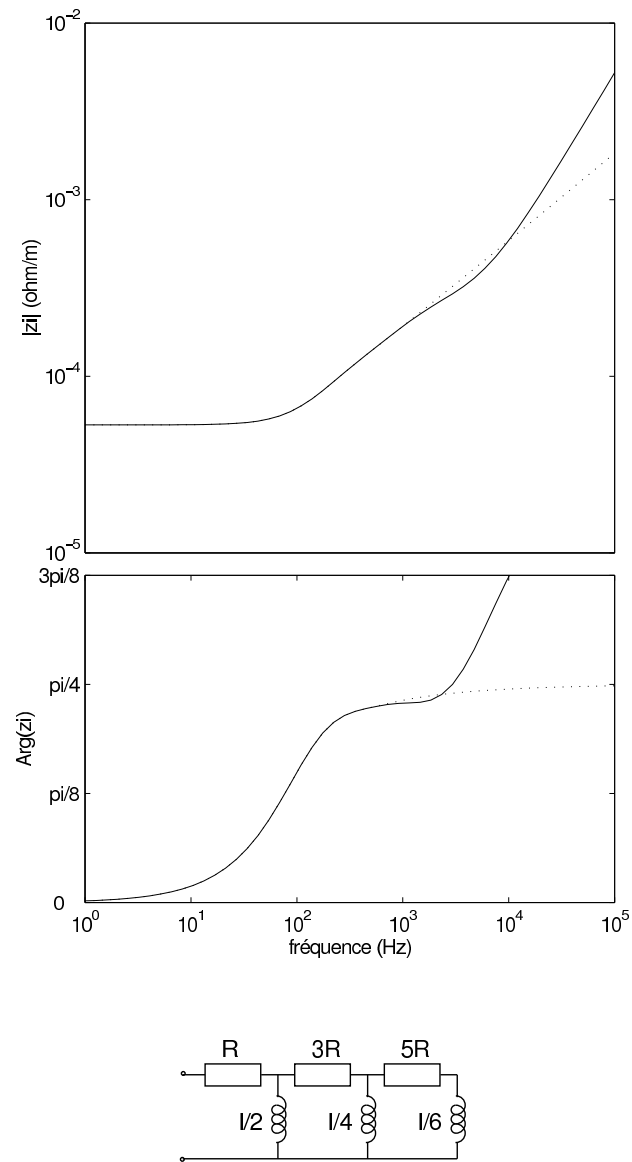


FIG. 4.23 – Module et argument du réseau en échelle obtenu par synthèse de Cauer.

4.3. Application à un câble avec effet de peau

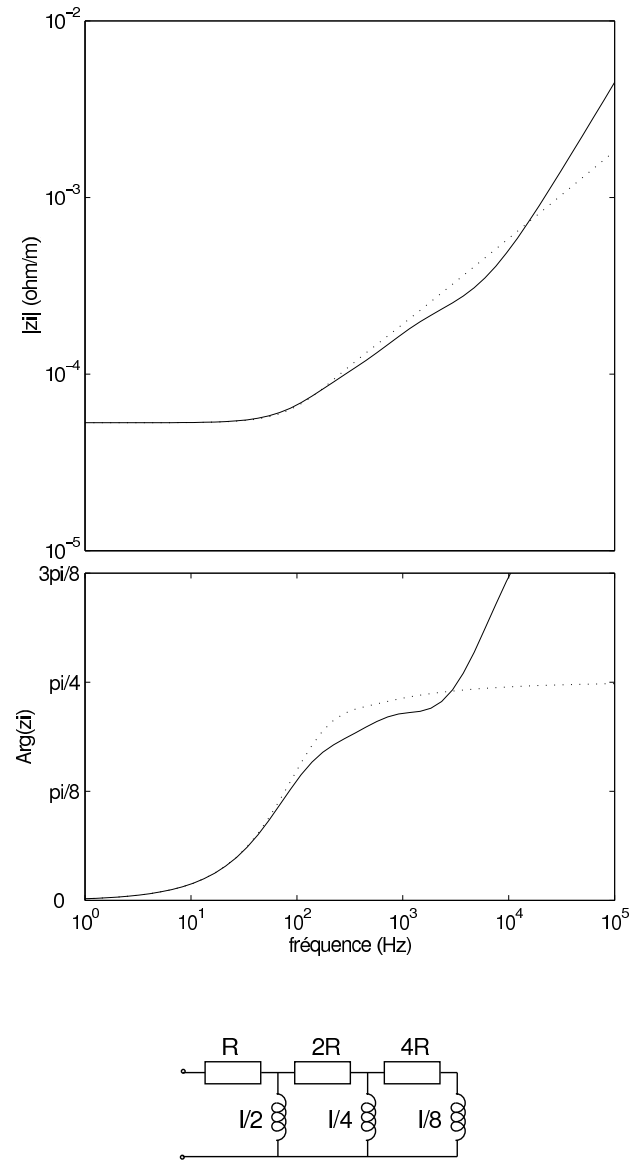


FIG. 4.24 – Module et argument du réseau en échelle obtenu par une modélisation de type non-entier [RM00].

4.3.6 Application à un câble réel

Nous disposons de données mesurées sur l'impédance d'un câble électrique rigide de type RO2V 3G1,5. Ces données proviennent de l'étude effectuée par M. Rossat-Mignod [RM00]. Les caractéristiques du câble mesurées sont reportées dans le tableau 4.2. Ces caractéristiques nous permettent de calculer l'impédance théorique du câble en court-circuit.

$$Z_{cc} = \sqrt{\frac{Z}{Y}} \tanh(\sqrt{ZY}) \quad (4.46)$$

Les expressions de Z et de Y sans prise en compte de l'effet de peau étant :

$$Z = (r + lp)X \quad Y = (g + cp)X \quad (4.47)$$

Nous pouvons comparer l'impédance calculée ainsi avec l'impédance mesurée (figure 4.25).

Paramètre	Symbole	Valeur
longueur	X	200 m
section	$\pi\rho_0^2$	1,5 mm ²
résistance linéique	r	1,215 10 ⁻² Ωm ⁻¹
inductance linéique	l	3,2 10 ⁻⁷ Hm ⁻¹
conductance linéique	g	3,87 10 ⁻⁸ Sm ⁻¹
capacité linéique	c	2,5 10 ⁻¹¹ Fm ⁻¹

TAB. 4.2 – Caractéristiques du câble mesurées à une fréquence d'environ 2 kHz.

Il est clair que les divergences sont très importantes. La modélisation choisie est donc trop simpliste car elle fait abstraction de l'effet de peau.

Modifions donc l'expression de Z pour y introduire l'impédance interne du câble. Nous faisons aussi apparaître l'inductance externe l_e ⁹.

$$Z = (z_i + l_e p)X \quad z_i = \frac{1}{2\pi\rho_0} \sqrt{\frac{\mu p}{\sigma}} \frac{I_0(\sqrt{\mu\sigma p\rho_0})}{I_1(\sqrt{\mu\sigma p\rho_0})} \quad (4.48)$$

Le calcul de l'impédance à l'aide de cette nouvelle expression fournit des valeurs en bien meilleur accord avec celles relevées (figure 4.26).

Il subsiste tout de même des écarts qui nous ont amenés à modifier les valeurs des paramètres du câble. Ces valeurs proviennent en effet d'une mesure faite à une unique fréquence et qui n'offre peut être donc pas toute la précision requise.

Afin d'ajuster les paramètres du câble, nous avons utilisé un algorithme d'optimisation qui minimise une fonction de coût quadratique de la différence entre les valeurs calculées et mesurées de l'impédance. La fonction de coût a été pondérée par $\frac{1}{|p|^2}$ afin d'accorder plus

⁹Pour le calcul de la conductivité et de l'inductance externe, nous utiliserons $\sigma = \frac{1}{r\pi\rho_0^2}$ et $l_e = l - l_i = l - \frac{\mu}{4\pi}$

4.3. Application à un câble avec effet de peau

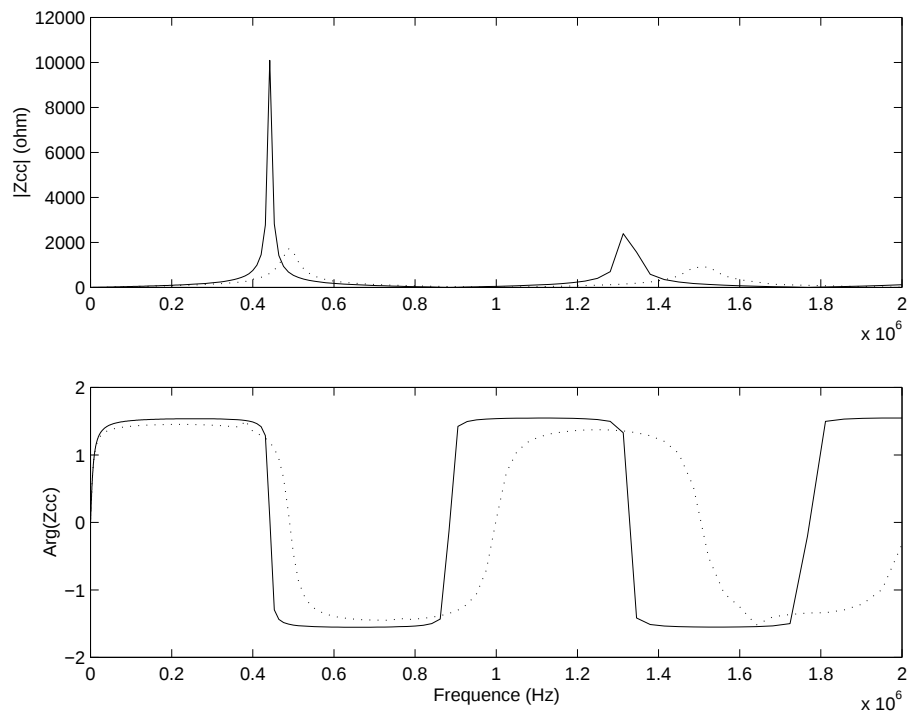


FIG. 4.25 – Impédance du câble en court-circuit mesurée (pointillés) et calculée sans effet de peau (trait plein).

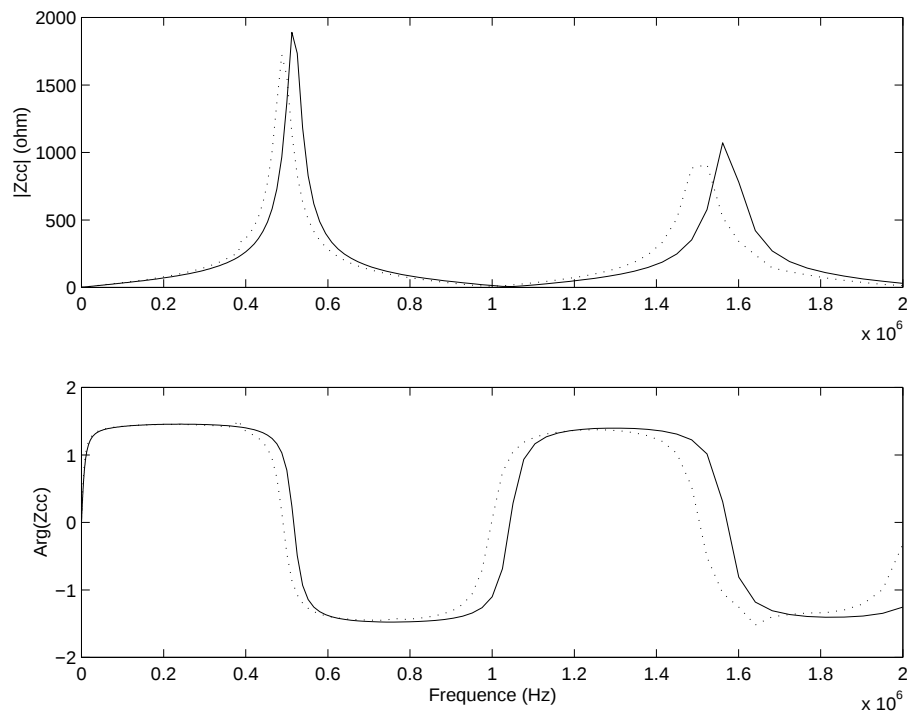


FIG. 4.26 – Impédance du câble en court-circuit mesurée (pointillés) et calculée en tenant compte de l'effet de peau (trait plein).

4.3. Application à un câble avec effet de peau

d'importance aux basses fréquences. Les nouvelles valeurs des constantes électriques sont reportées dans le tableau 4.3. Les corrections restent minimales, sauf en ce qui concerne la valeur de la conductance linéique g . Cette correction n'est pas surprenante puisque ce paramètre est réputé être le plus difficile à estimer et présente de plus une valeur extrêmement faible.

Paramètre	Symbole	Valeur corrigée
résistance linéique	r	$1,134 \cdot 10^{-2} \Omega \text{m}^{-1}$
inductance linéique	l	$3,530 \cdot 10^{-7} \text{Hm}^{-1}$
conductance linéique	g	$1,676 \cdot 10^{-6} \text{Sm}^{-1}$
capacité linéique	c	$2,430 \cdot 10^{-11} \text{Fm}^{-1}$

TAB. 4.3 – Caractéristiques corrigées du câble après optimisation.

Les valeurs corrigées des constantes du câble permettent de faire coïncider les variations de l'impédance calculée avec les mesures (figure 4.27).

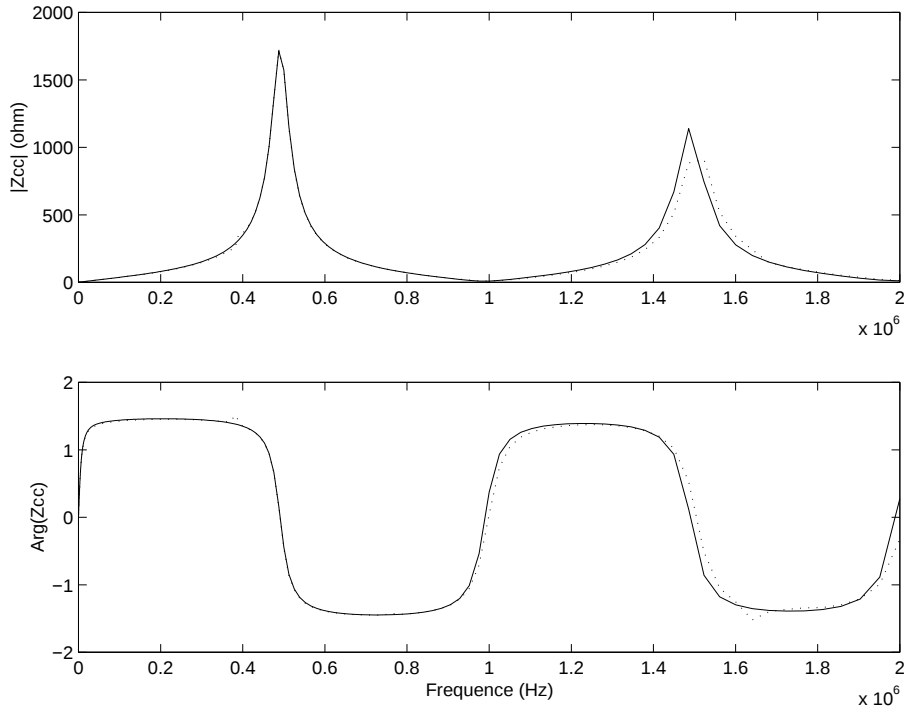


FIG. 4.27 – Impédance du câble en court-circuit mesurée (pointillés) et calculée avec les valeurs corrigées des paramètres (trait plein).

Notre modèle théorique étant maintenant validé, et les valeurs des paramètres ajustées, nous pouvons mettre à profit le développement en fractions continues pour générer un

réseau de Kirchhoff approchant la réponse fréquentielle du câble. L'impédance Z_{cc} est d'abord développée suivant la fonction tangente hyperbolique, puis l'impédance $z_i X$ est développée suivant le rapport $\frac{I_0}{I_1}$.

Le modèle approché obtenu est celui de la figure 4.28. Son impédance est représentée sur la même figure. Le réseau d'éléments discrets obtenu a donc une impédance équivalente très proche du câble réel en court-circuit.

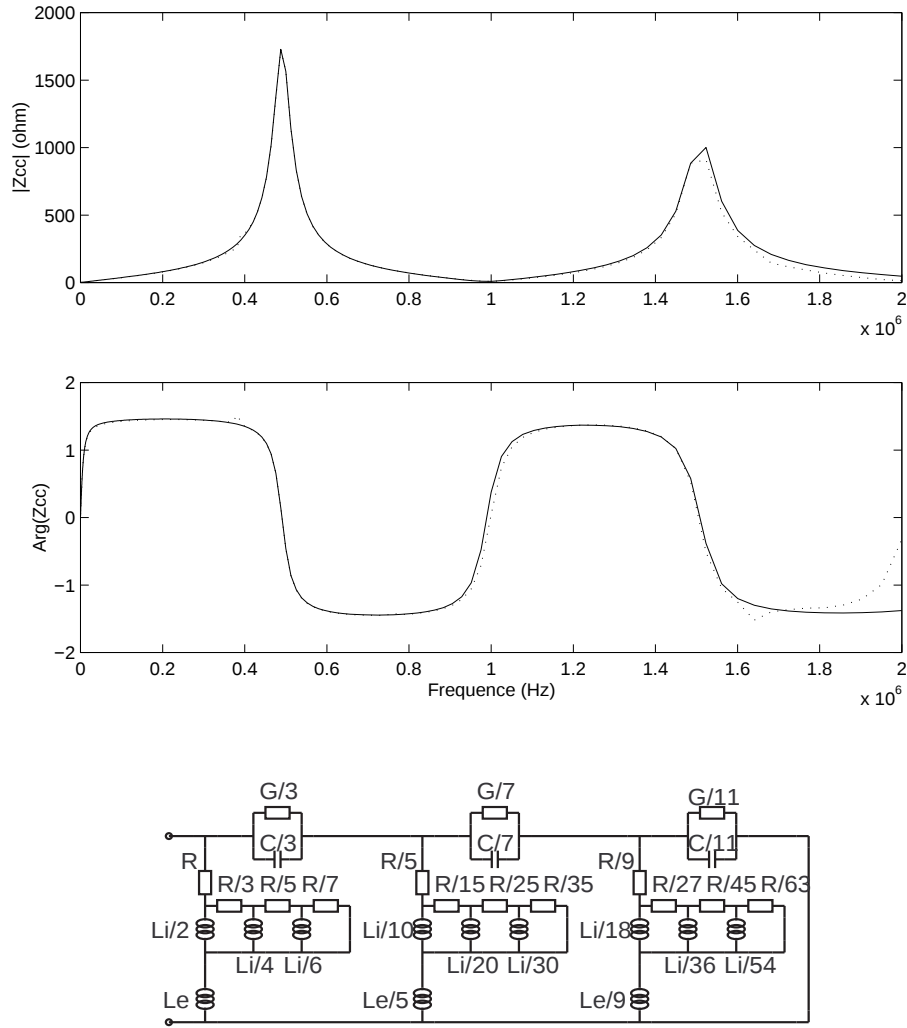


FIG. 4.28 – Modèle du câble en court-circuit avec effet de peau sous la forme d'un réseau de Kirchhoff. L'impédance du modèle (trait plein) est comparée à celle mesurée (pointillés).

Nous avons donc validé notre modèle d'effet de peau, et montré qu'il était plus pertinent que l'approche non-entière habituelle. Nous pouvons donc conclure que la synthèse de Cauey permet d'obtenir des modèles réalistes de câbles avec effet de peau, utilisables en temps réel ([SL04] dans [cem04]).

4.4. Application à une ailette de refroidissement

4.4 Application à une ailette de refroidissement

Nous avons vu à la section 2.3.1 que l'équation de la chaleur peut être vue comme un cas dégénéré de l'équation des télégraphistes pour lequel $l = 0$. Nous allons le vérifier en étudiant deux exemples de problèmes thermiques, puis nous utiliserons le développement d'impédances pour modéliser ces systèmes.

4.4.1 Le système physique étudié

Commençons par un système thermique simple. La modélisation d'une ailette faisant partie d'un dispositif de refroidissement nous permettra de comparer les modèles obtenus par ces différentes méthodes.

Considérons une ailette en aluminium homogène, de masse volumique ρ , de chaleur massique c et de conductivité thermique λ . Nous supposons que l'ailette est plongée dans un fluide à la température uniforme T_0 , et nous caractérisons les échanges thermiques avec l'extérieur par une liaison convective de coefficient h . Nous portons à T_0 la face arrière de l'ailette, et nous notons T_1 la température en entrée, où un flux thermique Φ est imposé (fig. 4.29).

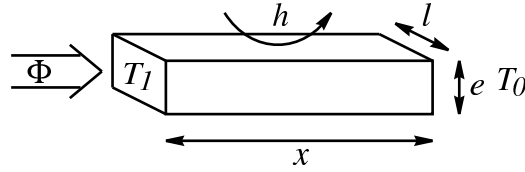


FIG. 4.29 – Le système thermique à modéliser : une ailette de refroidissement.

Vu de l'entrée, ce système est équivalent à une impédance d'ordre infini

$$Z = \sqrt{\frac{Z}{Y}} \tanh(\sqrt{ZY}) \quad (4.49)$$

([Mar90]) avec

$$Z = \frac{x}{\lambda e}, \quad Y = (h(l + e) + \rho c l e p)x. \quad (4.50)$$

Les différentes valeurs numériques utilisées pour la suite sont présentées dans le tableau 4.4.

4.4.2 La modélisation nodale classique

La modélisation nodale (vue dans la section 2.2.3) est la méthode la plus intuitive pour obtenir un système d'ordre réduit sous la forme d'un réseau de résistances et de capacités. Elle correspond à une discrétisation spatiale du système physique. Pour cela, nous calculons d'abord les constantes thermiques globales de l'ailette :

Symbole	Valeur
x	5 cm
l	3 cm
e	5 mm
ρ	2600 kg.m ⁻³
λ	204 W(mK) ⁻¹
c	879 J(kgK) ⁻¹
h	100 WK ⁻¹ m ⁻²

TAB. 4.4 – Valeurs numériques caractéristiques de l'ailette.

- la résistance thermique R ,
- la capacité thermique C ,
- la conductance fluidique G .

$$R = \frac{x}{\lambda e}, \quad C = \rho c x l e, \quad G = h(l + e)x. \quad (4.51)$$

Nous choisissons ensuite d'utiliser trois cellules élémentaires en Π , telles que celles de la figure 2.16, composées de fractions de ces constantes globales. En associant les cellules, et compte tenu de la température imposée à l'extrémité (analogue à un court-circuit électrique), nous obtenons le réseau de la figure 4.30.

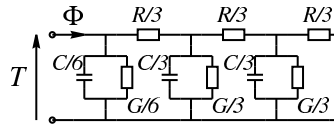


FIG. 4.30 – Modèle de l'ailette obtenu par la méthode nodale.

4.4.3 Les synthèses de Foster et Cauer

Utilisons les résultats de la section 4.1. Le réseau de la figure 4.3 nous fournit le deuxième modèle approché et les valeurs $Z = R$ et $Y = G + Cp$ nous permettent d'en expliciter les éléments pour obtenir le modèle de la figure 4.31. Ce modèle correspond à la première synthèse de Foster.

De la même façon, en utilisant le réseau de la figure 4.5, nous aboutissons au modèle de la figure 4.32 par la deuxième synthèse de Foster. Nous fixons encore une fois N de manière à conserver trois cellules.

Utilisons maintenant le développement de l'impédance formelle Z en fraction continue. En remplaçant Z et Y par les éléments correspondants, le réseau de la figure 4.7 nous fournit le quatrième modèle approché (fig. 4.33) qui est celui obtenu par synthèse de Cauer.

4.4. Application à une ailette de refroidissement

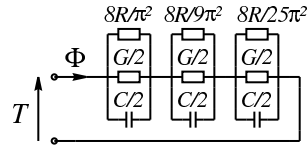


FIG. 4.31 – Modèle de l'ailette obtenu par la première synthèse de Foster.

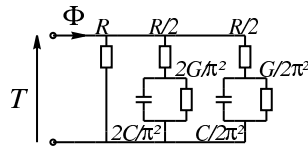


FIG. 4.32 – Modèle de l'ailette obtenu par la deuxième synthèse de Foster.

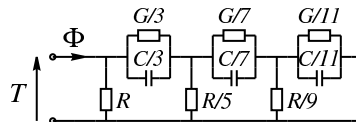


FIG. 4.33 – Modèle de l'ailette obtenu par la synthèse de Cauer.

4.4.4 Comparaison des modélisations

Les quatre réseaux présentés sont en tous points comparables :

- même nombre de composants,
- même nombre de nœuds,
- même ordre pour l'impédance approchée¹⁰.

Une première comparaison entre ces différents modèles est faite dans le plan de Nyquist. Sur la figure 4.34, nous avons donc tracé la réponse fréquentielle de chaque modèle ainsi que la réponse du modèle original (en pointillés).

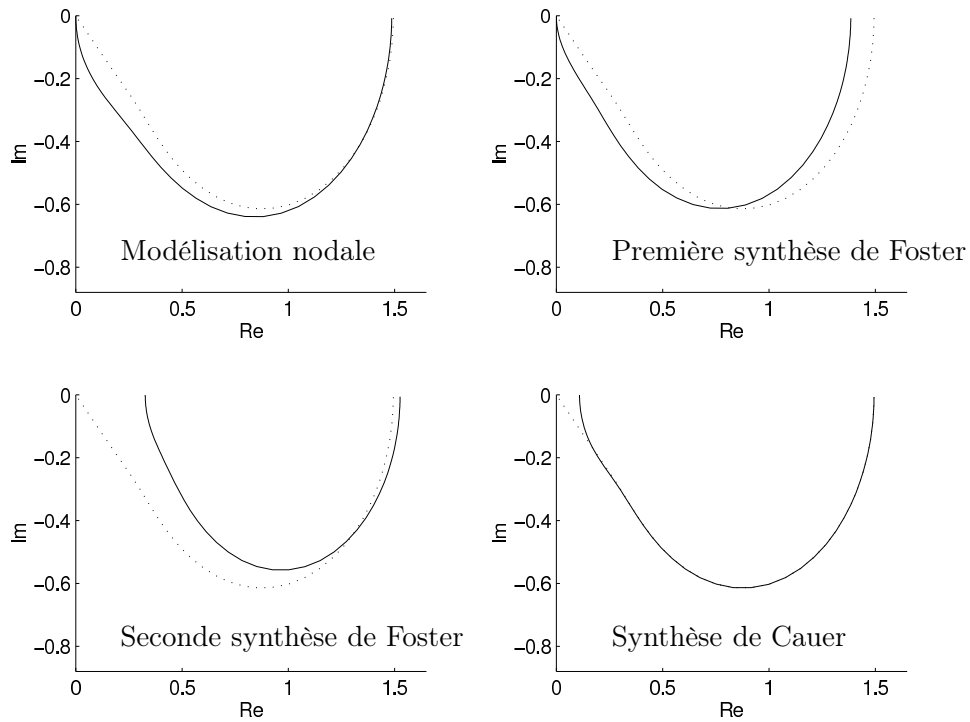


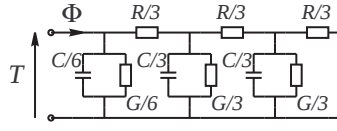
FIG. 4.34 – Diagrammes de Nyquist des différents modèles comparés à la réponse de référence (en pointillés).

Nous pouvons utiliser ces modèles comme précédemment avec un logiciel de simulation de réseau. Nous cherchons à déterminer la réponse temporelle en température $T = T_1 - T_0$ de l'ailette quand on lui impose un échelon de flux (réponse indicelle). Les réponses simulées sont comparées à la réponse de référence, calculée par une transformée de Fourier inverse, à partir de l'expression exacte de $\mathcal{Z}$. Ces résultats sont reportés sur la figure 4.35.

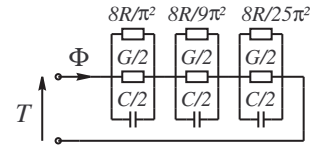
¹⁰Les impédances des réseaux présentés sont des approximations rationnelles de l'impédance d'ordre infini. Elles sont assimilables à des fonctions de transfert d'ordre fini égal à trois, sauf pour le modèle obtenu par la deuxième synthèse de Foster où trois cellules correspondent à un ordre deux, la première branche ne comportant qu'une résistance.

4.4. Application à une ailette de refroidissement

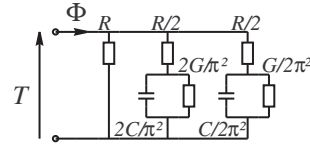
Modèles réseaux à trois nœuds



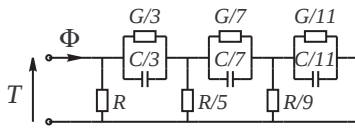
Synthèse nodale



Première synthèse de Foster



Deuxième synthèse de Foster



Synthèse de Cauer

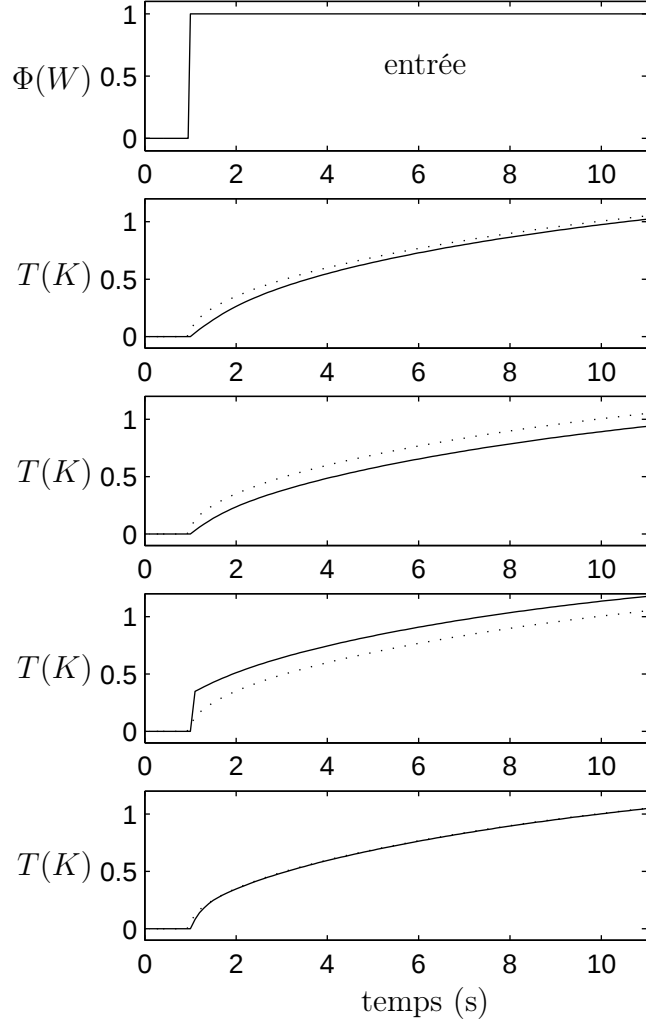


FIG. 4.35 – Réponses comparées des réseaux avec la réponse de référence (en pointillés).

Le modèle obtenu par développement en fraction continue, du fait de sa bande passante optimale, a un comportement bien meilleur que le modèle nodal durant les phases transitoires.

4.5 Application à la conduction thermique dans un cylindre

Nous allons étudier maintenant un dernier exemple d'application dans le domaine de la thermique en exploitant le développement de Cauer calculé pour modéliser l'effet de peau. Plaçons dans le cadre de la conduction thermique simple, c'est-à-dire, avec nos conventions de notation,

$$l = 0, \quad g = 0. \quad (4.52)$$

Soit un cylindre constitué d'un matériau homogène, de masse volumique ρ , de chaleur massique c et de conductivité thermique λ . Nous noterons T_0 la température au centre du cylindre. Nous supposons ensuite que la température est uniforme à la surface du cylindre : $T(r_0) = T_1$, ce qui est compatible avec la symétrie du problème (figure 4.36).

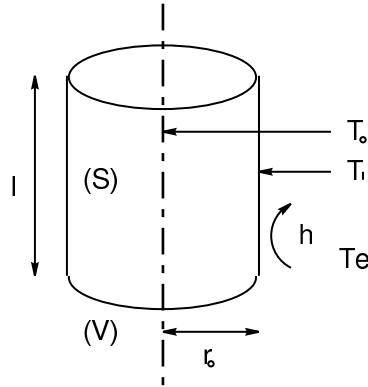


FIG. 4.36 – Les caractéristiques du cylindre.

Afin de se rapprocher d'un problème réel, nous considérerons par la suite que le cylindre est plongé dans un fluide à la température T_e et nous caractériserons les échanges par une liaison convective de coefficient h .

Nous allons proposer un modèle thermique du cylindre basé sur l'analogie électrique en nous limitant au cas où la diffusion est purement radiale (les extrémités sont adiabatiques). Les valeurs numériques utilisées lorsque cela sera nécessaire sont rappelées dans le tableau 4.5.

4.5. Application à la conduction thermique dans un cylindre

4.5.1 Modèle thermique d'ordre infini

Équation de la chaleur

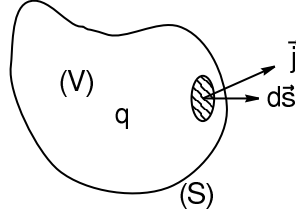


FIG. 4.37 – Volume de contrôle pour le bilan énergétique.

Soit un volume (V) quelconque de matériau conducteur de chaleur, délimité par une surface fermée (S) (figure 4.37). Notons q la densité de source de chaleur volumique. Le bilan énergétique sur (V) s'écrit :

$$\iiint_{(V)} q d\tau - \iiint_{(V)} \rho c \frac{\partial T}{\partial t} d\tau - \oint_{(S)} \mathbf{j} d\mathbf{s} = 0 \quad (4.53)$$

A l'aide de la formule d'Ostrogradski, nous pouvons transformer la dernière intégrale de façon à obtenir :

$$\iiint_{(V)} (q - \rho c \frac{\partial T}{\partial t} - \text{div} \mathbf{j}) d\tau = 0 \quad (4.54)$$

Cette relation étant valable quel que soit le volume de contrôle (V) , nous pouvons écrire de façon locale :

$$q - \rho c \frac{\partial T}{\partial t} - \text{div} \mathbf{j} = 0 \quad (4.55)$$

Exprimons maintenant le vecteur densité de chaleur à l'aide de la loi de Fourier, ce qui permet d'exprimer l'équation de diffusion de la chaleur sous sa forme classique :

$$\mathbf{j} = -\lambda \nabla T \quad (4.56)$$

$$\Delta T - \frac{\rho c}{\lambda} \frac{\partial T}{\partial t} = \frac{q}{\lambda} \quad (4.57)$$

Résolution à l'aide de la transformée de Laplace

Nous supposons que le cylindre ne contient pas de sources de chaleur. Après avoir effectué une transformation de Laplace sur l'équation, il reste :

$$\Delta T - \frac{\rho c}{\lambda} p T = 0; \quad 0 < r < r_0 \quad (4.58)$$

En utilisant l'expression du laplacien en coordonnées cylindriques, il vient :

$$\frac{\partial^2 T}{\partial r^2} + \frac{1}{r} \frac{\partial T}{\partial r} + \frac{1}{r^2} \frac{\partial^2 T}{\partial \theta^2} + \frac{\partial^2 T}{\partial z^2} - \frac{\rho c}{\lambda} p T = 0 \quad (4.59)$$

Et puisque nous nous sommes limité au cas de la diffusion radiale,

$$\frac{\partial^2 T}{\partial r^2} + \frac{1}{r} \frac{\partial T}{\partial r} - \frac{\rho c}{\lambda} p T = 0 \quad (4.60)$$

Nous reconnaissons ici l'équation différentielle de Bessel modifiée d'ordre 0 [SL00]. La solution générale $T(r)$ est donc la somme d'une fonction de Bessel modifiée de première espèce et d'une fonction de Bessel modifiée seconde espèce.

$$T = A I_0 \left(\sqrt{\frac{\rho c p}{\lambda}} r \right) + B K_0 \left(\sqrt{\frac{\rho c p}{\lambda}} r \right) \quad (4.61)$$

Mais les fonctions de Bessel modifiées de seconde espèce admettent une singularité à l'origine (figure 4.20). La température ne pouvant tendre vers l'infini au centre du cylindre, la seule solution physiquement acceptable est celle correspondant à $B = 0$. La condition à la limite $r = 0$ fixe alors $A = T_0$.

$$T = T_0 I_0 \left(\sqrt{\frac{\rho c p}{\lambda}} r \right) \quad (4.62)$$

Calcul de l'impédance équivalente

Calculons maintenant le flux à travers la surface $r = r_0$. La loi de Fourier donne :

$$\mathbf{j} = -\lambda \nabla T = -T_0 \sqrt{\rho c \lambda p} I_1 \left(\sqrt{\frac{\rho c p}{\lambda}} r \right) \mathbf{u}_r \quad (4.63)$$

Le flux est obtenu en intégrant $\mathbf{j}$ sur la surface du cylindre, dans le sens du flux « entrant ».

$$\Phi = - \iint_{(S)} \mathbf{j} ds = S j(r_0) \quad (4.64)$$

$$\Phi = 2\pi r_0 l T_0 \sqrt{\rho c \lambda p} I_1 \left(\sqrt{\frac{\rho c p}{\lambda}} r_0 \right) \quad (4.65)$$

Nous en déduisons l'expression de l'impédance thermique ([MAB⁺00]) équivalente du cylindre :

$$\mathcal{Z} = \frac{T_1}{\Phi} = \frac{1}{2\pi r_0 l \sqrt{\lambda \rho c p}} \frac{I_0 \left(\sqrt{\frac{\rho c p}{\lambda}} r_0 \right)}{I_1 \left(\sqrt{\frac{\rho c p}{\lambda}} r_0 \right)} \quad (4.66)$$

4.5. Application à la conduction thermique dans un cylindre

4.5.2 Modèles réduits

Nous allons chercher un modèle réduit du cylindre sous forme d'un réseau de Kirchhoff. Un tel modèle pourra être exploité dans le domaine temporel. Pour cela nous allons comparer deux méthodes. La première, consiste à discrétiser de façon régulière le cylindre, puis de remplacer chaque élément par une cellule électrique équivalente. La seconde, moins intuitive, se base sur le développement de l'impédance thermique en fraction continue. Les deux modèles obtenus seront ensuite comparés à la fois du point de vue fréquentiel et temporel.

Discrétisation régulière

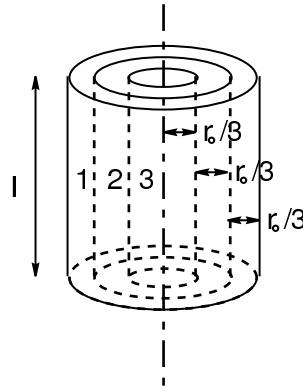


FIG. 4.38 – Découpage régulier du cylindre en trois.

Il s'agit de la méthode la plus « classique » pour obtenir un modèle thermique approché. Nous découpons de façon arbitraire le cylindre en trois, en respectant la symétrie cylindrique (figure 4.38). Nous attribuons à chaque segment une température moyenne. Il convient alors de calculer la capacité thermique de chaque segment, ainsi que les résistances inter-segments.

$$C_1 = \frac{5}{9}\pi\rho cr_0^2l \quad C_2 = \frac{1}{3}\pi\rho cr_0^2l \quad C_3 = \frac{1}{9}\pi\rho cr_0^2l \quad (4.67)$$

$$R_{12} = \frac{1}{4\pi\lambda l} \quad R_{23} = \frac{1}{2\pi\lambda l} \quad (4.68)$$

Nous pouvons alors définir le réseau en échelle équivalent (figure 4.39), en notant C , la capacité thermique totale du cylindre, et R la résistance thermique entre les segments 1 et 2 :

$$C = \pi\rho clr_0^2 \quad R = \frac{1}{4\pi\lambda l} \quad (4.69)$$

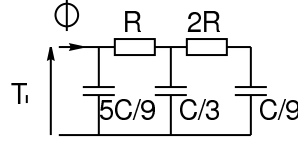


FIG. 4.39 – Réseau en échelle équivalent, méthode nodale.

Développement de l'impédance en fraction continue

Nous pouvons utiliser les définitions (4.69) de R et de C afin de réécrire l'impédance du cylindre (équation (4.66)) sous la forme :

$$\mathcal{Z} = \sqrt{\frac{R}{Cp}} \frac{I_0(2\sqrt{RCp})}{I_1(2\sqrt{RCp})} \quad (4.70)$$

Nous pouvons donc réutiliser les résultats de la section 4.3 pour donner une expression de $\mathcal{Z}$ sous forme de fraction continue. En tronquant la fraction continue à un certain ordre, on obtient un approximant de Padé de $\mathcal{Z}$, qui est aussi l'impédance du réseau en échelle de la figure 4.40.

$$\mathcal{Z} = \frac{1}{Cp} + \frac{1}{\frac{2}{R} + \frac{1}{\frac{3}{Cp} + \frac{1}{\frac{4}{R} + \dots}}} \quad (4.71)$$

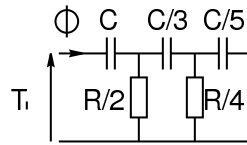


FIG. 4.40 – Réseau en échelle équivalent – Seconde méthode.

4.5.3 Comparaison des modèles

Diagrammes de Bode

Nous pouvons comparer les deux types de réseaux obtenus dans le domaine fréquentiel. Pour cela nous calculons l'impédance équivalente de chacun d'eux.

$$Z_1 = \frac{1 + \frac{2}{3}RCp + \frac{2}{27}R^2C^2p^2}{Cp + \frac{4}{9}C^2Rp^2 + \frac{10}{243}C^3R^2p^3} \quad (4.72)$$

4.5. Application à la conduction thermique dans un cylindre

$$Z_2 = \frac{1 + \frac{4}{5}RCp + \frac{3}{40}R^2C^2p^2}{Cp + \frac{3}{10}C^2Rp^2 + \frac{1}{120}C^3R^2p^3} \quad (4.73)$$

En remplaçant l'opérateur complexe de Laplace p par la variable imaginaire pure $2\pi i\nu$, nous pouvons tracer les diagrammes de Bode de ces deux impédances, et les comparer à l'impédance d'ordre infini (équation (4.66)). Le diagramme obtenu sur la figure 4.41 pour un cylindre de cuivre de rayon un centimètre et de longueur dix centimètres permet de conclure que la seconde méthode génère un modèle utilisable sur une plage de fréquence bien plus large que le « découpage » de la méthode nodale.

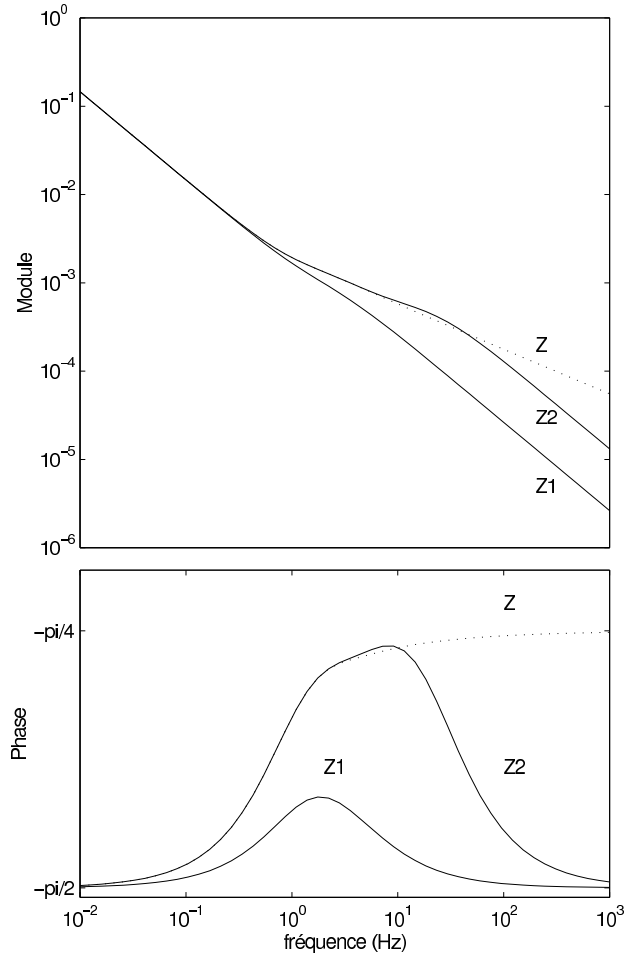


FIG. 4.41 – Diagrammes de Bode des impédances Z_1 et Z_2 comparées à Z .

Modèles complets en simulation temporelle

Nous obtenons les modèles complets en ajoutant en série une résistance thermique représentant les échanges convectifs entre la surface du cylindre et le fluide.

$$R_c = \frac{1}{2\pi r_0 l h} \quad (4.74)$$

Nous obtenons alors un montage électrique appelé « diviseur de tension » (figure 4.42). La relation qui relie alors T_1 , la température à la surface du cylindre à la température T_e du fluide est un résultat connu :

$$T_1 = \frac{Z}{R_c + Z} T_e \quad (4.75)$$

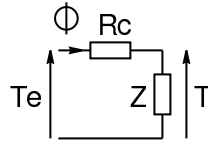


FIG. 4.42 – Modèle électrique complet.

Nous comparons alors la réponse en température du cylindre à un échauffement du milieu fluide (nous avons fixé h à $10 \text{ W K}^{-1} \text{ m}^{-2}$) obtenue à l'aide d'une transformée de Fourier inverse [MAB⁺00], à la réponse obtenue par simulation sous *Esacap* avec le premier puis le second modèle. Les résultats sont tracés sur la figure 4.43.

Paramètre	Symbole	Valeur
rayon	r_0	0,01 m
longueur	l	0,1 m
masse volumique	ρ	$8,8 \cdot 10^3 \text{ kg m}^{-3}$
chaleur massique	c	$394 \text{ J kg}^{-1} \text{ K}^{-1}$
conductivité thermique	λ	$384 \text{ W kg}^{-1} \text{ K}^{-1}$
coefficient d'échange fluidique	h	$10 \text{ W K}^{-1} \text{ m}^{-2}$

TAB. 4.5 – Récapitulatif des valeurs numériques employées.

Nous pouvons conclure que la méthode présentée permet d'obtenir des modèles utilisables dans le domaine temporel en simulation bien plus performants que des modèles nodaux « intuitifs » lors des phases transitoires.

4.5. Application à la conduction thermique dans un cylindre

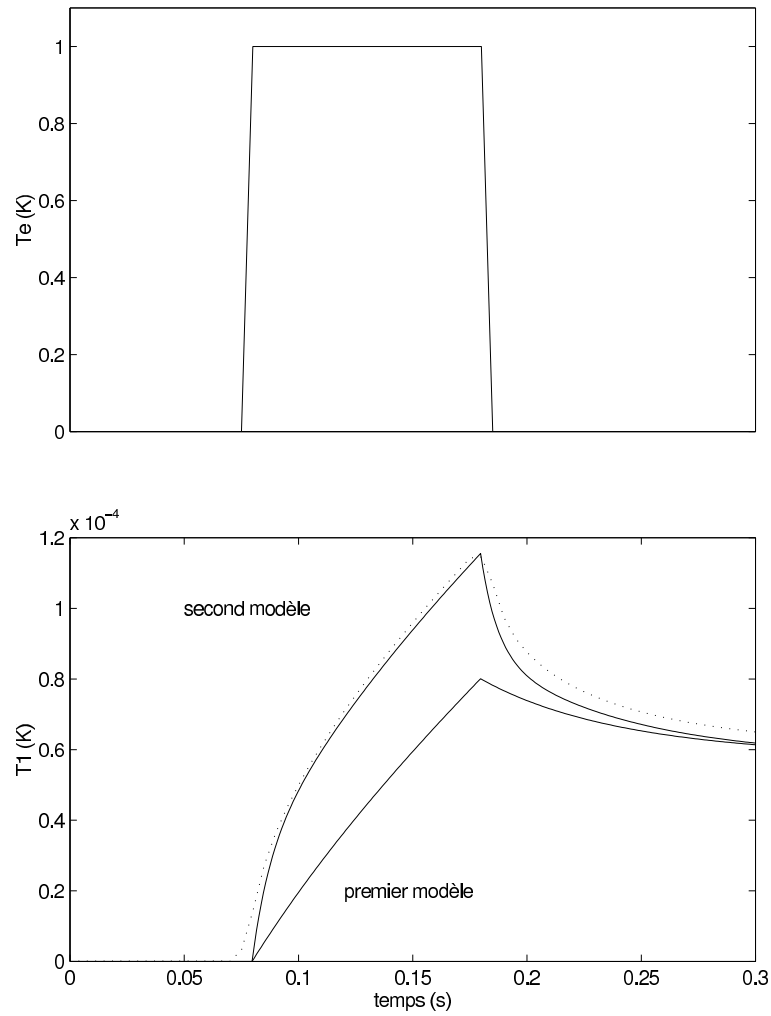


FIG. 4.43 – Réponses en température des deux modèles (simulation sous *Esacap*) comparés à la réponse du modèle d'ordre infini obtenue par transformée de Fourier inverse (en pointillés).

4.6 Conclusion sur les développements d'impédances

Au travers des différents exemples présentés dans ce chapitre, nous avons pu juger de la qualité des modèles d'ordres réduits obtenus par la synthèse de réseaux de Kirchhoff par développement d'impédances d'ordre infini. Il est remarquable que cette méthode permette de traiter de façon unifiée les phénomènes propagatifs (équation hyperbolique) et diffusifs (équation parabolique). Les applications dans le domaine de la thermique ont donné lieu à un DEA [Tri00] et à deux autres thèses actuellement en cours au sein du LET, ainsi qu'à plusieurs publications [LBBS01, TSL02]. Il serait sans aucun doute intéressant d'utiliser ces modèles pour d'autres applications, toujours par le biais d'analogies électriques (voir par exemple [Ber93] pour l'acoustique).

Le choix d'effectuer le développement selon des variables spatiales (X et ρ_0) et non fréquentielle (p), comme il est habituel de le faire pour l'approximation de modèles en automatique, se justifie par le fait que nous ne conservons qu'un faible nombre de paramètres quel que soit la complexité du modèle développé. En effet, les éléments constituant les réseaux sont toujours des fractions des valeurs des paramètres globaux (R , L , C , G) du système. Cet avantage est décisif dans une optique d'identification paramétrique (chapitre 6). Notons toutefois que le développement suivant les variables spatiales ou fréquentielles coïncident dans certains cas (notamment pour $\alpha = 0$ de 2.3.1, ou pour la ligne sans pertes). Pour les autres cas, l'approximation n'est pas parfaite au voisinage de la fréquence nulle.

Nous retiendrons que les modèles obtenus par développement d'impédances sont en général beaucoup plus efficaces, à complexité égale, que les modèles obtenus par des méthodes de discrétisation spatiale (méthode nodale). Enfin, les modèles obtenus par synthèse de Foster convergent plus lentement que ceux obtenus par synthèse de Cauer, ce qu'on peut expliquer par les propriétés des approximations de Padé. C'est pour cette raison que nous retiendrons la synthèse de Cauer comme méthode de modélisation pour la suite de cette étude.

L'approche que nous avons adoptée, repose cependant sur la résolution analytique des équations décrivant le système. Cela n'est possible en pratique que pour des géométries assez simples. Dans le cas de formes plus complexes, on peut utiliser l'interconnexion des réseaux pour modéliser le système par parties. Il est même possible de décrire ainsi des systèmes hétérogènes par couplage de modèles thermiques et électriques, comme par exemple dans [Bro00].

Une autre approche développée au LET par M. Nicolas Dollinger et M. Cédric Rouaud est la synthèse de réseaux par optimisation, à partir de la seule réponse fréquentielle du modèle. Les structures de Cauer donnent dans ce cas aussi de très bons résultats en simulation.

Mais pour utiliser le développement d'impédances en modélisation des lignes polyphasées, nous devons à présent généraliser cette méthode aux phénomènes propagatifs ou diffusifs couplés. Nous allons voir dans le chapitre suivant que cette extension est possible en utilisant un formalisme matriciel pour décrire les impédances.

4.6. Conclusion sur les développements d'impédances

Chapitre 5

Extension au cas des lignes polyphasées

Dans ce chapitre, nous allons démontrer que la synthèse de Cauer telle que nous l'avons exploitée au chapitre 4 peut se généraliser pour modéliser la propagation des ondes de courants et de tensions sur une ligne constituée de plusieurs conducteurs. Il peut s'agir d'une ligne triphasée, à un ou plusieurs ternes, ou d'un câble, les conducteurs étant alors l'âme et le blindage. Dans un cas encore plus général, on peut décider d'inclure les câbles de garde ou le retour par la terre dans la description de la ligne.

La principale différence avec le cas monophasé est l'existence de couplage entre les différents conducteurs. Nous les représenterons à l'aide d'un formalisme matriciel pour l'équation des télégraphistes que nous résoudrons pour obtenir des modèles de lignes composés d'*impédances matricielles*, auxquelles nous appliquerons la synthèse de Cauer.

5.1 Équation des télégraphistes matricielle

Commençons par exprimer, puis résoudre l'équation des télégraphistes pour une ligne à plusieurs conducteurs en adaptant la méthode exposée dans l'annexe B.1.

5.1.1 Modélisation de l'élément de ligne multifilaire

Considérons donc un élément de ligne à n conducteurs. Notons v_i et i_i la tension et le courant dans le conducteur i . Définissons alors les vecteurs

$$\mathbf{v} = \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix}, \quad \mathbf{i} = \begin{pmatrix} i_1 \\ i_2 \\ \vdots \\ i_n \end{pmatrix}. \quad (5.1)$$

La description de l'évolution de l'état électrique au niveau local se fait par un système d'équations aux dérivées partielles matricielles. Par la transformée de Laplace, nous

5.1. Équation des télégraphistes matricielle

obtenons

$$\begin{cases} \frac{d\mathbf{v}}{dx} = -\mathbf{z}\mathbf{i} \\ \frac{d\mathbf{i}}{dx} = -\mathbf{y}\mathbf{v} \end{cases} \quad (5.2)$$

Cette équation a la même forme que dans le cas monophasé, à la différence que les constantes réparties $\mathbf{z}$ et $\mathbf{y}$ sont des matrices $n \times n$ caractérisant la ligne et les couplages locaux [Joh98]. Les éléments de $\mathbf{z}$ ont la dimension d'une impédance par unité de longueur, les éléments de $\mathbf{y}$ celle d'une admittance par unité de longueur.

Par exemple, dans le cas d'une ligne triphasée parfaitement symétrique, nous pouvons représenter l'élément de ligne par la figure 5.1 [Cla94] et [Esc97, ch. 9]. Ce qui se traduit par

$$\mathbf{z} = \begin{pmatrix} r + lp & mp & mp \\ mp & r + lp & mp \\ mp & mp & r + lp \end{pmatrix}, \quad \mathbf{y} = \begin{pmatrix} g + cp & dp & dp \\ dp & g + cp & dp \\ dp & dp & g + cp \end{pmatrix}. \quad (5.3)$$

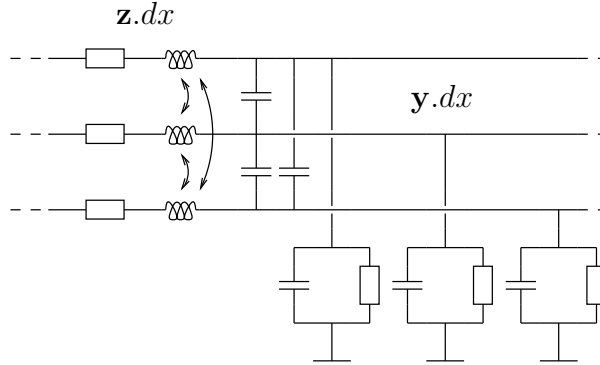


FIG. 5.1 – Élément de ligne triphasée (avec capacités interphases).

5.1.2 Passage en composantes modales

Pour résoudre l'équation des télégraphistes polyphasée, on utilise les composantes *modales*, c'est-à-dire les vecteurs propres pour diagonaliser les matrices [NR75, Gar76, Joh94].

Les équations de propagation pour une ligne multifilaire s'écrivent

$$\frac{d^2\mathbf{v}}{dx^2} = \mathbf{zyv}, \quad \frac{d^2\mathbf{i}}{dx^2} = \mathbf{yzv}. \quad (5.4)$$

Du fait de la réciprocité des inductances et des capacités mutuelles, les matrices $\mathbf{z}$ et $\mathbf{y}$ sont symétriques. Nous en déduisons que les matrices $\mathbf{zy}$ et $\mathbf{yz}$ sont transposées.

$$(\mathbf{zy})^T = \mathbf{y}^T \mathbf{z}^T = \mathbf{yz}. \quad (5.5)$$

Nous savons alors qu'elles possèdent les mêmes valeurs propres, que nous noterons γ_i^2 , pour $1 \leq i \leq n$. Nous supposons de plus que ces matrices sont diagonalisables et nous notons P_v et P_i les matrices de vecteurs propres correspondants.

$$\mathbf{z}\mathbf{y} = P_v \Gamma^2 P_v^{-1}, \quad \mathbf{y}\mathbf{z} = P_i \Gamma^2 P_i^{-1}, \quad \text{avec } \Gamma^2 = \text{Diag}(\gamma_i^2)_{1 \leq i \leq n}. \quad (5.6)$$

Le changement de base associé à chaque diagonalisation est une projection sur les vecteurs propres, encore appelés *modes*. Les composantes *modales* des courants et des tensions peuvent s'exprimer sous la forme de vecteurs, $\mathbf{v}_m$ et $\mathbf{i}_m$, pour lesquels nous réécrivons les équations de propagation.

$$\mathbf{v}_m = P_v^{-1} \mathbf{v}, \quad \mathbf{i}_m = P_i^{-1} \mathbf{i}. \quad (5.7)$$

$$\frac{d^2 \mathbf{v}_m}{dx^2} = \Gamma \mathbf{v}_m \quad \frac{d^2 \mathbf{i}_m}{dx^2} = \Gamma \mathbf{i}_m \quad (5.8)$$

5.1.3 Résolution dans l'espace modal

Puisque Γ est diagonale, la résolution des équations s'apparente à celle utilisée pour la ligne monophasée. Nous pouvons écrire, par exemple, pour chaque composante modale de la tension (voir section B.1)

$$v_{mi}(x) = \lambda_i e^{\gamma_i x} + \mu_i e^{-\gamma_i x}, \quad \text{pour } 1 \leq i \leq n. \quad (5.9)$$

ce qui se traduit plus simplement à l'aide des exponentielles de matrices $e^{\Gamma x} = \text{Diag}(e^{\gamma_i x})_{1 \leq i \leq n}$ et $e^{-\Gamma x} = \text{Diag}(e^{-\gamma_i x})_{1 \leq i \leq n}$.

$$\mathbf{v}_m = e^{\Gamma x} \lambda + e^{-\Gamma x} \mu \quad (5.10)$$

λ et μ étant les vecteurs composés des constantes d'intégration.

En utilisant les propriétés des exponentielles de matrices, nous pouvons exprimer la dérivée des tensions modales $\mathbf{v}_m$:

$$\frac{d\mathbf{v}_m}{dx} = \Gamma e^{\Gamma x} \lambda - \Gamma e^{-\Gamma x} \mu \quad (5.11)$$

D'autre part, à l'aide des équations (5.2) et (5.7), nous pouvons écrire :

$$\frac{d\mathbf{v}_m}{dx} = P_v^{-1} \frac{d\mathbf{v}}{dx} = -P_v^{-1} \mathbf{z}\mathbf{i} = -P_v^{-1} \mathbf{z} P_i \mathbf{i}_m \quad (5.12)$$

Supposons maintenant que la matrice Γ soit inversible, c'est-à-dire que les γ_i sont non nuls. Nous avons alors la relation suivante :

$$e^{\Gamma x} \lambda - e^{-\Gamma x} \mu = -\Gamma^{-1} P_v^{-1} \mathbf{z} P_i \mathbf{i}_m \quad (5.13)$$

Nous voyons apparaître dans cette dernière expression l'analogue de l'impédance caractéristique pour une ligne monophasée. Cette matrice qu'il convient de nommer *impédance caractéristique modale*, sera notée :

$$\mathbf{Z}_c = \Gamma^{-1} P_v^{-1} \mathbf{z} P_i \quad (5.14)$$

5.2. Développement d'impédances matricielles

Remarque L'hypothèse d'inversibilité de Γ est une hypothèse forte. Elle implique en particulier l'inversibilité de $\mathbf{z}$ et de $\mathbf{y}$. En effet,

$$\det(\Gamma) \neq 0 \Rightarrow \det(\Gamma^2) = \det(\mathbf{zy}) = \det(\mathbf{z}) \det(\mathbf{y}) \neq 0 \quad (5.15)$$

L'impédance caractéristique, en tant que produit de matrices inversibles, est donc elle aussi inversible.

Remarquons que la matrice Γ devient singulière pour des valeurs particulières de p (variable de Laplace) dans le plan complexe (voir la note 2).

5.1.4 Conditions aux limites

Supposons connus les tensions et les courants à l'origine. Les tensions et courants modaux sont alors connus eux aussi par les relations (5.7). Nous déduisons les vecteurs de constantes d'intégration λ et μ des équations (5.10) et (5.13) appliquées en $x = 0$,

$$\begin{cases} \lambda + \mu = \mathbf{v}_m(0) \\ \lambda - \mu = -\mathbf{Z}_c \mathbf{i}_m(0) \end{cases} \quad \begin{cases} \lambda = \frac{1}{2} (\mathbf{v}_m(0) - \mathbf{Z}_c \mathbf{i}_m(0)) \\ \mu = \frac{1}{2} (\mathbf{v}_m(0) + \mathbf{Z}_c \mathbf{i}_m(0)) \end{cases} \quad (5.16)$$

En réécrivant (5.10) et (5.13) en remplaçant λ et μ par leurs expressions, nous faisons apparaître des expressions analogues à celles de la représentation quadripolaire d'une ligne monophasée de longueur X .

$$\mathbf{v}_m(X) = \cosh(\Gamma X) \mathbf{v}_m(0) - \sinh(\Gamma X) \mathbf{Z}_c \mathbf{i}_m(0) \quad (5.17)$$

$$\mathbf{i}_m(X) = -\mathbf{Z}_c^{-1} \sinh(\Gamma X) \mathbf{v}_m(0) + \mathbf{Z}_c^{-1} \cosh(\Gamma X) \mathbf{Z}_c \mathbf{i}_m(0) \quad (5.18)$$

En notant dans cette dernière expression $\cosh(\Gamma X) = \text{Diag}(\cosh(\gamma_i X))_{1 \leq i \leq n}$ et de même, $\sinh(\Gamma X) = \text{Diag}(\sinh(\gamma_i X))_{1 \leq i \leq n}$.

5.2 Développement d'impédances matricielles

Nous utilisons encore une fois les relations (5.7) pour transformer les relations de type « quadripolaire » précédentes de façon à ne faire intervenir que les tensions et les courants de ligne (ceux qui sont directement mesurables).

$$\mathbf{v}(X) = P_v \cosh(\Gamma X) P_v^{-1} \mathbf{v}(0) - P_v \sinh(\Gamma X) \mathbf{Z}_c P_i^{-1} \mathbf{i}(0) \quad (5.19)$$

$$\mathbf{i}(X) = -P_i \mathbf{Z}_c^{-1} \sinh(\Gamma X) P_v^{-1} \mathbf{v}(0) + P_i \mathbf{Z}_c^{-1} \cosh(\Gamma X) \mathbf{Z}_c P_i^{-1} \mathbf{i}(0) \quad (5.20)$$

Nous allons chercher à représenter ce quadripôle « matriciel » par un réseau en Π . Pour cela, nous devons nous donner une définition rigoureuse des *impédances* (ou *admittances*) *matricielles*.

5.2.1 Définition d'une impédances matricielles

Nous appellerons impédances matricielles des éléments électriques à $2n$ bornes tels celui de la figure 5.2 pouvant être décrits par une relation du type

$$\mathbf{v} = \mathbf{Z}\mathbf{i}, \text{ avec } \mathbf{v} = \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix}, \text{ et } \mathbf{i} = \begin{pmatrix} i_1 \\ i_2 \\ \vdots \\ i_n \end{pmatrix}. \quad (5.21)$$

Nous utiliserons pour représenter un tel élément le schéma simplifié apparaissant sur la figure 5.2. De même une admittance¹ matricielle sera caractérisée par une relation telle que

$$\mathbf{i} = \mathbf{Y}\mathbf{v}. \quad (5.22)$$

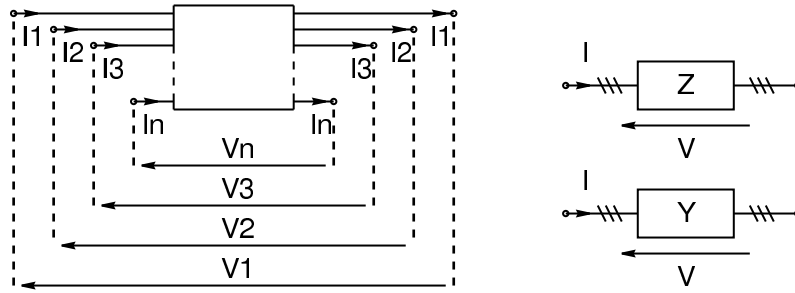


FIG. 5.2 – Élément électrique à $2n$ bornes.

Remarque Il est important de distinguer l'appellation « impédance matricielle » de la « matrice d'impédance » d'un réseau. On pourra se reporter à l'annexe E pour voir des exemples illustrant ces deux notions.

Ces éléments à $2n$ bornes obéissent à des lois de composition similaires à celles des dipôles. Nous allons le vérifier en calculant les impédances et les admittances équivalentes aux associations des figures 5.3 et 5.4. Pour une association en série, les relations suivantes sont vérifiées :

$$\mathbf{v} = \mathbf{v}_1 + \mathbf{v}_2 = \mathbf{Z}_1\mathbf{i} + \mathbf{Z}_2\mathbf{i} = \underbrace{(\mathbf{Z}_1 + \mathbf{Z}_2)}_{\mathbf{Z}_{eq}}\mathbf{i} \quad (5.23)$$

Et pour une association en parallèle :

$$\mathbf{i} = \mathbf{i}_1 + \mathbf{i}_2 = \mathbf{Y}_1\mathbf{v} + \mathbf{Y}_2\mathbf{v} = \underbrace{(\mathbf{Y}_1 + \mathbf{Y}_2)}_{\mathbf{Y}_{eq}}\mathbf{v} \quad (5.24)$$

¹Évidemment, si $\mathbf{Z}$ est inversible, nous pouvons écrire $\mathbf{Y} = \mathbf{Z}^{-1}$.

5.2. Développement d'impédances matricielles

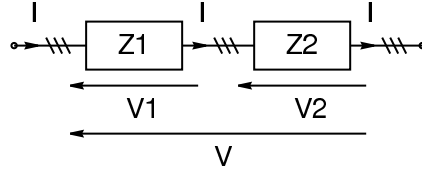


FIG. 5.3 – Associations en série d'impédances matricielles.

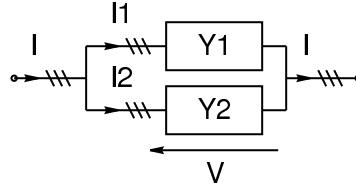
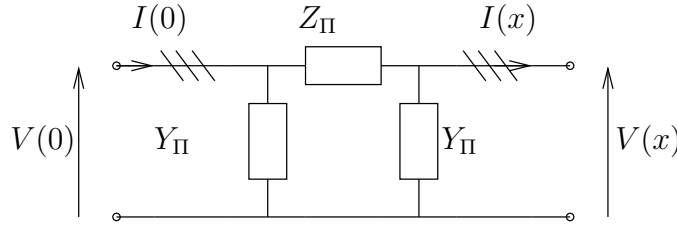


FIG. 5.4 – Associations en parallèle d'admittances matricielles.

5.2.2 Modèle en « Π » équivalent



$$\mathbf{Z}_{\Pi} = P_v \sinh(\Gamma X) \mathbf{Z}_c P_i^{-1} \quad (5.25)$$

$$\mathbf{Y}_{\Pi} = P_i \mathbf{Z}_c^{-1} \tanh\left(\Gamma \frac{X}{2}\right) P_v^{-1} \quad (5.26)$$

FIG. 5.5 – Schéma en Π multifilaire.

Considérons le schéma de la figure 5.5, constitué d'impédances matricielles. Nous déduisons les relations (matricielles) suivantes entre tensions et courants (vectoriels).

$$\mathbf{v}(X) = \mathbf{v}(0) - \mathbf{Z}_{\Pi}(\mathbf{i}(0) - \mathbf{Y}_{\Pi}\mathbf{v}(0)) = (\text{Id} + \mathbf{Z}_{\Pi}\mathbf{Y}_{\Pi})\mathbf{v}(0) - \mathbf{Z}_{\Pi}\mathbf{i}(0), \quad (5.27)$$

$$\mathbf{i}(X) = \mathbf{i}(0) - \mathbf{Y}_{\Pi}\mathbf{v}(0) - \mathbf{Y}_{\Pi}\mathbf{v} = -(2\mathbf{Y}_{\Pi} + \mathbf{Y}_{\Pi}\mathbf{Z}_{\Pi}\mathbf{Y}_{\Pi})\mathbf{v}(0) + (\text{Id} + \mathbf{Y}_{\Pi}\mathbf{Z}_{\Pi})\mathbf{i}(0). \quad (5.28)$$

Nous voulons que ce modèle soit équivalent à la ligne polyphasée. Il suffit pour cela de choisir $\mathbf{Y}_{\Pi}$ et $\mathbf{Z}_{\Pi}$ de façon à vérifier les équations (5.20). Rapprochons tout d'abord les équations donnant la valeur de $\mathbf{v}$. Par identification termes à termes nous obtenons :

$$\mathbf{Z}_{\Pi} = P_v \sinh(\Gamma X) \mathbf{Z}_c P_i^{-1}, \quad (5.29)$$

$$\text{Id} + \mathbf{Z}_{\Pi} \mathbf{Y}_{\Pi} = P_v \cosh(\Gamma X) P_v^{-1}. \quad (5.30)$$

Supposons maintenant que $\sinh(\Gamma X)$ soit inversible. Cela se traduit par le fait que $\sinh(\gamma_i X)$ est non nul, quel que soit i ; ou encore que nous nous plaçons sur un domaine du plan complexe ne contenant pas de zéros de ces fonctions². La matrice impédance $\mathbf{Z}_{\Pi}$ est donc inversible, d'inverse :

$$\mathbf{Z}_{\Pi}^{-1} = P_i \mathbf{Z}_c^{-1} (\sinh(\Gamma X))^{-1} P_v^{-1} \quad (5.31)$$

La matrice $(\sinh(\Gamma X))^{-1}$ étant la matrice diagonale $\text{Diag} \left(\frac{1}{\sinh(\gamma_i X)} \right)_{1 \leq i \leq n}$. La multiplication à gauche de (5.30) par $\mathbf{Z}_{\Pi}^{-1}$ nous fournit une expression de $\mathbf{Y}_{\Pi}$.

$$\begin{aligned} \mathbf{Y}_{\Pi} &= \mathbf{Z}_{\Pi}^{-1} P_v \cosh(\Gamma X) P_v^{-1} + \mathbf{Z}_{\Pi}^{-1} \\ &= P_i \mathbf{Z}_c^{-1} (\sinh(\Gamma X))^{-1} \cosh(\Gamma X) P_v^{-1} - P_i \mathbf{Z}_c^{-1} (\sinh(\Gamma X))^{-1} P_v^{-1} \\ &= P_i \mathbf{Z}_c^{-1} (\sinh(\Gamma X))^{-1} \cosh(\Gamma X) - \sinh(\Gamma X)^{-1} P_v^{-1}. \end{aligned} \quad (5.32)$$

Puisque les matrices $\sinh(\Gamma X)^{-1}$ et $\cosh(\Gamma X)$ sont diagonales, et par utilisation des formules de trigonométrie hyperbolique, il reste au final :

$$\mathbf{Y}_{\Pi} = P_i \mathbf{Z}_c^{-1} \tanh \left(\Gamma \frac{X}{2} \right) P_v^{-1}. \quad (5.33)$$

Il reste alors à vérifier que les expressions trouvées pour $\mathbf{Z}_{\Pi}$ et $\mathbf{Y}_{\Pi}$ fournissent bien l'expression de (5.20) pour le vecteur courant $\mathbf{i}$. Ce qui revient à vérifier les relations :

$$2\mathbf{Y}_{\Pi} + \mathbf{Y}_{\Pi} \mathbf{Z}_{\Pi} \mathbf{Y}_{\Pi} = P_i \mathbf{Z}_c^{-1} \sinh(\Gamma X) P_v^{-1}, \quad (5.34)$$

$$\text{Id} + \mathbf{Y}_{\Pi} \mathbf{Z}_{\Pi} = P_i \mathbf{Z}_c^{-1} \cosh(\Gamma X) \mathbf{Z}_c P_i^{-1}. \quad (5.35)$$

Cela se fait en utilisant les équations (5.30) pour l'expression de $\mathbf{Z}_{\Pi}$ et (5.33) pour l'expression de $\mathbf{Y}_{\Pi}$ et les relations trigonométriques :

$$2 \tanh \left(\Gamma \frac{X}{2} \right) + \sinh(\Gamma X) \tanh^2 \left(\Gamma \frac{X}{2} \right) = \sinh(\Gamma X), \quad (5.36)$$

$$\text{Id} + \sinh(\Gamma X) \tanh \left(\Gamma \frac{X}{2} \right) = \cosh(\Gamma X). \quad (5.37)$$

5.2.3 Fractions continues matricielles

Nous allons utiliser les résultats des développements en fractions continues des fonctions (chapitre 3) et l'adapter au cas des impédances et des admittances matricielles qui constituent le modèle en Π de la ligne multifilaire. En effet, considérons la matrice

² Physiquement, nous excluons les valeurs particulières de p correspondant aux résonances ou anti-résonances modales.

5.2. Développement d'impédances matricielles

$\sinh(\Gamma X)$. Nous pouvons développer chacun de ses termes en fraction continue, puis utiliser des relations d'inversion triviales pour des matrices diagonales.

$$\sinh(\Gamma X) = \text{Diag}(\sinh(\gamma_i X))_{1 \leq i \leq n} \quad (5.38)$$

$$\sinh(\Gamma X) = \text{Diag} \left(\frac{1}{\frac{1}{\gamma_i X} + \frac{1}{-\frac{6}{\gamma_i X} + \frac{1}{-\frac{10}{7\gamma_i X} + \dots}}} \right)_{1 \leq i \leq n} \quad (5.39)$$

$$\sinh(\Gamma X) = ((\Gamma X)^{-1} + (-6(\Gamma X)^{-1} + (-\frac{10}{7}(\Gamma X)^{-1} + \dots)^{-1})^{-1})^{-1} \quad (5.40)$$

Introduisons ce développement dans l'expression de $\mathbf{Z}_\Pi$ (équation (5.30)), puis explicitons $\mathbf{Z}_c$.

$$\begin{aligned} \mathbf{Z}_\Pi &= P_v((\Gamma X)^{-1} + (-6(\Gamma X)^{-1} + (-\frac{10}{7}(\Gamma X)^{-1} + \dots)^{-1})^{-1})^{-1} \mathbf{Z}_c P_i^{-1} \\ &= P_v((\Gamma X)^{-1} + (-6(\Gamma X)^{-1} + (-\frac{10}{7}(\Gamma X)^{-1} + \dots)^{-1})^{-1})^{-1} \Gamma^{-1} P_v^{-1} \mathbf{z} \end{aligned} \quad (5.41)$$

Par simple utilisation des règles d'inversion matricielle, nous pouvons simplifier cette dernière expression. Remarquons en effet la relation suivante :

$$\mathbf{z}^{-1} P_v \Gamma (\Gamma X)^{-1} P_v^{-1} = (\mathbf{z} X)^{-1} \quad (5.42)$$

Et d'autre part, en utilisant l'équation (5.6)

$$P_v(\Gamma X)^{-1} P_v^{-1} \mathbf{z} = (\mathbf{z}^{-1} P_v \Gamma^2 P_v^{-1} X)^{-1} = (\mathbf{z}^{-1} \mathbf{z} \mathbf{y} X)^{-1} = (\mathbf{y} X)^{-1} \quad (5.43)$$

Ces deux relations sont à exploiter alternativement, comme suit.

$$\begin{aligned} \mathbf{Z}_\Pi &= (\mathbf{z}^{-1} P_v \Gamma (\Gamma X)^{-1} P_v^{-1} + \mathbf{z}^{-1} P_v \Gamma (-6(\Gamma X)^{-1} + (-\frac{10}{7}(\Gamma X)^{-1} + \dots)^{-1})^{-1} P_v^{-1})^{-1} \\ &= ((\mathbf{z} X)^{-1} + (-6 P_v(\Gamma X)^{-1} \Gamma^{-1} P_v^{-1} \mathbf{z} + P_v(-\frac{10}{7}(\Gamma X)^{-1} + \dots)^{-1} \Gamma^{-1} P_v^{-1} \mathbf{z})^{-1})^{-1} \\ &= ((\mathbf{z} X)^{-1} + (-6(\mathbf{y} X)^{-1} + (-\frac{10}{7} \mathbf{z}^{-1} P_v \Gamma (\Gamma X)^{-1} P_v^{-1} + \dots)^{-1})^{-1})^{-1} \\ &= ((\mathbf{z} X)^{-1} + (-6(\mathbf{y} X)^{-1} + (-\frac{10}{7}(\mathbf{z} X)^{-1} + \dots)^{-1})^{-1})^{-1} \end{aligned} \quad (5.44)$$

De la même façon, nous pouvons développer la matrice admittance $\mathbf{Y}_\Pi$ sous la forme :

$$\begin{aligned} \mathbf{Y}_\Pi &= P_i \mathbf{Z}_c^{-1} (2(\Gamma X)^{-1} + (6(\Gamma X)^{-1} + (10(\Gamma X)^{-1} + \dots)^{-1})^{-1})^{-1} P_v^{-1} \\ &= (2(\mathbf{y} X)^{-1} + (6(\mathbf{z} X)^{-1} + (10(\mathbf{y} X)^{-1} + \dots)^{-1})^{-1})^{-1} \end{aligned} \quad (5.45)$$

Donc, au final, avec $\mathbf{Z} = \mathbf{z} X$ et $\mathbf{Y} = \mathbf{y} X$,

$$\mathbf{Z}_\Pi = (\mathbf{Z}^{-1} + (-6\mathbf{Y}^{-1} + (-\frac{10}{7}\mathbf{Z}^{-1} + \dots)^{-1})^{-1})^{-1}, \quad (5.46)$$

$$\mathbf{Y}_\Pi = (2\mathbf{Y}^{-1} + (6\mathbf{Z}^{-1} + (10\mathbf{Y}^{-1} + \dots)^{-1})^{-1})^{-1}. \quad (5.47)$$

Nous reconnaissons là des développements ayant une structure tout à fait similaire à celle d'une fraction continue. Nous pouvons donc les représenter par des réseaux en échelle constitués d'impédances matricielles.

5.3 Synthèse de Cauer matricielle

5.3.1 Réseau d'impédances matricielles

Nous allons exploiter les développements précédents pour obtenir des modèles sous forme de réseau. Vu les résultats de la section 5.2.1 nous pouvons identifier les développements précédents à l'impédance équivalente d'un « réseau multifilaire en échelle », c'est-à-dire ayant une structure telle que celle présentée sur la figure 5.6.

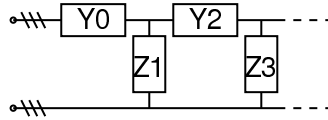


FIG. 5.6 – Réseau multifilaire en échelle.

Les modèles approchés seront obtenus à partir de cette structure en tronquant le réseau.

$$\begin{aligned} \mathcal{Z} &= \mathbf{Y}_0^{-1} + (\mathbf{Z}_1^{-1} + (\mathbf{Y}_2^{-1} + (\mathbf{Z}_3^{-1} + \dots)^{-1})^{-1})^{-1} \\ &\approx \mathbf{Y}_0^{-1} + (\mathbf{Z}_1^{-1} + (\mathbf{Y}_2^{-1} + (\mathbf{Z}_3^{-1} + \dots + \mathbf{Y}_N)^{-1})^{-1})^{-1} \end{aligned} \quad (5.48)$$

5.3.2 Représentation du couplage capacitif

La présence de capacités interphases ne permet pas d'exprimer les tensions en fonction de courants par une relation matricielle correspondant à une admittance matricielle. Nous définissons le couplage capacitif par *dualité* par rapport au couplage inductif (voir annexe E). Ceci revient à modifier le schéma de la figure 5.1 pour le remplacer par celui de la figure 5.7, ceci *sans modifier les relations* (5.2).

Nous pouvons alors appliquer les lois de composition « classiques » aux impédances et admittances matricielles de la figure 5.8 (association en série ou en parallèle).

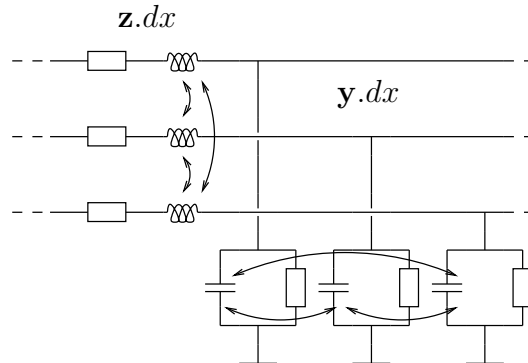


FIG. 5.7 – Élément de ligne triphasée avec couplage capacitif.

5.3. Synthèse de Cauer matricielle

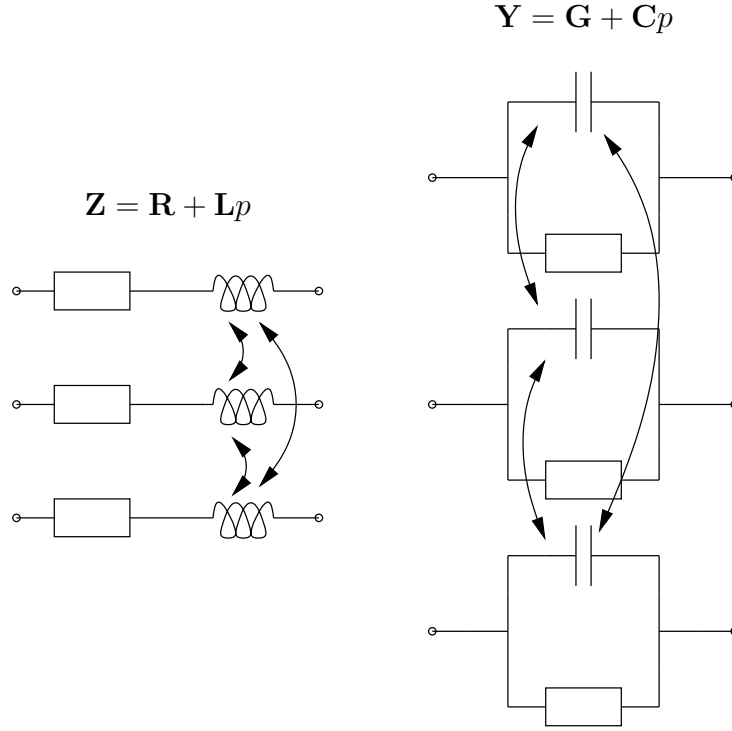


FIG. 5.8 – Impédances et admittances matricielles triphasées.

5.3.3 Représentation d'une résistance de défaut triphasée

Le schéma de la figure 5.9 permet de définir une résistance de charge quelconque d'une ligne triphasée. Elle se représente par sa matrice d'admittance (annexe E) $\mathbf{R}_{\text{def}}$ qui est une matrice 3×3 .

Selon la valeur de ses éléments, cette matrice permet de modéliser aussi bien :

- les défauts monophasés ou phase-neutre lorsque R_{10} , R_{20} ou R_{30} deviennent très faibles ;
- les défauts biphasés, lorsque R_{12} , ou bien R_{23} , ou bien R_{31} devient très faibles ;
- les défauts triphasés isolés, lorsque R_{12} , R_{23} , et R_{31} devient simultanément très faibles, ou triphasés à la terre lorsque de plus R_{10} , R_{20} et R_{30} deviennent très faibles.

L'impédance d'entrée du modèle en Π est alors

$$\mathcal{Z} = (\mathbf{Y}_{\Pi} + (\mathbf{Z}_{\Pi} + (\mathbf{Y}_{\Pi} + \mathbf{R}_{\text{def}}^{-1})^{-1})^{-1})^{-1}, \quad (5.49)$$

avec

$$\mathbf{R}_{\text{def}}^{-1} = \begin{pmatrix} \frac{1}{R_{10}} + \frac{1}{R_{12}} + \frac{1}{R_{31}} & -\frac{1}{R_{12}} & -\frac{1}{R_{31}} \\ -\frac{1}{R_{12}} & \frac{1}{R_{20}} + \frac{1}{R_{23}} + \frac{1}{R_{12}} & -\frac{1}{R_{23}} \\ -\frac{1}{R_{31}} & -\frac{1}{R_{23}} & \frac{1}{R_{30}} + \frac{1}{R_{31}} + \frac{1}{R_{23}} \end{pmatrix}. \quad (5.50)$$

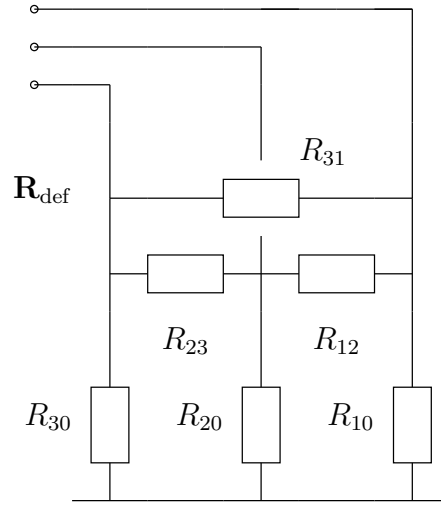


FIG. 5.9 – Résistance de défaut triphasée.

5.4 Exemple de simulation d'une ligne triphasée

Nous pouvons dès lors utiliser ces résultats pour obtenir un modèle approché d'une ligne triphasée utilisable avec un simulateur de réseaux. Considérons le cas (idéal) d'une ligne triphasée totalement symétrique. De plus, nous supposons pour simplifier que $G = 0$ et qu'il n'y a pas de couplage capacitif (le couplage capacitif n'existe pas sous *Spice* ou *Esacap*). Les matrices de constantes linéaires s'écrivent :

$$\mathbf{z} = \begin{pmatrix} r + lp & mp & mp \\ mp & r + lp & mp \\ mp & mp & r + lp \end{pmatrix} \quad \mathbf{y} = \begin{pmatrix} cp & 0 & 0 \\ 0 & cp & 0 \\ 0 & 0 & cp \end{pmatrix} \quad (5.51)$$

Les valeurs numériques utilisées pour la simulation sont données dans le tableau 5.1.

Paramètre	Symbole	Valeur
longueur	X	300 km
résistance linéique	r	$3 \cdot 10^{-2} \Omega \text{km}^{-1}$
inductance linéique	l	$1,05 \cdot 10^{-3} \text{Hkm}^{-1}$
inductance mutuelle linéique	m	$4 \cdot 10^{-4} \text{Hkm}^{-1}$
conductance linéique	g	$\sim 0 \text{Skm}^{-1}$
capacité linéique	c	$11 \cdot 10^{-9} \text{Fkm}^{-1}$

TAB. 5.1 – Valeurs numériques du modèle de ligne triphasée.

5.4. Exemple de simulation d'une ligne triphasée

5.4.1 Méthode modale

Dans ce cas simple, les produits de matrices $\mathbf{zy}$ et $\mathbf{yz}$ peuvent se diagonaliser à l'aide des matrices de changement de base

$$P_v = P_i = \begin{pmatrix} 1 & 1 & 1 \\ 1 & e^{-\frac{2i\Pi}{3}} & e^{\frac{2i\Pi}{3}} \\ 1 & e^{\frac{2i\Pi}{3}} & e^{-\frac{2i\Pi}{3}} \end{pmatrix} \quad (5.52)$$

$$P_v^{-1} = P_i^{-1} = \frac{1}{3} \begin{pmatrix} 1 & 1 & 1 \\ 1 & e^{\frac{2i\Pi}{3}} & e^{-\frac{2i\Pi}{3}} \\ 1 & e^{-\frac{2i\Pi}{3}} & e^{\frac{2i\Pi}{3}} \end{pmatrix} \quad (5.53)$$

Les composantes modales correspondantes sont connues sous le nom de *composantes symétriques* de Fortescue :

- composante *homopolaire* (indice $_0$),
- composante *directe* (indice $_d$),
- composante *inverse* (indice $_i$).

Remarquons que P_v (P_i) diagonalise également $\mathbf{z}$ et $\mathbf{y}$. Les équations (5.2) se réduisent alors dans la base modale à la forme :

$$\frac{d}{dx} \begin{pmatrix} v_0 \\ v_d \\ v_i \end{pmatrix} = - \begin{pmatrix} r + (l + 2m)p & 0 & 0 \\ 0 & r + (l - m)p & 0 \\ 0 & 0 & r + (l - m)p \end{pmatrix} \begin{pmatrix} i_0 \\ i_d \\ i_i \end{pmatrix} \quad (5.54)$$

$$\frac{d}{dx} \begin{pmatrix} i_0 \\ i_d \\ i_i \end{pmatrix} = - \begin{pmatrix} cp & 0 & 0 \\ 0 & cp & 0 \\ 0 & 0 & cp \end{pmatrix} \begin{pmatrix} v_0 \\ v_d \\ v_i \end{pmatrix} \quad (5.55)$$

C'est-à-dire que chaque composante modale peut être traitée à l'aide d'un modèle quadripolaire monophasé, à l'aide des valeurs

$$z_0 = r + (l + 2m)p, \quad z_d = z_i = r + (l - m)p, \quad (5.56)$$

$$y_0 = y_d = y_i = cp. \quad (5.57)$$

Cette remarque nous permet de calculer par transformée de Laplace inverse les courants de référence. Pour cela, nous supposons la ligne triphasée fermée sur des résistances de $1 \text{ k}\Omega$, choisie pour bien mettre en évidence les réflexions d'ondes.

$$R_{10} = R_{20} = R_{30} = 1 \text{ k}\Omega, \quad (5.58)$$

$$R_{12} = R_{23} = R_{31} = \infty. \quad (5.59)$$

Un seul conducteur étant alimenté par une source de tension, les deux autres étant mis à la terre. Nous calculons la réponse en courant pour chacun des modes par inversion numérique (IFFT). Les courants de lignes sont ensuite obtenus par multiplication par la matrice P_i . Ces résultats sont présentés sur la figure 5.11.

5.4.2 Modèle approché de la ligne polyphasée

Les développements des impédances $\mathbf{Z}_{\Pi}$ et $\mathbf{Y}_{\Pi}$ en fractions continues matricielles permettent d'obtenir le modèle approché de la figure 5.10, tronqué à l'ordre quatre. Ce modèle utilisé avec un simulateur de réseau (*Esacap*) permet d'obtenir les courants de la figure 5.12.

La comparaison de ces deux réponses temporelles nous permet de vérifier la qualité de notre modèle pour la simulation des courants absorbés. Les phénomènes de couplage sont bien reproduits, puisque les formes d'ondes correspondent, aussi bien pour le conducteur alimenté que pour les conducteurs mis à la terre.

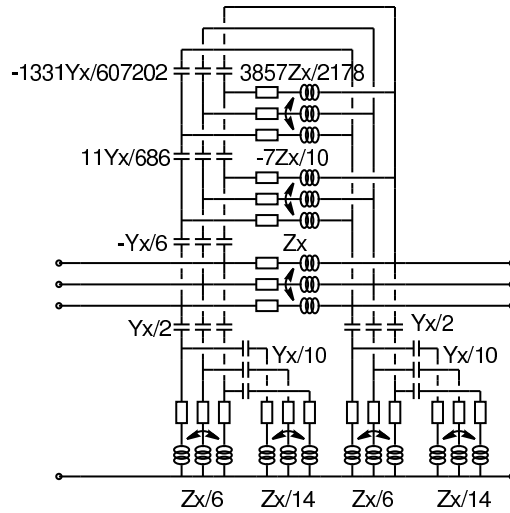


FIG. 5.10 – Modèle approché de la ligne triphasée.

5.5 Avantages du développement d'impédances matricielles

Une démarche se rapprochant de la notre peut être trouvée dans [LK92]. Il s'agit de calculer un approximant de Padé de chaque composante modale de l'impédance, puis de calculer par transformée inverse (annexe A) la réponse impulsionnelle approchée de chaque mode. La simulation de la réponse du système est alors effectuée de manière récursive par une convolution. Cette méthode, qui donne d'excellents résultats, a été implémentée dans le logiciel de simulation de réseau *Kspice*³, dérivé de *Spice* de l'université de Berkeley.

Cependant, comme cette méthode nécessite le calcul préalable des réponses impulsionnelles correspondant à chaque mode, elle est inadaptée à une brusque variation des

³<http://www.eecs.berkeley.edu/IPRO/Software/Catalog/Description/kspice.html>

5.5. Avantages du développement d'impédances matricielles

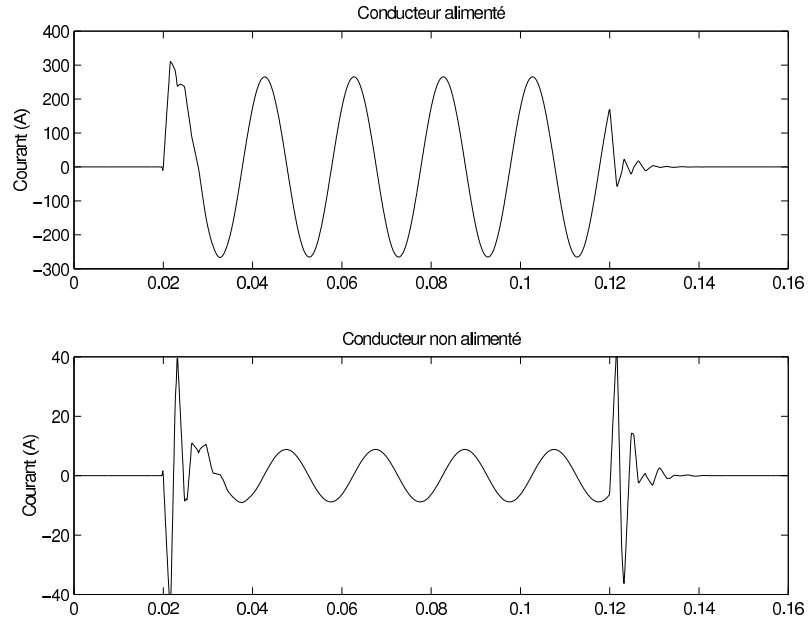


FIG. 5.11 – Courants de ligne obtenus par décomposition modale et transformée de Laplace inverse.

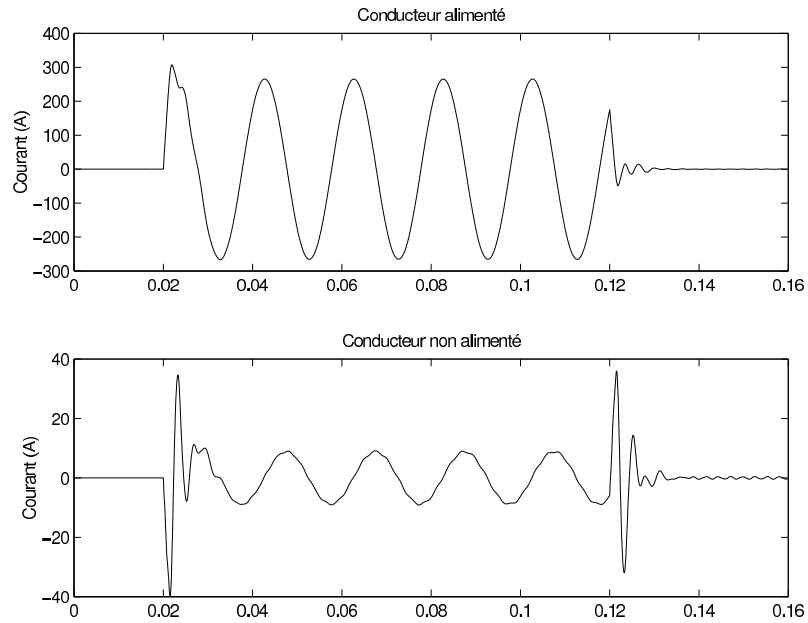


FIG. 5.12 – Courants de ligne obtenus par utilisation du modèle approché avec un simulateur de réseau.

paramètres dont dépend le calcul de ces modes. Notre méthode ne souffre pas de cette limitation puisque le modèle utilise directement les composantes de ligne et non les composantes modales des tensions et des courants.

Il est en effet remarquable que les composantes modales des impédances n'interviennent qu'en tant qu'étape dans le calcul du développement en fraction continue matricielle. Il n'est donc pas en pratique nécessaire d'effectuer de diagonalisation pour calculer le modèle. Seules sont nécessaires les valeurs des paramètres globaux de la ligne, c'est-à-dire les matrices $\mathbf{R}$, $\mathbf{L}$, $\mathbf{C}$ et $\mathbf{G}$, tout comme dans le cas monophasé.

C'est un avantage déterminant car le calcul des composantes modales peut être sujet à caution. Les composantes de Fortescue dépendent notamment d'hypothèses géométriques fausses dans le cas général pour les lignes de transmission, comme le fait remarquer M. Pélissier :

« Donc si les trois tensions imposées à la ligne sont des tensions sinusoïdales formant un système symétrique direct ou inverse ou homopolaire, [...] les courants ne pourront pas former un système semblable. Les systèmes de coordonnées symétriques sont des « modes de génération » mais ne constituent pas des « modes de propagation ». [Pél75, ch. 2] ».

Nous avons donc réussi à trouver une méthode innovante de modélisation des lignes triphasées qui fournit des modèles fidèles tout en restant simples, ce qui est un avantage pour l'utilisation en temps réel⁴. Les paramètres du système apparaissent explicitement dans les modèles, ce qui facilite leur identification, et enfin les modèles utilisent les composantes de lignes issues directement des mesures, ce qui est un avantage pour la mise en œuvre dans une protection numérique industrielle.

De tels modèles peuvent-ils être utilisés pour la détection de défauts? C'est ce que nous allons étudier au chapitre suivant.

⁴Remarquons, que sans cette contrainte de rapidité de calcul, et puisque l'on contrôle l'ordre du développement, on peut choisir des modèles avec des ordres bien plus grands pour simuler de façon encore plus précise la réponse des lignes; une application possible sont les logiciels de simulation de réseaux.

5.5. Avantages du développement d'impédances matricielles

Chapitre 6

Application à la détection de défauts

Nous avons vu dans le chapitre 1 qu'un algorithme de protection de distance pouvait se définir comme une méthode d'identification paramétrique d'une ligne de transmission. Nous allons donc utiliser les modèles approchés obtenus par la synthèse de Cauer pour identifier les paramètres d'un défaut à partir des mesures simulées des tensions et des courants au niveau d'un poste d'interconnexion.

6.1 Modèle numérique d'une ligne en défaut

Le modèle quadripolaire d'une ligne en défaut est celui qui découle de la résolution de l'équation des télégraphistes (section 5.2.2). En assimilant le défaut à une simple résistance, nous obtenons le modèle de référence de la figure 5.5.

6.1.1 Synthèse de Cauer

Par développement des impédances d'ordre infini et synthèse de Cauer, nous obtenons des modèles approchés à différents ordres (fig. 6.1 à 6.4).

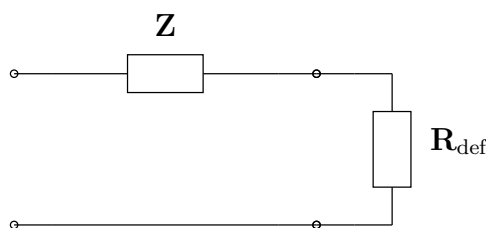


FIG. 6.1 – Modèle de ligne en défaut approché utilisé dans l'algorithme PXLN (*Alstom*).

Pour pouvoir utiliser ces modèles avec un algorithme d'optimisation programmé sous *Matlab*, il est nécessaire de procéder au calcul de la forme numérique des modèles, ou modèles ARMA¹.

¹Modèle auto-régressif à moyenne mobile (*Auto-regressive moving-average*)

6.1. Modèle numérique d'une ligne en défaut

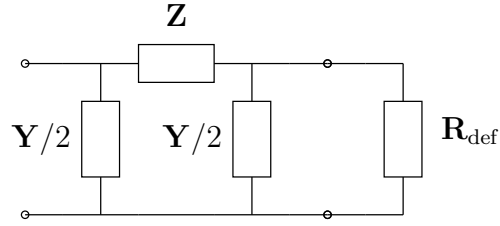


FIG. 6.2 – Modèle de ligne en défaut approché à l'ordre un (modèle de ligne courte).

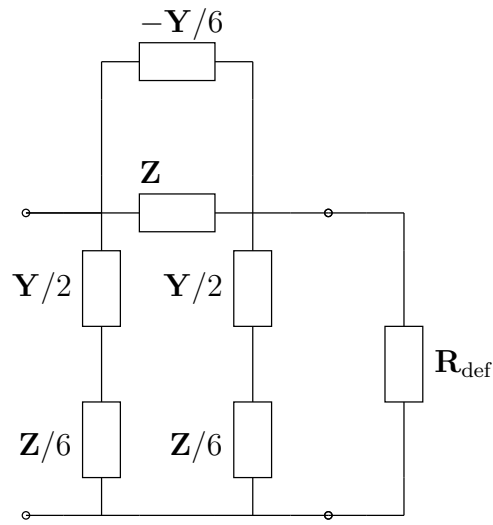


FIG. 6.3 – Modèle de ligne en défaut approché à l'ordre deux.

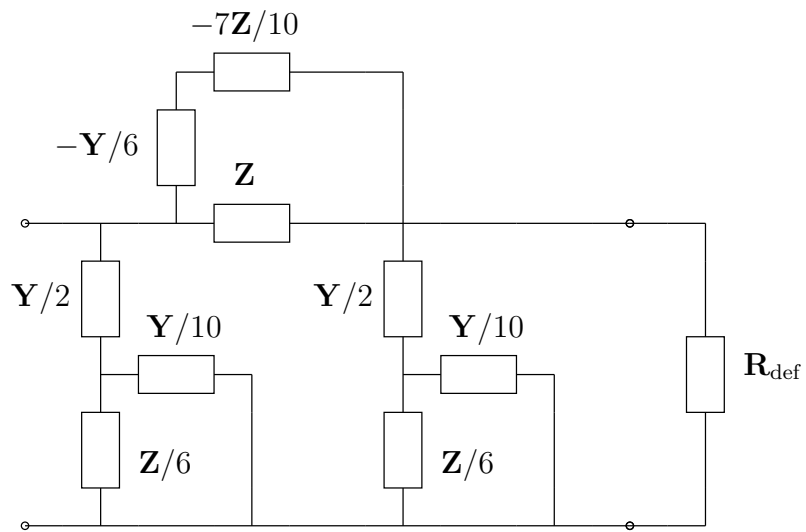


FIG. 6.4 – Modèle de ligne en défaut approché à l'ordre trois.

6.1.2 Modèle numérique monophasé

Pour simplifier la description de la méthode de numérisation nous commençons par utiliser un modèle monophasé.

L'entrée du modèle est le courant mesuré i et nous devons calculer numériquement la tension estimée $\hat{v}$ comme la sortie du modèle approché. La fonction de transfert s'identifie donc à l'impédance d'entrée du modèle de ligne $\hat{Z}$.

$$i(p) \rightarrow \boxed{\hat{Z}(p)} \rightarrow \hat{v}(p) = \hat{Z}i(p) \quad (6.1)$$

Les modèles obtenus par synthèse de Cauer sont d'ordre fini, du fait de la troncature des fractions continues. Ce sont des modèles d'ordres réduits. Leurs impédances équivalentes s'expriment donc par une fraction rationnelle, soit

$$\hat{Z}(p) = \frac{P(p)}{Q(p)} \quad (6.2)$$

où P et Q sont des polynômes.

Nous discrétisons ensuite ce modèle par *transformation bilinéaire*. Fixons la fréquence d'échantillonnage f_e ,

$$p \rightarrow 2f_e \frac{z-1}{z+1} \quad (6.3)$$

La transformée en z du modèle discret s'écrit donc

$$\frac{P(2f_e \frac{z-1}{z+1})}{Q(2f_e \frac{z-1}{z+1})} = \frac{b_0 + b_1 z^{-1} + \dots + b_n z^{-N}}{1 + a_1 z^{-1} + \dots + a_n z^{-N}} = \frac{v(z)}{i(z)} \quad (6.4)$$

La transformée en z (TZ) permet d'obtenir directement les coefficients du modèle ARMA. En effet, pour une suite quelconque (x_n) elle vérifie la propriété fondamentale

$$\text{TZ}(x_{n-N}) = z^{-N} x(z). \quad (6.5)$$

Ce qui fournit la structure du modèle sous forme auto-régressive :

$$v(n) = b_0 i(n) + b_1 i(n-1) + \dots - b_N i(n-N) - a_1 v(n-1) - \dots - a_N v(n-N) \quad (6.6)$$

Ou encore, sous forme vectorielle,

$$\mathbf{a} = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_N \end{pmatrix} \quad \mathbf{b} = \begin{pmatrix} b_0 \\ b_1 \\ \vdots \\ b_N \end{pmatrix} \quad \mathbf{v} = \begin{pmatrix} v(n-1) \\ v(n-2) \\ \vdots \\ v(n-N) \end{pmatrix} \quad \mathbf{i} = \begin{pmatrix} i(n) \\ i(n-1) \\ \vdots \\ i(n-N) \end{pmatrix} \quad (6.7)$$

$$v = \mathbf{b}^T \mathbf{i} - \mathbf{a}^T \mathbf{v}. \quad (6.8)$$

6.2. Méthode d'identification paramétrique

6.1.3 Modèle numérique polyphasé

La numérisation d'un modèle polyphasé ne diffère que peu du cas monphasé. Simplement, l'impédance matricielle s'identifie à la matrice de transfert d'un système MIMO². C'est-à-dire que $\hat{\mathbf{Z}}(p)$ est une matrice dont les éléments sont des fractions rationnelles de p (pour les modèles d'ordre finis).

On montre qu'on peut toujours exprimer la matrice de transfert du modèle approché sous la forme

$$\hat{\mathbf{Z}}(p) = \frac{\mathbf{P}(p)}{Q(p)} \quad (6.9)$$

où $\mathbf{P}$ est un polynôme de matrices et Q un polynôme scalaire, calculé comme le plus petit multiple commun (PPCM) des dénominateurs des éléments de $\hat{\mathbf{Z}}(p)$ [Wol74, Kai80, ch. 6].

Comme précédemment, ce modèle se discrétise par la *transformation bilinéaire*.

$$\frac{\mathbf{P}(2f_e \frac{z-1}{z+1})}{Q(2f_e \frac{z-1}{z+1})} = \frac{\mathbf{b}_0 + \mathbf{b}_1 z^{-1} + \dots + \mathbf{b}_N z^{-N}}{1 + a_1 z^{-1} + \dots + a_N z^{-N}} = \frac{v(z)}{i(z)} \quad (6.10)$$

D'où les coefficients $\mathbf{b}_k$ matriciels et a_k scalaires du modèle ARMA

$$\mathbf{v}(n) = \mathbf{b}_0 \mathbf{i}(n) + \mathbf{b}_1 \mathbf{i}(n-1) + \dots - \mathbf{b}_N \mathbf{i}(n-N) - a_1 \mathbf{v}(n-1) - \dots - a_N \mathbf{v}(n-N) \quad (6.11)$$

6.2 Méthode d'identification paramétrique

L'identification des paramètres du modèle de ligne est un problème d'optimisation non linéaire. C'est-à-dire que la réponse $\mathbf{v}$ n'est pas une fonction linéaire³ des paramètres $(X, \mathbf{R}_{\text{def}})$. L'algorithme de Levenberg-Marquardt est une méthode d'optimisation qui fonctionne très bien dans la pratique, est qui est devenu le standard pour les problèmes non linéaires [Men99]. Nous allons voir qu'il associe les avantages des méthodes du gradient et de Gauss-Newton [Mic92, SS89, WP94].

Ces méthodes d'optimisation visent à minimiser un critère quadratique $J(\theta)$ que nous prendrons égal au carré de la norme de la différence entre la tension mesurée $\mathbf{v}(n)$ (sortie du système réel) et la tension prédite par le modèle approché $\hat{\mathbf{v}}(n)$.

$$J(\theta) = (\hat{\mathbf{v}}(n) - \mathbf{v}(n))^T (\hat{\mathbf{v}}(n) - \mathbf{v}(n)). \quad (6.12)$$

6.2.1 Méthode du gradient

La méthode du gradient, ou de la plus forte pente consiste à faire varier les paramètres dans proportionnellement au gradient de la fonction de coût J , et dans la direction opposée à celui-ci.

$$\theta(n+1) = \theta(n) - \lambda \nabla J(\theta(n)) \quad (6.13)$$

²Plusieurs entrées, plusieurs sorties (*Multiple inputs multiple outputs*)

³Cependant le système *est* linéaire, dans le sens où la sortie est une fonction linéaire de l'entrée. C'est une des conditions nécessaires à l'emploi des modèles fréquentiels obtenus par les transformations de Laplace et Fourier.

Le « pas » de l'algorithme est donc donné par λ . Il peut être fixe, ou alors adaptatif, auquel cas, on l'augmente aux itérations où le critère J diminue, et on le diminue dans le cas contraire.

6.2.2 Méthode de Gauss-Newton

La méthode de Gauss Newton est basée sur une approximation à l'ordre deux de la fonction de coût à minimiser. Le calcul des paramètres se fait par

$$\theta(n+1) = \theta(n) - (\nabla^2 J(\theta(n)))^{-1} \nabla J(\theta(n)). \quad (6.14)$$

Dans le cas où J est effectivement une parabolioïde, l'algorithme converge en une itération. Il requiert cependant le calcul du Hessien $\nabla^2 J(\theta)$ et son inversion. Pour limiter les calculs, et améliorer la stabilité de l'algorithme, il est habituel de le remplacer par un « pseudo-hessien » en négligeant les termes du second ordre.

6.2.3 Méthode de Levenberg-Marquardt

La méthode de Levenberg-Marquardt repose aussi sur un algorithme itératif, défini par la relation

$$\theta(n+1) = \theta(n) - (\nabla^2 J(\theta(n)) + \mu \mathbf{I})^{-1} \nabla J(\theta(n)). \quad (6.15)$$

Remarquons qu'on peut adapter le comportement de l'algorithme en faisant varier le coefficient μ . En effet, si μ tend vers 0,

$$(\nabla^2 J(\theta) + \mu \mathbf{I}) \rightarrow \nabla^2 J(\theta), \quad (6.16)$$

et nous retrouvons l'algorithme de Gauss-Newton.

En revanche, si μ tend vers l'infini,

$$(\nabla^2 J(\theta) + \mu \mathbf{I}) \sim \mu \mathbf{I}. \quad (6.17)$$

L'algorithme se comporte alors comme la méthode du gradient, avec $\lambda = 1/\mu$.

D'un point de vue pratique, μ est augmenté aux itérations où le critère J augmente, ce qui permet d'utiliser le gradient lorsque l'on se trouve loin de l'optimum. Inversement, si J diminue, on se rapproche de l'optimum, et on accélère la convergence par un algorithme de type Gauss-Newton.

6.3 Mise en œuvre

D'après la section 5.3.3, la détection d'un défaut sur une ligne polyphasée est la détermination de sa distance X et des éléments de la résistance matricielle $\mathbf{R}_{\text{def}}$. La mise en œuvre de l'algorithme de Levenberg-Marquardt nécessite le calcul du gradient et du Hessien à chaque itération. Précisons comment il est réalisé en pratique.

6.3. Mise en œuvre

6.3.1 Calcul du gradient

Développons l'expression du gradient de la fonction de coût J .

$$\nabla J(\theta) = \begin{pmatrix} \frac{\partial J}{\partial \theta_1} \\ \vdots \\ \frac{\partial J}{\partial \theta_{n_p}} \end{pmatrix} \quad (6.18)$$

où $\theta_1 = X$ et $\theta_2, \dots, \theta_{n_p}$ sont les éléments distincts de la matrice symétrique $\mathbf{R}_{\text{def}}$ (il y en a $p(p+1)/2$ pour une matrice de dimension $p \times p$).

Seule la tension estimée dépend des paramètres θ , donc, à l'instant $t = n/f_e$,

$$\frac{\partial J}{\partial \theta_i}(n) = 2 \left(\frac{\partial \hat{\mathbf{v}}(n)}{\partial \theta_i}(n) \right)^T (\hat{\mathbf{v}}(n) - \mathbf{v}(n)). \quad (6.19)$$

Il apparaît donc une erreur de prédiction $(\hat{\mathbf{v}}(n) - \mathbf{v}(n))$ multipliée par une *fonction de sensibilité*

$$\sigma_i = \frac{\partial \hat{\mathbf{v}}(n)}{\partial \theta_i}. \quad (6.20)$$

6.3.2 Calcul du Hessian

Le terme général du Hessian est

$$\frac{\partial^2 J}{\partial \theta_i \partial \theta_j}(n) = 2\sigma_i^T \sigma_j + 2 \left(\frac{\partial^2 \hat{\mathbf{v}}}{\partial \theta_i \partial \theta_j}(n) \right)^T (\hat{\mathbf{v}}(n) - \mathbf{v}(n)), \quad (6.21)$$

soit, en négligeant les termes d'ordre deux,

$$\nabla^2 J(n) = \left(\frac{\partial^2 J}{\partial \theta_i \partial \theta_j}(n) \right)_{1 \leq i, j \leq n_p} \approx 2 (\sigma_i^T \sigma_j)_{1 \leq i, j \leq n_p}. \quad (6.22)$$

où n_p est le nombre de paramètres.

Le calcul des fonctions de sensibilité suffit donc à nous fournir à la fois le gradient et le Hessian de J .

6.3.3 Calcul des fonctions de sensibilité

Les fonctions de sensibilité peuvent être calculées de manière exacte à l'aide du logiciel de calcul symbolique tels que *Maple* ou *Maxima*⁴. En effet, nous pouvons reprendre l'expression (6.11) du modèle ARMA vu à la section 6.1

⁴<http://www.maplesoft.com>, <http://maxima.sourceforge.net>

Les a_k et les $\mathbf{b}_k$ sont des fonctions connues des θ_i fournies par le modèle de la ligne. Donc,

$$\begin{aligned} \frac{\partial \mathbf{v}(n)}{\partial \theta_i} = & \frac{\partial \mathbf{b}_0}{\partial \theta_i} \mathbf{i}(n) + \frac{\partial \mathbf{b}_1}{\partial \theta_i} \mathbf{i}(n-1) + \dots - \frac{\partial \mathbf{b}_N}{\partial \theta_i} \mathbf{v}(n-N) \\ & - \frac{\partial a_1}{\partial \theta_i} \mathbf{v}(n-1) - \dots - \frac{\partial a_N}{\partial \theta_i} \mathbf{v}(n-N). \end{aligned} \quad (6.23)$$

Mais dans la pratique, ces expressions deviennent rapidement inexploitable. Il ressort de nos essais qu'une évaluation numérique des fonctions de sensibilité donne d'aussi bon résultat.

$$\sigma_i \approx \frac{\hat{\mathbf{v}}(\theta_1, \dots, \theta_i + \varepsilon, \dots) - \hat{\mathbf{v}}(\theta_1, \dots, \theta_i, \dots)}{\varepsilon} \quad (6.24)$$

en prenant ε suffisamment petit.

6.3.4 Simulation pour une ligne monophasée

Pour tester cet algorithme, une simulation est effectuée sous *Matlab*, pour une ligne monophasée de longueur $X = 100$ km et $X = 300$ km pour un défaut de $R_{\text{def}} = 50 \Omega$. Les autres valeurs numériques sont celles du tableau 6.1. Les signaux de mesure sont simulés par la méthode de la transformée de Fourier inverse déjà employée dans cette étude.

Paramètre	Symbole	Valeur
résistance linéique	r	$3 \cdot 10^{-2} \Omega \text{km}^{-1}$
inductance linéique	l	$1,05 \cdot 10^{-3} \text{Hkm}^{-1}$
conductance linéique	g	$6,5 \cdot 10^{-9} \text{Skm}^{-1}$
capacité linéique	c	$11 \cdot 10^{-9} \text{Fkm}^{-1}$

TAB. 6.1 – Valeurs numériques du modèle de ligne utilisé pour l'identification.

Les coefficients des modèles ARMA sont calculés de façon symbolique par un programme *Maple*, qui écrit les fonctions correspondantes sous la forme de fichiers pouvant être interprétés par *Matlab*.

Les modèles utilisés sont le modèle classique des protections *Alstom* (fig. 6.1), et les modèles développés aux ordres un et deux (fig. 6.2 et 6.3). Pour des raisons pratiques, nous effectuons une seule itération de l'algorithme de Levenberg-Marquardt à chaque pas d'échantillonnage. Les résultats sont tracés sur les figures 6.5 à 6.10.

Pour une ligne de 100 km, nous pouvons conclure de ces résultats que le modèle des protections *Alstom* utilisé avec l'algorithme de Levenberg-Marquardt conduisent à sous-estimer la longueur de la ligne et la résistance de défaut. Cela s'explique bien par le fait qu'elle ne tiennent pas compte des effets capacitifs (c), qui ont tendance à faire diminuer le déphasage et des conductances réparties (g) que nous avons volontairement surestimées et qui font diminuer la résistance en régime permanent.

6.3. Mise en œuvre

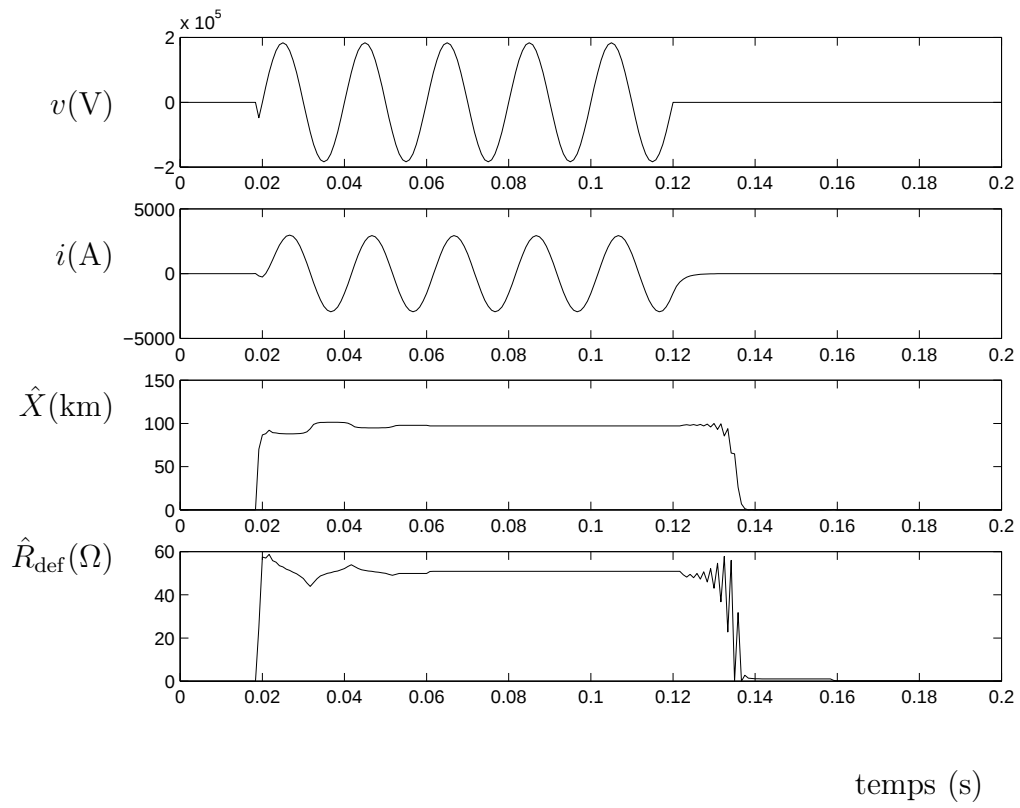


FIG. 6.5 – Identification par l’algorithme de Levenberg-Marquardt (ligne de 100 km, modèle *Alstom*).

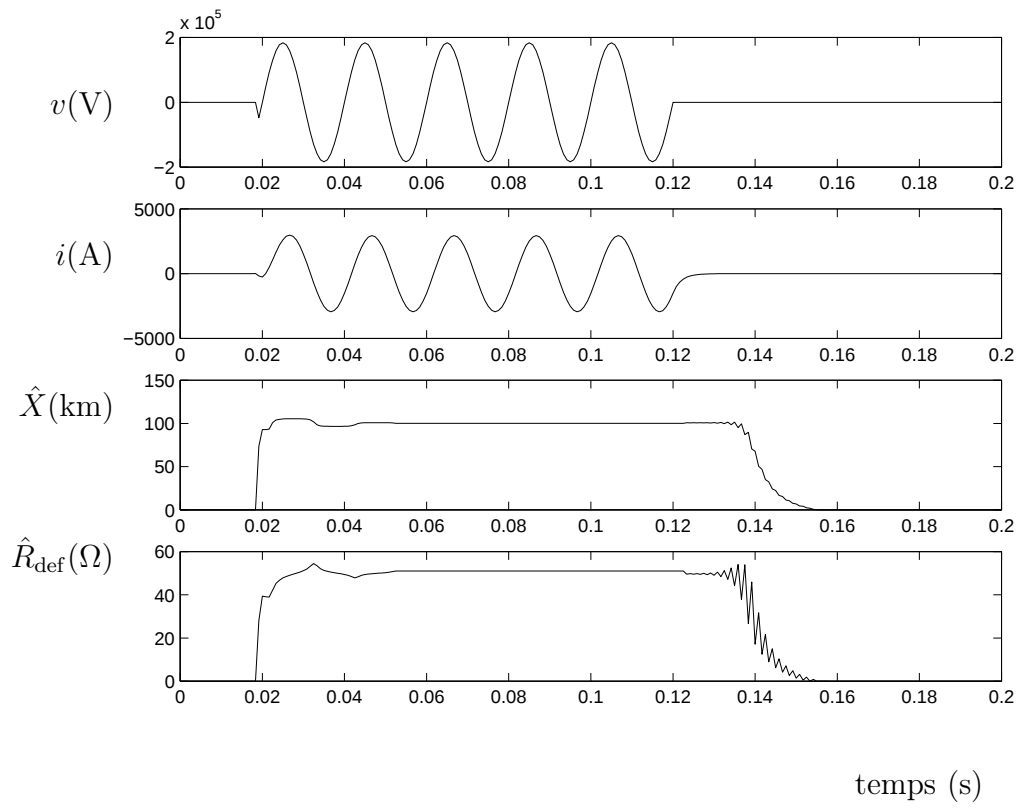


FIG. 6.6 – Identification par l’algorithme de Levenberg-Marquardt (ligne de 100 km, modèle développé à l’ordre 1).

6.3. Mise en œuvre

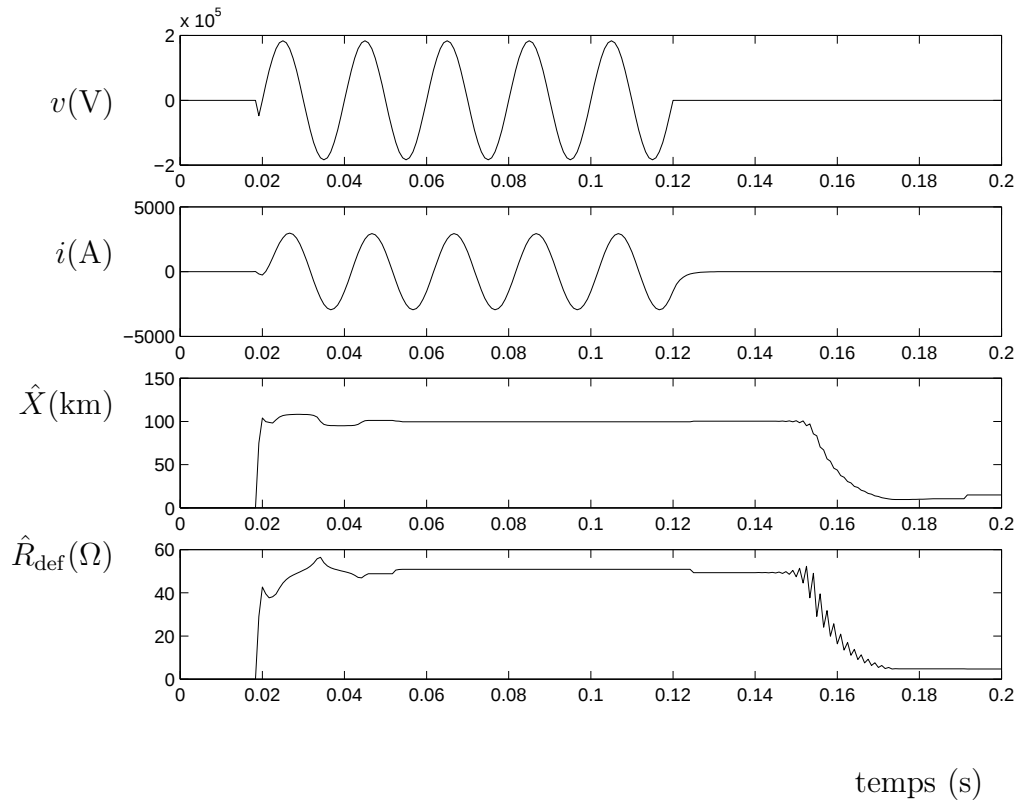


FIG. 6.7 – Identification par l’algorithme de Levenberg-Marquardt (ligne de 100 km, modèle développé à l’ordre 2).

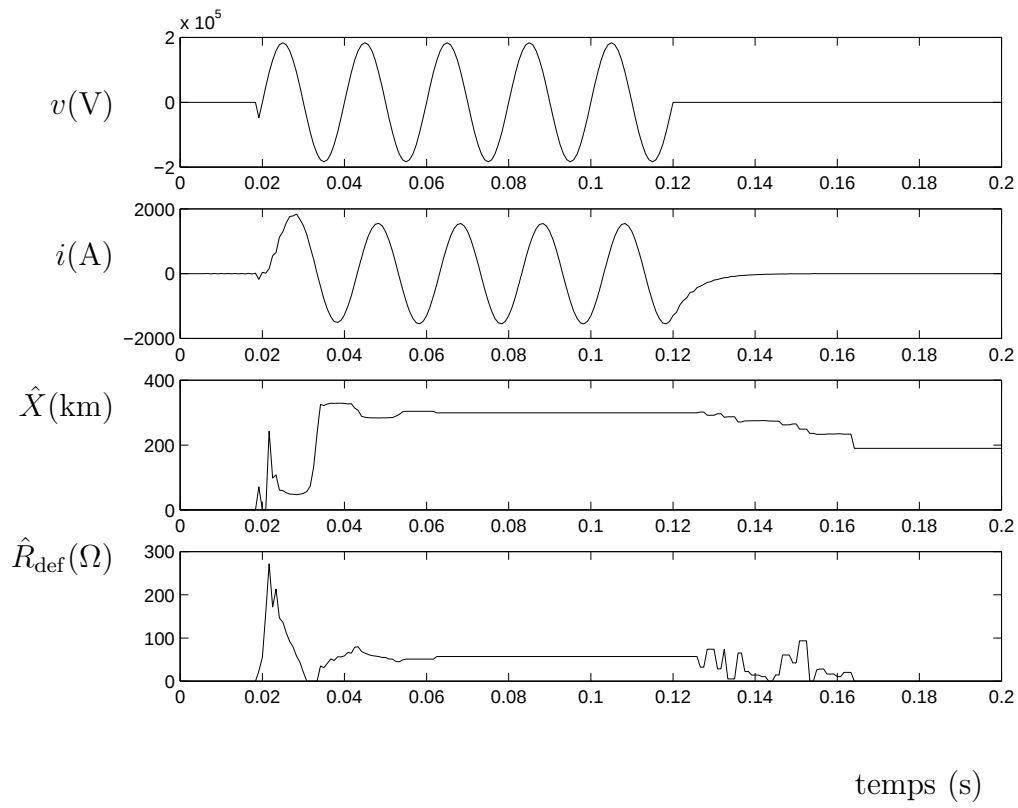


FIG. 6.8 – Identification par l’algorithme de Levenberg-Marquardt (ligne de 300 km, modèle *Alstom*).

6.3. Mise en œuvre

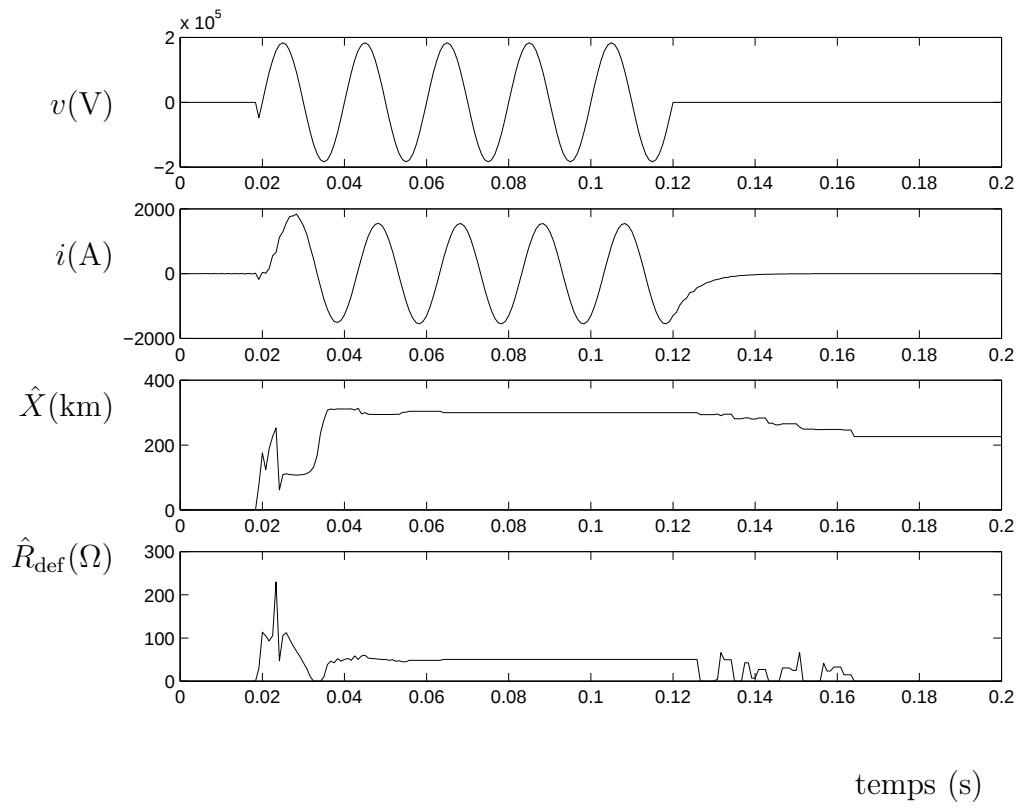


FIG. 6.9 – Identification par l’algorithme de Levenberg-Marquardt (ligne de 300 km, modèle développé à l’ordre 1).

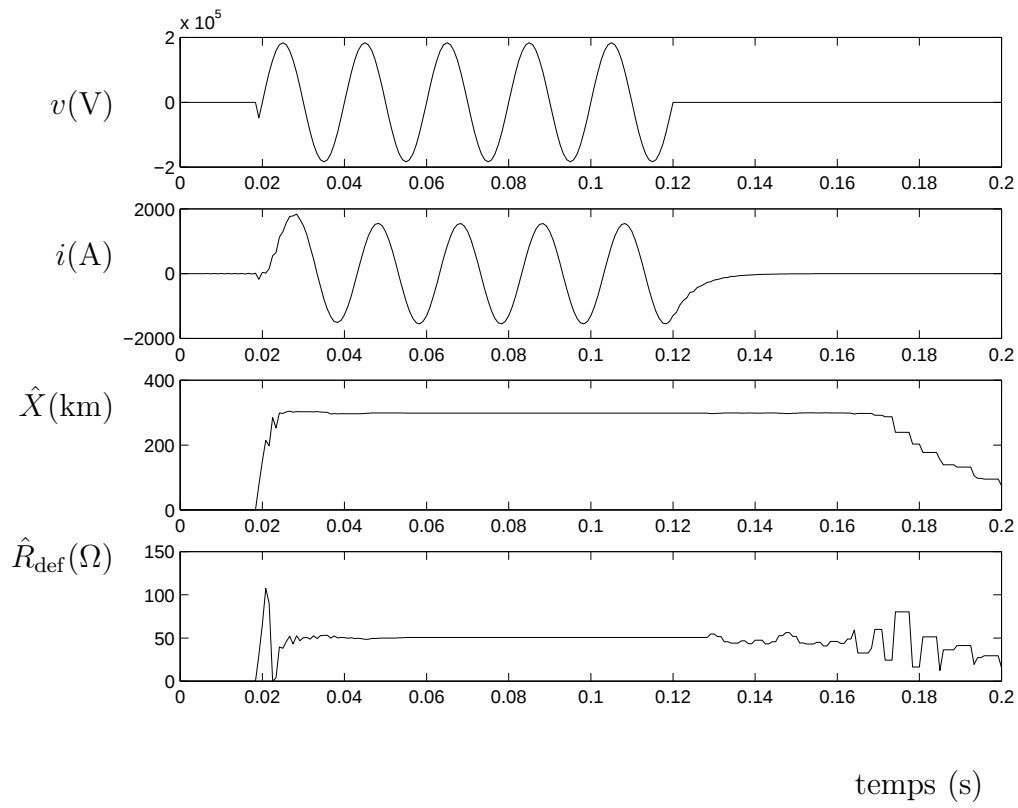


FIG. 6.10 – Identification par l'algorithme de Levenberg-Marquardt (ligne de 300 km, modèle développé à l'ordre 2).

6.3. Mise en œuvre

Les modèles développés sont quant à eux plus précis dans la détermination des paramètres, le modèle développé à l'ordre deux semblant un peu plus rapide que le modèle de ligne courte (dépassements) et plus stable quand l'amplitude des signaux devient faible. Toutefois, le gain de performances procuré par l'augmentation de l'ordre du modèle reste faible.

En revanche, pour une ligne de 300 km, les effets de la propagation sont plus sensibles et le modèle de ligne des protections *Alstom* ne peut pas en rendre compte de façon satisfaisante. Le gain en précision et en stabilité est notable en passant à un modèle de ligne courte (ordre un), et encore plus pour un modèle développé à l'ordre deux.

Malheureusement, dès que l'ordre du modèle dépasse trois (figure 6.4) ou que l'on utilise un modèle polyphasé, le protocole de simulation que nous avons adopté montre ses limites. En effet, les coefficients du modèle ARMA calculés par *Maple* ont des expressions qui deviennent trop complexes pour être traitées dans un temps acceptable, voire pour être simplement utilisables par *Matlab*, la taille des fichiers devenant trop importante.

Remarquons que l'algorithme de localisation de défauts que nous avons proposé ne nécessite que le calcul de la réponse du modèle approché. En effet, le gradient et le Hessien sont obtenus par les fonctions de sensibilité, et nous avons montré que l'on peut les estimer à l'aide du modèle direct. Il suffit donc de trouver une méthode d'implémentation des modèles plus efficace que la méthode ARMA utilisée, ce que nous nous proposons de faire dans le chapitre suivant.

Chapitre 7

Implémentation numérique des réseaux en échelle

Devant les problèmes rencontrés lors de la numérisation des modèles issus de la synthèse de Cauer, nous allons rechercher dans ce chapitre une méthode d'implémentation originale. Ces modèles approchés de ligne sont composés d'impédances développées sous la forme de réseau en échelle. Cette forme présente une régularité remarquable, que nous allons mettre à profit. Nous montrerons en effet que la réponse du modèle peut être simplement obtenue par une succession de filtrages d'ordres peu élevés, à l'aide d'une boucle.

7.1 Équivalence entre réseaux en échelle et filtres en treillis

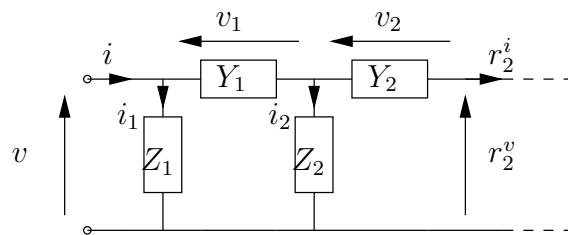


FIG. 7.1 – Réseau en échelle obtenu par la synthèse de Cauer généralisée.

Considérons tout d'abord le cas d'une impédance développée scalaire, qui correspond au cas monophasé. Nous allons exprimer les relations entre les tensions et les courants au sein d'un réseau en échelle.

7.1. Équivalence entre réseaux en échelle et filtres en treillis

7.1.1 Relations de récurrence

En plus des notations présentes sur la figure 7.1, nous définissons un *résidu de tension* r_n^v , ainsi qu'un *résidu de courant* r_n^i .

$$r_n^v = v - \sum_{i=1}^n v_i \quad r_n^i = i - \sum_{i=1}^n i_i \quad (7.1)$$

L'application des lois de Kirchhoff sur le réseau infini (avant troncature) nous permet d'écrire

$$v = \sum_{n=1}^{\infty} v_n, \quad i = \sum_{n=1}^{\infty} i_n. \quad (7.2)$$

La convergence de ces séries est donc équivalente à la convergence de la fraction continue, ce qui peut aussi s'exprimer par

$$\lim_{n \rightarrow \infty} r_n^i = \lim_{n \rightarrow \infty} r_n^v = 0. \quad (7.3)$$

Si nous posons de plus : $r_0^v = v$, $r_0^i = i$. À l'aide des lois de Kirchhoff appliquées au réseau en échelle nous aboutissons aux relations de récurrence suivantes :

$$\begin{cases} i_n = \frac{r_{n-1}^v}{Z_n} \\ r_n^i = r_{n-1}^i - i_n \\ v_n = \frac{r_n^i}{Y_n} \\ r_n^v = r_{n-1}^v - v_n \end{cases} \quad (7.4)$$

7.1.2 Filtres en treillis

Les fonctions de transfert $\frac{1}{Z_n}$ et $\frac{1}{Y_n}$ peuvent aussi être considérées comme des filtres passe-bas du premier ordre.

$$\frac{1}{Z_n} = \frac{1}{R_n + L_n p}, \quad \frac{1}{Y_n} = \frac{1}{G_n + C_n p}, \quad (7.5)$$

Les relations de récurrence précédentes correspondent alors à une succession de filtrages, que nous avons représentée sous la forme d'une association de filtres sous forme de treillis, selon le schéma de la figure 7.2.

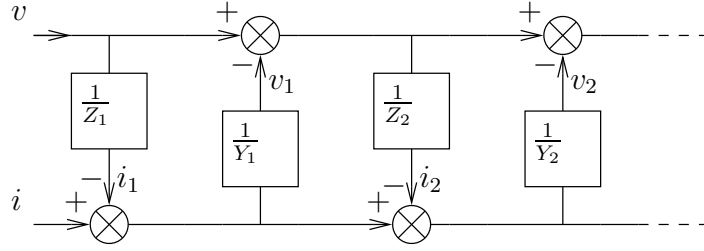


FIG. 7.2 – Filtre en treillis.

7.2 Synthèse du filtre numérique

7.2.1 Discrétisation

Pour numériser les filtres $\frac{1}{Z_n}$ et $\frac{1}{Y_n}$, nous avons recours à la transformation bilinéaire présentée au chapitre 6. Si nous écrivons un passe-bas sous la forme générale

$$F_{A,B}(p) = \frac{1}{A + Bp} \rightarrow F_{A,B}(z) = \frac{1 + z^{-1}}{(A + 2B) + (A - 2B)z^{-1}}, \quad (7.6)$$

alors par transformation bilinéaire,

$$\frac{1}{Z_n} \rightarrow F_{R_n, L_n}(z), \quad \frac{1}{Y_n} \rightarrow F_{G_n, C_n}(z). \quad (7.7)$$

Nous pouvons donc remplacer les impédances et les admittances du réseau par un même filtre numérique, en adaptant les paramètres A et B , ce qui présente une économie pour l'implémentation.

Le filtre en treillis est alors mis en œuvre en programmant une boucle parcourue N fois, N étant l'ordre de développement du modèle, en utilisant les relations de récurrence 7.4.

7.2.2 Troncature

Nous devons maintenant exprimer la troncature de la fraction continue qui apparaît dans la synthèse de Cauer. Soit un modèle de Cauer d'ordre N . Par rapport aux équations (7.2), cela se traduit par l'approximation

$$v(n) \approx \sum_{k=1}^N v_k(n-1), \quad i(n) \approx \sum_{k=1}^N i_k(n-1). \quad (7.8)$$

D'un point de vue pratique, pour synthétiser l'impédance (le courant est le signal d'entrée, et la tension est le signal de sortie) sous la forme d'un filtre en treillis, nous utilisons une somme et un retard pour réinjecter dans le réseau la tension calculée à l'instant $(n-1)$. Le principe est schématisé sur la figure 7.3.

7.2. Synthèse du filtre numérique

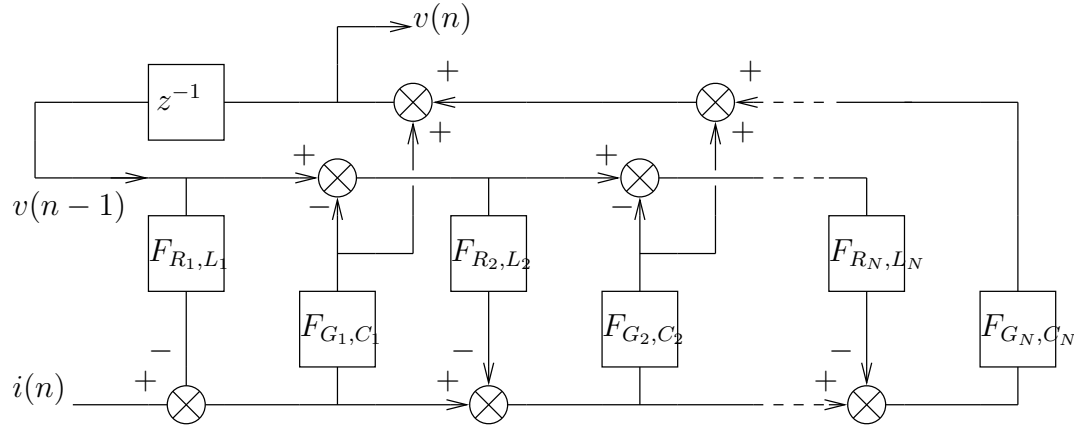


FIG. 7.3 – Structure du filtre en treillis.

7.2.3 Résultats de simulation

Pour vérifier cette méthode de simulation des modèles de lignes de transmission, nous avons appliqué la synthèse du filtre en treillis à l'impédance d'une ligne de transmission en court-circuit.

$$\mathcal{Z} = \sqrt{\frac{Z}{Y}} \tanh \sqrt{ZY} \quad (7.9)$$

Le modèle de Caueur de cette impédance est donné dans le chapitre 4, et correspond aux valeurs

$$Z_n = \frac{Z}{4n-3}, \quad Y_n = \frac{Y}{4n-1}. \quad (7.10)$$

Le filtre numérique passe-bas est implémenté sous *Matlab* et nous lui fournissons les valeurs successives des paramètres A et B dans une boucle correspondant aux équations de récurrence (7.4).

Nous prenons pour valeurs numériques

$$r = 1, \quad l = 1, \quad c = .25, \quad g = 1, \quad x = 1, \quad (7.11)$$

ce qui correspond à un cas général de l'équation réduite hyperbolique ($\alpha = 4$, cf. section 2.3.1).

Nous voyons sur la figure 7.4 la réponse du filtre en treillis à une arche de sinusoïde de courant. La tension de référence est calculée par une transformée de Fourier inverse.

Sur la figure 7.5, sont tracées les différentes composantes de la réponse en tension et de la décomposition du courant d'entrée par les différents étages du filtre. Ces signaux correspondent aux courants et aux tensions qui pourraient être mesurés dans le réseau de la figure 7.1.

Nous avons aussi tracé les sommes partielles successives de ces composantes, pour montrer leur convergence vers les signaux d'entrée (i) et de sortie (v).

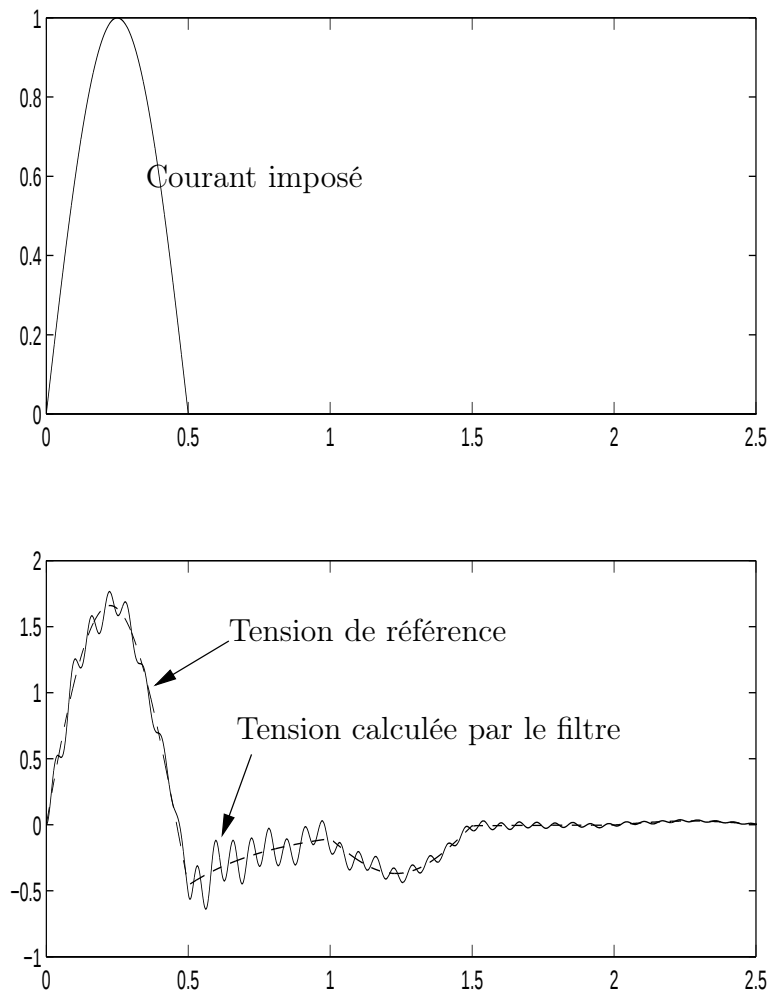


FIG. 7.4 – Réponse temporelle du filtre en treillis à un courant imposé.

7.2. Synthèse du filtre numérique

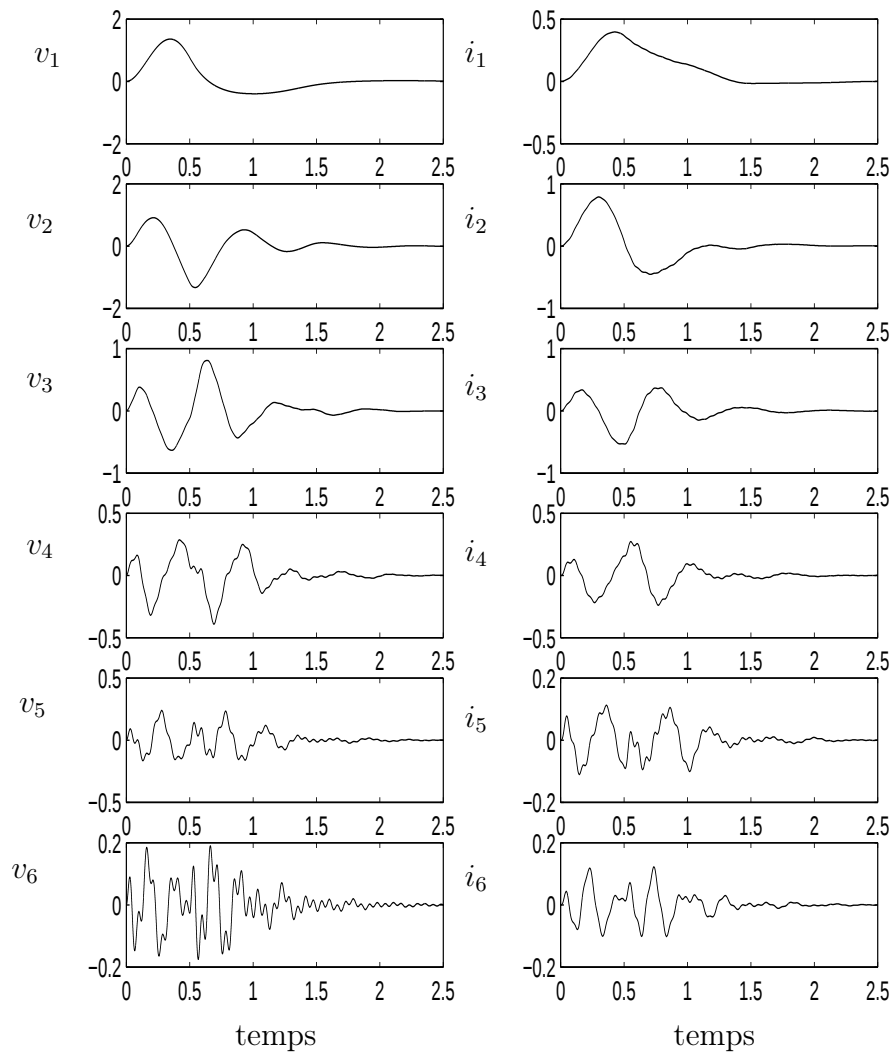


FIG. 7.5 – Composantes de la réponse temporelle du filtre en treillis.

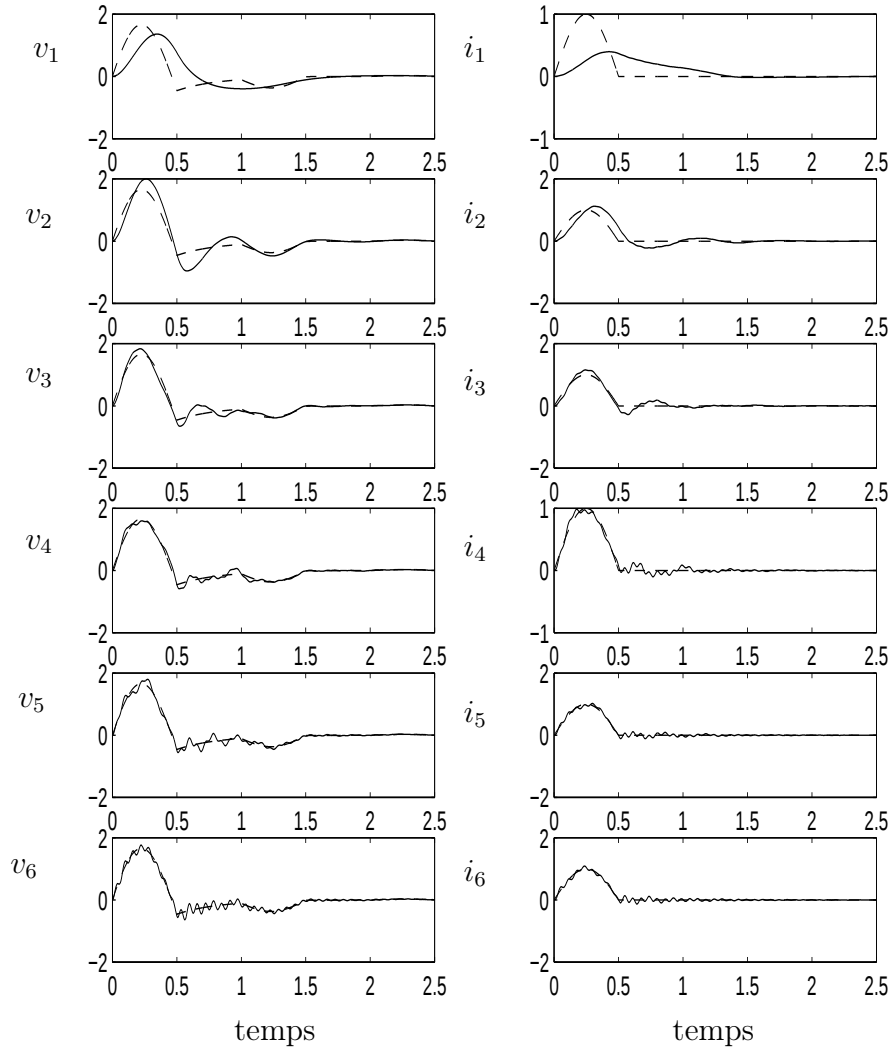


FIG. 7.6 – Sommes partielles des composantes de la réponse du filtre en treillis.

7.3. Remarques sur le cas d'une ligne sans perte

7.2.4 Extension au cas polyphasé

L'extension de la synthèse d'un filtre en treillis à la modélisation des impédances matricielles ne pose aucun problème particulier. La différence essentielle réside dans le fait que les passe-bas $\frac{1}{Z_n}$ et $\frac{1}{Y_n}$ sont remplacés par des matrices de transfert du type MIMO

$$\mathbf{Z}_n^{-1} \quad \text{et} \quad \mathbf{Y}_n^{-1}. \quad (7.12)$$

Ces filtres sont numérisés par transformation bilinéaire, comme décrit à la section 6.1. L'ordre du filtre sera donc celui de la matrice. Il s'agit de passe-bas MIMO d'ordre trois pour la modélisation d'une ligne triphasée.

7.3 Remarques sur le cas d'une ligne sans perte

Nous allons nous attarder maintenant sur le cas particulier d'une ligne de transmission « sans pertes » c'est à dire pour laquelle : $R = G = 0$ (voir section 2.3.1). En effet, en exploitant les récurrence mise en évidence pour la synthèse du filtre en treillis, nous allons montrer qu'il existe dans ce cas des relations intéressantes entre la synthèse de Cauer et les polynômes orthogonaux de Legendre.

7.3.1 Réponse impulsionnelle

Commençons par calculer la réponse impulsionnelle d'une ligne de transmission sans pertes en court-circuit. Exprimons l'impédance de court-circuit (7.9) en décomposant la fonction tangente hyperbolique.

$$\mathcal{Z} = \frac{v}{i} = \sqrt{\frac{Z}{Y}} \frac{1 - e^{-2\sqrt{ZY}}}{1 + e^{-2\sqrt{ZY}}}. \quad (7.13)$$

D'où l'expression de la tension

$$v = \sqrt{\frac{Z}{Y}} i - \sqrt{\frac{Z}{Y}} i e^{-2\sqrt{ZY}} - v e^{-2\sqrt{ZY}} \quad (7.14)$$

Puisque la ligne est sans pertes, nous avons $\sqrt{ZY} = \sqrt{LC}p = \tau p$ et $\sqrt{\frac{Z}{Y}} = \sqrt{\frac{L}{C}} = Z_c$ avec $(\tau, Z_c) \in \mathbb{R}^2$. Ce qui s'exprime dans le domaine temporel par des retards purs (voir l'annexe A).

$$v(t) = Z_c i(t) - Z_c i(t - 2\tau) - v(t - 2\tau) \quad (7.15)$$

Du point de vue de la propagation, 2τ est le temps mis par un onde progressive pour faire un aller retour sur la ligne après réflexion sur le court-circuit.

En posant $i(t) = \delta(t)$ (distribution de Dirac [Sch66]), nous identifions $v(t)$ à la réponse impulsionnelle de la ligne en court-circuit, que nous obtenons simplement à l'aide de

l'équation 7.15.

$$v(t) = Z_c \left(\delta(t) + 2 \sum_{n=1}^{\infty} (-1)^n \delta(t - 2n\tau) \right) \quad (7.16)$$

La réponse impulsionnelle est donc un « demi-peigne » de Diracs de poids alternés.

7.3.2 Décomposition de la réponse impulsionnelle

Nous considérons toujours une ligne en court-circuit sans pertes. Les filtres passe-bas de la section 7.1.2 sont dans ce cas des intégrateurs purs.

$$\frac{1}{Z_n} = \frac{1}{L_n p} = \frac{4n-3}{Lp} \quad (7.17)$$

$$\frac{1}{Y_n} = \frac{1}{C_n p} = \frac{4n-1}{Cp} \quad (7.18)$$

Les relations de récurrence de la section 7.1.1 s'expriment facilement dans le domaine temporel.

$$\begin{cases} i_n(t) = \frac{4n-3}{L} \int_0^t r_{n-1}^v(u) du \\ r_n^i(t) = r_{n-1}^i(t) - i_n(t) \\ v_n(t) = \frac{4n-1}{C} \int_0^t r_n^i(u) du \\ r_n^v(t) = r_{n-1}^v(t) - v_n(t) \end{cases} \quad (7.19)$$

Pour simplifier ce système, introduisons les changements de variables : $\tilde{v}_n = \sqrt{C}v_n$, $\tilde{i}_n = \sqrt{L}i_n$, $\tilde{r}_n^v = \sqrt{C}r_n^v$, $\tilde{r}_n^i = \sqrt{L}r_n^i$, $\tilde{t} = \frac{t}{\tau}$, $\tilde{u} = \frac{u}{\tau}$. Nous obtenons ainsi le système d'équations réduites (section 2.3.1)

$$\begin{cases} \tilde{i}_n(\tilde{t}) = (4n-3) \int_0^{\tilde{t}} \tilde{r}_{n-1}^v(\tilde{u}) d\tilde{u} \\ \tilde{r}_n^i(\tilde{t}) = \tilde{r}_{n-1}^i(\tilde{t}) - \tilde{i}_n(\tilde{t}) \\ \tilde{v}_n(\tilde{t}) = (4n-1) \int_0^{\tilde{t}} \tilde{r}_n^i(\tilde{u}) d\tilde{u} \\ \tilde{r}_n^v(\tilde{t}) = \tilde{r}_{n-1}^v(\tilde{t}) - \tilde{v}_n(\tilde{t}) \end{cases} \quad (7.20)$$

Plaçons-nous dans le cas où $\tilde{i} = \delta(\tilde{t})$. Alors, en utilisant les résultats de la section 7.3.1 sur la réponse impulsionnelle de la ligne,

$$\tilde{v} = \sqrt{C}v = \sqrt{C}z * i = \sqrt{C}z * \left(\frac{\tilde{i}}{\sqrt{L}} \right) = \frac{1}{Z_c} z \quad (7.21)$$

$$\tilde{v}(\tilde{t}) = \delta(\tilde{t}) + 2 \sum_{n=1}^{\infty} (-1)^n \delta(\tilde{t} - 2n) \quad (7.22)$$

Pour simplifier la relation de récurrence, posons pour $n \geq 1$,

$$a_{2n-1} = \tilde{i}_n \quad a_{2n} = \tilde{v}_n \quad (7.23)$$

7.3. Remarques sur le cas d'une ligne sans perte

$$r_{2n-1}^a = \tilde{r}_n^i \quad r_{2n}^a = \tilde{r}_n^v \quad (7.24)$$

Avec le changement de variable $2n \rightarrow n$, nous sommes simplement ramenés à la récurrence suivante :

$$\begin{cases} r_{-1}^a = \tilde{i} & r_0^a = \tilde{v} \\ a_n(\tilde{t}) = (2n-1) \int_0^{\tilde{t}} r_{n-1}^a(\tilde{u}) d\tilde{u}, & n \geq 1 \\ r_n^a(\tilde{t}) = r_{n-2}^a(\tilde{t}) - a_n(\tilde{t}), & n \geq 1 \end{cases} \quad (7.25)$$

Ces relations peuvent être utilisées pour calculer les premiers termes de la suite des a_n . Ainsi, pour les ordres $n = 1$ et $n = 2$, et en notant U la fonction de Heaviside (« échelon unité »), nous avons

$$\begin{aligned} a_1(\tilde{t}) &= \int_0^{\tilde{t}} r_0^a(\tilde{u}) d\tilde{u} \\ &= \int_0^{\tilde{t}} (\delta(\tilde{t}) + 2 \sum_{n=1}^{\infty} (-1)^n \delta(\tilde{t} - 2n)) d\tilde{u} \\ &= U(\tilde{t}) + 2 \sum_{n=1}^{\infty} (-1)^n U(\tilde{t} - 2n) \end{aligned}$$

Ce qui peut encore se lire :

$$\begin{cases} a_1(\tilde{t}) = 1 \text{ pour } t \in [4k, 4k+2[\\ a_1(\tilde{t}) = -1 \text{ pour } t \in [4k+2, 4k+4[\\ k \in \mathbb{N} \end{cases} \quad (7.26)$$

De même, a_2 se calcule par intégration.

$$\begin{aligned} a_2(\tilde{t}) &= 3 \int_0^{\tilde{t}} (r_{-1}^a(\tilde{u}) - a_1(\tilde{u})) d\tilde{u} \\ &= 3 \int_0^{\tilde{t}} (\delta(\tilde{u}) - a_1(\tilde{u})) d\tilde{u} \\ &= 3(U(\tilde{t}) - \int_0^{\tilde{t}} a_1(\tilde{u}) d\tilde{u}) \end{aligned}$$

C'est-à-dire

$$\begin{cases} a_2(\tilde{t}) = 3(4k+1-\tilde{t}) \text{ pour } \tilde{t} \in [4k, 4k+2[\\ a_2(\tilde{t}) = 3(\tilde{t}-4k-3) \text{ pour } \tilde{t} \in [4k+2, 4k+4[\\ k \in \mathbb{N} \end{cases} \quad (7.27)$$

$a_1 = \tilde{i}_1$ est une fonction constante par morceaux. Elle est formée de « créneaux » successifs de valeur ± 1 . $a_2 = \tilde{v}_1$ est continue et linéaire par morceaux (« dents de scie »). Les allures de ces deux premières composantes peuvent être observées sur la figure 7.7.

7.3.3 Expression à l'aide des polynômes de Legendre

Nous nous limitons dans un premier temps à l'intervalle $\tilde{t} \in [0, 2[$ allons démontrer que les composantes a_n (et donc $\tilde{i}_n$ et $\tilde{v}_n$) présentent la propriété suivante.

Propriété $\forall \tilde{t} \in [0, 2[$, $a_n(\tilde{t}) = (-1)^{n+1}(2n-1)P_{n-1}(\tilde{t}-1)$, les P_n étant les polynômes orthogonaux de Legendre (annexe C).

Démonstration Procédons par récurrence. Vérifions la relation pour les deux premiers termes. Soit $\tilde{t} \in [0, 2[$, d'après le calcul de la section 7.3.2, nous avons $a_1(\tilde{t}) = 1$ et $a_2(\tilde{t}) = 3(1-\tilde{t})$.

Or, les premiers polynômes de Legendre sont $P_0(x) = 1$, $P_1(x) = x$. La relation est donc vérifiée pour les ordres $n = 1$ et $n = 2$. Supposons maintenant la relation exacte jusqu'à $n \geq 2$. En utilisant les relations de récurrence (7.25), il vient

$$\begin{aligned} a_{n+1}(\tilde{t}) &= (2n+1) \int_0^{\tilde{t}} r_n^a(\tilde{u}) d\tilde{u} \\ &= (2n+1) \int_0^{\tilde{t}} (r_{n-2}^a(\tilde{u}) - a_n(\tilde{u})) d\tilde{u} \\ &= \frac{2n+1}{2n-3} a_{n-1}(\tilde{t}) - \int_0^{\tilde{t}} a_n(\tilde{u}) d\tilde{u} \end{aligned}$$

Puisque la relation est supposée vérifiée jusqu'à l'ordre n , nous pouvons écrire a_{n+1} à l'aide de polynômes de Legendre.

$$a_{n+1}(\tilde{t}) = (-1)^n(2n+1)P_{n-2}(\tilde{t}-1) - (-1)^{n+1}(2n+1)(2n-1) \int_0^{\tilde{t}} P_{n-1}(\tilde{u}-1) d\tilde{u} \quad (7.28)$$

Utilisons maintenant la propriété suivante des polynômes de Legendre (voir par exemple [San59], p.178).

$$P'_{n+1}(x) - P'_{n-1}(x) = (2n+1)P_n(x) \quad (7.29)$$

En remarquant que $P_{n+1}(-1) = P_{n-1}(-1) (= \pm 1, \text{ selon la parité de } n)$. Nous en déduisons l'identité

$$(2n+1) \int_0^{\tilde{t}} P_n(\tilde{u}-1) d\tilde{u} = P_{n+1}(\tilde{t}-1) - P_{n-1}(\tilde{t}-1) \quad (7.30)$$

Ce qui nous permet de conclure :

$$\begin{aligned} a_{n+1}(\tilde{t}) &= (-1)^{n+2}(2n+1) \left(P_{n-2}(\tilde{t}-1) + (2n-1) \int_0^{\tilde{t}} P_{n-1}(\tilde{u}-1) d\tilde{u} \right) \\ &= (-1)^{n+2}(2n+1) (P_{n-2}(\tilde{t}-1) + P_n(\tilde{t}-1) - P_{n-2}(\tilde{t}-1)) \\ &= (-1)^{n+2}(2n+1)P_n(\tilde{t}-1) \end{aligned}$$

La relation est donc bien vérifiée pour tout $n \geq 1$.

7.3. Remarques sur le cas d'une ligne sans perte

Périodicité

Montrons maintenant que l'on peut déduire $a_n(\tilde{t})$ pour tout $\tilde{t} \geq 0$ de son expression trouvée sur l'intervalle $\tilde{t} \in [0, 2[$. Pour cela, nous allons étudier les propriétés de symétrie et de périodicité des a_n .

Propriétés

- (i) $\forall \tilde{t} \geq 0, a_n(\tilde{t} + 2) = -a_n(\tilde{t})$ (symétrie) ;
- (ii) $\forall \tilde{t} \geq 0, a_n(\tilde{t} + 4) = a_n(\tilde{t})$ (périodicité).

Démonstration Raisonnons une nouvelle fois par récurrence. À l'aide des expressions de a_1 et a_2 calculées dans la section 7.3.2, nous vérifions immédiatement la propriété (i) pour les ordres $n = 1$ et $n = 2$.

Supposons dorénavant la relation vraie jusqu'à $n \geq 2$. Deux cas sont à distinguer suivant la parité de n . Si n est pair, notons $n = 2p, p \in \mathbb{N}$.

$$\begin{aligned}
 a_{n+1}(\tilde{t} + 2) &= a_{2p+1}(\tilde{t} + 2) \\
 &= (4p + 1) \int_0^{\tilde{t}+2} r_{2p}^a(\tilde{u}) d\tilde{u} \\
 &= (4p + 1) \int_0^{\tilde{t}+2} \left(r_0^a(\tilde{u}) - \sum_{k=1}^p a_{2k}(\tilde{u}) \right) d\tilde{u} \\
 &= (4p + 1) \left(\int_0^{\tilde{t}+2} r_0^a(\tilde{u}) d\tilde{u} - \sum_{k=1}^p \int_0^{\tilde{t}+2} a_{2k}(\tilde{u}) d\tilde{u} \right)
 \end{aligned}$$

Remarquons d'une part que

$$\int_0^{\tilde{t}+2} r_0^a(\tilde{u}) d\tilde{u} = a_1(\tilde{t} + 2) = -a_1(\tilde{t}) = - \int_0^{\tilde{t}} r_0^a(\tilde{u}) d\tilde{u} \quad (7.31)$$

D'autre part, puisque les a_n sont des polynômes de Legendre, à un coefficient multiplicatif près, nous pouvons écrire : $\int_0^2 a_1(\tilde{u}) d\tilde{u} = 2$ et $\int_0^2 a_n(\tilde{u}) d\tilde{u} = 0$ si $n \neq 1$ (propriété d'orthogonalité, annexe C). Nous pouvons en déduire, en utilisant de plus l'hypothèse de récurrence,

$$\int_0^{\tilde{t}+2} a_{2k}(\tilde{u}) d\tilde{u} = \int_0^2 a_{2k}(\tilde{u}) d\tilde{u} + \int_0^{\tilde{t}} a_{2k}(\tilde{u} + 2) d\tilde{u} = \int_0^{\tilde{t}} -a_{2k}(\tilde{u}) d\tilde{u} \quad (7.32)$$

Finalement,

$$\begin{aligned}
 a_{n+1}(\tilde{t} + 2) &= (4p + 1) \left(- \int_0^{\tilde{t}} r_0^a(\tilde{u}) d\tilde{u} - \sum_{k=1}^p \int_0^{\tilde{t}} -a_{2k}(\tilde{u}) d\tilde{u} \right) \\
 &= -(4p + 1) \int_0^{\tilde{t}+2} \left(r_0^a(\tilde{u}) - \sum_{k=1}^p a_{2k}(\tilde{u}) \right) d\tilde{u} \\
 &= -a_{2p+1}(\tilde{t}) = -a_{n+1}(\tilde{t})
 \end{aligned}$$

Il nous reste à étudier le cas impair : $n = 2p + 1$. En procédant comme précédemment, nous obtenons

$$\begin{aligned} a_{n+1}(\tilde{t} + 2) &= a_{2p+2}(\tilde{t} + 2) \\ &= (4p + 3) \left(\int_0^{\tilde{t}+2} r_{-1}^a(\tilde{u}) d\tilde{u} - \sum_{k=0}^p \int_0^{\tilde{t}+2} a_{2k+1}(\tilde{u}) d\tilde{u} \right) \end{aligned}$$

D'une part, comme $r_{-1}^a(\tilde{u}) = \delta(\tilde{u})$, $\int_0^{\tilde{t}+2} r_{-1}^a(\tilde{u}) d\tilde{u} = 1$, et d'autre part, si $k > 0$,

$$\int_0^{\tilde{t}+2} a_{2k+1}(\tilde{u}) d\tilde{u} = \int_0^2 a_{2k+1}(\tilde{u}) d\tilde{u} + \int_0^{\tilde{t}} a_{2k+1}(\tilde{u} + 2) d\tilde{u} = - \int_0^{\tilde{t}} a_{2k+1}(\tilde{u}) d\tilde{u} \quad (7.33)$$

Et si $k = 0$,

$$\int_0^{\tilde{t}+2} a_1(\tilde{u}) d\tilde{u} = \int_0^2 a_1(\tilde{u}) d\tilde{u} + \int_0^{\tilde{t}} a_1(\tilde{u} + 2) d\tilde{u} = 2 - \int_0^{\tilde{t}} a_1(\tilde{u}) d\tilde{u} \quad (7.34)$$

Au final, $\sum_{k=0}^p \int_0^{\tilde{t}+2} a_{2k+1}(\tilde{u}) d\tilde{u} = 2 - \sum_{k=0}^p \int_0^{\tilde{t}} a_{2k+1}(\tilde{u}) d\tilde{u}$, et donc

$$\begin{aligned} a_{n+1}(\tilde{t} + 2) &= (4p + 3) \left(-1 + \sum_{k=0}^p \int_0^{\tilde{t}} a_{2k+1}(\tilde{u}) d\tilde{u} \right) \\ &= -(4p + 3) \left(\int_0^{\tilde{t}} r_{-1}^a(\tilde{u}) d\tilde{u} - \sum_{k=0}^p \int_0^{\tilde{t}} a_{2k+1}(\tilde{u}) d\tilde{u} \right) \\ &= -a_{2p+2}(\tilde{t}) = -a_{n+1}(\tilde{t}) \end{aligned}$$

La propriété (i) est démontrée et (ii) résulte directement de (i).

Ces résultats nous permettent de tracer la décomposition de la réponse impulsionnelle d'une ligne de transmission sans pertes en court-circuit par la synthèse de Cauer. La figure 7.7 montre la forme réduite des composantes que nous avons exprimées comme des polynômes de Legendre. Les figures 7.8 et 7.9 montrent quant à elles les premières sommes partielles de ces composantes, pour observer la convergence, *au sens des distributions* [Sch66].

Perspectives

Il y a donc dans ce système particulier une équivalence entre la synthèse de Cauer et la transformée de Legendre. Les coefficients de Fourier (voir section 3.1.2) de la série sont les coefficients du développement en fraction continue.

Ce résultat assez inattendu qui a donné lieu à une publication [SDL03] permet d'envisager des applications pour la modélisation des lignes à retard en utilisant des polynômes de Legendre. Une autre voie d'application pourrait être la mesure de retards par autocorrélation. En effet, les propriétés d'orthogonalité des polynômes de Legendre font envisager un grand nombre de simplifications dans l'expression de cette autocorrélation.

7.3. Remarques sur le cas d'une ligne sans perte

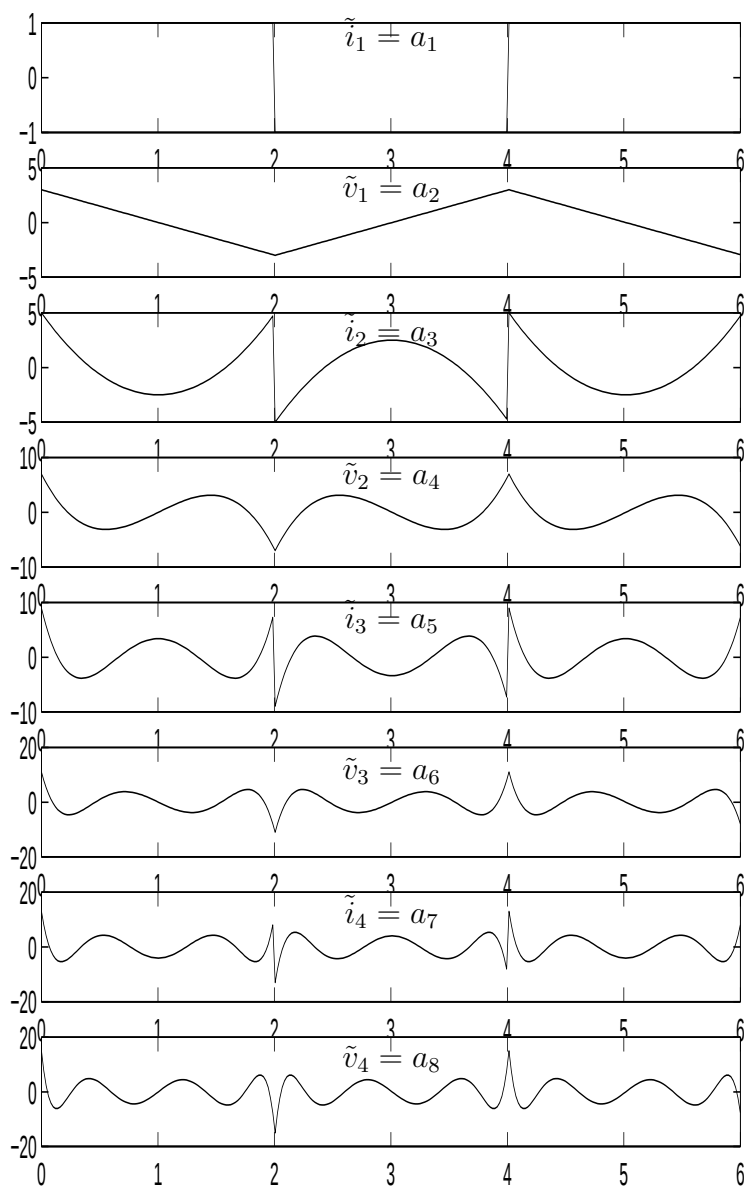


FIG. 7.7 – Premières composantes de la réponse impulsionnelle.

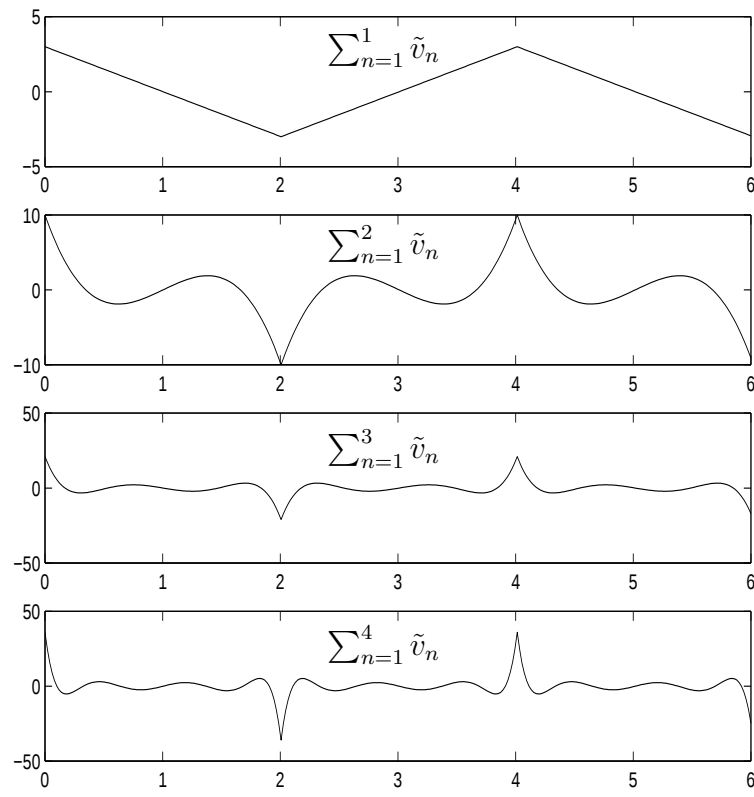


FIG. 7.8 – Sommes partielles de la décomposition de la tension.

7.3. Remarques sur le cas d'une ligne sans perte

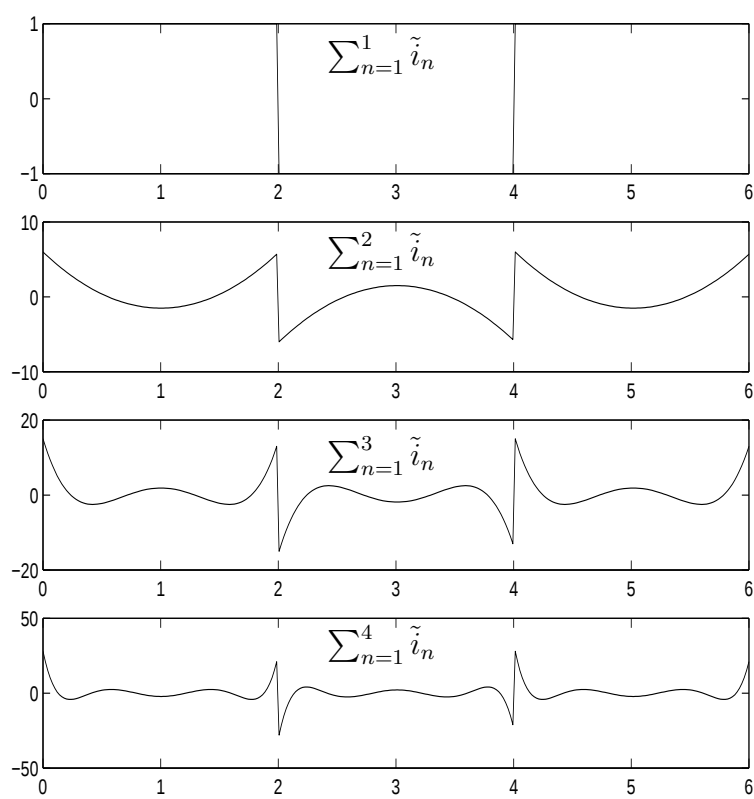


FIG. 7.9 – Sommes partielles de la décomposition du courant.

7.4 Conclusion sur filtrage récursif

Nous avons présenté une méthode de représentation des impédances développées sous forme de réseaux en échelle qui résout le problème de l'implémentation. Parallèlement, le travail sur la modélisation d'impédances par des filtres en treillis a permis de mettre en évidence de nouveaux résultats théoriques sur les relations entre la décomposition sur la base de Legendre et la synthèse de Cauer.

De plus, l'utilisation d'un filtre en treillis présente des avantages significatifs dans un contexte industriel :

- le filtre d'ordre élevé (correspondant à l'ordre de développement de l'impédance) est remplacé par une boucle ;
- la programmation d'un seul filtre du premier ordre (ou MIMO d'ordre trois pour le triphasé) est nécessaire ;
- l'expression des coefficients reste très simple puisqu'ils sont fonctions des constantes globales de la ligne ;

Toutes ces spécificités permettent d'envisager une utilisation optimale sur DSP. Notons que la programmation du calcul par une boucle permet d'envisager d'autres possibilités, comme l'adaptation de l'ordre de l'approximation en temps réel.

Dans l'optique d'une future industrialisation, il sera bien sûr nécessaire de faire suivre l'implémentation des modèles de lignes de tests exhaustifs avec des signaux réels issus d'enregistrement de défauts. Il sera ainsi possible d'effectuer le meilleur choix sur l'ordre de développement, qui sera notamment limité par la fréquence d'échantillonnage des convertisseurs analogiques-numériques, et sur la forme du modèle, en Π ou en T .

7.4. Conclusion sur filtrage récursif

Conclusion

La volonté des initiateurs de ce travail de thèse était de développer une *approche réseau* de la modélisation des lignes électriques, dans le but de trouver le juste compromis entre simplicité et précision. Cette contrainte était liée au contexte industriel et aux limites des technologies numériques. Cette démarche se démarquait à la fois de méthodes plus classiques comme l'analyse modale et la modélisation non entière, et, quoique risquée, se présentait comme un projet séduisant et motivant.

Pour mener à bien ce travail, des investigations dans des domaines scientifiques pluridisciplinaires furent nécessaires, comme :

- la modélisation des lignes mono et polyphasée en électrotechnique,
- les approximants de Padé, les fractions continues et les polynômes orthogonaux en mathématiques,
- la théorie des réseaux de Kirchhoff,
- la modélisation des systèmes à paramètres répartis en automatique,
- la synthèse de filtres numériques et l'identification paramétrique en traitement du signal.

Les principaux résultats scientifiques des travaux de recherche présentés dans ce mémoire concernent l'extension des synthèses de Cauer et Foster à l'ordre infini, et leur utilisation systématique pour modéliser les systèmes à paramètres répartis. L'originalité de notre approche réside dans le fait que les développements ne sont pas nécessairement faits selon la variable fréquentielle. Nous sacrifions de ce fait l'exactitude de la fonction de transfert pour la fréquence nulle, au profit de modèles dont les paramètres s'interprètent physiquement avec une simplicité remarquable.

Il convient cependant de relativiser cet inconvénient car la vitesse de convergence des approximants de Padé permet d'obtenir une large bande passante pour les modèles, y compris dans les basses fréquences. Leur représentation par des réseaux en échelle nous a poussé à retenir la synthèse de Cauer.

Notre démarche se distingue des travaux actuels sur les systèmes d'ordre infini en ce qu'elle s'attache à garder une interprétation physique des modèles. Si elle présente quelques ressemblances avec les méthodes basées sur la dérivation fractionnaire, elle s'en démarque toutefois par le fait que nous isolons la seule approximation par rapport à l'ordre infini comme étant la troncature d'un réseau infini. Cette approximation est parfaitement définie et analysée à l'aide de la théorie de l'approximation rationnelle.

Les excellents résultats obtenus en simulation laissent envisager de nombreuses appli-

7.4. Conclusion sur filtrage récursif

cations pour les phénomènes diffusifs ou propagatifs. Le formalisme matriciel mis en place pour les lignes polyphasées offre alors un outil de choix pour traiter tous les problèmes relatifs aux phénomènes couplés, y compris les systèmes hétérogènes. Ici encore, notre démarche se distingue particulièrement des travaux actuels, qui font essentiellement usage des composantes modales.

L'application de cette modélisation aux lignes de transmission permet pour un faible coût en calcul d'améliorer significativement la réponse fréquentielle des modèles. Il est ainsi possible de mettre à profit les composantes à hautes fréquences des signaux de mesures pour en extraire le maximum d'information durant les phases transitoires. L'utilisation des composantes de lignes pour les modèles polyphasés, ainsi que l'interprétation physique des paramètres permet de mettre au point des algorithmes très simples de détection et de localisation de défauts.

Enfin, nous avons montré comment la représentation des modèles sous forme de réseaux de Kirchhoff pouvait être utilisée pour implémenter ces algorithmes de manière pertinente, et de réduire le coût en charge de calcul des DSP. L'application à des simulateurs de grands réseaux électriques est également envisageable.

Certains résultats obtenus sont actuellement utilisés au laboratoire d'études thermiques dans le cadre de deux autres thèses de doctorat. L'une sur la thermique des semi-conducteurs, l'autre sur la modélisation des moteurs thermiques avec, pour toutes les deux, des conditions de refroidissement variables.

Bien qu'au cours de cette étude nous ayons fait appel à des notions issues de domaines variés, nous pensons avoir obtenu un ensemble et une démarche cohérents. Nous espérons enfin que ce travail ne restera pas sans suite, mais que nous pourrons voir un jour fonctionner des protections industrielles utilisant ces résultats.

« C'est parfois très simple de faire compliqué, mais c'est toujours très compliqué de faire simple » (P. Lagonotte).

Annexe A

La transformation de Laplace

La transformée de Laplace a été initialement utilisée par Heaviside, sous la forme de *calcul opérationnel*¹ pour résoudre des problèmes de type circuits électriques. Elle est maintenant un outil de choix pour le traitement du signal et l'automatique.

A.1 Définition et premières propriétés

Définition Soit f une fonction causale du temps ($f(t) = 0$ si $t < 0$) telle que

- f soit localement sommable sur $\mathbb{R}$;
- f soit continue par intervalle sur tout intervalle du type $0 < T_1 \leq t \leq T_2$;
- f est d'ordre exponentiel à l'infini, c'est-à-dire qu'il existe $M > 0$, $x_0 \in \mathbb{R}$ tel que $|f(t)| \leq Me^{x_0 t}$, pour t suffisamment grand ;

alors la transformée de Laplace est définie par

$$\phi(p) = \int_0^\infty e^{-pt} f(t) dt \quad (\text{A.1})$$

où p appartient à $\mathbb{C}$.

Notation On note

$$f(t) \supset \phi(p), \quad \phi(p) \sqsubset f(t). \quad (\text{A.2})$$

$\phi(p)$ existe quand l'intégral converge, dans le demi-plan $\Re(p) > x_0$. De plus, ϕ est continue et holomorphe dans le demi-plan.

Dérivation

$$\phi'(p) = \int_0^\infty -te^{-pt} f(t) dt \quad (\text{A.3})$$

¹Voir [Mik83] pour une forme moderne du calcul opérationnel.

A.2. Propriétés opérationnelles

Par conséquence,

$$(-t)^n f(t) \sqsubset \phi^{(n)}(p). \quad (\text{A.4})$$

Une multiplication par t sur la fonction f induit une dérivation sur la transformée de Laplace.

Inversement,

$$f^{(n)}(t) \sqsubset p^n \phi(p) - p^{n-1} f(0^+) - p^{n-2} f'(0^+) - \dots - f^{(n-1)}(0^+). \quad (\text{A.5})$$

Une dérivation dans le temps induit une multiplication algébrique sur les transformées de Laplace. C'est cette propriété qui est intéressante pour la résolution des équations différentielles.

Théorème de la valeur initiale En supposant que $|\arg(s)| \leq \alpha < \pi/2$, on a la limite

$$\lim_{|s| \rightarrow \infty} p\phi(p) = f(0^+). \quad (\text{A.6})$$

Théorème de la valeur finale Inversement,

$$\lim_{|s| \rightarrow 0} p\phi(p) = f(+\infty). \quad (\text{A.7})$$

Transformée de Laplace d'une primitive Soit F la primitive

$$F(t) = \int_0^t f(u) du. \quad (\text{A.8})$$

Si F admet une transformée de Laplace, alors

$$F(t) \sqsubset \frac{\phi(p)}{p}. \quad (\text{A.9})$$

A.2 Propriétés opérationnelles

Linéarité Si nous posons les transformées $f_i \sqsubset \phi_i$, pour $\Re(p) > x_i$, alors

$$\sum_{i=1}^n \alpha_i f_i \sqsubset \sum_{i=1}^n \alpha_i \phi_i, \quad \Re(p) > \sup(x_i). \quad (\text{A.10})$$

Facteur d'échelle

$$\begin{aligned} f(t) \sqsubset \phi(p), \quad \Re(p) > x_0, \\ \frac{1}{k} f\left(\frac{t}{k}\right) \sqsubset \phi(kp), \quad \Re(kp) > x_0. \end{aligned} \quad (\text{A.11})$$

Une « dilatation » sur f induit une « contraction » sur ϕ .

Retard sur la variable

$$\begin{aligned} f(t) &\sqsubset \phi(p), & \Re(p) &> x_0, \\ e^{-\lambda t} f(t) &\sqsubset \phi(p + \lambda), & \Re(p + \lambda) &> x_0. \end{aligned} \quad (\text{A.12})$$

Un « déphasage » sur f induit un « translation » sur ϕ .

Inversement,

$$\begin{aligned} f(t) &\sqsubset \phi(p), & \Re(p) &> x_0, \\ f(t - \lambda) &\sqsubset e^{-\lambda p} \phi(p), & \Re(p + \lambda) &> x_0. \end{aligned} \quad (\text{A.13})$$

Un « décalage » dans le temps induit un « déphasage » sur la transformée de Laplace.

Produit de convolution Le produit de convolution est défini par l'intégrale

$$f * g(x) = \int_{-\infty}^{+\infty} f(t)g(x - t)dt. \quad (\text{A.14})$$

Si ϕ est la transformée de f et ψ la transformée de g , alors,

$$f * g \sqsubset \phi\psi. \quad (\text{A.15})$$

A.3 Images de fonction usuelles

Les transformées de Laplace d'un grand nombre de fonctions usuelles sont tabulées dans les ouvrages tels que [Pet95] ou [Spi85].

A.4 Transformation de Laplace inverse

Si deux fonctions f et g admettent la même transformée de Laplace, alors f et g sont égales *presque partout*. Donc $f = g$ sur le domaine de continuité.

Soit $\phi(p)$ holomorphe sur $\Re(p) > x_0$. Le problème de la transformation inverse se réduit à trouver f continue, causale, admettant ϕ comme transformée de Laplace. Ce problème n'a pas toujours de solution.

Méthodes

- Il existe des tables de transformées de Laplace (voir section précédente) que l'on peut utiliser pour trouver la transformée inverse ;
- si ϕ est une fraction rationnelle, on peut la décomposer en éléments simples, puis utiliser les tables de transformées de Laplace ;
- si $\phi(p) = A(p)B(p)$, sous réserve d'existence, $f(t) = a(t) * b(t)$;
- formule de Mellin-Fourier :

$$f(t) = \frac{U(t)}{2i\pi} \int_D \phi(z)e^{zt}dz \quad (\text{A.16})$$

est la transformée inverse de ϕ qui est continue sur $\mathbb{R}_+$, avec (D) une droite verticale quelconque dans le domaine d'holomorphic de ϕ .

A.5 Lien avec la transformée de Fourier

La transformée de Fourier de f est définie par l'intégrale

$$\mathcal{F}(f)(\nu) = \int_{-\infty}^{+\infty} f(t) e^{-2i\pi\nu t} dt, \quad (\text{A.17})$$

sous réserve d'existence.

On vérifie immédiatement que si f est causale et que ϕ est définie sur l'axe imaginaire, alors

$$\mathcal{F}(f)(\nu) = \phi(p)|_{p=2i\pi\nu}. \quad (\text{A.18})$$

Cette dernière propriété permet d'obtenir rapidement la transformée de Fourier pour étudier la réponse harmonique d'un modèle représenté par sa fonction de transfert dans l'espace de Laplace.

Annexe B

Notes sur la résolution de l'équation des télégraphistes

Nous avons donné dans le chapitre 2 la solution de l'équation des télégraphistes. Bien que la résolution de cette équation ne pose pas de problèmes particuliers, nous tenons à la rappeler ici. Elle utilise les résultats classiques sur les équations différentielles linéaires du deuxième ordre.

Pour compléter cette démonstration, nous présenterons de plus deux autres méthodes de résolution qui peuvent être avantageuses à utiliser selon la forme de la solution recherchée. Il s'agit pour la première d'utiliser une équation de Riccati, et pour la seconde les exponentielles de matrices.

Rappelons l'expression du système d'équations couplées donnant les relations locales entre courant et tension (2.2)

$$\begin{cases} \frac{\partial v}{\partial x} = -ri - l \frac{\partial i}{\partial t}, \\ \frac{\partial i}{\partial x} = -gv - c \frac{\partial v}{\partial t}, \end{cases} \quad (\text{B.1})$$

et sa transformée de Laplace, (2.35),

$$\begin{cases} \frac{dv}{dx} + zi = 0, \\ \frac{di}{dx} + yv = 0. \end{cases} \quad (\text{B.2})$$

B.1 Résolution par découplage des équations

Dérivons par rapport à la variable x chacune des équations précédentes, nous découplons les variables, et nous obtenons donc les équations (2.20)

$$\frac{d^2v}{dx^2} - zyv = 0, \quad \frac{d^2i}{dx^2} - zyi = 0. \quad (\text{B.3})$$

B.1. Résolution par découplage des équations

L'équation différentielle du second degré (B.3) vérifiée par v et i est linéaire à coefficients constants (par rapport à la variable d'intégration x). Sa solution générale s'exprime donc comme une combinaison linéaire de fonctions exponentielles. Posons les « constantes d'intégration » λ et μ et notons γ une détermination de la racine carrée du produit zy .

$$\gamma = \sqrt{zy}. \quad (\text{B.4})$$

Pour la tension,

$$v(p, x) = \lambda e^{\gamma x} + \mu e^{-\gamma x} \quad (\text{B.5})$$

En utilisant l'équation (B.2) nous pouvons déduire l'expression du courant de celle de la tension.

$$i(p, x) = -\frac{1}{z}(\lambda \gamma e^{\gamma x} - \mu \gamma e^{-\gamma x}). \quad (\text{B.6})$$

Ou encore, en introduisant la quantité

$$Z_c = \sqrt{\frac{z}{y}}, \quad (\text{B.7})$$

$$i(p, x) = -\frac{\lambda}{Z_c} e^{\gamma x} + \frac{\mu}{Z_c} e^{-\gamma x}. \quad (\text{B.8})$$

On reconnaît γ le *facteur de propagation* et Z_c l'*impédance caractéristique* de la ligne définis au chapitre 2. Bien que cette dernière grandeur relie tension et courant en tout point, remarquons que dans le cas général $v \neq Z_c i$.

B.1.1 Conditions aux limites

Les constantes d'intégration sont fixées par la connaissance de la tension et du courant à l'origine ($x = 0$). On a en effet :

$$\begin{cases} v(p, 0) = \lambda + \mu \\ i(p, 0) = \frac{\mu}{Z_c} - \frac{\lambda}{Z_c} \end{cases} \quad (\text{B.9})$$

d'où

$$\begin{cases} \lambda = \frac{v(p, 0) - Z_c i(p, 0)}{2} \\ \mu = \frac{v(p, 0) + Z_c i(p, 0)}{2} \end{cases} \quad (\text{B.10})$$

Nous en déduisons les expressions de la tension et du courant en fonction de x dans l'espace de Laplace :

$$\begin{cases} v(p, x) = \cosh(\gamma x) v(p, 0) - Z_c \sinh(\gamma x) i(p, 0) \\ i(p, x) = -\frac{\sinh(\gamma x)}{Z_c} v(p, 0) + \cosh(\gamma x) i(p, 0) \end{cases} \quad (\text{B.11})$$

Nous reconnaissons bien la solution de l'équation des télégraphistes donnée au chapitre 2 sous la forme quadripolaire (2.21)

B.2 Utilisation de l'équation de Ricatti

B.2.1 Impédance locale

Nous nous proposons de calculer directement l'impédance d'une ligne électrique de transmission, c'est à dire sans faire appel aux calculs intermédiaires de la tension v et du courant i en tout point.

Pour cela, nous considérons l'impédance *locale* ζ définie par le rapport

$$\zeta = \frac{v}{i} \quad (\text{B.12})$$

Sa dérivée s'exprime par

$$\zeta' = \left(\frac{v}{i}\right)' = \frac{v'i - i'v}{i^2}, \quad (\text{B.13})$$

soit, d'après le système (B.2),

$$\zeta' = \frac{-zi^2 + yv^2}{i^2} = y \left(\frac{v}{i}\right)^2 - z. \quad (\text{B.14})$$

L'impédance locale ζ vérifie donc l'équation de Ricatti

$$\zeta' - y\zeta^2 + z = 0 \quad (\text{B.15})$$

B.2.2 Résolution

La résolution de cette équation se fait en se ramenant à une équation linéaire par un changement de variable. Commençons par exhiber une solution particulière de l'équation (B.15).

Il apparaît immédiatement que l'impédance *caractéristique* (B.7) $Z_c = \sqrt{\frac{z}{y}}$ de la ligne électrique est une telle solution particulière.

Le changement de variable consiste alors à poser

$$\zeta = \sqrt{\frac{z}{y}} + \frac{1}{\xi}. \quad (\text{B.16})$$

L'équation de Ricatti (B.15) devient

$$-\frac{\xi'}{\xi^2} - y \left(\sqrt{\frac{z}{y}} + \frac{1}{\xi} \right)^2 + z = 0, \quad (\text{B.17})$$

et après développement du carré et simplification,

$$\xi' + 2\sqrt{zy}\xi + y = 0. \quad (\text{B.18})$$

Résoudre l'équation différentielle linéaire du premier degré (B.18) ne pose aucun problème.

B.2. Utilisation de l'équation de Ricatti

– L'équation homogène $\xi' + 2\sqrt{zy}\xi = 0$ admet la solution générale : $\xi = Ke^{-2\sqrt{zy}x}$;
– L'équation complète $\xi' + 2\sqrt{zy}\xi + y = 0$ admet la solution particulière : $\xi = -\frac{1}{2}\sqrt{\frac{y}{z}}$.
La solution générale de l'équation (B.18) est donc

$$\xi = Ke^{-2\sqrt{zy}x} - \frac{1}{2}\sqrt{\frac{y}{z}}, \quad K \in \mathbb{C}. \quad (\text{B.19})$$

Et par suite, en effectuant le changement de variable inverse, la solution de l'équation de Ricatti (B.15) est donnée par

$$\zeta = \sqrt{\frac{z}{y}} + \frac{1}{Ke^{-2\sqrt{zy}x} - \frac{1}{2}\sqrt{\frac{y}{z}}} \quad (\text{B.20})$$

$$= \frac{Kze^{-2\sqrt{zy}x} + \frac{\sqrt{zy}}{2}}{K\sqrt{zy}e^{-2\sqrt{zy}x} - \frac{y}{2}} \quad (\text{B.21})$$

Cette expression se simplifie en substituant la constante d'intégration $K \rightarrow 2K$

$$\zeta = \frac{Kze^{-2\sqrt{zy}x} + \sqrt{zy}}{K\sqrt{zy}e^{-2\sqrt{zy}x} - y} \quad (\text{B.22})$$

B.2.3 Exemples de calcul d'impédances

L'impédance d'entrée de la ligne est donnée est obtenue par

$$\zeta(x=0) = \frac{Kz + \sqrt{zy}}{K\sqrt{zy} - y}. \quad (\text{B.23})$$

Supposons que la ligne soit en court-circuit en X . Nous avons la *condition à la limite* suivante

$$0 = \frac{Kze^{-2\sqrt{zy}X} + \sqrt{zy}}{K\sqrt{zy}e^{-2\sqrt{zy}X} - y}, \quad (\text{B.24})$$

dont nous déduisons la constante d'intégration

$$K = -\sqrt{\frac{y}{z}}e^{2\sqrt{zy}X} \quad (\text{B.25})$$

et l'impédance d'entrée

$$\begin{aligned} \zeta(x=0) &= \frac{-\sqrt{zy}e^{2\sqrt{zy}X} + \sqrt{zy}}{-ye^{2\sqrt{zy}X} - y} \\ &= \sqrt{\frac{z}{y}} \frac{e^{2\sqrt{zy}X} - 1}{e^{2\sqrt{zy}X} + 1} \\ &= \sqrt{\frac{z}{y}} \tanh(\sqrt{zy}X). \end{aligned} \quad (\text{B.26})$$

Nous retrouvons bien l'expression de Z_{cc} de la figure 2.14.

À l'inverse, si nous supposons que la ligne est en ouverte en X , la condition à la limite devient

$$\left(\frac{1}{\zeta}\right)_{x=X} = 0 = \frac{K\sqrt{zy}e^{-2\sqrt{zy}X} - y}{Kze^{-2\sqrt{zy}X} + \sqrt{zy}}. \quad (\text{B.27})$$

Et la constante d'intégration

$$K = \sqrt{\frac{y}{z}}e^{2\sqrt{zy}X}. \quad (\text{B.28})$$

Exprimons alors l'impédance d'entrée

$$\begin{aligned} \zeta(x=0) &= \frac{\sqrt{zy}e^{2\sqrt{zy}X} + \sqrt{zy}}{ye^{2\sqrt{zy}X} - y} \\ &= \sqrt{\frac{z}{y}} \frac{e^{2\sqrt{zy}X} + 1}{e^{2\sqrt{zy}X} - 1} \\ &= \sqrt{\frac{z}{y}} \coth(\sqrt{zy}X). \end{aligned} \quad (\text{B.29})$$

Cette fois encore, nous retrouvons l'expression de Z_0 de la figure 2.14.

Nous avons calculé l'expression de l'impédance locale en tout point d'une ligne électrique qui est rarement mise en évidence lors de la résolution classique de l'équation des télégraphistes. Grâce à ce résultat nous avons vu qu'il était possible de retrouver les expressions connues des impédances d'entrée lorsque la ligne est ouverte ou en court-circuit à son extrémité.

Remarquons enfin que cette façon de procéder, bien qu'équivalente mathématiquement à la méthode classique, « cache » l'aspect propagatif des ondes de tension et de courant pour se focaliser sur l'aspect local du couplage.

B.3 Utilisation d'une exponentielle de matrice

Écrivons sous forme matricielle l'équation des télégraphistes. Le système (B.2) donne

$$\frac{d}{dx} \begin{pmatrix} v \\ i \end{pmatrix} = - \begin{pmatrix} 0 & z \\ y & 0 \end{pmatrix} \begin{pmatrix} v \\ i \end{pmatrix} \quad (\text{B.30})$$

Cette équation différentielle du premier degré admet une solution qui s'exprime avec une « exponentielle de matrice ».

$$\begin{pmatrix} v \\ i \end{pmatrix} = e^{-Ax} \begin{pmatrix} v(x=0) \\ i(x=0) \end{pmatrix}; \quad A = \begin{pmatrix} 0 & z \\ y & 0 \end{pmatrix}. \quad (\text{B.31})$$

La matrice A étant clairement régulière, nous pouvons expliciter l'exponentielle à l'aide d'une diagonalisation. Les valeurs propres sont : $\sqrt{zy}$ et $-\sqrt{zy}$. La matrice de vecteurs propres associés est

$$P = \begin{pmatrix} \sqrt{z/y} & -\sqrt{z/y} \\ 1 & 1 \end{pmatrix}. \quad (\text{B.32})$$

B.3. Utilisation d'une exponentielle de matrice

Nous avons donc

$$\begin{pmatrix} v \\ i \end{pmatrix} = P \begin{pmatrix} e^{-\sqrt{zy}x} & 0 \\ 0 & e^{\sqrt{zy}x} \end{pmatrix} P^{-1} \begin{pmatrix} v_0 \\ i_0 \end{pmatrix}, \quad (\text{B.33})$$

et en effectuant les multiplications de matrices, nous retrouvons les relations qui sont caractéristiques du quadripôle équivalent (2.22).

$$\begin{pmatrix} v \\ i \end{pmatrix} = \frac{1}{2} \begin{pmatrix} e^{-\sqrt{zy}x} + e^{\sqrt{zy}x} & \sqrt{z/y}(e^{-\sqrt{zy}x} - e^{\sqrt{zy}x}) \\ \sqrt{y/z}(e^{-\sqrt{zy}x} - e^{\sqrt{zy}x}) & e^{-\sqrt{zy}x} + e^{\sqrt{zy}x} \end{pmatrix} \begin{pmatrix} v_0 \\ i_0 \end{pmatrix} \quad (\text{B.34})$$

$$\begin{pmatrix} v \\ i \end{pmatrix} = \begin{pmatrix} \cosh \sqrt{zy}x & -\sqrt{z/y}(\sinh \sqrt{zy}x) \\ -\sqrt{y/z}(\sinh \sqrt{zy}x) & \cosh \sqrt{zy}x \end{pmatrix} \begin{pmatrix} v_0 \\ i_0 \end{pmatrix} \quad (\text{B.35})$$

Remarque Nous pouvons faire apparaître deux « modes de courants » en posant

$$\begin{pmatrix} i_{m1} \\ i_{m2} \end{pmatrix} = P^{-1} \begin{pmatrix} v \\ i \end{pmatrix} = \frac{1}{2} \begin{pmatrix} -\sqrt{y/z}v + i \\ \sqrt{y/z}v + i \end{pmatrix} \quad (\text{B.36})$$

En effet, ces modes de propagation sont découplés. i_{m1} correspond à l'onde progressive et i_{m2} à l'onde rétrograde. Les valeurs propres de la matrice A s'interprètent alors comme les « constantes de propagation » de chacun de ces modes.

$$\frac{di_{m1}}{dx} = -\sqrt{zy}i_{m1}, \quad \frac{di_{m2}}{dX} = \sqrt{zy}i_{m2}. \quad (\text{B.37})$$

Annexe C

Polynômes orthogonaux

Nous présentons dans cette annexe quelques définitions et résultats mathématiques concernant les polynômes orthogonaux. Les polynômes orthogonaux font partie des familles de fonctions orthogonales qui sont classiquement utilisées pour résoudre des problèmes physiques.

Parmi les familles de fonctions orthogonales, on peut citer les fonctions sinusoïdales, utilisées par Fourier pour résoudre l'équation de la chaleur [Fou22], les ondelettes, qui permettent de représenter les signaux non-stationnaires par des méthodes *temps-échelle* [Fla98], et enfin, les polynômes orthogonaux qui sont depuis longtemps utilisés pour la résolution d'équations différentielles et l'approximation de fonctions.

C.1 Familles de fonctions orthogonales

La notion d'orthogonalité est liée au produit scalaire. Pour définir l'orthogonalité entre deux fonctions, commençons donc par expliciter un produit scalaire.

C.1.1 Produit scalaire et norme

Soit ρ une fonction positive, continue, sommable sur l'intervalle $[a, b]$. On appelle ρ fonction *poids* ou *mesure*. Soit alors $L^2_\rho(a, b)$ l'espace des fonctions à valeurs réelles telles que $\int_a^b f^2(x)\rho(x)dx$ converge.

L'application

$$\begin{aligned} (L^2_\rho(a, b))^2 &\longrightarrow \mathbb{R} \\ (f, g) &\longmapsto \langle f, g \rangle = \int_a^b f(x)g(x)\rho(x)dx \end{aligned} \tag{C.1}$$

définit un produit scalaire. La norme qui dérive de ce produit scalaire est donnée par

$$\|f\|^2 = \int_a^b f(x)^2 \rho(x)dx \tag{C.2}$$

C.1. Familles de fonctions orthogonales

On dit alors que f et g sont *orthogonales* dans $L^2_\rho(a, b)$ si

$$\langle f, g \rangle = 0. \quad (\text{C.3})$$

Une famille de fonctions $\{f_n\}_{n \in \mathbb{N}}$ est dite *orthonormée* relativement à un produit scalaire, si les fonction f_n sont orthogonales deux à deux et sont de plus *normées*.

$$\langle f_n, f_m \rangle = \begin{cases} 1 & \text{si } n = m, \\ 0 & \text{si } n \neq m. \end{cases} \quad (\text{C.4})$$

Exemple Soit $f_n(x) = \cos nx, n \in \mathbb{N}$. La famille $\{f_n\}_{n \in \mathbb{N}}$ est un système orthogonal si on prend pour définir le produit scalaire

$$a = 0, \quad b = 2\pi, \quad \rho(x) = \frac{1}{\pi}. \quad (\text{C.5})$$

C'est à dire

$$\langle f_n, f_m \rangle = \int_0^{2\pi} \cos(nx) \cos(mx) dx. \quad (\text{C.6})$$

Remarques

1. Dans les cas de fonctions à valeurs complexes, on utilise le produit scalaire *hermitien*,

$$\langle f, g \rangle = \int_a^b f(x) \overline{g(x)} \rho(x) dx. \quad (\text{C.7})$$

2. Il existe d'autres formes de produits scalaires que nous n'utiliserons pas pas la suite car ils ne sont pas liés à une forme linéaire. Par exemple

$$\langle f, g \rangle = \sum_{n=0}^{\infty} \left(\int \frac{\delta^{(n)}(x)}{n!} f(x) dx \right) \left(\int \frac{\delta^{(n)}(x)}{n!} g(x) dx \right) \quad (\text{C.8})$$

définit le produit scalaire pour lequel la base canonique des polynômes est orthogonale.

C.1.2 Forme linéaire et moments

Un produit scalaire défini par une expression de la forme (C.1) ne dépend que du produit $f(x)g(x)$. C'est donc un produit scalaire lié à la *forme linéaire*

$$c : (L^2_\rho(a, b)) \longrightarrow \mathbb{R} \quad (\text{C.9})$$

telle que

$$\langle f, g \rangle = c(f(x)g(x)) \quad (\text{C.10})$$

Cette forme linéaire est parfaitement définie par ses *moments*

$$c_n = c(x^n), \quad n \in \mathbb{N} \quad (\text{C.11})$$

L'étude des fonctions orthogonales est un vaste sujet, lié aux espaces de Hilbert. Nous nous limiterons pour la suite aux polynômes orthogonaux. Cependant des approfondissements pourront être trouvés par exemple dans [Wal94] et [San59].

C.2 Familles de polynômes orthogonaux

Soit un produit scalaire pouvant s'exprimer sous la forme de l'équation (C.1). Nous recherchons une famille de polynômes $\{P_n\}_{n \in \mathbb{N}}$ qui soit orthogonale par rapport à ce produit scalaire.

C.2.1 Calcul direct

Supposons que l'on écrive $P_n(x) = (a_0 + a_1x + \dots + a_{n-1}x^{n-1} + x^n)a_n$. P_n est orthogonal avec tous les P_k , pour $k < n$. Donc P_n est aussi orthogonal avec tous les x^k . Le calcul de P_n peut alors se faire directement en posant

$$c(x^k P_n(x)) = 0, \quad k = 0, \dots, n-1. \quad (\text{C.12})$$

C'est-à-dire que les a_k sont solutions du système linéaire

$$\begin{cases} c_0 a_0 + \dots + c_{n-1} a_{n-1} = -c_n \\ \vdots \\ c_{n-1} a_0 + \dots + c_{2n-2} a_{n-1} = -c_{2n-1}. \end{cases} \quad (\text{C.13})$$

La solution de ce système s'exprime en fonction du *déterminant de Hankel* [Ise88]

$$H_n = \begin{vmatrix} c_0 & \dots & c_{n-1} \\ \vdots & & \vdots \\ c_{n-1} & \dots & c_{2n-2} \end{vmatrix}. \quad (\text{C.14})$$

On pourra vérifier en développant le déterminant par rapport à sa dernière ligne que le polynôme P_n s'exprime alors par

$$P_n(x) = \frac{a_n}{H_n} \begin{vmatrix} c_0 & \dots & c_n \\ \vdots & & \vdots \\ c_{n-1} & \dots & c_{2n-1} \\ 1 & \dots & x^n \end{vmatrix}. \quad (\text{C.15})$$

Exemple Calculons le polynôme de Legendre d'ordre trois.

Les polynômes de Legendre sont orthogonaux par rapport au produit scalaire défini par

$$a = -1, \quad b = 1, \quad \rho(x) = 1. \quad (\text{C.16})$$

Les premiers moments de la forme linéaire associée sont

$$\begin{aligned} c_0 &= \int_{-1}^1 dx = 2, & c_1 &= \int_{-1}^1 x dx = 0, & c_2 &= \int_{-1}^1 x^2 dx = \frac{2}{3}, \\ c_3 &= \int_{-1}^1 x^3 dx = 0, & c_4 &= \int_{-1}^1 x^4 dx = \frac{2}{5}, & c_5 &= \int_{-1}^1 x^5 dx = 0. \end{aligned} \quad (\text{C.17})$$

C.2. Familles de polynômes orthogonaux

Le déterminant de Hankel d'ordre trois est donc

$$H_3 = \begin{vmatrix} 2 & 0 & \frac{2}{3} \\ 0 & \frac{2}{3} & 0 \\ \frac{2}{3} & 0 & \frac{2}{5} \end{vmatrix} = 32/135. \quad (\text{C.18})$$

Et

$$P_3(x) = \frac{135a_3}{32} \begin{vmatrix} 2 & 0 & \frac{2}{3} & 0 \\ 0 & \frac{2}{3} & 0 & \frac{2}{5} \\ \frac{2}{3} & 0 & \frac{2}{5} & 0 \\ 1 & x & x^2 & x^3 \end{vmatrix} = \frac{135a_3}{32} \left(-\frac{32}{225}x + \frac{32}{135}x^3 \right) = a_3 \left(x^3 - \frac{3}{5}x \right). \quad (\text{C.19})$$

Le choix du coefficient a_3 peut être fait pour normaliser le polynôme P_3 . Mais la forme classique des polynômes de Legendre est obtenue pour $a_3 = \frac{5}{2}$.

$$P_3(x) = \frac{5}{2}x^3 - \frac{3}{2}x. \quad (\text{C.20})$$

C.2.2 Récurrence à trois termes

La définition des polynômes orthogonaux est équivalente au fait qu'ils vérifient une récurrence à trois termes (théorème de Shohat-Favard 1936). Si les polynômes P_n (unitaires¹) sont orthogonaux par rapport à une forme c régulière, alors il existe deux suites de constantes $(\beta_n)_n, (\gamma_n)_n$ telle que

$$\begin{aligned} P_{-1} &= 0, \quad P_0 = 1, \\ P_{n+1} &= (x - \beta_n)P_n - \gamma_n P_{n-1}, \quad \forall n \geq 0, \gamma_n \neq 0. \end{aligned} \quad (\text{C.21})$$

On peut calculer les $(\beta_n)_n, (\gamma_n)_n$ en fonction de c

$$\beta_n = \frac{c(xP_n^2)}{c(P_n^2)}, \quad \gamma_n = \frac{c(P_n^2)}{c(P_{n-1}^2)}. \quad (\text{C.22})$$

Exemple Reprenons le produit scalaire défini par (C.16). Posons $P_0(x) = 1$ et $P_1(x) = x$. Nous avons

$$\beta_1 = \frac{c(x^3)}{c(x^2)} = \frac{c_3}{c_2} = 0, \quad \gamma_1 = \frac{c(x^2)}{c(1)} = \frac{c_2}{c_0} = \frac{1}{3}. \quad (\text{C.23})$$

Donc

$$P_2(x) = xP_1(x) - \frac{1}{3}P_0(x) = x^2 - \frac{1}{3}. \quad (\text{C.24})$$

C'est, à un coefficient multiplicateur près, le polynôme de Legendre d'ordre deux.

¹Un polynôme est unitaire si son coefficient dominant est pris égal à un : $a_n = 1$ si P_n est de degré n .

C.2.3 Fonctions génératrices

On appelle fonction génératrice d'une famille de polynôme $\{P_n\}_{n \in \mathbb{N}}$, une fonction $G(x, t)$ telle que

$$G(x, t) = \sum_{n=0}^{\infty} P_n(x) t^n. \quad (\text{C.25})$$

Exemple La fonction génératrice des polynômes de Legendre est

$$G(x, t) = \frac{1}{\sqrt{t^2 - 2tx + 1}} \quad (\text{C.26})$$

En effet, on a le développement en série de Taylor selon la variable t , au voisinage de zéro,

$$\frac{1}{\sqrt{t^2 - 2tx + 1}} = 1 + xt + \left(\frac{3}{2}x^2 - \frac{1}{2}\right)t^2 + \left(\frac{5}{2}x^3 - \frac{3}{2}x\right)t^3 + \left(\frac{3}{8} - \frac{15}{4}x^2 + \frac{35}{8}x^4\right)t^4 + \dots \quad (\text{C.27})$$

C.3 Polynômes orthogonaux usuels

C.3.1 Polynômes de Legendre

Nous avons déjà vu en exemple les polynômes de Legendre, orthogonaux par rapport au produit scalaire défini par

$$a = -1, \quad b = 1, \quad \rho(x) = 1. \quad (\text{C.28})$$

Rappelons la fonction génératrice

$$G(x, t) = \frac{1}{\sqrt{t^2 - 2tx + 1}} = \sum_{n=0}^{\infty} P_n(x) t^n. \quad (\text{C.29})$$

Les premiers polynômes de Legendre sont

$$\begin{aligned} P_0(x) &= 1, & P_1(x) &= x, & P_2(x) &= \frac{3}{2}x^2 - \frac{1}{2}, \\ P_3(x) &= \frac{5}{2}x^3 - \frac{3}{2}x, & P_4(x) &= \frac{35}{8}x^4 - \frac{15}{4}x^2 + \frac{3}{8}, & P_5(x) &= \frac{63}{8}x^5 - \frac{35}{4}x^3 + \frac{15}{8}x, \dots \end{aligned} \quad (\text{C.30})$$

Ils sont représentés sur la figure C.1.

C.3. Polynômes orthogonaux usuels

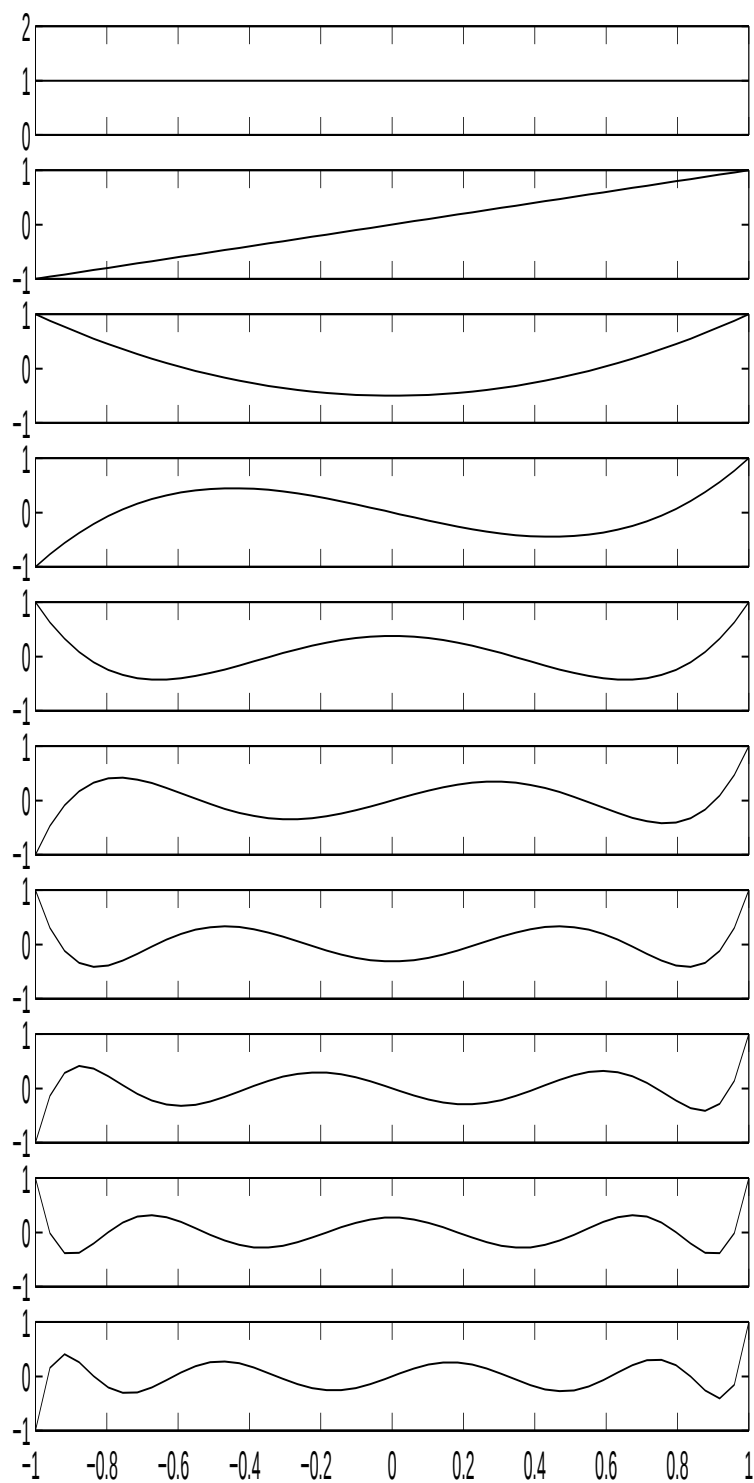


FIG. C.1 – Premiers polynômes de Legendre.

C.3.2 Polynômes de Tchebycheff de première espèce

Les polynômes de Tchebycheff de première espèce sont orthogonaux par rapport au produit scalaire défini par

$$a = -1, \quad b = 1, \quad \rho(x) = \frac{1}{\sqrt{1-x^2}}. \quad (\text{C.31})$$

Ils ont pour fonction génératrice

$$G(x, t) = \frac{1 - xt}{1 - 2xt + t^2} = \sum_{n=0}^{\infty} T_n(x) t^n. \quad (\text{C.32})$$

Les premiers polynômes de Tchebycheff de première espèce sont

$$\begin{aligned} T_0(x) &= 1, & T_1(x) &= x, & T_2(x) &= 2x^2 - 1, \\ T_3(x) &= 4x^3 - 3x, & T_4(x) &= 8x^4 - 8x^2 + 1, & T_5(x) &= 16x^5 - 20x^3 + 5x, \dots \end{aligned} \quad (\text{C.33})$$

Ils sont représentés sur la figure C.2.

C.3.3 Polynômes de Tchebycheff de deuxième espèce

Les polynômes de Tchebycheff de deuxième espèce sont orthogonaux par rapport au produit scalaire défini par

$$a = -1, \quad b = 1, \quad \rho(x) = \sqrt{1-x^2}. \quad (\text{C.34})$$

Ils ont pour fonction génératrice

$$G(x, t) = \frac{1}{1 - 2xt + t^2} = \sum_{n=0}^{\infty} U_n(x) t^n. \quad (\text{C.35})$$

Les premiers polynômes de Tchebycheff de deuxième espèce sont

$$\begin{aligned} U_0(x) &= 1, & U_1(x) &= 2x, & U_2(x) &= 4x^2 - 1, \\ U_3(x) &= 8x^3 - 4x, & U_4(x) &= 16x^4 - 12x^2 + 1, & U_5(x) &= 32x^5 - 32x^3 + 6x, \dots \end{aligned} \quad (\text{C.36})$$

Ils sont représentés sur la figure C.3.

C.3.4 Polynômes de Gegenbauer

Les polynômes de Gegenbauer, ou *ultra-sphériques* sont orthogonaux par rapport au produit scalaire défini par

$$a = -1, \quad b = 1, \quad \rho(x) = (1 - x^2)^{\lambda - \frac{1}{2}}. \quad (\text{C.37})$$

Ils généralisent donc les polynômes de

- Legendre, correspondant donc au cas particulier $\lambda = \frac{1}{2}$;
- Tchebycheff de première espèce, $\lambda = 0$;
- Tchebycheff de seconde espèce, $\lambda = 1$.

C.3. Polynômes orthogonaux usuels

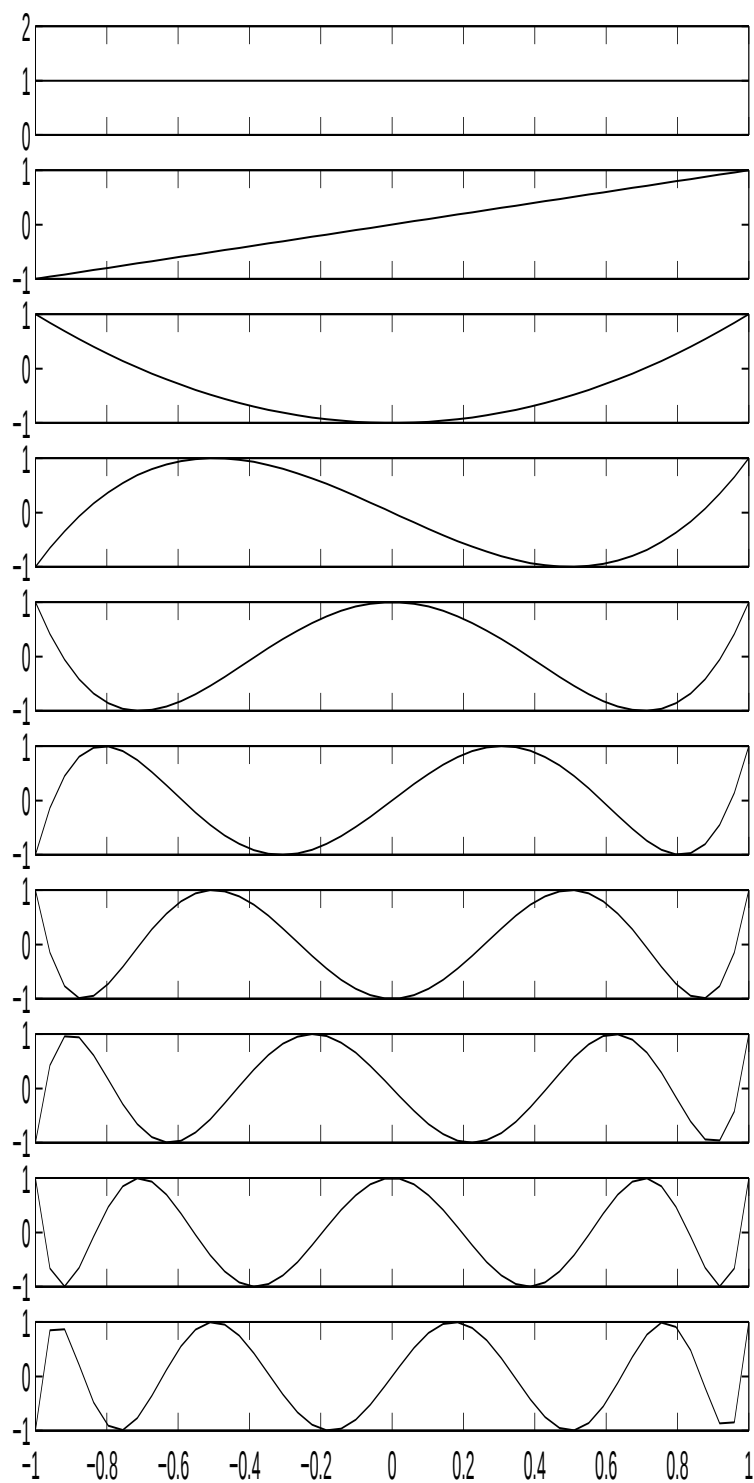


FIG. C.2 – Premiers polynômes de Tchebycheff de première espèce.

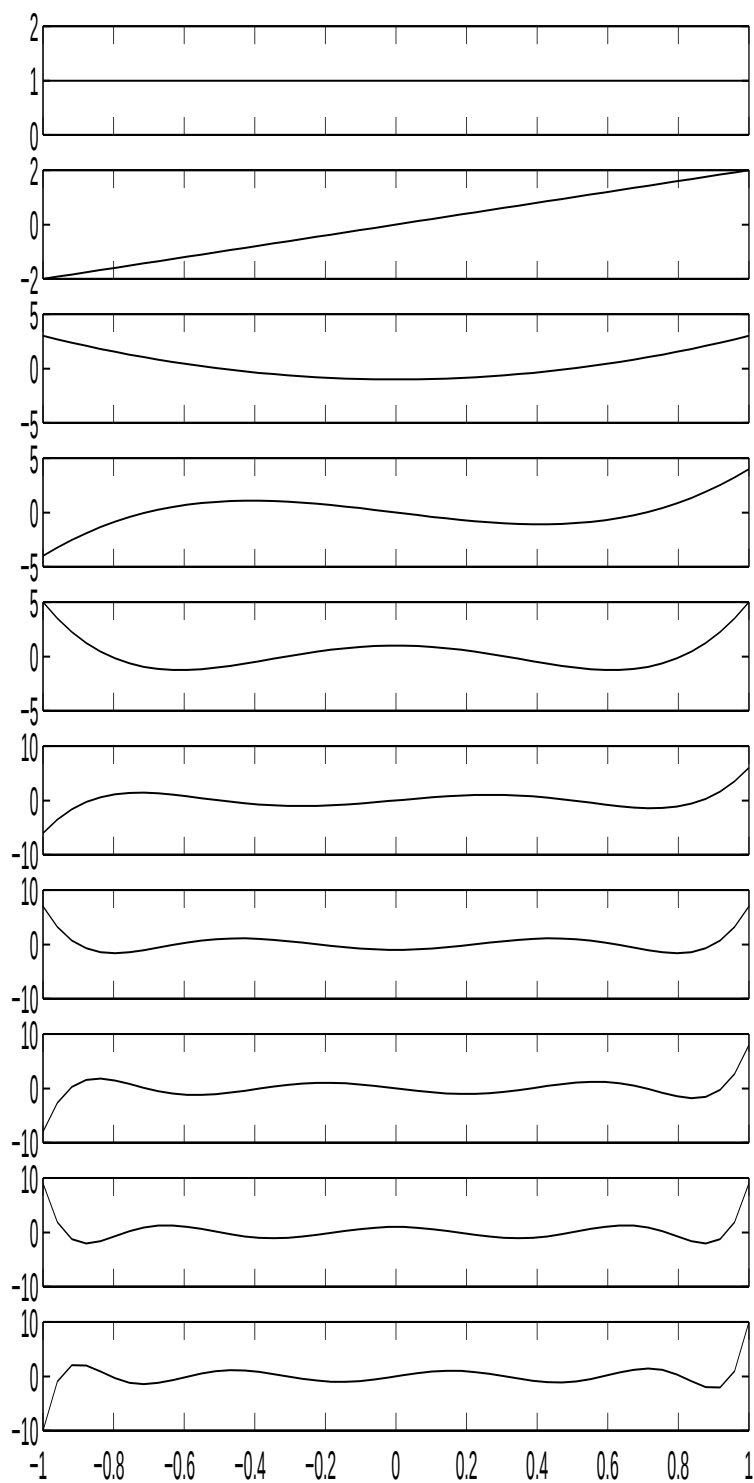


FIG. C.3 – Premiers polynômes de Tchebycheff de deuxième espèce.

C.3. Polynômes orthogonaux usuels

C.3.5 Polynômes de Jacobi

Les polynômes de Jacobi sont orthogonaux par rapport au produit scalaire défini par

$$a = -1, \quad b = 1, \quad \rho(x) = (1-x)^\alpha(1+x)^\beta. \quad (\text{C.38})$$

Ils généralisent donc les polynômes de

- Legendre, correspondant donc au cas particulier $\alpha = \beta = 0$;
- Tchebytscheff de première espèce, $\alpha = \beta = -\frac{1}{2}$;
- Tchebytscheff de seconde espèce, $\alpha = \beta = +\frac{1}{2}$;
- Gegenbauer, avec $\alpha = \beta = \lambda - \frac{1}{2}$.

C.3.6 Polynômes de Hermite

Les polynômes de Hermite sont orthogonaux par rapport au produit scalaire défini par

$$a = -\infty, \quad b = +\infty, \quad \rho(x) = e^{-x^2}. \quad (\text{C.39})$$

Ils ont pour fonction génératrice, à un coefficient multiplicatif $n!$ près,

$$G(x, t) = e^{-2xt-t^2} = \sum_{n=0}^{\infty} H_n(x) \frac{t^n}{n!}. \quad (\text{C.40})$$

Les premiers polynômes de Hermite sont

$$\begin{aligned} H_0(x) &= 1, & H_1(x) &= 2x, & H_2(x) &= 4x^2 - 2, \\ H_3(x) &= 8x^3 - 12x, & H_4(x) &= 16x^4 - 48x^2 + 12, & H_5(x) &= 32x^5 - 160x^3 + 120x, \dots \end{aligned} \quad (\text{C.41})$$

Ils sont représentés sur la figure C.4.

C.3.7 Polynômes de Laguerre

Les polynômes de Laguerre sont orthogonaux par rapport au produit scalaire défini par

$$a = 0, \quad b = \infty, \quad \rho(x) = t^\alpha e^{-t}. \quad (\text{C.42})$$

Ils ont pour fonction génératrice

$$G(x, t) = \frac{1}{(1-t)^{\alpha+1}} e^{\frac{-xt}{1-t}} = \sum_{n=0}^{\infty} L_{\alpha,n}(x) t^n. \quad (\text{C.43})$$

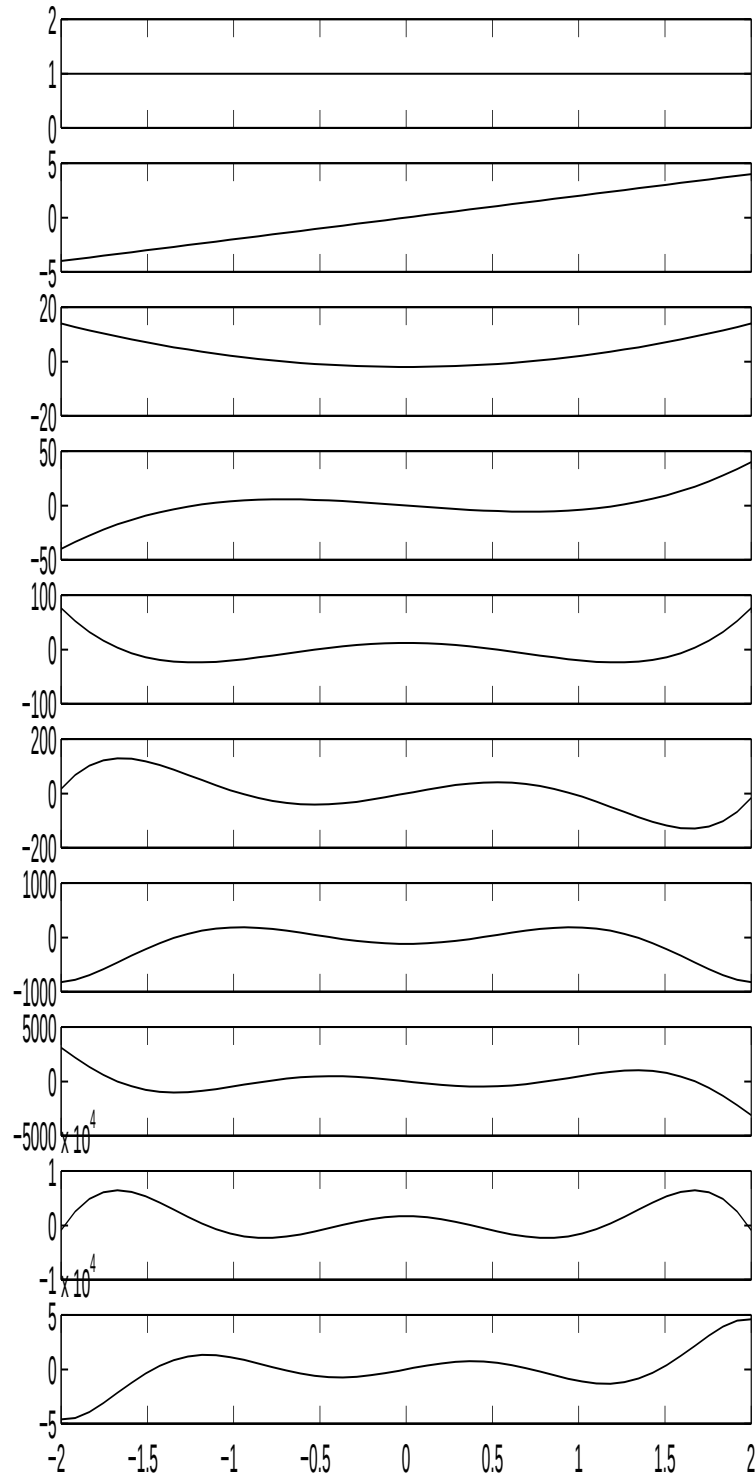


FIG. C.4 – Premiers polynômes de Hermite.

C.3. Polynômes orthogonaux usuels

Les premiers polynômes de Laguerre pour $\alpha = 0$ sont

$$\begin{aligned} L_{0,0}(x) &= 1, & L_{0,1}(x) &= 1 - x, \\ L_{0,2}(x) &= 1 - 2x + \frac{1}{2}x^2, & L_{0,3}(x) &= 1 - 3x + \frac{3}{2}x^2 - \frac{1}{6}x^3, \\ L_{0,4}(x) &= 1 - 4x + 3x^2 - \frac{2}{3}x^3 + \frac{1}{24}x^4, \\ L_{0,5}(x) &= 1 - 5x + 5x^2 - \frac{5}{3}x^3 + \frac{5}{24}x^4 - \frac{1}{120}x^5, \dots \end{aligned} \tag{C.44}$$

Ils sont représentés sur la figure C.5.

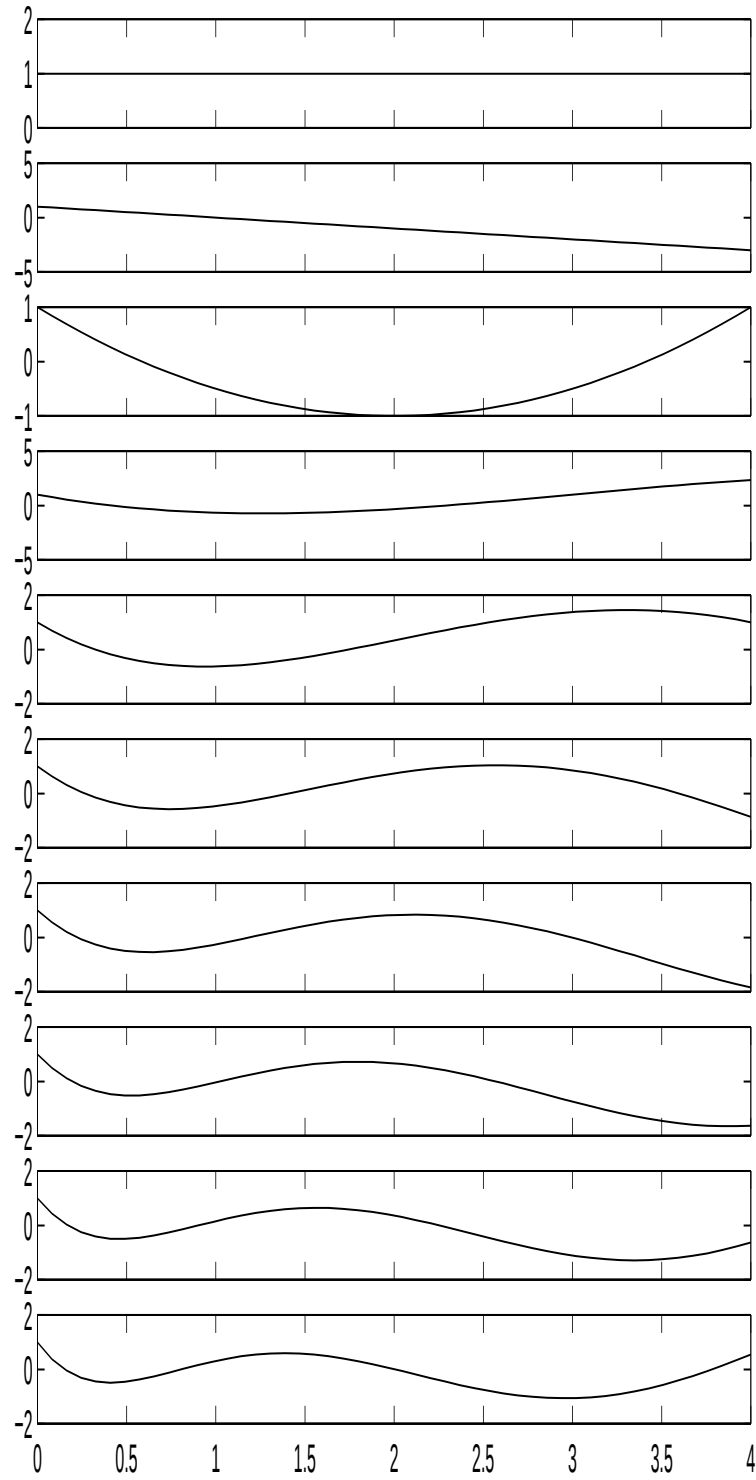


FIG. C.5 – Premiers polynômes de Laguerre.

C.3. Polynômes orthogonaux usuels

Annexe D

Caractéristiques électriques d'un conducteur cylindrique

Nous précisons ici les calculs permettant de définir les caractéristiques électriques (constantes réparties) d'un conducteur cylindrique. Ces expressions apparaissent en particulier dans la modélisation de l'effet de peau.

Nous supposons que le courant qui traverse le conducteur est invariant dans le temps (régime continu, ou basse fréquence). Les filets de courants sont alors uniformément répartis sur toute la section du conducteur.

$$j = \frac{i}{S} = \frac{i}{\pi \rho_0^2} \quad (\text{D.1})$$

La résistance totale s'exprime en fonction de la section S , de la longueur x du conducteur et de la conductivité σ du matériau.

$$R = \frac{x}{\sigma S} = \frac{x}{\pi \rho_0^2 \sigma} = r x \quad (\text{D.2})$$

Calculons maintenant le champ d'induction magnétique $\mathbf{B}$ à l'intérieur et à l'extérieur du cylindre. La symétrie nous permet de conclure que ce champ est purement orthoradial et uniquement fonction de ρ . Puisque nous considérons un régime statique, nous pouvons appliquer le théorème d'Ampère sur un circuit circulaire centré sur l'axe de symétrie du cylindre (figure D.1).

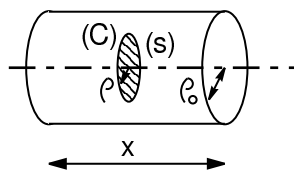


FIG. D.1 – Contour fermé (C) pour l'application du théorème d'Ampère.

$$\oint_{(C)} \mathbf{B} d\mathbf{l} = \mu \iint_{(s)} \mathbf{j} ds \quad (\text{D.3})$$

À l'intérieur du cylindre ($\rho < \rho_0$),

$$2\pi\rho B = \mu\pi\rho^2 j \text{ d'où } B = \frac{\mu i}{2\pi\rho_0^2}\rho \quad (\text{D.4})$$

À l'extérieur du cylindre ($\rho > \rho_0$),

$$2\pi\rho B = \mu\pi\rho_0^2 j \text{ d'où } B = \frac{\mu i}{2\pi\rho} \quad (\text{D.5})$$

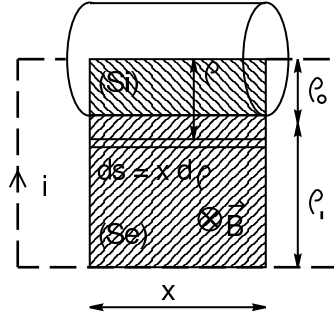


FIG. D.2 – Calcul du flux d'induction magnétique à l'intérieur et à l'extérieur du conducteur.

Nous pouvons alors calculer le flux du champ au travers du conducteur, mais uniquement du côté où se referme le circuit (figure D.2).

$$\Phi_i = \iint_{(S_i)} \mathbf{B} ds = \int_0^{\rho_0} \frac{\mu i}{2\pi\rho_0^2} \rho x d\rho = \frac{\mu x i}{4\pi} \quad (\text{D.6})$$

Calculons de façon similaire le flux à l'extérieur du conducteur :

$$\Phi_e = \iint_{(S_e)} \mathbf{B} ds = \int_{\rho_0}^{\rho_1} \frac{\mu i}{2\pi\rho} x d\rho = \frac{\mu x i}{2\pi} \ln\left(\frac{\rho_1}{\rho_0}\right) \quad (\text{D.7})$$

Nous pouvons alors définir les inductances « interne » et « externe » du conducteur cylindrique par les relations :

$$L_i = \frac{\Phi_i}{i} = \frac{\mu x}{4\pi} \quad (\text{D.8})$$

$$L_e = \frac{\Phi_e}{i} = \frac{\mu x}{2\pi} \ln\left(\frac{\rho_1}{\rho_0}\right) \quad (\text{D.9})$$

Annexe E

Représentation matricielle de réseaux élémentaires

Dans cette annexe, nous présentons brièvement la représentation sous forme matricielle des réseaux élémentaires monophasés, puis diphasés. Ce dernier cas faisant apparaître le problème de la représentation des couplages, il nous permet de définir la représentation d'une admittance matricielle capacitive, par dualité avec le couplage inductif. Les définitions et les propriétés utilisées pourront être trouvées dans [BN96] et [Fel86].

E.1 Impédances et admittances monophasées

E.1.1 Résistance longitudinale

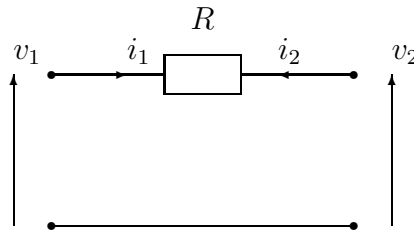


FIG. E.1 – Résistance longitudinale

Matrice d'impédance Elle est non définie.

Matrice d'admittance Elle est singulière (non inversible). En effet, les courants i_1 et i_2 sont liés ($i_1 = -i_2$).

$$\begin{pmatrix} i_1 \\ i_2 \end{pmatrix} = \begin{pmatrix} \frac{1}{R} & -\frac{1}{R} \\ -\frac{1}{R} & \frac{1}{R} \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} \quad (\text{E.1})$$

E.1. Impédances et admittances monophasées

Matrice de chaîne

$$\begin{pmatrix} v_1 \\ i_1 \end{pmatrix} = \begin{pmatrix} 1 & R \\ 0 & 1 \end{pmatrix} \begin{pmatrix} v_2 \\ -i_2 \end{pmatrix} \quad (\text{E.2})$$

E.1.2 Conductance transversale

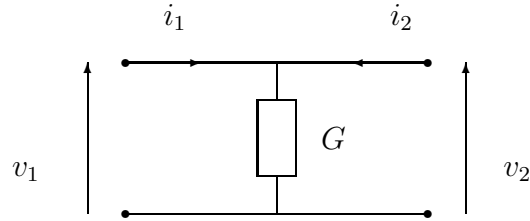


FIG. E.2 – Conductance transversale.

Matrice d'impédance Elle est singulière (non inversible). En effet, les tensions v_1 et v_2 sont liées ($v_1 = v_2$).

$$\begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \begin{pmatrix} \frac{1}{G} & \frac{1}{G} \\ \frac{1}{G} & \frac{1}{G} \end{pmatrix} \begin{pmatrix} i_1 \\ i_2 \end{pmatrix} \quad (\text{E.3})$$

Matrice d'admittance Elle est non définie.

Matrice de chaîne

$$\begin{pmatrix} v_1 \\ i_1 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ G & 1 \end{pmatrix} \begin{pmatrix} v_2 \\ -i_2 \end{pmatrix} \quad (\text{E.4})$$

E.1.3 Inductance longitudinale

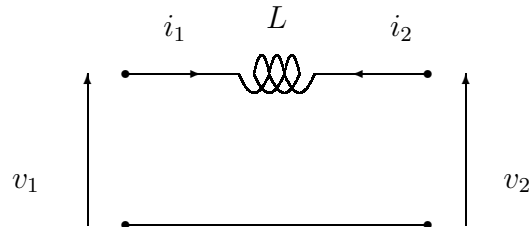


FIG. E.3 – Inductance longitudinale.

Matrice d'impédance Elle est non définie.

Matrice d'admittance Elle est singulière (non inversible). En effet, les courants i_1 et i_2 sont liés ($i_1 = -i_2$).

$$\begin{pmatrix} i_1 \\ i_2 \end{pmatrix} = \begin{pmatrix} \frac{1}{Lp} & -\frac{1}{Lp} \\ -\frac{1}{Lp} & \frac{1}{Lp} \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} \quad (\text{E.5})$$

Matrice de chaîne

$$\begin{pmatrix} v_1 \\ i_1 \end{pmatrix} = \begin{pmatrix} 1 & Lp \\ 0 & 1 \end{pmatrix} \begin{pmatrix} v_2 \\ -i_2 \end{pmatrix} \quad (\text{E.6})$$

E.1.4 Capacité transversale

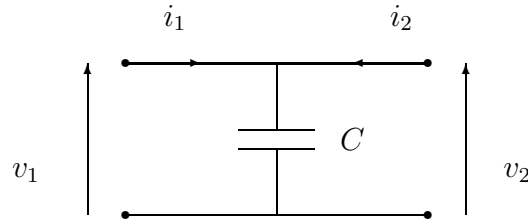


FIG. E.4 – Capacité transversale.

Matrice d'impédance Elle est singulière (non inversible). En effet, les tensions v_1 et v_2 sont liées ($v_1 = v_2$).

$$\begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \begin{pmatrix} \frac{1}{Cp} & \frac{1}{Cp} \\ \frac{1}{Cp} & \frac{1}{Cp} \end{pmatrix} \begin{pmatrix} i_1 \\ i_2 \end{pmatrix} \quad (\text{E.7})$$

Matrice d'admittance Elle est non définie.

Matrice de chaîne

$$\begin{pmatrix} v_1 \\ i_1 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ Cp & 1 \end{pmatrix} \begin{pmatrix} v_2 \\ -i_2 \end{pmatrix} \quad (\text{E.8})$$

E.2 Impédances et admittances biphasées

E.2.1 Inductances longitudinales couplées

Matrice d'impédance Elle est non définie.

E.2. Impédances et admittances biphasées

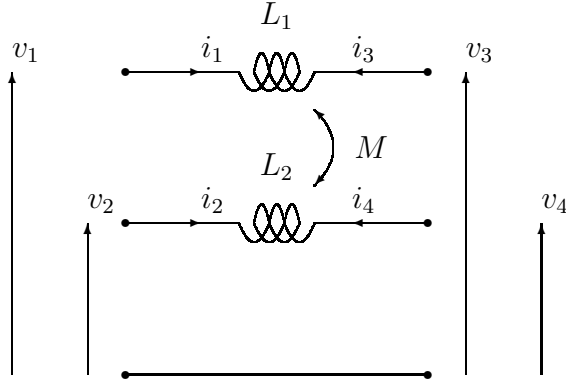


FIG. E.5 – Inductances longitudinales couplées.

Matrice d'admittance Les relations

$$\begin{cases} v_1 - v_3 = L_1 p i_1 + M p i_2 \\ v_2 - v_4 = M p i_1 + L_2 p i_2 \\ i_3 = -i_1 \\ i_4 = -i_2 \end{cases} \quad (\text{E.9})$$

Nous permettent de calculer la matrice d'admittance. Elle est singulière (non inversible).

$$\begin{pmatrix} i_1 \\ i_2 \\ i_3 \\ i_4 \end{pmatrix} = \frac{1}{(L_1 L_2 - M^2)p} \begin{pmatrix} L_2 & -M & -L_2 & M \\ -M & L_1 & M & -L_1 \\ -L_2 & M & L_2 & -M \\ M & -L_1 & -M & L_1 \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \\ v_3 \\ v_4 \end{pmatrix} \quad (\text{E.10})$$

Matrice de chaîne

$$\begin{pmatrix} v_1 \\ v_2 \\ i_1 \\ i_2 \end{pmatrix} = \left(\begin{array}{cc|cc} 1 & 0 & L_1 p & M p \\ 0 & 1 & M p & L_2 p \\ \hline 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{array} \right) \begin{pmatrix} v_3 \\ v_4 \\ -i_3 \\ -i_4 \end{pmatrix} \quad (\text{E.11})$$

E.2.2 Capacités transversales couplées

Matrice d'impédance Le couplage des capacités est défini par les relations duales du couplage des inductances.

$$\begin{cases} i_1 + i_3 = C_1 p v_1 + D p v_2 \\ i_2 + i_4 = D p v_1 + C_2 p v_2 \\ v_3 = v_1 \\ v_4 = v_2 \end{cases} \quad (\text{E.12})$$

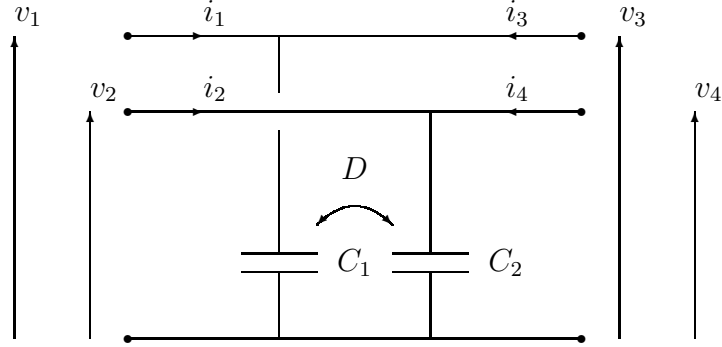


FIG. E.6 – Capacités transversales couplées.

D'où la matrice d'admittance, singulière (non inversible).

$$\begin{pmatrix} v_1 \\ v_2 \\ v_3 \\ v_4 \end{pmatrix} = \frac{1}{(C_1 C_2 - D^2)p} \begin{pmatrix} C_2 & -D & C_2 & -D \\ -D & C_1 & -D & C_1 \\ C_2 & -D & C_2 & -D \\ -D & C_1 & -D & C_1 \end{pmatrix} \begin{pmatrix} i_1 \\ i_2 \\ i_3 \\ i_4 \end{pmatrix} \quad (\text{E.13})$$

Matrice d'admittance Elle est non définie.

Matrice de chaîne

$$\begin{pmatrix} v_1 \\ v_2 \\ i_1 \\ i_2 \end{pmatrix} = \left(\begin{array}{cc|cc} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ \hline C_1 p & D p & 1 & 0 \\ D p & C_2 p & 0 & 1 \end{array} \right) \begin{pmatrix} v_3 \\ v_4 \\ -i_3 \\ -i_4 \end{pmatrix} \quad (\text{E.14})$$

Couplage par une capacité Exprimons la matrice de chaîne du schéma E.7.

Le schéma nous donne les relations

$$\begin{cases} v_1 = v_3 \\ v_2 = v_4 \\ i_1 + i_3 = C_{10} p v_3 + C_{12} p (v_2 - v_4) \\ i_2 + i_4 = C_{20} p v_4 + C_{12} p (v_4 - v_3) \end{cases} \quad (\text{E.15})$$

dont nous déduisons la matrice de chaîne

$$\begin{pmatrix} v_1 \\ v_2 \\ i_1 \\ i_2 \end{pmatrix} = \left(\begin{array}{cc|cc} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ \hline (C_{10} + C_{12})p & -C_{12}p & 1 & 0 \\ -C_{12}p & (C_{20} + C_{12})p & 0 & 1 \end{array} \right) \begin{pmatrix} v_3 \\ v_4 \\ -i_3 \\ -i_4 \end{pmatrix} \quad (\text{E.16})$$

E.3. Propriétés

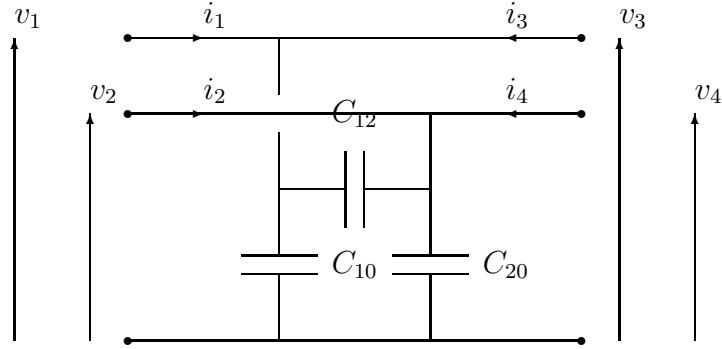


FIG. E.7 – Couplage par une capacité.

Nous pouvons donc identifier

$$\begin{cases} C_1 = C_{10} + C_{12} \\ C_2 = C_{20} + C_{12} \\ D = -C_{12} \end{cases} \quad (\text{E.17})$$

Les représentations par capacités couplées ou par capacité interphase sont donc équivalentes.

E.3 Propriétés

E.3.1 Passage entre les représentations

Regroupons les matrices de chaîne des impédances longitudinales sous la forme ($\mathbb{I}$ représente une matrice identité)

$$\begin{pmatrix} \mathbb{I} & Z \\ 0 & \mathbb{I} \end{pmatrix} \quad (\text{E.18})$$

Suivant le cas, Z peut être scalaire : R , Lp ,... ou bien matricielle : $\begin{pmatrix} L_{1p} & Mp \\ Mp & L_{2p} \end{pmatrix}$.

De même les admittances transversales

$$\begin{pmatrix} \mathbb{I} & 0 \\ Y & \mathbb{I} \end{pmatrix} \quad (\text{E.19})$$

Avec Y scalaire : G , Cp ,... ou bien matricielle : $\begin{pmatrix} C_{1p} & Dp \\ Dp & C_{2p} \end{pmatrix}$.

La matrice d'admittance correspondant à une impédance longitudinale est alors :

$$\begin{pmatrix} Z^{-1} & -Z^{-1} \\ -Z^{-1} & Z^{-1} \end{pmatrix} \quad (\text{E.20})$$

Et la matrice d'impédance correspondant à une admittance transversale :

$$\begin{pmatrix} Y^{-1} & Y^{-1} \\ Y^{-1} & Y^{-1} \end{pmatrix} \quad (\text{E.21})$$

Remarquons que l'assemblage d'éléments pour obtenir des réseaux en échelle se traduit par la multiplication de matrices de chaîne qui sont toujours définies.

E.3.2 Élément infinitésimal d'une ligne

Matrice de chaîne Il est admis qu'une section suffisamment petite d'une ligne de transmission admet une matrice de chaîne qui est la combinaison d'une impédance longitudinale et d'une admittance transversale.

$$\begin{pmatrix} \mathbb{I} & Z \\ Y & \mathbb{I} \end{pmatrix} \quad (\text{E.22})$$

La longueur de cet élément de ligne étant noté dx , les paramètres distribués (ou constantes réparties) z et y vérifient :

$$Z = zdx \quad Y = ydx \quad (\text{E.23})$$

Remarque sur la représentation L'élément de ligne est souvent représenté par un des schémas approchés

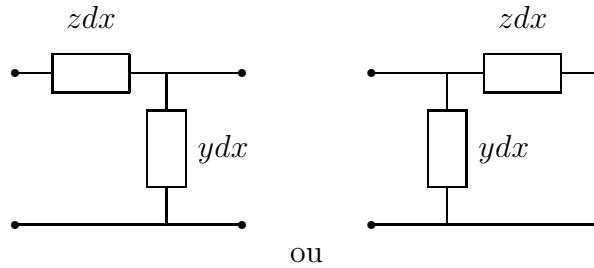


FIG. E.8 – Représentations approximatives d'un segment élémentaire de ligne.

Par multiplication matricielle, nous obtenons les matrices de chaînes (z et y peuvent être matricielles)

$$\begin{pmatrix} \mathbb{I} & zdx \\ 0 & \mathbb{I} \end{pmatrix} \begin{pmatrix} \mathbb{I} & 0 \\ ydx & \mathbb{I} \end{pmatrix} = \begin{pmatrix} \mathbb{I} + zy(dx)^2 & zdx \\ ydx & \mathbb{I} \end{pmatrix} \quad (\text{E.24})$$

ou bien

$$\begin{pmatrix} \mathbb{I} & 0 \\ ydx & \mathbb{I} \end{pmatrix} \begin{pmatrix} \mathbb{I} & zdx \\ 0 & \mathbb{I} \end{pmatrix} = \begin{pmatrix} \mathbb{I} & zdx \\ ydx & \mathbb{I} + yz(dx)^2 \end{pmatrix} \quad (\text{E.25})$$

C'est en *négligeant les termes d'ordre 2* que l'on retrouve l'expression de la matrice de chaîne de l'élément infinitésimal d'une ligne de transmission.

E.3.3 Mise en parallèle d'impédances longitudinales

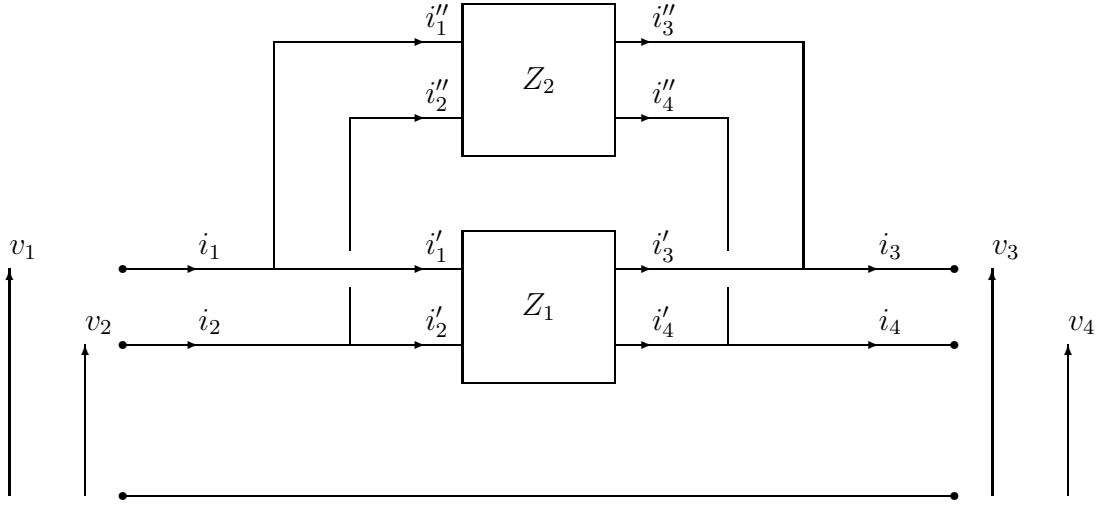


FIG. E.9 – Mise en parallèle d'impédances longitudinales.

D'après la définition des impédances longitudinales, nous avons les relations :

$$\begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \begin{pmatrix} v_3 \\ v_4 \end{pmatrix} + Z_1 \begin{pmatrix} -i'_3 \\ -i'_4 \end{pmatrix} \quad (\text{E.26})$$

$$\begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \begin{pmatrix} v_3 \\ v_4 \end{pmatrix} + Z_2 \begin{pmatrix} -i''_3 \\ -i''_4 \end{pmatrix} \quad (\text{E.27})$$

En supposant Z_1 et Z_2 inversibles, nous pouvons écrire

$$Z_1^{-1} \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = Z_1^{-1} \begin{pmatrix} v_3 \\ v_4 \end{pmatrix} + \begin{pmatrix} -i'_3 \\ -i'_4 \end{pmatrix} \quad (\text{E.28})$$

$$Z_2^{-1} \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = Z_2^{-1} \begin{pmatrix} v_3 \\ v_4 \end{pmatrix} + \begin{pmatrix} -i''_3 \\ -i''_4 \end{pmatrix} \quad (\text{E.29})$$

Additionnons ces deux équations et supposons que la somme $Z_1^{-1} + Z_2^{-1}$ est elle aussi inversible.

$$(Z_1^{-1} + Z_2^{-1}) \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = (Z_1^{-1} + Z_2^{-1}) \begin{pmatrix} v_3 \\ v_4 \end{pmatrix} + \begin{pmatrix} -i_3 \\ -i_4 \end{pmatrix} \quad (\text{E.30})$$

$$\begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \begin{pmatrix} v_3 \\ v_4 \end{pmatrix} + (Z_1^{-1} + Z_2^{-1})^{-1} \begin{pmatrix} -i_3 \\ -i_4 \end{pmatrix} \quad (\text{E.31})$$

La matrice chaîne équivalente est donc

$$\begin{pmatrix} v_1 \\ v_2 \\ i_1 \\ i_2 \end{pmatrix} = \begin{pmatrix} \mathbb{I} & (Z_1^{-1} + Z_2^{-1})^{-1} \\ 0 & \mathbb{I} \end{pmatrix} \begin{pmatrix} v_3 \\ v_4 \\ -i_3 \\ -i_4 \end{pmatrix} \quad (\text{E.32})$$

E.3.4 Mise en série d'impédances longitudinales

Par utilisation des matrices de chaîne, nous pouvons écrire

$$\begin{pmatrix} v_1 \\ v_2 \\ i_1 \\ i_2 \end{pmatrix} = \begin{pmatrix} \mathbb{I} & Z_1 \\ 0 & \mathbb{I} \end{pmatrix} \begin{pmatrix} \mathbb{I} & Z_2 \\ 0 & \mathbb{I} \end{pmatrix} \begin{pmatrix} v_3 \\ v_4 \\ -i_3 \\ -i_4 \end{pmatrix} \quad (\text{E.33})$$

$$\begin{pmatrix} v_1 \\ v_2 \\ i_1 \\ i_2 \end{pmatrix} = \begin{pmatrix} \mathbb{I} & Z_1 + Z_2 \\ 0 & \mathbb{I} \end{pmatrix} \begin{pmatrix} v_3 \\ v_4 \\ -i_3 \\ -i_4 \end{pmatrix} \quad (\text{E.34})$$

E.3.5 Couplage capacitif dans une impédance longitudinale

Considérons le réseau de la figure E.10.

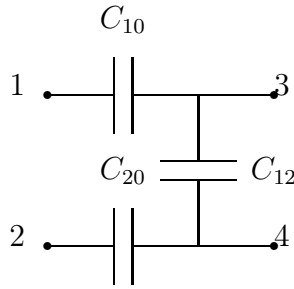


FIG. E.10 – Couplage par une capacité dans une impédance longitudinale.

La matrice d'admittance se déduit directement de la topologie.

$$\begin{pmatrix} i_1 \\ i_2 \\ i_3 \\ i_4 \end{pmatrix} = \begin{pmatrix} C_{10}p & 0 & -C_{10}p & 0 \\ 0 & C_{20}p & 0 & -C_{20}p \\ -C_{10}p & 0 & (C_{10} + C_{12})p & -C_{12}p \\ 0 & -C_{20}p & -C_{12}p & (C_{20} + C_{12})p \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \\ v_3 \\ v_4 \end{pmatrix} \quad (\text{E.35})$$

E.3. Propriétés

La structure de cette matrice d'admittance est différente de celle d'une impédance transversale (E.20).

$$\begin{pmatrix} i_1 \\ i_2 \\ i_3 \\ i_4 \end{pmatrix} = \begin{pmatrix} Z^{-1} & -Z^{-1} \\ -Z^{-1} & Z^{-1} \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \\ v_3 \\ v_4 \end{pmatrix} \quad (\text{E.36})$$

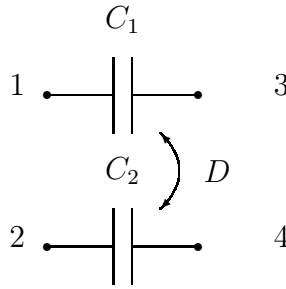


FIG. E.11 – Couplage capacitif dans une impédance longitudinale.

Par contre, avec des capacités couplées, telles que nous les avons définies précédemment (fig. E.11), la structure de la matrice d'admittance convient ($Z^{-1} = \begin{pmatrix} C_{1p} & Dp \\ Dp & C_{2p} \end{pmatrix} = Y$).

$$\begin{pmatrix} i_1 \\ i_2 \\ i_3 \\ i_4 \end{pmatrix} = \begin{pmatrix} C_{1p} & Dp & -C_{1p} & -Dp \\ Dp & C_{2p} & -Dp & -C_{2p} \\ -C_{1p} & -Dp & C_{1p} & Dp \\ -Dp & -C_{2p} & Dp & C_{2p} \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \\ v_3 \\ v_4 \end{pmatrix} \quad (\text{E.37})$$

Nous pouvons conclure sur cet exemple diphasé que la représentation du couplage capacitif doit se faire par des *capacités couplées* lorsque l'on manipule des impédances et admittances matricielles, et non par des *capacités interphases*. La généralisation au cas n -phasé est admise.

Bibliographie

- [AGV02] P. ABRY, P. GONÇALVES et J. L. VÉHEL (éds.) – *Lois d'échelle, fractales et ondelettes*, Hermès Science Publications, vol. 1, Lavoisier, Paris, 2002, ISBN 2-7462-0409-6.
- [AI87] M. AGUET et M. IANOZ – *Haute tension*, Traité d'électricité, d'électronique et d'électrotechnique, Dunod, 1987.
- [Bay54] M. BAYARD – *Théorie des réseaux de Kirchhoff, régime sinusoïdal et synthèse*, Collection technique et scientifique du CNET, Édition de la Revue d'optique, ministère des postes, télégraphes et téléphones, Paris, 1954.
- [Ber93] L. L. BERANEK – *Acoustics*, Acoustical Society of America (through the American Institute of Physics), 1993, First edition in 1954 at the Acoustic Laboratory, MIT.
- [BG83] A. BOUVIER et M. GEORGE – *Dictionnaire des Mathématiques*, 2^eéd., Presses Universitaires de France, Paris, 1983.
- [BI94] C. BREZINSKI et J. V. ISEGHEM – « Padé approximations », Techniques of Scientific Computing (Part 1), Numerical Methods for Solids (Part 1), Solution of Equations in R^n (Part 2) (P. G. Ciarlet et J. L. Lions, éds.), Handbook of Numerical Analysis, vol. 3, Elsevier Science, 1994.
- [BI95] — , « A Taste of Padé Approximation », *Acta Numerica* (1995), p. 53–103.
- [Blo20] A. BLONDEL – « Sur les méthodes modernes de calcul et sur le régime de fonctionnement des lignes de transmission d'énergie à haute tension », *Revue générale de l'électricité* **VIII** (1920), no. 5-6, p. 131–141(5), 163–174(6), 195–206(6), 227–236(6).
- [Blo37] — , « L'évolution des méthodes de calcul des phénomènes transitoires », *Revue générale de l'électricité* **XLI** (1937), no. 8-11, p. 227–240(8), 259–271(9), 298–311(10), 327–340(11), 579–596(19).
- [BN96] R. BOITE et J. NEIRYNCK – *Théorie des réseaux de kirchhoff*, Traité d'électricité de l'École Polytechnique Fédérale de Lausanne, vol. 4, Presses Polytechniques et Universitaires Romandes, CH-1015 Lausanne, 1996.
- [BPA02] J.-L. BATTAGLIA, L. PUIGSEGUR et A.KUSIAK – « Représentation non-entière du transfert de chaleur par diffusion. Utilité pour la caractérisation

Bibliographie

- et le contrôle non destructif thermique », *International journal of thermal science* (2002).
- [BR09] A. BLONDEL et C. L. ROY – « Calcul des lignes de transport d'énergie à courants alternatifs en tenant compte de la capacité et de la perdittance réparties », *La Lumière électrique* **VII (2)** (1909), no. 38-52, p. 355–363(38), 387–394(39), 99–103(43), 131–136(44), 387–395(52).
- [Bro00] M. BROUSSELY – « Réduction de modèles thermiques par la théorie des réseaux, application à la surveillance d'une machine asynchrone par couplage d'un modèle thermique réduit avec un schéma équivalent électrique », Thèse de doctorat, ENSMA, Poitiers, 2000.
- [Bru31] O. BRUNE – « Synthesis finite two terminal network whose driving point impedance is prescribed function frequency », *Journal of mathematics and physics* **10** (1931), p. 191–236.
- [Car93] P. A. CARRÉ – « Brève histoire des télécommunications. Du réseau simple aux réseaux pluriels », *Le réseau intelligent*, Les mémentos techniques des conférences France Télécom recherche, no. 1, octobre 1993, <http://www.rd.francetelecom.fr/fr/conseil/mento1/index.html>.
- [Cau26] W. CAUER – « Die Verwirklichung von Wechselstromwiderständen vorgeschriebener Frequenzabhängigkeit », *Archiv für Elektrotechnik* **17** (1926), p. 355–388.
- [cem04] *Colloque international de compatibilité électromagnétique* – Toulouse, mars 2004.
- [Cha96] M. CHARBIT – *Éléments de théorie du signal, aspects aléatoires*, Collection pédagogique de télécommunication, Ellipses, Paris, 1996, ISBN 2-7298-5640-4.
- [Cla94] P. R. CLAYTON – *Analysis of Multiconductor Transmission Lines*, Wiley series in Microwave and Optical Engineering, J. Wileys and sons, 1994.
- [CM95] E. CAUER et W. MATHIS – « Wilhelm Cauer (1900-1945) », *Archiv für Elektronik und Übertragungstechnik* **49** (1995), no. 5/6, p. 243–251.
- [Coh35] L. COHEN – *Théorie du circuit électrique de Heaviside. Application aux filtres électriques, aux câbles sous-marins, aux lignes de transmission d'énergie et aux lignes artificielles.*, Léon Eyrolles, Paris, 1935, Introduction de M. I. Pupin, traduit de l'anglais par F. Sarrat, avant-propos par Paul Janet.
- [Cor91] C. CORROYER – « Protection des réseaux, généralités », *Génie électrique*, no. D 4 800, Techniques de l'ingénieur, 1991.
- [CZ95] R. F. CURTAIN et H. ZWART – *An Introduction to Infinite-Dimensional Linear Systems Theory*, Texts in Applied Mathematics, Springer-Verlag, New-York, etc., 1995, ISBN 0-387-94475-3.

- [DC17] X. DEVAUX-CHARBONNEL – « La contribution des ingénieurs français à la téléphonie à grande distance : Vaschy et Barbarat », *Revue générale de l'électricité* **II** (1917), no. 8, p. 288–295.
- [DI87] A. DRAUX et P. V. INGELANDT – *Polynômes orthogonaux et approximations de Padé*, Collection Langages et Algorithmes de l'informatique, Éditions Technip, Paris, 1987.
- [Duv91] P. DUVAUT – *Traitement du signal*, Traité des nouvelles technologies, Hermès, Paris, 1991, ISBN 2-86601-278-X.
- [Esc97] J.-M. ESCANÉ – *Réseaux d'énergie électrique. Modélisation : lignes, câbles*, collection de la DER d'EDF, Édition Eyrolles, 1997.
- [Fal32] J. FALLOU – « Propagation des courants de haute fréquence polyphasés le long des lignes aériennes de transport d'énergie affectées de courts-circuits ou de défauts d'isolement. Application à la protection sélective des réseaux électriques. », Thèse, Faculté des sciences de l'université de Paris, juin 1932.
- [Fel86] M. FELDMANN – *Théorie des réseaux et systèmes linéaires*, deuxième éd., Collection technique et scientifique des télécommunications, Eyrolles, Paris, 1986.
- [Fla98] P. FLANDRIN – *Temps-fréquence*, deuxième éd., collection traitement du signal, Hermès, Paris, 1998, ISBN 2-86601-700-5.
- [Fos24a] R. M. FOSTER – « A reactance theorem », *The Bell system technical journal* (1924), p. 259–267.
- [Fos24b] — , « Theorems Regarding the Driving-Point Impedance of Two Mesh Circuits », *The Bell System Technical Journal* (1924).
- [Fou22] M. FOURIER – *Théorie analytique de la chaleur*, Firmin Didot, père et fils, Paris, 1822, Réédition Jacques Gabay, 1988.
- [Fou93] G. FOURNET – « Électromagnétisme – Application à l'Électrotechnique », *Traité de Génie Électrique*, vol. D1, Techniques de l'Ingénieur, 1993.
- [Fry29] T. FRY – « The use of continued fractions in the design of electrical networks », *Bull. Amer. Math. Soc.*, vol. 35, 1929, p. 463–498.
- [GABG70] J. GEORGE A. BAKER et J. L. GAMMEL (éds.) – *The Padé Approximant in Theoretical Physics*, Mathematics in Science and Engineering, Academic Press, New-York et Londres, 1970.
- [Gar64] P. R. GARABEDIAN – *Partial Differential Equations*, John Wiley & Sons, New York, London, Sydney, 1964.
- [Gar76] C. GARY – « Approche complète de la propagation multifilaire en haute fréquence par utilisation des matrices complexes », *Bulletin de la direction des études et recherches – série B, Réseaux électriques, matériels électriques* 3/4, EDF, 1976.

Bibliographie

- [Gau75] W. GAUTSCHI – « Computational Methods in Special Functions – A Survey », *Theory and Application of Special Functions* (R. A. Askey, éd.), Academic Press, 1975.
- [GEC] GEC Alsthom T&D – *Protection des réseaux de transport d'énergie électrique à haute tension*, Session technique de formation. FP5.
- [Har98] K. HARKER – *Power Commissioning and Maintenance Practice*, IEE Power Series, no. 24, The Institution of Electrical Engineers, London, 1998.
- [Ise88] J. V. ISEGHEM – « Polynômes orthogonaux : synthèse des présentations des polynômes orthogonaux classiques, extensions », PUB. IRMA, Lille, 1988.
- [Joh94] P. JOHANNET – « Les phénomènes de propagation – Généralités sur l'analyse modale », Direction des Études et Recherches – Service Matériel Électrique 94NR00066, EDF, 1994.
- [Joh97] — , « Lignes aériennes : chutes de tension », *Traité de génie électrique*, vol. D4, Techniques de l'ingénieur, 1997.
- [Joh98] — , « Constantes linéiques des lignes et des câbles : Formulaire pour le calcul des chutes de tension et des phénomènes de propagation », Direction des Études et Recherches – Service Matériel Électrique HM-76/98/022/A, EDF, 1998.
- [JT80] W. B. JONES et W. J. THRON – *Continued fractions : analytic theory and applications*, Encyclopedia of mathematics and its applications, vol. 11, Addison-Wesley Pub. Co., 1980.
- [Kai80] T. KAILATH – *Linear systems*, information and system sciences series, Prentice-Hall, 1980, ISBN : 0-13-536961-4.
- [Khi64] A. Y. KHINCHIN – *Continued Fractions*, 3^eéd., Phoenix Science Series, The University of Chicago Press, Chicago et Londres, 1964.
- [Lam78] C. LAMOUREUX – « Fonctions Cylindriques de Bessel », cours de première année d'étude de l'École Centrale des Arts et Manufactures, 1978.
- [LBBS01] P. LAGONOTTE, M. BROUSSELY, Y. BERTIN et J.-B. SAULNIER – « Improvement of Thermal Nodal Models with Negative Compensation Capacitors », *The European Physical Journal Applied Physics* **13** (2001), no. 3, p. 177–194.
- [LBS99] P. LAGONOTTE, Y. BERTIN et J.-B. SAULNIER – « Analyse de la qualité de modèles nodaux réduits à l'aide de la méthode des quadripôles », *International Journal of Thermal Sciences* **38** (1999), p. 51–65.
- [LK92] S. LIN et E. S. KUH – « Pade Approximation Applied to Transient Simulation of Lossy Coupled Transmission Lines », *Proc. IEEE Multi-Chip Module Conference*, 1992, p. 52–55.
- [LN00] P. LAGONOTTE et F. NDAGIJIMANA – « Modélisation de la propagation sur des lignes électriques », *Colloque international de compatibilité électromagnétique* (Clermont-Ferrand (France)), mars 2000, p. 110–115.

- [LPC00] P. LAGONOTTE, M. POLOUJADOFF et A. CALVAER – « Introduction to an Improved Modelling of Electric Line Propagation », *IEEE Power Engineering Letters* (2000), p. 49–51.
- [LSOH99] F. LEVRON, J. SABATIER, A. OUSTALOUP et L. HABSIEGER – « From partial differential equations of propagative recursive systems to non integer differentiation », *Fractional calculus and applied analysis* **2** (1999), no. 3, p. 245–264.
- [MAB⁺00] D. MAILLET, S. ANDRÉ, J. C. BATSALE, A. DEGIOVANNI et C. MOYNE – *Thermal Quadrupoles, Solving the Heat Equation through Integral Transforms*, John Wiley and Sons, Ltd, Chichester, Weinheim, New-York, Brisbane, Singapour, Toronto, 2000.
- [Mar90] J. MARTINET – *Thermocinétique approfondie*, Technique et Documentation, Lavoisier, Paris, 1990.
- [Mat94] D. MATIGNON – « Représentation en variables d'état de guides d'ondes avec dérivation fractionnaire », Thèse, université de Paris XI, novembre 1994.
- [Mat02] — , « Introduction au calcul fractionnaire », ch. 4, vol. 1 de Abry et al. [AGV02], 2002, ISBN 2-7462-0409-6.
- [Men99] MENSLER – « Analyse et étude comparative de méthodes d'identification des systèmes à représentation continue, développement d'une boîte à outils logicielle », Thèse de doctorat, université Henri Poincaré, Nancy I, Centre de recherche en automatique de Nancy, janvier 1999.
- [Mic92] F. MICHAUT – *Méthodes adaptatives pour le signal*, Traité des nouvelles technologies, série traitement du signal, Hermès, Paris, 1992.
- [Mik83] J. MIKUSIŃSKI – *Operational calculus*, second éd., International series of monographs in pure and applied mathematics, Pergamon Press, New-York, 1983, Translation of Rachunek Operatorów.
- [MK00] M. S. MAMIS et M. KÖKSAL – « Remark on the lumped parameter modeling of transmission lines », *Electric machines and power systems* **28** (2000), no. 6, p. 565–575.
- [Mor00] R. MORVAN – « Modélisation de circuits et systèmes de dimension infinie », thèse de doctorat, Université de Bretagne occidentale, Laboratoire d'électronique et systèmes de télécommunications, janvier 2000.
- [Mun90] S. M. MUNK – *Studies of the "PXLN LMS-algorithm"*, GEC-Alsthom, July 1990.
- [Niv56] I. NIVEN – *Irrational Numbers*, The Carus Mathematical Monographs, vol. 11, The Mathematical Association of America, 1956.
- [NR75] J.-P. NOUALY et G. L. ROY – « Réponse transitoire des câbles, première partie : étude des modes de propagation », Bulletin de la direction des études et Recherches – série B, Réseaux électriques, matériels électriques 1, EDF, 1975.

Bibliographie

- [Ous83] A. OUSTALOUP – *Systèmes asservis linéaires d’ordre fractionnaire, théorie et pratique*, 1983.
- [Ous95] — , *La dérivation non entière, théorie, synthèse et applications*, Traité des nouvelles technologies, série automatique, Hermès, Paris, 1995.
- [Ous98] — , « Les Mathématiques de l’amortissement », *Pour la Science* (1998), no. 253.
- [Pet95] R. PETIT – *L’outil mathématique*, 4 éd., Enseignement de la physique, Masson, Paris, Milan, Barcelone, 1995.
- [Pré98] C. PRÉVÉ – *Protection des réseaux électriques*, Hermès, Paris, 1998, ISBN 2-86601-688-2.
- [Pél75] R. PÉLISSIER – *Les Réseaux d’Énergie Électrique – Propagation des Ondes Électriques sur les Lignes d’Énergie*, vol. 5, Dunod, Bordas Paris Bruxelles Montréal, 1975.
- [Riu01] D. RIU – « Modélisation des courants induits dans les machines électriques par des systèmes d’ordre un demi », Thèse de doctorat, Institut national polytechnique de Grenoble, Laboratoire d’électrotechnique de Grenoble, Décembre 2001.
- [RM00] P. ROSSAT-MIGNOD – *Contribution à la Modélisation d’un réseau d’éclairage avec lampe à décharge type fluorescent, sous logiciel de type circuit SIMPLORER*, Mémoire, Conservatoire National des Arts et Métiers, Lyon, 2000.
- [Rou02] P. ROUCHON – « A Catalogue of "Flat" Partial Differential Systems », *International School in Automatic Control of Lille, Control of Distributed Parameter Systems : Theory and Applications* (École centrale de Lille), September 2002.
- [Rés02] Réseau de transport d’électricité – *Mémento de la sûreté du système électrique*, 2002, ISBN 2-912440-06-8, téléchargeable à partir du site <http://www.rte-france.com>.
- [Sab98] J. SABATIER – « La dérivation non entière en modélisation des systèmes à paramètres distribués récursifs et en commande robuste des procédés non stationnaires », thèse de doctorat, université Bordeaux I, laboratoire d’automatique et de productique, novembre 1998.
- [San59] G. SANSONE – *Orthogonal functions*, Pure and applied mathematics, vol. IX, Interscience publishers, inc., New York, 1959, Revised english edition, translated from the italian by A. H. Diamond.
- [Sas94] A. SASTRES – *Protection des réseaux HTA industriels et tertiaires*, Merlin Gérin, décembre 1994, CT 174.
- [Sch34] S. A. SCHELKUNOFF – « The electromagnetic theory of coaxial transmission lines and cylindrical shields », *Bell System Technical Journal (EU)* **13** (1934).
- [Sch66] L. SCHWARTZ – *Théorie des distributions*, Hermann, Paris, 1966, Nouveau tirage, 1998. ISBN 2 7056 5551 4.

- [SDL03] F. SOULIER, N. DOLLINGER et P. LAGONOTTE – « From Cauer synthesis to Legendre polynomials. The case of the lossless transmission line », *IFAC Workshop on Time-Delay Systems (TDS'03)* (Rocquencourt (France)), september 2003.
- [SL00] M. R. SPIEGEL et J. LIU – *Formules et Tables de Mathématique*, 2^eéd., Schaum, Mac Graw-Hill, Londres, etc., 2000.
- [SL02] F. SOULIER et P. LAGONOTTE – « Modeling Distributed Parameter Systems with Discrete Element Networks », *Proceedings of the Fifteenth International Symposium on Mathematical Theory of Networks and Systems (MTNS)* (South Bend, Indiana, USA) (D. S. Gilliam et J. Rosenthal, éd.), University of Notre Dame, August 2002, Article 25666.pdf, www.nd.edu/~mtns.
- [SL04] — , « Synthèse de Cauer bidimensionnelle pour la simulation en temps réel d'une ligne de transmission avec effet de peau », [cem04], soumis en octobre 2003.
- [Spi73] M. R. SPIEGEL – *Variables Complexes – cours et problèmes*, Schaum, Mac Graw-Hill, New-York, etc., 1973.
- [Spi85] — , *Transformées de Laplace : cours et problèmes*, 4e éd., McGraw-Hill, 1985, trad. française Romain Jacoud (Theory and Problems of Laplace Transforms).
- [SS89] T. SÖDERSTRÖM et P. STOICA – *System Identification*, Series in systems and control engineering, Prentice Hall, New York, etc., 1989.
- [Tou02] Y. TOURÉ – « Semigroup Formalism and Internal Model Control for a Heat Exchanger », *International School in Automatic Control of Lille, Control of Distributed Parameter Systems : Theory and Applications* (École centrale de Lille), September 2002.
- [Tri00] J.-F. TRIGEOL – *Nouvelles méthodes optimales de modélisation thermique d'ordre infini. Application à des modèles réduits analytiques*, Mémoire de dea, Université de Poitiers, Laboratoire d'Études Thermiques, 2000.
- [TSL02] J.-F. TRIGEOL, F. SOULIER et P. LAGONOTTE – « Reduction of analytical thermal models and their development in the form of networks », *The european physical journal applied physics* **20** (2002), no. 2, p. 105–119.
- [Val00] A. VALENTI – « Localisation de défauts monophasés à la terre dans les réseaux de distribution HTA », Thèse de doctorat, Université Paris 6, octobre 2000.
- [Wal48] H. S. WALL – *Analytic theory of continued fractions*, The university series in higher mathematics, D. Van Nostrand Co., New York, 1948.
- [Wal94] G. G. WALTER – *Wavelets and Other Orthogonal Systems with Applications*, CRC Press, 1994, ISBN 0-8493-7878-8.
- [WC97] D. J. WILCOX et D. J. CONDON – « A New Transmission-Line Model for Time-Domain Implementation », *The international Journal for Computation and Mathematics in Electrical and Electronic* **16** (1997), no. 4, p. 261–274.

Bibliographie

- [WC98] — , « Time-Domain Model for Multiconductor Power Transmission Lines », *IEEE Proceedings on Generation, Transmission and Distribution* **145** (1998), no. 4, p. 473–479.
- [Wol74] W. WOLOVICH – *Linear multivariable systems*, applied mathematical sciences, Springer Verlag, New York, etc., 1974, ISBN : 0-387-90101-9.
- [WP94] E. WALTER et L. PRONZATO – *Identification de modèles paramétriques : à partir de données expérimentales*, Modélisation, analyse, simulation, commande, Masson, Paris, Milan, Barcelone, 1994.
- [YFW82] C.-S. YEN, Z. FAZARINC et R. L. WHEELER – « Time-Domain Skin-Effect Model for Transient Analysis of Lossy Transmission Lines », *Proceedings of the IEEE*, vol. 70, juillet 1982.

Résumé

La modélisation des lignes électriques a donné lieu à de nombreux travaux. De nos jours, elle s'inscrit dans le contexte plus large de l'étude des systèmes à paramètres répartis. Cependant, au fur et à mesure que se raffinent les solutions analytiques, des modèles approchés doivent être développés pour des applications telles que la protection des lignes à haute tension.

Nous étudions donc des approximations rationnelles de fonctions grâce au théorème de Mittag-Leffler et au développement en fractions continues. Ces approches permettent de généraliser les synthèses de Foster et Cauer aux impédances d'ordre infini. Des exemples d'applications dans différents domaines de la physique sont donnés.

Nous montrons enfin que cette approche peut s'étendre aux fonctions matricielles et fournir une méthode peu contraignante de modélisation des lignes polyphasées. Ce dernier résultat permet de proposer un algorithme d'identification de défauts sur une ligne de transport d'énergie électrique.

Mots-clefs Équation des télégraphistes, diffusion, propagation, systèmes à paramètres répartis, théorème de Mittag-Leffler, fractions continues, approximations de Padé, synthèse de Foster-Cauer, modèles d'ordre réduit, impédance matricielle, protection des lignes polyphasées.

Abstract

Many researches have been made on the modeling of transmission lines. Nowadays, it takes place in the realm of distributed parameter systems. Nevertheless, while the analytical solutions get more sophisticated, approximate models need to be developed for applications such as the protection of high voltage lines.

Thus we study rational approximations of functions thanks to the Mittag-Leffler's theorem and to the continued fraction expansions. These methods allow the generalization of the Foster and Cauer's syntheses to infinite order impedances. Some examples are given as applications in different physical domains.

Eventually, we put in evidence how this approach can be extended to matrix functions, that gives a method to model polyphase lines with few constraints. This result allows to propose a fault identification algorithm on an electrical power transmission line.

Title Distributed parameter systems modeling by impedance expansions in the Laplace domain. Application to the digital protections and to fault location on three phase lines.

Laboratoire Laboratoire d'études thermiques, UMR CNRS 6608, ENSMA, 86 961 Futuroscope Chasseneuil.