



Modélisation réduite de la pile à combustible en vue de la surveillance et du diagnostic par spectroscopie d'impédance

Mohamad Safa

► To cite this version:

Mohamad Safa. Modélisation réduite de la pile à combustible en vue de la surveillance et du diagnostic par spectroscopie d'impédance. Mathématiques générales [math.GM]. Université Paris Sud - Paris XI, 2012. Français. NNT : 2012PA112239 . tel-00855160v2

HAL Id: tel-00855160

<https://theses.hal.science/tel-00855160v2>

Submitted on 2 Dec 2013

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.

UNIVERSITE PARIS-SUD

ÉCOLE DOCTORALE : STITS

Laboratoire: Centre de recherche Inria Paris-Rocquencourt, Equipe-projet SISYPHE

DISCIPLINE Physique- Mathématiques Appliquées

THÈSE DE DOCTORAT

Soutenue le 24/10/2012

par

Mohamad SAFA

MODÉLISATION REDUITE DE LA PILE À COMBUSTIBLE EN VUE DE LA SURVEILLANCE ET DU DIAGNOSTIC PAR SPECTROSCOPIE D'IMPÉDANCE
--

Directeur de thèse :
Co-directeur de thèse :

Pierre-Alexandre BLIMAN
Michel SORINE

Directeur de recherche INRIA
Directeur de recherche INRIA

Composition du jury :

Président du jury :
Rapporteurs :

Françoise LAMNABHI LAGARRIGUE
Daniel HISSEL
Jean-Claude VIVALDA
Alejandro FRANCO
André RAKOTONDRAINIBE

Directrice de recherche CNRS
Professeur des Universités Franche-Comté
Directeur de recherche INRIA
Chef de laboratoire CEA
Docteur ès sciences Ingénieur AREVA

Examineurs :

Remerciements

Tout d'abord, j'exprime ma plus grande gratitude à mon directeur de thèse, Pierre-Alexandre Bliman, pour son encadrement, sa rigueur scientifique et son soutien permanent. Je tiens à le remercier pour sa grande patience et sa disponibilité tout au long de ma thèse. J'exprime toute ma reconnaissance à Michel Sorine de m'avoir accueilli au sein de son équipe SISYPHE. Je tiens à le remercier pour m'avoir encadré, pour son suivi et ses conseils durant mes années de thèse.

Je remercie Monsieur Daniel Hissel et Monsieur Jean-Claude Vivalda d'avoir accepté d'évaluer mon travail et d'en avoir été les rapporteurs. Je remercie aussi Madame Françoise Lamnabhi-Lagarrigue, Monsieur Alejandro Franco et Monsieur André Rakotontrainibe de m'avoir fait l'honneur de participer à mon jury.

Je remercie tous les anciens et actuels membres du projet SYSIPHE à l'INRIA avec qui j'ai eu l'occasion d'interagir, ainsi que mes collègues : Alexandre, Arnaud, Pierre, Karima, Leila et Mohamed. Je remercie particulièrement Martine, assistante du projet SISYPHE pour sa gentillesse et sa serviabilité.

Je remercie Mehdi avec qui j'ai partagé le bureau la première année de ma thèse, Meriem pour tous ses conseils et nos discussions fructueuses.

Je ne pourrais pas oublier Hassan qui m'a toujours soutenu dans les moments difficiles. Pour son aide et ses conseils tout au long de ma thèse.

Je remercie tous mes amis avec qui j'ai passé d'agréables moments durant cette thèse.

Je remercie finalement ma famille. Mes parents Youssef et Mariam, mon oncle Ali, mes frères Karim, Nour et Hassan qui m'ont soutenu, m'ont encouragé sans cesse et dont l'affection m'a rendu la vie vraiment plus agréable.

Résumé

Cette thèse de doctorat porte sur la modélisation réduite des piles à combustible à membrane d'échange de protons (PEMFC), en vue de leur surveillance et de leur diagnostic par spectroscopie d'impédance. La première partie du document présente le principe de fonctionnement de ces piles, ainsi que l'état de l'art de la modélisation et des méthodes de surveillance et diagnostic. Le modèle physique multi échelle particulièrement détaillé publié en 2005 par A.A. Franco sert de point de départ. Il est simplifié de façon à aboutir à un système d'équations aux dérivées partielles en une unique dimension spatiale. La seconde partie est consacrée à l'analyse harmonique de la pile. En s'inspirant de travaux classiques sur l'analyse géométrique de réseaux de réactions électrochimiques, un modèle réduit compatible avec la thermodynamique est obtenu. Cette classe de systèmes dynamiques permet de déterminer, pour un tel réseau, une formule analytique de l'impédance de l'anode et de la cathode d'une pile PEMFC. Un modèle complet de la pile est obtenu en connectant ces éléments à des éléments représentant la membrane, les couches diffuses et les couches de diffusion des gaz. Les modèles précédents supposent la pile représentée par une cellule unique et homogène. Afin de permettre d'en décrire les possibles hétérogénéités spatiales, nous proposons finalement un résultat de modélisation réduite d'un réseau de cellules représentées par leur impédance. Ce modèle approxime l'impédance globale du réseau par une "cellule moyenne", connectée à deux cellules, "série" et "parallèle", représentatives d'écarts par rapport à la moyenne.

Mots-clés. pile à combustible, modélisation réduite, réseaux électrochimiques, spectroscopie d'impédance, balance harmonique.

Abstract. This PhD thesis focuses on reduced modeling of PEM fuel cell for supervision and diagnosis by impedance spectroscopy. The first part of the document presents the principle of the PEM fuel cell, as well as the state of the art of modeling and of the methods for supervision and diagnosis. The multiscale dynamic model published in 2005 by A.A. Franco is particularly detailed and serves as a starting point. It is simplified, in order to obtain a system of partial differential equations in a single spatial dimension. The second part is devoted to harmonic analysis of the PEM Fuel cell. Inspired by classical work on the geometric analysis of electrochemical reactions networks, a model compatible with thermodynamics is obtained. This class of dynamic systems allows establishing, for such a network, an analytical formula of the impedance of the anode and the cathode of the PEM fuel cell. A complete model of the cell is obtained by connecting these elements to the membrane, diffuse layers and gas diffusion layer. The previous models assumed the PEM Fuel cell represented by a single, homogeneous, cell. In order to describe the possible spatial heterogeneities, we finally propose a result of reduced modeling for the impedance of a cell network. This model approximates the overall impedance of the network by a "mean cell", connected to two cells, put in "serial" and "parallel", and representative of the deviations from the average.

Key words. PEM fuel cell, reduced modeling, electrochemical reaction networks, impedance spectroscopy, harmonic balance.

Table des matières

1	Introduction	5
I	État de l'art des modèles de piles à combustible PEM pour la surveillance/diagnostic et étude du modèle de d'A. Franco	8
2	Généralité et état de l'art sur la modélisation et la surveillance et diagnostic des piles à combustible PEM	9
2.1	Description générale de la pile PEM	9
2.1.1	Principe de fonctionnement	9
2.1.2	Structure interne	10
2.1.3	Empoisonnement par le CO	13
2.1.4	Noyage et assèchement de l'eau	13
2.2	Modélisation des phénomènes électrochimiques de la pile PEM	14
2.2.1	Potentiel électrique et formule de Butler-Volmer	14
2.2.2	Modèles de la cinétique électrochimique	16
2.2.3	Phénomène de la double couche électrochimique	18
2.3	Modélisation de l'activité de l'eau dans la pile	21
2.4	Modélisation au niveau du stack de cellules	23
2.5	Modélisation au niveau du système complet	24
2.6	Outils de caractérisation de la pile	25
2.6.1	Courbe de polarisation	25
2.6.2	Diagramme d'impédance	26
2.7	Travaux de surveillance/diagnostic	28
3	Réduction du modèle multi-échelles d'A. Franco	33
3.1	Le modèle d'A. Franco	33
3.1.1	Hypothèses physiques de la modélisation	33
3.1.2	Vision géométrique	33
3.1.3	Une vision modifiée du modèle d'A. Franco	35
3.1.4	Modèle microscopique de diffusion des gaz réactifs dans le Nafion	36
3.1.5	Modèle microscopique du transport des charges	37
3.1.6	Modèle nanoscopique du transport protonique	40
3.1.7	Modèle nanoscopique du mouvement des espèces dissoutes dans les couches diffuses	41
3.1.8	Modèle nanoscopique de la chimie dans la couche compacte	41
3.1.9	Connexion des modèles de demi-piles par la membrane	44
3.2	Exploitation de l'invariance en espace à l'interface micro/nano du modèle de la pile	45
3.2.1	Modèle de diffusion des gaz aux échelles microscopique et nanoscopique	45

3.2.2	Mouvement protonique et évolution du potentiel	45
3.2.3	Récapitulatif du modèle 1D (modèle A)	48
3.3	Réduction du modèle de transport des gaz	50
3.3.1	Principe de la réduction	50
3.3.2	Modèle B avec réduction du modèle de gaz	51
3.4	Réduction du modèle de transport des charges et modèle 0D	52
3.4.1	Résolution du modèle de membrane et simplification du modèle de demi-pile	53
3.4.2	Formulation variationnelle et semi-discrétisation du modèle de demi-pile	54
3.4.3	Modèle 0D ultime : Modèle C	55
3.5	Simulations numériques	56
3.5.1	Réalisation des simulations	56
3.5.2	Surtension et courbes de polarisation	57
3.5.3	Profils de concentration protonique et de potentiel électrique	58
3.5.4	Concentrations des gaz, taux d'occupation des sites catalytiques	60
3.5.5	Sensibilité aux vieillissement de la couche catalytique	62
3.6	Valeurs numériques utilisées dans les simulations	63
3.7	Conclusion	65

II Modélisation des réseaux électrochimiques. Applications à la modélisation harmonique des piles à combustible PEM 66

4	Réseaux électriques et chimiques. Modèle de pile	67
4.1	Réseaux électriques	67
4.1.1	Réseaux de dipôles; multipôle; N -ports	67
4.1.2	Représentations matricielles des 2-ports	69
4.1.3	Couplage de 2-ports en cascade	71
4.1.4	Branchements de dipôle et de quadripôle	72
4.1.5	Dipôles et 2-ports passif	72
4.1.6	Dipôle générateur et image par un 2-ports passif	73
4.2	Réseaux électrochimiques et analogie électrique	75
4.2.1	Point de vue géométrique sur les mélanges, les réactions et les réacteurs ouverts.	76
4.2.2	Energie interne, entropie, énergie de Gibbs et irréversibilité d'un sys- tème chimique ouvert	83
4.2.3	Représentations d'état d'un système chimique ouvert et dissipativité .	86
4.2.4	Cas d'un réseau de réactions électrochimiques	90
4.2.5	Représentation réduite du mécanisme réactionnel de Tafel-Heyrovsky- Volmer et analogue électrique	95
4.2.6	Mécanisme réactionnel de Damjanovic et Brusic	98
4.3	Modèle de la double couche et modèle de pile	103
4.3.1	Modèle de la double couche	103
4.3.2	Modèle de la pile	104
4.4	Conclusion	106
5	Comportement harmonique des réseaux électrochimiques	107
5.1	Méthode de la balance harmonique pour le calcul analytique d'une impédance	107
5.1.1	Principe de la méthode	107
5.1.2	Gain complexe équivalent (GCE)	107

5.1.3	Application du méthode du gain complexe équivalent sur les mesures d'impédance	109
5.2	Application à un réseau de réactions électrochimiques	109
5.3	Mécanisme réactionnel de Tafel-Heyrovsky-Volmer à l'anode de la pile	113
5.3.1	Courbes de polarisation	113
5.4	Mécanisme réactionnel de Damjanovic et Brusic à la cathode de la pile	114
5.4.1	Courbes de polarisation	115
5.5	Modèle d'impédance de double couche et de pile avec prise en compte la correction de Frumkin	116
5.5.1	Modèle de double couche	116
5.5.2	Modèle de la pile	117
5.6	Décomposition quadripolaire du GDL	118
5.6.1	Modèle quadripolaire	119
5.6.2	Formulation variationnelle	120
5.6.3	Comportement harmonique	121
5.6.4	Représentation dipolaire et impédance	130
5.7	Conclusion	134
6	Modélisation harmonique réduite d'un réseau de cellules	135
6.1	Modélisation réduite de l'impédance d'un réseau série-parallèle des cellules . .	135
6.1.1	Impédance du réseau	135
6.1.2	Modèle d'impédance réduit du réseau	136
6.1.3	Exemple d'un modèle simple d'impédance	139
6.2	Identification de l'impédance réduite du réseau	142
6.2.1	Méthodologie	142
6.2.2	Structure du problème d'identification	142
6.3	Conclusion	148
7	Conclusions générales	149
III	Annexes	151
A	Couches de Stern et de diffusion et formules de Butler-Volmer	152
A.1	Introduction	152
A.2	Les couches de Stern et de diffusion	152
A.2.1	Fonction et géométrie de l'interface entre la phase métallique et l'électrolyte	152
A.2.2	Le comportement électrique de la double couche	154
A.2.3	Correction de Frumkin pour la formule de Butler-Volmer	160
A.3	Construction élémentaire des formules de Butler-Volmer	163
A.4	Formule de Butler-Volmer et théorie de Marcus	165
A.4.1	Les bases du modèle de Marcus.	165
A.4.2	Calcul des énergies d'activation et du facteur de symétrie avec la théorie de Marcus	166
A.4.3	Formule de Butler-Volmer et coefficient de transfert de charge α . . .	168
A.4.4	Cas d'un système de réactions	170
B	Calcul de la surtension dans la couche compacte du modèle de Franco	172
C	Discrétisation des EDPs du modèle de Franco	175

D Outils généraux de surveillance et de diagnostic à base de modèles	187
D.1 Redondance physique et analytique	187
D.2 Approche à base d'observateur	189
D.3 Approche fréquentielle	193
D.4 Méthode locale pour la détection et la diagnostic de changement	194

Chapitre 1

Introduction

Avec les préoccupations croissantes au sujet de l'énergie durable et des questions environnementales, l'hydrogène et la technologie de pile à combustible attirent de plus en plus l'attention de la recherche académique, de l'industrie et des gouvernements. L'hydrogène, qui est le composant le plus présent dans l'écorce terrestre, est aujourd'hui considéré comme un vecteur énergétique non carboné et propre, hautement réactif et à forte densité énergétique. Sa combustion électrochimique dans une pile à combustible permet de générer un courant électrique de façon continue avec des rendements élevés tout en ne rejetant que de l'eau et de la chaleur. En effet, dans une pile à combustible alimentée en hydrogène pure, le rendement énergétique total (électrique+thermique) est de l'ordre de 80%, tandis que le rendement électrique dépasse actuellement les 40% .

Cependant, l'hydrogène n'existe pas à l'état naturel et doit être produit, en utilisant des sources indigènes (énergies renouvelables, nucléaires, biomasse, charbon ou gaz naturel), par électrolyse de l'eau ou par reformage des hydrocarbures.

Parmi toutes les familles existantes de piles à combustible, la pile à combustible à membrane électrolyte polymère (ou PEM selon l'acronyme des expressions anglaises Polymer Electrolyte Membrane) est la plus mûre et la plus prometteuse pour la production décentralisée d'énergie électrique. Elle permet de convertir l'énergie de l'hydrogène et d'oxygène directement en énergie électrique en générant une différence de potentiel de l'ordre d'un Volt (différence de potentiel qui est propre au couple redox que forment l'hydrogène et l'oxygène). Du fait de ce caractère naturel basse tension, les constructeurs assemblent plusieurs piles (ou cellules) en série pour former un empilement de plusieurs cellules (stack) afin d'obtenir une tension suffisamment élevée. Selon l'application, la tension de sortie d'un stack de pile à combustible PEM peut être comprise entre 6 V et 200 V, voire plus.

Les avantages de la pile à combustible PEM sont très nombreux, parmi eux sa faible température de fonctionnement ($< 100^{\circ}\text{C}$), sa densité de puissance élevée, son démarrage rapide et le plus important : son faible émission de gaz carbonique, qui induit de l'émission de gaz à effet de serre. Tous ces avantages font de la pile PEM le candidat idéal pour les applications stationnaires, automobiles et sous-marines, ainsi que des systèmes portatifs tels que les téléphones cellulaires ou les ordinateurs portables. Néanmoins, s'agissant d'une nouvelle technologie émergente, le développement de piles à combustible PEM est toujours confronté à de nombreux défis qui doivent être abordés avant la commercialisation généralisée, par exemple, le développement des outils de surveillance et de diagnostic, la réduction des coûts, la durabilité et la fiabilité.

Étant un système complexe, la pile à combustible PEM est susceptible comme tout processus industriel de présenter des dysfonctionnements et des défauts qui peuvent diminuer son rendement et sa durée de vie. Les principaux modes de dysfonctionnement sont l'empoisonnement des sites catalytiques par le monoxyde de carbone dû à l'impureté de l'hydrogène

utilisé pour alimenter la pile, l'engorgement en eau des couches de diffusion des gaz et l'assèchement de la membrane.

En effet, des traces du monoxyde de carbone peuvent se retrouver mélangées avec l'hydrogène lorsque ce dernier est fabriqué par reformage d'hydrocarbures. Le monoxyde de carbone rentre en concurrence avec l'hydrogène pour réagir sur les sites catalytiques, ce qui a pour effet de diminuer la surface active de la pile. Par ailleurs, il est nécessaire que la membrane de la pile soit bien hydratée pour une meilleure performance de la pile, mais l'excès de l'eau engorge la pile et empêche les gaz réactifs de diffuser jusqu'aux sites réactifs. Il est donc indispensable de mettre en place une procédure de surveillance et de diagnostic. Cette procédure nécessite d'une part, le développement des modèles qui aident à la compréhension du comportement dynamique des phénomènes physiques et électrochimiques au sein de la pile, et d'autre part, la mise en place de stratégies de diagnostic adapté, permettant la détection opportune des différents défauts.

Un des moyens non invasifs pour la surveillance et le diagnostic est d'utiliser comme données les courbes de polarisations, mais aussi la spectroscopie d'impédance. Les méthodes qui s'offrent sont celles basées sur des modèles boîte noire d'analyse résidus représentant l'écart entre les données expérimentales et les modèles. Ce travail cherche au contraire à utiliser des modèles physiques au lieu de modèles de boîte noire, ce qui à notre connaissance est nouveau en ce qui concerne la modélisation du comportement d'impédance.

Dans un premier temps, nous avons travaillé sur la modélisation instationnaire de différents phénomènes physique et électrochimique au sein de la pile. En s'inspirant de travaux classiques sur l'analyse géométrique des réseaux de réactions électrochimiques, un modèle réduit compatible avec la thermodynamique est obtenu. En vue de son utilisation dans les techniques de surveillance et de diagnostic, nous avons ensuite analysé son comportement harmonique en utilisant la méthode de la balance harmonique. Cette technique permet de calculer analytiquement l'impédance d'un modèle entrée/sortie formé d'équations différentielles ordinaires.

Les modèles précédents supposent la pile représentée par une cellule unique et homogène. Afin de permettre d'en décrire les hétérogénéités spatiales, nous proposons finalement un résultat de modélisation réduite d'un réseau de cellules représentées par leur impédance. Ce modèle approxime l'impédance globale du réseau par une "cellule moyenne", connectée à deux cellules "série" et "parallèle" représentatives d'écart par rapport à la moyenne. Nous étudions ensuite le problème lié à l'identification des quantités capables de décrire le comportement moyen et la dispersion, pour les utiliser comme des données d'alerte pour la surveillance et le diagnostic.

Différents modèles seront présentés dans cette étude. Dans le Chapitre 3, nous réduisons le modèle d'A. Franco en un modèle 0D. Dans le Chapitre 4, nous présentons un modèle générique d'évolution de réseaux des réactions électrochimiques sous forme d'un système dynamique. Ce modèle fournit par rapport au modèle de Franco un modèle détaillé de la couche compacte. Dans le Chapitre 5, nous présentons, moyennant la méthode de la balance harmonique, un modèle d'impédance issu du modèle obtenu dans le Chapitre 4. Ce modèle est ensuite enrichi en tenant compte de l'impédance de la couche de diffusion de gaz (GDL). Enfin, la modélisation réduite d'un réseau de cellule que nous présentons dans le Chapitre 6 permet de décrire l'hétérogénéité dans une cellule ou dans un stack. Cette méthode est testée sur un modèle simple d'impédance (du premier ordre).

Le contenu et la structure des travaux présentés dans ce rapport de thèse sont les suivants.

La première partie permet de décrire l'état de l'art de la modélisation de la pile à combustible de type PEM et les outils de surveillance et de diagnostic. Nous détaillons ensuite le modèle multi-échelle de Franco [37] et le réduisons en un modèle composé d'un

système d'équations algébro-différentielles.

Le Chapitre 2 rappelle des généralités sur le fonctionnement des piles à combustible de type PEM. Ce chapitre expose un état de l'art de la modélisation de différents phénomènes électrochimiques au sein de la pile et des principaux travaux de surveillance et diagnostic existants. L'objectif du Chapitre 3 est de présenter d'une manière détaillée le modèle multi-échelle de cellule développé par Franco [37] qui tient compte du phénomène de double couche électrique. Ce modèle est réduit en exploitant certaines de ses propriétés mathématiques.

L'apport principal de cette première partie est la réduction du modèle de Franco [37]. Ce modèle comporte un modèle microscopique 2D couplé avec un modèle nanoscopique 1D décrivant avec précision les phénomènes électrochimiques dans la couche active. Nous avons montré que le modèle microscopique était en fait invariant dans une de ses deux dimensions, qui peut donc être supprimée. La forme simplifiée du modèle alors obtenue reste multi-échelle, avec un modèle microscopique 1D couplé avec le modèle nanoscopique 1D. La semi-discrétisation des équations aux dérivées partielles en espace par la méthode des éléments finis conduit à un modèle composé d'un système d'équations algébro-différentiel 0D. Ce dernier, simple à implémenter dans un logiciel de calcul scientifique comme Matlab, permet de faire des simulations numériques relativement rapides avec une vingtaine de points de discrétisation.

La deuxième partie de la thèse est consacrée à la modélisation des réseaux électrochimiques sous forme de systèmes dynamiques, puis à la modélisation de leur comportement harmonique. Le Chapitre 4 présente un point de vue géométrique sur les mélanges, les réactions et les réacteurs ouverts. Ce chapitre fournit un modèle d'évolution de réseaux électrochimiques compatible avec la thermodynamique et un modèle de la pile incluant le phénomène de la double couche électrique. En se basant sur ces modèles, le Chapitre 5 décrit comment calculer analytiquement l'impédance d'un réseau de réactions électrochimiques et de la pile. On utilise pour cela la technique de la balance harmonique, qui permet d'approximer le comportement harmonique d'un système dynamique entrée-sortie non-linéaire.

Enfin, le Chapitre 6 est consacré à la modélisation en régime harmonique d'un réseau série-parallèle de cellules, à partir de la modélisation harmonique d'une cellule. Lorsque les cellules sont différentes, ces travaux permettent d'obtenir un modèle simple approchant le comportement du réseau de façon plus précise que ne le ferait une unique cellule moyenne. Ceci permet de rendre compte de l'inhomogénéité du réseau, et éventuellement d'identifier la nature de celle-ci.

La modélisation des réseaux électrochimiques dans le Chapitre 4 contient des points originaux par rapport à d'autres travaux dans le domaine [64, 85]. L'obtention d'une forme analytique pour l'impédance des réseaux électrochimiques, basée sur le modèle précédent (Chapitre 5), est une contribution originale de cette thèse, de même que la décomposition quadripolaire des éléments de la pile et le résultat de modélisation réduite d'un réseau obtenu au Chapitre 6.

Première partie

État de l'art des modèles de piles à combustible PEM pour la surveillance/diagnostic et étude du modèle de d'A. Franco

Chapitre 2

Généralité et état de l'art sur la modélisation et la surveillance et diagnostic des piles à combustible PEM

Dans ce chapitre, nous rappelons brièvement le principe de fonctionnement de la pile à combustible de type PEM et sa structure interne, par une description générale de différents phénomènes physiques et électrochimiques au sein de la pile. Nous présentons ensuite, un état de l'art sur la modélisation des différents phénomènes électrochimique de la pile. Enfin, nous rappelons quelques travaux liés à la surveillance et au diagnostic de la pile ainsi que les outils de caractérisation les plus couramment utilisés : courbes de polarisation et diagrammes d'impédance.

2.1 Description générale de la pile PEM

2.1.1 Principe de fonctionnement

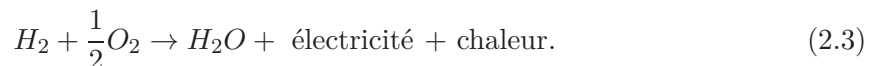
Une pile à combustible à membrane polymère PEM est un dispositif électrochimique qui permet de convertir l'énergie chimique des gaz hydrogène et oxygène en énergie électrique. La cellule de la pile est constituée de deux électrodes séparées par une membrane polymère. L'hydrogène est oxydé à l'anode pour donner des protons et des électrons selon la réaction d'oxydation suivante :



Les électrons sont acheminés par le circuit électrique externe de la pile alors que les protons traversent la membrane, imperméable aux gaz et aux électrons, où ils se combinent à l'oxygène et aux électrons pour produire de l'eau, selon la réaction de réduction suivante :



la réaction globale quant à elle s'écrit :



Cette réaction crée une différence de potentiel entre les électrodes de l'ordre d'un Volt, différence de potentiel qui est propre au couple rédox que forment H_2 et O_2 . Du fait de ce

caractère naturel très basse tension, les constructeurs assemblent plusieurs cellules en série afin d'obtenir une tension suffisamment élevée avec un rendement satisfaisant. Une pile à combustible est alors faite d'un empilement de générateurs électrochimiques élémentaires appelés cellules entre deux plaques terminales servant à l'alimentation et qui constituent les bornes électriques de la pile, comme le montre la Figure 2.1. Les piles sont dites propres

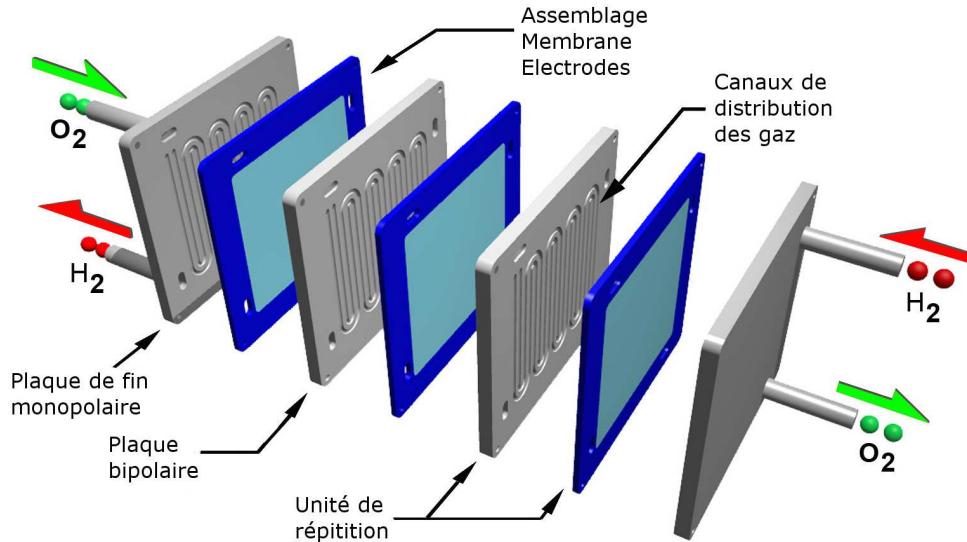


FIGURE 2.1 – Pile à combustible de type PEM alimentée en hydrogène et en oxygène

puisque elles fonctionnent sur le principe inverse de l'électrolyse de l'eau et ne rejettent alors que de l'eau et les gaz réactifs non consommés.

2.1.2 Structure interne

Chaque cellule de la pile est constituée de deux plaques bipolaires et d'un Assemblage Membrane Électrodes (MEA). Les deux plaques bipolaires assurent la distribution des gaz réactifs jusqu'aux les électrodes et collectent le courant électrique produit. L'assemblage électrode membrane est le coeur de la cellule, il est constitué d'une membrane polymère conductrice de protons et imperméable aux gaz et aux électrons ; de deux couches actives AL où est dispersé le catalyseur ; et de deux couches de diffusion de gaz GDL, comme représenté sur la Figure 2.2. Les gaz réactifs circulant dans les plaques bipolaires se diffusent dans les GDL avant d'atteindre les couches actives où ont lieu les réactions chimiques. D'un côté, l'oxydation de l'hydrogène à l'anode produit des protons et des électrons qui génèrent le courant électrique, de l'autre côté, la réduction de l'oxygène à la cathode produit de l'eau qui hydrate la membrane. La production d'eau liquide se fait donc uniquement du côté cathodique. En pratique néanmoins, de l'eau est susceptible d'être présente à l'anode à cause de la diffusion à travers la membrane.

Les plaques bipolaires. Elles permettent l'alimentation en gaz réactifs : O_2 à la cathode d'une cellule et H_2 à l'anode de la cellule juxtaposée, ce qui justifie la désignation de plaque "bipolaire". Elles assurent la liaison électrique entre deux cellules juxtaposées (la liaison électrique avec le circuit externe est assurée par les plaques terminales), la tenue mécanique et l'évacuation de la chaleur produite. Les gaz non consommés et l'eau produite par la réaction sont évacués par les canaux. Les plaques sont typiquement conçues en graphite non poreux ou en acier inoxydable, avec une bonne conductivité électronique.

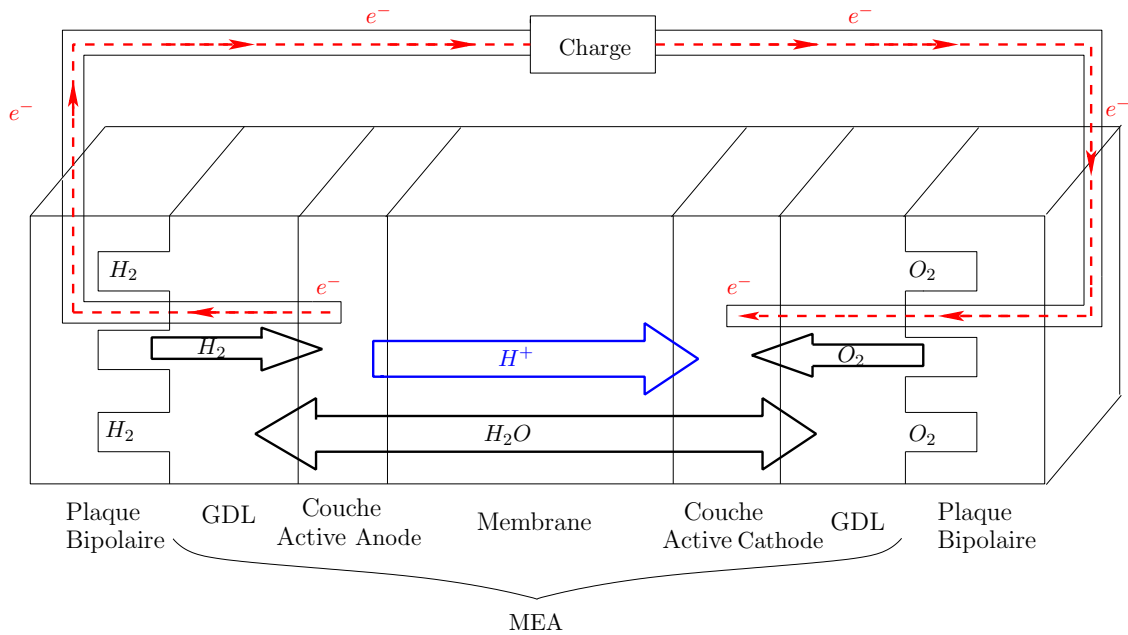


FIGURE 2.2 – Cellule d'une pile

Les couches de diffusion des gaz GDL. Elles sont constituées d'un matériau poreux, hydrophobe et conducteur électronique et thermique. Elles assurent la diffusion des réactifs jusqu'à la couche active, comme son nom l'indique, et le transport des électrons entre la couche active et les plaques bipolaires. Elles sont généralement constituées de papier ou de tissu de carbone. Grâce à leur caractère hydrophobe, elles assurent une bonne évacuation de l'eau, car si l'eau s'accumule dans les GDL, une large proportion du catalyseur n'est plus accessible au gaz réactif.

Les couches actives AL. Les couches actives constituent le lieu des réactions chimiques, elles sont disposées directement sur la membrane et elles contiennent le catalyseur (nanoparticules de platine supportées par des particules de Carbone assurant la conduction électronique), du polymère Nafion qui assure la conduction protonique et une phase poreuse qui permet l'accès des gaz aux sites catalytiques et l'évacuation de l'eau. La couche active constitue l'endroit où le polymère Nafion, le catalyseur et le gaz se rencontrent dans des volumes très faibles (le point triple de la réaction), comme le montre la Figure 2.3.

La membrane. La membrane constitue l'électrolyte de la pile et assure la conduction protonique de l'anode vers la cathode sous l'influence d'un champ électrique. Elle doit être imperméable aux gaz pour séparer les réactifs entre les deux électrodes et un bon isolant électronique. Les matériaux le plus couramment employés pour fabriquer les membranes sont les polymères perfluorosulfonés de type Nafion® (société Du pont de Nemours). Ce dernier a une bonne stabilité chimique et possède une bonne conductivité ionique lorsqu'il est bien hydraté. Le polymère est formé d'un squelette Fluoro-carboné et de groupements acides sulfoniques comme représenté sur la Figure 2.4, il présente un fort caractère hydrophile dû aux groupements sulfoniques terminaux des chaînes pendantes, ce qui permet le déplacement des protons hydratés. La conductivité protonique dépend donc fortement de la quantité d'eau adsorbée.

En effet, quand la membrane est sèche, les charges ne sont pas dissociées, chaque proton est donc lié à son groupe SO_3^- . Lorsque la membrane est hydratée, les molécules d'eau ne

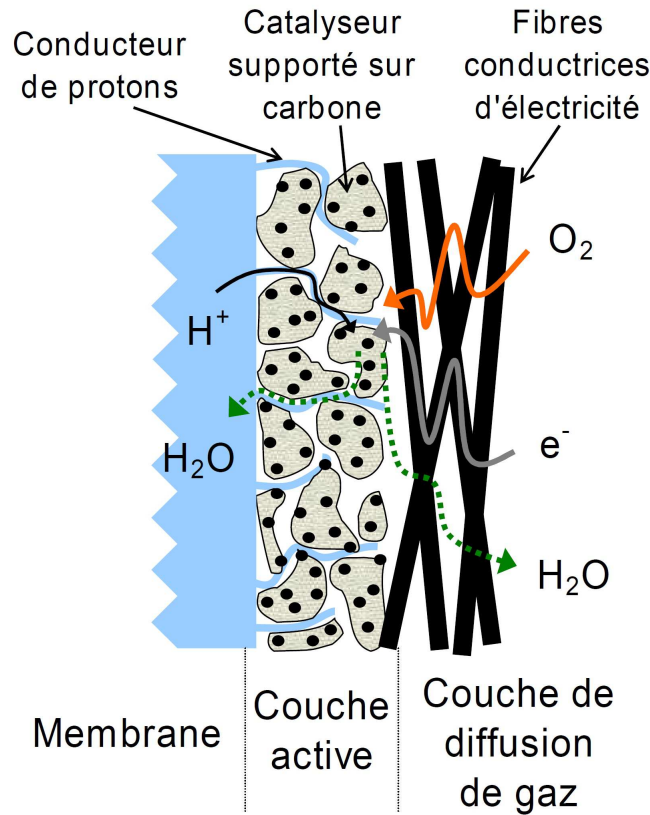


FIGURE 2.3 – Diffusion et réaction de l'oxygène à la cathode. Figure extraite de [58].

peuvent pas pénétrer dans les chaînes de polymères, elles diffusent lentement entre les fibres jusqu'à solvater les groupes SO_3^- , c'est la phase dite d'hydratation. Au cours de cette phase, les protons commencent à se dissocier et ils sautent d'une molécule d'eau à sa voisine, c'est le processus local de diffusion. Les nouvelles molécules d'eau arrivant remplissent les espaces entre les fibres, ceci crée un réseau dans lequel les protons peuvent désormais diffuser à plus longue distance. Tant que l'hydratation est partielle les protons ne peuvent diffuser que dans les régions qui contiennent de l'eau, et lorsque la membrane devient complètement hydratée, un proton peut la traverser de part en part et la pile peut alors fonctionner.

Les épaisseurs typiques des couches présentées ci-dessus sont données dans le tableau suivant :

Composant	Épaisseur caractéristique
Plaqué bipolaire	$1000 \times 10^{-6} \text{ m}$
Couche de diffusion des gaz GDL	$100 - 400 \times 10^{-6} \text{ m}$
Couche active AL	$5 - 10 \times 10^{-6} \text{ m}$
Membrane	$20 - 150 \times 10^{-6} \text{ m}$

TABLE 2.1 – Épaisseurs caractéristiques des différents composants d'une cellule de la pile PEM, [56].

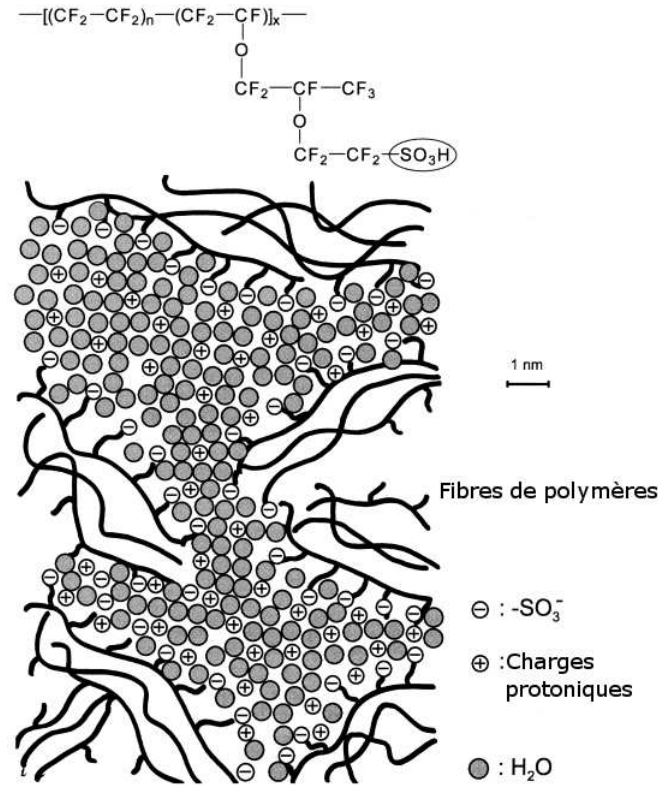


FIGURE 2.4 – Structure chimique de la membrane Nafion : chaque fibre est composée d'une chaîne de polymères organisées de manière que les groupements sulfonates soient disposés à l'extérieur du cœur de la fibre. Figure extraite de [57].

2.1.3 Empoisonnement par le CO

Les matériaux précieux, notamment le Platine Pt ou le Ruthénium Ru, utilisés comme catalyseur dans la couche active de la cellule, sont très sensibles à la présence d'impuretés dans les gaz réactifs. Les impuretés comme le CO se présentent lors de l'utilisation d'un reformer pour produire l'hydrogène. Les expériences montrent [5] que la présence de quelques centaines de ppm¹ du CO peut endommager la pile, en bloquant les sites catalytiques utilisés par l'hydrogène pour la réaction chimique.

2.1.4 Noyage et assèchement de l'eau

Dans le fonctionnement normal de la pile, l'eau produite dans la couche active de la cathode est soit utilisée pour bien humidifier la membrane polymère, soit elle est transportée vers l'extérieur de la pile par le flux des gaz sous forme de vapeur d'eau. La conductivité protonique de la membrane dépend fortement de son taux d'hydratation. En effet, elle est multipliée par 10 entre l'état sec et l'état gonflé.

Les problèmes d'engorgement en eau peuvent se localiser à deux niveaux, dans les électrodes (à la fois au niveau de la couche active et le GDL) et dans les canaux d'alimentation. Au niveau de l'électrode, l'eau produite dans la zone active doit être rapidement évacuée sinon elle peut s'accumuler et gêner la diffusion d'oxygène vers les sites catalytiques. Au niveau des canaux de distribution, si l'eau liquide s'accumule dans un canal, elle risque de

1. partie par million

l'obstruer et de le rendre inopérant. La surface d'électrode active diminue alors, car il y a des zones mortes qui ne produisent pas d'électricité.

2.2 Modélisation des phénomènes électrochimiques de la pile PEM

2.2.1 Potentiel électrique et formule de Butler-Volmer

Potentiel électrique. Une cellule de la pile à combustible réalise la transformation d'une énergie chimique en énergie électrique. Le travail électrique W_{ele} fourni par la pile correspond au déplacement des charges électriques entre les deux niveaux de potentiel où se situent les électrodes [56] :

$$W_{\text{ele}} = -nFE,$$

dans la formule précédente, E est la force électromotrice à l'équilibre (circuit ouvert), F est la constante de Faraday et n le nombre des électrons produits par mole de combustible oxydé (ici $n=2$). Dans le cas d'une réaction réversible, comme celle de la pile, W_{ele} est égal à la variation d'enthalpie libre ΔG , appelée aussi "variation d'enthalpie libre de Gibbs" au cours de la réaction chimique :

$$W_{\text{ele}} = \Delta G,$$

ceci signifie que le potentiel théorique de la pile en circuit ouvert vaut :

$$E = -\frac{\Delta G}{2F}. \quad (2.4)$$

D'après la deuxième loi de la thermodynamique on a :

$$\Delta G = \Delta H - T\Delta S, \quad (2.5)$$

où ΔH est la variation d'enthalpie de la réaction, elle caractérise l'énergie chimique susceptible d'être transformée en énergie électrique. Le reste de cette variation énergétique est libérée sous forme de la chaleur, représentée ici par le terme $T\Delta S$, où T est la température et ΔS la variation d'entropie du système. Toutes ces grandeurs dépendent de la température et de la pression. Par exemple, pour une pression de 1 bar et une température de 80°C, la variation d'enthalpie libre de la réaction (2.3) vaut $-226,1 \text{ kJ.mol}^{-1}$, si l'eau est obtenue sous forme vapeur et $-228,2 \text{ kJ.mol}^{-1}$ si l'eau est sous forme liquide. Dans ce cas le potentiel théorique E varie entre 1.17 et 1.18 V. En pratique, la variation d'enthalpie libre ΔG dépend des conditions de la réaction, et plus précisément des activités des espèces. Dans le cas de la réaction d'oxydoréduction entre H_2 et O_2 , nous avons d'après la loi de Nernst :

$$\Delta G = \Delta G^0 - RT \ln \left(\frac{a_{H_2} a_{O_2}^{\frac{1}{2}}}{a_{H_2O}} \right) \quad (2.6)$$

où ΔG^0 est la variation d'enthalpie libre standard de la réaction à la température T , *i.e.* définie par rapport à la pression de référence de 1 bar. Par définition, l'état physique d'un corps pur sous la pression 1 bar, à la température T , est appelé état standard du corps pur à cette température.

Remarquons d'après (2.4) et (2.6) que si l'activité des réactifs augmente, ΔG diminue et plus d'énergie électrique est libérée. D'autre part, lorsque l'activité du produit augmente, ΔG augmente et moins d'énergie est libérée. Pour voir cet effet sur la tension de la cellule,

nous pouvons remplacer (2.6) par sa valeur dans (2.4) et obtenir :

$$E = -\frac{\Delta G^0}{2F} + \frac{RT}{2F} \ln \left(\frac{a_{H_2} a_{O_2}^{\frac{1}{2}}}{a_{H_2O}} \right) \quad (2.7)$$

sachant que d'après (2.5) on a :

$$\Delta G^0 = \Delta H^0 - T \Delta S^0$$

où ΔH^0 et ΔS^0 sont respectivement les variations d'enthalpie et d'entropie standard. Ils sont calculés expérimentalement et diffèrent légèrement d'une référence à l'autre.

Formule de Butler-Volmer. Dans cette partie, nous décrivons, d'un point de vue macroscopique, le phénomène d'activation dû à la cinétique des réactions électrochimiques dans la pile. Dans le Chapitre 4 nous étudions dans un cadre plus général ce phénomène. En général, les réactions électrochimiques qui se produisent dans la pile sont :



à l'anode et



à la cathode. On note par $v_{o,i}$ et $v_{r,i}$ les constantes de vitesse de réaction d'oxydation et de réduction, à l'anode pour $i = a$ et à la cathode pour $i = c$. Ces derniers s'écrivent d'après la théorie de la cinétique [74], sous la forme :

$$v_{o,i} = k_i \exp \left(-\frac{\Delta G_{ox,i}}{RT} \right), \quad v_{r,i} = k_i \exp \left(-\frac{\Delta G_{red,i}}{RT} \right), \quad i = a, c \quad (2.9)$$

où $\Delta G_{ox,i}$ resp. $\Delta G_{red,i}$ est la variation de l'énergie de Gibbs d'activation d'oxydation resp. de réduction. k_i est un paramètre cinétique de la réaction. Nous verrons dans le Chapitre 4 que d'après [74], les variations de l'énergie de Gibbs s'écrivent pour $i = a, c$:

$$\Delta G_{ox,i} = \Delta G_{ox,i}^0 - \alpha n F \eta_i, \quad (2.10)$$

$$\Delta G_{red,i} = \Delta G_{red,i}^0 + (1 - \alpha) n F \eta_i, \quad (2.11)$$

où η_a resp. η_c est la surtension anodique, resp. cathodique. $\Delta G_{ox,i}^0$ est l'énergie de Gibbs à l'équilibre, n est le nombre des électrons échangés (ici $n = 2$), α_i est le coefficient de transfert, qui indique dans quel sens la réaction est favorisée.

On a alors d'après (2.9) :

$$v_{o,i} = k_{o,i} \exp \left(\frac{2\alpha_i F \eta_i}{RT} \right), \quad v_{r,i} = k_{r,i} \exp \left(-\frac{2(1 - \alpha_i) F \eta_i}{RT} \right), \quad (2.12)$$

où

$$k_{o,i} = k_i \exp \left(-\frac{\Delta G_{ox,i}^0}{RT} \right), \quad k_{r,i} = k_i \exp \left(-\frac{\Delta G_{red,i}^0}{RT} \right), \quad (2.13)$$

En utilisant la loi d'action de masse, les courants anodique et cathodique I_a et I_c s'écrivent sous la forme :

$$\begin{aligned} I_a &= F \left(a_{H_2} v_{o,a} - a_{H^+}^2 v_{r,a} \right) \\ &= F \left(k_{o,a} a_{H_2} \exp \left(\frac{2\alpha_a F}{RT} \eta_a \right) - k_{r,a} a_{H^+}^2 \exp \left(-\frac{2(1-\alpha_a)F}{RT} \eta_a \right) \right) \end{aligned} \quad (2.14)$$

$$\begin{aligned} I_c &= F \left(a_{O_2}^{\frac{1}{2}} a_{H^+}^2 v_{r,c} - a_{H_2O} v_{o,c} \right) \\ &= F \left(k_{r,c} a_{O_2}^{\frac{1}{2}} a_{H^+}^2 \exp \left(\frac{2(1-\alpha_c)F}{RT} \eta_c \right) - k_{o,c} a_{H_2O} \exp \left(-\frac{2\alpha_c F}{RT} \eta_c \right) \right) \end{aligned} \quad (2.15)$$

où a_i , $i \in \{H_2, O_2, H_2O, H^+\}$ est l'activité chimique de l'espèce i . En notant par I_a^0 et I_c^0 les courants d'échanges à l'équilibre, définis par :

$$I_a^0 = F k_{o,a} a_{H_2} = F k_{r,a} a_{H^+}^2 \quad (2.16)$$

$$I_c^0 = F k_{r,c} a_{O_2}^{\frac{1}{2}} a_{H^+}^2 = F k_{o,c} a_{H_2O} \quad (2.17)$$

les équations (2.14) et (2.15) s'écrivent sous la forme :

$$I_a = I_a^0 \left(\exp \left(\frac{2\alpha_a F}{RT} \eta_a \right) - \exp \left(-\frac{2(1-\alpha_a)F}{RT} \eta_a \right) \right) \quad (2.18)$$

$$I_c = I_c^0 \left(\exp \left(\frac{2(1-\alpha_c)F}{RT} \eta_c \right) - \exp \left(-\frac{2\alpha_c F}{RT} \eta_c \right) \right) \quad (2.19)$$

Ces équations sont usuellement connues sous le nom de Butler-Volmer.

2.2.2 Modèles de la cinétique électrochimique

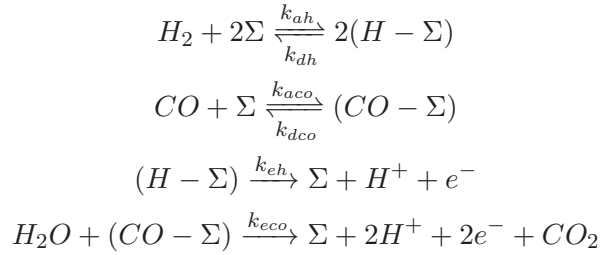
Dans une pile à combustible, les phénomènes électrochimique ont lieu dans la couche active. On peut distinguer deux approches traitant l'électrochimie dans la couche active de la pile : les approches macroscopique et microscopique. Les modèles macroscopiques sont généralement définis par les conditions globales de pression, de température et de débit, [82], [55] et [91], où la consommation des gaz réactifs et la production de l'eau s'expriment en fonction de la densité du courant, en négligeant la cinétique chimique sur les sites catalytiques. Les modèles considérés comme microscopiques, [81], [5], [6], [79] et [39] utilisent les principes complexes d'électrochimie, de la thermodynamique et de la mécanique de fluides. Ils tiennent compte des variations spatiales des paramètres et des performances locales et ils traitent plus finement l'adsorption, la désorption et l'électro-réduction des espèces réactives sur les sites catalytiques. Dans les deux approches, la densité du courant est exprimée en fonction de la surtension à l'électrode par la formule de Butler-Volmer.

Par ailleurs, les matériaux précieux utilisés comme catalyseur dans la couche active de la cellule, notamment le Platine Pt ou le Ruthenium Ru, sont très sensibles à la présence du CO. Les expériences montrent que la présence de quelques centaines de ppm² du CO peut endommager la pile, en bloquant les sites catalytiques utilisés par l'hydrogène pour la réaction d'électroréduction.

L'étude de l'empoisonnement par le CO nécessite donc une modélisation fine de la dynamique de l'occupation des sites catalytiques par les différents gaz au sein de la couche active. Parmi les modèles existants qui traitent le mécanisme cinétique de l'empoisonnement par le CO, on peut citer les travaux de Springer et al. [81], Baschuk et Li [5, 6], Bhatia et Wang [15], Shah et al. [79] et Nwoga et Van Zee [61]. Le mécanisme de l'empoisonnement

par le CO proposé par Springer et al. [81] a été intégré dans un modèle statique de la couche active de l'anode. Baschuk et Li [7] utilisent le même mécanisme dans un modèle statique 1D et ils ont étendu leur modèle dans [6] pour inclure l'effet de l'"air bleed". Bhatia et Wang [15] ont étudié le problème dynamique, en simplifiant le mécanisme cinétique proposé par Springer et al. [81] et en négligeant les phénomènes de transport. Chu et al. [30] ont présenté un modèle similaire à celui de [15] mais qui inclut la diffusion des gaz réactifs. Shah et al. [79] ont présenté un modèle dynamique 1D, non-isotherme, qui tient compte du transport de la matière, de l'empoisonnement par le CO et de l'utilisation de l'"air bleed".

Travaux de Springer et al. Selon Springer et al. [81] le schéma réactionnel à l'anode en présence du H_2 , du CO et de l'eau s'écrit avec nos notations :



Le symbole Σ désigne un site catalytique libre, et on indique par $(X - \Sigma)$ la forme adsorbée de l'espèce X , avec :

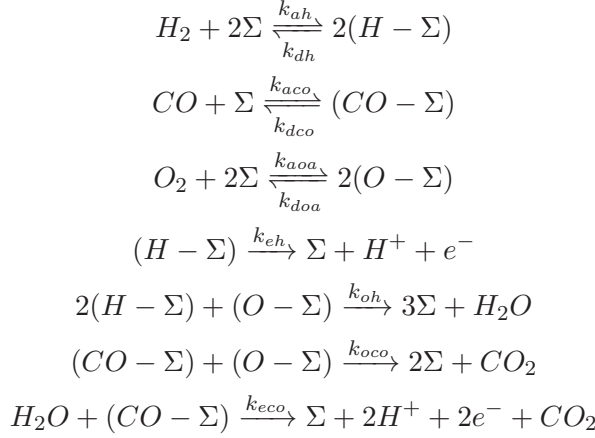
- k_{aco}, k_{ah} : constantes d'adsorption du CO, du H_2 .
- k_{eco}, k_{eh} : constantes d'électro-oxydation du CO, du H_2 .
- k_{dco}, k_{dh} : constantes de désorption du CO, du H_2 .

L'évolution des fractions des sites catalytiques occupés par l'hydrogène θ_h et le monoxyde de Carbone θ_{co} s'écrivent :

$$\begin{aligned} \rho_a \frac{d\theta_h}{dt} &= k_{ah}P_h(1 - \theta_h - \theta_{co})^n - k_{dh}\theta_h^n - 2k_{eh}\theta_h \sinh\left(\frac{\eta_a}{b_h}\right) \\ \rho_a \frac{d\theta_{co}}{dt} &= k_{aco}P_{co}(1 - \theta_h - \theta_{co}) - k_{dco}\theta_{co} - k_{eco}\theta_{co} \exp\left(\frac{\eta_a}{b_{co}}\right) \end{aligned}$$

où P_h et P_{co} sont les pressions partielles d'hydrogène et de CO, η_a la surtension anodique. n est l'ordre de la réaction de l'hydrogène sur les sites catalytiques, b_h, b_{co} sont respectivement les pentes de Tafel pour l' H_2 et le CO. ρ_a est la capacité de stockage par unité de volume de catalyseur à l'anode.

Travaux de Baschuk et Li. Dans [6], les auteurs étendent le mécanisme réactionnel proposé par Springer et al. [81] en ajoutant les réactions liées à l'oxygène de l'"air bleed" :



avec :

- k_{aoa}, k_{doa} : constantes d'adsorption et de désorption du O_2 à l'anode.
- k_{oh}, k_{oco} : constantes d'oxydation du H_2 , du CO .

L'évolution des fractions des sites catalytiques occupés par l'hydrogène θ_h et le monoxyde de Carbone θ_{co} et l'oxygène θ_{oa} s'écrivent :

$$\begin{aligned}
 \rho_a \frac{d\theta_h}{dt} &= k_{ah}C_h(1 - \theta_h - \theta_{co} - \theta_{oa})^2 - k_{dh}\theta_h^2 - 2k_{eh}\theta_h \sinh\left(\frac{\eta_a}{2RT/F}\right) - 2k_{oh}\theta_{oa}\theta_h^2 \\
 \rho_a \frac{d\theta_{co}}{dt} &= k_{aco}C_{co}(1 - \theta_h - \theta_{co} - \theta_{oa}) \exp\left(\frac{\beta r \theta_{co}}{RT}\right) - k_{dco}\theta_{co} \exp\left(\frac{(1 - \beta)r\theta_{co}}{RT}\right) \\
 &\quad - 2k_{eco}\theta_{co} \exp\left(\frac{\eta_a}{2RT/F}\right) - k_{oco}\theta_{oa}\theta_{co} \\
 \rho_a \frac{d\theta_{oa}}{dt} &= k_{aoa}C_{oa}(1 - \theta_h - \theta_{co} - \theta_{oa})^2 - k_{doa}\theta_{oa}^2 - k_{oco}\theta_{oa}\theta_{co} - k_{oh}\theta_{oa}\theta_h^2
 \end{aligned}$$

où C_h, C_{co} et C_{oa} désignent les concentrations de l' H_2 , du CO et de l' O_2 à l'anode. β et r sont respectivement les coefficients d'échange et d'interaction anodique.

2.2.3 Phénomène de la double couche électrochimique

À l'interface métal/électrolyte de la couche active, le transfert de charge est couplé à un processus interfacial : l'excès des électrons dans la phase métallique est compensé par une accumulation de protons dans la phase électrolyte. Il existe donc une double couche de charges, de signes opposés et qui se comporte en première approximation comme un condensateur. On peut donc observer en régime dynamique le passage d'un courant capacitif dû à la charge ou décharge de ce condensateur. Ce courant capacitif dépend de la valeur de la capacité de double couche.

Dans les modèles analogiques (constitués par des circuits électriques), la double couche électrique est modélisée par un condensateur en parallèle avec la résistance liée à la réaction chimique (voir Barsoukov et al. [3], Diard et al. [33]). Cette modélisation suppose donc que la capacité de double couche est découplée des réactions sur les sites catalytiques, bien qu'elle soit fortement liée aux réactions chimiques de la couche active [37].

Afin de modéliser, en régime transitoire, le passage d'un courant capacitif dû à la charge ou la décharge de ce condensateur, on établit un bilan de charge suivant [33] :

$$I(t) = I_f(t) + C_{dc} \frac{d\eta}{dt} \quad (2.20)$$

où I est le courant total, I_f est le courant faradique exprimé par la formule de Butler-Volmer, détaillée dans la section 2.2.1. C_{dc} est la capacité de double couche, η est la différence de potentiel entre les deux phases métallique et électrolyte. La relation (2.20) signifie que le courant total, lorsqu'une réaction électrochimique se déroule à l'interface métal/électrolyte, est égal à la somme du courant faradique I_f et du courant capacitif $C_{dc} \frac{d\eta}{dt}$.

Dans l'approche de type circuit électrique analogique, l'équation (2.20) est modélisée par un condensateur de capacité C_{dc} en parallèle à la résistance associée à la réaction électrochimique, c-à-d il est supposé que la capacité de double couche est découplé des réactions chimiques. Cependant, le phénomène de double couche est fortement couplé aux réactions chimiques, comme nous le verrons en détail dans le Chapitre 3.

Le concept et le modèle de la double couche ont été présentés dans les travaux de Helmholtz sur les interfaces des suspensions colloïdales et ont ensuite été étendus à des surfaces d'électrodes en métal par Gouy, Chapman et Stern, et puis dans les travaux de Grahame et de Bockris [18].

Modèle de Helmholtz : Le modèle le plus ancien de la double couche a été proposé en 1853 par Helmholtz qui a suggéré que l'excès des charges électriques sur le métal attire à l'interface une quantité équivalente d'ions de signes opposés, Figure 2.5. Les deux couches sont séparées par une distance de l'ordre de grandeur du rayon atomique ou de la moitié de la taille d'une molécule, donc de l'ordre de l'Angström. Le plan dans lequel se trouvent les centres des ions est appelé plan de Helmholtz (PH) et le reste de l'électrolyte est considéré électriquement neutre. Il en résulte une double couche d'épaisseur fixe, qui ressemble à un condensateur électrique constitué de deux plaques parallèles, dont la capacité C_{dc} est fonction des constantes diélectriques et ne dépend pas du potentiel de l'électrode.

Modèle de Gouy-Chapman : Afin de justifier la variation expérimentale de la valeur de la capacité différentielle C_{dc} en fonction de la différence de potentiel, Gouy et Chapman (1913) ont considéré les ions dans l'électrolyte comme des charges ponctuelles, distribuée statistiquement selon la loi de Boltzmann. Ils constatent alors la formation d'une couche diffuse de charges, où la concentration en ions est maximale à la proximité de l'électrode et diminue progressivement jusqu'à une distribution homogène au sein de l'électrolyte, Figure 2.5. Ce modèle permet de décrire quantitativement la façon dont la capacité C_{dc} varie avec la différence de potentiel et la concentration des ions en solution.

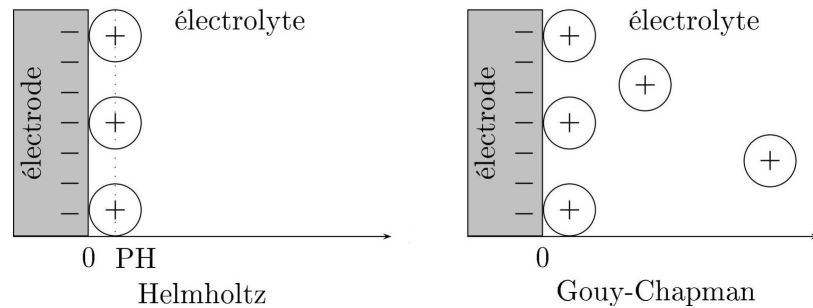


FIGURE 2.5 – Représentation de charges et d'ions selon les modèles de Helmholtz et Gouy-Chapman pour la double couche. PH= Plan de Helmholtz, [33]

Modèle de Gouy-Chapman-Stern : Dans le modèle de Gouy-Chapman les ions de l'électrolyte peuvent s'approcher jusqu'à une distance nulle de l'électrode et la capacité de double

couche tend vers l'infini lorsque le potentiel de l'électrode est très élevé. Stern (1924) a combiné les approches de Helmholtz et celle de Gouy Chapman en considérant que les ions de l'électrolyte sont solvatés, c'est-à-dire entourés par une couche de molécules d'eau, et ne pouvait pas s'approcher de l'électrode au point de la toucher. Le lieu des centres des ions solvatés définit un plan de moindre approche des espèces, appelé le plan externe de Helmholtz (PEH). La double couche est composée d'un côté, d'une partie rigide appelée couche de Stern ou de Helmholtz comprise entre l'électrode et le plan externe de Helmholtz (PEH), et d'un autre côté d'une couche de diffuse des ions qui s'étendent du PEH vers la masse de l'électrolyte, voir la Figure 2.6. Le condensateur de double couche se comporte comme deux condensateurs branchés en série. L'un a pour valeur la capacité du condensateur de Helmholtz et l'autre celle de Gouy-Chapman.

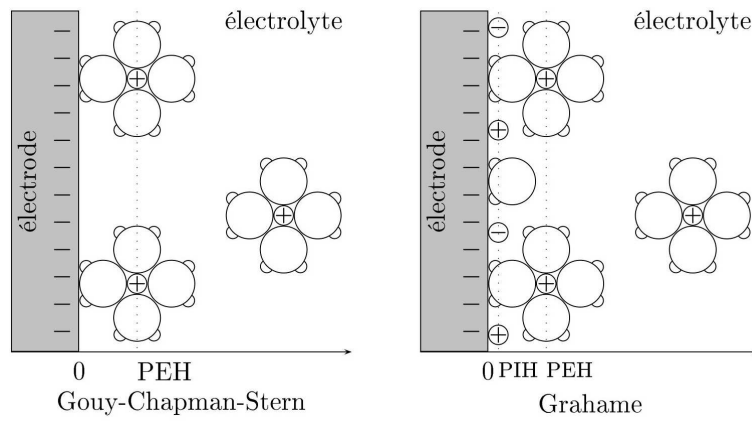


FIGURE 2.6 – Représentation de charges et d'ions selon les modèles de Stern et Grahame. PEH= Plan Externe de Helmholtz, PIH =Plan Interne de Helmholtz. [33]

Modèle de Grahame : Grahame (1947) a introduit non seulement les ions solvatés qui se trouvent sur le plan de Helmholtz externe et les ions libres dans la couche diffuse, mais aussi l'adsorption spécifique de certains ions à la surface de l'électrode. Cela pourrait se réaliser si l'ion ne possédait pas de couche de solvation, ou si la couche de solvation a été perdue lorsque l'ion s'est rapproché de l'électrode. Les ions en contact direct ont été considérés comme "spécifiquement adsorbés". Le lieu du centre des espèces spécifiquement adsorbées définit le plan interne de Helmholtz PIH. La double couche électrochimique comprend alors trois couches, l'une constituée d'espèces adsorbées spécifiquement au PIH, l'autre d'espèces adsorbées non spécifiquement au PEH, et de la couche diffuse, Figure 2.6.

La double couche électrique a également été étudiée dans le cadre du supercondensateur qui permet de stocker l'électricité par accumulation de charges sur les deux électrodes, servant de collecteurs lorsqu'on impose un potentiel entre celles-ci. Plusieurs modèles mathématiques ont été développés pour analyser la performance de la double couche électrique dans le cadre d'un supercondensateur. Johnson et Newman [52] ont développé un modèle pour décrire la double couche de charge dans une cellule électrochimique afin d'estimer la densité de l'énergie spécifique et la puissance du condensateur électrochimique. Pillay et al. [70] ont modélisé l'influence des réactions secondaires sur la performance des condensateurs électrochimiques. Srinivassan et Weidner [83] ont présenté une solution analytique de la capacité de double couche qui a été utilisée pour étudier l'importance de la résistance ionique et électronique dans la conception des supercondensateurs. Verbrugge et Liu [87] ont présenté un modèle mathématique de la double couche électrique d'un supercondensateur en utilisant les équations de conservation des charges électriques. Ce modèle permet d'étudier

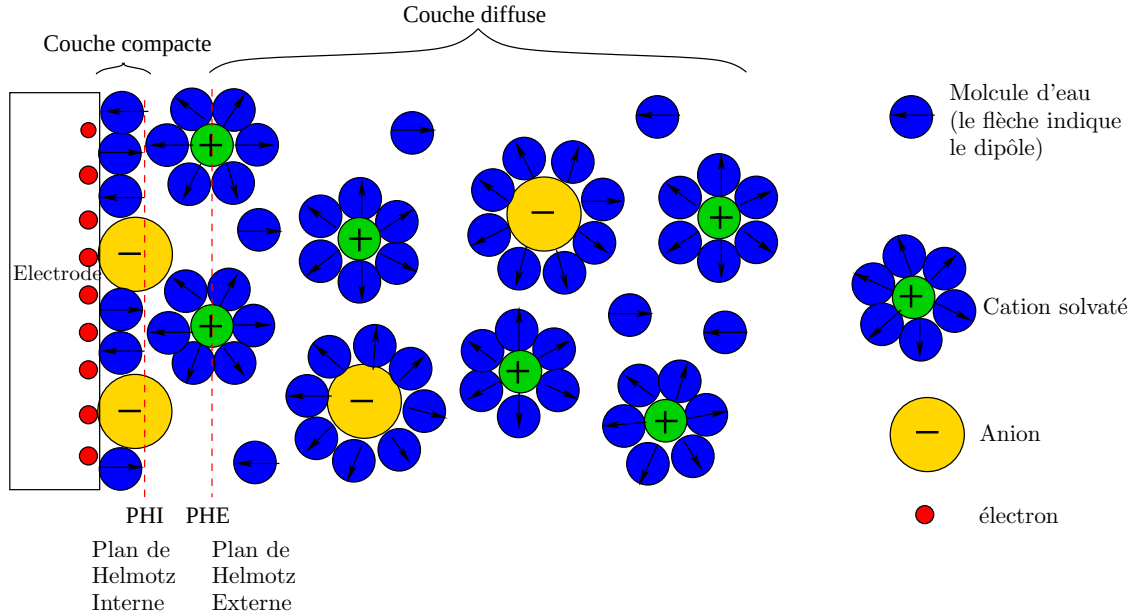


FIGURE 2.7 – Structure de la double couche selon le modèle de Bockris et al [18].

la dynamique de chargement et de déchargement du condensateur.

Cependant, ces modèles ne sont pas adaptés à la modélisation de la double couche électrique dans une pile PEM. En effet, la capacité de double couche est fortement couplée à la réaction chimique dans la couche active de la pile.

Franco et al. [37], [38] proposent un modèle multiéchelle de la pile, qui modélise finement le phénomène de la double couche électrique dans une pile PEM, en détaillant les flux électroniques et protoniques, ainsi que les effets dipolaires liés à l'adsorption de l'eau dans la couche compacte. Ce modèle est détaillé dans le chapitre suivant.

2.3 Modélisation de l'activité de l'eau dans la pile

Une voie très importante dans la recherche sur les piles à combustible PEM vise la modélisation détaillée des phénomènes liés à la gestion de l'eau. Au début des années 90, le développement des modèles traitant du bilan d'eau, est particulièrement centrée sur des modèles isothermes, stationnaires et unidimensionnels pour une cellule [82],[12],[13]. Springer et al. [82] et Bernardi et Verbrugge [12, 13] ont modélisé les phénomènes du transport dans le MEA (assemblage électrodes membrane) et dans les canaux des gaz, l'intensité du courant est supposée constante dans toute la pile, et la chimie sur les sites catalytiques a été négligée. Le modèle de [82] se concentre sur le transport d'eau dans la membrane et il prend en compte l'électro-osmose, la rétro diffusion de l'eau de la cathode vers l'anode et la résistance de la membrane. Dans le modèle de Bernardi et Verbrugge [12, 13], la couche active est considérée comme un milieu poreux d'épaisseur non nulle composé d'une phase poreuse de Pt/C et d'une phase de polymère. Le transport des charges et des gaz sont modélisés plus finement que dans le modèle de Springer et al. [82] qui néglige l'épaisseur de la couche active. Cependant, Bernardi et Verbrugge ne prennent pas en compte la gestion de l'eau dans la membrane supposée parfaitement hydratée. Un modèle 1D, statique et isotherme a été développé dans [4], il intègre l'essentiel des processus physiques et électrochimiques survenant dans la pile. La particularité de ce modèle provient du fait qu'il introduit des degrés variable de noyage dans la couche active de la cathode. Cet effet est modélisé en permettant aux gaz réactifs, à l'eau liquide d'occuper le milieu poreux de la couche active. Wang et al. [90] modélise la

coexistence de zones sèches et de zones humides dans la cellule. Ceci induit un phénomène de diffusion, une distribution de courant et une distribution thermique non homogène. Benziger et al. [26] ont développé un modèle dynamique de la pile de type “Stirred Tank reactor PEM fuel cell” pour étudier l’effet de l’inondation dans le GDL de la cathode et analyser l’existence d’équilibres multiples, dus à la nature auto catalytique du transport de l’eau couplée à la production d’eau dans la pile. C’est un modèle 0D pour la cellule, mais la mise en série de plusieurs cellules permet d’étudier des phénomènes plus compliqués comme la diffusion. Stefanopoulou et al. [53] ont étudié la gestion de l’eau dans la cellule, leur approche emploie deux variables d’état, les quantités d’eau à l’anode et à la cathode, desquelles sont déduites les activités à chacune des électrodes, puis l’activité dans la membrane.

Approche de Springer et al. Dans le modèle de Springer et al. [82], le transport d’eau dans la membrane est gouverné, d’une part, par un flux diffusif de type Fick :

$$N_{w,\text{diff}} = -D_\lambda \frac{\rho_{\text{dry}}}{M_m} \frac{d\lambda}{dz}$$

et d’autre part, par un flux électro-osmotique traduisant la présence d’un cortège de molécules d’eau emportées par chaque proton lors de sa traversée du Nafion, exprimé en fonction de la densité du courant I de la cellule par la relation :

$$N_{w,\text{drag}} = n_{\text{drag}} \frac{I}{F}$$

où λ est le contenu en eau, déterminé expérimentalement en fonction de l’activité de l’eau et de la température. z est la variable d’espace suivant le transport des protons (c’est-à-dire selon l’axe anode/cathode, perpendiculairement aux plaques bipolaires et à la membrane). D_λ est le coefficient de diffusion de l’eau, il dépend de λ et est donné par une corrélation expérimentale. n_{drag} est le coefficient d’électro-osmose qui désigne le nombre de molécules d’eau emportées par chaque proton, donné expérimentalement par la relation :

$$n_{\text{drag}} = \frac{2.5\lambda}{22}$$

ρ_{dry} et M_m sont respectivement la densité et la masse molaire de la membrane sec. La conductivité ionique de la membrane est exprimée en fonction de λ et de la température, la relation est obtenue expérimentalement sur des membranes en hydratation uniforme :

$$\sigma_m(T, \lambda) = \exp \left[1268 \left(\frac{1}{303} - \frac{1}{T} \right) \right] (0.005139\lambda - 0.00326).$$

La résistance R_m de la membrane est calculée en intégrant la résistivité (inverse de la conductivité) sur l’épaisseur de la membrane e_m :

$$R_m = \int_0^{e_m} \frac{dz}{\sigma_m(T, \lambda)}.$$

Approche de Benziger et al. Benziger et al. [26] considèrent les quantités d’eau dans quatre compartiments : anode et cathode, membrane, et couche de diffusion côté cathode (les phénomènes se déroulant dans la couche de diffusion à l’anode étant considérés comme négligeables). Avec nos notations, les variables du modèle sont les pressions partielles de l’hydrogène et de l’eau à l’anode, P_h et P_w^a , l’activité de l’eau dans la membrane a_w , le nombre de moles d’eau dans le GDL cathodique $n_{w,\text{gdl}}$ et les pressions partielles de l’oxygène

et de l'eau à la cathode P_o et P_w^c . Les équations suivantes décrivent l'évolution des variables d'états du modèle :

$$\frac{V_g^a}{RT} \frac{dP_h}{dt} = F_a^{in} \frac{P_h^{in}}{RT} - F_a \frac{P_h}{RT} - A \frac{I}{2F}$$

$$\frac{V_g^c}{RT} \frac{dP_o}{dt} = F_c^{in} \frac{P_o^{in}}{RT} - F_c \frac{P_h}{RT} - A \frac{I}{4F}$$

$$N_{SO_3} \frac{d\lambda}{da_w} \frac{da_w}{dt} = k' A (a_{gdl} - a_w) - k A \left(a_w - \frac{P_w^a}{P_{sat}} \right)$$

$$\frac{dn_{w,gdl}}{dt} = A \frac{I}{2F} - k' A (a_{gdl} - a_w) - k'' A \left(a_{gdl} - \frac{P_w^c}{P_{sat}} \right) + n_{drag} A \frac{I}{F}$$

Le contenu en eau dans la membrane λ est exprimé en fonction de l'activité de l'eau a_w dans la membrane par une relation empirique :

$$\lambda(a_w) = 14.9a_w - 44.7a_w^2 + 70a_w^3 - 29.5a_w^4 - 0.446a_w^5.$$

Dans les équations précédentes, N_{SO_3} est le nombre d'ions sulfonates fixes dans le Nafion, k' est le coefficient de transfert de masse du GDL cathodique vers la membrane, k'' est le coefficient de transfert de masse du GDL cathodique vers le canal d'écoulement de la cathode, k est le coefficient de transfert de masse de la membrane vers l'anode. A est la surface active de l'électrode, V_g^a et V_g^c sont respectivement les volumes des canaux d'alimentation à l'anode et la cathode, F_a^{in} et F_c^{in} sont respectivement les débits volumiques à l'entrée de l'anode et la cathode. a_{gdl} est l'activité de l'eau dans le GDL liée à $n_{w,gdl}$ par la relation :

$$a_{gdl} = \frac{n_{w,gdl} RT}{V_{gdl} P_{sat}} \leq 1$$

où V_{gdl} est le volume de GDL et P_{sat} la pression de saturation de la vapeur d'eau. P_w^a et P_w^c sont donnés par :

$$P_w^a = P_T - P_h, \quad P_w^c = P_T - P_o$$

où P_T est la pression totale dans les canaux d'alimentation des gaz. F_a et F_c sont respectivement les débits volumiques à l'intérieur de l'anode et de la cathode, ils sont exprimés en fonction des autres variables du modèle. D'une manière différente de celle de Springer et al. [82], le coefficient d'électro-osmose n_{drag} est calculé expérimentalement par la relation :

$$n_{drag} = 2 (a_{w,a})^4,$$

où $a_{w,a}$ est l'activité de l'eau à l'anode. La résistance de la membrane est donnée par la formule :

$$R_m(a_w) = 10^5 \exp(-14a_w^{0.2}) \frac{e_m}{A}.$$

2.4 Modélisation au niveau du stack de cellules

La modélisation au niveau du stack de cellule est souvent destinée à l'étude du couplage thermique et à la gestion de l'eau. Parmi les modèles intéressants qui traitent le stack de cellule, on trouve ceux de Shan et al. [80] et Park et al. [65]. Shan et al. [80] ont développé un modèle dynamique 2D d'un stack composé de deux cellules pour étudier la gestion thermique.

Le modèle inclut le transport des gaz, de l'eau et de la température. Ce modèle a été étendu plus tard par Park et al. [65] pour un stack de 20 cellules en considérant la température et l'effet diphasique. Philipps et Ziegler [69] ont présenté un modèle dynamique non isotherme d'un stack basé sur les bilans de masse et d'énergie. Les équations sont résolues pour une cellule et le résultat est mis à l'échelle du stack. Dans l'article [26], Benziger et al. ont développé un modèle pour une cellule unique, ils ont ensuite étendu ce modèle pour une stack formé de plusieurs cellules en séries. Leur modèle tient compte essentiellement de la gestion d'eau et du transport des gaz. M. Marchand [59] a étudiée les problèmes de la gestion de l'eau dans un stack, en prenant en considération le transfert d'eau dans les canaux sous ces deux formes, c'est-à-dire un transfert d'eau exclusivement sous forme liquide et un transfert d'eau exclusivement sous forme gazeuse. Elle décrit les équations hydrauliques et thermodynamiques qui régissent la pile.

2.5 Modélisation au niveau du système complet

La pile à combustible nécessite un grand nombre de dispositifs auxiliaires. Les gaz doivent être contrôlés en termes de pression, d'hydratation et de régulation de débit avant d'entrer dans les électrodes de la pile. Les réactions électrochimiques sont de nature exothermique et comme la pile doit fonctionner dans une plage de températures bien précise, le contrôle de température est essentiel. D'autre part, l'eau produite par les réactions chimiques peut être évacuée avec les gaz sortant, pour être réutilisée pour hydrater les gaz entrants. Le contrôle global du système constitue une application importante des travaux de modélisation de la pile.

Les modèles développés dans ce cadre utilisent généralement des approches macroscopiques pour traiter les bilans de matière. Il s'agit souvent de modèles dynamique 0D de la pile, couplés à des systèmes auxiliaires tels que le groupe moto-compresseur, le système d'humidification et le système d'air. Les modèles proposés par Pukrushpan et al. [71], [72] et Golbert et al. [45] rentrent dans ce cadre et sont principalement dédiés à la commande du système. Golbert et al. [45] ont développé un modèle pour le contrôle et la régulation de la pile. Il tient compte de la dépendance spatiale de la tension, du courant, des flux de matière et de la température. Seule l'évolution de la température est considérée en régime dynamique, les autres bilans étant supposés en régime stationnaire. La simplification des équations conduit à un modèle de contrôle prédictif non linéaire qui permet d'utiliser le contrôle optimal pour améliorer la performance de la pile. Pukrushpan et al. [72] ont présenté un modèle dynamique d'un système pile approprié pour l'étude de la commande. Il intègre la dynamique d'un motocompresseur et un modèle statique d'humidification. Le modèle utilise le bilan massique pour évaluer la masse des réactants dans le coeur de la pile. Pathapati et al. [66] ont développé un modèle dynamique 0D avec une capacité de double couche constante, ajoutée dans la dynamique de la tension de la pile. Le système pile développé, intègre la dynamique du flux, de la pression et du transfert thermique.

D'autres approches modélisent la pile à travers des représentations graphiques. On trouve parmi elles le Bond Graph, le graphe informationnel causal (GIC) et la représentation énergétique macroscopique (REM)

La méthode du Bond Graph est une approche graphique pour représenter des systèmes complexes grâce à un nombre limité de blocs standardisées qui sont connectés par des paramètres qui représentent un couple (flux, effort). Cette approche causale est basée sur la première loi de la thermodynamique. Le Bond Graph a été utilisé pour décrire des systèmes piles à combustible [73], [77] mais ces travaux n'ont pas été utilisés pour développer leur commande.

La méthode du graphe informationnel causal (GIC) est une méthodologie de modélisation

graphique, fonctionnelle et causale mais l'approche n'est pas énergétique [47] et [48]. Elle est basée sur l'idée principale que chaque système peut être décrit par un ensemble de sous-systèmes avec une entrée et une sortie.

Le GIC est bien adapté pour le développement de la structure de commande mais ne prend pas en compte l'aspect énergétique et peut devenir très volumineux pour des systèmes grands et complexes. La méthode du représentation énergétique macroscopique (REM) est un autre approche développée par le L2EP, Université de Lille, France depuis les années 2000. La REM est une combinaison des aspects du développement de la structure de commande de GIC avec des aspects de la modélisation énergétique [23]. Elle est basée sur la première loi de la thermodynamique et elle a été étendue pour modéliser les piles à combustible par Hissel et al. [49], Chrenko et al. [29, 27, 28] et Boulon et al. [22].

2.6 Outils de caractérisation de la pile

2.6.1 Courbe de polarisation

La courbe de polarisation (ou courbe tension-courant) est un outil couramment utilisé dans la caractérisation de la performance de la pile PEM. Elle est obtenue en balayant la plage du courant désirée en régime stabilisé, et en mesurant simultanément la tension délivrée par la pile et le courant qui lui est imposé. Elle a typiquement le profil montré dans la figure 2.8 : à chaque valeur de courant on fait correspondre le potentiel de cellule une fois l'état stationnaire atteint.

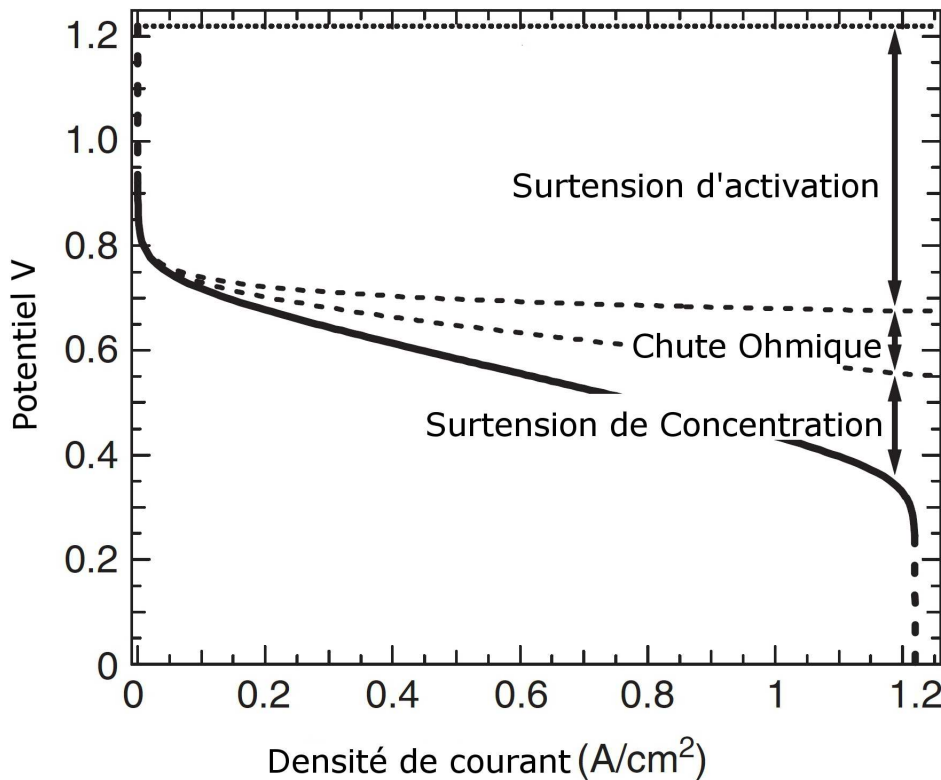


FIGURE 2.8 – Courbe de polarisation d'une pile PEM

La courbe de polarisation traduit les pertes de tension entraînées par différents phéno-

mènes physiques et électrochimiques dissipant l'énergie dans la pile. Il existe trois types de surtension : activation, ohmique et concentration. Selon la densité de courant à laquelle on travaille, on peut distinguer trois zones sur les courbes.

Pour les faibles densités de courant, la perte d'activation ou la surtension d'activation est une conséquence de la nécessité de provoquer le transfert d'électrons et de rompre (respectivement former) les liaisons chimiques à l'anode (respectivement la cathode). La majorité des pertes est attribuée à la cinétique de l'électrode cathodique, puisque la réduction de l'oxygène est plus compliquée et lente que celle de l'hydrogène. Cependant, si du CO est présent à l'anode, la perte de tension anodique est également importante.

Pour des densités de courant moyennes, les pertes d'activation augmentent à un rythme plus lent et la pile est dominée par les pertes ohmiques. La surtension ohmique découle de la résistance à la migration des protons dans la membrane et des électrons dans les GDLs et les couches actives. Ces pertes se traduisent par une zone linéaire légèrement convexe à cause des phénomènes de diffusion des espèces, notamment les protons.

Pour des densités de courant élevées, la tension aux bornes de la cellule chute très rapidement à cause de la surtension de concentration qui est due à des limitations de transport de masse. En fait, les taux des réactions électrochimiques dans les couches actives sont entravés par un manque de réactifs puisque l'apport des gaz n'est plus assez rapide par rapport à la cinétique des réactions. Les limitations de transport de masse sont dues à des limitations diffusionnelles dans les GDL et les couches actives. Notons que le phénomène de noyage dans le GDL, particulièrement côté cathode où l'eau est produite, réduit l'espace disponible pour une diffusion rapide des gaz et augmente la résistance au transport des gaz.

Cependant, les phénomènes physiques au sein de la pile sont fortement couplés : transport des charges, mécanique des fluides, réactions électrochimiques, migration électrique, etc. Les courbes de polarisation ne fournissent pas un aperçu sur le comportement électrochimique en dynamique et ne permettent pas de distinguer les facteurs impliqués dans les pertes. Par exemple, une dégradation de nature ohmique peut être due à une mauvaise hydratation de la membrane, à une corrosion des plaques bipolaires, mais aussi à un décollement de la membrane par rapport à la couche active sous l'effet de cycles thermiques.

2.6.2 Diagramme d'impédance

La spectroscopie d'impédance est une technique de mesure très utilisée dans le domaine de l'électrochimie. Elle consiste à imposer au système étudié une intensité de courant (respectivement une tension) périodique, autour d'un point de fonctionnement pour une fréquence donnée, et de mesurer le déphasage et l'amplitude de la tension (respectivement de l'intensité du courant) pour la même fréquence. L'amplitude de l'excitation (typiquement on prend des perturbations entre 1% et 10%) est faible afin de garantir un comportement linéaire du système étudié, voire la figure 2.9. La mesure du spectre d'impédance à un point de fonctionnement et l'identification d'un modèle d'impédance approprié permettent une description du comportement dynamique de l'électrochimie de la pile.

Dans le cadre de la pile PEM, l'emploi d'une excitation sinusoïdale du courant comme entrée est plus fréquente parce qu'il est plus aisé de contrôler le courant que la tension. De plus, les cellules d'un stack fonctionnent dans des conditions différentes, du fait des hétérogénéités au sein de ces cellules, il adviendrait qu'imposer une tension sur le stack conduise à produire un courant élevé que certaines cellules "dégradées" ne pourraient fournir.

Pour un courant sinusoïdal injecté à l'entrée de la pile sous la forme :

$$I(t) = I_0 + I_1 \sin(\omega t)$$

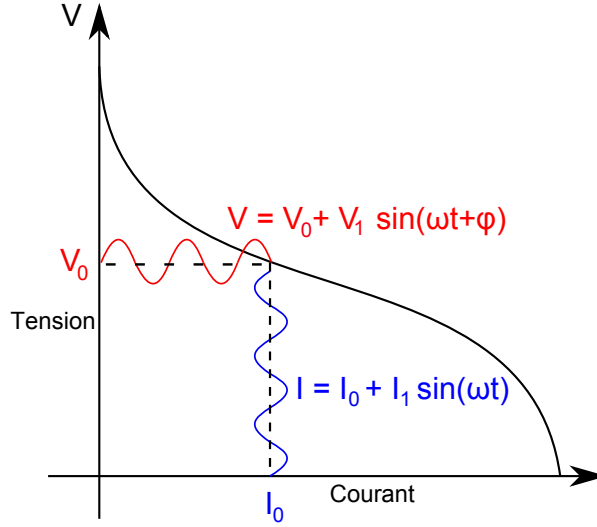


FIGURE 2.9 – Principe de la spectroscopie d'impédance : réponse linéaire en potentiel à une excitation sinusoïdal en courant de faible amplitude autour d'une valeur stationnaire I_0 .

la réponse de la pile à ce type de stimulus s'écrit :

$$V(t) = V_0 + V_1 \sin(\omega t + \varphi_\omega)$$

où ω est la fréquence du courant électrique, I_0 et V_0 sont respectivement les composantes continues du courant et de la tension, I_1 et V_1 sont les amplitudes des composantes alternatives du courant et de la tension tel que $I_1 \ll I_0$ (dans les expérience on utilise une amplitude comprise entre 1% et 10% de la composante continue I_0). φ_ω est le déphasage entre le courant et la tension pour chaque harmonique ω . De façon générale, le courant d'excitation $I(t)$ et sa réponse $V(t)$ sont reliés par $Z(\omega)$, considérée comme une fonction de transfert, telle que :

$$\mathcal{F}(V(t)) = Z(\omega)\mathcal{F}(I(t))$$

où $\mathcal{F}$ est l'opérateur de transformation de Fourier.

La spectroscopie d'impédance électrochimique permet de caractériser l'état de la pile et de distinguer un certain nombre de phénomènes dégradants qui peuvent paralyser son fonctionnement. La contamination des sites catalytiques par le CO , le noyage et l'assèchement sont les défauts les plus courants. Dans cette section, nous présentons un état de l'art sur l'utilisation de la spectroscopie d'impédance dans la littérature comme un outil de surveillance/diagnostic.

Dans les conditions nominales, l'impédance de la pile est formée de deux "bosses" ou boucles qui apparaissent dans les différentes bandes de fréquences, comme le montre la Figure 2.10. Schnider et al. [76] ont étudié la contribution des différents phénomènes survenant dans la pile sur le spectre d'impédance. Dans la bande de fréquence la plus élevée du spectre, l'intersection de l'axe réel avec le spectre est attribuée à la résistance ohmique. La distance entre les deux points d'intersection du spectre avec l'axe réel caractérise la somme de la résistance de transfert des charges anodique et cathodique.

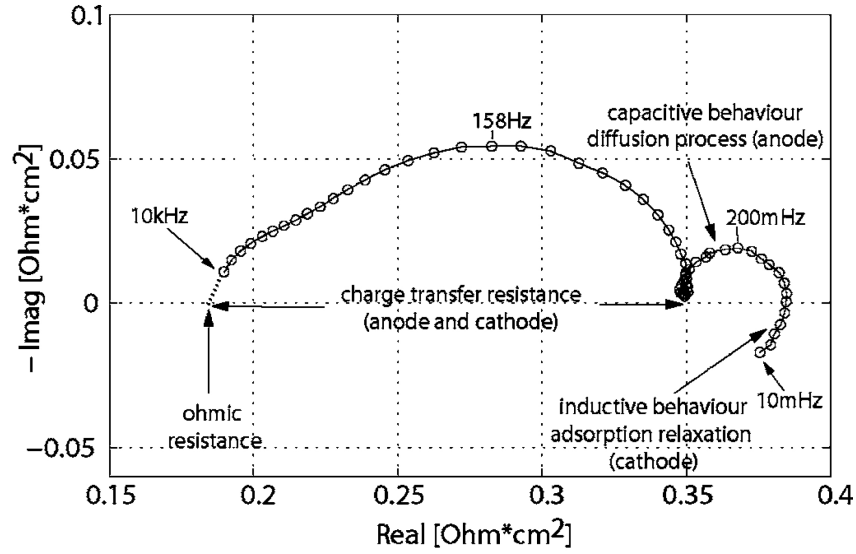


FIGURE 2.10 – [75] Relation entre les différentes parties du spectre et les processus limitatifs, bande de fréquence de 10^{-2} à 10^4 Hz.

2.7 Travaux de surveillance/diagnostic

Dans la littérature, un large spectre de techniques de diagnostic existe, car une communauté scientifique pluridisciplinaire y apporte sa contribution (automaticiens, informaticiens, statisticiens,...). Le choix de la méthode de diagnostic à mettre en œuvre est dicté par la quantité et le type de connaissances dont on dispose, ainsi que par les contraintes du problème de diagnostic posé. Dans l'Annexe D nous présentons un état de l'art des outils généraux de surveillance et de diagnostic à base de modèles.

Plusieurs approches de diagnostic de la pile basé sur des méthodes de statistique existent dans la littérature. On peut citer les méthodes basées sur la génération de résidus grâce à un modèle boîte noire du type réseau de neurones dynamique [50], [51], [36] et [94] ; celles qui concernent le diagnostic par reconnaissance de forme basée sur l'extraction de caractéristiques à partir de différents types de mesures [63] et celles qui mettent en œuvre une analyse temps échelle par paquets d'ondelettes [84].

Les méthodes temps-échelles telles que l'analyse en ondelettes constituent des outils adaptés à l'analyse de signaux non stationnaires (combinant des phénomènes se manifestant à des résolutions temporelles et fréquentielles différentes). L'analyse temps-échelle basée sur une décomposition en ondelettes est adaptée à un système dans lequel différents phénomènes ou défauts se produisent à différentes échelles de temps et de fréquence.

Dans [84], une méthode peu coûteuse de diagnostic de l'inondation dans une pile à combustible PEM est décrite. Cette méthode n'utilise que la tension de pile et elle consiste en une extraction de caractéristiques de signal par décomposition multiéchelle utilisant la transformée en ondelettes discrète, suivie par l'identification et la classification de défaut.

Dans leur travail, Hissel et al. [50] ont développé un modèle orienté pour la sur-

veillance/diagnostic de la pile PEM en utilisant la logique floue et l'analyse des résidus comme vecteurs de diagnostic pour deux défauts : l'assèchement de la membrane et l'accumulation d'eau ou d'azote à l'anode de la pile. Les résidus sont calculés en considérant la différence entre le point de fonctionnement expérimental courant-tension et un point de fonctionnement obtenu par un modèle boîte noire (à base de réseaux de neurones). Dans [63], les auteurs utilisent une méthode de diagnostic de la pile PEM basée sur la reconnaissance de formes. Avec cette méthode, à la fois des informations extraites de courbes de polarisation et de spectres d'impédance peuvent être utilisées. Pour les courbes de polarisation, un modèle empirique est exploité pour assurer l'extraction de caractéristiques, tandis que pour les spectres d'impédance, des connaissances d'experts et la modélisation paramétrique sont utilisées pour l'extraction de caractéristiques.

Dans [93], les auteurs proposent une procédure de diagnostic de l'inondation de la pile basée sur un modèle boîte noire. Les entrées du modèle sont des variables qui sont essentielles pour la gestion de l'eau dans la pile, alors que la sortie est une variable qui peut être contrôlée d'une manière non-intrusive et peut être utilisée pour détecter les inondations (chute de pression dans la cathode par exemple). La procédure de diagnostic est basée sur l'analyse d'un résidu obtenu à partir de la comparaison entre l'expérimentation et une chute de pression estimée. L'estimation de cette dernière est assurée par un réseau de neurones formé avec les données d'une pile saine. La détection de défauts est obtenue moyennant une analyse résiduelle.

Effets de l'empoisonnement par le CO sur les courbes de polarisation. Oetjen et al. [62] ont examiné l'influence de CO mélangés avec l'hydrogène sur la performance de la pile. Ils ont constaté que la courbe de polarisation d'une cellule alimentée par l'hydrogène avec de CO se distingue par une pente négative comparée à celle alimentée par l'hydrogène pur. Cette chute de potentiel s'explique par une augmentation de la surtension anodique due à l'empoisonnement par le CO , comme le montre la Figure 2.11.

Effets de l'assèchement et du noyage sur les courbes de polarisation. Shah et al. [78] ont étudié l'effet de l'activité de l'eau dans la membrane et les GDL sur la performance de la pile. Les courbes de polarisation du modèle et des résultats expérimentaux montrent que l'augmentation de l'hydratation de la membrane conduit à une amélioration de la performance de la pile, ceci dans le rang de la densité du courant moyen et faible. Cependant, lorsque la densité du courant augmente, la production de l'eau dans la couche active de la cathode s'accroît et par conséquent, une chute de potentiel est observée. Cette chute de potentiel est généralement attribuée à la limitation de la diffusion de l'oxygène à cause du noyage de la couche active cathodique.

La technique de la spectroscopie d'impédance, détaillée dans la section précédente, permet de séparer et d'identifier différents processus dynamiques qui se produisent dans la pile et de mieux distinguer la contribution de ces différents phénomènes. La séparation est possible puisque les phénomènes physiques au sein de la pile possèdent des constantes de temps bien distinctes.

Effets de l'empoisonnement par le CO sur le spectre d'impédance. Wagner et al. [89] ont mesuré expérimentalement l'influence du CO sur l'impédance de la pile. Ils ont observé un élargissement du spectre d'impédance et l'apparition d'un pseudo inductance en basse fréquence, qui s'accroît avec la durée de l'expérience. Les mesures faites en mode galvanostatique (lorsque le courant est perturbé à l'entrée de la pile) montrent que durant un empoisonnement progressif de la pile, la partie réelle et la partie imaginaire de l'impédance

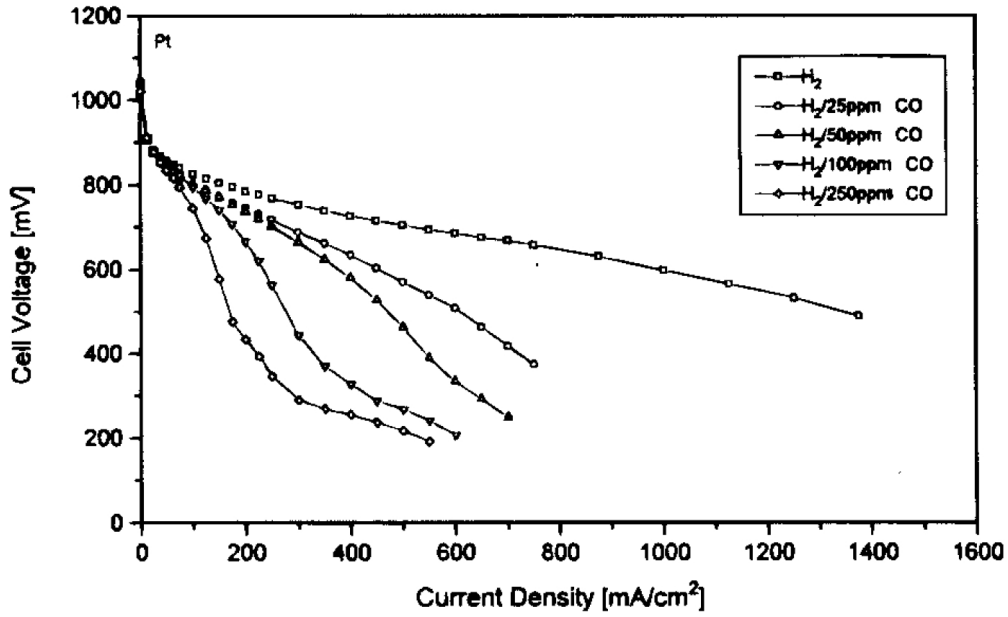


FIGURE 2.11 – Illustration de l'effet du CO sur les courbes de polarisation de la pile PEM (Oetjen et al. [62]).

augmentent et les deux boucles ne peuvent plus être différenciée visuellement, comme le montre la Figure 2.12.

Grâce à cette expérimentation, Wagner et al. [89] proposent un modèle électrique simplifié de la pile qui tient compte de la présence du CO comme présenté sur la Figure 2.13. Selon ce modèle, la dégradation du spectre d'impédance est dominée par les réactions à l'anode. Le schéma électrique de la demi-cellule cathodique (réduction de l'oxygène) est représenté par une résistance de transfert de charge notée R_{ct,O_2} mise en parallèle avec un élément de constante de phase (CPE). Cependant, le schéma de l'anode est plus compliqué à cause des processus électrochimiques dans la couche active qui sont influencés par la présence du CO . Le schéma de l'anode est modélisé en utilisant une électrode poreuse [44]. L'impédance de l'interface électrode/pore est donnée par une capacité double couche $C_{dl,A}$ mise en parallèle avec une impédance faradique Z_F qui contient l'impédance de relaxation de surface [3]. L'impédance de relaxation est liée à la présence du CO et elle explique le développement du comportement pseudo-inductif en basse fréquence du spectre. L'expression de l'impédance faradique Z_F s'écrit :

$$Z_F = \frac{R_{ct} + Z_C}{1 + R_{ct}/Z_k}$$

où R_{ct} est la résistance de transfert des charges à l'anode, Z_C est l'impédance de diffusion finie (impédance de Nernst) et Z_k l'impédance de relaxation donnée par :

$$Z_k = \frac{1 + i\omega\tau_K}{I_F d \left(\ln \tau_k^{-1} \right) / d\varepsilon}$$

où I_F est le courant faradique, τ_k est la constante de temps de relaxation, ε est le potentiel à l'anode. R_k et L_k sont respectivement la résistance et la pseudo-inductance de la relaxation

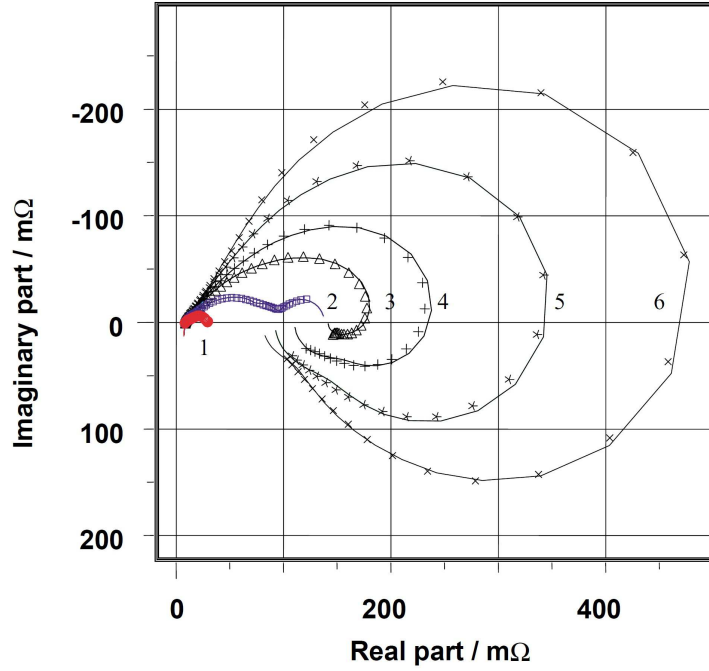


FIGURE 2.12 – [89] Impédance de la pile alimentée par $H_2 + 100$ ppm du CO dans le plan de Nyquist à des moments différents. Les numéros de 1 à 6 représentent l'évolution du spectre avec le temps. Lorsque la durée de l'alimentation de la pile par un mélange du H_2 et du CO augmente, l'impédance anodique s'élargit et devient de plus en plus dominante, elle masque l'impédance cathodique et présente un comportement pseudo inductif en basse fréquence.

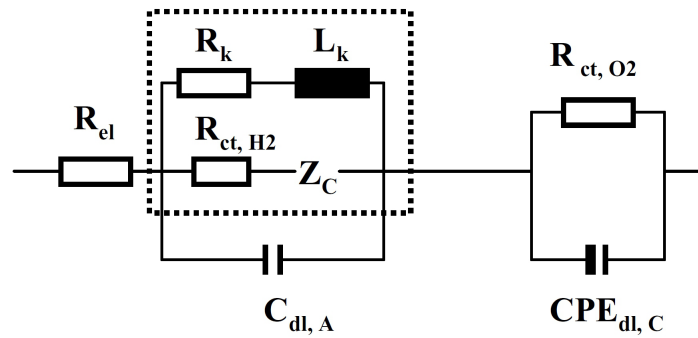


FIGURE 2.13 – [89] Circuit électrique équivalent pour l'évaluation des spectres d'impédance mesurés pour une pile empoisonnée à l'anode par le CO

données par :

$$R_k = \frac{1}{I_F d \left(\ln \tau_k^{-1} \right) / d\varepsilon}, \quad L_k = \tau_k R_k. \quad (2.21)$$

Selon les Wagner et al. [89], les mesures des spectres d'impédances réalisés montre que les éléments R_k , τ_k et R_{ct,H_2} sont les plus sensibles à la présence du CO à l'anode.

Effets de l'assèchement et le noyage sur le spectre d'impédance. Le Canut et al. [24] ont étudié l'effet de l'assèchement et le noyage de la membrane de la pile sur le spectre d'impédance. Dans le cas de l'assèchement de la membrane, Le Canut et al. [24] ont observé une augmentation de la magnitude et de la phase de l'impédance pour toute la plage de fréquence étudiée ($1Hz$ jusqu'au 10^4Hz). L'effet de l'assèchement se traduit par la translation du spectre le long de l'axe réel, ce comportement s'explique par l'augmentation de la résistance de la membrane. Le noyage de la membrane de la pile conduit à une augmentation dans la magnitude de l'impédance en basse fréquence ($< 10Hz$) et à la diminution de la phase pour une fréquence inférieure à $100Hz$.

Chapitre 3

Réduction du modèle multi-échelles d'A. Franco

Dans ce chapitre nous détaillons dans un premier temps le modèle d'A. Franco de la pile [37], qui est de type modèle à pores cylindriques et nous le réduisons ensuite, en exploitant certaines de ses propriétés mathématiques. Nous présentons tout d'abord son choix de modélisation, en détaillant les phénomènes physiques et électrochimiques aux différentes échelles (microscopique et nanoscopique) de son modèle. Nous simplifions ensuite le modèle microscopique de transport des charges et nous démontrons un fait fondamental relatif au modèle de Franco : les réactions et transferts de charges sont homogènes en tout point de l'interface des pores et le modèle est invariant en cette dimension. La forme simplifiée du modèle obtenue reste multiéchelle, avec un modèle microscopique 1D couplé avec un modèle nanoscopique 1D. La semi-discrétisation en espace des équations aux dérivées partielles par la méthode des éléments finis conduit finalement à un modèle composé d'un système d'équations algébro-différentielles 0D.

3.1 Le modèle d'A. Franco

3.1.1 Hypothèses physiques de la modélisation

Les principales hypothèses physiques du modèle d'A. Franco sont les suivantes :

- La pile est alimentée en hydrogène et en oxygène purs, saturés avec de la vapeur d'eau.
- Les mélanges gaz réactifs/vapeur d'eau anodiques et cathodiques se comportent comme des gaz parfaits.
- Les transferts thermiques sont négligés : la température de la pile est supposée constante et uniforme.
- La concentration des protons est constante dans la membrane (condition d'électroneutralité), mais elle est distribuée en espace dans la double couche électrochimique, à l'interface électrolyte/métal.
- Le Nafion est supposé bien hydraté pour assurer un bon transport protonique.

3.1.2 Vision géométrique

A. Franco [37] s'inspire d'un modèle de type modèle à pores cylindriques pour développer son modèle physique multiéchelle de la dynamique électrochimique dans la pile. Ce modèle est un assemblage de trois modèles 1D (Figures 3.1 et 3.2) :

- Un modèle microscopique de transport des gaz réactifs à travers le Nafion imprégné, suivant la coordonnée d'espace z .

- Un modèle microscopique de transport protonique dans l'épaisseur de l'électrode et dans la membrane, et de transport électronique dans l'épaisseur de l'électrode, suivant la coordonnée d'espace r .
- Un modèle nanoscopique de transport décrivant les phénomènes électrochimiques à l'interface, suivant la coordonnée d'espace x .

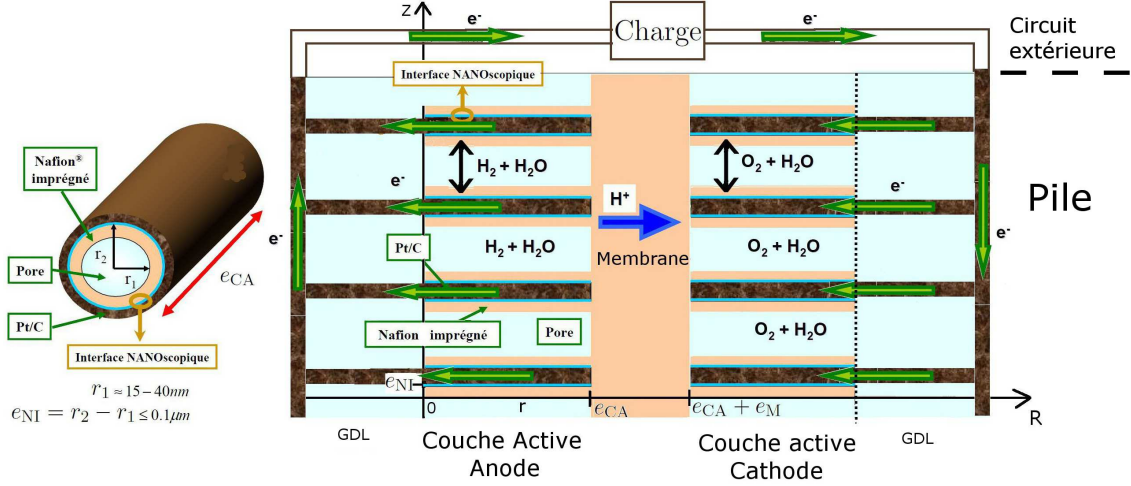


FIGURE 3.1 – Modèle à pores cylindriques, figure inspirée d'A. Franco [39]. Les cylindres sont constitués de trois phases : les pores, le Nafion imprégné et la phase métallique Pt/C. L'axe du cylindre est perpendiculaire à la membrane et aux couches de diffusion des gaz (GDL).

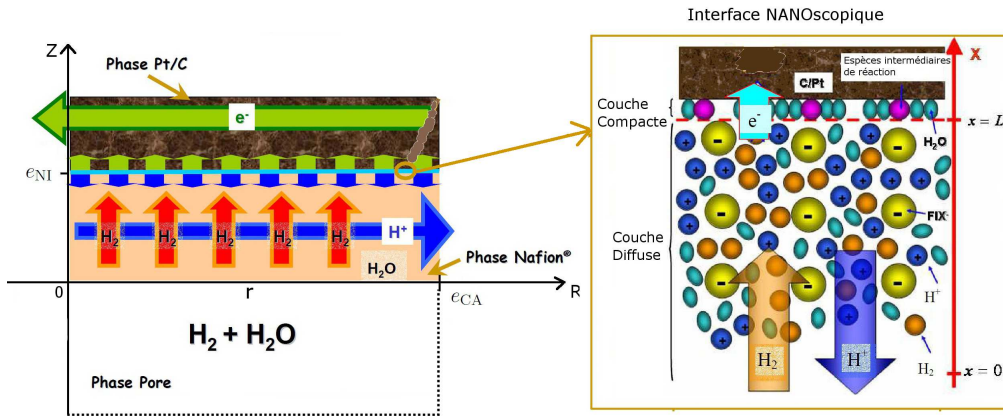


FIGURE 3.2 – À gauche, les échelles microscopiques du modèle d'A. Franco et l'emplacement du modèle nanoscopique en bleu clair, $z = e_{NI}$. À droite, l'échelle nanoscopique décrivant la double couche électrochimique, où x est la variable d'espace utilisée à cette échelle et $L = 1 \text{ nm}$ représente l'épaisseur de la couche diffuse. Figure extraite de [37].

Les domaines géométriques représentés sur les Figures 3.1 et 3.2 peuvent être schématisés dans la figure suivante :

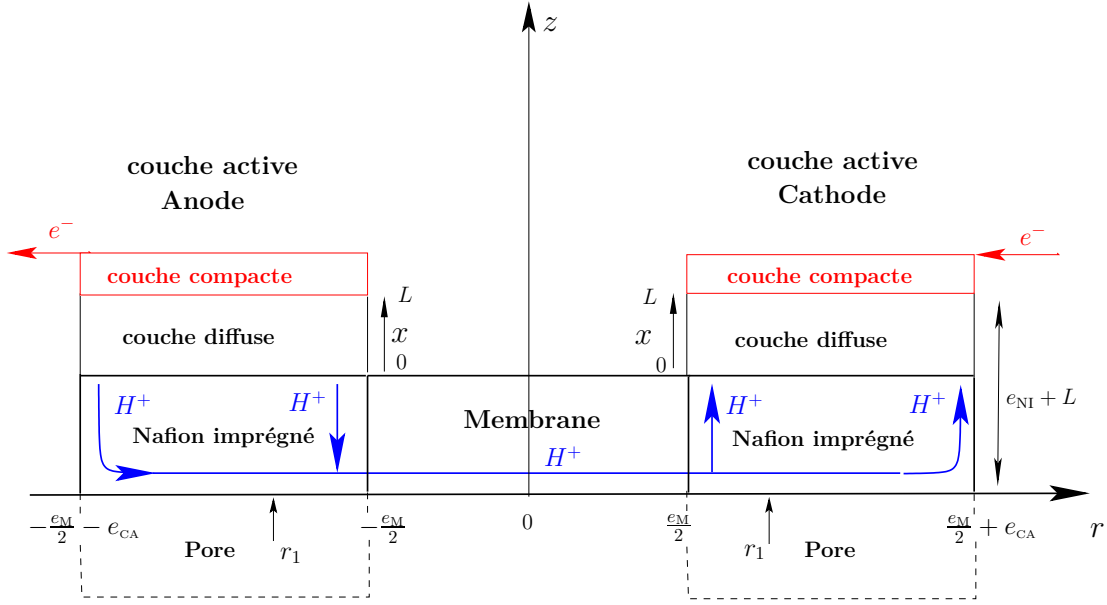


FIGURE 3.3 – Schéma du modèle d'A. Franco de la pile.

L'ordre de grandeur des dimensions est donné par :

r_1	largeur du pore	entre 15 et 40 nm
e_{NI}	épaisseur du Nafion imprégné	100 nm
e_M	épaisseur de la membrane	50×10^3 nm
e_{CA}	profondeur du pore	15×10^3 nm
L	épaisseur de la couche diffuse	1 nm
d	épaisseur de la couche compacte	0.2 nm

On a donc :

$$e_M \gg e_{NI} \gg L.$$

3.1.3 Une vision modifiée du modèle d'A. Franco

Dans la suite, nous utiliserons un schéma simplifié du modèle d'A. Franco donné par la Figure 3.4. Dans ce schéma, compte tenu du fait que la pression des gaz est constante dans la phase poreuse, l'épaisseur du pore a disparu : le pore est spatialement situé en $\pm \frac{e_M}{2}$, zone de contact "virtuelle" du gaz avec le Nafion imprégné. La couche active cathodique (resp. anodique) est représentée avec une rotation d'angle $\frac{\pi}{2}$ (resp. $-\frac{\pi}{2}$).

Les divers zones de la pile sont décrites de la façon suivante.

- En $\pm z \in]\frac{e_M}{2}, \frac{e_M}{2} + e_{NI}[$, on trouve la phase du Nafion imprégné.
- En $z = \pm(e_{NI} + \frac{e_M}{2})$, on trouve l'interface métal/Nafion imprégné, où les charges électriques sont produites (anode) et consommées (cathode).
- En $z = \pm(e_{NI} + \frac{e_M}{2})$, $x = L$ la couche compacte.
- En $z = \pm(e_{NI} + \frac{e_M}{2})$, $x \in]0, L[$ la couche diffuse.
- En $z = \pm(e_{NI} + \frac{e_M}{2})$, $x = 0$ la frontière couche diffuse/Nafion imprégné.
- En $z \in]-\frac{e_M}{2}, \frac{e_M}{2}[$, on trouve la membrane, dans laquelle les protons sont transportés vers la cathode.
- En $z = \pm \frac{e_M}{2}$, $r \in [0, e_{CA}]$ le lieu "ponctuel" du pore.

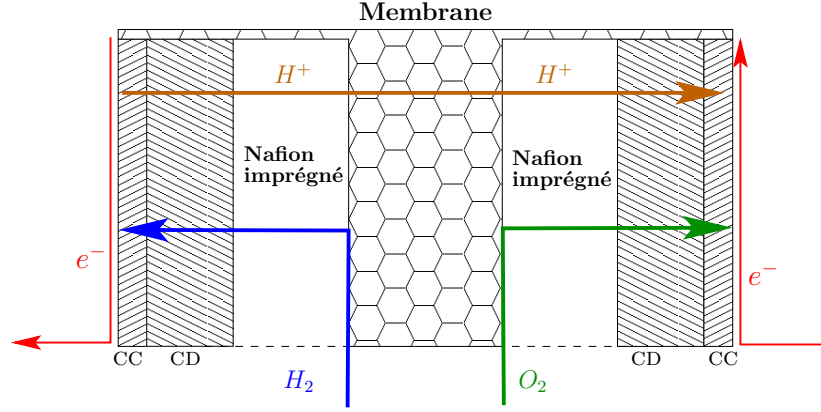


FIGURE 3.4 – Schéma simplifié du modèle d'A. Franco de la pile. A l'échelle microscopique, la pile est un système distribué à deux dimensions z et r . En $z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2})$ où se trouve l'interface métal/Nafion, une échelle plus fine à une dimension, x , permet de décrire les réactions électrochimiques. En $z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2})$, $x = \pm L$ on trouve la couche compacte CC, siège des réactions chimiques ; la couche diffuse CD est située en $z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2})$, $0 < \pm x < L$. Les mêmes variables z et r sont utilisées de l'anode à la cathode.

- En $r = 0$, l'interface couche active/GDL.
- En $r = e_{\text{CA}}$, l'interface couche active/membrane.

Le tableau suivant montre la correspondance géométrique entre le modèle de Franco représenté sur la Figure 3.3 et le modèle modifié selon la Figure 3.4.

Zone	Modèle de Franco (Figure 3.3)	Modèle de Franco modifié selon la Figure 3.4
“Axe du pore”	$\pm r \in [\frac{e_{\text{M}}}{2}, \frac{e_{\text{M}}}{2} + e_{\text{CA}}]$	$r \in [0, e_{\text{CA}}]$
Nafion imprégné	$z \in [0, e_{\text{NI}}]$	$\pm z \in [\frac{e_{\text{M}}}{2}, \frac{e_{\text{M}}}{2} + e_{\text{NI}}]$
Couche compacte	$x = L$	$x = \pm L$
Couche diffuse	$x \in]0, L[$	$\pm x \in]0, L[$
Membrane	$r \in] - \frac{e_{\text{M}}}{2}, \frac{e_{\text{M}}}{2} [$	$z \in] - \frac{e_{\text{M}}}{2}, \frac{e_{\text{M}}}{2} [$
Pore	$z \in] - r_1, 0 [$	$z = \pm \frac{e_{\text{M}}}{2}$

La vision résumée par la Figure 3.4 est à la base des équations développées plus bas.

3.1.4 Modèle microscopique de diffusion des gaz réactifs dans le Nafion

Le premier modèle microscopique décrit la diffusion des gaz réactifs vers les sites catalytiques à travers la phase du Nafion imprégné ($\pm z \in [\frac{e_{\text{M}}}{2}, \frac{e_{\text{M}}}{2} + e_{\text{NI}}]$). L'évolution des concentrations C_{H_2} et C_{O_2} en hydrogène à l'anode et en oxygène à la cathode est donné de la façon suivante.

• A l'anode

$$\frac{\partial C_{H_2,m}}{\partial t}(r, z, t) = D_{H_2} \frac{\partial^2 C_{H_2,m}}{\partial z^2}(r, z, t), \quad z \in]-\frac{e_M}{2} - e_{NI}, -\frac{e_M}{2}[, \quad r \in]0, e_{CA}[\quad (3.1a)$$

$$C_{H_2,m}(r, z = -\frac{e_M}{2}, t) = \frac{P_{H_2}}{H_{H_2}}, \quad r \in]0, e_{CA}[\quad (3.1b)$$

• A la cathode

$$\frac{\partial C_{O_2,m}}{\partial t}(r, z, t) = D_{O_2} \frac{\partial^2 C_{O_2,m}}{\partial z^2}(r, z, t), \quad r \in]\frac{e_M}{2}, \frac{e_M}{2} + e_{NI}[, \quad r \in]0, e_{CA}[\quad (3.2a)$$

$$C_{O_2,m}(r, z = \frac{e_M}{2}, t) = \frac{P_{O_2}}{H_{O_2}}, \quad r \in]0, e_{CA}[\quad (3.2b)$$

• Relations de transmission micro/nano

$$C_{H_2,m}(r, z = -\frac{e_M}{2} - e_{NI}, t) = C_{H_2,n}(r, x = 0, t), \quad r \in]0, e_{CA}[\quad (3.3a)$$

$$\frac{\partial C_{H_2,m}}{\partial z}(r, z = -\frac{e_M}{2} - e_{NI}, t) = \frac{\partial C_{H_2,n}}{\partial x}(r, x = 0, t), \quad r \in]0, e_{CA}[\quad (3.3b)$$

$$C_{O_2,m}(r, z = \frac{e_M}{2} + e_{NI}, t) = C_{O_2,n}(r, x = 0, t), \quad r \in]0, e_{CA}[\quad (3.3c)$$

$$\frac{\partial C_{O_2,m}}{\partial z}(r, z = \frac{e_M}{2} + e_{NI}, t) = \frac{\partial C_{O_2,n}}{\partial x}(r, x = 0, t), \quad r \in]0, e_{CA}[\quad (3.3d)$$

Les équations (3.1a) et (3.2a) sont des équations de diffusion. La condition (3.1b) indique que la concentration de l'hydrogène en $z = -\frac{e_M}{2}$ est supposée être celle dans le milieu poreux, où C_{H_2} et D_{H_2} sont respectivement la concentration et le coefficient de diffusion de l'hydrogène, P_{H_2} est la pression de l'hydrogène supposée constante et considérée comme une entrée du modèle, H_{H_2} est la constante de Henry qui dépend de la température. Il en est de même pour l'oxygène.

(3.3a)+(3.3b) et (3.3c)+(3.3d) sont les conditions de transmission des flux de gaz à l'interface $\pm z = \frac{e_M}{2} + e_{NI}$ entre les deux modèles microscopique et nanoscopique, où x est la variable d'espace du modèle à l'échelle nanoscopique. L'indice n dans les termes de droite est introduit pour rappeler que les fonctions des variables x et t sont des quantités nanoscopiques, contrairement aux fonctions des termes de gauche, définies à l'échelle microscopique (et notées par l'indice m quand le contexte l'exige). Voir la référence [37] pour les détails.

3.1.5 Modèle microscopique du transport des charges

Le deuxième modèle microscopique décrit d'une part le transport des électrons dans la phase métallique (Pt/C) de la couche active, et d'autre part le transport des protons dans la phase du Nafion imprégné de la couche active et dans la membrane.

• **Mouvement électronique dans le domaine $r \in [0, e_{CA}]$, $z = \pm(e_{NI} + \frac{e_M}{2})$ (Interface métal/Nafion)**

Pour les électrons, on a :

$$\frac{\partial i(r, \pm(e_{NI} + \frac{e_M}{2}), t)}{\partial r} = \gamma S_{EME} J(r, \pm(e_{NI} + \frac{e_M}{2}), t), \quad r \in]0, e_{CA}[\quad (3.4a)$$

$$i(r, \pm(e_{NI} + \frac{e_M}{2}), t) = -g_{e-} S_{EME} \frac{\partial \psi}{\partial r}(r, \pm(e_{NI} + \frac{e_M}{2}), t), \quad r \in]0, e_{CA}[\quad (3.4b)$$

$$i(r = e_{CA}, \pm(e_{NI} + \frac{e_M}{2}), t) = 0 \quad (3.4c)$$

où $i(r, \pm(e_{NI} + \frac{e_M}{2}), t)$ est le courant électronique local (microscopique) ; $J(r, \pm(e_{NI} + \frac{e_M}{2}), t)$ est la densité de courant issue de l'échelle nanoscopique ; $\psi(r, t)$ est le potentiel électrique dans la phase métallique ; S_{EME} est la surface de l'électrode et γ la surface spécifique, et enfin g_{e-} la conductivité électronique du milieu Pt/C. Notons que la surface de la couche active est S_{EME} , son volume est $e_{CA} S_{EME}$, et la surface active est $\gamma e_{CA} S_{EME}$, par définition de γ .

Le courant électronique local $i(r, \pm(e_{NI} + \frac{e_M}{2}), t)$ obéit à la loi de conservation selon l'équation (3.4a), il est proportionnel au gradient du potentiel électronique $\psi(r, t)$ selon la loi d'Ohm classique (équation (3.4b)). La condition (3.4c) provient du fait que les électrons ne circulent pas dans la membrane.

La résolution du problème (3.4a)+(3.4c), permet de déduire les valeurs

$$I_-(t) \doteq -i(0, -(e_{NI} + \frac{e_M}{2}), t), \quad I_+(t) \doteq -i(0, (e_{NI} + \frac{e_M}{2}), t). \quad (3.5)$$

$I_{\pm}$ représentent les courants électriques aux bornes de chacune des électrodes. On remarque que si $J(r, \pm(e_{NI} + \frac{e_M}{2}), t)$ est indépendant de r , i.e. $J(r, -(e_{NI} + \frac{e_M}{2}), t) \equiv J_-(t)$, $J(r, e_{NI} + \frac{e_M}{2}, t) \equiv J_+(t)$, alors

$$i(r, \pm(e_{NI} + \frac{e_M}{2}), t) = \gamma S_{EME} J_{\pm}(t)(r - e_{CA}), \quad r \in [0, e_{CA}] \quad (3.6)$$

d'où

$$I_-(t) = \gamma S_{EME} J_-(t) e_{CA}, \quad I_+(t) = \gamma S_{EME} J_+(t) e_{CA}. \quad (3.7)$$

On définit la sortie (macroscopique) qui nous intéresse, à savoir la différence de potentiel aux bornes de la pile, comme :

$$U_{cell}(t) = \psi(r = 0, z = e_{NI} + \frac{e_M}{2}, t) - \psi(r = 0, z = -(e_{NI} + \frac{e_M}{2}), t). \quad (3.8)$$

• **Mouvement protonique dans le domaine $r \in [0, e_{CA}]$, $z \in]-(e_{NI} + \frac{e_M}{2}), e_{NI} + \frac{e_M}{2}[$ (Nafion imprégné et membrane)**

Le courant protonique dans la phase du Nafion s'écrit sous la forme :

$$i_{H^+,m}(r, z, t) = F D_{+,M} \left(\nabla C_{H^+,m} + \frac{F}{RT} C_{H^+,m} \nabla \phi_m \right) (r, z, t) \quad (3.9)$$

Dans l'équation précédente, $i_{H^+,m}(r, t)$ et $\phi_m(r, t)$ sont respectivement le courant protonique et le potentiel électrique dans la phase du Nafion. $D_{+,M}$ est le coefficient de diffusion protonique dans la membrane. On fait maintenant quatre hypothèses.

1. Le gradient de la concentration $C_{H^+,m}$ est supposé négligeable en dehors de l'interface métal/Nafion ($z = \pm(e_{NI} + \frac{e_M}{2})$). Les variations de $C_{H^+,m}$ sont représentées par celles de $C_{H^+,n}$.

2. Le potentiel électrique est homogène en r .
3. $C_{H^+,m}(r, z, t) = C_{FIX}, z \in]-e_{NI} - \frac{e_M}{2}, e_{NI} + \frac{e_M}{2}[$. Le Nafion est donc considéré comme un milieu parfaitement conducteur et électroneutre.
4. Le courant protonique est discontinu aux bords microscopiques $z = \pm(e_{NI} + \frac{e_M}{2})$. La continuité apparaîtra en détaillant le comportement dans les couches limites $x \in [0, L]$.

En appliquant ces hypothèses, il reste :

$$\frac{\partial i_{H^+,m}}{\partial z}(r, z, t) = 0, \quad r \in [0, e_{CA}], \quad z \in]-e_{NI} - \frac{e_M}{2}, e_{NI} + \frac{e_M}{2}[\quad (3.10a)$$

$$i_{H^+,m}(r, z, t) = -g_{H^+} S_{EME} \frac{e_{CA}}{e_{NI}} \frac{\partial \phi_m}{\partial z}(r, z, t), \quad r \in [0, e_{CA}], \quad z \in]-e_{NI} - \frac{e_M}{2}, e_{NI} + \frac{e_M}{2}[\quad (3.10b)$$

$$C_{H^+,m}(r, z, t) \equiv C_{FIX}, \quad r \in [0, e_{CA}], \quad z \in]-e_{NI} - \frac{e_M}{2}, e_{NI} + \frac{e_M}{2}[\quad (3.10c)$$

où g_{H^+} est la conductivité protonique du Nafion, elle est liée à $D_{+,M}$ par l'identité :

$$g_{H^+} S_{EME} = \frac{F^2 D_{+,M} C_{FIX}}{RT}.$$

Le terme $\frac{e_{CA}}{e_{NI}}$ dans l'équation (3.10b) est introduit en passant de la vision symbolisée par la Figure 3.3 à celle en Figure 3.4 pour tenir compte de la différence entre les longueurs parcourues par les protons dans chacun des deux cas.

On en déduit que le courant $i_{H^+,m}(r, z, t)$ dans la membrane ne dépend pas de z , et on note par définition $i_{H^+,M}(r, t)$ cette valeur :

$$i_{H^+,m}(r, z, t) = i_{H^+,M}(r, t), \quad r \in [0, e_{CA}], \quad z \in]-e_{NI} - \frac{e_M}{2}, e_{NI} + \frac{e_M}{2}[. \quad (3.11a)$$

On déduit alors de (3.10b) la relation entre les potentiels microscopiques en $\pm(e_{NI} + \frac{e_M}{2})$:

$$\begin{aligned} \phi_m(r, z = (e_{NI} + \frac{e_M}{2}), t) - \phi_m(r, z = -e_{NI} - \frac{e_M}{2}, t) &= -\frac{e_M + 2e_{NI}}{g_{H^+} S_{EME}} \frac{e_{NI}}{e_{CA}} i_{H^+,M}(r, t) \\ &\approx -\frac{e_M e_{NI}}{g_{H^+} S_{EME} e_{CA}} i_{H^+,M}(r, t), \quad r \in [0, e_{CA}]. \end{aligned} \quad (3.11b)$$

• Conditions de transmission micro/nano

Les conditions de transmission entre les milieux microscopique et nanoscopique sont données par :

$$C_{H^+,m}(r, \pm(e_{NI} + \frac{e_M}{2}), t) = C_{FIX} = C_{H^+,n}(r, \pm(e_{NI} + \frac{e_M}{2}), x = 0, t), \quad r \in]0, e_{CA}[\quad (3.12a)$$

$$i_{H^+,m}(r, \pm(e_{NI} + \frac{e_M}{2}), t) = F i_{H^+,ca,\pm}(r, t), \quad r \in [0, e_{CA}] \quad (3.12b)$$

$$\phi_m(r, z = \pm(e_{NI} + \frac{e_M}{2}), t) = \phi_n(r, z = \pm(e_{NI} + \frac{e_M}{2}), x = 0, t), \quad r \in [0, e_{CA}] \quad (3.12c)$$

où les flux molaires protoniques $i_{H^+,ca,\pm}$, issus de la description nanoscopique, seront définis au paragraphe suivant, cf. (3.15).

On déduit de (3.11a), (3.12b) :

$$i_{H^+,ca,\pm}(r, t) = \frac{1}{F} i_{H^+,M}(r, t), \quad r \in [0, e_{CA}]. \quad (3.13)$$

3.1.6 Modèle nanoscopique du transport protonique

Le modèle nanoscopique décrit les phénomènes interfaciaux métal/Nafion de la couche active, située en $z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2})$, $r \in [0, e_{\text{CA}}]$, où ont lieu les réactions électrochimiques ainsi que le phénomène de la double couche électrique. La variable d'espace utilisée pour décrire les phénomènes à cette échelle est x , et le modèle nanoscopique est divisé en deux parties : la couche diffuse ($0 < x < L$) et la couche compacte (en $x = L$).

- La couche compacte (en $x = L$) : C'est l'endroit où ont lieu les réactions chimiques. C'est un domaine 0D d'adsorption de molécule d'eau et des espèces intermédiaires des réactions électrochimiques, et de transport des électrons entre les phases électrolytes et métallique. Son épaisseur est notée d , elle est de l'ordre du diamètre d'une molécule d'eau ($2 \times 10^{-10} \text{ m}$).
- La couche diffuse ($0 < x < L$) : Son épaisseur L est de l'ordre d'un nanomètre. C'est un domaine 1D, non électriquement neutre, dans lequel les protons produits (resp. consommés) dans la couche compacte sont transportés vers (resp. extraits de) la membrane par diffusion et migration sous l'effet d'un champ électrique. Cette couche est un isolant électronique, elle contient les gaz ayant diffusé dans le Nafion imprégné, de l'eau, des protons libres qui ne sont pas liés au réseau d'ions sulfonate fixes.

Les équations du transport protonique et du potentiel électrique dans la couche diffuse sont données par :

$$\begin{aligned} \frac{\partial C_{H^+,n}}{\partial t}(r, \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x, t) &= D_+ \frac{\partial}{\partial x} \left(\frac{\partial C_{H^+,n}}{\partial x}(r, \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x, t) \right. \\ &\quad \left. + \frac{F}{RT} C_{H^+,n}(r, \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x, t) \frac{\partial \phi_n}{\partial x}(r, \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x, t) \right), \quad r \in]0, e_{\text{CA}}[, \pm x \in]0, L[\end{aligned} \quad (3.14a)$$

$$\frac{\partial^2 \phi_n}{\partial x^2}(r, \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x, t) = -\frac{F}{\varepsilon_{cd}} (C_{H^+,n}(r, \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x, t) - C_{\text{FIX}}), \quad r \in]0, e_{\text{CA}}[, \pm x \in]0, L[\quad (3.14b)$$

$$D_+ \left(\frac{\partial C_{H^+,n}}{\partial x} + \frac{F}{RT} C_{H^+,n} \frac{\partial \phi_n}{\partial x} \right) (r, z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x = \pm L, t) = -\gamma_{e_{\text{CA}}} j_{H^+,cc,\pm}(r, t), \quad r \in]0, e_{\text{CA}}[\quad (3.14c)$$

$$C_{H^+,n}(r, z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x = 0, t) = C_{\text{FIX}} \quad (3.14d)$$

$$\phi_n(r, x = 0, t) = \phi_m(r, z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), t), \quad r \in [0, e_{\text{CA}}] \quad (3.14e)$$

$$\frac{\partial \phi}{\partial x}(r, z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x = \pm L, t) = -\frac{\sigma_{\pm}(t)}{\varepsilon_{cc}} \quad (3.14f)$$

Les courants mentionnés dans les conditions de transmission (3.12b) peuvent maintenant être définis, par les formules :

$$-i_{H^+,ca,\pm}(r, t) \doteq S_{\text{EME}} D_+ \left(\frac{\partial C_{H^+,n}}{\partial x} + \frac{F}{RT} C_{H^+,n} \frac{\partial \phi_n}{\partial x} \right) (r, z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x = 0, t), \quad r \in]0, e_{\text{CA}}[\quad (3.15)$$

On déduit alors de (3.13) la relation :

$$S_{\text{EME}} F D_+ \left(\frac{\partial C_{H^+,n}}{\partial x} + \frac{F}{RT} C_{H^+,n} \frac{\partial \phi_n}{\partial x} \right) (r, z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), x = 0, t) = -i_{H^+,M}(r, t) \quad r \in]0, e_{\text{CA}}[\quad (3.16)$$

Les équations précédentes modélisent le transport des protons dans la couche diffuse par diffusion et migration, selon l'équation de Nernst-Planck donnée par (3.14a). L'équation (3.14b), de type Poisson, décrit le couplage entre le potentiel électrique ϕ_n et la concentration protonique C_{H^+} . En $x = L$, les densités du courant de charge $j_{H^+,cc,\pm}(r, t)$ mentionnés en (3.14c) sont donnés par la chimie, cf. plus bas (3.20a) et (3.24d). Les équations (3.14d) et (3.14e) indiquent la continuité des concentrations et des potentiels entre les échelles nanoscopique et microscopique. Les équations (3.14f) représente les conditions aux bords en utilisant la loi de Gauss dans la couche compacte (voir [37]). Le produit sans dimension γe_{CA} dans (3.14c) intervient pour prendre en compte la surface active $\gamma e_{CA} S_{EME}$, et non la surface apparente S_{EME} .

Remarque 3.1. *La solution des équations précédentes ne dépend pas de r dès que les conditions aux bords n'en dépendent pas.*

3.1.7 Modèle nanoscopique du mouvement des espèces dissoutes dans les couches diffuses

L'évolution des gaz réactifs dans les couches diffuses est donné par les équations :

$$\frac{\partial C_{H_2,n}}{\partial t}(r, x, t) = D_{H_2} \frac{\partial^2 C_{H_2,n}}{\partial x^2}(r, x, t), \quad r \in] -\frac{e_M}{2} - e_{CA}, -\frac{e_M}{2}[, \quad x \in]0, L[\quad (3.17a)$$

en $x = 0$: (3.3a), (3.3b), et en $x = L$:

$$D_{H_2} \frac{\partial C_{H_2,n}}{\partial x}(r, L, t) = i_{H_2}(r, t), \quad r \in] -\frac{e_M}{2} - e_{CA}, \frac{e_M}{2}[\quad (3.17b)$$

$$\frac{\partial C_{O_2,n}}{\partial t}(r, x, t) = D_{O_2} \frac{\partial^2 C_{O_2,n}}{\partial x^2}(r, x, t), \quad r \in]\frac{e_M}{2}, \frac{e_M}{2} + e_{CA}[, \quad x \in]0, L[\quad (3.17c)$$

en $x = 0$: (3.3c), (3.3d), et en $x = L$:

$$D_{O_2} \frac{\partial C_{O_2,n}}{\partial x}(r, L, t) = -i_{O_2}(r, t), \quad r \in]\frac{e_M}{2}, \frac{e_M}{2} + e_{CA}[\quad (3.17d)$$

Les courants i_{H_2}, i_{O_2} seront donnés par la chimie, cf. (3.20b), (3.24i) plus bas.

3.1.8 Modèle nanoscopique de la chimie dans la couche compacte

• **A l'anode** , $r \in] -\frac{e_M}{2} - e_{CA}, \frac{e_M}{2}[$

Les cinétiques chimiques utilisées par A. Franco correspondent aux étapes élémentaires d'un mécanisme réactionnel sur les sites catalytiques de la couche compacte écrit selon le schéma de "Tafel-Heyrovsky-Volmer", [46] :



Dans les équations précédentes, le symbole s désigne un site catalytique libre, et on indique par Hs la forme adsorbée de l'hydrogène (il s'agit ici d'une adaptation de l'article [39] d'A. Franco). L'évolution des espèces est donnée par :

$$\frac{n^{\max}}{N_A} \frac{\partial \theta_{Hs}}{\partial t}(r, t) = -j_{H^+,cc,-}(r, t) + 2i_{H_2}(r, t) \quad (3.19a)$$

$$\frac{\partial \sigma_-(r, t)}{\partial t} = -J(r, e_{NI}, t) + Fj_{H^+,cc,-}(r, t) \quad (3.19b)$$

où les termes cinétiques sont définis par (on omet l'indice n dans C_{H_2})

$$j_{H^+,cc,-}(r, t) \doteq v_{HEY}(r, t) + v_{VOL}(r, t) \quad (3.20a)$$

$$i_{H_2}(r, t) \doteq v_{TAF}(r, t) + v_{HEY}(r, t) \quad (3.20b)$$

$$v_{TAF}(r, t) \doteq k_{TAF}\theta_{s,-}^2(r, t)C_{H_2}(r, -L, t) - k_{-TAF}\theta_{Hs}^2(r, t) \quad (3.20c)$$

$$v_{HEY}(r, t) \doteq k_{HEY}\theta_{s,-}(r, t)C_{H_2}(r, -L, t) \exp\left(\frac{\eta_-(r, t)}{2RT/F}\right) \\ - k_{-HEY}\theta_{Hs}(r, t)C_{H^+}(r, -L, t) \exp\left(-\frac{\eta_-(r, t)}{2RT/F}\right) \quad (3.20d)$$

$$v_{VOL}(r, t) \doteq k_{VOL}\theta_{Hs}(r, t) \exp\left(\frac{\eta_-(r, t)}{2RT/F}\right) \\ - k_{-VOL}\theta_{s,-}(r, t)C_{H^+}(r, -L, t) \exp\left(-\frac{\eta_-(r, t)}{2RT/F}\right) \quad (3.20e)$$

$$\theta_{s,-}(r, t) \doteq 1 - \theta_{Hs}(r, t) - \vec{\theta}_w(r, t) - \overleftarrow{\theta}_w(r, t) = \frac{1 - \theta_{Hs}(r, t)}{1 + a_w K \cosh(X_-(r, t))} \quad (3.20f)$$

$$\vec{\theta}_w(r, t) \doteq \frac{1}{2} K a_w \theta_s(r, t) \exp(-X_-(r, t)), \quad \overleftarrow{\theta}_w(r, t) \doteq \frac{1}{2} K a_w \theta_s(r, t) \exp(X_-(r, t)) \quad (3.20g)$$

$$\eta_-(r, t) \doteq -\left(1 + \frac{d^2 n^{\max}}{\varepsilon_{cd} A}\right) \frac{d}{\varepsilon_{cc}} \sigma_-(r, t) + \frac{k_B T d^3 n^{\max}}{\varepsilon_{cd} A \mu} X_-(r, t) \quad (3.20h)$$

$$K a_w \theta_{s,-}(r, t) \sinh(X_-(r, t)) = \frac{d^3}{\varepsilon_{cc} A \mu} \sigma_-(r, t) - \frac{k_B T d^3}{A \mu^2} X_-(r, t). \quad (3.20i)$$

Ici, $\theta_{Hs}(r, t)$ (respectivement θ_s) est le taux surfacique de recouvrement par l'hydrogène (respectivement par des sites catalytiques libres), $\vec{\theta}_w$ (respectivement $\overleftarrow{\theta}_w$) est la fraction de recouvrement en dipôles orientés vers le métal (respectivement opposé au métal), $\sigma_-(r, t)$ est la densité surfacique des charges électroniques. La différence de potentiel interfaciale $\eta_-(r, t)$ entre l'électrolyte et le métal (cf. (3.20h)) résulte de la superposition de deux phénomènes : l'effet capacitif dû à la présence de charges positives et négatives en regard ; et un phénomène dû à la nature dipolaire des molécules d'eau. Le calcul conduisant à l'obtention de l'expression de η_- (et de même pour η_+ à la cathode) est détaillé dans l'annexe B.

L'équation (3.19a) résulte de l'occupation et de la libération des sites catalytiques telles que modélisées par (3.18). Le bilan des charges (3.19b) indique que les charges électroniques stockées dans la double couche sont celles créées à la couche compacte qui ne participent pas à la densité de courant $J(r, t)$, cf. [37] pour plus de détails sur cette analyse, qui conduit aux formules (3.20f), (3.20h) et (3.20i).

La formule (3.20h) qui donne la surtension η , est la relation entre la différence de potentiel aux bords de la couche compacte et les différents types de charges : charges électroniques correspondant à l'effet capacitif usuel (σ) et charges dipolaires (X). L'équation algébrique

(3.20i) permet de déterminer la valeur de la concentration dipolaire $X_-(r, t)$ en fonction de $\sigma_-(r, t)$ et $\theta_{s,-}(r, t)$.

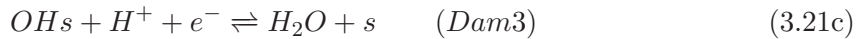
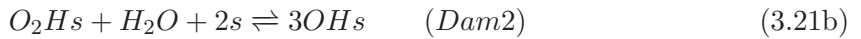
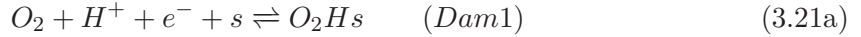
Les formules (3.20c) à (3.20f) donnent les quantités $v_{VOL}(r, t)$, $v_{HEY}(r, t)$, $v_{TAF}(r, t)$, $\theta_{s,-}(r, t)$ comme fonctions uniquement de $\theta_{Hs}(r, t)$ et $\sigma_-(r, t)$.

$j_{H^+,cc,\pm}$ sont des densités du courant de charge, i_{H_2} ainsi que plus loin i_{O_2} désignent des flux molaires de même que $i_{H^+,ca,-}$ et $i_{H^+,ca,+}$ introduits précédemment. k_i et k_{-i} , $i \in \{TAF, HEY, VOL\}$ sont les constantes cinétiques standards pour chaque étape, T est la température, a_w est l'activité de l'eau dans la couche diffuse, μ est le moment dipolaire d'une molécule d'eau. ε_{cd} (resp. ε_{cc}) est la permittivité électrique dans la couche diffuse (resp. la couche compacte), $A \doteq 1.2/2\pi\varepsilon_{cd}$, k_B la constante de Boltzmann, d est l'épaisseur d'une molécule d'eau. $K \doteq 2 \exp(-\Delta G_{c,a}^0/RT)$ (ce paramètre est noté a dans [37]) où $\Delta G_{c,a}^0$ est l'énergie chimique d'adsorption. Le tableau suivant récapitule les notations utilisées pour les courant et les flux molaires.

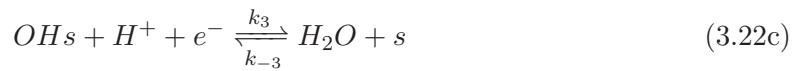
courants de charges	$i_{H^+,m}, i_{H^+,M}$	$A = C.s^{-1}$
courants molaires	$i_{H^+,ca,\pm}$	$mol.s^{-1}$
flux molaires	i_{H_2}, i_{O_2}	$mol.m^{-2}.s^{-1}$
flux de charges	$j_{H^+,cc,\pm}$	$mol.m^{-2}.s^{-1}$

• **A la cathode**, $r \in]\frac{e_M}{2}, \frac{e_M}{2} + e_{CA}[$

Le mécanisme réactionnel de Damjanovic et Brusic [32] qui permet de décrire le mécanisme de réduction de l'oxygène est adopté par Franco [37]. Ce mécanisme réactionnel s'écrit sous la forme :



Les réactions mises en jeu à la cathode sont modélisées par



où, comme pour l'anode, les s désignent des sites catalytiques et des espèces activées. La cinétique est donnée par :

$$\frac{n^{\max}}{N_A} \frac{\partial \theta_{O_2Hs}}{\partial t}(r, t) = v_1(r, t) - v_2(r, t), \quad (3.23a)$$

$$\frac{n^{\max}}{N_A} \frac{\partial \theta_{OHs}}{\partial t}(r, t) = 3v_2(r, t) - v_3(r, t), \quad (3.23b)$$

$$\frac{\partial \sigma_+(r, t)}{\partial t} = J(r, e_{NI}, t) - Fj_{H^+,cc,+}(r, t), \quad (3.23c)$$

où

$$v_1(r, t) \doteq k_1 \theta_{s,+}(r, t) C_{H+}(r, L, t) C_{O_2}(r, L, t) \exp\left(\frac{-\alpha_1 \eta_+(r, t)}{RT/F}\right) - k_{-1} \theta_{O_2 Hs}(r, t) \exp\left(\frac{(1 - \alpha_1) \eta_+(r, t)}{RT/F}\right) \quad (3.24a)$$

$$v_2(r, t) \doteq k_2 \theta_{O_2 Hs}(r, t) a_w \theta_{s,+}^2(r, t) - k_{-2} \theta_{OHs}^3(r, t) \quad (3.24b)$$

$$v_3(r, t) \doteq k_3 \theta_{OHs}(r, t) C_{H+}(r, L, t) \exp\left(\frac{-\alpha_3 \eta_+(r, t)}{RT/F}\right) - k_{-3} \theta_{s,+}(r, t) a_w \exp\left(\frac{(1 - \alpha_3) \eta_+(r, t)}{RT/F}\right) \quad (3.24c)$$

$$j_{H+,cc,+}(r, t) \doteq v_1(r, t) + v_3(r, t) \quad (3.24d)$$

$$\begin{aligned} \theta_{s,+}(r, t) &\doteq 1 - \theta_{O_2 Hs}(r, t) - \theta_{OHs}(r, t) - \vec{\theta}_w(r, t) - \overleftarrow{\theta}_w(r, t) \\ &= \frac{1 - \theta_{OHs}(r, t) - \theta_{O_2 Hs}(r, t)}{1 + a_w K \cosh(X_-(r, t))} \end{aligned} \quad (3.24e)$$

$$\vec{\theta}_w(r, t) \doteq \frac{1}{2} K a_w \theta_s(r, t) \exp(-X_+(r, t)), \quad \overleftarrow{\theta}_w(r, t) \doteq \frac{1}{2} K a_w \theta_s(r, t) \exp(X_+(r, t)) \quad (3.24f)$$

$$\eta_+(r, t) \doteq - \left(1 + \frac{d^2 n^{\max}}{\varepsilon_{cd} A}\right) \frac{d}{\varepsilon_{cc}} \sigma_+(r, t) + \frac{k_B T d^3 n^{\max}}{\varepsilon_{cd} A \mu} X_+(r, t) \quad (3.24g)$$

$$K a_w \theta_{s,+}(r, t) \sinh(X_+(r, t)) = \frac{d^3}{\varepsilon_{cc} A \mu} \sigma_+(r, t) - \frac{k_B T d^3}{A \mu^2} X_+(r, t) \quad (3.24h)$$

$$i_{O_2}(r, t) \doteq v_1(r, t). \quad (3.24i)$$

Dans les équations précédentes, θ_{OHs} et $\theta_{O_2 Hs}$ sont les taux de recouvrement surfacique par OHs et $O_2 Hs$. On définit d'une manière analogue à l'anode les variables $\sigma_+, \theta_{s,+}, \vec{\theta}_w, \overleftarrow{\theta}_w, \eta_+$ et X_+ . Les constantes k_i $i \in \{1, -1, 2, -2, 3, -3\}$ sont les constantes cinétiques standards pour chaque étape réactionnelle.

3.1.9 Connexion des modèles de demi-piles par la membrane

Connecter les deux demi-piles par la membrane définit le comportement de la charge extérieure à la pile, et permet donc de définir la relation entre $I_-(t)$ et $I_+(t)$ et ainsi de résoudre le système. Si par exemple la charge extérieure est purement résistive, alors $I_- = I_+$.

La réunion des sous-modèles microscopiques et nanoscopique constitue le modèle de la pile, qui fournit la tension U_{cell} à partir des pressions P_{H_2}, P_{O_2} et des courants I_-, I_+ . On a vu que chacun des sous-modèles agit de façon indépendante de la valeur de r . Nous avons ainsi démontré un fait fondamental relatif au modèle d'A. Franco : les réactions et transferts de charges ne dépendent en fait pas de la variable r , dès lors que les pressions des réactifs P_{H_2} et P_{O_2} sont elles-mêmes uniformes — ce qui est le cas. Ainsi, tout se passe comme si les phénomènes à cette échelle étaient localisés en deux points situés aux extrémités de la membrane, plutôt que dans les segments $]0, e_{CA}[\times\{z = e_{NI} + \frac{e_M}{2}\}$ et $]0, e_{CA}[\times\{z = -(e_{NI} + \frac{e_M}{2})\}$.

On peut donc supprimer pour ces quantités l'indication de r . Lorsque nécessaire, on introduira deux indices, a et c , permettant de différencier anode et cathode.

3.2 Exploitation de l'invariance en espace à l'interface micro/nano du modèle de la pile

On remarque que, du fait de l'invariance en r , les interfaces microscopique/nanoscopique à chaque électrode se réduisent à des interfaces entre domaines à la même échelle, les deux “nanoscopiques” étant simplement plus petits que le microscopique et représentant des couches limites pour les phénomènes qui se produisent à l'anode et à la cathode.

L'état microscopique du système est constitué des fonctions de z $C_{H_2}(z, t)$, $C_{O_2}(z, t)$, $z \in [-(e_{NI} + \frac{e_M}{2}), -\frac{e_M}{2}] \cup [\frac{e_M}{2}, e_{NI} + \frac{e_M}{2}]$. L'état nanoscopique contient les fonctions $C_{H_2}(z = -(e_{NI} + \frac{e_M}{2}), x, t)$, $C_{O_2}(z = e_{NI} + \frac{e_M}{2}, x, t)$, $C_{H^+}(z = \pm(e_{NI} + \frac{e_M}{2}), x, t)$, $\pm x \in [0, L]$; la fonction $\phi(z = \pm(e_{NI} + \frac{e_M}{2}), x, t)$, $\pm x \in [0, L]$; et les fonctions $\theta_{H_s}(t)$, $\theta_{OH_s}(t)$, $\theta_{O_2H_s}(t)$, $\sigma_-(t)$, $\sigma_+(t)$.

On agrège maintenant les milieux microscopique et nanoscopique, et on représente l'ensemble du milieu sur un axe unique, indexé par la variable $z \in [-(e_{NI} + \frac{e_M}{2} + L), e_{NI} + \frac{e_M}{2} + L]$. Les variables seront maintenant notées sans indice m/n .

3.2.1 Modèle de diffusion des gaz aux échelles microscopique et nanoscopique

On obtient le modèle suivant. On a utilisé la formule (3.7) pour retirer les notations J_- , J_+ au profit de I_- , I_+ .

$$\frac{\partial C_{H_2}}{\partial t}(z, t) = D_{H_2} \frac{\partial^2 C_{H_2}}{\partial z^2}(z, t), \quad \frac{\partial C_{O_2}}{\partial t}(z, t) = D_{O_2} \frac{\partial^2 C_{O_2}}{\partial z^2}(z, t), \quad \pm z \in]\frac{e_M}{2}, e_{NI} + \frac{e_M}{2} + L[\quad (3.25a)$$

$$C_{H_2}(z = -\frac{e_M}{2}, t) = \frac{P_{H_2}}{H_{H_2}}, \quad C_{O_2}(z = \frac{e_M}{2}, t) = \frac{P_{O_2}}{H_{O_2}} \quad (3.25b)$$

$$D_{H_2} \frac{\partial C_{H_2}}{\partial z}(-e_{NI} - \frac{e_M}{2} - L, t) = i_{H_2}(t), \quad D_{O_2} \frac{\partial C_{O_2}}{\partial z}(e_{NI} + \frac{e_M}{2} + L, t) = -i_{O_2}(t) \quad (3.25c)$$

3.2.2 Mouvement protonique et évolution du potentiel

$$\frac{\partial C_{H^+}}{\partial t}(z, t) = D_+ \frac{\partial}{\partial z} \left(\frac{\partial C_{H^+}}{\partial z}(z, t) + \frac{F}{RT} C_{H^+}(z, t) \frac{\partial \phi}{\partial z}(z, t) \right), \quad \pm z \in]e_{NI} + \frac{e_M}{2}, e_{NI} + \frac{e_M}{2} + L[\quad (3.26a)$$

$$\frac{\partial^2 \phi}{\partial z^2}(z, t) = -\frac{F}{\varepsilon_{cd}} (C_{H^+}(z, t) - C_{FIX}), \quad \pm z \in]e_{NI} + \frac{e_M}{2}, e_{NI} + \frac{e_M}{2} + L[\quad (3.26b)$$

$$\frac{\partial^2 \phi}{\partial z^2}(z, t) = 0, \quad \pm z \in]-e_{NI} - \frac{e_M}{2}, e_{NI} + \frac{e_M}{2}[\quad (3.26c)$$

$$C_{H^+}(z, t) = C_{FIX}, \quad z \in [-e_{NI} - \frac{e_M}{2}, e_{NI} + \frac{e_M}{2}] \quad (3.26d)$$

$$D_+ \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \frac{\partial \phi}{\partial z} \right) (\pm(e_{NI} + \frac{e_M}{2} + L), t) = -\gamma_{e_{CA}} j_{H^+, cc, \pm}(t) \quad (3.26e)$$

$$\frac{\partial \phi}{\partial z}(z = \pm(e_{NI} + \frac{e_M}{2} + L), t) = -\frac{\sigma_{\pm}(t)}{\varepsilon_{cc}} \quad (3.26f)$$

$$F D_+ \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \frac{\partial \phi}{\partial z} \right) (\pm(e_{NI} + \frac{e_M}{2}), t) = -\frac{i_{H^+, M}(t)}{S_{EME}} \quad (3.26g)$$

Le courant circulant dans la membrane est alors donné, d'après (3.11b), par :

$$i_{H^+,M}(t) = -\frac{e_{CA}g_{H^+}S_{EME}}{e_M e_{NI}} \left(\phi(e_{NI} + \frac{e_M}{2}, t) - \phi(-(e_{NI} + \frac{e_M}{2}), t) \right). \quad (3.27)$$

- **Chimie dans la couche compacte**, $z = \pm(e_{NI} + \frac{e_M}{2})$, $x = L$

À l'anode

$$\begin{aligned} \frac{n^{\max}}{N_A} \frac{d\theta_{Hs}}{dt}(t) = & -j_{H^+,cc,-}(\theta_{s,-}(t), C_{H_2}(-(e_{NI} + \frac{e_M}{2} + L), t), C_{H^+}(-(e_{NI} + \frac{e_M}{2} + L), t), \sigma_-(t)) \\ & + 2i_{H_2}(\theta_{s,-}(t), C_{H_2}(-(e_{NI} + \frac{e_M}{2} + L), t), C_{H^+}(-(e_{NI} + \frac{e_M}{2} + L), t), \sigma_-(t)) \end{aligned} \quad (3.28a)$$

$$\begin{aligned} \frac{d\sigma_-(t)}{dt} = & -\frac{I_-(t)}{\gamma S_{EME} e_{CA}} \\ & + Fj_{H^+,cc,-}(\theta_{s,-}(t), C_{H_2}(-(e_{NI} + \frac{e_M}{2} + L), t), C_{H^+}(-(e_{NI} + \frac{e_M}{2} + L), t), \sigma_-(t)) \end{aligned} \quad (3.28b)$$

où les termes cinétiques sont donnés par

$$v_{TAF}(t) = k_{TAF}\theta_{s,-}^2(t)C_{H_2}(-(e_{NI} + \frac{e_M}{2} + L), t) - k_{-TAF}\theta_{Hs}^2(t) \quad (3.29a)$$

$$v_{HEY}(t) = k_{HEY}\theta_{s,-}(t)C_{H_2}(-(e_{NI} + \frac{e_M}{2} + L), t) \exp\left(\frac{\eta_-(t)}{2RT/F}\right) \quad (3.29b)$$

$$-k_{-HEY}\theta_{Hs}(t)C_{H^+}(-(e_{NI} + \frac{e_M}{2} + L), t) \exp\left(-\frac{\eta_-(t)}{2RT/F}\right) \quad (3.29c)$$

$$v_{VOL}(t) = k_{VOL}\theta_{Hs}(t) \exp\left(\frac{\eta_-(t)}{2RT/F}\right) \quad (3.29d)$$

$$-k_{-VOL}\theta_{s,-}(t)C_{H^+}(-(e_{NI} + \frac{e_M}{2} + L), t) \exp\left(-\frac{\eta_-(t)}{2RT/F}\right) \quad (3.29e)$$

$$j_{H^+,cc,-}(\theta_{s,-}(t), C_{H_2}(-(e_{NI} + \frac{e_M}{2} + L), t), C_{H^+}(-(e_{NI} + \frac{e_M}{2} + L), t), \sigma_-(t)) \doteq v_{HEY}(t) + v_{VOL}(t) \quad (3.29f)$$

$$i_{H_2}(\theta_{s,-}(t), C_{H_2}(-(e_{NI} + \frac{e_M}{2} + L), t), C_{H^+}(-(e_{NI} + \frac{e_M}{2} + L), t), \sigma_-(t)) \doteq v_{TAF}(t) + v_{HEY}(t) \quad (3.29g)$$

À la cathode

$$\frac{n^{\max}}{N_A} \frac{d\theta_{O_2Hs}}{dt}(t) = v_1(t) - v_2(t), \quad (3.30a)$$

$$\frac{n^{\max}}{N_A} \frac{d\theta_{OHs}}{dt}(t) = 3v_2(t) - v_3(t), \quad (3.30b)$$

$$\frac{d\sigma_+(t)}{dt} = \frac{I_+(t)}{\gamma S_{EME} e_{CA}} - Fj_{H^+,cc,+}(\theta_{s,+}(t), C_{O_2}(e_{NI} + \frac{e_M}{2} + L, t), C_{H^+}(e_{NI} + \frac{e_M}{2} + L, t), \sigma_+(t)) \quad (3.30c)$$

avec

$$v_1(t) = k_1 \theta_{s,+}(t) C_{H^+}(e_{\text{NI}} + \frac{e_M}{2} + L, t) C_{O_2}(e_{\text{NI}} + L, t) \exp\left(\frac{-\alpha_1 \eta_+(t)}{RT/F}\right) - k_{-1} \theta_{O_2 H s}(t) \exp\left(\frac{(1 - \alpha_1) \eta_+(t)}{RT/F}\right) \quad (3.31a)$$

$$v_2(t) = k_2 \theta_{O_2 H s}(t) a_w \theta_{s,+}^2(t) - k_{-2} \theta_{O H s, c}^3(t) \quad (3.31b)$$

$$v_3(t) = k_3 \theta_{O H s}(t) C_{H^+}(e_{\text{NI}} + \frac{e_M}{2} + L, t) \exp\left(\frac{-\alpha_3 \eta_+(t)}{RT/F}\right) - k_{-3} \theta_{s,+}(t) a_w \exp\left(\frac{(1 - \alpha_3) \eta_+(t)}{RT/F}\right) \quad (3.31c)$$

$$j_{H^+, cc, +}(\theta_{s,+}(t), C_{O_2}(e_{\text{NI}} + \frac{e_M}{2} + L, t), C_{H^+}(e_{\text{NI}} + \frac{e_M}{2} + L, t), \sigma_+(t)) \doteq v_1(t) + v_3(t) \quad (3.31d)$$

Les formules suivantes sont valables pour les deux électrodes.

$$\theta_{s,-}(t) \doteq \frac{1 - \theta_{H s}(t)}{1 + K a_w \cosh(X_-(t))}, \quad \theta_{s,+}(t) \doteq \frac{1 - \theta_{O H s}(t) - \theta_{O_2 H s}(t)}{1 + K a_w \cosh(X_+(t))} \quad (3.32a)$$

$$\eta_{\pm}(t) \doteq - \left(1 + \frac{d^2 n^{\max}}{\varepsilon_{cd} A}\right) \frac{d}{\varepsilon_{cc}} \sigma_{\pm}(t) + \frac{kT d^3 n^{\max}}{\varepsilon_{cd} A \mu} X_{\pm}(t) \quad (3.32b)$$

$$K a_w \theta_{s, \pm}(t) \sinh(X_{\pm}(t)) = \frac{d^3}{\varepsilon_{cc} A \mu} \sigma_{\pm}(t) - \frac{kT d^3}{A \mu^2} X_{\pm}(t) \quad (3.32c)$$

$$i_{H_2}(t) \doteq v_{TAF}(t) + v_{HEY}(t), \quad i_{O_2}(t) \doteq v_1(t) \quad (3.32d)$$

• **Tension de sortie** En utilisant les résultats trouvés précédemment, on trouve

$$\begin{aligned} U_{cell}(t) &= \psi(z = e_{\text{NI}} + \frac{e_M}{2} + L, t) - \psi(z = -(e_{\text{NI}} + \frac{e_M}{2} + L), t) \\ &= \eta_+(t) + \phi(e_{\text{NI}} + \frac{e_M}{2} + L, t) - \phi(-(e_{\text{NI}} + \frac{e_M}{2} + L), t) - \eta_-(t) \end{aligned} \quad (3.33)$$

La valeur de $\phi(e_{\text{NI}} + \frac{e_M}{2} + L, t) - \phi(-(e_{\text{NI}} + \frac{e_M}{2} + L), t)$ est obtenue comme sortie du système d'état, cf. (3.26). On peut définir une valeur de référence du potentiel, par exemple $\phi(-(e_{\text{NI}} + \frac{e_M}{2}), t) = 0$ pour reprendre la même convention que A. Franco, ou remarquer que le potentiel solution est défini à une constante près, puisque le système d'équations ne fait intervenir que le gradient $\frac{\partial \phi}{\partial z}$ et non ϕ lui-même.

En écrivant (cf.(3.27))

$$\begin{aligned} \phi(e_{\text{NI}} + \frac{e_M}{2} + L, t) - \phi(-(e_{\text{NI}} + \frac{e_M}{2} + L), t) &= \phi(e_{\text{NI}} + \frac{e_M}{2} + L, t) - \phi(e_{\text{NI}} + \frac{e_M}{2}, t) - \frac{e_M + 2e_{\text{NI}}}{g_{H^+} S_{\text{EME}}} i_{H^+, M}(t) \\ &\quad + \phi(-(e_{\text{NI}} + \frac{e_M}{2}), t) - \phi(-(e_{\text{NI}} + \frac{e_M}{2} + L), t) \end{aligned} \quad (3.34)$$

la tension de sortie U_{cell} apparaît, comme dans les modèles réduits usuels, égale à la somme de termes correspondants aux surtensions de chaque électrode, plus des termes liés aux pertes ohmiques et aux effets capacitifs des doubles couches.

3.2.3 Récapitulatif du modèle 1D (modèle A)

On récapitule le modèle obtenu précédemment, en spatialisant les phénomènes. En remarquant que seules les variations du potentiel ϕ interviennent, on utilisera dorénavant

$$\phi_z \doteq \frac{\partial \phi}{\partial z} \quad (3.35)$$

comme nouvelle variable. On a :

$$\frac{\partial C_{H^+}}{\partial t} = D_+ \frac{\partial}{\partial z} \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right), \quad \pm z \in]e_{NI} + \frac{e_M}{2}, e_{NI} + \frac{e_M}{2} + L[\quad (3.36a)$$

$$\frac{\partial \phi_z}{\partial z} = -\frac{F}{\varepsilon_{cd}} (C_{H^+} - C_{FIX}), \quad \pm z \in]e_{NI} + \frac{e_M}{2}, e_{NI} + \frac{e_M}{2} + L[\quad (3.36b)$$

$$\frac{\partial C_{H_2}}{\partial t}(z, t) = D_{H_2} \frac{\partial^2 C_{H_2}}{\partial z^2}(z, t), \quad -z \in]\frac{e_M}{2}, e_{NI} + \frac{e_M}{2} + L[\quad (3.36c)$$

$$\frac{\partial C_{O_2}}{\partial t}(z, t) = D_{O_2} \frac{\partial^2 C_{O_2}}{\partial z^2}(z, t), \quad z \in]\frac{e_M}{2}, e_{NI} + \frac{e_M}{2} + L[\quad (3.36d)$$

$$C_{H^+}(z, t) \equiv C_{FIX}, \quad z \in]-e_{NI} - \frac{e_M}{2}, e_{NI} + \frac{e_M}{2}[\quad (3.36e)$$

$$\frac{\partial \phi_z}{\partial z} \equiv 0, \quad z \in]-e_{NI} - \frac{e_M}{2}, e_{NI} + \frac{e_M}{2}[\quad (3.36f)$$

$$\frac{d\mathcal{X}_-}{dt} = \mathcal{F}_- \left(\mathcal{X}_-(t), C_{H^+}(-e_{NI} + \frac{e_M}{2} + L, t), C_{H_2}(-e_{NI} + \frac{e_M}{2} + L, t), I_-(t) \right) \quad (3.37a)$$

$$\frac{d\mathcal{X}_+}{dt} = \mathcal{F}_+ \left(\mathcal{X}_+(t), C_{H^+}(e_{NI} + \frac{e_M}{2} + L, t), C_{O_2}(e_{NI} + \frac{e_M}{2} + L, t), I_+(t) \right) \quad (3.37b)$$

$$D_+ \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right) (\pm(e_{NI} + \frac{e_M}{2} + L), t) = \alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) C_N(\pm z_{N+1}, t) \quad (3.37c)$$

$$\phi_z(\pm(e_{NI} + \frac{e_M}{2} + L), t) = -\frac{\sigma_{\pm}(t)}{\varepsilon_{cc}} \quad (3.37d)$$

$$FD_+ \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right) (\pm(e_{NI} + \frac{e_M}{2}), t) = -\frac{i_{H^+,M}(t)}{S_{EME}} \quad (3.37e)$$

$$C_{H_2}(z = -\frac{e_M}{2}, t) = \frac{P_{H_2}}{H_{H_2}}, \quad C_{O_2}(z = \frac{e_M}{2}, t) = \frac{P_{O_2}}{H_{O_2}} \quad (3.37f)$$

$$D_{H_2} \frac{\partial C_{H_2}}{\partial z}(-e_{NI} - \frac{e_M}{2} - L, t) = i_{H_2}(t), \quad D_{O_2} \frac{\partial C_{O_2}}{\partial z}(e_{NI} + \frac{e_M}{2} + L, t) = -i_{O_2}(t), \quad (3.37g)$$

$$\mathcal{X}_-(t) \doteq \begin{pmatrix} \theta_{Hs}(t) \\ \sigma_-(t) \end{pmatrix}, \quad \mathcal{X}_+(t) \doteq \begin{pmatrix} \theta_{OHs}(t) \\ \theta_{O_2Hs}(t) \\ \sigma_+(t) \end{pmatrix} \quad (3.37h)$$

où $\mathcal{F}_{\pm}(\mathcal{X}_{\pm})$ sont définis par (3.28) et (3.30), en interprétant les coordonnées de $\mathcal{X}_{\pm}$ comme dans (3.37h).

En utilisant les formules (3.29f), (3.29b), (3.29d) pour l'anode ; (3.31d), (3.31a), (3.31c) pour la cathode, nous avons écrit dans (3.37c), et nous écrirons par la suite, les densités du courant de charge $j_{H^+,cc,\pm}$ de la façon suivante :

$$\gamma e_{CA} j_{H^+,cc,\pm}(\mathcal{X}_{\pm}, C_{H^+}(\pm z_{N+1})) = \alpha_{\pm}(\mathcal{X}_{\pm}) + \beta_{\pm}(\mathcal{X}_{\pm}) C_{H^+}(\pm z_{N+1}) \quad (3.38a)$$

$$\alpha_{-}(t) \doteq -\gamma e_{CA} \left(k_{HEY} \theta_{s,-}(t) C_{H_2}(-z_{N+1}, t) + k_{VOL} \theta_{Hs}(t) \right) \exp \left(\frac{\eta_{-}(t)}{2RT/F} \right) \quad (3.38b)$$

$$\alpha_{+}(t) \doteq \gamma e_{CA} \left(k_{-1} \theta_{O_2 Hs}(t) \exp \left(\frac{(1-\alpha_1)\eta_{+}(t)}{RT/F} \right) + k_{-3} \theta_{s,+}(t) a_w \exp \left(\frac{(1-\alpha_3)\eta_{+}(t)}{RT/F} \right) \right) \quad (3.38c)$$

et

$$\beta_{-}(t) \doteq \gamma e_{CA} \left(k_{-HEY} \theta_{Hs}(t) + k_{-VOL} \theta_{s,-}(t) \right) \exp \left(\frac{-\eta_{-}(t)}{2RT/F} \right) \quad (3.38d)$$

$$\beta_{+}(t) \doteq -\gamma e_{CA} \left(k_1 \theta_{s,+}(t) C_{O_2}(z_{N+1}, t) \exp \left(\frac{-\alpha_1 \eta_{+}(t)}{RT/F} \right) + k_3 \theta_{OHs}(t) \exp \left(\frac{-\alpha_3 \eta_{+}(t)}{RT/F} \right) \right). \quad (3.38e)$$

• **Sorties.** Les sorties d'intérêt sont le courant membranaire, la tension aux bornes des couches diffuses, la tension aux bornes de la pile, donnés respectivement par (3.39), (3.40) et (3.41).

$$i_{H^+,M}(t) = -S_{EME} F D_{+} \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right) (\pm(e_{NI} + \frac{e_M}{2}), t) \quad (3.39)$$

$$U_{\pm}(t) = \int_{\pm e_{NI} + \frac{e_M}{2}}^{\pm(e_{NI} + \frac{e_M}{2}) + L} \phi_z(z, t) \cdot dz \quad (3.40)$$

$$\begin{aligned} U_{cell}(t) &= \eta_{+}(\mathcal{X}_{+}(t)) + \int_{-(e_{NI} + \frac{e_M}{2}) + L}^{e_{NI} + \frac{e_M}{2} + L} \phi_z(z, t) \cdot dz - \eta_{-}(\mathcal{X}_{-}(t)) \\ &= \eta_{+}(\mathcal{X}_{+}(t)) + U_{+}(t) - \frac{e_M e_{NI}}{g_{H^+} S_{EME} e_{CA}} i_{H^+,M}(t) - (U_{-}(t) + \eta_{-}(\mathcal{X}_{-}(t))) \end{aligned} \quad (3.41)$$

où $\eta_{\pm}(\mathcal{X}_{\pm})$ sont définis par (3.32b). Le **modèle A** est défini par le système d'équations : (3.36), (3.37) et de (3.39) à (3.41).

Lemme 3.1. On considère le modèle A avec un courant $I_{\pm}$ constant. À l'équilibre, on a :

$$i_{H^+,M} = I_{\pm}.$$

Démonstration. En prenant dans $\mathcal{X}_{\pm}$ les équations qui concernent $\sigma_{\pm}$, définies par (3.28b) et (3.30c), on obtient à l'équilibre les égalités suivantes :

$$I_{\pm} = \gamma S_{EME} e_{CA} F j_{H^+,cc,\pm}.$$

En intégrant (3.26a) en z entre $\pm(e_{NI} + \frac{e_M}{2})$ et $\pm(e_{NI} + \frac{e_M}{2} + L)$, on obtient à l'équilibre :

$$\begin{aligned} 0 &= \frac{d}{dt} \int_{\pm(\frac{e_M}{2} + e_{NI})}^{\pm(\frac{e_M}{2} + e_{NI} + L)} C_{H^+}(z, t) dz = \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \frac{\partial \phi}{\partial z} \right) (\pm(\frac{e_M}{2} + e_{NI} + L)) \\ &\quad - \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \frac{\partial \phi}{\partial z} \right) (\pm(\frac{e_M}{2} + e_{NI})) \end{aligned} \quad (3.42)$$

En utilisant les conditions aux bords (3.26e) et (3.26g), on a d'après l'égalité précédentes :

$$\gamma S_{\text{EMBE}} e_{\text{CA}} F j_{H^+,cc,\pm} = i_{H^+,M},$$

d'où $I_{\pm} = i_{H^+,M}$, ce qui achève la démonstration. $\square$

3.3 Réduction du modèle de transport des gaz

3.3.1 Principe de la réduction

Il est possible d'approcher les équations (3.25), décrivant la diffusion des gaz dans le Nafion imprégné, en considérant la projection de la solution sur la première fonction propre du problème de Sturm-Liouville associé. On utilise pour cela le lemme suivant.

Lemme 3.2. *Soit $C(z, t)$ vérifiant*

$$\frac{\partial C}{\partial t} = D \frac{\partial^2 C}{\partial z^2}, \quad z \in]a, b[, \quad (3.43a)$$

$$C(a, t) = c_a(t), \quad \frac{\partial C}{\partial z}(b, t) = c_b(t) \quad (3.43b)$$

pour $D > 0$ et $a \neq b$. Alors la fonction

$$c(t) \doteq \int_a^b \sin \alpha(z - a) C(z, t) \cdot dz, \quad \alpha \doteq \frac{\pi}{2(b - a)}$$

vérifie

$$\frac{dc}{dt} + D\alpha^2 c = D(\alpha c_a(t) + c_b(t)) \quad (3.44)$$

Démonstration. La preuve est directe, moyennant deux intégrations par parties. De (3.43a), on tire :

$$\begin{aligned} \frac{dc}{dt} &= D \int_a^b \sin \alpha(z - a) \frac{\partial^2 C}{\partial z^2} \cdot dz \\ &= D \left(\left[\sin \alpha(z - a) \frac{\partial C}{\partial z} \right]_a^b - \alpha \int_a^b \cos \alpha(z - a) \frac{\partial C}{\partial z} \cdot dz \right) \\ &= D c_b(t) - D \alpha \left([\cos \alpha(z - a) C]_a^b + \alpha \int_a^b \sin \alpha(z - a) C \cdot dz \right) \\ &= D(c_b(t) + \alpha c_a(t)) - D \alpha^2 \int_a^b \sin \alpha(z - a) C \cdot dz, \end{aligned}$$

d'où (3.44) et la conclusion. $\square$

Le principe de l'approximation revient à approcher la fonction $C(z, t)$ solution de (3.43) par la fonction

$$\tilde{C}(z, t) \doteq \frac{2}{b - a} \left(\int_a^b \sin \alpha(z' - a) C(z', t) \cdot dz' \right) \sin \alpha(z - a),$$

qui est donc solution de l'équation différentielle ordinaire

$$\frac{\partial \tilde{C}(z, t)}{\partial t} + D \alpha^2 \tilde{C}(z, t) = \frac{2D}{b - a} (\alpha c_a(t) + c_b(t)) = \frac{D\pi}{(b - a)^2} c_a(t) + \frac{2D}{b - a} c_b(t) \quad (3.45)$$

En revenant maintenant aux notations adoptées précédemment et en considérant la valeur terminale (c'est-à-dire en $z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2})$ et $z = \pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L)$), on obtient le modèle 0D suivant du transport des gaz, dérivé du précédent en remplaçant (3.36c) et (3.36d) par

$$\frac{dC_{H_2}}{dt}(-e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L, t) = -\frac{\pi^2 D_{H_2}}{4(e_{\text{NI}} + L)^2} \left(C_{H_2}(-e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L, t) - \frac{4}{\pi} \frac{P_{H_2}}{H_{H_2}} \right) - \frac{2i_{H_2}(t)}{e_{\text{NI}} + L} \quad (3.46a)$$

$$\frac{dC_{O_2}}{dt}(e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L, t) = -\frac{\pi^2 D_{O_2}}{4(e_{\text{NI}} + L)^2} \left(C_{O_2}(e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L, t) - \frac{4}{\pi} \frac{P_{O_2}}{H_{O_2}} \right) - \frac{2i_{O_2}(t)}{e_{\text{NI}} + L}. \quad (3.46b)$$

3.3.2 Modèle B avec réduction du modèle de gaz

En appliquant l'idée précédemment exposée, on obtient en regroupant les variables qui sont définies par une simple EDO dans les vecteurs $\mathcal{X}_{\pm}$, le modèle suivant¹ :

$$\frac{\partial C_{H^+}}{\partial t} = D_+ \frac{\partial}{\partial z} \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right), \quad \pm z \in]e_{\text{NI}} + \frac{e_{\text{M}}}{2}, e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L[\quad (3.47a)$$

$$\frac{\partial \phi_z}{\partial z} = -\frac{F}{\varepsilon_{cd}} (C_{H^+} - C_{\text{FIX}}), \quad \pm z \in]e_{\text{NI}} + \frac{e_{\text{M}}}{2}, e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L[\quad (3.47b)$$

$$C_{H^+}(z, t) \equiv C_{\text{FIX}}, \quad z \in]-e_{\text{NI}} - \frac{e_{\text{M}}}{2}, e_{\text{NI}} + \frac{e_{\text{M}}}{2}[\quad (3.47c)$$

$$\frac{\partial \phi_z}{\partial z} \equiv 0, \quad z \in]-e_{\text{NI}} - \frac{e_{\text{M}}}{2}, e_{\text{NI}} + \frac{e_{\text{M}}}{2}[\quad (3.47d)$$

$$\frac{d\mathcal{X}_{\pm}}{dt} = \mathcal{F}_{\pm}(\mathcal{X}_{\pm}, I_{\pm}, C_{H^+}(e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L), \mathcal{X}_{in, \pm}) \quad (3.48a)$$

$$D_+ \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right) (\pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L), t) = \alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) C_N(\pm z_{N+1}, t) \quad (3.48b)$$

$$\phi_z(\pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2} + L), t) = -\frac{\sigma_{\pm}(t)}{\varepsilon_{cc}} \quad (3.48c)$$

$$FD_+ \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right) (\pm(e_{\text{NI}} + \frac{e_{\text{M}}}{2}), t) = -\frac{i_{H^+, M}(t)}{S_{\text{EME}}} \quad (3.48d)$$

1. Par abus de notation, on note de la même façon les variables $\mathcal{X}_{\pm}$ que dans le paragraphe 3.2.3, même si les définitions diffèrent légèrement, cf. (3.37h) et (3.48e). De même pour les fonctions $\mathcal{F}_{\pm}$, $\alpha_{\pm}$ et $\beta_{\pm}$.

$$\mathcal{X}_-(t) \doteq \begin{pmatrix} C_{H_2}(-(e_{NI} + \frac{e_M}{2} + L), t) \\ \theta_{Hs}(t) \\ \sigma_-(t) \end{pmatrix}, \quad \mathcal{X}_+(t) \doteq \begin{pmatrix} C_{O_2}(e_{NI} + \frac{e_M}{2} + L, t) \\ \theta_{OHs}(t) \\ \theta_{O_2Hs}(t) \\ \sigma_+(t) \end{pmatrix} \quad (3.48e)$$

où $\mathcal{F}_-(\mathcal{X}_-)$ (resp. $\mathcal{F}_+(\mathcal{X}_+)$) sont définis par (3.28) et (3.46a) (resp. (3.30) et (3.46b)), en interprétant les coordonnées de $\mathcal{X}_\pm$ comme dans (3.48e). $\mathcal{X}_{in,-} \doteq P_{H_2}$, $\mathcal{X}_{in,+} \doteq P_{O_2}$ et $\alpha_\pm, \beta_\pm$ sont définis par (3.38).

• **Sorties.** Les sorties d'intérêt sont le courant membranaire, la tension aux bornes des couches diffuses, la tension aux bornes de la pile, donnés respectivement par (3.50), (3.49) et (3.51).

$$-i_{H^+,M}(t) = S_{EME} F D_+ \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right) (\pm(e_{NI} + \frac{e_M}{2}), t) \quad (3.49)$$

$$U_\pm(t) = \int_{\pm e_{NI} + \frac{e_M}{2}}^{\pm(e_{NI} + \frac{e_M}{2} + L)} \phi_z(z, t) \cdot dz \quad (3.50)$$

$$\begin{aligned} U_{cell}(t) &= \eta_+(\mathcal{X}_+(t)) + \int_{-(e_{NI} + \frac{e_M}{2} + L)}^{e_{NI} + \frac{e_M}{2} + L} \phi_z(z, t) \cdot dz - \eta_-(\mathcal{X}_-(t)) \\ &= \eta_+(\mathcal{X}_+(t)) + U_+(t) - \frac{e_M e_{NI}}{g_{H^+} S_{EME} e_{CA}} i_{H^+,M}(t) - (U_-(t) + \eta_-(\mathcal{X}_-(t))) \end{aligned} \quad (3.51)$$

où $\eta_\pm(\mathcal{X}_\pm)$ sont définies par (3.32b). **Le modèle B** est constitué par définition des équations (3.47), (3.48) et de (3.49) à (3.51).

Remarque 3.2. Le résultat obtenu dans le Lemme 3.1 reste valable pour le modèle B.

3.4 Réduction du modèle de transport des charges et modèle 0D

Afin de parvenir finalement à un unique modèle 0D, l'équation aux dérivées partielles donnant l'évolution du potentiel et de la concentration de protons dans la couche diffuse est semi discrétisée en espace par la méthode des éléments finies. La répartition des points de discrétisation est un paramètre de réglage, on imagine a priori que leur densité doit être plus forte près de la couche compacte, lieu de production (à l'anode) et de consommation (à la cathode) des protons.

En posant $z_0 = e_{NI} + \frac{e_M}{2}$ et $z_{N+1} = e_{NI} + \frac{e_M}{2} + L$, les équations décrivant le transport des charges dans les couches limites et dans la membrane sont données par :

Modèles de chacune des deux demi-piles.

Dans le domaine.

$$\frac{\partial C_{H^+}}{\partial t} = D_+ \frac{\partial}{\partial z} \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right), \quad \pm z \in]z_0, z_{N+1}[\quad (3.52a)$$

$$\frac{\partial \phi_z}{\partial z} = -\frac{F}{\varepsilon_{cd}} (C_{H^+} - C_{FIX}), \quad \pm z \in]z_0, z_{N+1}[\quad (3.52b)$$

Aux bords du domaine.

$$\frac{d\mathcal{X}_{\pm}}{dt} = \mathcal{F}_{\pm}(\mathcal{X}_{\pm}, I_{\pm}, C_{H+}(\pm z_{N+1}), \mathcal{X}_{in,\pm}) \quad (3.52c)$$

$$D_+ \left(\frac{\partial C_{H+}}{\partial z} + \frac{F}{RT} C_{H+} \phi_z \right) (\pm z_{N+1}) = \alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) C_N(\pm z_{N+1}) \quad (3.52d)$$

$$\phi_z(\pm z_{N+1}) = -\frac{\sigma_{\pm}}{\varepsilon_{cc}} \quad (3.52e)$$

$$C_{H+}(\pm z_0) = C_{FIX} \quad (3.52f)$$

Les conditions initiales sont données par :

$$\mathcal{X}_{\pm}(0) = \mathcal{X}_{\pm}^0, \quad C_{H+}(z, 0) = C_{H+}^0(z), \quad \pm z \in]z_0, z_{N+1}[.$$

On rappelle que les entrées $\mathcal{X}_{in,\pm}$, qui représentent les quantités de gaz réactifs (P_{H_2} et P_{O_2}) et le courant I sont les uniques entrées du système global. Les surtensions sont fonctions des variables d'état $\mathcal{X}_{\pm}$.

Modèle de la membrane. Les équations dans membrane sont données par :

$$C_{H+}(z) = C_{FIX}, \quad z \in [-z_0, z_0] \quad (3.53a)$$

$$\frac{\partial \phi_z}{\partial z} = 0, \quad z \in]-z_0, z_0[. \quad (3.53b)$$

Le courant membranaire est définit par :

$$i_{H+,M} \doteq S_{EME} F D_+ \left(\frac{\partial C_{H+}}{\partial z} + \frac{F}{RT} C_{H+} \phi_z \right) (\pm z_0) \quad (3.53c)$$

Sorties de chaque demi-pile. Surtension et tension anodique et cathodique :

$$\eta_{\pm}, \quad U_{\pm}$$

Courants dans la membrane :

$$i_{H+,M}.$$

3.4.1 Résolution du modèle de membrane et simplification du modèle de demi-pile

Normalisation des concentrations par translation. On translate les concentrations de la valeur C_{FIX} , pour se ramener à l'étude de

$$\tilde{C}_{H+}(z, t) \doteq C_{H+}(z, t) - C_{FIX}. \quad (3.54)$$

En intégrant (3.52b) en z , entre $\pm z_{N+1}$ et z , on obtient :

$$\phi_z(z) = \phi_z(\pm z_{N+1}) + \frac{F}{\varepsilon_{cd}} \int_z^{\pm z_{N+1}} \tilde{C}_{H+}(z') dz'. \quad (3.55a)$$

On définit la fonctionnelle (intégro-différentielle) de C_{H+} suivante :

$$\Phi_{\pm}(\tilde{C}_{H+}, \mathcal{X}_{\pm}) = \mp \frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{\varepsilon_{cd}} \int_z^{\pm z_{N+1}} \tilde{C}_{H+}(z') dz'. \quad (3.55b)$$

On a, d'après (3.55a) et (3.52e), $\phi_z(z) = \Phi_\pm(\tilde{C}_{H^+}, \mathcal{X}_\pm)(z)$ pour tout $\pm z \in [e_{\text{NI}} + \frac{e_M}{2}, e_{\text{NI}} + \frac{e_M}{2} + L]$. On en déduit que le système décrivant chaque demi-pile s'écrit sous la forme² :

$$\frac{\partial \tilde{C}_{H^+}}{\partial t} = D_+ \frac{\partial^2 \tilde{C}_{H^+}}{\partial z^2} + \frac{D_+ F}{RT} \frac{\partial}{\partial z} \left((\tilde{C}_{H^+} + C_{FIX}) \Phi_\pm(\tilde{C}_{H^+}, \mathcal{X}_\pm) \right), \quad \pm z \in]z_0, z_{N+1}[\quad (3.56a)$$

$$\frac{d\mathcal{X}_\pm}{dt} = \mathcal{F}_\pm(\mathcal{X}_\pm, I_\pm, \tilde{C}_{H^+}(\pm z_{N+1}), \mathcal{X}_{in,\pm}) \quad (3.56b)$$

$$D_+ \left(\frac{\partial \tilde{C}_{H^+}}{\partial z} + \frac{F}{RT} (\tilde{C}_{H^+} + C_{FIX}) \Phi_\pm(\tilde{C}_{H^+}, \mathcal{X}_\pm) \right) (\pm z_{N+1}) = \alpha_\pm(\mathcal{X}_\pm(t)) + \beta_\pm(\mathcal{X}_\pm(t)) C_N(\pm z_{N+1}) \quad (3.56c)$$

$$\tilde{C}_{H^+}(\pm z_0) = 0 \quad (3.56d)$$

avec

$$\Phi_\pm(\tilde{C}_{H^+}, \mathcal{X}_\pm)(z) \doteq \mp \frac{\sigma_\pm}{\varepsilon_{cc}} + \frac{F}{\varepsilon_{cd}} \int_z^{\pm z_{N+1}} \tilde{C}_{H^+}(z') dz'.$$

3.4.2 Formulation variationnelle et semi-discrétisation du modèle de demi-pile

Formulation variationnelle. Pour toute fonction $\varphi \in C(\pm[z_0, z_{N+1}]; \mathbb{R})$ telle que $\varphi(\pm z_0) = 0$, on a :

$$\begin{aligned} \frac{d}{dt} \int_{\pm z_0}^{\pm z_{N+1}} \tilde{C}_{H^+}(t) \varphi dz &= \left(\alpha_\pm(\mathcal{X}_\pm(t)) + \beta_\pm(\mathcal{X}_\pm(t)) C_N(\pm z_{N+1}, t) \right) \varphi(\pm z_{N+1}) \\ &\quad - D_+ \int_{\pm z_0}^{\pm z_{N+1}} \left(\frac{\partial \tilde{C}_{H^+}}{\partial z} + \frac{F}{RT} (\tilde{C}_{H^+} + C_{FIX}) \Phi_\pm(\tilde{C}_{H^+}, \mathcal{X}_\pm) \right) \frac{d\varphi}{dz} dz. \end{aligned} \quad (3.57)$$

Principe de la semi-discrétisation. On utilise la méthode des éléments finis $\mathbb{P}_1$, c'est-à-dire avec des fonctions affines par morceaux :

- On choisit une collection de points $e_{\text{NI}} = z_0 < z_1 < \dots < z_N < z_{N+1} = e_{\text{NI}} + L$ et on note :

$$d_j = z_{j+1} - z_j, \quad \forall j \in \{0, \dots, N\}.$$

- On approxime $\tilde{C}_{H^+}(z, t)$ par $\tilde{C}_{H^+,N}(z, t)$:

$$\tilde{C}_{H^+,N}(z, t) = \sum_{j=1}^{N+1} \tilde{C}_{H^+,N}(z_j, t) \varphi_j(z) \quad (3.58)$$

où les fonctions test φ_j sont affines sur chaque intervalle $\pm[z_{j-1}, z_j]$ et $\varphi_j(\pm z_i) = \delta_{ij}$. En d'autres termes, pour tout $j = 1, \dots, N$ on a :

$$\varphi_j(\pm z) = \begin{cases} \frac{z - z_{j-1}}{d_{j-1}} & \text{si } z \in [z_{j-1}, z_j], \\ \frac{z_{j+1} - z}{d_j} & \text{si } z \in [z_j, z_{j+1}], \\ 0 & \text{si } z \leq z_{j-1} \text{ ou } z \geq z_{j+1}, \end{cases} \quad (3.59)$$

2. Par abus, on note $\mathcal{F}_\pm(\mathcal{X}_\pm, I_\pm, \tilde{C}_{H^+}(\pm z_{N+1}), \mathcal{X}_{in,\pm})$ au lieu de $\mathcal{F}_\pm(\mathcal{X}_\pm, I_\pm, \tilde{C}_{H^+}(\pm z_{N+1}) + C_{FIX}, \mathcal{X}_{in,\pm})$. Il faut en fait ici adapter les définitions des fonctions en question.

et

$$\varphi_{N+1}(\pm z) = \begin{cases} \frac{z - z_N}{d_N} & \text{si } z \in [z_N, z_{N+1}], \\ 0 & \text{si } z \leq z_N. \end{cases}$$

Par simplicité, on écrira simplement $\tilde{C}_N(z, t)$ à la place de $\tilde{C}_{H^+, N}(z, t)$. L'entier positif N en indice se réfère au nombre des points de discrétisation.

• L'approximation du problème (3.57) que nous considérerons consiste à chercher $\tilde{C}_N(\pm z_j, t)$, $j \in \{1, \dots, N+1\}$, tels que pour tout $j \in \{1, \dots, N+1\}$:

$$\begin{aligned} \frac{d}{dt} \int_{\pm z_{j-1}}^{\pm z_{j+1}} \tilde{C}_N(z, t) \varphi_j(z) dz &= \left(\alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) C_N(\pm z_{N+1}, t) \right) \varphi_j(\pm z_{N+1}) \\ &\quad - D_+ \int_{\pm z_{j-1}}^{\pm z_{j+1}} \left(\frac{\partial \tilde{C}_N}{\partial z} + \frac{F}{RT} (\tilde{C}_N + C_{FIX}) \Phi_+(\tilde{C}_N, \mathcal{X}_{\pm}) \right) \frac{d\varphi_j}{dz} dz \end{aligned} \quad (3.60)$$

où

$$\tilde{C}_N(z, t) = \frac{(\pm z - z_{j-1}) \tilde{C}_N(\pm z_j, t) - (\pm z - z_j) \tilde{C}_N(\pm z_{j-1}, t)}{d_{j-1}}, \quad \pm z \in [z_{j-1}, z_j]. \quad (3.61)$$

Écriture du système différentiel 0D. Par simplicité on met les calculs permettant d'écrire sous forme matricielle l'équation (3.57) sont détaillés dans l'annexe C. Les formules correspondantes et portant un numéro précédé d'un C figurent dans l'annexe C. On note :

$$\mathcal{C}_{\pm}(t) \doteq \begin{pmatrix} C_N(\pm z_1, t) \\ \vdots \\ C_N(\pm z_{N+1}, t) \end{pmatrix}, \quad \tilde{\mathcal{C}}_{\pm}(t) \doteq \begin{pmatrix} \tilde{C}_N(\pm z_1, t) \\ \vdots \\ \tilde{C}_N(\pm z_{N+1}, t) \end{pmatrix} = \begin{pmatrix} C_N(\pm z_1, t) - C_{FIX} \\ \vdots \\ C_N(\pm z_{N+1}, t) - C_{FIX} \end{pmatrix} \in \mathbb{R}^{N+1} \quad (3.62)$$

D'après (C.17), le système des équations données par (3.60) s'écrit sous la forme :

$$\frac{d\tilde{\mathcal{C}}_{\pm}(t)}{dt} = \mathcal{G}_{\pm}(\mathcal{X}_{\pm}(t), \mathcal{C}_{\pm}(t)) \quad (3.63)$$

où

$$\begin{aligned} \mathcal{G}_{\pm}(\mathcal{X}_{\pm}, \mathcal{C}_{\pm}) &= M_1^{-1} \left[M_2 \tilde{\mathcal{C}}_{\pm} + M_3 \begin{pmatrix} C_{FIX} \\ \mathcal{C}_{\pm} \end{pmatrix} \cdot (M_4 \tilde{\mathcal{C}}_{\pm} + M_5^{\pm} \mathcal{X}_{\pm}) + M_6 \mathcal{C}_{\pm} \cdot M_7 \tilde{\mathcal{C}}_{\pm} \right. \\ &\quad \left. + M_8 \mathcal{C}_{\pm} \cdot M_9 \tilde{\mathcal{C}}_{\pm} + M_{10} \begin{pmatrix} C_{FIX} \\ \mathcal{C}_{\pm} \end{pmatrix} \cdot \tilde{\mathcal{C}}_{\pm} \pm (\alpha_{\pm}(\mathcal{X}_{\pm}) + \beta_{\pm}(\mathcal{X}_{\pm}) v_{N+1}^{\top} \mathcal{C}_{\pm}) v_{N+1} \right] \end{aligned} \quad (3.64)$$

où “ $\cdot$ ” représente le produit terme à terme de vecteurs. Les matrices $\{M_i\}_{i=1}^9$ sont définies par (C.16) et le vecteur v_{N+1} par (C.15). On vérifie que la matrice M_1 est inversible.

3.4.3 Modèle 0D ultime : Modèle C

En utilisant les formules (C.29), (C.21) dans l'annexe C, le système global d'équations s'écrit sous la forme :

$$\frac{d\tilde{\mathcal{C}}_{\pm}(t)}{dt} = \mathcal{G}_{\pm}(\mathcal{X}_{\pm}(t), \mathcal{C}_{\pm}(t)) \quad (3.65a)$$

$$\frac{d\mathcal{X}_{\pm}(t)}{dt} = \mathcal{F}_{\pm}(\mathcal{X}_{\pm}(t), I, v_{N+1}^{\top} \tilde{\mathcal{C}}_{\pm}(t), \mathcal{X}_{in,\pm}(t)) \quad (3.65b)$$

$$U_{\pm}(t) = \left(-\frac{\sigma_{\pm}}{\varepsilon_{cc}} L + \frac{F}{2\varepsilon_{cd}} w^{\top} \right) \tilde{\mathcal{C}}_{\pm}(t) \quad (3.65c)$$

$$i_{H^+,M}(t) = FS_{EME} \left(\alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) v_{N+1}^{\top} \mathcal{C}_{\pm}(t) - \frac{1}{2} v^{\top} \mathcal{G}_{\pm}(\mathcal{X}_{\pm}(t), \mathcal{C}_{\pm}(t)) \right) \quad (3.65d)$$

$$U_{cell}(t) = \eta_+(\mathcal{X}_+(t)) + U_+(t) - \frac{e_M}{g_{H^+} S_{EME}} i_{H^+,M}(t) - (U_-(t) + \eta_-(\mathcal{X}_-(t))) \quad (3.65e)$$

où $\mathcal{C}(t)$ et $\tilde{\mathcal{C}}(t)$ sont définies par (3.62), $\mathcal{F}_-(\mathcal{X}_-)$ (resp. $\mathcal{F}_+(\mathcal{X}_+)$) sont définis par (3.28) et (3.46a) (resp. (3.30) et (3.46b)), en interprétant les coordonnées de $\mathcal{X}_{\pm}$ comme dans (3.48e). $\mathcal{G}_{\pm}$ est défini par (3.64). $\mathcal{X}_{in,-} \doteq P_{H_2}$, $\mathcal{X}_{in,+} \doteq P_{O_2}$ et $\alpha_{\pm}, \beta_{\pm}$ sont définis par (3.38) et $\eta_{\pm}(\mathcal{X}_{\pm})$ par (3.32b). Les vecteurs v et w sont définis respectivement dans (C.28) et (C.22).

Le modèle C est défini par les équations (3.65). Les équations (3.65a) et (3.65b) constituent les équations d'état. Les équations (3.65c), (3.65d) et (3.65e) sont les équations de sortie du modèle.

On obtient alors un système d'équations algébro-différentielles de dimension $2N + 9$ où N est le nombre de points de discrétisation à chaque couche diffuse de la pile. Le vecteur $\tilde{\mathcal{C}}_{\pm}$ permet de rendre compte de ce qui se passe dans les couches diffuses, tandis que les vecteurs $\mathcal{X}_{\pm}$ sont liés aux réactions électrochimiques dans les couches compactes.

3.5 Simulations numériques

Les modèles de demi-piles sont découplés, ce qui permet des simulations indépendantes obtenues en choisissant le courant I .

3.5.1 Réalisation des simulations

Afin de résoudre les équations algébro-différentielles du modèle 0D obtenu dans la section 3.4.3, nous avons utilisé le solveur IDA de SUNDIALS (SUite of Nonlinear and Differential/ALgebraic equation Solvers)³ développé à Lawrence Livermore National Laboratory. La méthode d'intégration utilisée est à pas multiple appelée *fixed-leading coefficient form of BDF* (Backward differentiation formula). Le tracé et l'exploitation des résultats de calcul (profils de concentration protonique, courbes de polarisation, etc) a été réalisé sous Matlab.

L'algorithme que nous avons développé a comme entrées :

$I(t)$	le courant de la pile fixé par une charge extérieure
P_a	la pression totale à l'anode
P_c	la pression totale à la cathode
T	la température

et il permet de calculer les quantités suivantes :

3. <https://computation.llnl.gov/casc/sundials/main.html>

$\eta_{\pm}(t)$	les surtensions de Frumkin anodique et cathodique
$U_{cell}(t)$	le potentiel de cellule
$i_{H^+,M}(t)$	le courant dans la membrane
$\sigma_{\pm}(t)$	les densités de charge électronique anodique et cathodique
$C_{H_2}(t)$	la concentration d'hydrogène à l'interface anodique
$C_{O_2}(t)$	la concentration d'oxygène à l'interface cathodique
$\mathcal{C}(t)$	les concentrations protoniques
$\theta_{Hs}(t)$	taux de recouvrement par l'intermédiaire Hs (anode)
$\theta_{OHs}(t)$	taux de recouvrement par l'intermédiaire OHs (cathode)
$\theta_{O_2Hs}(t)$	taux de recouvrement par l'intermédiaire O_2Hs (cathode)
$\vec{\theta}_w(t), \overleftarrow{\theta}_w(t)$	taux de recouvrement par les dipôles
$\theta_{s,\pm}(t)$	proportion de sites libres anodiques ou cathodiques

L'état du système est de dimension $2N + 9$ où N est le nombre de points de discrétisation à chaque électrode ($N + 4$ à l'anode, et $N + 5$ à la cathode). Pour obtenir les valeurs d'équilibre, une simulation est effectuée jusqu'à la convergence vers les valeurs asymptotiques. La simulation est très rapide (moins d'une seconde de calcul sur un PC CORE i7, 8 Go de RAM). Les simulations sont effectuées pour plusieurs valeurs du courant I , les valeurs initiales choisies sont les valeurs asymptotiques obtenues pour la précédente valeur du courant. On obtient d'ailleurs un comportement différent selon que I est balayé en croissant ou en décroissant, cf. plus bas le paragraphe 3.5.2.

La durée des simulations est extrêmement courte. Avec un PC CORE i7, 8 Go de RAM, le calcul des courbes de polarisation avec une quarantaine de valeurs différentes du courant I dans le plage $0 - 10A$ nécessite moins d'une minute avec $N = 10$.

3.5.2 Surtension et courbes de polarisation

Sur les figures suivantes, nous montrons les surtensions anodique et cathodique $\eta_{\pm}$, ainsi que le potentiel de cellule U_{cell} en fonction du courant stationnaire I qui varie entre 0 et 10 A. Les courbes de polarisation obtenues présentent un effet d'hystérésis lorsque le sens de balayage du courant croît ou décroît. Par ailleurs, sur la Figure 3.6 on remarque que l'on obtient des tensions négatives, correspondant à un fonctionnement en électrolyse. Ces tensions apparaissent pour des densité de courant élevées, et en fait trop élevées pour être observées expérimentalement.

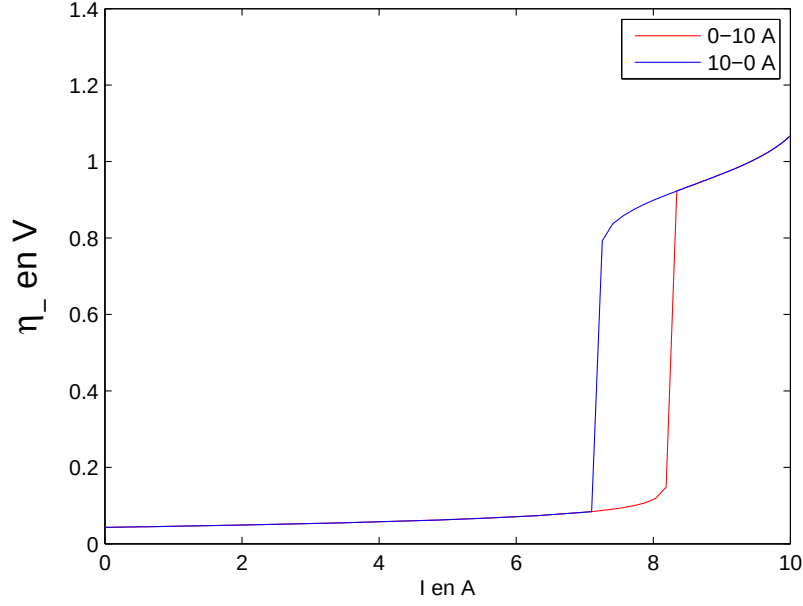


FIGURE 3.5 – Surtension anodique en fonction du courant nominal

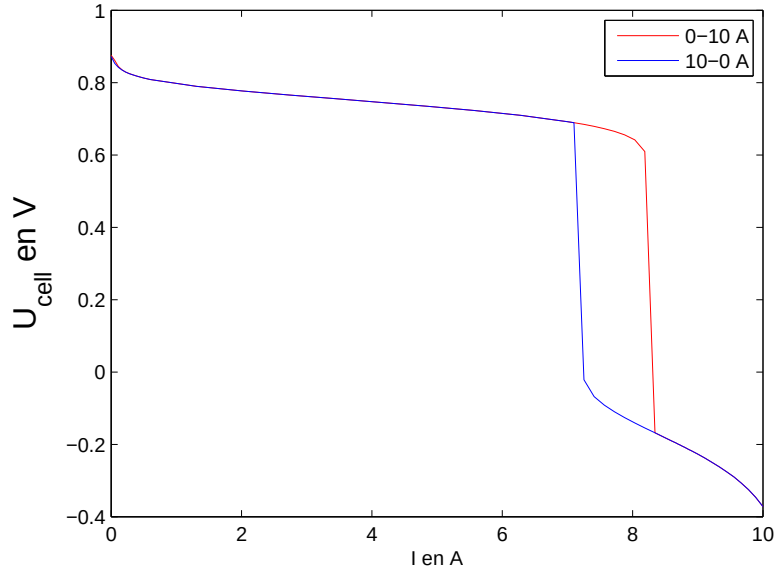


FIGURE 3.6 – Courbe de polarisation de la pile

3.5.3 Profils de concentration protonique et de potentiel électrique

La Figure 3.7 montre les profils de concentration protonique, à travers la pile, en fonction du courant. On remarque un surplus de protons à l'anode et un déficit à la cathode, par rapport à la valeur C_{FIX} caractérisant l'électroneutralité du milieu. Dans tous les profils représentés ici, la membrane, beaucoup plus étendue que les couche diffuses, n'est pas représentée à l'échelle.

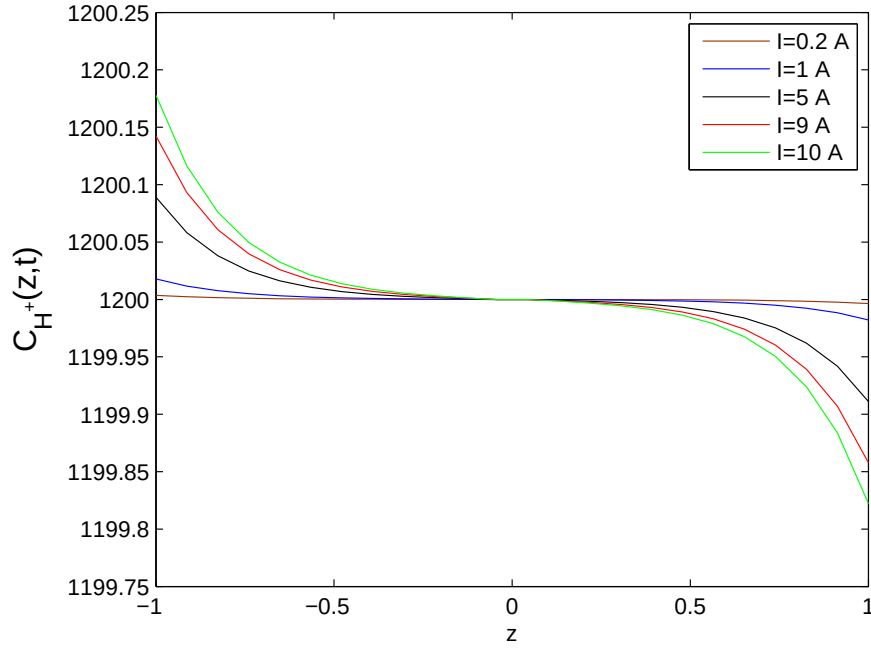


FIGURE 3.7 – Profils de la concentration protonique en fonction du courant I , pour $N = 10$.

La Figure 3.8 montre le profil de potentiel électrique à travers la pile.

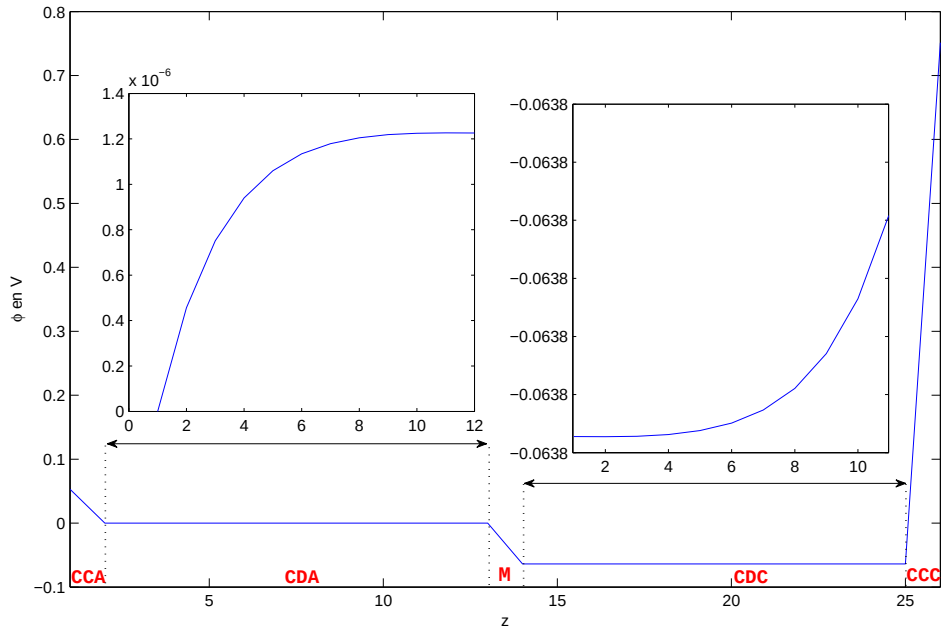


FIGURE 3.8 – Profil de potentiel électrique à travers la pile, pour $I = 3A$. CCA (resp. CCC) : couche compacte à l'anode (resp. à la cathode) ; CDA (resp. CDC) : couche diffuse à l'anode (resp. à la cathode) ; M : membrane.

3.5.4 Concentrations des gaz, taux d'occupation des sites catalytiques

Les Figures 3.9 et 3.10 montrent les concentrations d'hydrogène à l'interface anodique et de l'oxygène à l'interface cathodique en fonction du courant I . Comme obtenu par Franco ([37] page 132), les concentrations des gaz diminuent de façon linéaire avec l'augmentation du courant. Les Figures 3.11 et 3.12 montrent le taux de recouvrement en fonction du courant. La Figure 3.13 montre les densités surfaciques de charge électronique σ_- et σ_+ en fonction du courant.

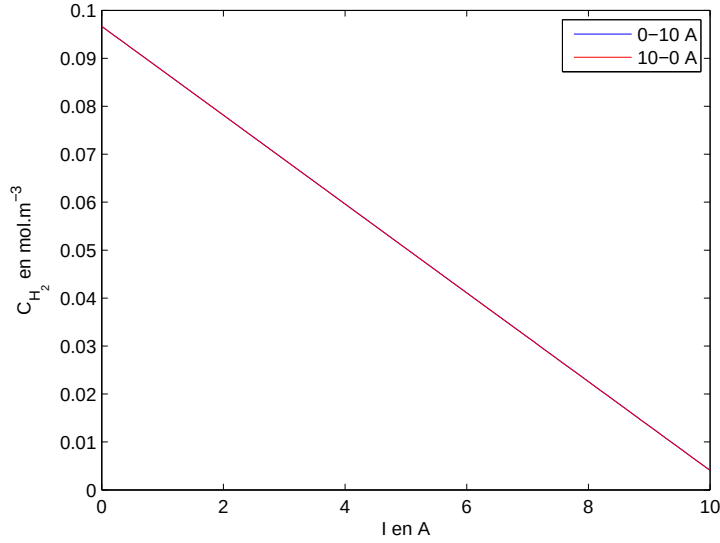


FIGURE 3.9 – Concentration d'hydrogène à l'interface anodique en fonction du courant

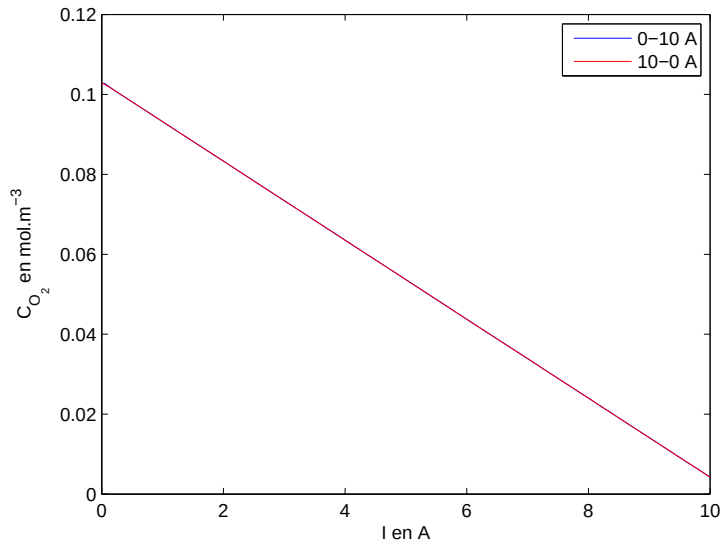


FIGURE 3.10 – Concentration d'oxygène à l'interface cathodique en fonction du courant

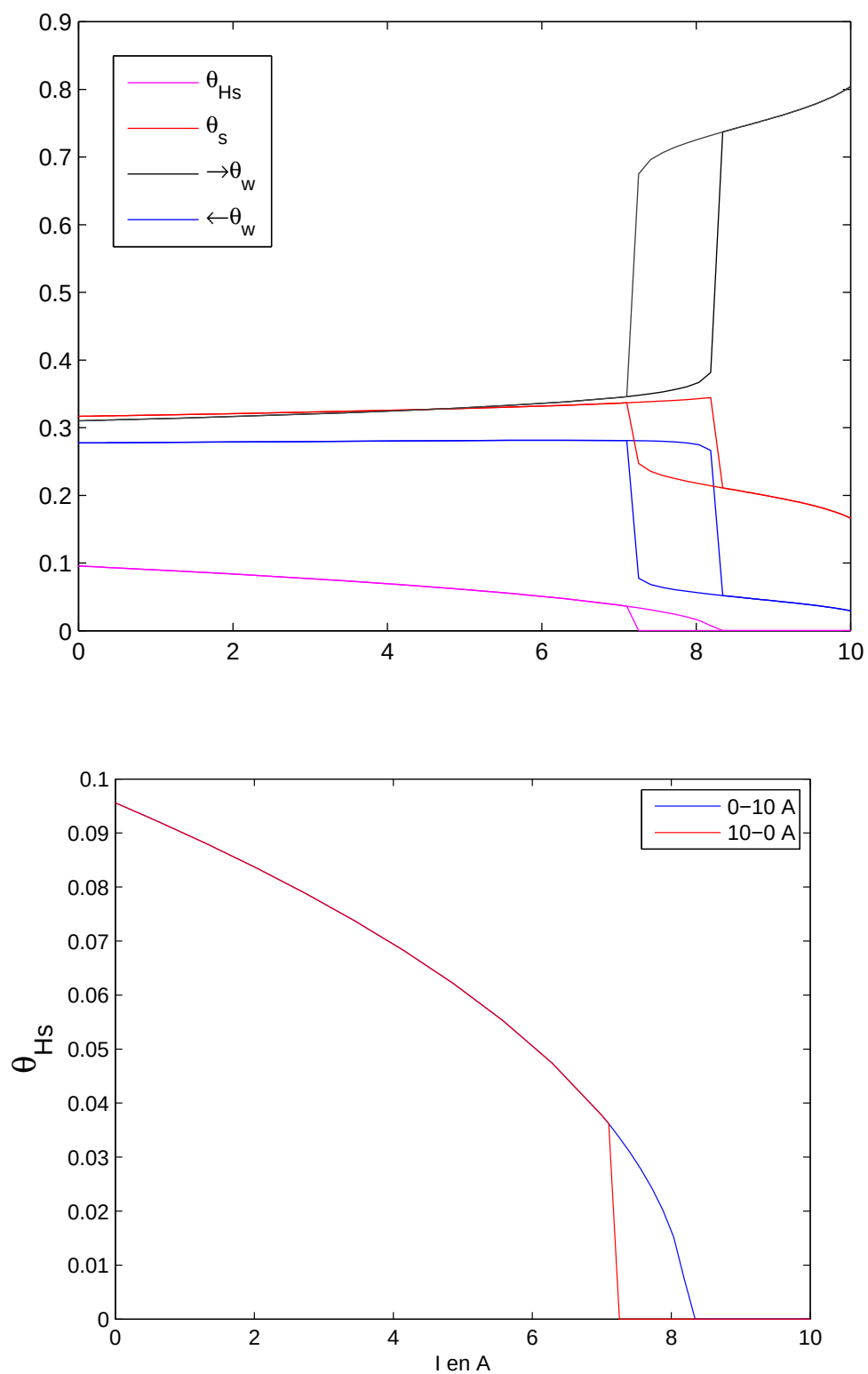


FIGURE 3.11 – Taux de recouvrement à l'anode en fonction du courant

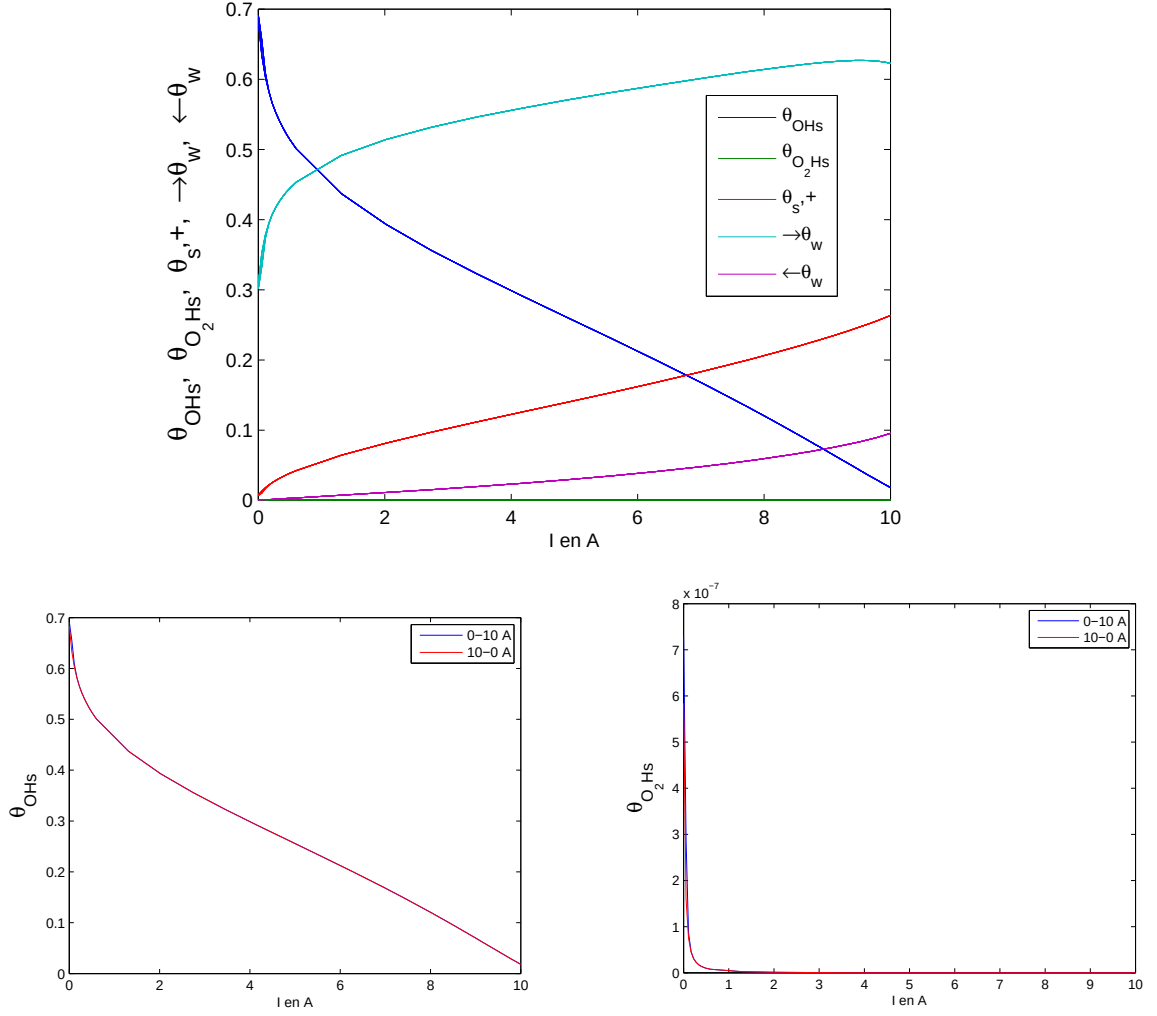


FIGURE 3.12 – Taux de recouvrement à la cathode en fonction du courant. θ_{O_2Hs} prend des valeurs inférieures à 10^{-6} , invisibles sur la figure.

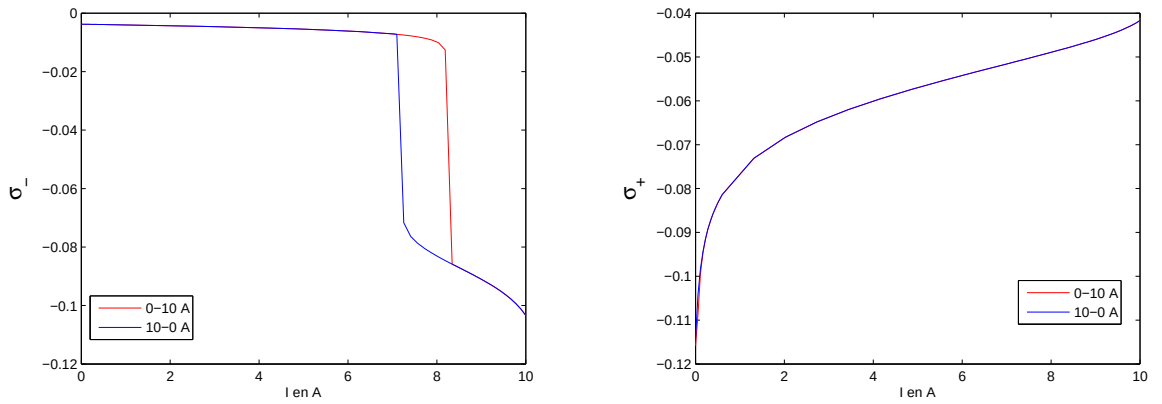


FIGURE 3.13 – Densité de charge électronique à l'anode (à gauche) et à la cathode (à droite).

3.5.5 Sensibilité aux vieillissement de la couche catalytique

Les courbes de polarisation varient aussi en fonction de la surface spécifique γ . Sur les Figures 3.14 et 3.15 nous montrons les courbes de polarisation, ainsi que les surtensions

anodiques et cathodiques pour différentes valeurs de surface spécifique γ et e_{NI} considérées par Franco.

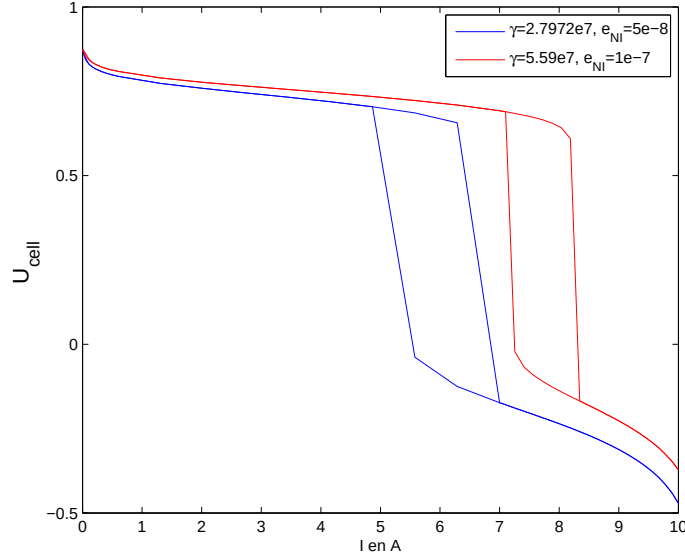


FIGURE 3.14 – Courbes de polarisation

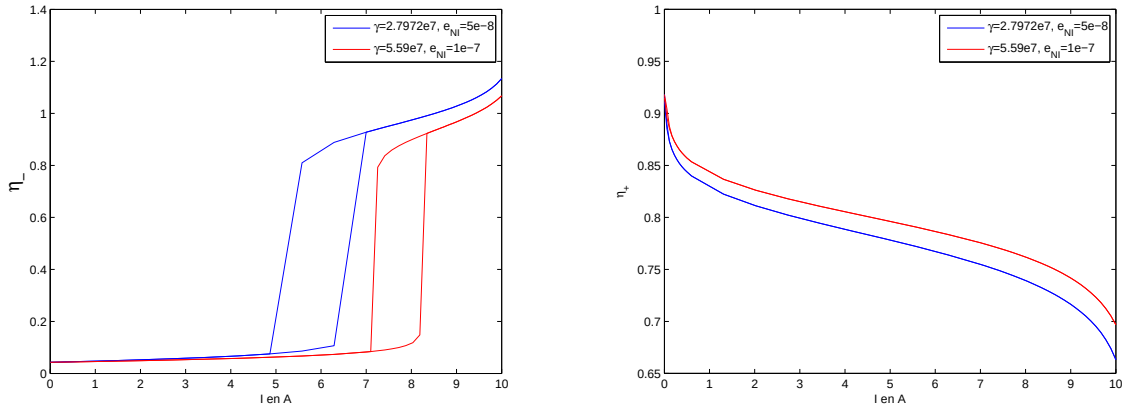


FIGURE 3.15 – Surtensions anodique (à gauche) et cathodique (à droite).

3.6 Valeurs numériques utilisées dans les simulations

Les valeurs numériques utilisées pour faire les simulations dans la section précédente sont tirées de Franco [37]. Elles sont données dans le tableau suivant :

Paramètre	valeur	unité
γ	5.59×10^7	m^{-1}
e_{CA}	15×10^{-6}	m
S_{EME}	2.1×10^{-4}	m^2
T	353	K
R	8.314	$J.mol^{-1}.K^{-1}$
k_B	1.38×10^{-23}	$J.K^{-1}$
F	96485	$C.m^{-1}$
ε_0	8.8541×10^{-12}	$C.V^{-1}.m^{-1}$
ε_{cc}	$6 \times \varepsilon_0$	$C.V^{-1}.m^{-1}$
ε_{cd}	$20 \times \varepsilon_0$	$C.V^{-1}.m^{-1}$
N_A	6.022×10^{23}	mol^{-1}
d	2×10^{-10}	m
λ	14	sans unité
a_w	$1.435 + 0.0022\lambda - \frac{2.75}{\lambda} - 0.13 \ln \lambda$	sans unité
μ	0.62×10^{-29}	$C.m$
ΔG^0	1	$J.mol^{-1}$
K	$2 \exp\left(-\frac{\Delta G^0}{RT}\right)$	sans unité
P_{H_2}	1.0293×10^5	Pa
P_{O_2}	1.0293×10^5	Pa
D_{H_2}	$3.921 \times 10^{-14} \times \exp(0.025 \times T)$	$m^2.s^{-1}$
H_{H_2}	$1.09 \times 10^6 \times \exp(77/T)$	$Pa.mol^{-1}.m^3$
H_{O_2}	$5.08 \times 10^6 \times \exp(-498/T)$	$Pa.mol^{-1}.m^3$
A	$\frac{1.2}{2\pi\varepsilon_{cd}}$	$C^{-1}.V.m$
$n^{\max}$	$\frac{1}{d^2}$	m^{-2}
k_{HEY}	10^{-8}	$m.s^{-1}$
k_{-HEY}	10^{-6}	$mol.m^{-2}.s^{-1}$
k_{VOL}	10^{-2}	$mol.m^{-2}.s^{-1}$
k_{-VOL}	10^{-5}	$mol.m^{-2}.s^{-1}$
k_{TAF}	0.1	$m.s^{-1}$
k_{-TAF}	0.1	$mol.m^{-2}.s^{-1}$

Paramètre	valeur	unité
e_{NI}	10^{-7}	m
e_M	50×10^{-6}	m
L	10^{-9}	m
k_1	10^4	$m^4.mol^{-1}.s^{-1}$
k_{-1}	10^{-5}	$mol.m^{-2}.s^{-1}$
k_2	10^8	$mol.m^{-2}.s^{-1}$
k_{-2}	10^{-2}	$mol.m^{-2}.s^{-1}$
k_3	10^8	$m.s^{-1}$
k_{-3}	1	$mol.m^{-2}.s^{-1}$
α_1	0.8	sans unité
α_3	0.99	sans unité
D_{H^+}	$2.97 \times 10^{-11} - [1.33 \times (293 - T) - 0.001 \times (293 - T)^2] / (T - 168) \times T$	$m^2.s^{-1}$

3.7 Conclusion

Dans ce chapitre nous avons détaillé le modèle d'A. Franco de la pile [37], qui est de type modèle à pores cylindriques et nous l'avons simplifié de façon à aboutir à un système d'équations aux dérivées partielles en une unique dimension spatiale. Nous avons ensuite discrétisé en espace les EDPs issues de ce modèle et nous avons obtenu un système d'équations algébro-différentielles en 0D. Ce dernier permet de faire des simulations numériques relativement rapides avec une vingtaine de points de discrétisation.

Par ailleurs, il est possible d'étendre le modèle 1D obtenu afin de tenir compte du transfert d'eau au cœur de la pile. Ce transfert d'eau est produit par deux mécanismes (cf. Section 2.3) : la diffusion et l'électro-osmose qui est liée au mouvement des protons dans la membrane. L'extension du modèle afin de tenir compte de ces deux phénomènes est donc possible moyennant une variable d'état représentant la concentration en eau en chaque point du segment joignant les deux couches compactes.

Deuxième partie

Modélisation des réseaux électrochimiques. Applications à la modélisation harmonique des piles à combustible PEM

Chapitre 4

Réseaux électriques et chimiques. Modèle de pile

Dans la première section de ce chapitre, nous présentons tout d'abord les 2-port qui sont des composantes élémentaires utilisées dans les circuits électriques et nécessaires pour la modélisation quadripolaire de la pile, comme nous le verrons dans le Chapitre 5. Nous étudions ensuite leurs représentations matricielles ainsi que leurs propriétés de passivité. La deuxième section présente un point de vue géométrique sur les mélanges, les réactions et les réacteurs ouverts, et fournit un modèle de réseaux électrochimiques compatible avec la thermodynamique. Comme application, nous utilisons ensuite le mécanisme réactionnel de Tafel-Heyrovsky-Volmer pour modéliser l'oxydation de l'hydrogène sur les sites catalytiques à l'anode de la pile. Côté cathode, nous utilisons le mécanisme réactionnel de Damjanovic et Brusica pour modéliser la réduction de l'oxygène. Nous présentons ensuite un modèle complet de la pile dans le cadre d'une réaction élémentaire à l'anode et à la cathode. Ce modèle prend en compte d'une manière précise l'effet de la double couche électrique.

4.1 Réseaux électriques

4.1.1 Réseaux de dipôles ; multipôle ; N -ports

Un réseau est un dispositif physique consistant en un assemblage d'éléments dont les pôles sont reliés à des nœuds [92]. Les éléments de réseau sont connectés entre eux selon un certain schéma. Un réseau a donc une structure topologique de graphe, que nous supposons orienté. À chaque arête du graphe est associé un courant I , une différence de potentiel V entre les deux sommets de l'arête et une loi de comportement. Le courant et la tension sont des fonctions scalaires, que l'on peut considérer comme dépendant du temps ou de la fréquence. La loi de comportement permet pour une valeur donnée du courant traversant l'arête, de calculer la différence de potentiel correspondante et inversement.

L'hypothèse fondamentale de la théorie des circuits [92] repose sur le fait que les tensions resp. les courants obéissent à la loi de Kirchhoff des mailles, resp. des nœuds. La loi de Kirchhoff des mailles (KVL) stipule que la somme algébrique des différences de potentiel le long d'une maille est nulle. La loi de Kirchhoff aux nœuds (KCL) stipule que la somme algébrique des intensités des courants arrivant à un nœud est nulle.

Multipôle : Un multipôle est un réseau avec un nombre fini des pôles. On y adjoint par convention un pôle correspondant à la terre. Dans la suite on respecte la convention de signe suivante : le sens positif du courant est le sens entrant dans un pôle.

Multipôle bien posé : Pour un multipôle on peut écrire l'équation déduite de KVL pour chaque boucle et celle déduite de KCL pour chaque nœud. L'ensemble de ces équations avec les lois de comportement sur les pôles constituent les équations du multipôle. On appelle multipôle bien posé un multipôle pour lequel la connaissance de la moitié des variables de tension et de courant aux pôles permet de déduire de façon unique la valeur de toutes les autres variables.

Dipôles et quadripôles : Parmi les multipôles, on peut en particulier distinguer les dipôles et les quadripôles, pourvus respectivement de deux et quatre pôles. Le couplage série ou parallèle de deux dipôles est bien évidemment un dipôle. La Figure 4.1 représente un quadripôle Q .

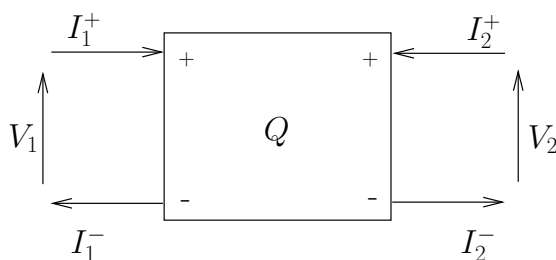


FIGURE 4.1 – Représentation d'un quadripôle

On supposera que le dipôle a une propriété particulière : l'intensité entrant à un pôle est égale à l'intensité sortant par l'autre pôle. Le dipôle est le composant élémentaire de tout réseau : tout réseaux est un assemblage de dipôles et le comportement des pôles externes de certains réseaux peut s'agréger sous la forme de dipôles équivalents.

Cette propriété est caractéristique d'une classe plus large de multipôles : la classe des n -ports.

N-ports : On appelle n -ports les multipôles qui ont la particularité d'avoir un nombre pair $2n$ de pôles, appariés entre eux en n ports respectant la propriété suivante : le courant quittant un pôle est égal au courant rentrant dans l'autre pôle du même port. On peut ainsi associer à tout n -port le vecteur des n différences de potentiel V_i entre chaque pôle du port i , et celui des n courants orientés I_i associés à chacun des ports. Ainsi, tout dipôle est également un 1-port.

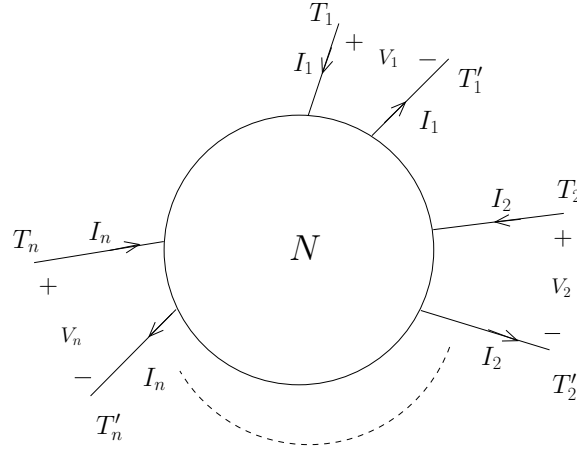


FIGURE 4.2 – Représentation d'un n -port

Remarque 4.3. On peut associer à chacun des deux pôles des ports d'un n -port un unique courant et une différence de potentiel. On peut donc lui attribuer une puissance, obtenue comme produit de ces deux quantités, et analyser les transferts d'énergie dans un réseau. C'est l'approche choisie par la méthode dites du Bond Graph. Celle-ci consiste à décomposer un système physique en sous-systèmes pour obtenir un réseau d'échange d'énergie. Les éléments composant ce réseau sont liés par des liens de puissance. A chaque lien sont associés deux types de variables, équivalents à un courant et un potentiel, dont le produit est une puissance. Nous n'utilisons pas le Bond Graph dans la modélisation de la pile parce que la connexion, du point de vue électrique, entre les couches actives (anodique et cathodique) et la membrane sera rompue par la connexion avec le circuit extérieur (voir Figure 4.5). D'où la nécessité du formalisme des réseaux électriques, introduite ci-dessus, qui permet rompre les connexions (par un prélèvement de courant par exemple) tout en respectant les lois de Kirchhoff.

4.1.2 Représentations matricielles des 2-ports

Représentation d'impédance : Elle permet d'exprimer les tensions (entrée et sortie) en fonction des courants (entrée et sortie). Sous forme matricielle, on peut écrire :

$$\begin{pmatrix} V_1 \\ V_2 \end{pmatrix} = Z(Q) \begin{pmatrix} I_1 \\ I_2 \end{pmatrix} \quad (4.1)$$

où $Z(Q)$ est par définition la matrice d'impédance du 2-ports.

Représentation de transmission : Elle permet d'exprimer les grandeurs d'entrée (tension et courant) en fonction des grandeurs de sortie (tension et courant). Sous forme matricielle, on peut écrire :

$$\begin{pmatrix} V_1 \\ I_1 \end{pmatrix} = T(Q) \begin{pmatrix} V_2 \\ -I_2 \end{pmatrix} \quad (4.2)$$

où $T(Q)$ est par définition la matrice de transmission du 2-ports Q .

Transformateur idéal : Un transformateur idéal est un 2-port particulier permettant de modifier les valeurs de tension et d'intensité du courant délivrées par une source d'énergie

électrique, en un système de tension et de courant de valeurs différentes, mais de même fréquence et de même forme sans aucune perte. Pour un transformateur idéal 1 : r on a la relation suivante :

$$\begin{pmatrix} V_2 \\ I_1 \end{pmatrix} = \begin{pmatrix} 0 & r \\ -\frac{1}{r} & 0 \end{pmatrix} \begin{pmatrix} I_2 \\ V_1 \end{pmatrix}$$

où V_k , resp. I_k , est la tension, resp. le courant, sur le k -ème port pour $k = 1, 2$. Comme on néglige les pertes, la puissance est transmise intégralement, on a :

$$V_2 I_2 = -r V_1 \frac{I_1}{r} = -V_1 I_1.$$

Passage d'une représentation d'impédance vers une représentation de transmission et inversement.

Soit $Z = \begin{pmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{pmatrix}$ la matrice d'impédance d'un 2-ports et $T = \begin{pmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \end{pmatrix}$ sa matrice de transmission. On a d'après [20] les relations suivantes :

$$T_{11} = \frac{Z_{11}}{Z_{21}}, \quad T_{12} = \frac{Z_{11}Z_{22} - Z_{12}Z_{21}}{Z_{21}}, \quad T_{21} = \frac{1}{Z_{21}}, \quad T_{22} = \frac{Z_{22}}{Z_{21}}. \quad (4.3)$$

et

$$Z_{11} = \frac{T_{11}}{T_{21}}, \quad Z_{12} = \frac{T_{11}T_{22} - T_{12}T_{21}}{T_{21}}, \quad Z_{21} = \frac{1}{T_{21}}, \quad Z_{22} = \frac{T_{22}}{T_{21}}. \quad (4.4)$$

Réseaux 2-ports réciproques. Un réseau 2-ports est dit réciproque s'il est possible d'invertir les quantités des ports 1 et 2 : $\begin{pmatrix} V_1 \\ V_2 \end{pmatrix} = \begin{pmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{pmatrix} \begin{pmatrix} I_1 \\ I_2 \end{pmatrix}$ et $\begin{pmatrix} V_2 \\ V_1 \end{pmatrix} = \begin{pmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{pmatrix} \begin{pmatrix} I_2 \\ I_1 \end{pmatrix}$. En d'autres termes le réseau est réciproque si et seulement si :

$$Z_{11} = Z_{22}, \quad Z_{12} = Z_{21}.$$

De la même manière, en utilisant (4.4), le réseau écrit sous forme $T = \begin{pmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \end{pmatrix}$ est réciproque si et seulement si :

$$T_{11} = T_{22}, \quad T_{11}T_{22} - T_{12}T_{21} = 1.$$

Impédance caractéristique du 2-ports. La représentation transmission d'un 2-ports s'écrit sous la forme :

$$\begin{pmatrix} V_1 \\ I_1 \end{pmatrix} = \begin{pmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \end{pmatrix} \begin{pmatrix} V_2 \\ -I_2 \end{pmatrix}. \quad (4.5)$$

Soit Z_c l'impédance caractéristique qui vérifie par définition :

$$Z_c = \frac{V_1}{I_1} = -\frac{V_2}{I_2}$$

on a alors :

$$\begin{pmatrix} V_1 \\ I_1 \end{pmatrix} = - \begin{pmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \end{pmatrix} \begin{pmatrix} Z_c \\ 1 \end{pmatrix} I_2$$

d'où

$$Z_c = \frac{V_1}{I_1} = \frac{T_{11}Z_c + T_{12}}{T_{21}Z_c + T_{22}}$$

on en déduit que :

$$T_{21}Z_c^2 + (T_{22} - T_{11})Z_c - T_{12} = 0.$$

Si T est réciproque on a $T_{22} = T_{11}$, d'où

$$Z_c = \sqrt{\frac{T_{12}}{T_{21}}}. \quad (4.6)$$

On en déduit d'après (4.3) :

$$Z_c = \sqrt{Z_{11}^2 - Z_{12}^2}. \quad (4.7)$$

4.1.3 Couplage de 2-ports en cascade

Il existe plusieurs façon de coupler des quadripôles. Le couplage en cascade de deux quadripôles Q_1 et Q_2 , qui nous intéressera particulièrement est noté par le quadripôle $Q_1; Q_2$. Il est défini par la Figure 4.3

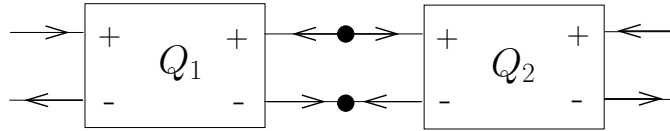


FIGURE 4.3 – Deux quadripôles Q_1 et Q_2 en cascade.

Definition 4.1. Pour tout couple de dipôles (Q_1, Q_2) , on appelle couplage en cascade de Q_1 et Q_2 le 2-ports $Q_1; Q_2$ qui vérifie :

$$T(Q_1; Q_2) = T(Q_1) \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} T(Q_2). \quad (4.8)$$

Cette définition est consistante avec l'application des lois de Kirchhoff aux nœuds et de (4.2). En effet, en notant $\begin{pmatrix} V \\ I \end{pmatrix}_{i,1}$, resp. $\begin{pmatrix} V \\ I \end{pmatrix}_{i,2}$, les grandeurs d'entrée, resp. de sortie, du quadripôle $i, i = 1, 2$, on a

$$\begin{pmatrix} V \\ I \end{pmatrix}_{i,1} = T(Q_i) \begin{pmatrix} V \\ -I \end{pmatrix}_{i,2}$$

Par ailleurs, les lois de Kirchhoff indiquent que :

$$V_{1,2} = V_{2,1}, \quad I_{1,2} = -I_{2,1}$$

ce qui fournit (4.8).

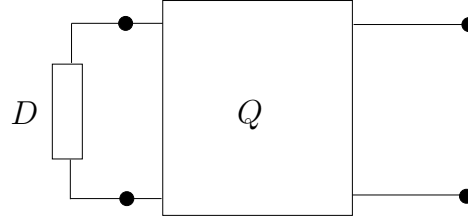


FIGURE 4.4 – Dipôle D branché à un quadripôle Q

4.1.4 Branchements de dipôle et de quadripôle

On appelle dipôle image le résultat du branchement d'un dipôle D sur un quadripôle Q . Le dipôle image est bien un dipôle.

Dans le cadre de la décomposition quadripolaire de la pile, que nous traitons dans le Chapitre 4, la pile sera décomposée sous une forme schématisée par la Figure 4.5. On peut s'apercevoir à travers ce schéma de l'importance du formalisme développé dans la Section 4.1.1.

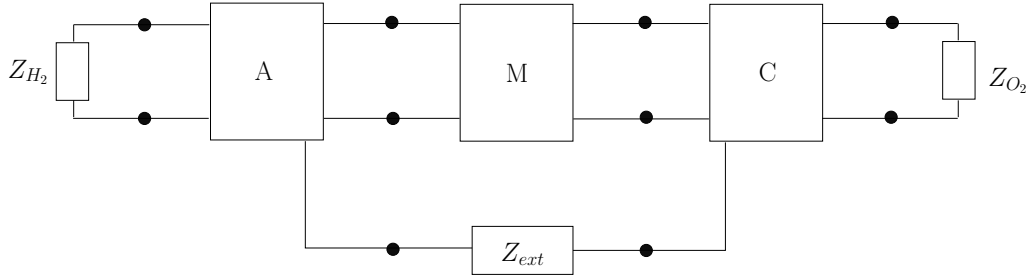


FIGURE 4.5 – Assemblage électrique de la pile. Les 3-ports A et C représentent respectivement l'anode et la cathode, le 2-ports M représente la membrane. Le dipôle Z_{H_2} , resp. Z_{O_2} représente la source d'hydrogène, resp. d'oxygène à la cathode. Z_{ext} est la charge qui relie les deux électrodes de la pile par le circuit extérieur.

4.1.5 Dipôles et 2-ports passif

Un dipôle traversé par un courant d'intensité I et dont la tension aux bornes est V met en jeu une puissance P telle que $P = VI$. Les dipôles passifs ne peuvent que consommer de la puissance électrique (cette puissance est dissipée par effet Joule).

Definition 4.2. Une représentation sous forme d'impédance $Z(j\omega)$ est dite passive si $Z(j\omega) + Z^*(j\omega) \geq 0$ est semi-défini positif pour tout $\omega \in \mathbb{R}$ ou de façon identique :

$$\forall I, \quad \text{Re}(I^*ZI) \geq 0 \quad (4.9)$$

Ici et désormais,

- $*$ dénote la transconjuguée
- Re s'applique également aux matrices carrées, avec la définition

$$\text{Re } Z = \frac{1}{2}(Z + Z^*).$$

Critère de passivité : D'après [21], la matrice $Z = \begin{pmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{pmatrix}$ est passive si et seulement si :

$$\operatorname{Re} Z_{11} \geq 0 \quad (4.10a)$$

$$\operatorname{Re} Z_{22} \geq 0 \quad (4.10b)$$

$$4 \operatorname{Re} Z_{11} \operatorname{Re} Z_{22} \geq |Z_{12} + Z_{21}^*|^2 \quad (4.10c)$$

En termes de T , cette propriété s'applique de façon équivalente :

$$\operatorname{Re} \frac{T_{11}}{T_{21}} \geq 0 \quad (4.11a)$$

$$\operatorname{Re} \frac{T_{22}}{T_{21}} \geq 0 \quad (4.11b)$$

$$4 \operatorname{Re} \frac{T_{11}}{T_{21}} \operatorname{Re} \frac{T_{22}}{T_{21}} \geq \left| \frac{T_{11}T_{22} - T_{12}T_{21}}{T_{21}} + \frac{1}{T_{21}^*} \right|^2 \quad (4.11c)$$

Passivité d'un réseau 2-ports réciproque : Un réseau réciproque est passif si et seulement si :

$$\operatorname{Re} Z_{11} \geq |\operatorname{Re} Z_{12}| \quad (4.12a)$$

En termes de T , on a :

$$\operatorname{Re} \left(\frac{T_{11}}{T_{21}} \right) \geq \left| \operatorname{Re} \frac{1}{T_{21}} \right| \quad (4.13a)$$

On peut citer ici des dipôles classiques dont on connaît le comportement vis-à-vis de la passivité : une résistance ; un condensateur et un transformateur, qui sont conservatifs pour toute fréquence (on peut donner la définition).

4.1.6 Dipôle générateur et image par un 2-ports passif

On appelle dipôle générateur tout dipôle qui se décompose lui-même comme la somme d'un dipôle passif et d'une source de tension contrôlée. Si on branche un dipôle générateur G sur un port du 2-ports Q , comme le montre la Figure 4.6, on a alors :

$$\begin{pmatrix} V_1 \\ V_2 \end{pmatrix} = Z(Q) \begin{pmatrix} I_1 \\ I_2 \end{pmatrix} \quad (4.14a)$$

$$V_1 = V_G - Z_G I_1 \quad (4.14b)$$

où

$$Z(Q) \doteq \begin{pmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{pmatrix}$$

On en déduit que :

$$V_G = (Z_{11} + Z_G)I_1 + Z_{12}I_2 \quad (4.15a)$$

$$V_2 = Z_{21}I_1 + Z_{22}I_2 \quad (4.15b)$$

d'où

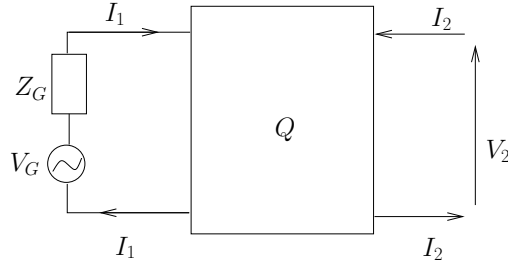


FIGURE 4.6 – Représentation d'un 2-ports branché sur un générateur

$$I_1 = \frac{\begin{vmatrix} V_G & Z_{12} \\ V_2 & Z_{22} \end{vmatrix}}{(Z_{11} + Z_G)Z_{22} - Z_{12}Z_{21}} = \frac{V_G Z_{22} - Z_{12} V_2}{(Z_{11} + Z_G)Z_{22} - Z_{12}Z_{21}}$$

et

$$I_2 = \frac{\begin{vmatrix} Z_{11} + Z_G & V_G \\ Z_{21} & V_2 \end{vmatrix}}{(Z_{11} + Z_G)Z_{22} - Z_{12}Z_{21}} = \frac{Z_{11} + Z_G}{(Z_{11} + Z_G)Z_{22} - Z_{12}Z_{21}} V_2 - \frac{Z_{21}}{(Z_{11} + Z_G)Z_{22} - Z_{12}Z_{21}} V_G$$

On a alors d'après l'équation précédente :

$$V_2 = Z(G; Q)I_2 + \frac{Z_{21}}{Z_{11} + Z_G} V_G \quad (4.16a)$$

où

$$Z(G; Q) \doteq Z_{22} - \frac{Z_{12}Z_{21}}{Z_{11} + Z_G} \quad (4.16b)$$

est l'image par le 2-ports Q du dipôle passif Z_G . Dans le cas où le 2-ports est réciproque, on a :

$$Z(G; Q) = Z_{11} - \frac{Z_{12}^2}{Z_{11} + Z_G}. \quad (4.17)$$

Le terme $\frac{Z_{21}}{Z_{11} + Z_G} V_G$ dans le membre de droite de (4.16a) représente l'image par le 2-ports Q du terme V_G .

Théorème 4.3. *L'image d'un générateur par un 2-ports passif est un générateur.*

Démonstration. Pour un 2-ports dont l'impédance s'écrit sous la forme :

$$Z = \begin{pmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{pmatrix}$$

nous avons montré précédemment que si on branche une impédance passive Z_d sur un port, l'impédance de dipôle résultant Z' s'écrit sous la forme :

$$Z' \doteq Z_{22} - \frac{Z_{12}Z_{21}}{Z_{11} + Z_d}. \quad (4.18)$$

Pour Z_d donné, Z' est dissipatif si et seulement si

$$\operatorname{Re} Z_{22} - \operatorname{Re} \left(\frac{Z_{12}Z_{21}}{Z_{11} + Z_d} \right) \geq 0, \quad (4.19)$$

On a sous la condition que $\operatorname{Re} Z_{22} > 0$:

$$\begin{aligned} (4.19) &\iff \operatorname{Re} Z_{22} \geq \frac{1}{2} \left(\frac{Z_{12}Z_{21}}{Z_{11} + Z_d} + \frac{Z_{12}^*Z_{21}^*}{Z_{11}^* + Z_d^*} \right) \\ &\iff |Z_{11} + Z_d|^2 \geq \frac{1}{2 \operatorname{Re} Z_{22}} \left(Z_{12}Z_{21}(Z_{11}^* + Z_d^*) + Z_{12}^*Z_{21}^*(Z_{11} + Z_d) \right) \\ &\iff \left(Z_{11} + Z_d - \frac{1}{2 \operatorname{Re} Z_{22}} Z_{12}Z_{21} \right) \left(Z_{11} + Z_d - \frac{1}{2 \operatorname{Re} Z_{22}} Z_{12}Z_{21} \right)^* \geq \frac{|Z_{12}|^2 |Z_{21}|^2}{4(\operatorname{Re} Z_{22})^2} \end{aligned} \quad (4.20)$$

On en déduit que Z' est dissipatif si et seulement si Z_d se trouve à l'extérieur d'un cercle $\mathcal{C}$ de centre $-Z_{11} + \frac{1}{2 \operatorname{Re} Z_{22}} Z_{12}Z_{21}$ et de rayon $\frac{1}{2 \operatorname{Re} Z_{22}} |Z_{12}| |Z_{21}|$.

On veut maintenant garantir cette propriété pour tout Z_d dissipatif. Or, la partie réelle du point le plus à droite du cercle $\mathcal{C}$ précédemment introduit s'écrit :

$$-\operatorname{Re}(Z_{11}) + \frac{1}{2 \operatorname{Re} Z_{22}} \operatorname{Re}(Z_{12}Z_{21}) + \frac{1}{2 \operatorname{Re} Z_{22}} |Z_{12}| |Z_{21}| \quad (4.21)$$

Z' est donc dissipatif pour tout Z_d dissipatif si et seulement si cette dernière quantité est négative ou nulle.

En utilisant le fait que :

$$|Z_{12} + Z_{21}^*|^2 = |Z_{12}|^2 + |Z_{21}|^2 + 2 \operatorname{Re}(Z_{12}Z_{21}) \geq 2|Z_{12}| |Z_{21}| + 2 \operatorname{Re}(Z_{12}Z_{21}) \quad (4.22)$$

on obtient d'après (4.10c) la relation suivante :

$$\frac{1}{2} \left(\operatorname{Re}(Z_{12}Z_{21}) + |Z_{12}| |Z_{21}| \right) \leq \operatorname{Re} Z_{11} \operatorname{Re} Z_{22} \quad (4.23)$$

Pour $\operatorname{Re} Z_{22} = 0$, on a d'après (4.10c) : $Z_{12} + Z_{21}^* = 0$. D'où d'après (4.22)

$$\operatorname{Re}(Z_{12}Z_{21}) = -|Z_{12}| |Z_{21}|$$

Le membre de gauche et le membre de droite de (4.23) sont tous deux nuls, et donc égaux.

On en déduit que le point le plus à droite du cercle, et donc le cercle lui-même, se trouve alors dans le demi-plan gauche, et par suite : $\forall Z_d, \operatorname{Re} Z_d \geq 0$ on a $Z_d \notin \mathcal{C}$. $\square$

4.2 Réseaux électrochimiques et analogie électrique

Nous rappelons tout d'abord ici quelques éléments de chimie et de cinétique chimique (sections 4.2.1 et 4.2.3). Dans la section 4.2.2 nous étudions l'aspect thermodynamique d'un réacteur chimique ouvert et ses propriétés de dissipativité. Dans la section 4.2.4 nous traitons le cas d'un réseau de réactions électrochimiques. Un réseau chimique est un ensemble d'espèces chimiques interconnectées par des réactions. D'après Oster et al. [64], les réseaux chimiques sont mathématiquement équivalents à une classe de réseaux multiports constitués de condensateurs, de résistances et de transformateurs idéaux reliés entre eux d'une manière prescrite. Cet aspect est présenté à la section 4.2.5.

4.2.1 Point de vue géométrique sur les mélanges, les réactions et les réacteurs ouverts.

Notations. On définit pour tous vecteurs $x, z \in \mathbb{R}^N$, le vecteur $x \cdot z \in \mathbb{R}^N$, resp. le vecteur $\frac{x}{z}$ comme le produit, resp. la division élément par élément de x et z . On introduit les applications $\text{Ln} : \mathbb{R}_+^N \rightarrow \mathbb{R}^N$ et $\text{Exp} : \mathbb{R}_+^N \rightarrow \mathbb{R}^N$, $x \mapsto \text{Ln}(x)$ et $x \mapsto \text{Exp}(x)$ définies comme étant les applications dont l' i -ème composante est donnée par : $(\text{Ln}(x))_i = \ln(x_i)$ et $(\text{Exp}(x))_i = \exp(x_i)$. On note aussi $\mathbb{1} = (1, \dots, 1)^\top \in \mathbb{R}^N$. Pour un vecteur $u = (u_1, \dots, u_N)$, on note $u \geq 0$ si $u_k \geq 0$, $k = 1, \dots, N$.

Mélanges et réactions, aspects géométriques. On notera par $X_k, k = 1, \dots, N$ un système de N espèces chimiques. Pour la représentation, nous suivons l'approche géométrique d'Aris [1] intéressante pour la puissance des notations. On considère l'espace vectoriel des réactions $\mathcal{R} = \left\{ \sum_{k=1}^N \alpha_k X_k / \alpha_k \in \mathbb{R} \right\}$ et on note $\mathcal{R}^+ = \left\{ \sum_{k=1}^N \alpha_k X_k / \alpha_k \geq 0 \right\}$ le cône positif des mélanges. Un mélange s'écrit sous la forme :

$$X = \sum_{k=1}^N \alpha_k X_k, \quad X \in \mathcal{R}^+. \quad (4.24)$$

Inversement on supposera les espèces X_k indépendantes, au sens où aucune n'est un mélange des autres. Cela permet de définir le produit scalaire $\langle \rangle$:

$$\langle X_k, X_{k'} \rangle = \delta_{kk'}, \quad k, k' = 1, \dots, N. \quad (4.25)$$

avec $\delta_{kk'}$ est le symbole de Kronecker défini par : $\delta_{kk'} = 1$ si $k = k'$ et 0 sinon. En particulier, pour une réaction $X \in \mathcal{R}$ donnée par (4.24), on a $\alpha_k = \langle X, X_k \rangle$. On notera $n_k = \alpha_k$ le nombre de moles de X_k dans X si X est un mélange.

Les réactions élémentaires considérées plus loin sont représentées par des équations de droites dans $\mathcal{R}$:

$$X = X_0 + X_s \xi \quad (4.26)$$

où $\xi \in \mathbb{R}$ (abscisse sur la droite) est le degré d'avancement de la réaction et $X_s = \sum_{k=1}^N \nu_k X_k \in \mathcal{R}$. Les ν_k peuvent être de signe quelconque. On considère les décompositions (non uniques) de réactions en réactants et produits de la forme $X_s = X^P - X^R$ avec :

$$X^R = \sum_{k=1}^N \nu_k^R X_k \in \mathcal{R}^+, \quad X^P = \sum_{k=1}^N \nu_k^P X_k \in \mathcal{R}^+, \quad (4.27)$$

de sorte que la réaction s'écrit :

$$X = X_0 + (X^P - X^R)\xi, \quad X^P, X^R \in \mathcal{R}^+ \quad (4.28)$$

En notant

$$\nu_k = \nu_k^P - \nu_k^R, \quad k = 1, \dots, N \quad (4.29)$$

on a sous forme différentielle, avec (4.27)

$$dX = (X^P - X^R)d\xi = \sum_{k=1}^N \nu_k d\xi X_k \quad (4.30)$$

En particulier :

$$\frac{dX}{d\xi} = X^P - X^R = \sum_{k=1}^N \nu_k X_k \quad (4.31)$$

Remarquons que pour $X \in \mathcal{R}$ on a $\nu_k = \langle \frac{dX}{d\xi}, X_k \rangle$. En particulier on a :

$$d \langle X, X_k \rangle = \nu_k d\xi, \quad k = 1, \dots, N \quad (4.32)$$

Formellement, à l'équilibre on a $X^R = X^P$ ou encore $\sum_{k=1}^N (\nu_k^P - \nu_k^R) X_k = 0$, que l'on note aussi :

$$\sum_{k=1}^N \nu_k^R X_k \rightleftharpoons \sum_{k=1}^N \nu_k^P X_k \quad (4.33)$$

Systèmes de réactions. En général, un système de M réactions faisant intervenir N espèces chimiques est représenté par le schéma suivant où les ν_{kj}^R, ν_{kj}^P , $k = 1, \dots, N$, $j = 1, \dots, M$, sont positifs [88] :

$$\sum_{k=1}^N \nu_{kj}^R X_k \rightleftharpoons \sum_{k=1}^N \nu_{kj}^P X_k, \quad j = 1, \dots, M. \quad (4.34)$$

Utilisant les notations précédentes, on introduit la j -ième réaction X^j en fonction de $\xi_j \in \mathbb{R}$ son degré d'avancement, donnée par :

$$X^{P,j} = \sum_{k=1}^N \nu_{kj}^P X_k \in \mathcal{R}^+, \quad X^{R,j} = \sum_{k=1}^N \nu_{kj}^R X_k \in \mathcal{R}^+ \quad (4.35a)$$

$$\frac{dX^j}{d\xi_j} = X_s^j, \quad X_s^j = X^{P,j} - X^{R,j}, \quad j = 1, \dots, M \quad (4.35b)$$

où $X^{P,j}, X^{R,j} \in \mathcal{R}^+$ car ν_{kj}^R et ν_{kj}^P sont positifs. En particulier, en notant $\nu_{kj} = \nu_{kj}^P - \nu_{kj}^R$, $j = 1, \dots, M, k = 1, \dots, N$ avec (4.32), l'avancement infinitésimal de la j -ème réaction chimique de l'espèce X_k est donné par :

$$d \langle X^j, X_k \rangle = \nu_{kj} d\xi_j, \quad j = 1, \dots, M, \quad k = 1, \dots, N \quad (4.36)$$

Sous l'hypothèse actuelle de réacteur fermé, cela conduit à l'expression suivante de la variation du nombre de moles de l'espèce X_k due à l'ensemble des réactions chimiques :

$$d_i n_k = \sum_{j=1}^M \nu_{kj} d\xi_j, \quad k = 1, \dots, N \quad (4.37)$$

On a construit ici la forme différentielle $d_i n_k$ utilisée par Prigogine et al. dans [54] (Formule 4.1.4, page 107). Dans ce cas d'un système fermé, on a simplement $d_i n_k = dn_k$, ce que l'on peut intégrer en notant n_k^0 et ξ_j^0 les origines dans l'espace des n_k et sur les droites des réactions, on a :

$$n_k = n_k^0 + \sum_{j=1}^M \nu_{kj} (\xi_j - \xi_j^0), \quad (4.38)$$

On prendra par la suite $\xi_j^0 = 0$, $j = 1, \dots, M$. On définit par ailleurs les matrices

$$\nu = (\nu_{kj}) \in \mathbb{R}^{N \times M}, \quad \nu^R = (\nu_{kj}^R) \in \mathbb{R}^{N \times M}, \quad \nu^P = (\nu_{kj}^P) \in \mathbb{R}^{N \times M} \quad (4.39)$$

avec $\nu = \nu^R - \nu^P$ et les vecteurs

$$n = \begin{pmatrix} n_1 \\ \vdots \\ n_N \end{pmatrix} \in \mathbb{R}^N, \quad n^0 = \begin{pmatrix} n_1^0 \\ \vdots \\ n_N^0 \end{pmatrix} \in \mathbb{R}^N, \quad \xi = \begin{pmatrix} \xi_1 \\ \vdots \\ \xi_M \end{pmatrix} \in \mathbb{R}^M \quad (4.40)$$

On a alors avec (4.38) :

$$n = \nu \xi + n^0. \quad (4.41)$$

Système de réactions réduit. Soient X_s^j , $j = 1, \dots, M \in \mathcal{R}$ et M réactions de $\mathcal{R}$ associées comme dans (4.35b), notées $X^j = X_0^j + X_s^j \xi_j$. On supposera que chaque espèce de $\mathcal{R}$ participe au moins à une réaction, au sens où :

$$\forall k = 1, \dots, N, \exists j = 1, \dots, M : \nu_{kj} = \langle X_s^j, X_k \rangle > 0.$$

Autrement dit, la matrice ν n'a pas de lignes de zéros. La matrice ν n'est pas nécessairement de rang plein. En particulier les colonnes peuvent être liées. Il est alors intéressant de faire apparaître un système libre de réactions $\tilde{X}_s^1, \dots, \tilde{X}_s^{M_\nu}$ qui soit une base de l'espace engendré par $X_s^1, \dots, X_s^M$. Le rang M_ν de ce système est le rang du système d'origine. On a $M_\nu \leq N$ et $M_\nu \leq M$.

On peut réduire le système des réactions en utilisant des règles de stœchiométrie que l'on va préciser [35] avant de les illustrer avec le système réactionnel de THV. On introduit la matrice atomique Z des coordonnées des espèces dans la base des "atomes" dans laquelle on inclut q la charge positive élémentaire. Par définition la matrice $Z \in \mathbb{Z}^{A \times N}$, où A est le nombre d'atomes distincts, charge positive incluse, est telle que pour tout $i = 1, \dots, A$, $k = 1, \dots, N$, l'espèce X_k contient Z_{ik} atomes A_i .

On en déduit le lemme suivant :

Lemme 4.4. Notons $[X] = (X_1 \ X_2 \ \dots \ X_N)^T \in \mathbb{R}^N$ le vecteur des coordonnées de la réaction X . Le nombre de réactions indépendantes X ayant la propriété de conservation $Z[X^P] = Z[X^R]$ est égale à $\dim \ker(Z) = N - \text{rank}(Z)$.

L'interprétation de ce résultat est la suivante : il est possible d'écrire $\text{rank}(Z)$ équations de conservation d'espèces atomiques et le système réactionnel peut être représenté sous la forme réduite d'équations, paramétrées par $N - \text{rank}(Z)$ variables. On a donc $M_\nu = N - \text{rank}(Z)$.

D'un point de vue algébrique, il existe $\tilde{\nu} \in \mathbb{R}^{N \times M_\nu}$ et $\Pi \in \mathbb{R}^{M_\nu \times M}$ de rang M_ν telles que :

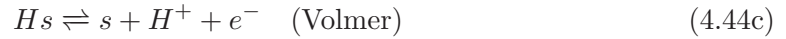
$$\nu = \tilde{\nu} \Pi \quad (4.42)$$

On peut alors écrire $\nu \xi = \tilde{\nu} \tilde{\xi}$ en définissant :

$$\tilde{\xi} = \Pi \xi \quad (4.43)$$

qui représente le vecteur des taux d'avancement d'un ensemble maximal de réactions indépendantes.

Exemple du mécanisme réactionnel de Tafel-Heyrovsky-Volmer. Pour le mécanisme réactionnel de Tafel-Heyrovsky-Volmer, nous introduisons, dans cette section et par la suite, l'indice a (comme anode) dans les matrices et les formules utilisées précédemment. De la même manière, nous introduirons plus bas l'indice c (comme cathode) pour le mécanisme réactionnel à la cathode. Le mécanisme réactionnel de Tafel-Heyrovsky-Volmer utilisé dans le modèle de Franco (Chapitre 3) et retenu par la suite s'écrit sous la forme [46] :



Nous sommes donc en présence de trois réactions écrites entre cinq constituants indépendants. L'espèce électronique joue un rôle particulier, qui nécessite de la considérer à part. On a $N = 4, M = 3$. On note par $n_k, k \in \{H_2, Hs, s, H^+\}$ le nombre de moles de l'espèce k et par $\xi_j, j \in \{T, H, V\}$ l'avancement de la j -ème réaction. On note :

$$n_a = \begin{pmatrix} n_{H_2} \\ n_{H^+} \\ n_{Hs} \\ n_s \end{pmatrix}, \quad \xi_a = \begin{pmatrix} \xi_T \\ \xi_H \\ \xi_V \end{pmatrix} \quad (4.45)$$

Les composantes n_{H^+} et n_{H_2} sont ici associées respectivement à un bilan de protons dans la couche diffuse, et un bilan d'hydrogène dans le *GDL*.

Le système (4.44) s'écrit sous la forme :

$$X_i = X_i^P - X_i^R, \quad i = 1, 2, 3 \quad (4.46)$$

avec

$$\begin{aligned} X_1^R &= H_2 + 2s & X_1^P &= 2Hs \\ X_2^R &= H_2 + s & X_2^P &= Hs + H^+ \\ X_3^R &= Hs & X_3^P &= s + H^+ \end{aligned} \quad (4.47)$$

La matrice stœchiométrique ν est la matrice des coordonnées des réactions X_1, X_2, X_3 dans la base H_2, H^+, Hs, s . Elle s'écrit :

$$\nu_a = \nu_{THV} = \begin{pmatrix} X_1 & X_2 & X_3 \\ -1 & -1 & 0 \\ 0 & 1 & 1 \\ 2 & 1 & -1 \\ -2 & -1 & 1 \end{pmatrix} \left| \begin{array}{l} H_2 \\ H^+ \\ Hs \\ s \end{array} \right. \quad (4.48)$$

On a aussi

$$\nu_{THV} = \nu_a^P - \nu_a^R \text{ avec } \nu_a^P = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 1 & 1 \\ 2 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}, \nu_a^R = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \\ 2 & 1 & 0 \end{pmatrix} \quad (4.49)$$

La matrice Z est donnée par

$$Z = \left(\begin{array}{cccc} H_2 & H^+ & Hs & s \\ \hline 2 & 1 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \end{array} \right) \left| \begin{array}{c} H \\ q \\ s \end{array} \right. \quad (4.50)$$

et on vérifie que $\text{rank}(Z) = 3$ et $\text{rank}(\nu) = 2 = N + 1 - \text{rank}(Z)$, ($N + 1$ est l'ancien N du Lemme 4.4).

L'électroneutralité permet de retrouver la stœchiométrie des électrons. En notant ν_{ej} , $j = 1, \dots, M$ les coefficients correspondants pour chacune des M réactions, on a :

$$\nu_{ej} = \sum_{k=1}^N \nu_{kj} z_k \quad (4.51)$$

où z_k est la charge (algébrique) des espèces chargées. Ici les protons constituent la seule espèce chargée, et on a donc : $\nu_{ej} = \nu_{2j}$, $j = 1, \dots, M$, H^+ étant attaché à la deuxième composante de n .

On peut réduire le système de réactions en considérant par exemple les 2 réactions indépendantes $\tilde{X}_1, \tilde{X}_2$ données par :

$$\tilde{X}_1 = -H_2 - 2s + 2Hs \quad (4.52a)$$

$$\tilde{X}_2 = -Hs + s + H^+ \quad (4.52b)$$

Ceci revient à prendre :

$$\Pi = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \end{pmatrix} \quad (4.53)$$

On vérifie que :

$$\tilde{\nu}_a \Pi = \nu_{THV} \quad (4.54)$$

où

$$\tilde{\nu}_a = \begin{pmatrix} -1 & 0 \\ 0 & 1 \\ 2 & -1 \\ -2 & 1 \end{pmatrix}, \quad \tilde{\nu}_a \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \tilde{X}_1, \quad \tilde{\nu}_a \begin{pmatrix} 0 \\ 1 \end{pmatrix} = \tilde{X}_2. \quad (4.55)$$

On a donc ici :

$$\tilde{\xi}_a = \Pi \begin{pmatrix} \xi_T \\ \xi_H \\ \xi_V \end{pmatrix} = \begin{pmatrix} \xi_{TH} \\ \xi_{HV} \end{pmatrix} \quad (4.56)$$

où on pose $\xi_{TH} = \xi_T + \xi_H$ et $\xi_{HV} = \xi_H + \xi_V$.

Un autre exemple consiste à prendre

$$\tilde{X}_1 = -H_2 - \frac{3}{2}s + \frac{1}{2}H^+ + \frac{3}{2}Hs, \quad (4.57)$$

$$\tilde{X}_2 = -Hs + H^+ + s \quad (4.58)$$

correspondant à $\Pi = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$, $\tilde{X}_1 = \tilde{\nu}_a \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ et $\tilde{X}_2 = \tilde{\nu}_a \begin{pmatrix} 0 \\ 1 \end{pmatrix}$, avec

$$\tilde{\nu}_a = \begin{pmatrix} -1 & 0 \\ \frac{1}{2} & 1 \\ \frac{3}{2} & -1 \\ -\frac{3}{2} & 1 \end{pmatrix}.$$

Échanges dans un réacteur ouvert et conservation de la matière. La variation du nombre total de moles de l'espèce X_k pendant un temps dt , s'écrit :

$$dn_k = d_i n_k + d_e n_k, \quad k = 1, \dots, N \quad (4.59)$$

avec

$$d_e n_k = r_k dt, \quad k = 1, \dots, N \quad (4.60)$$

qui est la quantité d'espèce X_k , $k = 1, \dots, N$ échangée par le réacteur avec l'extérieur pendant dt où $r_k dt$ est le flux de l'espèce X_k , $k = 1, \dots, N$ échangé avec l'extérieur.

Remarque 4.4. L'expression (4.59) est à la base de l'application de la thermodynamique des systèmes hors équilibre. Elle est analogue à l'équation 4.1.4 de Prigogine [54]. L'indice i est pour irréversible, ce qui sera précisé plus loin.

On a alors en utilisant (4.37) et (4.60) :

$$dn_k = \sum_{j=1}^M \nu_{kj} d\xi_j + r_k dt, \quad k = 1, \dots, N \quad (4.61)$$

ou encore

$$\dot{n} = \nu \dot{\xi} + r \quad (4.62)$$

En faisant apparaître comme précédemment un ensemble maximal de réactions indépendantes, on a aussi bien :

$$\dot{n} = \tilde{\nu} \dot{\xi} + r \quad (4.63)$$

M_ν étant le rang de $\tilde{\nu}$, on peut supposer sans perte de généralité que les M_ν premières lignes de $\tilde{\nu}$ sont indépendantes, et on écrit :

$$n = \begin{pmatrix} n_c \\ n_{\bar{c}} \end{pmatrix}, \quad n_c = \begin{pmatrix} n_1 \\ \vdots \\ n_{M_\nu} \end{pmatrix}, \quad n_{\bar{c}} = \begin{pmatrix} n_{M_\nu+1} \\ \vdots \\ n_N \end{pmatrix} \quad (4.64)$$

En partitionnant $\tilde{\nu} \in \mathbb{R}^{N \times M_\nu}$ selon :

$$\tilde{\nu} = \begin{pmatrix} \nu_c \\ \nu_{\bar{c}} \end{pmatrix}, \quad \nu_c \in \mathbb{R}^{M_\nu \times M_\nu}, \quad \nu_{\bar{c}} \in \mathbb{R}^{(N-M_\nu) \times M_\nu} \quad (4.65)$$

on a donc

$$\nu_c \text{ inversible.}$$

En posant $r_c = \nu_c r$, $r_{\bar{c}} = \nu_{\bar{c}} r$, on a alors :

$$\dot{n}_{\bar{c}} = \nu_{\bar{c}} \dot{\xi} + r_{\bar{c}} \quad (4.66a)$$

d'où l'équation de conservation qui s'écrit sous la forme :

$$\dot{n}_c = \nu_c \dot{\xi} + r_c \quad (4.67a)$$

$$\dot{n}_{\bar{c}} = \nu_{\bar{c}} \nu_c^{-1} (\dot{n}_c - r_c) + r_{\bar{c}} \quad (4.67b)$$

En introduisant le changement de variable d'état suivant :

$$\xi_c = \nu_c^{-1} n_c, \quad \tilde{n}_{\bar{c}} = n_{\bar{c}} - \nu_{\bar{c}} \xi_c \quad (4.68)$$

le système (4.67) s'écrit sous la forme :

$$\dot{\xi}_c = \dot{\xi} + \nu_c^{-1} r_c \quad (4.69a)$$

$$\dot{\tilde{n}}_{\bar{c}} = r_{\bar{c}} - \nu_{\bar{c}} \nu_c^{-1} r_c \quad (4.69b)$$

Exemple du mécanisme réactionnel de Tafel-Heyrovsky-Volmer (suite). On a :

$$n_{c,a} = \begin{pmatrix} n_{H_2} \\ n_{H^+} \end{pmatrix}, \quad r_{c,a} = \begin{pmatrix} r_{H_2} \\ r_{H^+} \end{pmatrix}, \quad r_{\bar{c},a} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad (4.70)$$

$$\nu_{c,a} = \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix}, \quad \nu_{\bar{c},a} = \begin{pmatrix} 2 & -1 \\ -2 & 1 \end{pmatrix} \quad (4.71)$$

d'où

$$\nu_{\bar{c},a} \nu_{c,a}^{-1} = \begin{pmatrix} 2 & -1 \\ -2 & 1 \end{pmatrix} \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} -2 & -1 \\ 2 & 1 \end{pmatrix} \quad (4.72)$$

Le changement de variable (4.69) correspond donc ici à changer n pour les nouvelles variables :

$$\xi_{c,a} = \begin{pmatrix} -n_{H_2} \\ n_{H^+} \end{pmatrix}, \quad \tilde{n}_{\bar{c},a} = \begin{pmatrix} n_{H_2} + 2n_{H^+} \\ n_s - 2n_{H_2} - n_{H^+} \end{pmatrix} \quad (4.73)$$

On a alors avec (4.56) et (4.69a) :

$$\dot{\xi}_{c,a} = \begin{pmatrix} -\dot{n}_{H_2} \\ \dot{n}_{H^+} \end{pmatrix} = \begin{pmatrix} \dot{\xi}_{TH} \\ \dot{\xi}_{HV} \end{pmatrix} + \begin{pmatrix} -r_{H_2} \\ r_{H^+} \end{pmatrix}$$

dont on déduit

$$\dot{n}_{H_2} = -\dot{\xi}_{TH} + r_{H_2} \quad (4.74a)$$

$$\dot{n}_{H^+} = \dot{\xi}_{HV} + r_{H^+} \quad (4.74b)$$

On a d'autre part, d'après (4.68) :

$$\dot{n}_{\bar{c},a} = \dot{n}_{\bar{c},a} - \nu_{\bar{c},a} \dot{\xi}_{c,a} = \begin{pmatrix} \dot{n}_{Hs} + 2\dot{n}_{H_2} + \dot{n}_{H^+} \\ \dot{n}_s - 2\dot{n}_{H_2} - \dot{n}_{H^+} \end{pmatrix} = r_{\bar{c},a} - \nu_{\bar{c},a} \nu_{c,a}^{-1} r_{c,a} = \begin{pmatrix} 2r_{H_2} + r_{H^+} \\ -2r_{H_2} - r_{H^+} \end{pmatrix}$$

On en déduit les deux équations

$$\dot{n}_{Hs} + 2\dot{n}_{H_2} + \dot{n}_{H^+} = r_{H^+} + 2r_{H_2} \quad (4.74c)$$

$$\dot{n}_s - 2\dot{n}_{H_2} - \dot{n}_{H^+} = -r_{H^+} - 2r_{H_2} \quad (4.74d)$$

Le système (4.74) est ici l'analogue de (4.69). La formule (4.74c) s'interprète comme un bilan de matière sur les protons. Par ailleurs, en sommant (4.74c) et (4.74d), on fait apparaître l'équation :

$$\dot{n}_{Hs} + \dot{n}_s = 0$$

qui s'interprète comme le bilan des sites catalytiques.

4.2.2 Energie interne, entropie, énergie de Gibbs et irréversibilité d'un système chimique ouvert

Energie interne, entropie, énergie de Gibbs d'un mélange en phase homogène.

L'énergie libre de Gibbs G d'un système thermodynamique est définie par [54] :

$$G = U + pV - TS \quad (4.75)$$

où U est l'énergie interne du système, T est la température, S est l'entropie, p et V sont respectivement la pression et le volume du mélange. Lorsque la température et la pression sont maintenues constantes ($dT = 0, dp = 0$), on obtient :

$$dG = dU + pdV - TdS \quad (4.76)$$

L'énergie libre d'un système constitué par un mélange en phase homogène d'espèces chimiques non chargées X_k , $k = 1, \dots, N$, peut s'écrire en supposant G fonction de n sous la forme [54] :

$$dG = \mu^T(n)dn \quad (4.77)$$

où μ est le vecteur des potentiels chimiques qui vérifie :

$$\mu(n) = \begin{pmatrix} \mu_c(n_c) \\ \mu_{\bar{c}}(n_{\bar{c}}) \end{pmatrix} = \begin{pmatrix} \frac{\partial G}{\partial n_1} \\ \vdots \\ \frac{\partial G}{\partial n_N} \end{pmatrix} \in \mathbb{R}^N. \quad (4.78)$$

Dans (4.78), les potentiels $\mu_c(n_c)$ et $\mu_{\bar{c}}(n_{\bar{c}})$ correspondent respectivement aux espèces rassemblées dans n_c et dans $n_{\bar{c}}$.

Affinités. On en déduit de (4.62) et (4.77) la formule suivante :

$$dG = \mu^\top(n) (\nu d\xi + r dt) \quad (4.79)$$

En définissant l'affinité chimique par ¹

$$A(n) = -\nu^\top \mu(n) = \Delta G^R(n) - \Delta G^P(n) \in \mathbb{R}^M \quad (4.80)$$

où ΔG^R , ΔG^P sont définies par :

$$\Delta G^R(n) = (\nu^R)^\top \mu(n), \quad \Delta G^P(n) = (\nu^P)^\top \mu(n) \quad (4.81)$$

l'équation (4.79) s'écrit

$$dG = -A^\top(n) d\xi + \mu^\top(n) r dt \quad (4.82)$$

ou encore

$$\frac{dG}{dt} = -A^\top(n) \frac{d\xi}{dt} + \mu^\top(n) r. \quad (4.83)$$

Comme on le voit par la définition (4.80), le vecteur des affinités est le vecteur des différences entre énergies libres de formation des réactants et des produits, ou encore des énergies libres des réactions.

On déduit aussi, à partir de (4.79), la forme réduite :

$$\frac{dG}{dt} = \mu^\top(n) \tilde{\nu} \frac{d\xi}{dt} + \mu^\top(n) r \quad (4.84)$$

En utilisant le fait que : $d\xi_c = d\tilde{\xi} + \nu_c^{-1} r_c dt$, on obtient :

$$\frac{dG}{dt} = \mu^\top(n) \tilde{\nu} \left(\frac{d\xi_c}{dt} - \nu_c^{-1} r_c \right) + \mu^\top(n) r \quad (4.85)$$

Variations de l'entropie, seconde loi de la thermodynamique et irréversibilité.

D'après (4.76) et (4.77), la variation d'entropie vaut :

$$dS = \frac{1}{T} (dU + p dV - \mu^\top dn) \quad (4.86)$$

De Donder [54] a distingué deux termes dans la variation de l'entropie dS :

$$dS = d_i S + d_e S \quad (4.87)$$

où $d_e S$ est la variation d'entropie due aux échanges réversibles de matière et d'énergie avec l'extérieur, et $d_i S$ est la variation d'entropie produite par les processus irréversibles à l'intérieur du système. Les deux quantités $d_i S$ et $d_e S$ sont alors définies par [54] :

$$d_e S = \frac{dU + p dV}{T} - \frac{1}{T} \mu^\top d_e n, \quad (4.88a)$$

$$d_i S = -\frac{1}{T} \mu^\top d_i n \quad (4.88b)$$

1. Comme écrit dans [54], page 111, "l'affinité d'une réaction : $X + Y \rightleftharpoons 2Z$ définie comme $A = \mu_X + \mu_Y - 2\mu_Z$, peut être interprétée comme l'opposé du changement d'énergie libre de Gibbs quand 1 mole de X et 1 mole de Y réagissent pour produire 2 moles de Z ". Ainsi l'affinité est vue tantôt comme une énergie, tantôt comme une énergie par unité de matière. Au sens strict, il faudrait multiplier le membre de droite de (4.80) par "1 mole".

La seconde loi de la thermodynamique s'exprime par la relation suivante :

$$d_i S = -\frac{1}{T} \mu^\top d_i n \geq 0 \quad (4.89)$$

ce qui s'écrit encore avec (4.60) et (4.59) lorsque les dérivés sont définies :

$$\frac{d_i S}{dt} = -\frac{1}{T} \mu^\top \left(\frac{dn}{dt} - r \right) \geq 0. \quad (4.90)$$

Cette équation représente la propriété de dissipativité du système. La compatibilité thermodynamique de notre réacteur conduit à écrire, par analogie avec (4.89) :

$$\frac{d_i S}{dt} = \frac{A^\top(n)}{T} \frac{d\xi}{dt} \geq 0 \quad (4.91)$$

avec

$$A^\top(n) \frac{d\xi}{dt} \geq 0. \quad (4.92)$$

On peut aussi exprimer ceci en fonction de la variable ξ_c :

$$-\mu^\top(n) \tilde{\nu} \left(\frac{d\xi_c}{dt} - \nu_c^{-1} r_c \right) \geq 0. \quad (4.93)$$

Formule de Gibbs-Duhem et potentiels isothermes de gaz parfaits. On déduit de (4.76) et (4.77) la formule :

$$dS = \frac{1}{T} dU + \frac{p}{T} dV - \frac{1}{T} \mu^\top dn \quad (4.94)$$

On utilise maintenant le fait que l'entropie est une fonction d'état. En choisissant comme variable d'état (U, V, n) , on en déduit de (4.94) :

$$\left(\frac{\partial S}{\partial U} \right)_{V,n} = \frac{1}{T}, \quad \left(\frac{\partial S}{\partial V} \right)_{U,n} = \frac{p}{T}, \quad \left(\frac{\partial S}{\partial n} \right)_{U,V} = -\frac{\mu}{T} \quad (4.95)$$

Par ailleurs, l'entropie est une grandeur extensive, on a donc pour tout réel λ positif :

$$S(\lambda U, \lambda V, \lambda n) = \lambda S(U, V, n) \quad (4.96)$$

En différenciant cette formule en $\lambda = 1$, on déduit de cette propriété d'homogénéité l'identité suivante :

$$S = \left(\frac{\partial S}{\partial U} \right)_{V,n} U + \left(\frac{\partial S}{\partial V} \right)_{U,n} V + \left(\frac{\partial S}{\partial n} \right)_{U,V}^\top n \quad (4.97)$$

En combinant (4.95) et (4.97), on obtient la relation fondamentale

$$S = \frac{U}{T} + \frac{pV}{T} - \frac{1}{T} \mu^\top n \quad (4.98)$$

La formule précédente implique

$$U = TS - pV + \mu^\top n. \quad (4.99)$$

En différenciant cette relation et en comparant à (4.76), on obtient l'équation suivante, connue sous le nom d'équation de Gibbs-Duhem :

$$SdT - Vdp + n^\top d\mu = 0. \quad (4.100)$$

On se place maintenant dans le cas d'évolutions isothermes : $dT = 0$. Considérons les variations accompagnant un changement uniquement de la quantité n_k de moles d'espèce X_k pour k fixé et supposons que tous les composants vérifient la loi des gaz parfaits. On déduit alors de (4.100) :

$$n_k d\mu_k = V dp = V dp_k = RT dn_k \quad (4.101)$$

où p_k est la pression partielle de l'espèce gazeuse X_k , et où on a utilisé la loi des gaz parfaits. On a donc :

$$d(\mu_k - RT \ln(n_k)) = 0$$

En intégrant la formule précédente, le long d'une isotherme, on obtient donc :

$$\mu_k(n_k) = \mu_k^0(T) + RT \ln\left(\frac{n_k}{V}\right), \quad i = 1, \dots, N \quad (4.102)$$

où par commodité on a introduit la concentration $\frac{n_k}{V}$ et μ_k^0 est le potentiel de référence de l'espèce X_k . On a jusqu'ici supposé que seul n_k variait. Mais on sait par ailleurs que le potentiel μ_k ne dépend que de n_k .

La formule précédente est donc vraie en toute généralité. On l'écrira en regroupant les différentes formules scalaires sous la forme vectorielle :

$$\mu(n) = \mu^0 + RT \ln\left(\frac{n}{V}\right) \quad (4.103)$$

En utilisant (4.102), ΔG^R et ΔG^P , définis par (4.81), s'écrivent sous la forme :

$$\Delta G^R(n) = (\nu^R)^\top \mu(n) = (\nu^R)^\top \left(\mu^0 + RT \ln\left(\frac{n}{V}\right) \right) \quad (4.104a)$$

$$\Delta G^P(n) = (\nu^P)^\top \mu(n) = (\nu^P)^\top \left(\mu^0 + RT \ln\left(\frac{n}{V}\right) \right) \quad (4.104b)$$

4.2.3 Représentations d'état d'un système chimique ouvert et dissipativité

Cinétique chimique et propriété de dissipativité. On considère qu'au cours d'une réaction, un système chimique passe une "barrière de potentiel", le maximum d'énergie séparant les minima locaux du mélange des réactants et du mélange des produits (voir la Figure 4.7).

Les énergies libre d'activation des réactions directe et rétrograde sont définies par :

$$\Delta G^{\ddagger R} = \Delta G^* - \Delta G^R \quad (4.105a)$$

$$\Delta G^{\ddagger P} = \Delta G^* - \Delta G^P \quad (4.105b)$$

où ΔG^* est l'énergie libre du complexe activé de la réaction. L'évolution des réactions dépend des énergies d'activation :

$$\dot{\xi} = \Lambda(\Delta G^{\ddagger R}(n), \Delta G^{\ddagger P}(n)) \quad (4.106)$$

d'où l'on déduit la représentation d'état

$$\dot{n} = \nu \Lambda(\Delta G^{\ddagger R}(n), \Delta G^{\ddagger P}(n)) + r \quad (4.107)$$

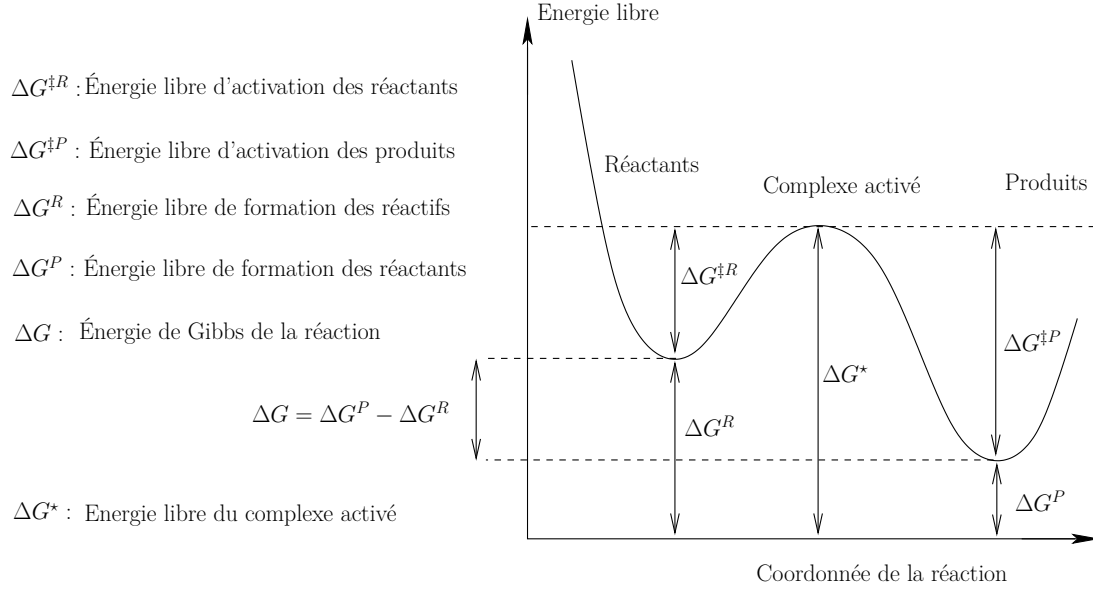


FIGURE 4.7 – Les différentes énergies libres utilisées pour décrire les réactions. Toutes les grandeurs sont signées.

où $\Lambda : \mathbb{R}^M \times \mathbb{R}^M \rightarrow \mathbb{R}^M$ reste à déterminer. On verra plus loin dans le cas des systèmes électrochimiques que nous examinons, que r peut être considéré comme une fonction de n et d'entrées exogènes. Il est nécessaire de vérifier la compatibilité thermodynamique de ce modèle c'est-à-dire d'après (4.91) :

$$\forall G', G'' \in \mathbb{R}^M, \quad (G' - G'')^\top \Lambda(G', G'') \leq 0. \quad (4.108)$$

En particulier, pour des raisons de continuité on a :

$$\forall G \in \mathbb{R}^M, \quad \Lambda(G, G) = 0. \quad (4.109)$$

Un choix possible conduisant à respecter (4.108) et (4.109) consiste à définir

$$\Lambda(G', G'') = \mathcal{N}(G') - \mathcal{N}(G'') \quad (4.110)$$

où $-\mathcal{N}$ est un *opérateur positif*, ou en d'autres termes :

$$\forall G', G'' \in \mathbb{R}^M, \quad (G' - G'')^\top (\mathcal{N}(G') - \mathcal{N}(G'')) \leq 0. \quad (4.111)$$

En effet, d'après (4.105), l'équation (4.110) implique alors :

$$\Delta G^\top \Lambda(\Delta G^{\ddagger R} - \Delta G^{\ddagger P}) = (\Delta G^{\ddagger P} - \Delta G^{\ddagger R})^\top \Lambda(\Delta G^{\ddagger R} - \Delta G^{\ddagger P}) \geq 0.$$

Il en est clairement ainsi pour l'exponentiel $x \mapsto e^{-x}$, et également pour l'opérateur $G \mapsto \text{Exp}(-G)$ à valeurs vectorielles.

La loi d'action de masse couramment introduite (cf. paragraphe A.4 d'Oster et Perelson [64] et les références citées) pour décrire la cinétique des réactions, revient justement à prendre :

$$\forall G', G'' \in \mathbb{R}^M, \quad \Lambda(G', G'') = \kappa \cdot \left(\text{Exp}\left(-\frac{G'}{RT}\right) - \text{Exp}\left(-\frac{G''}{RT}\right) \right) \quad (4.112)$$

où $\kappa = (\kappa_1, \dots, \kappa_M)^\top \in \mathbb{R}^M$, $\kappa_j \geq 0, j = 1, \dots, M$ est le vecteur des facteurs pré-exponentiels. Nous rappelons que la notation “ $\cdot$ ” représente le produit composante par composante de deux vecteurs. On a alors d’après (4.106) :

$$\dot{\xi}(t) = \kappa \cdot \left(\text{Exp} \left(\frac{-\Delta G^{\ddagger R}(n)}{RT} \right) - \text{Exp} \left(\frac{-\Delta G^{\ddagger P}(n)}{RT} \right) \right) \quad (4.113)$$

En utilisant (4.104) et (4.105), on obtient :

$$\dot{\xi}(t) = \kappa \cdot \text{Exp} \left(-\frac{\Delta G^*}{RT} \right) \cdot \left[\text{Exp} \left((\nu^R)^\top \left(\frac{\mu_0}{RT} + \text{Ln} \left(\frac{n}{V} \right) \right) \right) - \text{Exp} \left((\nu^P)^\top \left(\frac{\mu_0}{RT} + \text{Ln} \left(\frac{n}{V} \right) \right) \right) \right] \quad (4.114)$$

De façon alternative, on peut aussi écrire

$$\begin{aligned} \dot{\xi}(t) = \kappa \cdot \text{Exp} \left(-\frac{\Delta G^*}{RT} \right) \cdot \left[\text{Exp} \left(\frac{(\nu^R)^\top \mu_0}{RT} \right) \cdot \text{Exp} \left((\nu^R)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \right. \\ \left. - \text{Exp} \left(\frac{(\nu^P)^\top \mu_0}{RT} \right) \cdot \text{Exp} \left((\nu^P)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \right] \end{aligned} \quad (4.115)$$

En notant :

$$k_R = \kappa \cdot \text{Exp} \left(-\frac{\Delta G^*}{RT} \right) \cdot \text{Exp} \left(\frac{(\nu^R)^\top \mu_0}{RT} \right), \quad k_P = \kappa \cdot \text{Exp} \left(-\frac{\Delta G^*}{RT} \right) \cdot \text{Exp} \left(\frac{(\nu^P)^\top \mu_0}{RT} \right) \quad (4.116)$$

on a

$$\dot{\xi} = k_R \cdot \text{Exp} \left((\nu^R)^\top \text{Ln} \left(\frac{n}{V} \right) \right) - k_P \cdot \text{Exp} \left((\nu^P)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \quad (4.117)$$

En utilisant (4.62), la représentation d’état (4.107) s’écrit alors de la même manière que [85] sous la forme suivante :

$$\dot{n} = \nu \left[k_R \cdot \text{Exp} \left((\nu^R)^\top \text{Ln} \left(\frac{n}{V} \right) \right) - k_P \cdot \text{Exp} \left((\nu^P)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \right] + r, \quad n(0) = n^0 \geq 0 \quad (4.118)$$

On montre que le modèle (4.112) est dissipatif, au sens de la proposition suivante :

Proposition 4.5. *Pour tout n^0 donné, on a :*

$$\forall t \geq 0, \quad \left(\Delta G^R(n) - \Delta G^P(n) \right)^\top \dot{\xi} \geq 0 \quad (4.119)$$

pour tout solution n de l’équation différentielle (4.106).

La démonstration découle directement du fait de la positivité de l’opérateur $G \mapsto \text{Exp}(G)$. Pour l’énergie libre G , on déduit immédiatement de (4.85) et de la dissipativité de la loi d’action de masse la proposition suivante :

Proposition 4.6. *L’énergie libre G vérifie la propriété suivante :*

$$\frac{dG}{dt} \leq \mu^\top r. \quad (4.120)$$

L'énergie libre sert de fonction de Lyapunov pour (4.118) pour assurer la stabilité asymptotique globale lorsque les apports r sont nuls.

On peut utiliser le changement de variable (4.68) dans la représentation d'état du système. En utilisant la décomposition (4.69) et le fait que le vecteur réduit $\tilde{\xi}$ des taux d'avancement vérifie (cf. (4.56)) $\tilde{\xi} = \Pi\xi$, l'équation (4.118) peut se réécrire sous la forme :

$$\dot{\xi}_c = \Pi\mathcal{L}(\xi_c, \tilde{n}_{\bar{c}}) + \nu_c^{-1}r_c \quad (4.121a)$$

$$\dot{\tilde{n}}_{\bar{c}} = r_{\bar{c}} - \nu_{\bar{c}}\nu_c^{-1}r_c \quad (4.121b)$$

où on pose, pour tout $\xi'_c \in \mathbb{R}^{M_\nu}$, $\tilde{n}'_{\bar{c}} \in \mathbb{R}^{N-M_\nu}$:

$$\mathcal{L}(\xi'_c, \tilde{n}'_{\bar{c}}) = \Lambda(\Delta G^{\ddagger R}(n'), \Delta G^{\ddagger P}(n')), \quad n' = \begin{pmatrix} \nu_c \xi'_c \\ \nu_{\bar{c}} \xi'_c + \tilde{n}'_{\bar{c}} \end{pmatrix} \quad (4.122)$$

Dissipativité et représentation réduite. On suppose que l'opérateur Λ associé au système (4.106) vérifie (4.108). On suppose qu'il existe pour le système de réactions réduit exhibé en 4.2.1 une décomposition en produits et réactants telle que :

$$\nu^R = \tilde{\nu}^R \Pi \quad (4.123a)$$

$$\nu^P = \tilde{\nu}^P \Pi \quad (4.123b)$$

où la matrice Π est telle que $\nu = \tilde{\nu}\Pi$ et $\tilde{\nu}$ est de rang plein.

D'après (4.81) :

$$\Delta G^{\ddagger R} = \Delta G^* - \Delta G^R = \Delta G^* - (\nu^R)^\top \mu = \Delta G^* - \Pi^\top (\tilde{\nu}^R)^\top \mu \quad (4.124a)$$

$$\Delta G^{\ddagger P} = \Delta G^* - \Delta G^P = \Delta G^* - (\nu^P)^\top \mu = \Delta G^* - \Pi^\top (\tilde{\nu}^P)^\top \mu \quad (4.124b)$$

En définissant pour $\tilde{G}', \tilde{G}'' \in \mathbb{R}^{M_\nu}$

$$\tilde{\Lambda}(\tilde{G}', \tilde{G}'') = \Pi\Lambda(\Delta G^* - \Pi^\top \tilde{G}', \Delta G^* - \Pi^\top \tilde{G}'') \quad (4.125)$$

on a :

$$\dot{\tilde{\xi}} = \tilde{\Lambda}((\tilde{\nu}^R)^\top \mu(n), (\tilde{\nu}^P)^\top \mu(n)) \quad (4.126)$$

à comparer avec (4.106).

Lemme 4.7. Soit Π de rang plein tel que $\tilde{\nu}^R, \tilde{\nu}^P \geq 0$ vérifient $\nu^R = \Pi\tilde{\nu}^R$, $\nu^P = \Pi\tilde{\nu}^P$. Si

$$\forall G', G'' \in \mathbb{R}^M, \quad (G' - G'')^\top \Lambda(G', G'') \leq 0,$$

alors

$$\forall \tilde{G}', \tilde{G}'' \in \mathbb{R}^{M_\nu}, \quad (\tilde{G}' - \tilde{G}'')^\top \tilde{\Lambda}(\tilde{G}', \tilde{G}'') \leq 0.$$

Démonstration.

$$\begin{aligned} (\tilde{G}' - \tilde{G}'')^\top \tilde{\Lambda}(\tilde{G}', \tilde{G}'') &= (\tilde{G}' - \tilde{G}'')^\top \Pi\Lambda(\Pi^\top \tilde{G}', \Pi^\top \tilde{G}'') \\ &= \langle \Pi^\top \tilde{G}' - \Pi^\top \tilde{G}'', \Lambda(\Pi^\top \tilde{G}', \Pi^\top \tilde{G}'') \rangle \leq 0 \end{aligned}$$

□

La représentation réduite obtenue au paragraphe 4.2.1 hérite donc de la propriété de compatibilité thermodynamique si $\tilde{\nu}^R, \tilde{\nu}^P$ vérifient (4.123).

Exemple du mécanisme réactionnel de Tafel-Heyrovsky-Volmer (suite). L'exemple du mécanisme réactionnel de THV montre que le système réduit ne vérifie pas toujours les hypothèses du Lemme 4.7. En effet pour ce système, aucune représentation réduite $(\tilde{\nu}^R, \tilde{\nu}^P)$ ne vérifient ces hypothèses. En effet, on a $\text{rank}(\nu) = 2$ mais $\text{rank}(\nu^P) = 3$. Ainsi, pour toutes matrices $\Pi \in \mathbb{R}^{2 \times 3}$ et $\tilde{\nu} \in \mathbb{R}^{4 \times 2}$ telle que $\tilde{\nu}\Pi = \nu$, on a $\tilde{\nu}^P\Pi \neq \nu^P$, car $\text{rank}(\nu^P) = 3 > 2 \geq \text{rank}(\tilde{\nu}^P)\Pi$. Ceci est vrai pour toute décomposition $\tilde{\nu} = \tilde{\nu}^P - \tilde{\nu}^R$.

4.2.4 Cas d'un réseau de réactions électrochimiques

Dans cette section, nous adaptons l'étude faite précédemment à un réseau de réactions électrochimiques dont les espèces X_k , $k = 1, \dots, N$ peuvent être chargées électriquement.

Potentels et affinités électrochimiques. Le potentiel électrochimique de l'espèce X_k , noté $\tilde{\mu}_k$, est une extension du potentiel chimique et s'écrit sous la forme² [74] :

$$\tilde{\mu}_k(\Delta\phi_S, n) = \left(\frac{\partial \tilde{G}}{\partial n_k} \right)_{p,T} = \mu_k(n) + z_k F \Delta\phi_S \quad (4.127)$$

où $\tilde{G}(\Delta\phi_S, n)$ est l'énergie libre du système, $\mu_k(n)$, resp. z_k est le potentiel chimique, resp. la valence de l'espèce X_k et $\Delta\phi_S$ est la différence de potentiel aux bornes de la couche compacte (couche de Stern). On note :

$$\tilde{\mu} = \begin{pmatrix} \tilde{\mu}_1 \\ \vdots \\ \tilde{\mu}_N \end{pmatrix} \in \mathbb{R}^N. \quad (4.128)$$

L'affinité électrochimique $\tilde{A}$ est définie par :

$$\tilde{A}(\Delta\phi_S, n) = -\nu^\top \tilde{\mu}(\Delta\phi_S, n) = \Delta\tilde{G}^R(\Delta\phi_S, n) - \Delta\tilde{G}^P(\Delta\phi_S, n) \quad (4.129)$$

où

$$\Delta\tilde{G}^R(\Delta\phi_S, n) = (\nu^R)^\top \tilde{\mu}(\Delta\phi_S, n), \quad \Delta\tilde{G}^P(\Delta\phi_S, n) = (\nu^P)^\top \tilde{\mu}(\Delta\phi_S, n). \quad (4.130)$$

Les réactions électrochimiques que nous considérons ont lieu à l'interface entre un milieu électrolyte et le métal, conducteur des électrons. Elles donnent lieu à un transfert de charges électroniques de part et d'autre de l'interface (plus précisément) de la couche compacte ou couche de Stern. Ce transfert entre deux milieux situés à des potentiels différents est pris en compte dans l'écriture de l'énergie libre de la réaction électrochimique (4.129).

On peut ainsi écrire l'énergie libre de la réaction comme :

$$\Delta\tilde{G} = \Delta\tilde{G}^P - \Delta\tilde{G}^R = \Delta G + zF\Delta\phi_S \quad (4.131)$$

où z est la valence des charges transférées (les électrons dans notre cas).

Cependant, ce sont les énergies d'activation qui interviennent dans la cinétique, cf. (4.113). Pour les réactions électrochimiques, on identifie l'avancement de la réaction et des charges d'un milieu à l'autre. Ainsi, on peut considérer que le complexe activé est soumis à une variation de potentiel intermédiaire entre 0 et $\Delta\phi_S$. On définit pour une réaction [74] :

$$\alpha = \frac{1}{Fz} \frac{\partial}{\partial \Delta\phi_S} \Delta\tilde{G}^{\ddagger R}, \quad \beta = -\frac{1}{Fz} \frac{\partial}{\partial \Delta\phi_S} \Delta\tilde{G}^{\ddagger P} \quad (4.132)$$

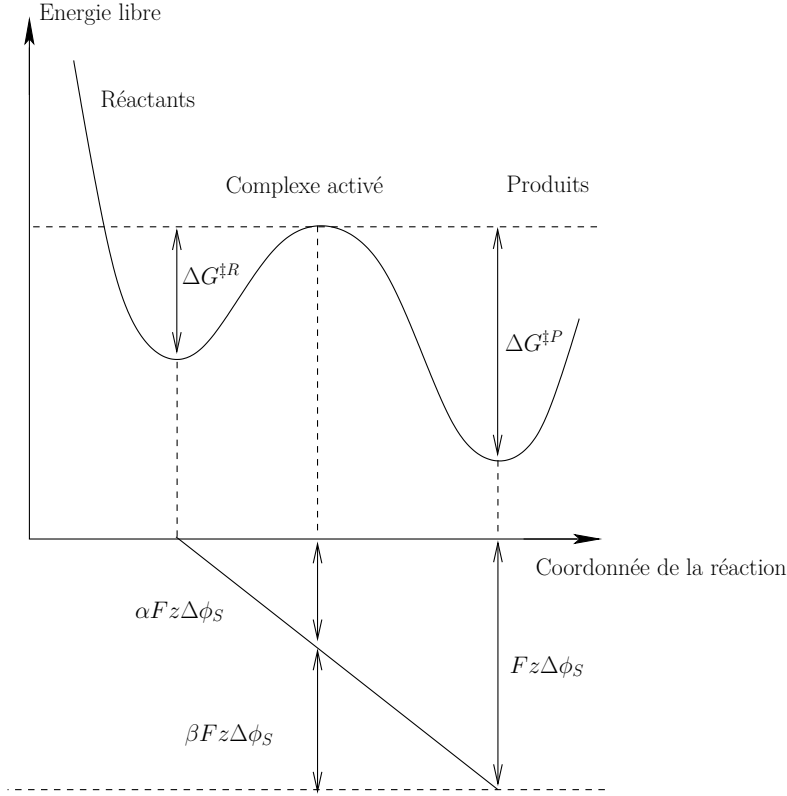


FIGURE 4.8 – Effet du potentiel électrique sur la cinétique électrochimique. La réaction directe est favorisée par un potentiel électrique $\Delta\phi_S < 0$ puisque alors $\Delta\tilde{G}^{\ddagger R}(\Delta\phi_S) < \Delta G^{\ddagger R}$. La réaction rétrograde est, elle, contrariée puisque $\Delta\tilde{G}^{\ddagger P}(\Delta\phi_S) > \Delta G^{\ddagger P}$.

On a :

$$\begin{aligned}\alpha + \beta &= \frac{1}{Fz} \frac{\partial}{\partial \Delta\phi_S} (\Delta\tilde{G}^{\ddagger R} - \Delta\tilde{G}^{\ddagger P}) \\ &= \frac{1}{Fz} \frac{\partial}{\partial \Delta\phi_S} (\Delta\tilde{G}^P - \Delta\tilde{G}^R) \\ &= 1\end{aligned}$$

On écrira donc :

$$\Delta\tilde{G}^{\ddagger R}(\Delta\phi_S) = \Delta G^{\ddagger R} - \alpha z F \Delta\phi_S \quad (4.133)$$

$$\Delta\tilde{G}^{\ddagger P}(\Delta\phi_S) = \Delta G^{\ddagger P} + \beta z F \Delta\phi_S \quad (4.134)$$

En utilisant la définition du potentiel chimique, donné par (4.103), on a alors :

$$\Delta G^* - \Delta\tilde{G}^{\ddagger R}(\Delta\phi_S, n) = (\nu^R)^\top \left(\mu^0 + RT \ln \left(\frac{n}{V} \right) + \alpha F \Delta\phi_S \right) \quad (4.135a)$$

$$\Delta G^* - \Delta\tilde{G}^{\ddagger P}(\Delta\phi_S, n) = (\nu^P)^\top \left(\mu^0 + RT \ln \left(\frac{n}{V} \right) - \beta F \Delta\phi_S \right) \quad (4.135b)$$

2. L'utilisation de tilde ici est traditionnelle [14, 74]. Ne pas confondre ici avec l'usage différent des quantités liées à la réduction du modèle.

Représentation d'état. De la même manière que (4.115), la cinétique de la réaction s'écrit sous la forme :

$$\begin{aligned}\dot{\xi}(t) &= \kappa \cdot \left(\text{Exp} \left(-\frac{\Delta \tilde{G}^{\ddagger R}(\Delta \phi_S, n)}{RT} \right) - \text{Exp} \left(-\frac{\Delta \tilde{G}^{\ddagger P}(\Delta \phi_S, n)}{RT} \right) \right) \\ &= \kappa \cdot \text{Exp} \left(-\frac{\Delta G^*}{RT} \right) \cdot \left[\text{Exp} \left((\nu^R)^\top \left(\frac{\mu^0}{RT} + \frac{\alpha F \Delta \phi_S}{RT} + \text{Ln} \left(\frac{n}{V} \right) \right) \right) \right. \\ &\quad \left. - \text{Exp} \left((\nu^P)^\top \left(\frac{\mu^0}{RT} - \frac{\beta F \Delta \phi_S}{RT} + \text{Ln} \left(\frac{n}{V} \right) \right) \right) \right] \quad (4.136)\end{aligned}$$

et

$$\begin{aligned}\dot{n}(t) &= \nu \left[\kappa \cdot \text{Exp} \left((\nu^R)^\top \left(\frac{\mu^0}{RT} + \frac{\alpha F \Delta \phi_S}{RT} + \text{Ln} \left(\frac{n}{V} \right) \right) \right) \right. \\ &\quad \left. - \kappa \cdot \text{Exp} \left((\nu^P)^\top \left(\frac{\mu^0}{RT} - \frac{\beta F \Delta \phi_S}{RT} + \text{Ln} \left(\frac{n}{V} \right) \right) \right) \right] + r \quad (4.137)\end{aligned}$$

En utilisant les définitions de k_R et k_P dans (4.116) :

$$k_R = \kappa \cdot \text{Exp} \left(-\frac{\Delta G^*}{RT} \right) \cdot \text{Exp} \left(\frac{(\nu^R)^\top \mu^0}{RT} \right), \quad k_P = \kappa \cdot \text{Exp} \left(-\frac{\Delta G^*}{RT} \right) \cdot \text{Exp} \left(\frac{(\nu^P)^\top \mu^0}{RT} \right)$$

la formule (4.136) s'écrit aussi sous la forme :

$$\begin{aligned}\dot{\xi} &= k_R \cdot \text{Exp} \left((\nu^R)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \cdot \text{Exp} \left(\frac{(\nu^R)^\top \alpha F \Delta \phi_S}{RT} \right) \\ &\quad - k_P \cdot \text{Exp} \left((\nu^P)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \cdot \text{Exp} \left(-\frac{(\nu^P)^\top \beta F \Delta \phi_S}{RT} \right) \quad (4.138)\end{aligned}$$

d'où

$$\begin{aligned}\dot{n} &= \nu \left[k_R \cdot \text{Exp} \left((\nu^R)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \cdot \text{Exp} \left(\frac{(\nu^R)^\top \alpha F \Delta \phi_S}{RT} \right) \right. \\ &\quad \left. - k_P \cdot \text{Exp} \left((\nu^P)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \cdot \text{Exp} \left(-\frac{(\nu^P)^\top \beta F \Delta \phi_S}{RT} \right) \right] + r \quad (4.139)\end{aligned}$$

Les débits exogènes r dépendent de concentrations n non électroniques. On peut considérer que cette dernière dépendance est affine (de type loi d'Ohm). On notera donc

$$r(n, I_M) = B_n(n^{ext} - n) + B_I I_M \quad (4.140)$$

où $B_n \in \mathbb{R}^{N \times N}$, $B_I \in \mathbb{R}^N$; et où le vecteur n^{ext} représente les concentrations extérieures en gaz hydrogène et oxygène dans les GDL (il s'agit d'un modèle élémentaire de GDL, que l'on peut enrichir) et I_M le courant dans la membrane.

La tension $\Delta \phi_S$ aux bornes de la couche de Stern est proportionnelle à la charge contenue dans la couche diffuse (cf. Annexe A), elle-même égale à une constante près à la quantité de protons dans la couche diffuse. On écrira donc, en notant C_{S0} la capacité associée à la couche de Stern :

$$\Delta \phi_S = -\frac{1}{F C_{S0}} c^\top n \quad (4.141)$$

pour un vecteur $c \in \mathbb{R}^N$ choisi de façon adéquate. Cette formule reste valable dans le cas où d'autres anions et cations sont présents, il faut alors également compter leurs charges.

On a par ailleurs :

$$r_{H^+} = B_I I_M = \frac{1}{F} I_M. \quad (4.142)$$

On a alors la représentation d'état suivante :

$$\dot{n} = \nu \left[k_R \cdot \text{Exp} \left((\nu^R)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \cdot \text{Exp} \left(\frac{(\nu^R)^\top \alpha F \Delta \phi_S}{RT} \right) - k_P \cdot \text{Exp} \left((\nu^P)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \cdot \text{Exp} \left(-\frac{(\nu^P)^\top \beta F \Delta \phi_S}{RT} \right) \right] + r(n, I_M) \quad (4.143a)$$

$$r(n, I_M) = B_n(n^{ext} - n) + B_I I_M \quad (4.143b)$$

$$\Delta \phi_S = -\frac{1}{F C_{S0}} c^\top n \quad (4.143c)$$

Représentations réduites. Comme on l'a fait pour les systèmes de réaction chimiques (cf. le passage de (4.118) à (4.121)), on peut écrire de façon plus structurée l'équation d'état (4.143a). En utilisant de nouveau le changement de variable (4.68), on a :

$$\dot{\xi}_c = \Pi \mathcal{L}(\xi_c, \tilde{n}_{\bar{c}}) + \nu_c^{-1} r_c \quad (4.144a)$$

$$\dot{\tilde{n}}_{\bar{c}} = r_{\bar{c}} - \nu_{\bar{c}} \nu_c^{-1} r_c \quad (4.144b)$$

où $\mathcal{L}$ est ici définie par :

$$\mathcal{L}(\xi'_c, \tilde{n}'_{\bar{c}}) = \Lambda \left((\nu^R)^\top \tilde{\mu}(n'), (\nu^P)^\top \tilde{\mu}(n') \right), \quad n' = \begin{pmatrix} \nu_c \xi'_c \\ \nu_{\bar{c}} \xi'_c + \tilde{n}'_{\bar{c}} \end{pmatrix}$$

c'est-à-dire

$$\mathcal{L}(\xi'_c, \tilde{n}'_{\bar{c}}) = k_R \cdot \text{Exp} \left((\nu^R)^\top \left(-\frac{\alpha c^\top n'}{C_{S0} RT} + \text{Ln} \left(\frac{n'}{V} \right) \right) \right) - k_P \cdot \text{Exp} \left((\nu^P)^\top \left(\frac{\beta c^\top n'}{C_{S0} RT} + \text{Ln} \left(\frac{n'}{V} \right) \right) \right), \quad n' = \begin{pmatrix} \nu_c \xi'_c \\ \nu_{\bar{c}} \xi'_c + \tilde{n}'_{\bar{c}} \end{pmatrix}. \quad (4.145)$$

En utilisant les variables $\xi_c, \tilde{n}_{\bar{c}}$ introduites en (4.68) plutôt que n pour exprimer la dépendance des débits r (cf.(4.140)), on écrira aussi bien :

$$r(\xi_c, \tilde{n}_{\bar{c}}, I) = B_n \begin{pmatrix} \nu_c & 0_{M_\nu \times (N-M_\nu)} \\ \nu_{\bar{c}} & I_{N-M_\nu} \end{pmatrix} \begin{pmatrix} \xi_c^{ext} - \xi_c \\ \tilde{n}_{\bar{c}}^{ext} - \tilde{n}_{\bar{c}} \end{pmatrix} + B_I I. \quad (4.146)$$

On exprimera aussi $\Delta \phi_S$ comme fonction de ξ_c et $\tilde{n}_{\bar{c}}$:

$$\Delta \phi_S = -\frac{1}{F C_{S0}} c^\top \begin{pmatrix} \nu_c & 0_{M_\nu \times (N-M_\nu)} \\ \nu_{\bar{c}} & I_{N-M_\nu} \end{pmatrix} \begin{pmatrix} \xi_c \\ \tilde{n}_{\bar{c}} \end{pmatrix} \quad (4.147)$$

En introduisant maintenant dans (4.143) le changement de variable (4.68), le système d'état s'écrit alors finalement :

$$\begin{aligned} \begin{pmatrix} \dot{\xi}_c \\ \dot{\tilde{n}}_{\bar{c}} \end{pmatrix} &= \begin{pmatrix} \Pi \mathcal{L}(\xi_c, \tilde{n}_{\bar{c}}) \\ 0_{N-M_\nu} \end{pmatrix} + \begin{pmatrix} \nu_c^{-1} & 0_{M_\nu \times (N-M_\nu)} \\ -\nu_{\bar{c}} \nu_c^{-1} & I_{N-M_\nu} \end{pmatrix} \left[B_n \begin{pmatrix} \nu_c & 0_{M_\nu \times (N-M_\nu)} \\ \nu_{\bar{c}} & I_{N-M_\nu} \end{pmatrix} \begin{pmatrix} \xi_c^{ext} - \xi_c \\ \tilde{n}_{\bar{c}}^{ext} - \tilde{n}_{\bar{c}} \end{pmatrix} + B_I I_M \right] \\ &\quad (4.148a) \\ \Delta \phi_S &= -\frac{1}{FC_{S0}} c^T \begin{pmatrix} \nu_c & 0_{M_\nu \times (N-M_\nu)} \\ \nu_{\bar{c}} & I_{N-M_\nu} \end{pmatrix} \begin{pmatrix} \xi_c \\ \tilde{n}_{\bar{c}} \end{pmatrix} \\ &\quad (4.148b) \end{aligned}$$

Exemple du mécanisme réactionnel de Tafel-Heyrovsky-Volmer (fin). On va écrire la représentation d'état du système de réactions (4.44) sous la forme (4.143). On rappelle que, pour ce système, n_a et ξ_a ont été définis par (4.45) et ν_a, ν_a^R et ν_a^P par (4.49).

Les seules entrées non nulles de $r(n, I_M)$ ici concernant l'hydrogène et les protons. On a :

$$r_a(n_a, I_M) = \begin{pmatrix} r_{H_2} \\ r_{H^+} \\ 0 \\ 0 \end{pmatrix} \quad (4.149)$$

En supposant que ces échanges sont essentiellement dûs à de la diffusion (depuis le GDL pour l'hydrogène), on écrira par exemple

$$r_a(n_a, I_M) = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix} S_{CA} D_{H_2}^{cc} (n_{H_2}^{GDL} - n_{H_2}) - \frac{1}{F} \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \end{pmatrix} I_M \quad (4.150)$$

où $D_{H_2}^{cc}$ est le coefficient de diffusion, dans la couche compacte, positif ou nul correspondant. La formule (4.143a) est ici valable, où les matrices $B_{n,a}, B_{I,a}$ sont données par

$$B_{n,a} = -S_{CA} \text{diag}\{D_{H_2}^{cc}; 0; 0; 0\}, \quad B_{I,a} = -\frac{1}{F} \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \end{pmatrix} \quad (4.151)$$

où S_{CA} est la surface d'échange qui vaut :

$$S_{CA} = \gamma S_{EME} e_{CA}. \quad (4.152)$$

On rappelle que γ est la surface active spécifique, S_{EME} est la surface géométrique de l'assemblage électrodes-membrane (EME) et e_{CA} est l'épaisseur de la couche active.

D'autre part, on a

$$c_a^T = \begin{pmatrix} 0 & 1 & 0 & 0 \end{pmatrix} \quad (4.153)$$

Enfin, les vecteurs α, β vérifient :

$$(\nu_a^R)^\top \alpha = \begin{pmatrix} 0 \\ \alpha_H \\ \alpha_V \end{pmatrix}, \quad (\nu_a^P)^\top \beta = \begin{pmatrix} 0 \\ 1 - \alpha_H \\ 1 - \alpha_V \end{pmatrix}, \quad \alpha_H, \alpha_V \in [0, 1] \quad (4.154)$$

où α_H , resp. α_V est le coefficient de transfert électronique de l'étape réactionnelle de Heyrovsky, resp. de Volmer.

On décomposera $\mathcal{L}$ défini par (4.145) selon la formule :

$$\mathcal{L}(n_a) = \begin{pmatrix} \mathcal{L}_T(n_a) \\ \mathcal{L}_H(n_a) \\ \mathcal{L}_V(n_a) \end{pmatrix} \quad (4.155)$$

où $\mathcal{L}_T, \mathcal{L}_H$ et $\mathcal{L}_V$ s'écrivent d'après (4.138) sous la forme :

$$\begin{pmatrix} \mathcal{L}_T(n_a) \\ \mathcal{L}_H(n_a) \\ \mathcal{L}_V(n_a) \end{pmatrix} = k_{Ra} \cdot \text{Exp} \left(-\frac{(\nu_a^R)^\top \alpha c_a^\top n_a}{C_{S0}RT} \right) \cdot \text{Exp} \left((\nu_a^R)^\top \text{Ln} \left(\frac{n_a}{V} \right) \right) \\ - k_{Pa} \cdot \text{Exp} \left(\frac{(\nu_a^P)^\top \beta c_a^\top n_a}{C_{S0}RT} \right) \cdot \text{Exp} \left((\nu_a^P)^\top \text{Ln} \left(\frac{n_a}{V} \right) \right) \quad (4.156)$$

4.2.5 Représentation réduite du mécanisme réactionnel de Tafel-Heyrovsky-Volmer et analogue électrique

Représentation réduite. On reprend le mécanisme réactionnel réduit de Tafel-Heyrovsky-Volmer donné par (4.52) :



Comme on l'a vu précédemment, ceci conduit à considérer le nouveau vecteur des taux d'avancement (cf.(4.56)) :

$$\begin{pmatrix} \xi_{TH} \\ \xi_{HV} \end{pmatrix} = \Pi \begin{pmatrix} \xi_T \\ \xi_H \\ \xi_V \end{pmatrix} = \begin{pmatrix} \xi_T + \xi_H \\ \xi_H + \xi_V \end{pmatrix} \quad (4.158)$$

où Π est donné par (4.53). En exploitant les propriétés introduites dans la section précédente, on écrit d'après (4.74) et (4.144) :

$$\begin{pmatrix} -\dot{n}_{H_2} \\ \dot{n}_{H^+} \end{pmatrix} = \begin{pmatrix} \mathcal{L}_{TH}(n_a) \\ \mathcal{L}_{HV}(n_a) \end{pmatrix} + \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} r_{H_2} \\ -\frac{I_M}{F} \end{pmatrix} \quad (4.159a)$$

$$\begin{pmatrix} \dot{n}_{Hs} + 2\dot{n}_{H_2} + \dot{n}_{H^+} \\ \dot{n}_s - 2\dot{n}_{H_2} - \dot{n}_{H^+} \end{pmatrix} = \begin{pmatrix} 2 & 1 \\ -2 & -1 \end{pmatrix} \begin{pmatrix} r_{H_2} \\ -\frac{I_M}{F} \end{pmatrix} \quad (4.159b)$$

où on a posé

$$\begin{pmatrix} \mathcal{L}_{TH}(n_a) \\ \mathcal{L}_{HV}(n_a) \end{pmatrix} = \Pi \begin{pmatrix} \mathcal{L}_T(n_a) \\ \mathcal{L}_H(n_a) \\ \mathcal{L}_V(n_a) \end{pmatrix} = \begin{pmatrix} \mathcal{L}_T(n_a) + \mathcal{L}_H(n_a) \\ \mathcal{L}_H(n_a) + \mathcal{L}_V(n_a) \end{pmatrix}$$

On effectue maintenant les calculs matriciels pour simplifier la forme prise ici par (4.148). $\nu_{c,a}, \nu_{\bar{c},a}$ étant définis par (4.71) et $B_{n,a}, B_{I,a}$ par (4.151), on trouve :

$$\begin{pmatrix} \nu_{c,a}^{-1} & 0_{2 \times 2} \\ -\nu_{\bar{c},a} \nu_{c,a}^{-1} & I_2 \end{pmatrix} B_{n,a} \begin{pmatrix} \nu_{c,a} & 0_{2 \times 2} \\ \nu_{\bar{c},a} & I_2 \end{pmatrix} = S_{CA} \begin{pmatrix} - \begin{pmatrix} D_{H_2}^{cc} & 0 \\ 0 & 0 \end{pmatrix} & 0_{2 \times 2} \\ \nu_{\bar{c},a} \begin{pmatrix} D_{H_2}^{cc} & 0 \\ 0 & 0 \end{pmatrix} & 0_{2 \times 2} \end{pmatrix} \\ = S_{CA} \begin{pmatrix} -D_{H_2}^{cc} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 2D_{H_2}^{cc} & 0 & 0 & 0 \\ -2D_{H_2}^{cc} & 0 & 0 & 0 \end{pmatrix}$$

et

$$\begin{pmatrix} \nu_{c,a}^{-1} & 0_{2 \times 2} \\ -\nu_{\bar{c},a} \nu_{c,a}^{-1} & I_2 \end{pmatrix} B_{I,a} = -\frac{I_M}{F} \begin{pmatrix} \nu_{c,a}^{-1} \begin{pmatrix} 0 \\ 1 \end{pmatrix} \\ -\nu_{\bar{c},a} \nu_{c,a}^{-1} \begin{pmatrix} 0 \\ 1 \end{pmatrix} \end{pmatrix} = -\frac{I_M}{F} \begin{pmatrix} 0 \\ 1 \\ 1 \\ -1 \end{pmatrix}.$$

Enfin

$$c_a^T \begin{pmatrix} \nu_{c,a} & 0_{2 \times 2} \\ \nu_{\bar{c},a} & I_2 \end{pmatrix} = \begin{pmatrix} 0 & 1 & 0 & 0 \end{pmatrix} \begin{pmatrix} \nu_{c,a} & 0_{2 \times 2} \\ \nu_{\bar{c},a} & I_2 \end{pmatrix} = \begin{pmatrix} 0 & 1 & 0 & 0 \end{pmatrix}.$$

Le système (4.159) s'écrit donc d'après (4.148) sous la forme :

$$\dot{\xi}_{c,a} = \begin{pmatrix} \mathcal{L}_{TH}(\xi_c, \tilde{n}_{\bar{c},a}) \\ \mathcal{L}_{HV}(\xi_{c,a}, \tilde{n}_{\bar{c},a}) \end{pmatrix} - S_{CA} \begin{pmatrix} D_{H_2}^{cc} & 0 \\ 0 & 0 \end{pmatrix} (\xi_{c,a}^{ext} - \xi_{c,a}) - \frac{I_M}{F} \begin{pmatrix} 0 \\ 1 \end{pmatrix} \quad (4.160a)$$

$$\dot{\tilde{n}}_{\bar{c},a} = \begin{pmatrix} 2D_{H_2}^{cc} & 0 \\ -2D_{H_2}^{cc} & 0 \end{pmatrix} (\xi_{c,a}^{ext} - \xi_{c,a}) - \frac{1}{F} \begin{pmatrix} 1 \\ -1 \end{pmatrix} I_M \quad (4.160b)$$

$$\Delta\phi_{Sa} = -\frac{1}{FC_{S0}} \begin{pmatrix} 0 & 1 & 0 & 0 \end{pmatrix} \begin{pmatrix} \xi_{c,a} \\ \tilde{n}_{\bar{c},a} \end{pmatrix} \quad (4.160c)$$

On rappelle (cf. (4.73)) que :

$$\xi_{c,a} = \begin{pmatrix} -n_{H_2} \\ n_{H^+} \end{pmatrix}, \quad \tilde{n}_{\bar{c},a} = \begin{pmatrix} n_{Hs} + 2n_{H_2} + n_{H^+} \\ n_s - 2n_{H_2} - n_{H^+} \end{pmatrix}$$

et

$$\xi_{c,a}^{ext} = \begin{pmatrix} -n_{H_2}^{ext} \\ 0 \end{pmatrix}, \quad \tilde{n}_{\bar{c},a}^{ext} = \begin{pmatrix} 2n_{H_2}^{ext} \\ -2n_{H_2}^{ext} \end{pmatrix}$$

Analogie électrique. Le schéma électrique représentant le système précédent est donné, d'après [64], par la Figure 4.9. On va maintenant le vérifier.

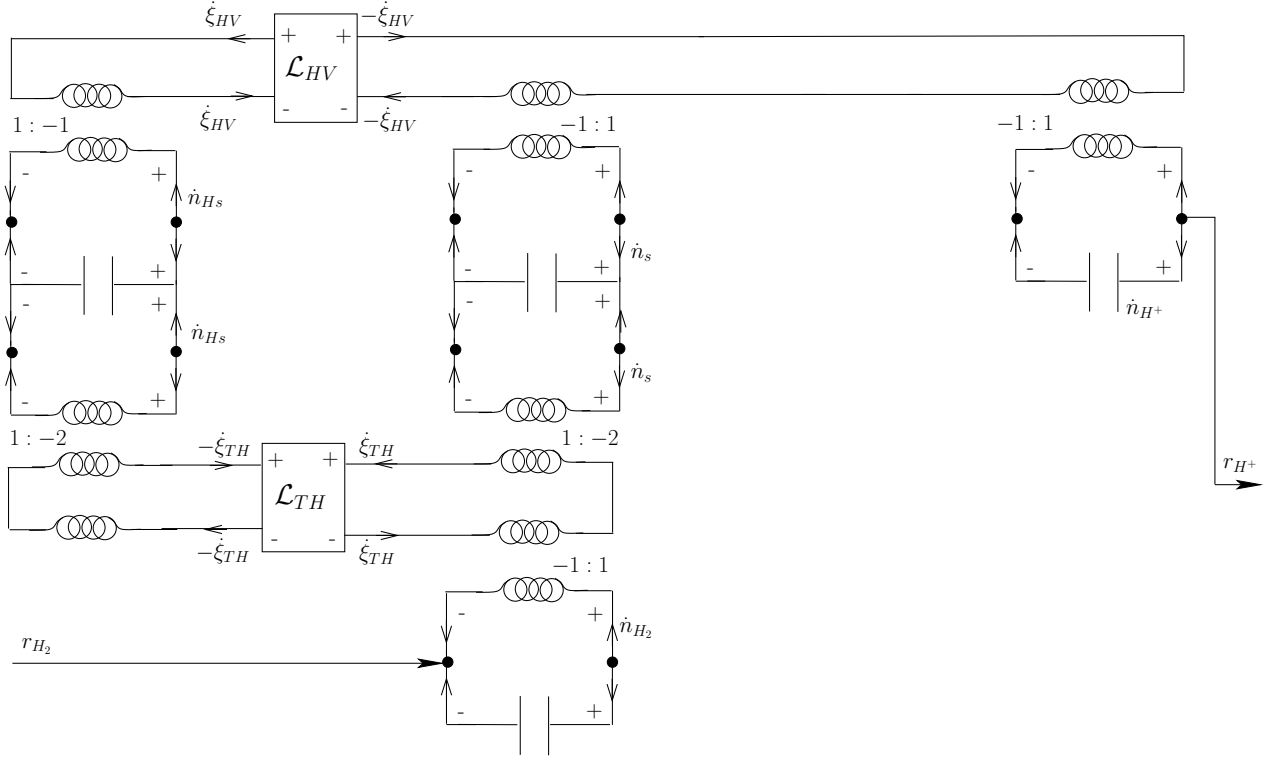


FIGURE 4.9 – Schéma électrique représentant les réactions (4.157).

La consommation ou la production de chacune des espèces est définie par les coefficients stœchiométriques. Ceux-ci sont introduits à travers des transformateurs idéaux. Si l'espèce k participe à la j -ème réaction comme étant un réactant ou bien un produit, le condensateur de k est connecté en parallèle avec l'enroulement primaire du transformateur idéal associé à k . Cette connexion est effectuée avec le ratio $1 : -\nu_{kj}^R$ si k est un réactant ou bien $-\nu_{kj}^P : 1$ si k est un produit, où ν_{kj}^R et ν_{kj}^P sont les coefficients stœchiométriques de k .

Le couplage entre les 2-ports est effectué en appliquant la loi de Kirchhoff aux nœuds c'est-à-dire la somme algébrique des intensités des courants arrivant à un nœud est nulle, tout en respectant la convention de signe : le sens positif du courant est le sens entrant.

Dans ce cadre [64], chaque espèce est représentée par un condensateur modélisant son stock. Ce stock est altéré par la réaction qui est dissipative comme nous l'avons montré dans la proposition 4.5. Les espèces H_2 et H^+ échangent avec l'extérieur à travers r_{H_2} , et r_{H^+} . On déduit en effet de la Figure 4.9

$$\begin{pmatrix} \dot{n}_{H_2} \\ \dot{n}_{H^+} \\ \dot{n}_{H_s} \\ \dot{n}_s \end{pmatrix} = \begin{pmatrix} -1 & 0 \\ 0 & 1 \\ 2 & -1 \\ -2 & 1 \end{pmatrix} \begin{pmatrix} \dot{\xi}_{TH} \\ \dot{\xi}_{HV} \end{pmatrix} + \begin{pmatrix} r_{H_2} \\ -\frac{I_M}{F} \\ 0 \\ 0 \end{pmatrix}$$

qui est équivalent à (4.160). Chaque réaction est alors représentée par un 2-port noté $\mathcal{L}$ sur la Figure 4.9. Le courant qui traverse le port qui relie $\mathcal{L}$ aux condensateurs modélisant le stock des réactants, resp. des produits, est $\dot{\xi}$ resp. $-\dot{\xi}$. Les tensions aux bornes des ports de $\mathcal{L}$ sont $\Delta\tilde{G}^R$ et $\Delta\tilde{G}^P$.

Le schéma analogique obtenu dans la Figure 4.9 représente d'une manière détaillée le condensateur attribué usuellement à la couche de Stern. En effet, sur la Figure 4.10, les condensateurs dans l'encadré noir représente d'une manière détaillé la condensateur de Stern C_S . Tandis que le condensateur dans l'encadré rouge est attribué au condensateur dans la couche diffuse.

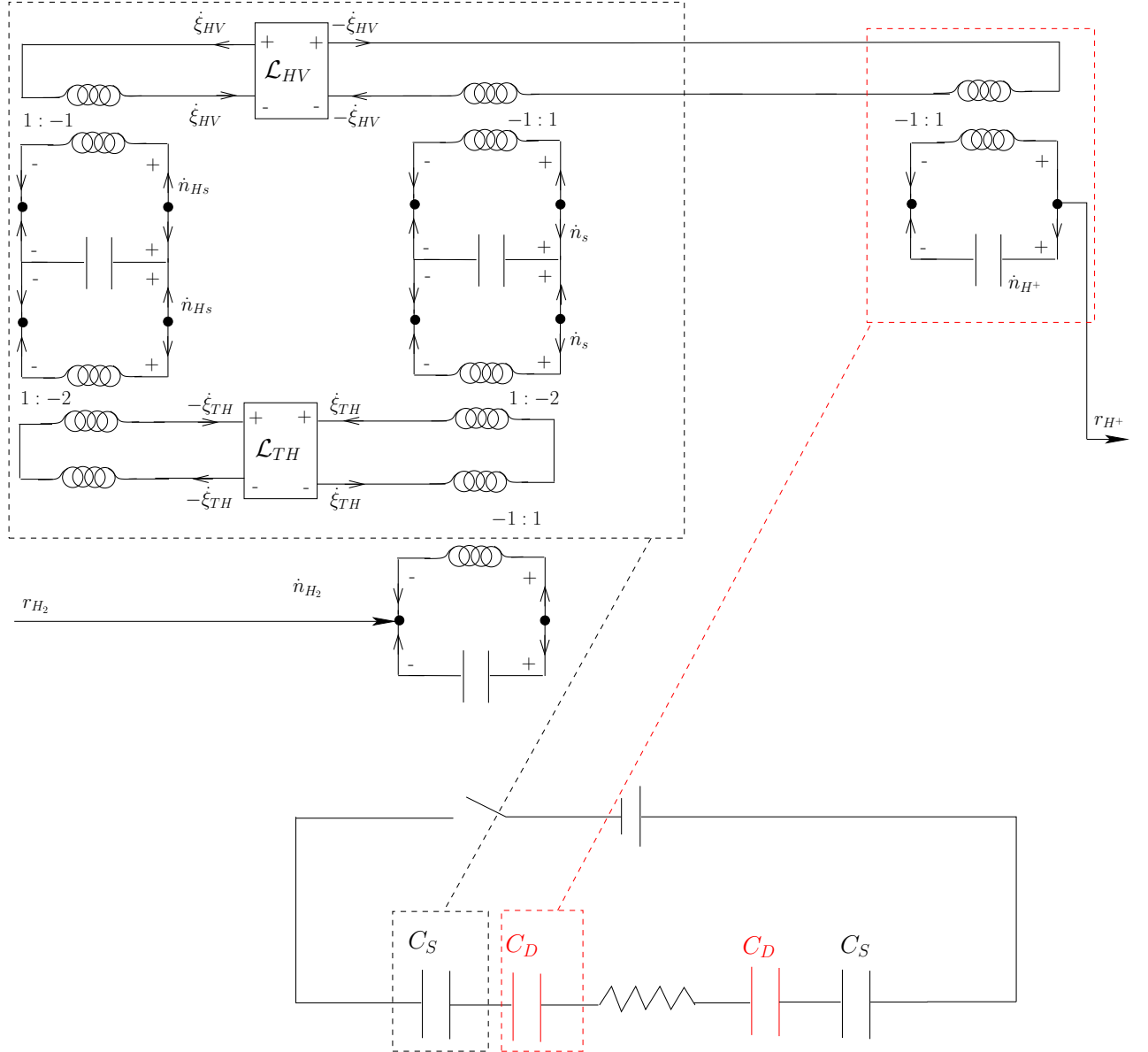
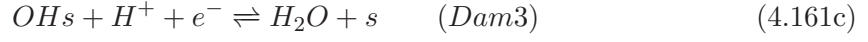
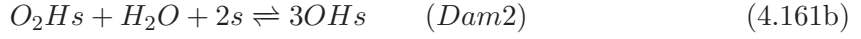
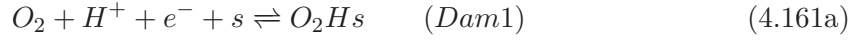


FIGURE 4.10 – En haut le schéma analogique détaillé du mécanisme réactionnel de Tafel-Heyrovsky-Volmer présenté sur la Figure 4.9. En bas l'emplacement de la couche de Stern et de la couche diffuse.

4.2.6 Mécanisme réactionnel de Damjanovic et Brusic

Le mécanisme réactionnel de Damjanovic et Brusic [32] permet de décrire le mécanisme de réduction de l'oxygène. Nous utilisons par la suite ce choix de mécanisme réactionnel, qui est adopté par Franco [37] (voir le Chapitre 3), pour le modèle de la couche compacte à la

cathode. Ce mécanisme réactionnel s'écrit sous la forme :



Nous sommes donc en présence de trois réactions écrites entre sept constituants indépendants. On note par $n_k, k \in \{O_2, O_2Hs, OHs, H_2O, s, H^+\}$ le nombre de moles de l'espèce k et par $\xi_j, j \in \{1, 2, 3\}$ l'avancement de la j -ème réaction. On a alors :

$$n_c = \begin{pmatrix} n_{O_2} \\ n_{O_2Hs} \\ n_{OHs} \\ n_{H_2O} \\ n_{H^+} \\ n_s \end{pmatrix}, \quad \xi_c = \begin{pmatrix} \xi_1 \\ \xi_2 \\ \xi_3 \end{pmatrix} \quad (4.162)$$

$$\nu_c = \begin{pmatrix} -1 & 0 & 0 \\ 1 & -1 & 0 \\ 0 & 3 & -1 \\ 0 & -1 & 1 \\ -1 & 0 & -1 \\ -1 & -2 & 1 \end{pmatrix}, \quad \nu_c^R = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \\ 1 & 2 & 0 \end{pmatrix}, \quad \nu_c^P = \begin{pmatrix} 0 & 0 & 0 \\ 1 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad (4.163)$$

Seules les espèces O_2, H^+, H_2O échangent avec l'extérieur. On suppose, comme pour l'anode que cela se fait par diffusion. En introduisant les vecteurs :

$$r_c = \begin{pmatrix} r_{O_2} \\ 0 \\ 0 \\ r_{H_2O} \\ -\frac{I_M}{F} \\ 0 \end{pmatrix}, \quad c_c^\top = (0 \quad 0 \quad 0 \quad 0 \quad 1 \quad 0) \quad (4.164)$$

le système d'état s'écrit donc, de la même manière que pour l'anode, sous la forme :

$$\dot{n}_c(t) = \nu_c \left[k_{Rc} \cdot \text{Exp} \left((\nu_c^R)^\top \left(-\frac{\alpha_c c_c^\top n_c}{C_{S0} RT} + \text{Ln} \left(\frac{n_c}{V} \right) \right) \right) - k_{Pc} \cdot \text{Exp} \left((\nu_c^P)^\top \left(\frac{\beta_c c_c^\top n_c}{C_{S0} RT} + \text{Ln} \left(\frac{n_c}{V} \right) \right) \right) \right] + r_c(n_c, I_M) \quad (4.165a)$$

$$r_c(n_c, I_M) = B_{n,c}(n_c^{ext} - n_c) + B_{I,c} I_M \quad (4.165b)$$

$$\Delta \phi_{Sc} = -\frac{1}{F C_{S0}} c_c^\top n_c \quad (4.165c)$$

où les vecteur α_c et β_c vérifient :

$$(\nu_c^R)^\top \alpha_c = \begin{pmatrix} \alpha_1 \\ 0 \\ \alpha_3 \end{pmatrix}, \quad (\nu_c^P)^\top \beta_c = \begin{pmatrix} 1 - \alpha_1 \\ 0 \\ 1 - \alpha_3 \end{pmatrix}$$

où α_1 , resp. α_3 est le coefficient de transfert électronique de l'étape réactionnelle de Dam1, resp. de Dam3 et où $B_{n,c}$ et $B_{I,c}$ vérifient :

$$B_{n,c} = -S_{CA} \text{diag} \{D_{O_2}^{cc}; 0; 0; D_{H_2O}^{cc}; 0; 0\}, \quad B_{I,c} = \frac{1}{F} \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 1 \\ 0 \end{pmatrix}. \quad (4.166)$$

On vérifie que la matrice stochioémétrique ν_c est de rang plein : $\text{rank}(\nu_c) = 3$. Par conséquent, il n'est pas possible de réduire le système des réactions. On a alors $\tilde{\nu}_c = \nu_c$, $N = 6$ et $M_\nu = 3$. On note :

$$n_{c,c} = \begin{pmatrix} n_{O_2} \\ n_{O_2Hs} \\ n_{OHs} \end{pmatrix}, \quad n_{\bar{c},c} = \begin{pmatrix} n_{H_2O} \\ n_{H^+} \\ n_s \end{pmatrix}, \quad \nu_{c,c} = \begin{pmatrix} -1 & 0 & 0 \\ 1 & -1 & 0 \\ 0 & 3 & -1 \end{pmatrix}, \quad \nu_{\bar{c},c} = \begin{pmatrix} 0 & -1 & 1 \\ -1 & 0 & -1 \\ -1 & -2 & 1 \end{pmatrix} \quad (4.167)$$

En introduisant le changement de variable d'état suivant :

$$\xi_{c,c} = \nu_{c,c}^{-1} n_{c,c}, \quad \tilde{n}_{\bar{c},c} = n_{\bar{c},c} - \nu_{\bar{c},c} \xi_{c,c} \quad (4.168)$$

nous écrivons, d'après (4.148), le système (4.165) sous la forme :

$$\begin{pmatrix} \dot{\xi}_{c,c} \\ \dot{\tilde{n}}_{\bar{c},c} \end{pmatrix} = \begin{pmatrix} \mathcal{L}(\xi_{c,c}, \tilde{n}_{\bar{c},c}) \\ 0_{3 \times 1} \end{pmatrix} + \begin{pmatrix} \nu_{c,c}^{-1} & 0_{3 \times 3} \\ -\nu_{\bar{c},c} \nu_{c,c}^{-1} & I_3 \end{pmatrix} \left[B_{nc} \begin{pmatrix} \nu_{c,c} & 0_{3 \times 3} \\ \nu_{\bar{c},c} & I_3 \end{pmatrix} \begin{pmatrix} \xi_{c,c}^{ext} - \xi_{c,c} \\ \tilde{n}_{\bar{c},c}^{ext} - \tilde{n}_{\bar{c},c} \end{pmatrix} + B_{I,c} I_M \right] \quad (4.169a)$$

$$\Delta \phi_{Sc} = -\frac{1}{FC_{S0}} c_c^\top \begin{pmatrix} \nu_{c,c} & 0_{3 \times 3} \\ \nu_{\bar{c},c} & I_3 \end{pmatrix} \begin{pmatrix} \xi_{c,c} \\ \tilde{n}_{\bar{c},c} \end{pmatrix} \quad (4.169b)$$

où

$$\xi_{c,c} = \begin{pmatrix} -n_{O_2} \\ -n_{O_2} - n_{O_2Hs} \\ -3n_{O_2} - 3n_{O_2Hs} - n_{OHs} \end{pmatrix}, \quad \tilde{n}_{\bar{c},c} = \begin{pmatrix} n_{H_2O} + 2n_{O_2} + 2n_{O_2Hs} + n_{OHs} \\ n_{H^+} - 4n_{O_2} - 3n_{O_2Hs} - n_{OHs} \\ n_s + n_{O_2Hs} + n_{OHs} \end{pmatrix} \quad (4.170)$$

Remarque 4.5. On déduit de (4.170) :

$$\begin{pmatrix} n_{H_2O} + 2n_{O_2} + 2n_{O_2Hs} + n_{OHs} \\ 2n_{H_2O} + n_{O_2Hs} + n_{OHs} \\ n_s + n_{O_2Hs} + n_{OHs} \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 2 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \tilde{n}_{\bar{c},c}$$

Dans la matrice précédente, la première, resp. la deuxième ligne représente la conservation de l'oxygène, resp. de l'hydrogène. Enfin la troisième ligne représente la conservation des sites catalytiques.

On effectue maintenant les calculs matriciels pour simplifier la forme prise ici par (4.169). $\nu_{c,c}, \nu_{\bar{c},c}$ étant définis par (4.168) et $B_{n,c}, B_{I,c}$ par (4.166), on trouve :

$$\begin{aligned}
 & \begin{pmatrix} \nu_{c,c}^{-1} & 0_{3 \times 3} \\ -\nu_{\bar{c},c} \nu_{c,c}^{-1} & I_3 \end{pmatrix} B_{n,c} \begin{pmatrix} \nu_{c,c} & 0_{3 \times 3} \\ \nu_{\bar{c},c} & I_3 \end{pmatrix} \\
 &= S_{CA} \left(\begin{array}{c|c} \nu_{c,c}^{-1} \begin{pmatrix} D_{O_2}^{cc} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} & 0_{3 \times 3} \\ \hline -\nu_{\bar{c},c} \nu_{c,c}^{-1} \begin{pmatrix} D_{O_2}^{cc} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} & \begin{pmatrix} D_{H_2O}^{cc} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \end{array} \right) \begin{pmatrix} \nu_{c,c} & 0_{3 \times 3} \\ \nu_{\bar{c},c} & I_3 \end{pmatrix} \\
 &= S_{CA} \left(\begin{array}{c|c} \nu_{c,c}^{-1} \begin{pmatrix} D_{O_2}^{cc} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \nu_{c,c} & 0_{3 \times 3} \\ \hline -\nu_{\bar{c},c} \nu_{c,c}^{-1} \begin{pmatrix} D_{O_2}^{cc} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \nu_{c,c} + \begin{pmatrix} D_{H_2O}^{cc} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \nu_{\bar{c},c} & \begin{pmatrix} D_{H_2O}^{cc} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \end{array} \right) \\
 &= S_{CA} \begin{pmatrix} D_{O_2}^{cc} & 0 & 0 & 0 & 0 & 0 \\ D_{O_2}^{cc} & 0 & 0 & 0 & 0 & 0 \\ 3D_{O_2}^{cc} & 0 & 0 & 0 & 0 & 0 \\ -2D_{O_2}^{cc} & -D_{H_2O}^{cc} & D_{H_2O}^{cc} & D_{H_2O}^{cc} & 0 & 0 \\ 4D_{O_2}^{cc} & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}
 \end{aligned}$$

et

$$\begin{pmatrix} \nu_{c,c}^{-1} & 0_{3 \times 3} \\ -\nu_{\bar{c},c} \nu_{c,c}^{-1} & I_3 \end{pmatrix} B_{I,c} = \frac{I_M}{F} \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ -1 \\ 0 \end{pmatrix}.$$

Enfin

$$\begin{aligned} c_c^\top \begin{pmatrix} \nu_{c,c} & 0_{3 \times 3} \\ \nu_{\bar{c},c} & I_3 \end{pmatrix} &= \begin{pmatrix} 0 & 0 & 0 & 0 & 1 & 0 \end{pmatrix} \begin{pmatrix} \nu_{c,c} & 0_{3 \times 3} \\ \nu_{\bar{c},c} & I_3 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 & 0 \end{pmatrix} \begin{pmatrix} \nu_{\bar{c},c} & I_3 \end{pmatrix} = \begin{pmatrix} 0 & -1 & 1 & 0 & 0 & 0 \end{pmatrix} \end{aligned}$$

Le système (4.169) s'écrit donc sous la forme :

$$\dot{\xi}_{c,c} = \begin{pmatrix} \mathcal{L}_1(\xi_{c,c}, \tilde{n}_{\bar{c},c}) \\ \mathcal{L}_2(\xi_{c,c}, \tilde{n}_{\bar{c},c}) \\ \mathcal{L}_3(\xi_{c,c}, \tilde{n}_{\bar{c},c}) \end{pmatrix} + S_{CA} \begin{pmatrix} D_{O_2}^{cc} & 0 & 0 \\ D_{O_2}^{cc} & 0 & 0 \\ 3D_{O_2}^{cc} & 0 & 0 \end{pmatrix} (\xi_{c,c}^{ext} - \xi_{c,c}) \quad (4.171a)$$

$$\begin{aligned} \dot{\tilde{n}}_{\bar{c},c} &= S_{CA} \begin{pmatrix} -2D_{O_2}^{cc} & -D_{H_2O}^{cc} & D_{H_2O}^{cc} \\ 4D_{O_2}^{cc} & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} (\xi_{c,c}^{ext} - \xi_{c,c}) \\ &\quad + S_{CA} \begin{pmatrix} D_{H_2O}^{cc} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} (\tilde{n}_{\bar{c},c}^{ext} - \tilde{n}_{\bar{c},c}) - \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \frac{I_M}{F} \quad (4.171b) \end{aligned}$$

$$\Delta\phi_{Sc} = -\frac{1}{FC_{S0}} \begin{pmatrix} 0 & -1 & 1 & 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} \xi_{c,c} \\ \tilde{n}_{\bar{c},c} \end{pmatrix} \quad (4.171c)$$

On rappelle (cf. (4.170)) que :

$$\xi_{c,c} = \begin{pmatrix} -n_{O_2} \\ -n_{O_2} - n_{O_2Hs} \\ -3n_{O_2} - 3n_{O_2Hs} - n_{OHS} \end{pmatrix}, \quad \tilde{n}_{\bar{c},c} = \begin{pmatrix} n_{H_2O} + 2n_{O_2} + 2n_{O_2Hs} + n_{OHS} \\ n_{H^+} - 4n_{O_2} - 3n_{O_2Hs} - n_{OHS} \\ n_s + n_{O_2Hs} + n_{OHS} \end{pmatrix},$$

et

$$\xi_{c,c}^{ext} = -\begin{pmatrix} 1 \\ 1 \\ 3 \end{pmatrix} n_{O_2}^{GDL}, \quad \tilde{n}_{\bar{c},c}^{ext} = \begin{pmatrix} n_{H_2O}^{GDL} + 2n_{O_2}^{GDL} \\ -4n_{O_2}^{GDL} \\ 0 \end{pmatrix}.$$

où $n_{O_2}^{GDL}$, resp $n_{H_2O}^{GDL}$ est le nombre de moles d'hydrogène, resp. le nombre de moles d'eau à l'interface GDL/CC.

4.3 Modèle de la double couche et modèle de pile

On utilise ici le modèle de double couche avec anions fixes développé en Annexe A. Ce modèle est présenté avec de simples réactions d'oxydo-réduction, du type $H_2 \rightleftharpoons 2H^+ + 2e^-$ à l'anode et $O_2 + 4H^+ \rightleftharpoons 2H_2O$ à la cathode.

4.3.1 Modèle de la double couche

On utilise les notations de l'Annexe A. Le modèle de double couche s'obtient à partir des relations (A.5), (A.4), (A.22) et (A.25).

En définissant $\Delta\phi_{dl} = \Delta\phi_S + \Delta\phi_D$, on obtient la :

Propriété 4.1. *Le modèle de la double couche s'écrit sous la forme suivante :*

$$C_{S0} \frac{\partial}{\partial t} \Delta\phi_S = -\tilde{\mathcal{R}}(c_R, \Delta\phi_S) + J_M \quad (4.172a)$$

$$\Delta\phi_{dl} = \Delta\phi_S + \Phi_D(\Delta\phi_S) \quad (4.172b)$$

$$J = \tilde{\mathcal{R}}(c_R, \Delta\phi_S) \quad (4.172c)$$

où

$$\tilde{\mathcal{R}}(c_R, \Delta\phi_S) = k_{OCR} \exp\left(\frac{\alpha_O F \Delta\phi_S}{RT}\right) - k_{RC\infty} \left(1 + \frac{1}{2} \left(\frac{F \Delta\phi_S}{\delta RT}\right)^2\right) \exp\left(-\frac{\alpha_R F \Delta\phi_S}{RT}\right) \quad (4.173)$$

$$\Phi_D(\Delta\phi_S) = \frac{RT}{zF} \ln \left(1 + \frac{1}{2} \left(\frac{zF \Delta\phi_S}{\delta RT}\right)^2\right) \operatorname{sgn} \Delta\phi_S \quad (4.174)$$

D'après la Propriété A.5, les fonctions J et $\Delta\phi_{dl}$ dépendent continument des entrées $J_M \geq 0$ et $c_R > 0$.

Pour J_M suffisamment petit, ce modèle a un point d'équilibre localement stable.

On rappelle le sens des variables : $\Delta\phi_S$, resp. $\Delta\phi_D$ représentent la différence de potentiel aux bornes de la couche de Stern, resp. de la couche diffuse; c_R est la concentration de réducteur (H_2) à l'anode.

Par rapport aux modèles développés dans la Section 4.2, le terme de correction de Frumkin est ici pris en compte. Pour intégrer ce dernier dans modèles électrochimiques plus élaborés de la Section 4.2, il suffit :

- de remplacer la concentration de protons par l'expression : $c_\infty \left(1 + \frac{1}{2} \left(\frac{F \Delta\phi_S}{\delta RT}\right)^2\right)$
- de considérer la différence de potentiel total $\Delta\phi_S + \Delta\phi_D = \Delta\phi_S + \Phi_D(\Delta\phi_S)$ à la place de $\Delta\phi_S$ dans l'évaluation de la tension aux bornes de la double couche.
- de prendre en compte l'évolution de la différence de potentiel aux bornes de la couche de Stern, cf. (4.172a).

Remarque : le fait que ce modèle soit causal en prenant comme entrées c_R et J_M , permet de représenter par exemple des essais harmoniques ou le relevé de la courbe de polarisation de cet élément car on peut calculer les sorties utiles que sont alors J et $\Delta\phi_{dl}$. Avec le choix fait, $\Delta\phi_S$ est une variable d'état mais d'autres variables peuvent être nécessaires pour réaliser la fonctionnelle $\tilde{\mathcal{R}}(c_R, \Delta\phi_S)$ si on prend un mécanisme réactionnel plus complexe.

4.3.2 Modèle de la pile

On l'obtient en instanciant le modèle à l'anode et à la cathode et en fermant le circuit extérieur sur une impédance Z_{ext} pouvant être variable afin de disposer d'une entrée sur le circuit extérieur.

Dans ce cas, ni J_M ni J ne sont connus : les modèles des doubles couches ne sont que des relations "courants (J, J_M) - tension ($\Delta\phi_{dl}$)". Les courants seront déterminés à l'aide de deux relations supplémentaires traduisant que :

1/ le courant membrane est supposé le même pour chaque double couche ;

2/ le circuit extérieur impose une relation entre U et J . On la supposera linéaire, de la forme $U = Z_{ext} * J$, qui se réduit à un produit usuel pour une simple résistance : $U = R_{ext}J$.

En pratique, on pourra régler le circuit extérieur (ici choisir Z_{ext} par exemple) afin d'avoir le courant J désiré dans certaines limites. Avant de déterminer les courants atteignables, il faut vérifier que pour tout Z_{ext} "raisonnable" donné, $J, J_M, \dots$ existent.

Pour simplifier les notations, définissons les fonctions

$$\Phi_{dli}(\Delta\phi) = \Delta\phi + \Phi_{Di}(\Delta\phi), \quad \text{pour } i = a, c \quad (4.175)$$

Cela donne :

$$-C_{S0a} \frac{\partial}{\partial t} \Delta\phi_{Sa} = \tilde{\mathcal{R}}_a(c_{Ra}, \Delta\phi_{Sa}) - J_M \quad (4.176a)$$

$$C_{S0c} \frac{\partial}{\partial t} \Delta\phi_{Sc} = -\tilde{\mathcal{R}}_c(c_{Oc}, \Delta\phi_{Sc}) + J_M \quad (4.176b)$$

$$U = \Phi_{dlc}(\Delta\phi_{Sc}) - R_M J_M - \Phi_{dla}(\Delta\phi_{Sa}) = Z_{ext} * J \quad (4.176c)$$

$$J = \tilde{\mathcal{R}}_a(c_{Ra}, \Delta\phi_{Sa}) = \tilde{\mathcal{R}}_c(c_{Oc}, \Delta\phi_{Sc}) \quad (4.176d)$$

Les entrées sont les alimentations de la pile représentées ici par c_{Ra} et c_{Oc} et les sorties U et J . Vérifions qu'il y a bien une relation causale entre elles. Ajoutant (4.176a) et (4.176b), on déduit que $C_{S0c} \Delta\phi_{Sc} - C_{S0a} \Delta\phi_{Sa}$ est constant, ce qui conduit à introduire la capacité totale des couches de Stern et les facteurs de symétrie suivants :

$$C_{S0} = \frac{C_{S0a} C_{S0c}}{C_{S0a} + C_{S0c}}, \quad \beta_a = \frac{C_{S0a}}{C_{S0a} + C_{S0c}}, \quad \beta_c = \frac{C_{S0c}}{C_{S0a} + C_{S0c}} \quad (4.177)$$

et à écrire cette constante ainsi : $-\beta_a \Delta\phi_{Sa} + \beta_c \Delta\phi_{Sc} = \Delta\phi_{S0}$.

Multipliant maintenant (4.176a) par C_{S0c} et (4.176b) par C_{S0a} , soustrayant et divisant par $C_{S0a} + C_{S0c}$, on obtient

$$C_{S0} \frac{\partial}{\partial t} (\Delta\phi_{Sa} + \Delta\phi_{Sc}) = -\tilde{\mathcal{R}}_a(c_{Ra}, \Delta\phi_{Sa}) + J_M = -\tilde{\mathcal{R}}_c(c_{Oc}, \Delta\phi_{Sc}) + J_M$$

ce qui conduit à définir la nouvelle variable : $\Delta\phi_S = \Delta\phi_{Sc} + \Delta\phi_{Sa}$. Le changement de variables proposé s'écrit donc :

$$\Delta\phi_{S0} = -\beta_a \Delta\phi_{Sa} + \beta_c \Delta\phi_{Sc} \quad (4.178a)$$

$$\Delta\phi_S = \Delta\phi_{Sc} + \Delta\phi_{Sa} \quad (4.178b)$$

Le changement de variable réciproque s'écrit :

$$\Delta\phi_{Sa} = -\Delta\phi_{S0} + \beta_c \Delta\phi_S \quad (4.179)$$

$$\Delta\phi_{Sc} = \Delta\phi_{S0} + \beta_a \Delta\phi_S \quad (4.180)$$

on a :

$$C_{S0} \frac{\partial}{\partial t} \Delta \phi_S = -\tilde{\mathcal{R}}_a(c_{Ra}, -\Delta \phi_{S0} + \beta_c \Delta \phi_S) + J_M = -\tilde{\mathcal{R}}_c(c_{Oc}, \Delta \phi_{S0} + \beta_a \Delta \phi_S) + J_M \quad (4.181)$$

Nous allons maintenant étudier l'équation suivante en $\Delta \phi_S$ pour un J donné :

$$\left(\tilde{\mathcal{R}}_a(c_{Ra}, -\Delta \phi_{S0} + \beta_c \Delta \phi_S) - J \right)^2 + \left(\tilde{\mathcal{R}}_c(c_{Oc}, \Delta \phi_{S0} + \beta_a \Delta \phi_S) - J \right)^2 = 0 \quad (4.182)$$

On doit trouver un $\Delta \phi_{S0}$ tel que pour tout J suffisamment petit, elle ait une solution $\Delta \phi_S$.

Pour $J = 0$, on a $\tilde{\mathcal{R}}_a(c_{Ra}, -\Delta \phi_{S0} + \beta_c \Delta \phi_S) = \tilde{\mathcal{R}}_c(c_{Oc}, \Delta \phi_{S0} + \beta_a \Delta \phi_S) = 0$ et nous avons vu (Propriété A.5) que l'on a alors

$$\Delta \phi_{S0a} = -\frac{RT}{zF} L_{\delta_a} \left(\frac{k_{Oa} c_{Ra}}{k_{Ra} c_{\infty}} \right), \quad \Delta \phi_{S0c} = -\frac{RT}{zF} L_{\delta_c} \left(\frac{k_{Rc} c_{Oc}}{k_{Oc} c_{\infty}} \right)$$

On a donc nécessairement $\Delta \phi_{S0}$ donné par

$$\Delta \phi_{S0} = -\beta_a \Delta \phi_{S0a} + \beta_c \Delta \phi_{S0c} = -\frac{RT}{zF} \left(-\beta_a L_{\delta_a} \left(\frac{k_{Oa} c_{Ra}}{k_{Ra} c_{\infty}} \right) + \beta_c L_{\delta_c} \left(\frac{k_{Rc} c_{Oc}}{k_{Oc} c_{\infty}} \right) \right) \quad (4.183)$$

et il vérifie

$$\tilde{\mathcal{R}}_a(c_{Ra}, -\Delta \phi_{S0} + \beta_c \Delta \phi_S^0) = \tilde{\mathcal{R}}_c(c_{Oc}, \Delta \phi_{S0} + \beta_a \Delta \phi_S^0) = 0$$

où

$$\Delta \phi_S^0 = \Delta \phi_{Sc0} + \Delta \phi_{Sa0} = -\frac{RT}{zF} \left(L_{\delta_a} \left(\frac{k_{Oa} c_{Ra}}{k_{Ra} c_{\infty}} \right) + L_{\delta_c} \left(\frac{k_{Rc} c_{Oc}}{k_{Oc} c_{\infty}} \right) \right) \quad (4.184)$$

de sorte que $\Delta \phi_{Sa0} = -\Delta \phi_{S0} + \beta_c \Delta \phi_S^0$ et $\Delta \phi_{Sc0} = \Delta \phi_{S0} + \beta_a \Delta \phi_S^0$.

On a donc montré que, $\Delta \phi_{S0}$ étant donné par (4.183), pour $J = 0$ l'équation (4.182) a une solution $\Delta \phi_S^0$. Pour ce même paramètre $\Delta \phi_{S0}$, (4.182) a donc encore une solution $\Delta \phi_S$ pour J suffisamment petit. On peut aussi interpréter ce résultat ainsi :

$\Delta \phi_{S0}$ étant donné par (4.183), pour $\Delta \phi_S$ dans un voisinage de $\Delta \phi_S^0$ donné par (4.184), on a

$$\tilde{\mathcal{R}}_a(c_{Ra}, -\Delta \phi_{S0} + \beta_c \Delta \phi_S) = \tilde{\mathcal{R}}_c(c_{Oc}, \Delta \phi_{S0} + \beta_a \Delta \phi_S)$$

Cette valeur commune définit une fonction $\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S)$.

Par ailleurs, avec (4.176c), on a $\Phi_{dlc}(\Delta \phi_{Sc}) - R_M J_M - \Phi_{dla}(\Delta \phi_{Sa}) = Z_{ext} * J$, c'est à dire :

$$R_M J_M = \Phi_{dlc}(\Delta \phi_{Sc}) - \Phi_{dla}(\Delta \phi_{Sa}) - Z_{ext} * J$$

soit encore, en fonction de la variable $\Delta \phi_S$:

$$R_M J_M = \Phi_{dlc}(\Delta \phi_{S0} + \beta_a \Delta \phi_S) - \Phi_{dla}(-\Delta \phi_{S0} + \beta_c \Delta \phi_S) - Z_{ext} * \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S)$$

Finalement, on a donc montré le résultat suivant :

Propriété 4.2. *Le système pile vérifie :*

i) $\Delta \phi_{S0}$ étant une fonction de c_{Ra} et c_{Oc} donnée par (4.183), pour $\Delta \phi_S$ dans un voisinage de $\Delta \phi_S^0$, lui aussi fonction de c_{Ra} et c_{Oc} donnée par (4.184), on a

$$\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S) := \tilde{\mathcal{R}}_a(c_{Ra}, -\Delta \phi_{S0} + \beta_c \Delta \phi_S) = \tilde{\mathcal{R}}_c(c_{Oc}, \Delta \phi_{S0} + \beta_a \Delta \phi_S)$$

ii) Le système pile s'écrit sous la forme causale suivante dans laquelle les entrées sont c_{Ra} et c_{Oc} et les sorties J et U :

$$C_{S0} \frac{\partial}{\partial t} \Delta\phi_S = -\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S) + J_M \quad (4.185a)$$

$$J_M = \frac{1}{R_M} [\Phi_{dlc}(\Delta\phi_{S0} + \beta_a \Delta\phi_S) - \Phi_{dla}(-\Delta\phi_{S0} + \beta_c \Delta\phi_S) - U_{cell}] \quad (4.185b)$$

$$J = \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S) \quad (4.185c)$$

$$U_{cell} = Z_{ext} * \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S) \quad (4.185d)$$

On a de plus

$$\Delta\phi_{Sc} = \Delta\phi_{S0} + \beta_a \Delta\phi_S, \quad \Delta\phi_{Sa} = -\Delta\phi_{S0} + \beta_c \Delta\phi_S \quad (4.185e)$$

$$\Delta\phi_{Di} = \Phi_{Di}(\Delta\phi_{Si}), \quad \text{pour } i = a, c. \quad (4.185f)$$

On a donc

$$\begin{aligned} C_{S0} \frac{\partial}{\partial t} \Delta\phi_S &= -\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S) + J_M \\ &= -\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S) + \frac{1}{R_M} \left(\Phi_{dlc}(\Delta\phi_{S0} + \beta_a \Delta\phi_S) - \Phi_{dla}(-\Delta\phi_{S0} + \beta_c \Delta\phi_S) - U_{cell} \right) \\ &= -\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S) \\ &\quad + \frac{1}{R_M} \left(\Phi_{dlc}(\Delta\phi_{S0} + \beta_a \Delta\phi_S) - \Phi_{dla}(-\Delta\phi_{S0} + \beta_c \Delta\phi_S) - Z_{ext} * \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S) \right) \end{aligned}$$

d'où :

$$\begin{aligned} C_{S0} \frac{\partial}{\partial t} \Delta\phi_S &= -\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S) \\ &\quad + \frac{1}{R_M} \left(\Phi_{dlc}(\Delta\phi_{S0} + \beta_a \Delta\phi_S) - \Phi_{dla}(-\Delta\phi_{S0} + \beta_c \Delta\phi_S) - Z_{ext} * \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S) \right) \end{aligned} \quad (4.186)$$

On a vu que $\Delta\phi_{S0}$ est constant : l'équation précédente définit donc complètement l'évolution de $\Delta\phi_S$ (dès qu'une condition initiale est donnée). Les sorties sont alors données par (4.185c) et (4.185d).

4.4 Conclusion

Compte tenu du fait que les mécanismes réactionnels dans les couches compactes de la pile mettent en jeu plusieurs étapes intermédiaires, la modélisation de réseaux de réactions électrochimiques est nécessaire dans le cadre de notre étude. Nous nous sommes inspirés de travaux classiques sur l'analyse géométrique de réseaux de réactions électrochimiques pour construire un modèle de réseaux électrochimiques réduit et compatible avec la thermodynamique. Nous avons utilisé ensuite le mécanisme réactionnel de Tafel-Heyrovsky-Volmer pour modéliser l'oxydation de l'hydrogène sur les sites catalytiques à l'anode de la pile. Côté cathode, nous avons utilisé le mécanisme réactionnel de Damjanovic et Brusic pour modéliser la réduction de l'oxygène. La surtension dans la couche diffuse très souvent omise dans la modélisation de la pile à été incluse dans notre modèle moyennant la correction de Frumkin dans la formule de Butler-Volmer. Ce calcul à été détaillé dans l'Annexe A.

Chapitre 5

Comportement harmonique des réseaux électrochimiques

L'objectif de ce chapitre est d'utiliser la technique de la balance harmonique pour calculer l'impédance du réseau de réactions électrochimiques et le modèle de la pile obtenus au Chapitre 4. Dans un premier temps, nous introduisons le principe de cette technique et indiquons comment calculer l'impédance d'un système dynamique non linéaire entrée-sortie. Nous appliquons ensuite cette méthode au modèle du réseau de réactions électrochimiques et le modèle de la pile. Ceci permet d'obtenir l'impédance du réseau et de la pile avec leur expression analytique présentée dans la Section 5.2. Nous calculons finalement dans la Section 5.6 l'impédance du GDL et nous fournissons le schéma électrique associé moyennant une décomposition quadripolaire.

5.1 Méthode de la balance harmonique pour le calcul analytique d'une impédance

5.1.1 Principe de la méthode

La méthode de la balance harmonique [43] est basée sur une analyse qui permet d'approcher la sortie périodique d'un système non-linéaire soumis à une entrée périodique, en se limitant aux premières harmoniques. Dans le cadre de l'analyse d'un système de commande non linéaire, l'idée de base de cette méthode est de remplacer la non-linéarité d'un système par son gain complexe équivalent (GCE). Le GCE est défini comme étant le fondamental de la sortie divisé par le fondamental de l'entrée. Dans ce contexte, le GCE est considéré comme l'équivalent de la fonction de transfert dans le cas des systèmes non linéaires.

Dans le cas de mesures d'impédance, au cours desquelles l'effet d'un faible signal sinusoïdal est observé et analysé, cet outil de "linéarisation" est tout à fait adapté.

5.1.2 Gain complexe équivalent (GCE)

L'application de la méthode de la balance harmonique suppose que le système étudié puisse être décomposé en une partie non linéaire suivie d'une partie linéaire se comportant comme un filtre passe-bas efficace. Dans ces conditions, pour une entrée sinusoïdale, la sortie du système en générale n'est pas sinusoïdale. On admet toutefois qu'elle est périodique de même période que l'excitation. Ainsi, la sortie renferme, en plus du premier harmonique, des harmoniques d'ordres supérieurs supposés parfaitement éliminés par le filtre passe-bas. Cette idée est explorée dans les lignes qui suivent.

On considère l'élément non linéaire suivant, dont l'entrée est $x(t) = X \sin(\omega t)$ (Figure

5.1). Sa sortie $y(t) = f(X \sin(\omega t))$, étant un signal périodique de période $T = \frac{2\pi}{\omega}$, où f est

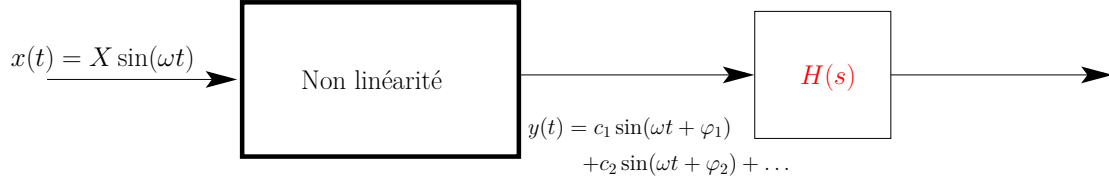


FIGURE 5.1 – Élément non linéaire suivi d'un système linéaire $H(s)$ (filtre passe-bas).

une fonction non linéaire, peut être décomposée en série de Fourier :

$$\begin{aligned} y(t) &= \frac{a_0}{2} + \sum_{n=1}^{\infty} (a_n \cos(n\omega t) + b_n \sin(n\omega t)) \\ &= \frac{a_0}{2} + \sum_{n=1}^{\infty} c_n \sin(n\omega t + \phi_n) \end{aligned}$$

avec

$$\begin{aligned} a_n &= \frac{\omega}{\pi} \int_0^T y(t) \cos(n\omega t) dt, \quad n = 0, 1, 2, \dots \\ b_n &= \frac{\omega}{\pi} \int_0^T y(t) \sin(n\omega t) dt, \quad n = 1, 2, \dots \\ c_n &= \sqrt{a_n^2 + b_n^2}, \quad \phi_n = \arctan \frac{a_n}{b_n}, \quad n = 1, 2, \dots \end{aligned}$$

où $T = \frac{2\pi}{\omega}$. Partant du principe que les harmoniques d'ordre 2, 3, ... sont filtrées par la partie linéaire (notée $H(s)$ dans la Figure 5.1), seul le premier harmonique $a_1 \cos(\omega t) + b_1 \sin(\omega t) = c_1 \sin(\omega t + \phi_1)$ est conservé. On remarque que ce premier harmonique est, par rapport à l'entrée $x(t) = X \sin(\omega t)$, atténué par le facteur $\frac{c_1}{X}$ et déphasé de la quantité ϕ_1 . Nous sommes donc amené à définir la *fonction de transfert généralisée* ou le *gain complexe équivalent (GCE)* comme suit :

$$N(X, \omega) = \frac{c_1}{X} \exp(i\phi_1). \quad (5.1)$$

Le schéma fonctionnel de la Figure 5.1 est par conséquent remplacé par celui de la Figure 5.2.

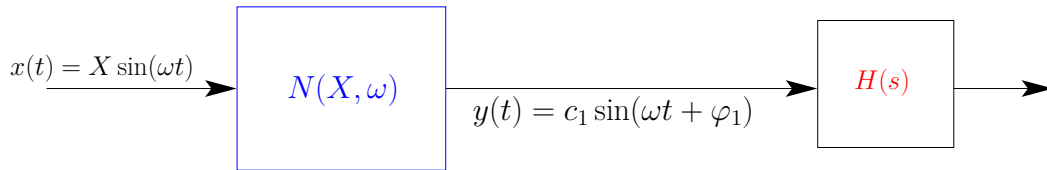


FIGURE 5.2 – Approximation d'une non linéarité par la méthode du premier harmonique.

On peut alors conclure que le GCE est l'équivalent de la fonction de transfert dans le cas de système non linéaire tant que l'entrée est sinusoïdale.

5.1.3 Application du méthode du gain complexe équivalent sur les mesures d'impédance

La méthode du premier harmonique permet d'approcher la sortie périodique d'un système différentiel non-linéaire soumis à une entrée périodique. En se limitant aux premières harmoniques, on peut alors calculer le GCE analytiquement et le tracer dans le plan de Nyquist. Cette idée est illustrée dans les lignes qui suivent.

On considère le système différentiel suivant :

$$\dot{x}(t) = F(x(t), u(t)), \quad y(t) = G(x(t), u(t)) \quad (5.2)$$

où x est la variable d'état et y la sortie. F et G sont deux fonctions non linéaires. u est l'entrée du système supposée sinusoïdale de la forme :

$$u(t) = u_0 + u_1 \exp(j\omega t), \quad |u_1| \ll |u_0|.$$

Dans le cadre de l'approximation par les premières harmoniques, on cherche les signaux solutions sous la forme

$$\begin{cases} x(t) & \approx x_0 + x_1 \exp(j\omega t + j\phi_x) \\ y(t) & \approx y_0 + y_1 \exp(j\omega t + j\phi_y) \end{cases}$$

selon la méthode de la balance harmonique, présentée dans la section précédente, on a :

$$x_0 = \frac{\omega}{2\pi} \int_0^{\frac{2\pi}{\omega}} x(t) dt, \quad x_1 \exp(j\phi_x) = \frac{\omega}{\pi} \int_0^{\frac{2\pi}{\omega}} x(t) \exp(-j\omega t) dt$$

et d'une manière similaire on a y_0 et y_1 . Par conséquent, le développement limité d'ordre un du système (5.2), permet d'exprimer analytiquement l'expression du GCE, $Z(j\omega)$, comme suit :

$$F(x_0, u_0) = 0 \quad (5.3a)$$

$$y_0 = G(x_0, u_0) \quad (5.3b)$$

$$Z(j\omega) \doteq \frac{y_1 \exp(j\phi_y)}{u_1} = \nabla_x G(x_0, u_0) (j\omega I_d - \nabla_x F(x_0, u_0))^{-1} \nabla_u F(x_0, u_0) + \nabla_u G(x_0, u_0). \quad (5.3c)$$

La formule (5.3c) peut tout aussi bien être vue comme l'impédance entrée sortie. Contrairement à la fonction de transfert usuelle, elle est liée au point de fonctionnement particulier envisagé.

5.2 Application à un réseau de réactions électrochimiques

Le système d'équations différentielles obtenu dans la Section 4.2.4 du Chapitre 4 (cf. (4.143a)) s'écrit sous la forme :

$$\begin{aligned} \dot{n} = \nu & \left[k_R \cdot \text{Exp} \left((\nu^R)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \cdot \text{Exp} \left(\frac{(\nu^R)^\top \alpha F \Delta \phi_S}{RT} \right) \right. \\ & \left. - k_P \cdot \text{Exp} \left((\nu^P)^\top \text{Ln} \left(\frac{n}{V} \right) \right) \cdot \text{Exp} \left(-\frac{(\nu^P)^\top \beta F \Delta \phi_S}{RT} \right) \right] + r(n, I_M) \quad (5.4) \end{aligned}$$

ou encore (cf. (4.148a))

$$\begin{pmatrix} \dot{\xi}_c \\ \dot{\tilde{n}}_{\bar{c}} \end{pmatrix} = \begin{pmatrix} \Pi\mathcal{L}(\xi_c, \tilde{n}_{\bar{c}}) \\ 0_{N-M_\nu} \end{pmatrix} + \begin{pmatrix} \nu_c^{-1} & 0_{M_\nu \times (N-M_\nu)} \\ -\nu_{\bar{c}}\nu_c^{-1} & I_{N-M_\nu} \end{pmatrix} \left[B_n \begin{pmatrix} \nu_c & 0_{M_\nu \times (N-M_\nu)} \\ \nu_{\bar{c}} & I_{N-M_\nu} \end{pmatrix} \begin{pmatrix} \xi_c^{ext} - \xi_c \\ \tilde{n}_{\bar{c}}^{ext} - \tilde{n}_{\bar{c}} \end{pmatrix} + B_I I_M \right] \quad (5.5)$$

où $\mathcal{L}$ est donné par (4.145).

Courbe de polarisation. Dans le cas stationnaire on fixe $I = I_0$, $\xi_c^{ext} = \xi_{c0}^{ext}$, $\tilde{n}_{\bar{c}}^{ext} = \tilde{n}_{\bar{c}0}^{ext}$ où $I_0, \xi_{c0}^{ext}, \tilde{n}_{\bar{c}0}^{ext}$ sont des constantes de $\mathbb{R}$, $\mathbb{R}^{M_\nu}$ et $\mathbb{R}^{N-M_\nu}$. En utilisant (5.5) on doit avoir à l'équilibre $\forall t \geq 0$:

$$\dot{\xi}_c = 0 \quad (5.6a)$$

$$\dot{\tilde{n}}_{\bar{c}} = 0 \quad (5.6b)$$

On cherche à déterminer les états d'équilibre $\begin{pmatrix} \xi_{c0} \\ \tilde{n}_{\bar{c}0} \end{pmatrix} \in \mathbb{R}^N$. Ceux-ci sont donc les solutions du système :

$$\begin{pmatrix} \Pi\mathcal{L}(\xi_{c0}, \tilde{n}_{\bar{c}0}) \\ 0_{N-M_\nu} \end{pmatrix} = - \begin{pmatrix} \nu_c^{-1} & 0_{M_\nu \times (N-M_\nu)} \\ -\nu_{\bar{c}}\nu_c^{-1} & I_{N-M_\nu} \end{pmatrix} \left[B_n \begin{pmatrix} \nu_c & 0_{M_\nu \times (N-M_\nu)} \\ \nu_{\bar{c}} & I_{N-M_\nu} \end{pmatrix} \begin{pmatrix} \xi_{c0}^{ext} - \xi_{c0} \\ \tilde{n}_{\bar{c}0}^{ext} - \tilde{n}_{\bar{c}0} \end{pmatrix} + B_I I_{M0} \right] \quad (5.7)$$

Théorème 5.8. Pour tout $I_{M0}, \xi_{c0}^{ext}, \tilde{n}_{\bar{c}0}^{ext}$, les états d'équilibres sont, s'ils existent, les solutions $(\xi_{c0}, \tilde{n}_{\bar{c}0})$ de l'équation (5.7) telles que toutes les composantes de

$$n_0 = \begin{pmatrix} \nu_c \xi_{c0} \\ \nu_{\bar{c}} \xi_{c0} + \tilde{n}_{\bar{c}0} \end{pmatrix}$$

sont positives ou nulles.

Impédance. On supposera ici que l'entrée courant $I(t)$ est égale à

$$I(t) = I_{M0} + I_{M1} \exp(j\omega t), \quad t \geq 0 \quad (5.8)$$

pour I_{M0} et I_{M1} fixés et que les valeurs de $\xi_c^{ext}, \tilde{n}_{\bar{c}}^{ext}$ restent constantes :

$$\xi_c^{ext}(t) = \xi_{c0}^{ext}, \quad \tilde{n}_{\bar{c}}^{ext}(t) = \tilde{n}_{\bar{c}0}^{ext}, \quad t \geq 0$$

pour $\xi_{c0}^{ext} \in \mathbb{R}^{M_\nu}$ et $\tilde{n}_{\bar{c}0}^{ext} \in \mathbb{R}^{N-M_\nu}$ donnés. Dans le cadre de l'approximation par les premières

harmoniques, on suppose que $n_0 = \begin{pmatrix} \nu_c \xi_{c0} \\ \nu_{\bar{c}} \xi_{c0} + \tilde{n}_{\bar{c}0} \end{pmatrix}$ est un état d'équilibre et on cherche la solution sous la forme approchée :

$$n(t) = n_0 + n_1 \exp(j\omega t). \quad (5.9)$$

L'impédance que l'on cherche s'écrit sous la forme :

$$Z(j\omega) = \frac{\Delta\phi_{S1}}{I_1} \quad (5.10)$$

où la première harmonique $\Delta\phi_{S1}$ vérifie :

$$\Delta\phi_{S1} = -\frac{1}{FC_{S0}}c^\top n_1. \quad (5.11)$$

Les calculs peuvent être menés indifféremment à partir de (5.5) ou de (5.4). Pour plus de simplicité, on part ici de (5.4). Le principe de la méthode de la balance harmonique consiste, en remplaçant (5.9) et (5.8) dans (5.4), à écrire :

$$j\omega n_1 = \frac{1}{2\pi}\nu \int_0^{2\pi} \Lambda\left(\Delta\tilde{G}^R(n_0 + n_1 \exp(jt)), \Delta\tilde{G}^P(n_0 + n_1 \exp(jt))\right) \exp(-jt)dt \\ - B_n n_1 + B_I I_{M1} \quad (5.12)$$

On pose maintenant par simplicité :

$$F^R(n) = k_R \cdot \text{Exp}\left[(\nu^R)^\top \left(-\frac{\alpha c^\top n}{C_{S0}RT} + \text{Ln}\left(\frac{n}{V}\right)\right)\right] \quad (5.13a)$$

$$F^P(n) = k_P \cdot \text{Exp}\left[(\nu^P)^\top \left(\frac{\beta c^\top n}{C_{S0}RT} + \text{Ln}\left(\frac{n}{V}\right)\right)\right] \quad (5.13b)$$

L'intégrale contenue dans (5.12) s'écrit encore :

$$\frac{1}{2\pi} \int_0^{2\pi} \left[F^R(n_0 + n_1 \exp(jt)) - F^P(n_0 + n_1 \exp(jt))\right] \exp(-jt)dt \quad (5.14)$$

on l'approxime par :

$$\left(\nabla F^R\Big|_{n_0} - \nabla F^P\Big|_{n_0}\right) n_1 \quad (5.15)$$

On montre sans difficulté les deux résultats suivants :

Lemme 5.9. Soit $V \in \mathbb{R}, A \in \mathbb{R}^{M \times N}, B \in \mathbb{R}^M$. Pour tout $n_0, dn_0 \in \mathbb{R}^N$, on a :

$$\text{Ln}\left(\frac{n_0 + dn_0}{V}\right) = \text{Ln}\left(\frac{n_0}{V}\right) + \frac{dn_0}{n_0} + o(dn_0) \quad (5.16)$$

$$\text{Exp}(A(n_0 + dn_0) + B) = \text{Exp}(An_0 + B) + (Adn_0) \cdot \text{Exp}(An_0 + B) + o(dn_0) \quad (5.17)$$

où le terme o est tel que :

$$\left|\frac{o(dn_0)}{dn_0}\right| \xrightarrow{\omega \rightarrow +\infty} 0.$$

On rappelle que dans (5.16) et (5.17), les quotients des vecteurs sont calculés terme à terme, de même que les fonctions Ln et Exp. On déduit du lemme précédent :

$$F^R(n_0 + dn_0) = F^R(n_0) - F^R(n_0) \cdot \frac{\alpha F}{C_{S0}RT} c^\top dn_0 + F^R(n_0) \cdot \left((\nu^R)^\top \frac{dn_0}{n_0}\right) + o(dn_0) \quad (5.18)$$

Lemme 5.10. Pour tout $A \in \mathbb{R}^{n \times p}$, $u \in \mathbb{R}^p$ et $v \in \mathbb{R}^n$, on a :

$$Au \cdot v = \text{diag}\{v\}Au$$

$$\text{où par définition } \text{diag}\{v\} = \begin{pmatrix} v_1 & & \\ & \ddots & \\ & & v_n \end{pmatrix}.$$

On déduit de (5.18) et du Lemme 5.10

$$\begin{aligned} \nabla F^R \Big|_{n_0} n_1 &= F^R(n_0) \cdot \left((\nu^R)^\top \left(-\frac{1}{C_{S_0}RT} \alpha c^\top n_1 + \frac{n_1}{n_0} \right) \right) \\ &= F^R(n_0) \cdot \left((\nu^R)^\top \left(-\frac{1}{C_{S_0}RT} \alpha c^\top n_1 + \text{diag} \left\{ \frac{\mathbb{1}_N}{n_0} \right\} n_1 \right) \right) \\ &= \text{diag} \left\{ F^R(n_0) \right\} (\nu^R)^\top \left(-\frac{1}{C_{S_0}RT} \alpha c^\top + \text{diag} \left\{ \frac{\mathbb{1}_N}{n_0} \right\} \right) n_1 \end{aligned} \quad (5.19)$$

La formule (5.12) peut maintenant s'écrire :

$$j\omega n_1 = An_1 - B_n n_1 + B_I I_1 \quad (5.20)$$

où

$$\begin{aligned} A &= \nu \left(\text{diag} \left\{ F^R(n_0) \right\} (\nu^R)^\top \left(-\frac{1}{C_{S_0}RT} \alpha c^\top + \text{diag} \left\{ \frac{\mathbb{1}_N}{n_0} \right\} \right) \right. \\ &\quad \left. - \text{diag} \left\{ F^P(n_0) \right\} (\nu^P)^\top \left(\frac{1}{C_{S_0}RT} \beta c^\top + \text{diag} \left\{ \frac{\mathbb{1}_N}{n_0} \right\} \right) \right) \end{aligned}$$

On obtient d'après (5.20) :

$$n_1 = (j\omega I_N - A + B_n)^{-1} B_I I_1 \quad (5.21)$$

L'impédance $Z(j\omega)$ s'écrit finalement sous la forme :

$$Z(j\omega) = C_n (j\omega I_N - A + B_n)^{-1} B_I \quad (5.22a)$$

où

$$\begin{aligned} A &= \nu \text{diag} \left\{ F^R(n_0) \right\} (\nu^R)^\top \left(-\frac{1}{C_{S_0}RT} \alpha c^\top + \text{diag} \left\{ \frac{\mathbb{1}_N}{n_0} \right\} \right) \\ &\quad - \nu \text{diag} \left\{ F^P(n_0) \right\} (\nu^P)^\top \left(-\frac{1}{C_{S_0}RT} \beta c^\top + \text{diag} \left\{ \frac{\mathbb{1}_N}{n_0} \right\} \right) \end{aligned} \quad (5.22b)$$

et F^P, F^R sont définis par (5.13).

5.3 Mécanisme réactionnel de Tafel-Heyrovsky-Volmer à l'anode de la pile

Le système d'état dans la couche compacte de l'anode de la pile est donné par (cf. (4.160)) :

$$\dot{\xi}_{c,a} = \begin{pmatrix} \mathcal{L}_{TH}(\xi_{c,a}, \tilde{n}_{\bar{c},a}) \\ \mathcal{L}_{HV}(\xi_{c,a}, \tilde{n}_{\bar{c},a}) \end{pmatrix} - S_{CA} \begin{pmatrix} D_{H_2}^{cc} & 0 \\ 0 & 0 \end{pmatrix} (\xi_{c,a}^{ext} - \xi_{c,a}) - \frac{I_M}{F} \begin{pmatrix} 0 \\ 1 \end{pmatrix} \quad (5.23a)$$

$$\dot{\tilde{n}}_{\bar{c},a} = S_{CA} \begin{pmatrix} 2D_{H_2}^{cc} & 0 \\ -2D_{H_2}^{cc} & 0 \end{pmatrix} (\xi_{c,a}^{ext} - \xi_{c,a}) - \frac{I_M}{F} \begin{pmatrix} 1 \\ -1 \end{pmatrix} \quad (5.23b)$$

$$\Delta\phi_{Sa} = -\frac{1}{FC_{S0}} \begin{pmatrix} 0 & 1 & 0 & 0 \end{pmatrix} \begin{pmatrix} \xi_{c,a} \\ \tilde{n}_{\bar{c},a} \end{pmatrix} \quad (5.23c)$$

On rappelle (cf. (4.73)) que :

$$\xi_{c,a} = \begin{pmatrix} -n_{H_2} \\ n_{H^+} \end{pmatrix}, \quad \tilde{n}_{\bar{c},a} = \begin{pmatrix} n_{Hs} + 2n_{H_2} + n_{H^+} \\ n_s - 2n_{H_2} - n_{H^+} \end{pmatrix}$$

et (cf. (4.156), (4.158))

$$\begin{pmatrix} \mathcal{L}_{TH}(\xi_{c,a}, \tilde{n}_{\bar{c},a}) \\ \mathcal{L}_{HV}(\xi_{c,a}, \tilde{n}_{\bar{c},a}) \end{pmatrix} = \Pi \left[k_{Ra} \cdot \text{Exp} \left((\nu_a^R)^\top \left(-\frac{\alpha F c_a^\top n_a}{C_{S0} RT} + \text{Ln} \left(\frac{n_a}{V} \right) \right) \right) \right. \\ \left. - k_{Pa} \cdot \text{Exp} \left((\nu_a^P)^\top \left(\frac{\beta F c_a^\top n_a}{C_{S0} RT} + \text{Ln} \left(\frac{n_a}{V} \right) \right) \right) \right], \quad n_a = \tilde{\nu} \xi_{c,a} + \begin{pmatrix} 0_{2 \times 1} \\ \tilde{n}_{\bar{c},a} \end{pmatrix} \quad (5.24)$$

En s'inspirant de Franco [37], on suppose que :

$$\xi_{c,a}^{ext} = \begin{pmatrix} -n_{H_2}^{GDL} \\ 0 \end{pmatrix} = \begin{pmatrix} -\frac{P_{H_2}}{H_{H_2}} V \\ 0 \end{pmatrix} \quad (5.25)$$

où $n_{H_2}^{GDL}$ est le nombre de mole de l'hydrogène à l'interface GDL/CC.

5.3.1 Courbes de polarisation

À l'équilibre, (5.23b) est équivalent à :

$$-S_{CA} \begin{pmatrix} 2D_{H_2}^{cc} \\ -2D_{H_2}^{cc} \end{pmatrix} (n_{H_2,0}^{ext} - n_{H_2,0}) = \begin{pmatrix} 1 \\ -1 \end{pmatrix} \frac{I_M}{F} \quad (5.26)$$

d'où

$$n_{H_2,0}^{ext} - n_{H_2,0} = -\frac{1}{2S_{CA}D_{H_2}} \frac{I_M}{F} \quad (5.27)$$

et on cherche à calculer $n_{H^+,0}$, $n_{Hs,0}$ et $n_{s,0}$. La valeur de $n_{Hs} + n_s$ est égale au nombre total de sites :

$$n_{Hs} + n_s = n_{\max} \quad (5.28)$$

On peut alors exprimer $n_{s,0}$ en fonction de $n_{Hs,0}$ et il nous reste à déterminer $n_{H^+,0}$ et $n_{Hs,0}$. À l'équilibre, on a d'après (5.23a)

$$\begin{pmatrix} \mathcal{L}_{TH}(\xi_{c,a}, \tilde{n}_{\bar{c},a}) \\ \mathcal{L}_{HV}(\xi_{c,a}, \tilde{n}_{\bar{c},a}) \end{pmatrix} = S_{CA} \begin{pmatrix} D_{H_2}^{cc} & 0 \\ 0 & 0 \end{pmatrix} (\xi_{c,a}^{ext} - \xi_{c,a}) + \frac{I_M}{F} \begin{pmatrix} 0 \\ 1 \end{pmatrix} \quad (5.29)$$

En éliminant $n_{H_2,0}$ et $n_{s,0}$, ce système se résume donc à un système de deux équations en deux inconnues (en n_{H^+} et n_{Hs}).

La résolution de ce système, permet de trouver les valeurs correspondante n_{a0} des concentrations d'équilibre. La différence de potentiel à l'équilibre s'écrit alors sous la forme :

$$\Delta\phi_{Sa,0} = -\frac{1}{FS_{C0}} \begin{pmatrix} 0 & 1 & 0 & 0 \end{pmatrix} \begin{pmatrix} \xi_{c,a0} \\ \tilde{n}_{\bar{c},a0} \end{pmatrix} \quad (5.30)$$

Cette formule permet de tracer les courbes de polarisation.

5.4 Mécanisme réactionnel de Damjanovic et Brusic à la cathode de la pile

Le système d'état dans la couche compacte de la cathode de la pile est donné par (cf. (4.171))

$$\dot{\xi}_{c,c} = \begin{pmatrix} \mathcal{L}_1(\xi_{c,c}, \tilde{n}_{\bar{c},c}) \\ \mathcal{L}_2(\xi_{c,c}, \tilde{n}_{\bar{c},c}) \\ \mathcal{L}_3(\xi_{c,c}, \tilde{n}_{\bar{c},c}) \end{pmatrix} + S_{CA} \begin{pmatrix} D_{O_2} & 0 & 0 \\ D_{O_2} & 0 & 0 \\ 3D_{O_2} & 0 & 0 \end{pmatrix} (\xi_{c,c}^{ext} - \xi_{c,c}) \quad (5.31a)$$

$$\begin{aligned} \dot{\tilde{n}}_{\bar{c},c} = S_{CA} \begin{pmatrix} -2D_{O_2} & -D_{H_2O} & D_{H_2O} \\ 4D_{O_2} & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} (\xi_{c,c}^{ext} - \xi_{c,c}) \\ + S_{CA} \begin{pmatrix} D_{H_2O} & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} (\tilde{n}_{\bar{c},c}^{ext} - \tilde{n}_{\bar{c},c}) + \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \frac{I_M}{F} \end{aligned} \quad (5.31b)$$

$$\Delta\phi_{Sc} = -\frac{1}{FS_{C0}} \begin{pmatrix} 0 & -1 & 1 & 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} \xi_{c,c} \\ \tilde{n}_{\bar{c},c} \end{pmatrix} \quad (5.31c)$$

On rappelle (cf. (4.170)) que :

$$\xi_{c,c} = \begin{pmatrix} -n_{O_2} \\ -n_{O_2} - n_{O_2Hs} \\ -3n_{O_2} - 3n_{O_2Hs} - n_{OHs} \end{pmatrix}, \quad \tilde{n}_{\bar{c},c} = \begin{pmatrix} n_{H_2O} + 2n_{O_2} + 2n_{O_2Hs} + n_{OHs} \\ n_{H^+} - 4n_{O_2} - 3n_{O_2Hs} - n_{OHs} \\ n_s + n_{O_2Hs} + n_{OHs} \end{pmatrix}, \quad (5.32)$$

En s'inspirant de Franco [37], on suppose que :

$$\xi_{c,c}^{ext} = - \begin{pmatrix} 1 \\ 1 \\ 3 \end{pmatrix} n_{O_2}^{ext} = - \begin{pmatrix} 1 \\ 1 \\ 3 \end{pmatrix} \frac{P_{O_2}}{H_{O_2}} V \quad (5.33a)$$

et

$$\tilde{n}_{\bar{c},c}^{ext} = \begin{pmatrix} 2\frac{P_{O_2}}{H_{O_2}} V + n_{H_2O}^{GDL} \\ -4\frac{P_{O_2}}{H_{O_2}} V \\ 0 \end{pmatrix} \quad (5.33b)$$

5.4.1 Courbes de polarisation

À l'équilibre, (5.31b) est équivalent à :

$$\begin{aligned} S_{CA} \begin{pmatrix} -2D_{O_2} & -D_{H_2O} & D_{H_2O} \\ 4D_{O_2} & 0 & 0 \end{pmatrix} (\xi_{c,c0}^{ext} - \xi_{c,c0}) \\ = -S_{CA} \begin{pmatrix} D_{H_2O} & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} (\tilde{n}_{\bar{c},c0}^{ext} - \tilde{n}_{\bar{c},c0}) + \begin{pmatrix} 0 \\ 1 \end{pmatrix} \frac{I_M}{F} \end{aligned} \quad (5.34)$$

On déduit d'après (5.32) :

$$2D_{O_2}(n_{O_2,0}^{ext} - n_{O_2,0}) + D_{H_2O}(n_{H_2O,0}^{ext} - n_{H_2O,0}) = 0 \quad (5.35a)$$

$$4D_{O_2}(n_{O_2,0}^{ext} - n_{O_2,0}) = -\frac{I_M}{FS_{CA}} \quad (5.35b)$$

d'où

$$n_{H_2O,0}^{ext} - n_{H_2O,0} = \frac{I_M}{2FD_{H_2O}S_{CA}} \quad (5.36a)$$

$$n_{O_2,0}^{ext} - n_{O_2,0} = -\frac{I_M}{4FD_{O_2}S_{CA}} \quad (5.36b)$$

et on cherche à calculer $n_{H^+,0}$, $n_{O_2Hs,0}$, $n_{OHs,0}$ et $n_{s,0}$. La valeur de $n_{O_2Hs} + n_{OHs} + n_s$ est égale au nombre total de sites :

$$n_{O_2Hs} + n_{OHs} + n_s = n_{\max}. \quad (5.37)$$

On peut alors exprimer $n_{s,0}$ en fonction de $n_{O_2Hs,0}$ et $n_{OHs,0}$, il nous reste alors à déterminer : $n_{H^+,0}$, $n_{O_2Hs,0}$ et $n_{OHs,0}$.

D'après (5.31a), on a à l'équilibre :

$$\begin{pmatrix} \mathcal{L}_1(\xi_{c,c0}, \tilde{n}_{\bar{c},c0}) \\ \mathcal{L}_2(\xi_{c,c0}, \tilde{n}_{\bar{c},c0}) \\ \mathcal{L}_3(\xi_{c,c0}, \tilde{n}_{\bar{c},c0}) \end{pmatrix} = -S_{CA} \begin{pmatrix} D_{O_2} & 0 & 0 \\ D_{O_2} & 0 & 0 \\ 3D_{O_2} & 0 & 0 \end{pmatrix} (\xi_{c,c0}^{ext} - \xi_{c,c0}) \quad (5.38)$$

En éliminant $n_{O_2,0}$, $n_{H_2O,0}$ et $n_{s,0}$, le système (5.38) se résume donc à un système de trois équations scalaires en trois inconnues scalaires en $n_{H^+,0}$, $n_{O_2Hs,0}$ et $n_{OHs,0}$. La résolution de ce système permet de déterminer ces inconnus et par conséquent on obtient $\xi_{c,c0}$ et $\tilde{n}_{\bar{c},c0}$. La différence de potentiel à l'équilibre s'écrit alors sous la forme :

$$\Delta\phi_{Sc,0} = -\frac{1}{FS_{C0}} \begin{pmatrix} 0 & -1 & 1 & 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} \xi_{c,c0} \\ \tilde{n}_{\bar{c},c0} \end{pmatrix} \quad (5.39)$$

Cette formule permet de tracer les courbes de polarisation.

5.5 Modèle d'impédance de double couche et de pile avec prise en compte la correction de Frumkin

5.5.1 Modèle de double couche

On considère une demi pile soumise à un courant :

$$J_M(t) = J_{M0} + J_{M1} \cos(\omega t) \quad (5.40)$$

et à une concentration d'entrée constante c_{R0} . On suppose qu'il existe une tension d'équilibre $\Delta\phi_{S0}$ telle que $J_{M0} = \tilde{\mathcal{R}}(c_R, \Delta\phi_{S0})$. Au vu de (4.172a), on a au premier ordre :

$$\Delta\phi_{S,1} = \left(j\omega C_{S0} + \partial_2 \tilde{\mathcal{R}}(c_R, \Delta\phi_{S,0}) \right)^{-1} J_{M1} \quad (5.41a)$$

$$\Delta\phi_{dl,1} = (1 + \partial\Phi_D(\Delta\phi_{S0})) \Delta\phi_{S,1} \quad (5.41b)$$

En notant par $Z_{dl}(\omega) = \frac{\Delta\phi_{dl,1}}{J_{M1}}$ l'impédance de la double couche, on a alors d'après (5.41) :

$$Z_{dl}(\omega) = (1 + \partial\Phi_D(\Delta\phi_{S0})) \left(j\omega C_{S0} + \partial_2 \tilde{\mathcal{R}}(c_R, \Delta\phi_{S,0}) \right)^{-1} \quad (5.42)$$

Ce qui fournit un modèle d'impédance de la double couche.

5.5.2 Modèle de la pile

On considère maintenant le modèle de pile donné par (4.172). Comme on l'a vu, il peut être transformé en (4.186) donné par :

$$C_{S0} \frac{\partial}{\partial t} \Delta \phi_S = -\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S) + \frac{1}{R_M} \left(\Phi_{dlc}(\Delta \phi_{S0} + \beta_a \Delta \phi_S) - \Phi_{dla}(-\Delta \phi_{S0} + \beta_c \Delta \phi_S) - Z_{ext} * \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S) \right) \quad (5.43)$$

Circuit utilisé pour une expérience d'impédance. Il faut introduire un courant oscillant. On suppose qu'on met *en série avec la charge extérieure* un courant $J_{imp}(t)$. Seul le courant $J - J_{imp}$ passe maintenant dans la charge extérieure, et on a donc au lieu de (4.185d)

$$U_{cell} = Z_{ext} * \left(\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S) - J_{imp} \right) \quad (5.44)$$

Le reste des équations (4.176) est inchangé, et l'équation (4.186) est maintenant à remplacer par :

$$C_{S0} \frac{\partial}{\partial t} \Delta \phi_S = -\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S) + \frac{1}{R_M} \left(\Phi_{dlc}(\Delta \phi_{S0} + \beta_a \Delta \phi_S) - \Phi_{dla}(-\Delta \phi_{S0} + \beta_c \Delta \phi_S) - Z_{ext} * \left(\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S) - J_{imp} \right) \right) \quad (5.45)$$

Equilibre. On suppose J_{imp} stationnaire : $J_{imp} \equiv J_{imp}^0$. L'équilibre $\Delta \phi_S^0$ doit vérifier

$$R_M \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S^0) = \Phi_{dlc}(\Delta \phi_{S0} + \beta_a \Delta \phi_S^0) - \Phi_{dla}(-\Delta \phi_{S0} + \beta_c \Delta \phi_S^0) - Z_{ext}(0) \left(\tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S^0) - J_{imp}^0 \right) \quad (5.46)$$

où $Z_{ext}(0) = Z_{ext}(j\omega)|_{\omega=0}$. Les tensions et courant d'équilibre sont donnés par :

$$J^0 = \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S^0) \quad (5.47a)$$

$$U_{cell}^0 = Z_{ext}(0) \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S^0) \quad (5.47b)$$

Impédance. On suppose maintenant : J_{imp} oscillant autour de l'équilibre précédent : $J_{imp} \equiv J_{imp}^0 + J_{imp}^1 \cos \omega t$. En utilisant l'exposant 1 pour les composantes sinusoïdales des signaux, on a

$$\begin{aligned} R_M \left(j\omega C_{S0} + \partial_3 \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S^0) \right) \Delta \phi_S^1 \\ = \left(\beta_a \partial \Phi_{dlc}(\Delta \phi_{S0} + \beta_a \Delta \phi_S^0) - \beta_c \partial \Phi_{dla}(-\Delta \phi_{S0} + \beta_c \Delta \phi_S^0) \right) \Delta \phi_S^1 \\ - Z_{ext}(j\omega) \left(\partial_3 \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S^0) \Delta \phi_S^1 - J_{imp}^1 \right) \end{aligned}$$

La notation ∂_3 désigne la dérivation par rapport à la troisième variable. On a donc :

$$\left(j\omega + a + \partial_3 \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta \phi_S^0) \frac{Z_{ext}(j\omega)}{R_M C_{S0}} \right) \Delta \phi_S^1 = \frac{Z_{ext}(j\omega)}{R_M C_{S0}} J_{imp}^1 \quad (5.48)$$

où la constante a est définie par :

$$a = \frac{1}{R_M C_{S0}} \left(\partial_3 \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S^0) - \beta_a \partial \Phi_{dlc}(\Delta\phi_{S0} + \beta_a \Delta\phi_S^0) + \beta_c \partial \Phi_{dla}(-\Delta\phi_{S0} + \beta_c \Delta\phi_S^0) \right) \quad (5.49)$$

On déduit de (5.48) le transfert

$$Z(j\omega) = \frac{\Delta\phi_S^1}{J_{imp}^1} = \left(j\omega + a + \partial_3 \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S^0) \frac{Z_{ext}(j\omega)}{R_M C_{S0}} \right)^{-1} \frac{Z_{ext}(j\omega)}{R_M C_{S0}}$$

et on déduit finalement de ce transfert le modèle d'impédance de la pile complète notée Z_{pile} que nous recherchons. En effet,

$$Z_{pile} = \frac{U^1}{J_{imp}^1} = \partial_3 \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S^0) Z_{ext}(j\omega) Z(j\omega)$$

On en déduit donc

$$Z_{pile} = \partial_3 \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S^0) Z_{ext}(j\omega) \left(j\omega + a + \partial_3 \tilde{\mathcal{R}}_{cell}(c_{Ra}, c_{Oc}, \Delta\phi_S^0) \frac{Z_{ext}(j\omega)}{R_M C_{S0}} \right)^{-1} \frac{Z_{ext}(j\omega)}{R_M C_{S0}} \quad (5.50)$$

5.6 Décomposition quadripolaire du GDL

L'objectif de cette section est d'obtenir un schéma électrique analogique de la couche de diffusion des gaz (DGL). Tout d'abord, nous calculons ici l'impédance de la couche de diffusion des gaz (GDL) et en se basant sur une décomposition quadripolaire nous obtenons le schéma électrique associé. On considère par exemple l'anode, le GDL à la cathode étant modélisée de la même façon.

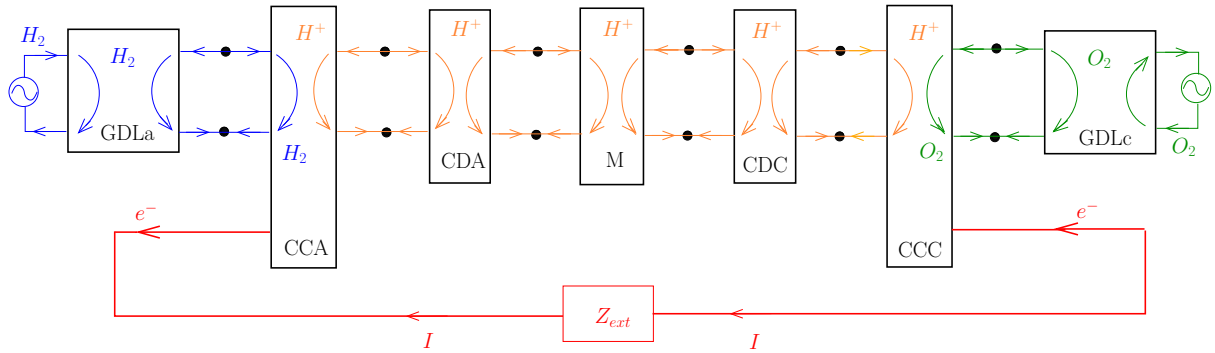


FIGURE 5.3 – Schéma électrique de la pile. La couleur bleue indique le chemin de l'hydrogène. Le couleur rouge (resp. orange) indique le déplacement des électrons (resp. des protons). La couleur verte indique le chemin de l'oxygène.

Remarque 5.6. Lorsque les sources d'hydrogène et d'oxygène sont branchées, la pile est un dipole électrique comme on le voit Figure 5.3. Sans ces sources, elle possède naturellement 3 ports : un port pour l'hydrogène, un port pour l'oxygène et un port électrique. Chaque demi-pile possède donc trois des six pôles de départ plus un port à l'endroit de la coupure : l'anode (resp. la cathode) a un port d'hydrogène (resp. d'oxygène), un port de protons et le pôle négatif (resp. positif) de la pile. On retrouve cette organisation avec 2 ports et un pôle

électrique pour les couches compactes, les autres éléments étant des dipôles ou quadripôles classiques

Le sens positif du courant est le sens entrant pour chaque port. La somme des courants est nulle à chaque nœud.

5.6.1 Modèle quadripolaire



FIGURE 5.4 – GDL anode

	quantité	unité	quantité	unité
port 1 gazeux	$C_{H_2,gdla,1}$	$mol.m^{-3}$	$q_{H_2,gdla,1}$	$mol.s^{-1}$
port 2 gazeux	$C_{H_2,gdla,2}$	$mol.m^{-3}$	$q_{H_2,gdla,2}$	$mol.s^{-1}$

TABLE 5.1 – Variables des ports du GDL anode

L'évolution de la concentration C_{H_2} en hydrogène à l'anode s'écrit sous la forme :

$$S_{CA} \frac{\partial C_{H_2}}{\partial t}(z, t) = -\frac{\partial q_{H_2}}{\partial z}(z, t), \quad z \in]z_1, z_2[\quad (5.51a)$$

$$q_{H_2}(z, t) = -S_{CA} D_{H_2} \frac{\partial C_{H_2}}{\partial z}(z, t), \quad z \in]z_1, z_2[\quad (5.51b)$$

$$q_{H_2,gdla,i} = (-1)^i S_{CA} D_{H_2} \frac{\partial C_{H_2}}{\partial z}(z_i), \quad i = 1, 2 \quad (5.51c)$$

$$\frac{C_{H_2,gdla,i} - C_{H_2}(z_i)}{R_i} = (-1)^i S_{CA} D_{H_2} \frac{\partial C_{H_2}}{\partial z}(z_i), \quad i = 1, 2 \quad (5.51d)$$

où S_{CA} est la surface d'échange, qui vaut :

$$S_{CA} = \gamma S_{EME} e_{CA}. \quad (5.52)$$

En regroupant les équations (5.51a) et (5.51b), le système (5.51) s'écrit sous la forme :

$$\frac{\partial C_{H_2}}{\partial t}(z, t) = D_{H_2} \frac{\partial^2 C_{H_2}}{\partial z^2}(z, t), \quad z \in]z_1, z_2[\quad (5.53a)$$

$$q_{H_2,gdla,i} = (-1)^i S_{CA} D_{H_2} \frac{\partial C_{H_2}}{\partial z}(z_i), \quad i = 1, 2 \quad (5.53b)$$

$$C_{H_2,gdla,i} = C_{H_2}(z_i) + R_i q_{H_2,gdla,i}, \quad i = 1, 2 \quad (5.53c)$$

On posera

$$z_1 = 0, \quad z_2 = e_{gdl a}$$

On a d'après (5.51a) et (5.51b) :

$$S_{CA} \frac{d}{dt} \int_0^{e_{gdl a}} C_{H_2}(z, t) dz = q_{H_2, gdl a, 1} + q_{H_2, gdl a, 2} \quad (5.54a)$$

et d'après (5.51c) et (5.51d) :

$$\frac{1}{e_{gdl a}} \int_0^{e_{gdl a}} q_{H_2}(z, t) dz = R_{gdl a}^{-1} (C_{H_2, gdl a, 1} - R_1 q_{H_2, gdl a, 1} - C_{H_2, gdl a, 2} + R_1 q_{H_2, gdl a, 2}) \quad (5.54b)$$

où

$$R_{gdl a} = \frac{e_{gdl a}}{S_{CA} D_{H_2}} \quad (5.54c)$$

Remarque 5.7. *Le GDL dans notre modèle représente le Nafion imprégné dans le modèle de Franco.*

Remarque 5.8. (5.53a) représente l'équation de la diffusion dans le GDL. Les équations (5.53b) définissent des notations compatibles avec la théorie des quadripôles : on a effectivement

$$q_{H_2, gdl a, 1} = q_{H_2}(z_1), \quad q_{H_2, gdl a, 2} = -q_{H_2}(z_2) \quad (5.55)$$

Pour toute valeur donnée de $C_{H_2, gdl a, 1}$, on peut définir l'opérateur Neumann-to-Dirichlet NtD_a , qui vérifie par construction :

$$C_{H_2, gdl a, 2} = NtD_a \cdot q_{H_2, gdl a, 2} = \mathcal{R}_{W a} * q_{H_2, gdl a, 2} + \mathcal{G}_a * C_{H_2, gdl a, 1}. \quad (5.56)$$

5.6.2 Formulation variationnelle

En multipliant (5.53a) par $\varphi \in H^1([z_1, z_2])$ et en intégrant en z entre z_1 et z_2 , on obtient :

$$\begin{aligned} \frac{d}{dt} \int_{z_1}^{z_2} C_{H_2}(z, t) \varphi(z) dz &= D_{H_2} \int_{z_1}^{z_2} \frac{\partial^2 C_{H_2}}{\partial z^2}(z, t) \varphi(z) dz \\ &= -D_{H_2} \int_{z_1}^{z_2} \frac{\partial C_{H_2}}{\partial z}(z, t) \frac{\partial \varphi}{\partial z}(z) dz + D_{H_2} \frac{\partial C_{H_2}}{\partial z}(z_2) \varphi(z_2) - D_{H_2} \frac{\partial C_{H_2}}{\partial z}(z_1) \varphi(z_1) \end{aligned} \quad (5.57)$$

En utilisant les conditions aux bords (5.53b) et (5.53c) on obtient :

$$\begin{aligned} \frac{d}{dt} \int_{z_1}^{z_2} C_{H_2}(z, t) \varphi(z) dz &= -D_{H_2} \int_{z_1}^{z_2} \frac{\partial C_{H_2}}{\partial z}(z, t) \frac{\partial \varphi}{\partial z}(z) dz + \frac{C_{H_2, gdl a, 2} - C_{H_2}(z_2)}{S_{CA} R_2} \varphi(z_2) \\ &\quad + \frac{C_{H_2, gdl a, 1} - C_{H_2}(z_1)}{S_{CA} R_1} \varphi(z_1) \end{aligned} \quad (5.58)$$

On note :

$$A(C, \varphi) = -D_{H_2} \left(\frac{\partial C}{\partial z}, \frac{\partial \varphi}{\partial z} \right)_{L^2} - \frac{1}{S_{CA}} \sum_{i=1}^2 \frac{1}{R_i} C(z_i) \varphi(z_i) \quad (5.59)$$

où $(\cdot, \cdot)_{L^2}$ est le produit scalaire dans L^2 défini par :

$$(C, \varphi)_{L^2} = \int_{z_1}^{z_2} C(z, t) \varphi(z) dz$$

On a alors

$$\left(\frac{\partial C_{H_2}}{\partial t}, \varphi\right)_{L^2} = A(C_{H_2}, \varphi) + \frac{1}{S_{CA}} \sum_{i=1}^2 \frac{1}{R_i} C_{H_2, gdl a, i} \gamma_{z_i}^*(\varphi) \quad (5.60)$$

où $\gamma_{z_i}^*$ est l'opérateur trace, défini par :

$$\gamma_{z_i}^*(\varphi) = \varphi(z_i).$$

5.6.3 Comportement harmonique

On cherche une solution de (5.51) sous la forme :

$$C_{H_2}(z, t) = C(z) \exp(j\omega t), \quad q_{H_2}(z, t) = q(z) \exp(j\omega t). \quad (5.61)$$

On a d'après (5.61), (5.51c) et (5.51d) :

$$q_{H_2, gdl a, 1} = q(0) \exp(j\omega t), \quad q_{H_2, gdl a, 2} = -q(e_{gdl a}) \exp(j\omega t) \quad (5.62a)$$

$$C_{H_2, gdl a, 1} = (C(0) + R_1 q(0)) \exp(j\omega t), \quad C_{H_2, gdl a, 2} = (C(e_{gdl a}) - R_2 q(e_{gdl a})) \exp(j\omega t) \quad (5.62b)$$

D'après (5.51a), on a :

$$S_{CA} j\omega C(z) = -\frac{\partial q(z)}{\partial z}, \quad q(z) = -S_{CA} D_{H_2} \frac{\partial C(z)}{\partial z}. \quad (5.63)$$

On obtient l'équation suivante :

$$\frac{\partial^2 q}{\partial z^2}(z) = -S_{CA} j\omega \frac{\partial C}{\partial z} = \frac{j\omega}{D_{H_2}} q(z). \quad (5.64)$$

En introduisant la constante de temps de diffusion suivante

$$\tau_{gdl a} = \frac{e_{gdl a}^2}{D_{H_2}} \quad (5.65)$$

l'équation précédente s'écrit sous la forme :

$$\frac{\partial^2 q}{\partial z^2}(z) = \frac{j\omega \tau_{gdl a}}{e_{gdl a}^2} q(z). \quad (5.66)$$

La solution de (5.66) est donnée par :

$$q(z) = \alpha \exp\left(\frac{\sqrt{j\omega \tau_{gdl a}}}{e_{gdl a}} z\right) + \beta \exp\left(-\frac{\sqrt{j\omega \tau_{gdl a}}}{e_{gdl a}} z\right). \quad (5.67)$$

On a alors :

$$q(0) = \alpha + \beta \quad (5.68a)$$

$$q(e_{gdl a}) = \alpha \exp\left(\sqrt{j\omega \tau_{gdl a}}\right) + \beta \exp\left(-\sqrt{j\omega \tau_{gdl a}}\right) \quad (5.68b)$$

On obtient

$$\begin{aligned} \alpha &= -\frac{1}{2 \sinh \sqrt{j\omega \tau_{gdl a}}} \times \begin{vmatrix} q(0) & 1 \\ q(e_{gdl a}) & \exp(-\sqrt{j\omega \tau_{gdl a}}) \end{vmatrix} \\ &= \frac{q(e_{gdl a}) - q(0) \exp(-\sqrt{j\omega \tau_{gdl a}})}{2 \sinh \sqrt{j\omega \tau_{gdl a}}} \end{aligned} \quad (5.69a)$$

et

$$\begin{aligned}\beta &= -\frac{1}{2 \sinh \sqrt{j\omega\tau_{gdl a}}} \times \begin{vmatrix} 1 & q(0) \\ \exp(\sqrt{j\omega\tau_{gdl a}}) & q(e_{gdl a}) \end{vmatrix} \\ &= \frac{q(0) \exp(\sqrt{j\omega\tau_{gdl a}}) - q(e_{gdl a})}{2 \sinh \sqrt{j\omega\tau_{gdl a}}}\end{aligned}\quad (5.69b)$$

Par ailleurs, on a d'après (5.63) et (5.67) :

$$\begin{aligned}C(z) &= -\frac{1}{S_{CA}j\omega} \frac{\partial q(z)}{\partial z} = -\frac{1}{S_{CA}e_{gdl a}} \sqrt{\frac{\tau_{gdl a}}{j\omega}} \left[\alpha \exp\left(\frac{\sqrt{j\omega\tau_{gdl a}}}{e_{gdl a}} z\right) - \beta \exp\left(-\frac{\sqrt{j\omega\tau_{gdl a}}}{e_{gdl a}} z\right) \right] \\ &= -\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \left[\alpha \exp\left(\frac{\sqrt{j\omega\tau_{gdl a}}}{e_{gdl a}} z\right) - \beta \exp\left(-\frac{\sqrt{j\omega\tau_{gdl a}}}{e_{gdl a}} z\right) \right]\end{aligned}\quad (5.70a)$$

d'où, en utilisant (5.69) :

$$\begin{aligned}C(0) &= -\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} (\alpha - \beta) \\ &= -\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \frac{q(e_{gdl a}) - q(0) \exp(-\sqrt{j\omega\tau_{gdl a}}) - q(0) \exp(\sqrt{j\omega\tau_{gdl a}}) + q(e_{gdl a})}{2 \sinh \sqrt{j\omega\tau_{gdl a}}} \\ &= \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \frac{q(0) \cosh \sqrt{j\omega\tau_{gdl a}} - q(e_{gdl a})}{\sinh \sqrt{j\omega\tau_{gdl a}}}\end{aligned}\quad (5.71a)$$

$$\begin{aligned}C(e_{gdl a}) &= -\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \left(\alpha \exp(\sqrt{j\omega\tau_{gdl a}}) - \beta \exp(-\sqrt{j\omega\tau_{gdl a}}) \right) \\ &= -\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \frac{q(e_{gdl a}) \exp \sqrt{j\omega\tau_{gdl a}} - 2q(0) + q(e_{gdl a}) \exp(-\sqrt{j\omega\tau_{gdl a}})}{2 \sinh \sqrt{j\omega\tau_{gdl a}}} \\ &= \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \frac{q(0) - q(e_{gdl a}) \cosh \sqrt{j\omega\tau_{gdl a}}}{\sinh \sqrt{j\omega\tau_{gdl a}}}\end{aligned}\quad (5.71b)$$

En utilisant (5.62), on en déduit que :

$$C_{H_2, gdl a, 1} = \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \frac{1}{\sinh \sqrt{j\omega\tau_{gdl a}}} \left(q_{H_2, gdl a, 2} + q_{H_2, gdl a, 1} \cosh \sqrt{j\omega\tau_{gdl a}} \right) + R_1 q_{H_2, gdl a, 1}\quad (5.72a)$$

$$C_{H_2, gdl a, 2} = \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \frac{1}{\sinh \sqrt{j\omega\tau_{gdl a}}} \left(q_{H_2, gdl a, 1} + q_{H_2, gdl a, 2} \cosh \sqrt{j\omega\tau_{gdl a}} \right) + R_2 q_{H_2, gdl a, 2}\quad (5.72b)$$

d'où la relation d'impédance suivante :

$$\begin{pmatrix} C_{H_2, gdl a, 1} \\ C_{H_2, gdl a, 2} \end{pmatrix} = Z(Q_{gdl a}) \begin{pmatrix} q_{H_2, gdl a, 1} \\ q_{H_2, gdl a, 2} \end{pmatrix}\quad (5.73a)$$

où

$$Z(Q_{gdl a}) = \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \begin{pmatrix} \frac{1}{\tanh \sqrt{j\omega\tau_{gdl a}}} & \frac{1}{\sinh \sqrt{j\omega\tau_{gdl a}}} \\ \frac{1}{\sinh \sqrt{j\omega\tau_{gdl a}}} & \frac{1}{\tanh \sqrt{j\omega\tau_{gdl a}}} \end{pmatrix} + \begin{pmatrix} R_1 & 0 \\ 0 & R_2 \end{pmatrix} \quad (5.73b)$$

On va maintenant écrire le système (5.72) en utilisant une représentation sous forme Transmission. En utilisant par exemple [42], si $Z(Q_{gdl a}) = \begin{pmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{pmatrix}$, on a $T(Q_{gdl a}) = \begin{pmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \end{pmatrix}$ où

$$T_{11} = \frac{Z_{11}}{Z_{21}}, \quad T_{12} = \frac{Z_{11}Z_{22} - Z_{12}Z_{21}}{Z_{21}}, \quad T_{21} = \frac{1}{Z_{21}}, \quad T_{22} = \frac{Z_{22}}{Z_{21}}.$$

On a donc

$$T_{ii} = \left(\frac{1}{\tanh \sqrt{j\omega\tau_{gdl a}}} + \frac{R_i \sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \right) \sinh \sqrt{j\omega\tau_{gdl a}}, \quad i = 1, 2 \quad (5.74a)$$

$$T_{21} = \frac{\sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \sinh \sqrt{j\omega\tau_{gdl a}} \quad (5.74b)$$

$$\begin{aligned} T_{12} &= \frac{\left(\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \tanh \sqrt{j\omega\tau_{gdl a}} + R_1 \right) \left(\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \tanh \sqrt{j\omega\tau_{gdl a}} + R_2 \right) - \frac{R_{gdl a}^2}{j\omega\tau_{gdl a} \sinh^2 \sqrt{j\omega\tau_{gdl a}}}}{\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}} \sinh \sqrt{j\omega\tau_{gdl a}}}} \\ &= \left(\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} + \frac{R_1 R_2 \sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \right) \sinh \sqrt{j\omega\tau_{gdl a}} + (R_1 + R_2) \cosh \sqrt{j\omega\tau_{gdl a}} \end{aligned} \quad (5.74c)$$

d'où la relation de transmission suivante :

$$\begin{pmatrix} C_{H_2, gdl a, 1} \\ q_{H_2, gdl a, 1} \end{pmatrix} = T(Q_{gdl a}) \begin{pmatrix} C_{H_2, gdl a, 2} \\ -q_{H_2, gdl a, 2} \end{pmatrix} \quad (5.75)$$

où

$$T(Q_{gdl a}) = \begin{pmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \end{pmatrix} \quad (5.76a)$$

$$T_{ii} = \left(\frac{1}{\tanh \sqrt{j\omega\tau_{gdl a}}} + \frac{R_i \sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \right) \sinh \sqrt{j\omega\tau_{gdl a}}, \quad i = 1, 2 \quad (5.76b)$$

$$T_{12} = \left(\frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} + \frac{R_1 R_2 \sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \right) \sinh \sqrt{j\omega\tau_{gdl a}} + (R_1 + R_2) \cosh \sqrt{j\omega\tau_{gdl a}} \quad (5.76c)$$

$$T_{21} = \frac{\sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \sinh \sqrt{j\omega\tau_{gdl a}} \quad (5.76d)$$

où $R_{gdl a}$ et $\tau_{gdl a}$ sont donnés en (5.54c) et (5.65).

Calcul de l'impédance caractéristique. On note par Z_c l'impédance caractéristique, qui vérifie par définition les relations suivantes :

$$C_{H_2, gdl a, 2} = -Z_c(j\omega)q_{H_2, gdl a, 2} \quad (5.77a)$$

$$C_{H_2, gdl a, 1} = Z_c(j\omega)q_{H_2, gdl a, 1} \quad (5.77b)$$

D'après (5.76) et (4.6), on obtient :

$$\begin{aligned} Z_c &= \sqrt{\frac{T_{12}}{T_{21}}} = \sqrt{\frac{R_{gdl a}^2}{j\omega\tau_{gdl a}} + R_1 R_2 + \frac{(R_1 + R_2)R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}} \tanh \sqrt{j\omega\tau_{gdl a}}} } \\ &= \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \sqrt{1 + \frac{R_1 R_2}{R_{gdl a}^2} j\omega\tau_{gdl a} + \frac{R_1 + R_2}{R_{gdl a}} \frac{\sqrt{j\omega\tau_{gdl a}}}{\tanh \sqrt{j\omega\tau_{gdl a}}}} \end{aligned} \quad (5.78)$$

Etude de la positivité de $Z(Q_{gdl a})$. On a clairement :

$$Z(Q_{gdl a}) \geq \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \begin{pmatrix} \frac{1}{\tanh \sqrt{j\omega\tau_{gdl a}}} & \frac{1}{\sinh \sqrt{j\omega\tau_{gdl a}}} \\ \frac{1}{\sinh \sqrt{j\omega\tau_{gdl a}}} & \frac{1}{\tanh \sqrt{j\omega\tau_{gdl a}}} \end{pmatrix}$$

au sens des matrices définies positives, cette dernière valeur étant obtenue pour $R_1 = R_2 = 0$. On va montrer que $Z(Q_{gdl a})$ est déjà dissipatif dans ce cas limite.

Lemme 5.11. *Le quadripôle défini par $Z(Q_{gdl a})$ est dissipatif lorsque $R_1 = R_2 = 0$.*

Démonstration. On introduit l'égalité suivante :

$$\sqrt{j\omega\tau_{gdl a}} = \alpha(1 + j), \quad \alpha \doteq \frac{\sqrt{\omega\tau_{gdl a}}}{\sqrt{2}}$$

on a alors :

$$\begin{aligned} \sinh \sqrt{j\omega\tau_{gdl a}} &= \sinh(\alpha(1 + j)) = \frac{\exp(\alpha)(\cos \alpha + j \sin \alpha) - \exp(-\alpha)(\cos \alpha - j \sin \alpha)}{2} \\ &= \sinh \alpha \cos \alpha + j \cosh \alpha \sin \alpha \end{aligned}$$

$$\begin{aligned} \cosh \sqrt{j\omega\tau_{gdl a}} &= \cosh(\alpha(1 + j)) = \frac{\exp(\alpha)(\cos \alpha + j \sin \alpha) + \exp(-\alpha)(\cos \alpha - j \sin \alpha)}{2} \\ &= \cosh \alpha \cos \alpha + j \sinh \alpha \sin \alpha \end{aligned}$$

d'où

$$\begin{aligned} \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}} \sinh \sqrt{j\omega\tau_{gdl a}}} &= \frac{R_{gdl a}}{\alpha(1 + j)} \frac{\sinh \alpha \cos \alpha - j \cosh \alpha \sin \alpha}{\sinh^2 \alpha \cos^2 \alpha + \cosh^2 \alpha \sin^2 \alpha} \\ &= \frac{R_{gdl a}}{2\alpha} \frac{(1 - j)(\sinh \alpha \cos \alpha - j \cosh \alpha \sin \alpha)}{\sinh^2 \alpha \cos^2 \alpha + \cosh^2 \alpha \sin^2 \alpha} \\ &= \frac{R_{gdl a}}{2\alpha} \frac{(\sinh \alpha \cos \alpha - \cosh \alpha \sin \alpha) - j(\sinh \alpha \cos \alpha + \cosh \alpha \sin \alpha)}{\sinh^2 \alpha \cos^2 \alpha + \cosh^2 \alpha \sin^2 \alpha}. \end{aligned}$$

On en déduit que :

$$\begin{aligned} \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}} \tanh \sqrt{j\omega\tau_{gdl a}}} &= \frac{R_{gdl a}}{2\alpha} \frac{(\sinh \alpha \cos \alpha - \cosh \alpha \sin \alpha) - j(\sinh \alpha \cos \alpha + \cosh \alpha \sin \alpha)}{\sinh^2 \alpha \cos^2 \alpha + \cosh^2 \alpha \sin^2 \alpha} \\ &\quad \times (\cosh \alpha \cos \alpha + j \sinh \alpha \sin \alpha) \\ &= \frac{R_{gdl a}}{2\alpha} \frac{(\sinh \alpha \cosh \alpha - \cos \alpha \sin \alpha) - j(\sinh \alpha \cosh \alpha + \cos \alpha \sin \alpha)}{\sinh^2 \alpha \cos^2 \alpha + \cosh^2 \alpha \sin^2 \alpha} \end{aligned}$$

On a donc

$$Z_{11} + Z_{11}^*, Z_{22} + Z_{22}^* \geq \frac{R_{gdl a}}{2\alpha} \frac{\sinh \alpha \cosh \alpha - \cos \alpha \sin \alpha}{\sinh^2 \alpha \cos^2 \alpha + \cosh^2 \alpha \sin^2 \alpha} \geq 0, \quad \forall \alpha \geq 0.$$

d'où

$$\operatorname{Re}(Z_{11}) \geq 0, \quad \operatorname{Re}(Z_{22}) \geq 0. \quad (5.79)$$

D'autre part,

$$\begin{aligned} &(Z_{11} + Z_{11}^*)(Z_{22} + Z_{22}^*) - (Z_{12} + Z_{21}^*)(Z_{21} + Z_{12}^*) \\ &\geq \left(\frac{R_{gdl a}}{\alpha}\right)^2 \frac{(\sinh \alpha \cosh \alpha - \cos \alpha \sin \alpha)^2 - (\sinh \alpha \cos \alpha - \cosh \alpha \sin \alpha)^2}{(\sinh^2 \alpha \cos^2 \alpha + \cosh^2 \alpha \sin^2 \alpha)^2} \\ &= \left(\frac{R_{gdl a}}{\alpha}\right)^2 \frac{(\sinh \alpha - \sin \alpha)(\cosh \alpha + \cos \alpha)(\sinh \alpha + \sin \alpha)(\cosh \alpha - \cos \alpha)}{(\sinh^2 \alpha \cos^2 \alpha + \cosh^2 \alpha \sin^2 \alpha)^2} \\ &= \left(\frac{R_{gdl a}}{\alpha}\right)^2 \frac{(\sinh^2 \alpha - \sin^2 \alpha)(\cosh^2 \alpha - \cos^2 \alpha)}{(\sinh^2 \alpha \cos^2 \alpha + \cosh^2 \alpha \sin^2 \alpha)^2} \geq 0, \quad \alpha \geq 0. \end{aligned}$$

d'où

$$(Z_{11} + Z_{11}^*)(Z_{22} + Z_{22}^*) \geq (Z_{12} + Z_{21}^*)(Z_{21} + Z_{12}^*). \quad (5.80)$$

□

Factorisation de la matrice de transmission : Soit le quadripôle représenté sur la Figure 5.5. On rappelle que la relation de transmission correspondant à la Figure 5.5 est

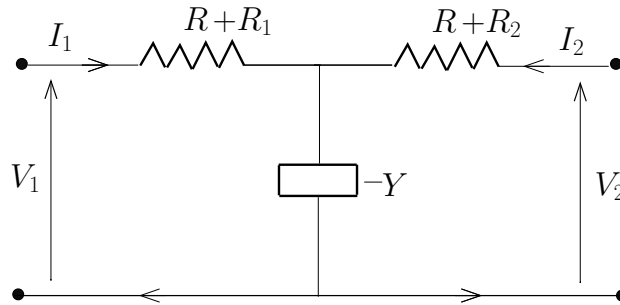


FIGURE 5.5 – Schéma électrique d'un quadripôle

alors :

$$\begin{pmatrix} V_1 \\ I_1 \end{pmatrix} = T(Q_1; Q_2; Q_3) \begin{pmatrix} V_2 \\ -I_2 \end{pmatrix} \quad (5.81)$$

où

$$T(Q_1) = \begin{pmatrix} 1 & R + R_1 \\ 0 & 1 \end{pmatrix}, \quad T(Q_2) = \begin{pmatrix} 1 & 0 \\ -Y & 1 \end{pmatrix}, \quad T(Q_3) = \begin{pmatrix} 1 & R + R_1 \\ 0 & 1 \end{pmatrix}.$$

On va maintenant factoriser $T_{gdl a}$ en utilisant le lemme suivant.

Proposition 5.12. *La matrice de transmission $T(Q_{gdl a})$ définie par (5.76) se décompose de la manière suivante :*

$$T(Q_{gdl a}) = T(Q_1; Q_2; Q_3) \quad (5.82)$$

avec

$$T(Q_1) = \begin{pmatrix} 1 & R_{Wa} + R_1 \\ 0 & 1 \end{pmatrix}, \quad T(Q_2) = \begin{pmatrix} 1 & 0 \\ -Y_{Wa} & 1 \end{pmatrix}, \quad T(Q_3) = \begin{pmatrix} 1 & R_{Wa} + R_1 \\ 0 & 1 \end{pmatrix}$$

où

$$R_{Wa}(j\omega) = \frac{R_{gdl a}}{\sqrt{j\omega\tau_{gdl a}}} \tanh\left(\frac{\sqrt{j\omega\tau_{gdl a}}}{2}\right) \quad (5.83a)$$

$$Y_{Wa}(j\omega) = \frac{\sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \sinh\sqrt{j\omega\tau_{gdl a}} \quad (5.83b)$$

et $R_{gdl a}$, $\tau_{gdl a}$ sont données par (5.54c) et (5.65). De plus cette décomposition est unique.

Démonstration. Montrons d'abord (5.82).

$$\begin{aligned} T(Q_1; Q_2; Q_3) &= \begin{pmatrix} 1 & R_{Wa} + R_1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ -Y_{Wa} & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} 1 & R_{Wa} + R_2 \\ 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 1 & R_{Wa} + R_1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ Y_{Wa} & 1 \end{pmatrix} \begin{pmatrix} 1 & R_{Wa} + R_2 \\ 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 1 + (R_{Wa} + R_1)Y_{Wa} & R_{Wa} + R_1 \\ Y_{Wa} & 1 \end{pmatrix} \begin{pmatrix} 1 & R_{Wa} + R_2 \\ 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 1 + (R_{Wa} + R_1)Y_{Wa} & (R_{Wa} + R_2)(1 + (R_{Wa} + R_1)Y_{Wa}) + R_{Wa} + R_1 \\ Y_{Wa} & 1 + (R_{Wa} + R_2)Y_{Wa} \end{pmatrix} \end{aligned}$$

On a d'après (5.83a)-(5.83b) :

$$\begin{aligned} 1 + (R_{Wa} + R_i)Y_{Wa} &= 1 + \frac{\sinh\sqrt{\frac{j\omega\tau_{gdl a}}{2}}}{\cosh\sqrt{\frac{j\omega\tau_{gdl a}}{2}}} \sinh\sqrt{j\omega\tau_{gdl a}} + \frac{R_i\sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \sinh\sqrt{j\omega\tau_{gdl a}} \\ &= 1 + 2 \sinh^2\left(\frac{\sqrt{j\omega\tau_{gdl a}}}{2}\right) + \frac{R_i\sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \sinh\sqrt{j\omega\tau_{gdl a}} \\ &= \cosh\sqrt{j\omega\tau_{gdl a}} + \frac{R_i\sqrt{j\omega\tau_{gdl a}}}{R_{gdl a}} \sinh\sqrt{j\omega\tau_{gdl a}}, \quad i = 1, 2 \end{aligned}$$

On en déduit que :

$$\begin{aligned}
 & (R_{W_a} + R_2)(1 + (R_{W_a} + R_1)Y_{W_a}) \\
 &= \left(\frac{R_{gdl_a}}{\sqrt{j\omega\tau_{gdl_a}}} \tanh\left(\frac{\sqrt{j\omega\tau_{gdl_a}}}{2}\right) + R_2 \right) \left(\cosh\sqrt{j\omega\tau_{gdl_a}} + \frac{R_1\sqrt{j\omega\tau_{gdl_a}}}{R_{gdl_a}} \sinh\sqrt{j\omega\tau_{gdl_a}} \right) \\
 &= \left(\frac{R_{gdl_a}}{\sqrt{j\omega\tau_{gdl_a}}} \tanh\left(\frac{\sqrt{j\omega\tau_{gdl_a}}}{2}\right) + R_2 \right) \left(2 \cosh^2\left(\frac{\sqrt{j\omega\tau_{gdl_a}}}{2}\right) - 1 + \frac{R_1\sqrt{j\omega\tau_{gdl_a}}}{R_{gdl_a}} \sinh\sqrt{j\omega\tau_{gdl_a}} \right) \\
 &= \frac{R_{gdl_a}}{\sqrt{j\omega\tau_{gdl_a}}} \sinh\sqrt{j\omega\tau_{gdl_a}} + 2R_2 \cosh^2\left(\frac{\sqrt{j\omega\tau_{gdl_a}}}{2}\right) + R_1 \tanh\left(\frac{\sqrt{j\omega\tau_{gdl_a}}}{2}\right) \sinh\sqrt{j\omega\tau_{gdl_a}} \\
 &\quad + \frac{R_1 R_2 \sqrt{j\omega\tau_{gdl_a}}}{R_{gdl_a}} \sinh\sqrt{j\omega\tau_{gdl_a}} - \frac{R_{gdl_a}}{\sqrt{j\omega\tau_{gdl_a}}} \tanh\left(\frac{\sqrt{j\omega\tau_{gdl_a}}}{2}\right) - R_2 \\
 &= \left(\frac{R_{gdl_a}}{\sqrt{j\omega\tau_{gdl_a}}} + \frac{R_1 R_2 \sqrt{j\omega\tau_{gdl_a}}}{R_{gdl_a}} \right) \sinh\sqrt{j\omega\tau_{gdl_a}} + R_2 \cosh\sqrt{j\omega\tau_{gdl_a}} \\
 &\quad + R_1 \tanh\left(\frac{\sqrt{j\omega\tau_{gdl_a}}}{2}\right) \sinh\sqrt{j\omega\tau_{gdl_a}} - \frac{R_{gdl_a}}{\sqrt{j\omega\tau_{gdl_a}}} \tanh\left(\frac{\sqrt{j\omega\tau_{gdl_a}}}{2}\right)
 \end{aligned}$$

on a alors :

$$\begin{aligned}
 & (R_{W_a} + R_2)(1 + (R_{W_a} + R_1)Y_{W_a}) + R_{W_a} + R_1 \\
 &= \left(\frac{R_{gdl_a}}{\sqrt{j\omega\tau_{gdl_a}}} + \frac{R_1 R_2 \sqrt{j\omega\tau_{gdl_a}}}{R_{gdl_a}} \right) \sinh\sqrt{j\omega\tau_{gdl_a}} + (R_1 + R_2) \cosh\sqrt{j\omega\tau_{gdl_a}}
 \end{aligned}$$

d'où (5.82). Réciproquement, supposons que (5.82) est vérifiée, pour des valeurs de R_{W_a} , Y_{W_a} à déterminer. On remarque que le problème d'identification revient à résoudre un système de deux équations à deux inconnues en R_{W_a} et Y_{W_a} , d'où l'unicité de la solution. $\square$

En utilisant la Proposition 5.12, on obtient le schéma électrique suivant du quadripôle GDL à l'anode :

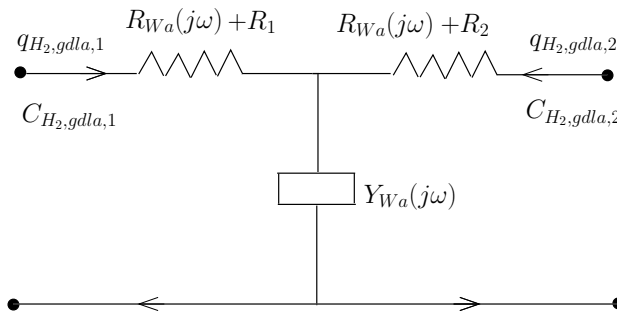


FIGURE 5.6 – Schéma électrique équivalent du GDL à l'anode

Avec cette factorisation, on vérifie d'après (5.83) et (5.73b) que :

$$Z_{ii} = R_{W_a} + \frac{1}{Y_{W_a}} + R_i, \quad i = 1, 2 \quad (5.85a)$$

et

$$Z_{12} = Z_{21} = \frac{1}{Y_{W_a}}. \quad (5.85b)$$

Lemme 5.13. Propriétés de R_{Wa} et de Y_{Wa} .

• Propriétés de R_{Wa}

$$\operatorname{Re}(R_{Wa}(j\omega)) \geq 0, \quad \forall \omega \in \mathbb{R} \quad (5.86a)$$

$$\lim_{\omega \rightarrow 0} R_{Wa}(j\omega) = \frac{R_{gdla}}{2} \quad (5.86b)$$

$$\lim_{\omega \rightarrow \infty} R_{Wa}(j\omega) = 0 \quad (5.86c)$$

$$\lim_{\omega \rightarrow \infty} \frac{\operatorname{Im}(R_{Wa}(j\omega))}{\operatorname{Re}(R_{Wa}(j\omega))} = -1 \quad (5.86d)$$

$$\lim_{\omega \rightarrow 0} \frac{\operatorname{Im}(R_{Wa}(j\omega))}{\operatorname{Re}(R_{Wa}(j\omega))} = -\infty \quad (5.86e)$$

• Propriétés de $\frac{1}{Y_{Wa}(j\omega)}$

$$-\frac{R_{gdla}}{6} < \operatorname{Re}\left(\frac{1}{Y_{Wa}(j\omega)}\right) < 0, \quad \forall \omega \in \mathbb{R} \quad (5.87a)$$

$$\lim_{\omega \rightarrow 0} \operatorname{Re}\left(\frac{1}{Y_{Wa}(j\omega)}\right) = -\frac{R_{gdla}}{6} \quad (5.87b)$$

$$\lim_{\omega \rightarrow 0} \operatorname{Im}\left(\frac{1}{Y_{Wa}(j\omega)}\right) = -\infty \quad (5.87c)$$

$$\lim_{\omega \rightarrow \infty} \frac{1}{Y_{Wa}(j\omega)} = 0. \quad (5.87d)$$

Les figures suivants montrent le tracé de R_{Wa} et $\frac{1}{Y_{Wa}}$ dans le plan de Nyquist.

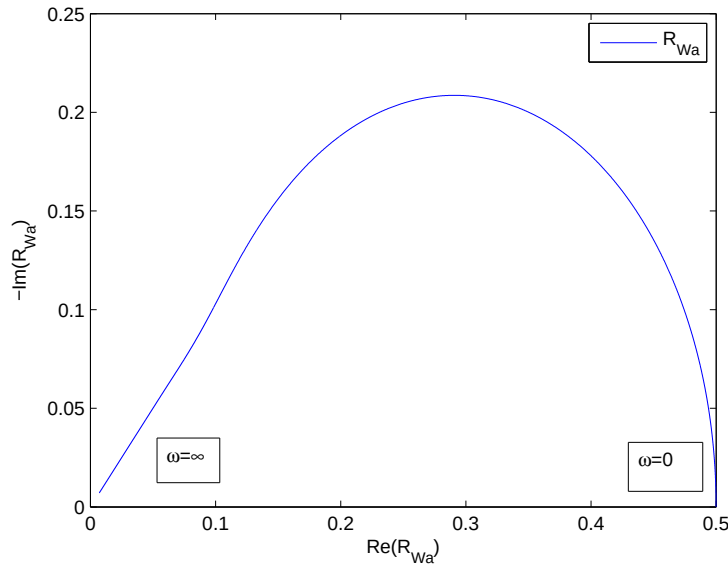
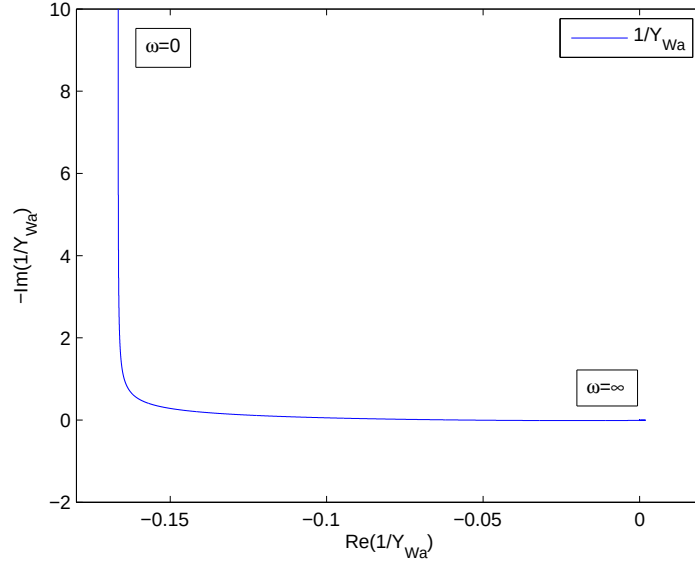


FIGURE 5.7 – Impédance R_{Wa} définie par (5.83a)

Démonstration. En utilisant l'égalité suivante :

$$\frac{\sqrt{j\omega\tau_{gdla}}}{2} = \alpha(1+j), \quad \alpha \doteq \frac{\sqrt{\omega\tau_{gdla}}}{2\sqrt{2}}$$


 FIGURE 5.8 – Impédance $\frac{1}{Y_{Wa}}$

on a :

$$\begin{aligned} \sinh \frac{\sqrt{j\omega\tau_{gdl}a}}{2} &= \sinh(\alpha(1+j)) = \sinh \alpha \cos \alpha + j \cosh \alpha \sin \alpha \\ \cosh \frac{\sqrt{j\omega\tau_{gdl}a}}{2} &= \cosh(\alpha(1+j)) = \cosh \alpha \cos \alpha + j \sinh \alpha \sin \alpha \end{aligned}$$

On a alors d'après (5.83a)

$$\begin{aligned} R_{Wa}(j\omega) &= \frac{R_{gdl}a}{2\alpha(1+j)} \tanh(\alpha(1+j)) = \frac{R_{gdl}a}{2\alpha(1+j)} \frac{\sinh \alpha \cos \alpha + j \cosh \alpha \sin \alpha}{\cosh \alpha \cos \alpha + j \sinh \alpha \sin \alpha} \\ &= \frac{R_{gdl}a}{4\alpha} \frac{\sinh \alpha \cos \alpha + \cosh \alpha \sin \alpha + j(\cosh \alpha \sin \alpha - \sinh \alpha \cos \alpha)}{\cosh \alpha \cos \alpha + j \sinh \alpha \sin \alpha} \\ &= \frac{R_{gdl}a}{4\alpha} \frac{(\sinh \alpha \cos \alpha + \cosh \alpha \sin \alpha + j(\cosh \alpha \sin \alpha - \sinh \alpha \cos \alpha))}{4\alpha(\cosh^2 \alpha \cos^2 \alpha + \sinh^2 \alpha \sin^2 \alpha)} \\ &\quad \times (\cosh \alpha \cos \alpha - j \sinh \alpha \sin \alpha) \\ &= \frac{R_{gdl}a}{4\alpha} \frac{\sinh \alpha \cosh \alpha + \sin \alpha \cos \alpha + j(\sin \alpha \cos \alpha - \sinh \alpha \cosh \alpha)}{\cosh^2 \alpha \cos^2 \alpha + \sinh^2 \alpha \sin^2 \alpha} \end{aligned} \quad (5.88)$$

De même, on a d'après (5.83b) :

$$\begin{aligned} \frac{1}{Y_{Wa}(j\omega)} &= \frac{R_{gdl}a}{\sqrt{j\omega\tau_{gdl}a} \sinh \sqrt{j\omega\tau_{gdl}a}} = \frac{R_{gdl}a}{2\alpha(1+j) \left(\sinh(2\alpha) \cos(2\alpha) + j \cosh(2\alpha) \sin(2\alpha) \right)} \\ &= \frac{R_{gdl}a}{2\alpha \left[\sinh(2\alpha) \cos(2\alpha) - \cosh(2\alpha) \sin(2\alpha) + j(\sinh(2\alpha) \cos(2\alpha) + \cosh(2\alpha) \sin(2\alpha)) \right]} \\ &= \frac{R_{gdl}a}{4\alpha} \frac{\sinh(2\alpha) \cos(2\alpha) - \cosh(2\alpha) \sin(2\alpha) - j(\sinh(2\alpha) \cos(2\alpha) + \cosh(2\alpha) \sin(2\alpha))}{\sinh^2(2\alpha) \cos^2(2\alpha) + \cosh^2(2\alpha) \sin^2(2\alpha)} \end{aligned} \quad (5.89)$$

Le développement en série entière de

$$\frac{\sinh \alpha \cosh \alpha + \sin \alpha \cos \alpha}{\alpha} = \frac{\sinh(2\alpha) + \sin(2\alpha)}{2\alpha}$$

s'écrit sous la forme :

$$\sum_{n=0}^{\infty} \frac{(2\alpha)^{2n}}{(1+2n)!} + \sum_{n=0}^{\infty} (-1)^n \frac{(2\alpha)^{2n}}{(1+2n)!}$$

On en déduit que :

$$\frac{\sinh \alpha \cosh \alpha + \sin \alpha \cos \alpha}{\alpha} \geq 0$$

d'où (5.86a).

En utilisant le développement limité au voisinage de $\alpha = 0$ des quantités suivantes :

$$\sinh \alpha \cosh \alpha = \left(\alpha + \frac{\alpha^3}{6} \right) \left(1 + \frac{\alpha^2}{2} \right) + O(\alpha^5) = \alpha + \frac{2}{3}\alpha^3 + O(\alpha^5)$$

$$\sin \alpha \cos \alpha = \left(\alpha - \frac{\alpha^3}{6} \right) \left(1 - \frac{\alpha^2}{2} \right) + O(\alpha^5) = \alpha - \frac{2}{3}\alpha^3 + O(\alpha^5)$$

$$\sinh \alpha \cos \alpha = \left(\alpha + \frac{\alpha^3}{6} \right) \left(1 - \frac{\alpha^2}{2} \right) + O(\alpha^5) = \alpha - \frac{\alpha^3}{3} + O(\alpha^5)$$

$$\cosh \alpha \sin \alpha = \left(1 + \frac{\alpha^2}{2} \right) \left(\alpha - \frac{\alpha^3}{6} \right) + O(\alpha^5) = \alpha + \frac{\alpha^3}{3} + O(\alpha^5)$$

$$\cosh^2 \alpha \cos^2 \alpha = (1 + \alpha^2) (1 - \alpha^2) + O(\alpha^6) = 1 - \alpha^4 + O(\alpha^6)$$

$$\sinh^2 \alpha \sin^2 \alpha = \left(\alpha^2 + \frac{\alpha^4}{3} \right) \left(\alpha^2 - \frac{\alpha^4}{3} \right) + O(\alpha^6) = \alpha^4 + O(\alpha^6)$$

$$\sinh^2 \alpha \cos^2 \alpha = \left(\alpha^2 + \frac{\alpha^4}{3} \right) (1 - \alpha^2) + O(\alpha^6) = \alpha^2 - \frac{2}{3}\alpha^4 + O(\alpha^6)$$

$$\cosh^2 \alpha \sin^2 \alpha = (1 + \alpha^2) \left(\alpha^2 - \frac{\alpha^4}{3} \right) + O(\alpha^6) = \alpha^2 + \frac{2}{3}\alpha^4 + O(\alpha^6)$$

on en déduit (5.86b) et (5.87). Dans (5.88) le dénominateur tend vers l'infini comme $e^{2\alpha}$ lorsque α tend vers l'infini, alors que le numérateur tend vers l'infini comme e^α d'où (5.86c). (5.86d) est obtenu par le même type de raisonnement. $\square$

5.6.4 Représentation dipolaire et impédance

Le schéma électrique correspondant au système (5.73) avec un terme source à l'entrée du port 1 est le dipôle image représenté sur la Figure 5.9.

On a alors en utilisant (5.73), le système suivant :

$$\begin{pmatrix} C_{H_2,gdla,1} \\ C_{H_2,gdla,2} \end{pmatrix} = Z(Q_{gdla}) \begin{pmatrix} q_{H_2,gdla,1} \\ q_{H_2,gdla,2} \end{pmatrix} \quad (5.91a)$$

$$C_{H_2,gdla,1} = -Z_G q_{H_2,gdla,1} + V_G \quad (5.91b)$$

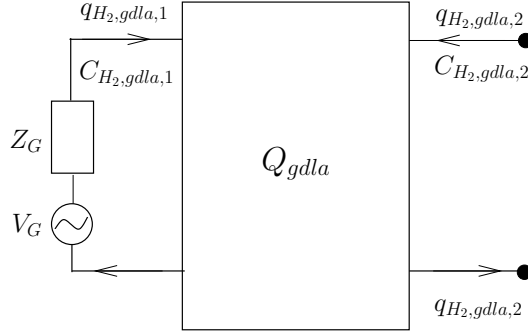


FIGURE 5.9 – Schéma électrique du dipôle image du GDL à l'anode

où $Z_{gdl a} = \begin{pmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{pmatrix}$ est définie par (5.73b). On en déduit d'après les considérations de la Section 5.6 que :

$$C_{H_2,gdla,2} = Z(G; Q_{gdl a}) q_{H_2,gdla,2} + \frac{Z_{21}}{Z_{11} + Z_G} V_G \quad (5.92a)$$

où

$$Z(G; Q_{gdl a}) \doteq Z_{22} - \frac{Z_{12} Z_{21}}{Z_{11} + Z_G}. \quad (5.92b)$$

En utilisant les égalités suivantes :

$$Z_{12} = Z_{21} = \frac{1}{Y_{W_a}}, \quad Z_{ii} = R_{W_a} + \frac{1}{Y_{W_a}} + R_i, \quad i = 1, 2$$

(5.92) s'écrit sous la forme :

$$C_{H_2,gdla,2} = Z(G; Q_{gdl a}) q_{H_2,gdla,2} + \frac{\frac{1}{Y_{W_a}}}{R_{W_a} + \frac{1}{Y_{W_a}} + R_1 + Z_G} V_G \quad (5.93a)$$

où

$$Z(G; Q_{gdl a}) = R_{W_a} + \frac{1}{Y_{W_a}} + R_2 - \frac{\frac{1}{Y_{W_a}^2}}{R_{W_a} + \frac{1}{Y_{W_a}} + R_1 + Z_G} \quad (5.93b)$$

On revient maintenant à (5.51), en exploitant la structure mise en évidence dans la section 5.6.3. Le schéma électrique correspondant au système (5.51) avec un terme source à l'entrée du port 1 est le dipôle image représenté sur la Figure 5.10. On a alors :

$$C_{H_2,gdla,1} - \frac{P_{H_2}}{H_{H_2}} + (R_{W_a}(j\omega) + R_1) q_{H_2,gdla,1} = -\frac{1}{Y_{W_a}(j\omega)} (q_{H_2,gdla,1} + q_{H_2,gdla,2}) \quad (5.94a)$$

$$C_{H_2,gdla,2} = C_{H_2,gdla,1} - \frac{P_{H_2}}{H_{H_2}} + (R_{W_a}(j\omega) + R_1) q_{H_2,gdla,1} - (R_{W_a}(j\omega) + R_2) q_{H_2,gdla,2} \quad (5.94b)$$

$$C_{H_2,gdla,1} - \frac{P_{H_2}}{H_{H_2}} = Z_s q_{H_2,gdla,1} \quad (5.94c)$$

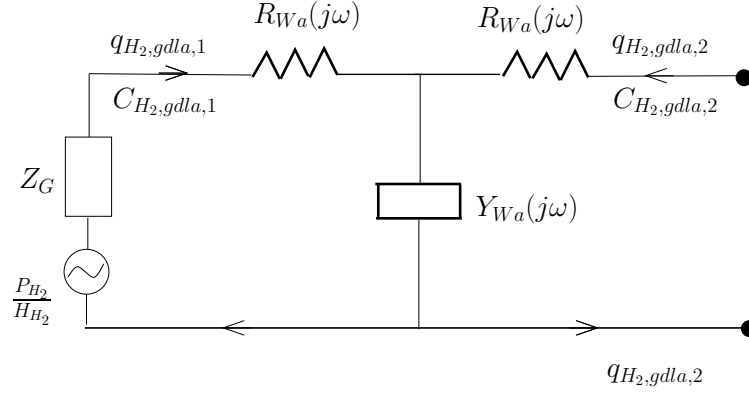


FIGURE 5.10 – Schéma électrique du dipôle image du GDL à l'anode

En utilisant les égalités suivantes :

$$Z_{12} = Z_{21} = \frac{1}{Y_{Wa}}, \quad Z_{ii} = R_{Wa} + \frac{1}{Y_{Wa}} + R_i, \quad i = 1, 2$$

on vérifie avec (4.16) que le schéma électrique représenté sur la Figure 5.10 est équivalent au schéma représenté sur la Figure 5.11, où V_G est défini par :

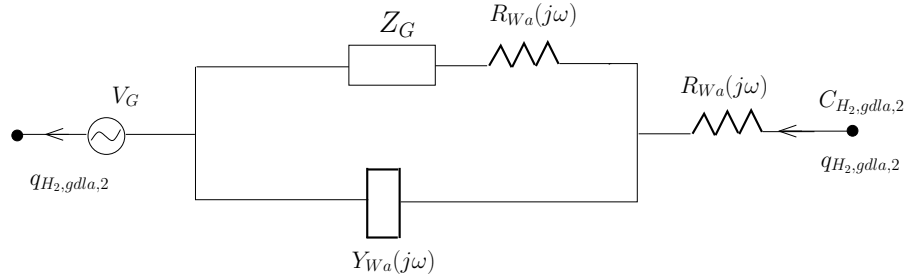


FIGURE 5.11 – Schéma électrique équivalent du dipôle image du GDL à l'anode

$$V_G = \frac{Z_{21}}{Z_{11} + Z_G} \frac{P_{H_2}}{H_{H_2}} = \frac{\frac{1}{Y_{Wa}}}{R_{Wa} + \frac{1}{Y_{Wa}} + R_1 + Z_G} \frac{P_{H_2}}{H_{H_2}}. \quad (5.95)$$

L'impédance du dipôle représenté sur la Figure 5.11 vaut (cf.(4.16)) :

$$Z(G; Q_{gdla}) = Z_{22} - \frac{Z_{12}Z_{21}}{Z_{11} + Z_G} = R_{Wa} + \frac{1}{Y_{Wa}} + R_2 - \frac{\frac{1}{Y_{Wa}^2}}{R_{Wa} + \frac{1}{Y_{Wa}} + R_1 + Z_G}$$

En particulier cette impédance vaut :

$$Z(\infty; Q_{gdla}) = Z_{22} = \frac{1}{Y_{Wa}(j\omega)} + R_{Wa}(j\omega) + R_2, \quad \text{lorsque } Z_G = \infty \quad (5.96)$$

et

$$Z(0; Q_{gdla}) = Z_{22} - \frac{Z_{12}Z_{21}}{Z_{11}} = \left(Y_{Wa}(j\omega) + \frac{1}{R_{Wa}(j\omega) + R_1} \right)^{-1} + R_{Wa}(j\omega) + R_2, \quad \text{lorsque } Z_G = 0. \quad (5.97)$$

On vérifie graphiquement que, dans un cas comme l'autre, $Z(G; Q_{gdl a})$ est passif. Ceci est en fait vrai pour tout générateur G , par suite du théorème 4.3.

Les Figures 5.12 et 5.13 montrent le tracé de $Z(\infty; Q_{gdl a})$ et $Z(0; Q_{gdl a})$ dans le plan de Nyquist pour $R_1 = R_2 = 0$.

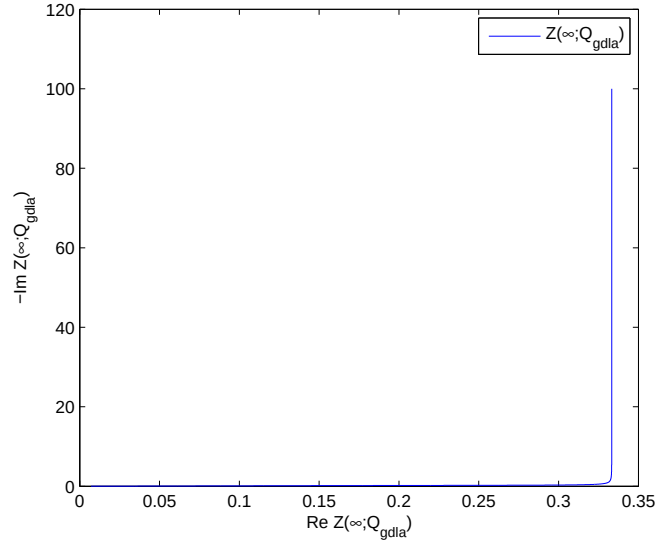


FIGURE 5.12 – Impédance $Z(\infty; Q_{gdl a})$

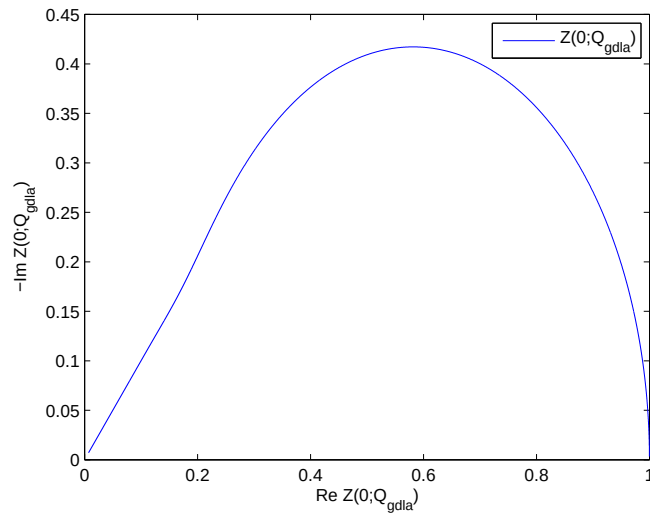


FIGURE 5.13 – Impédance $Z(0; Q_{gdl a})$

Lemme 5.14. Propriétés de $\frac{1}{Y_{Wa}(j\omega)} + R_{Wa}(j\omega)$

$$\operatorname{Re} \left(R_{Wa} + \frac{1}{Y_{Wa}} \right) > 0, \quad \forall \omega \in \mathbb{R} \quad (5.98a)$$

$$\lim_{\omega \rightarrow 0} \operatorname{Re} \left(R_{Wa} + \frac{1}{Y_{Wa}} \right) = \frac{R_{gdl a}}{3} \quad (5.98b)$$

$$\lim_{\omega \rightarrow 0} \operatorname{Im} \left(R_{Wa} + \frac{1}{Y_{Wa}} \right) = -\infty \quad (5.98c)$$

$$\lim_{\omega \rightarrow \infty} \operatorname{Im} \left(R_{Wa} + \frac{1}{Y_{Wa}} \right) = 0 \quad (5.98d)$$

$$\lim_{\omega \rightarrow \infty} \frac{\operatorname{Im} \left(R_{Wa} + \frac{1}{Y_{Wa}} \right)}{\operatorname{Re} \left(R_{Wa} + \frac{1}{Y_{Wa}} \right)} = 0 \quad (5.98e)$$

5.7 Conclusion

Dans ce chapitre nous avons calculé analytiquement l'impédance de réseaux de réactions électrochimiques et de la pile en nous basant sur les modèles obtenus dans le Chapitre 4. Compte tenu de la nature non linéaire de ces modèles, nous avons utilisé la technique de la balance harmonique qui permet d'approximer le comportement harmonique d'un système dynamique entrée-sortie non linéaire. Nous avons obtenu un modèle d'impédance de la pile incluant le phénomène de la double couche. Dans ce modèle les gaz sont supposés connectés directement à la source d'alimentation. Afin d'enrichir ce modèle nous avons ensuite calculé l'impédance de la couche de diffusion des gaz (GDL) et nous avons obtenu le schéma électrique associé.

Chapitre 6

Modélisation harmonique réduite d'un réseau de cellules

Dans ce chapitre, nous étudions la modélisation réduite d'un réseau d'impédances. Le réseau en question est constitué d'un ensemble de sous-cellules approximativement identiques, couplées par des liaisons électriques série-parallèle. Afin de détecter et d'identifier l'apparition de disparités dans le comportement de ce dernier (provenant, par exemple, de vieillissement ou de dégradation), la description du comportement électrique de l'ensemble du système par une cellule moyenne n'est pas suffisante. La principale contribution de ce chapitre est de fournir une meilleure approximation de la fonction d'impédance de l'ensemble du réseau, en introduisant des termes correctifs qui caractérisent la dispersion par rapport au comportement moyen. Nous étudions aussi le problème lié à l'identification des quantités capables de décrire le comportement moyen et la dispersion, pour les utiliser comme des données d'alerte pour la surveillance et le diagnostic. À titre d'exemple illustratif, le problème d'identification correspondant pour le modèle réduit est étudié plus en détail dans le cas où les fonctions d'impédances individuelles sont des transferts du premier ordre.

6.1 Modélisation réduite de l'impédance d'un réseau série-parallèle des cellules

6.1.1 Impédance du réseau

L'impédance électrique d'un réseau série-parallèle de sous cellules est notée par $Z_{\text{mod}}(s)$ et donnée par :

$$Z_{\text{mod}}(s) \doteq \int_{\eta} \left(\int_{\theta} z(\theta; s)^{-1} d\mu_{\eta}(\theta) \right)^{-1} d\nu(\eta) \quad (6.1)$$

où les mesures $d\mu_{\eta}$ et $d\nu$ sont positives et normalisées.

$$\int_{\theta} d\nu_{\eta}(\theta) = 1, \quad \int_{\eta} d\nu(\eta) = 1$$

pour tout η . Dans cette formule, $z(\theta; s)$ représente l'impédance d'une sous cellule dont le vecteur de paramètre vaut θ (s est la variable de Laplace, que désormais on omettra), la variable η permet d'indexer les différents types de cellules placées en série : la mesure $d\nu(\eta)$ permet de représenter la répartition de différents types. Pour avoir la possibilité de décrire les disparités internes dans la distribution des paramètres à l'intérieur de chaque cellule de la pile, on introduit la mesure $d\mu_{\eta}(\theta)$. Ces mesures peuvent ou non contenir des poids concentrés. Ainsi, chaque cellule est vue comme une population de sous cellules placées électriquement en

parallèle : de cette manière, $d\mu_\eta(\theta)$ doit être interprété comme “la proportion de sous-cellules dans la cellule η dont la valeur de paramètre est entre θ et $\theta + d\theta$ ”.

On supposera, par exemple ici, que $z(\theta; s)$ est pour tout $\theta \in \mathbb{R}^p$, une fonction de transfert réelle propre. $p \in \mathbb{N}$ est donc le nombre des paramètres scalaires permettant d'écrire la classe des fonctions de transfert $z(\theta; s)$. Dans la suite, nous considérons essentiellement un réseau de m cellules, chacune d'entre elles est subdivisée en n sous cellules, comme illustré dans la Figure 6.1.

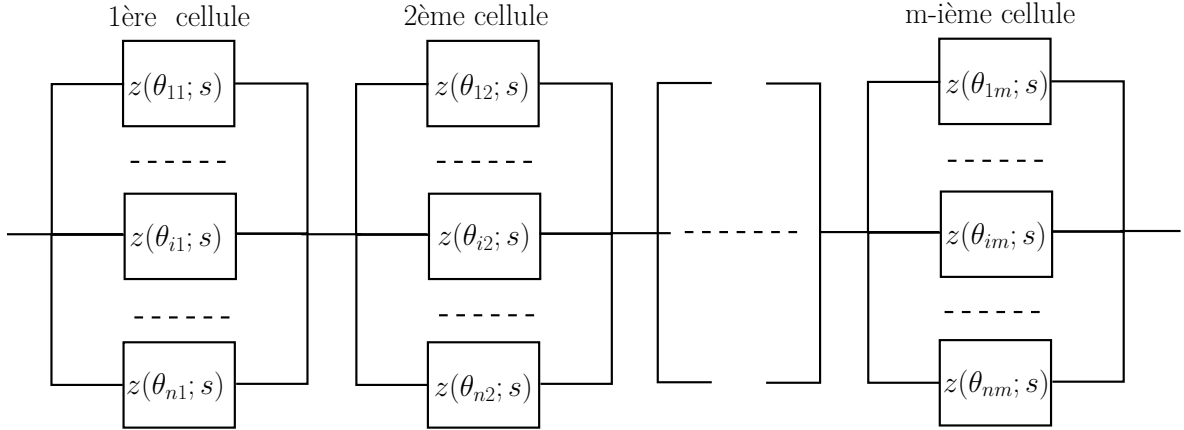


FIGURE 6.1 – Réseau série-parallèle correspondant à $Z_{\text{mod}}(s)$ dans (6.2)

La formule (6.1) s'écrit :

$$Z_{\text{mod}}(s) = \frac{1}{m} \sum_{j=1}^m \left(\frac{1}{n} \sum_{i=1}^n z(\theta_{ij}; s)^{-1} \right)^{-1} \quad (6.2)$$

ce cas correspond aux sommes des fonctions de Dirac pour les mesures $d\mu_\eta$, $d\nu$, plus précisément :

$$d\mu_\eta(\theta) = \frac{1}{n} \sum_{i=1}^n \delta_{\theta_{in}}(\theta), \quad d\nu(\eta) = \frac{1}{m} \sum_{j=1}^m \delta_j(\eta). \quad (6.3)$$

On remarque dans la formule (6.1) que, bien que ce ne soit pas explicite dans les notations, $Z_{\text{mod}}(s)$ dépend de l'ensemble de la distribution des paramètres, à travers les mesures μ_η et ν (comme dit précédemment, celles-ci indiquent la répartition des valeurs des paramètres).

On supposera qu'il existe pour tout $\theta \in \mathbb{R}^p$ et pour tout $s \in \mathbb{C}$ qui ne soit ni pôle ni zéro de $z(\theta; s)$, les gradients $\nabla_\theta z(\theta; s) \in \mathbb{R}^p$ et $\nabla_\theta z^{-1}(\theta; s) \in \mathbb{R}^p$ et la matrice hessienne $H(z) \in \mathbb{R}^{p \times p}$ (respectivement $H(z^{-1})$) de $z(\theta)$ (respectivement $z(\theta)^{-1}$) par rapport à $\theta \in \mathbb{R}^p$. Autrement dit :

$$H(z)(\theta; s) = (\partial_{kl} z)_{k,l}, \quad \partial_{kl} z = \frac{\partial^2 z(\theta; s)}{\partial \theta_k \partial \theta_l}, \quad k, l = 1, \dots, p \quad (6.4)$$

6.1.2 Modèle d'impédance réduit du réseau

Afin d'obtenir un modèle d'impédance réduit du réseau, nous effectuons dans la formule (6.1) un développement limité à des termes du 3ème ordre près au voisinage d'une valeur $\theta^* \in \mathbb{R}^p$ quelconque.

En utilisant le fait que pour une fonction f deux fois différentiable en $a \in \mathbb{R}^p$ à valeur dans $\mathbb{R}$, on peut écrire :

$$f(x) = f(a) + (x - a)^\top \nabla f(a) + \frac{1}{2}(x - a)^\top H(f)(a)(x - a) + o(\|x - a\|^2).$$

où ∇f est le gradient de f et $H(f)(a)$ est sa matrice hessienne évaluée en a , on obtient :

$$\begin{aligned} \left(\int_{\theta} z^{-1}(\theta) d\mu_{\eta}(\theta) \right)^{-1} &\approx \left(\int_{\theta} \left(z^{-1}(\theta^*) + (\theta - \theta^*)^\top \nabla z^{-1}(\theta^*) \right. \right. \\ &\quad \left. \left. + \frac{1}{2}(\theta - \theta^*)^\top H(z^{-1})(\theta^*)(\theta - \theta^*) \right) d\mu_{\eta}(\theta) \right)^{-1}. \end{aligned}$$

En prenant $\theta^* = \bar{\theta}_{\eta}$ tel que $\int_{\theta} (\theta - \bar{\theta}_{\eta}) d\mu_{\eta}(\theta) = 0$, c'est-à-dire

$$\bar{\theta}_{\eta} = \int_{\theta} \theta d\mu_{\eta}(\theta) \quad (6.5)$$

$\bar{\theta}_{\eta}$ est la valeur moyenne de θ dans la cellule η , on obtient donc à des termes d'ordre au moins égal à 3 près :

$$\begin{aligned} \left(\int_{\theta} z^{-1}(\theta) d\mu_{\eta}(\theta) \right)^{-1} &\approx \left(\int_{\theta} \left(z^{-1}(\bar{\theta}_{\eta}) + \frac{1}{2}(\theta - \bar{\theta}_{\eta})^\top H(z^{-1})(\bar{\theta}_{\eta})(\theta - \bar{\theta}_{\eta}) \right) d\mu_{\eta}(\theta) \right)^{-1} \\ &\approx z(\bar{\theta}_{\eta}) - \frac{1}{2}z^2(\bar{\theta}_{\eta}) \int_{\theta} (\theta - \bar{\theta}_{\eta})^\top H(z^{-1})(\bar{\theta}_{\eta})(\theta - \bar{\theta}_{\eta}) d\mu_{\eta}(\theta). \end{aligned}$$

On a donc, toujours en négligeant les termes d'ordre 3 :

$$\int_{\eta} \left(\int_{\theta} z^{-1}(\theta) d\mu_{\eta}(\theta) \right)^{-1} d\nu(\eta) \approx \int_{\eta} \left(z(\bar{\theta}_{\eta}) - \frac{1}{2}z^2(\bar{\theta}_{\eta}) \int_{\theta} (\theta - \bar{\theta}_{\eta})^\top H(z^{-1})(\bar{\theta}_{\eta})(\theta - \bar{\theta}_{\eta}) d\mu_{\eta}(\theta) \right) d\nu(\eta) \quad (6.6)$$

On approxime le première terme de l'expression précédente de la façon suivante. Pour tout $\theta^* \in \mathbb{R}^p$,

$$\int_{\eta} z(\bar{\theta}_{\eta}) d\nu(\eta) \approx \int_{\eta} \left(z(\theta^*) + (\bar{\theta}_{\eta} - \theta^*)^\top \nabla z(\theta^*) + \frac{1}{2}(\bar{\theta}_{\eta} - \theta^*)^\top H(z)(\theta^*)(\bar{\theta}_{\eta} - \theta^*) \right) d\nu(\eta) \quad (6.7)$$

En prenant $\theta^* = \bar{\theta}$ défini par

$$\bar{\theta} = \int_{\eta} \bar{\theta}_{\eta} d\nu(\eta) = \int_{\eta} \int_{\theta} \theta d\mu_{\eta}(\theta) d\nu(\eta) \quad (6.8)$$

$\bar{\theta}$ est la valeur moyenne de θ dans l'ensemble de réseau, on a alors

$$\int_{\eta} z(\bar{\theta}_{\eta}) d\nu(\eta) \approx z(\bar{\theta}) + \frac{1}{2} \int_{\eta} (\bar{\theta}_{\eta} - \bar{\theta})^\top H(z)(\bar{\theta})(\bar{\theta}_{\eta} - \bar{\theta}) d\nu(\eta)$$

En introduisant cette formule dans (6.6), on peut donc écrire

$$\begin{aligned} \int_{\eta} \left(\int_{\theta} z^{-1}(\theta) d\mu_{\eta}(\theta) \right)^{-1} d\nu(\eta) &\approx z(\bar{\theta}) + \frac{1}{2} \int_{\eta} (\bar{\theta}_{\eta} - \bar{\theta})^\top H(z)(\bar{\theta})(\bar{\theta}_{\eta} - \bar{\theta}) d\nu(\eta) \\ &\quad - \frac{1}{2} \int_{\eta} z^2(\bar{\theta}_{\eta}) \int_{\theta} (\theta - \bar{\theta}_{\eta})^\top H(z^{-1})(\bar{\theta}_{\eta})(\theta - \bar{\theta}_{\eta}) d\mu_{\eta}(\theta) d\nu(\eta) \end{aligned}$$

En approximant dans le dernier terme $z^2(\bar{\theta}_\eta)$ par $z^2(\bar{\theta})$, et $H(z^{-1})(\bar{\theta}_\eta)$ par $H(z^{-1})(\bar{\theta})$, on obtient finalement

$$\begin{aligned} \int_\eta \left(\int_\theta z^{-1}(\theta) d\mu_\eta(\theta) \right)^{-1} d\nu(\eta) &\approx z(\bar{\theta}) + \frac{1}{2} \int_\eta (\bar{\theta}_\eta - \bar{\theta})^\top H(z)(\bar{\theta})(\bar{\theta}_\eta - \bar{\theta}) d\nu(\eta) \\ &\quad - \frac{1}{2} z^2(\bar{\theta}) \int_\eta \int_\theta (\theta - \bar{\theta}_\eta)^\top H(z^{-1})(\bar{\theta})(\theta - \bar{\theta}_\eta) d\mu_\eta(\theta) d\nu(\eta) \end{aligned} \quad (6.9)$$

Par ailleurs, on a, pour $|z_1|^{-1} \gg |z_2|^{-1}$,

$$(z_1^{-1} + z_2^{-1})^{-1} + z_3 \sim z_1 - z_1^2 z_2^{-1} + z_3.$$

En identifiant ces expressions, on voit que l'impédance du stack est donnée par le montage dans la Figure 6.2. Il est donc possible d'approximer (6.1) par l'expression à trois termes suivante :

$$\begin{aligned} Z_{\text{red}}(s) &\doteq \left(z(\bar{\theta})^{-1} + \frac{1}{2} \int_\eta \int_\theta (\theta - \bar{\theta}_\eta)^\top H(z^{-1})(\bar{\theta})(\theta - \bar{\theta}_\eta) d\mu_\eta(\theta) d\nu(\eta) \right)^{-1} \\ &\quad + \frac{1}{2} \int_\eta (\bar{\theta} - \bar{\theta}_\eta)^\top H(z)(\bar{\theta})(\bar{\theta} - \bar{\theta}_\eta) d\nu(\eta) \end{aligned} \quad (6.10)$$

où, par simplicité, la dépendance de $z(\theta; s)$ et $z(\theta; s)^{-1}$ en s est omis. Le modèle réduit est présenté dans la Figure 6.2.

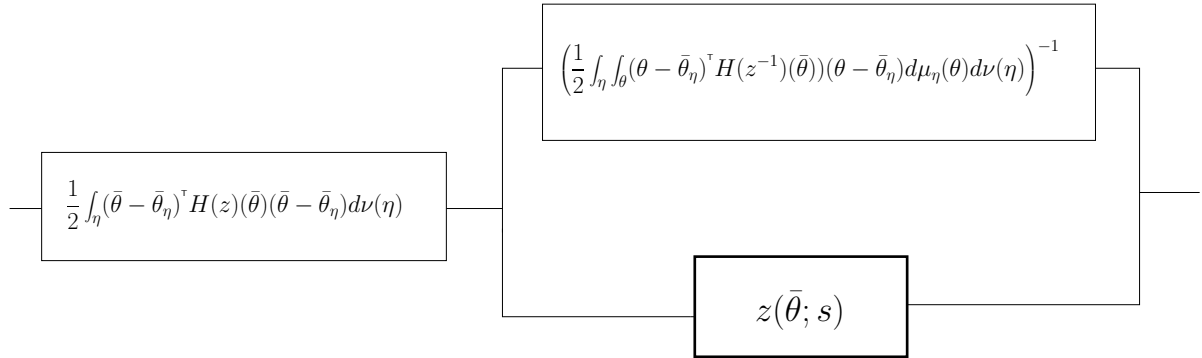


FIGURE 6.2 – Principe du modèle réduit (6.10) de (6.1). L'impédance approximée est égale à l'impédance d'une sous-cellule correspondant aux valeurs moyennes des paramètres, en parallèle et en série avec deux termes correctifs.

Dans le cas où $d\nu$, $d\mu_\eta$ sont des fonctions de Dirac, comme dans (6.3), avec $Z_{\text{red}}(s)$ écrite comme dans (6.10), on a alors :

$$\begin{aligned} Z_{\text{red}}(s) &= \frac{1}{2m} \sum_{j=1}^m (\bar{\theta} - \bar{\theta}_j)^\top H(z)(\bar{\theta})(\bar{\theta} - \bar{\theta}_j) \\ &\quad + z(\bar{\theta}) \left(1 + \frac{z(\bar{\theta})}{2mn} \sum_{j=1}^m \sum_{i=1}^n (\theta_{ij} - \bar{\theta}_j)^\top H(z^{-1})(\bar{\theta})(\theta_{ij} - \bar{\theta}_j) \right)^{-1} \end{aligned} \quad (6.11a)$$

avec

$$\bar{\theta}_j \doteq \frac{1}{n} \sum_{i=1}^n \theta_{ij}, \quad \bar{\theta} \doteq \frac{1}{m} \sum_{j=1}^m \bar{\theta}_j. \quad (6.11b)$$

Bien sûr lorsque $m = n = 1$, $Z_{\text{red}}(s) = Z_{\text{mod}}(s) = z(\bar{\theta}, s)$. La relation formelle entre l'impédance globale et son approximation est donnée par la proposition suivante.

Proposition 6.15. *Supposons que z est une fonction C^∞ par rapport à θ . Les fonctions $Z_{\text{mod}}(s)$ et $Z_{\text{red}}(s)$ sont égales à des termes près de la forme*

$$\int_{\eta} \left(\int_{\theta} (\theta - \bar{\theta})^{q_1} d\mu_{\eta}(\theta) (\bar{\theta}_j - \bar{\theta})^{q_2} \right)^{q_3} d\nu(\eta), \quad q_1, q_2, q_3 \in \mathbb{N}, \quad (q_1 + q_2)q_3 \geq 3.$$

La Proposition 6.15 résulte du développement des deux fonctions. La démonstration est lourde, mais ne présente aucun intérêt ou difficulté particulier.

6.1.3 Exemple d'un modèle simple d'impédance

Exemple 6.1 (Impédance du premier ordre). *On considère la classe d'impédances correspondant à des circuits électriques RC, données par les fonctions de transfert du premier ordre :*

$$z(\theta; s) \doteq \theta_1 + \frac{\theta_2}{s + \theta_3} \quad (6.12)$$

où $\theta_3 > 0$ (pour des raisons de stabilité) et $\theta_1, \theta_2 \geq 0$ (passivité). Ici, les paramètres θ sont a priori choisis comme un ensemble de variables permettant de représenter les fonctions d'impédance du premier ordre. De façon générale, les paramètres peuvent également être choisis directement comme des quantités physiques entrant d'une manière complexe dans certains des coefficients de la fonction d'impédance. On a donc $p = 3$ paramètres.

Pour alléger les notations, notons

$$\tilde{s} \doteq s + \theta_3 \quad (6.13)$$

on a alors

$$\nabla_{\theta} z(\theta) = \begin{pmatrix} 1 \\ \frac{1}{\tilde{s}} \\ -\frac{\theta_2}{\tilde{s}^2} \end{pmatrix}, \quad (6.14)$$

et la matrice hessienne $H(z)$ de $z(\theta; s)$ est :

$$H(z)(s) = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -\frac{1}{\tilde{s}^2} \\ 0 & -\frac{1}{\tilde{s}^2} & \frac{2\theta_2}{\tilde{s}^3} \end{pmatrix}. \quad (6.15)$$

L'inverse de $z(\theta; s)$ est :

$$z(\theta)^{-1} = \left(\theta_1 + \frac{\theta_2}{s + \theta_3} \right)^{-1} = \frac{\tilde{s}}{\theta_1 \tilde{s} + \theta_2} \quad (6.16)$$

d'où :

$$\nabla_{\theta} z(\theta)^{-1} = \begin{pmatrix} -\frac{\tilde{s}^2}{(\theta_1 \tilde{s} + \theta_2)^2} \\ -\frac{\tilde{s}}{(\theta_1 \tilde{s} + \theta_2)^2} \\ \frac{\theta_2}{(\theta_1 \tilde{s} + \theta_2)^2} \end{pmatrix} \quad (6.17)$$

et la matrice hessienne $H(z^{-1})$ de $z(\theta)^{-1}$ est :

$$H(z^{-1})(\theta) = \begin{pmatrix} \frac{2\bar{s}^3}{(\theta_1\bar{s}+\theta_2)^3} & \frac{2\bar{s}^2}{(\theta_1\bar{s}+\theta_2)^3} & -\frac{2\theta_2\bar{s}}{(\theta_1\bar{s}+\theta_2)^3} \\ \frac{2\bar{s}^2}{(\theta_1\bar{s}+\theta_2)^3} & \frac{2\bar{s}}{(\theta_1\bar{s}+\theta_2)^3} & \frac{\theta_1\bar{s}-\theta_2}{(\theta_1\bar{s}+\theta_2)^3} \\ -\frac{2\theta_2\bar{s}}{(\theta_1\bar{s}+\theta_2)^3} & \frac{\theta_1\bar{s}-\theta_2}{(\theta_1\bar{s}+\theta_2)^3} & -\frac{2\theta_1\theta_2}{(\theta_1\bar{s}+\theta_2)^3} \end{pmatrix}. \quad (6.18)$$

Comme nous l'avons vu, l'impédance totale $Z_{\text{mod}}(s)$ est approximée par (6.11a). Nous définissons, pour $k, l = 1, 2, 3$, les scalaires :

$$\forall \eta, \bar{\theta}_{k\eta} = \int_{\theta} \theta_k d\mu_{\eta}(\theta), \quad k = 1, 2, 3 \quad (6.19a)$$

$$\bar{\theta}_k = \int_{\eta} \bar{\theta}_{k\eta} d\nu(\eta), \quad k = 1, 2, 3 \quad (6.19b)$$

$$\alpha_{kl} \doteq \int_{\eta} (\bar{\theta}_k - \bar{\theta}_{k\eta})(\bar{\theta}_l - \bar{\theta}_{l\eta}) d\nu(\eta), \quad k = 1, 2, 3 \quad (6.19c)$$

$$\beta_{kl} \doteq \int_{\eta} \int_{\theta} (\theta_k - \bar{\theta}_{k\eta})(\theta_l - \bar{\theta}_{l\eta}) d\mu_{\eta}(\theta) d\nu(\eta), \quad k = 1, 2, 3 \quad (6.19d)$$

On a alors :

$$\begin{aligned} z(\bar{\theta})^{-1} &+ \frac{1}{2} \int_{\eta} \int_{\theta} (\theta - \bar{\theta}_{\eta})^T H(z^{-1})(\bar{\theta})(\theta - \bar{\theta}_{\eta}) d\mu_{\eta}(\theta) d\nu(\eta) \\ &= \frac{\bar{s}}{\bar{\theta}_1\bar{s} + \bar{\theta}_2} + \frac{\beta_{11}\bar{s}^3 + \beta_{22}\bar{s} - \bar{\theta}_1\bar{\theta}_2\beta_{33} + 2\beta_{12}\bar{s}^2 - 2\bar{\theta}_2\beta_{13}\bar{s} + \beta_{23}(\bar{\theta}_1\bar{s} - \bar{\theta}_2)}{(\bar{\theta}_1\bar{s} + \bar{\theta}_2)^3} \\ &= \frac{\bar{\theta}_1^2\bar{s}^3 + 2\bar{\theta}_1\bar{\theta}_2\bar{s}^2 + \bar{\theta}_2^2\bar{s} + \beta_{11}\bar{s}^3 + \beta_{22}\bar{s} - \bar{\theta}_1\bar{\theta}_2\beta_{33} + 2\beta_{12}\bar{s}^2 - 2\bar{\theta}_2\beta_{13}\bar{s} + \beta_{23}(\bar{\theta}_1\bar{s} - \bar{\theta}_2)}{(\bar{\theta}_1\bar{s} + \bar{\theta}_2)^3} \\ &= \frac{(\bar{\theta}_1^2 + \beta_{11})\bar{s}^3 + 2(\bar{\theta}_1\bar{\theta}_2 + \beta_{12})\bar{s}^2 + (\bar{\theta}_2^2 + \beta_{22} - 2\bar{\theta}_2\beta_{13} + \bar{\theta}_1\beta_{23})\bar{s} - (\bar{\theta}_1\bar{\theta}_2\beta_{33} + \bar{\theta}_2\beta_{23})}{(\bar{\theta}_1\bar{s} + \bar{\theta}_2)^3} \end{aligned}$$

On en déduit que, pour le modèle d'impédance donné par (6.12) :

$$\begin{aligned} Z_{\text{red}}(s) &= \frac{\bar{\theta}_2\alpha_{33}}{\bar{s}^3} - \frac{\alpha_{23}}{\bar{s}^2} \\ &+ \frac{(\bar{\theta}_1\bar{s} + \bar{\theta}_2)^3}{(\bar{\theta}_1^2 + \beta_{11})\bar{s}^3 + 2(\bar{\theta}_1\bar{\theta}_2 + \beta_{12})\bar{s}^2 + (\bar{\theta}_2^2 + \beta_{22} - 2\bar{\theta}_2\beta_{13} + \bar{\theta}_1\beta_{23})\bar{s} - (\bar{\theta}_1\bar{\theta}_2\beta_{33} + \bar{\theta}_2\beta_{23})} \end{aligned} \quad (6.20)$$

Remarquons que, en éliminant $\bar{s}$ via la formule $z(\bar{\theta}; s) = \bar{\theta}_1 + \frac{\bar{\theta}_2}{s + \bar{\theta}_3}$ qui est $\frac{\theta_1\bar{s} + \bar{\theta}_2}{\bar{s}}$ on peut écrire aussi :

$$\begin{aligned} Z_{\text{red}}(s) &= \frac{\alpha_{33}}{\bar{\theta}_2^2} z(\bar{\theta})^3 - \left(3\frac{\bar{\theta}_1}{\bar{\theta}_2^2} \alpha_{33} + \frac{1}{\bar{\theta}_2^2} \alpha_{23} \right) z(\bar{\theta})^2 + \left(3\frac{\bar{\theta}_1^2}{\bar{\theta}_2^2} \alpha_{33} + 2\frac{\bar{\theta}_1}{\bar{\theta}_2^2} \alpha_{23} \right) z(\bar{\theta}) - \frac{\bar{\theta}_1^3}{\bar{\theta}_2^2} \alpha_{33} - \frac{\bar{\theta}_1^2}{\bar{\theta}_2^2} \alpha_{23} \\ &+ z(\bar{\theta}) \left[1 + \frac{\beta_{22}}{\bar{\theta}_2^2} - 2\frac{\beta_{13}}{\bar{\theta}_2} + 4\frac{\bar{\theta}_1\beta_{23}}{\bar{\theta}_2^2} + 3\frac{\bar{\theta}_1^2\beta_{33}}{\bar{\theta}_2^2} - \left(\frac{\bar{\theta}_1\beta_{33} + \beta_{23}}{\bar{\theta}_2^2} \right) z(\bar{\theta}) \right. \\ &\quad \left. + \left(2\frac{\beta_{12}}{\bar{\theta}_2} - 2\frac{\bar{\theta}_1}{\bar{\theta}_2^2} (\beta_{22} - 2\bar{\theta}_2\beta_{13}) - 5\frac{\bar{\theta}_1^2}{\bar{\theta}_2^2} \beta_{23} - 3\frac{\bar{\theta}_1^3}{\bar{\theta}_2^2} \beta_{33} \right) z(\bar{\theta})^{-1} \right. \\ &\quad \left. + \left(\beta_{11} + \frac{\bar{\theta}_1^2}{\bar{\theta}_2^2} (\beta_{22} - 2\bar{\theta}_2\beta_{13}) + \frac{\bar{\theta}_1^3}{\bar{\theta}_2^2} (\bar{\theta}_1\beta_{33} + 2\beta_{23}) \right) z(\bar{\theta})^{-2} \right]^{-1} \end{aligned}$$

Cette expression qui dépend de $z(\bar{\theta}; s)$ rend évident le fait que $Z_{\text{red}}(s) = z(\bar{\theta}; s)$ en l'absence de dispersion, quand tous les coefficients α, β sont nuls.

La Figure 6.3 donne des résultats obtenues pour un réseau “10 × 10” (c.-à-d. $m = n = 10$), avec $\theta_{1,ij} = 9 \pm 30\%$, $\theta_{2,ij} = 6 \pm 30\%$, $\theta_{3,ij} = 1 \pm 30\%$, les variations étant générées aléatoirement. Dans ce cas, les identités (6.19) s'écrivent simplement par :

$$\alpha_{kl} \doteq \frac{1}{m} \sum_{j=1}^m (\bar{\theta}_k - \bar{\theta}_{k,j})(\bar{\theta}_l - \bar{\theta}_{l,j}), \quad \beta_{kl} \doteq \frac{1}{mn} \sum_{j=1}^m \sum_{i=1}^n (\theta_{k,ij} - \bar{\theta}_{k,j})(\theta_{l,ij} - \bar{\theta}_{l,j}), \quad k, l = 1, 2, 3 \quad (6.21a)$$

$$\bar{\theta}_{k,j} \doteq \frac{1}{n} \sum_{i=1}^n \theta_{k,ij}, \quad \bar{\theta}_k = \frac{1}{m} \sum_{j=1}^m \bar{\theta}_{k,j}, \quad k, l = 1, 2, 3 \quad (6.21b)$$

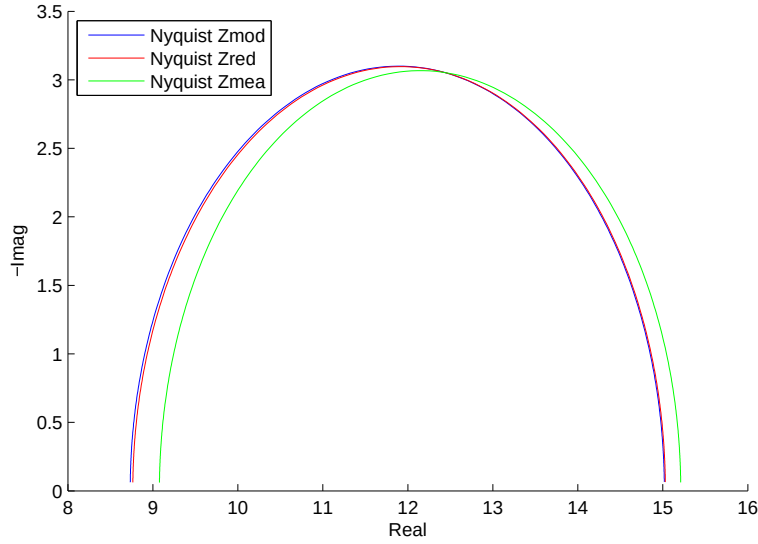


FIGURE 6.3 – Les tracés dans le plan de Nyquist de Z_{mod} , Z_{red} et $z(\bar{\theta}, s)$

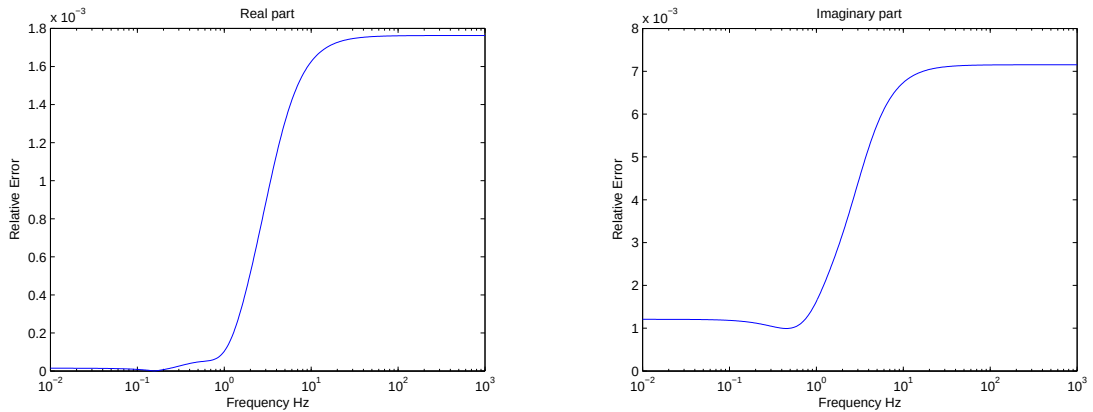


FIGURE 6.4 – Erreur relative entre Z_{mod} et Z_{red} en partie réelle et partie imaginaire

La Figure 6.3 montre dans le plan de Nyquist la courbe d'impédance exacte (6.2), en *bleu* ; la courbe d'impédance obtenue en utilisant uniquement le terme moyen $z(\bar{\theta}; s)$ ($\bar{\theta}$ est calculé

exactement), en *vert*; et l'impédance du modèle réduit (6.11) (qui pour une impédance d'une sous cellule du premier ordre, s'écrit comme (6.20), et pour les valeurs exactes de $\bar{\theta}$, α , β), en *rouge*. L'approximation de $Z_{\text{mod}}(s)$ par $Z_{\text{red}}(s)$ est meilleure (numériquement, on vérifie que l'erreur maximale en partie réelle et partie imaginaire est inférieure à 10^{-2} dans toute la plage de fréquences étudiée) que lorsqu'on considère le terme moyen $z(\bar{\theta}, s)$ seul.

6.2 Identification de l'impédance réduite du réseau

6.2.1 Méthodologie

On suppose que des observations à partir des mesures de spectroscopie d'impédance sont disponibles, essentiellement le diagramme de Nyquist $\omega \mapsto Z_{\text{obs}}(j\omega)$, ou plus généralement un nombre fini des points $Z_{\text{obs}}(j\omega)$, $\omega \in \Omega$ de ce diagramme. Notre objectif est d'identifier les valeurs du paramètre moyen $\bar{\theta}$, ainsi que les dispersions de ce paramètre données par α et β (voir (6.19)). Nous proposons une procédure d'identification en deux étapes.

Tout d'abord, les mesures sont utilisées pour ajuster de manière optimale les paramètres d'une fonction d'impédance $Z_{\text{par}}(s)$ choisie dans une classe appropriée. Ce problème (ajustement de façon optimale d'une courbe paramétrée à des mesures) renvoie à des questions classiques dans l'identification et n'est pas spécifiquement exploré dans ce travail.

La deuxième étape consiste à considérer l'identification de $\bar{\theta}$, α et β eux-mêmes : c'est la question de l'existence et l'unicité des valeurs de ces variables pour lesquelles la fonction correspondante $Z_{\text{red}}(s)$ est égale à $Z_{\text{par}}(s)$. Bien entendu, la classe des fonctions choisies pour $Z_{\text{par}}(s)$ est choisie en fonction de la structure de $Z_{\text{red}}(s)$, et donc finalement du modèle d'impédance à l'échelle des sous-cellules.

Dorénavant, on fait apparaître la fonction d'impédance comme une fonction rationnelle, en écrivant :

$$z(\theta; s) \doteq \frac{N(\theta; s)}{D(\theta; s)}, \quad (6.22)$$

où $N(\theta; s)$ et $D(\theta; s)$ sont des polynômes par rapport à la variable de Laplace s .

6.2.2 Structure du problème d'identification

Nous exposons tout d'abord un résultat technique, dont la démonstration est omise.

Lemme 6.16. *La hessienne d'un rapport de fonctions $\frac{N}{D}$, est donnée par la formule*

$$H\left(\frac{N}{D}\right) = \frac{1}{d}H(N) - \frac{N}{D^2}H(D) - \frac{1}{D^2}(\nabla N \cdot \nabla D^\top + \nabla D \cdot \nabla N^\top) + 2\frac{N}{D^3}\nabla D \cdot \nabla D^\top.$$

Nous pouvons alors déduire le résultat suivant.

Proposition 6.17. *Pour $z(s)$ comme dans (6.22), $Z_{\text{red}}(s)$ définie par (6.10) s'écrit :*

$$Z_{\text{red}}(s) = N^3 \left(N^2 + \frac{1}{2} \sum_{k,l} \beta_{kl} \left\{ N^2 \partial_{kl} D - N D \partial_{kl} N - N (\partial_k D \partial_l N + \partial_k N \partial_l D) + 2 D \partial_k N \partial_l N \right\} \right)^{-1} \\ + \frac{1}{2} \sum_{k,l} \alpha_{kl} \left\{ \frac{1}{D} \partial_{kl} D - \frac{N}{D^2} \partial_{kl} D - \frac{1}{D^2} (\partial_k D \partial_l N + \partial_k N \partial_l D) + 2 \frac{N}{D^3} \partial_k D \partial_l D \right\} \quad (6.23)$$

où les fonctions N et D ainsi que leurs dérivées dans (6.23) sont considérées au point $\bar{\theta}$.

En gardant à l'esprit que les inconnues du problème d'identification sont le vecteur des paramètres moyens $\bar{\theta}$ et les coefficients α, β , on remarque que pour les fonctions rationnelles comme dans (6.22), le numérateur de (6.23) ne dépendra que de $\bar{\theta}$ et α , tandis que dénominateur ne dépendra que de $\bar{\theta}$ et β . Cette propriété est illustrée et expliquée maintenant, à partir de l'exemple 6.1.

Exemple 6.1 (suite). En réduisant au même dénominateur, (6.20) conduit à :

$$Z_{\text{red}}(s) = \frac{N(s)}{D(s)}$$

avec

$$\begin{aligned} D(s) &\doteq (\bar{\theta}_1^2 + \beta_{11})\tilde{s}^3 + 2(\bar{\theta}_1\bar{\theta}_2 + \beta_{12})\tilde{s}^2 + (\bar{\theta}_2^2 + \beta_{22} - 2\bar{\theta}_2\beta_{13} + \bar{\theta}_1\beta_{23})\tilde{s} - (\bar{\theta}_1\bar{\theta}_2\beta_{33} + \bar{\theta}_2\beta_{23}) \\ N(s) &\doteq (\bar{\theta}_2\alpha_{33} - \alpha_{23}\tilde{s}) d(s) + \tilde{s}^3(\bar{\theta}_1\tilde{s} + \bar{\theta}_2)^3. \end{aligned}$$

Le dénominateur $Z_{\text{red}}(s)$ est un polynôme de degré $\deg D = 3$ en $\tilde{s}$ (ainsi qu'en s), et le numérateur est un polynôme de degré $\deg N = 6$. Il semble ainsi raisonnable d'essayer d'identifier ces coefficients à une expression donnée de la forme :

$$Z_{\text{par}}(s) \doteq \frac{\tilde{s}^6 + y_1\tilde{s}^5 + y_2\tilde{s}^4 + y_3\tilde{s}^3 + y_4\tilde{s}^2 + y_5\tilde{s} + y_6}{\tilde{s}^3(x_0\tilde{s}^3 + x_1\tilde{s}^2 + x_2\tilde{s} + x_3)}. \quad (6.24)$$

Ici les constantes x_0, x_1, x_2, x_3 et $y_1, y_2, y_3, y_4, y_5, y_6$ sont supposées avoir été déduites des mesures de spectroscopie d'impédance, de même que $\bar{\theta}_3$, la racine triple du dénominateur (rappelons que la variable $\tilde{s} = s + \bar{\theta}_3$ a été introduite à la place de s). Rappelons que cette étape est supposée être la première étape de la procédure d'identification.

Correspondant aux dix coefficients dans (6.24), dix identités peuvent ainsi être déduites, rendant évidente la propriété de découplage entre α et β mentionnée ci-dessus. Pour le dénominateur :

$$\frac{1}{\bar{\theta}_1} + \frac{\beta_{11}}{\bar{\theta}_1^3} = x_0 \quad (6.25a)$$

$$2 \left(\frac{1}{\bar{\theta}_1} \frac{\bar{\theta}_2}{\bar{\theta}_1} + \frac{\beta_{12}}{\bar{\theta}_1^3} \right) = x_1 \quad (6.25b)$$

$$\frac{1}{\bar{\theta}_1} \left(\frac{\bar{\theta}_2}{\bar{\theta}_1} \right)^2 + \frac{\beta_{22}}{\bar{\theta}_1^3} - 2 \frac{\bar{\theta}_2}{\bar{\theta}_1} \frac{\beta_{13}}{\bar{\theta}_1^2} + \frac{\beta_{23}}{\bar{\theta}_1^2} = x_2 \quad (6.25c)$$

$$\frac{\bar{\theta}_2}{\bar{\theta}_1} \frac{\beta_{33}}{\bar{\theta}_1} + \frac{\bar{\theta}_2}{\bar{\theta}_1} \frac{\beta_{23}}{\bar{\theta}_1^2} + x_3 = 0 \quad (6.25d)$$

et pour le numérateur :

$$3 \frac{\bar{\theta}_2}{\bar{\theta}_1} = y_1 \quad (6.26a)$$

$$3 \left(\frac{\bar{\theta}_2}{\bar{\theta}_1} \right)^2 - \frac{\alpha_{23}}{\bar{\theta}_1^3} x_0 = y_2 \quad (6.26b)$$

$$\left(\frac{\bar{\theta}_2}{\bar{\theta}_1} \right)^3 + \frac{\bar{\theta}_2}{\bar{\theta}_1} \frac{\alpha_{33}}{\bar{\theta}_1^2} x_0 - \frac{\alpha_{23}}{\bar{\theta}_1^3} x_1 = y_3 \quad (6.26c)$$

$$\frac{\bar{\theta}_2}{\bar{\theta}_1} \frac{\alpha_{33}}{\bar{\theta}_1^2} x_1 - \frac{\alpha_{23}}{\bar{\theta}_1^3} x_2 = y_4 \quad (6.26d)$$

$$\frac{\bar{\theta}_2}{\bar{\theta}_1} \frac{\alpha_{33}}{\bar{\theta}_1^2} x_2 - \frac{\alpha_{23}}{\bar{\theta}_1^3} x_3 = y_5 \quad (6.26e)$$

$$\frac{\bar{\theta}_2}{\bar{\theta}_1} \frac{\alpha_{33}}{\bar{\theta}_1^2} x_3 = y_6. \quad (6.26f)$$

En se référant à (6.23), on voit clairement que l'expression de $Z_{\text{red}}(s)$ est très structurée. Le dénominateur (commun) est

$$\left(N^2 + \frac{1}{2} \sum_{k,l} \beta_{kl} \left\{ N^2 \partial_{kl} D - N D \partial_{kl} N - N (\partial_k D \partial_l N + \partial_k N \partial_l D) + 2 D \partial_k N \partial_l N \right\} \right) D(s)^3$$

et le numérateur

$$N^3 + \frac{1}{2} \left(N^2 + \frac{1}{2} \sum_{k,l} \beta_{kl} \left\{ N^2 \partial_{kl} D - N D \partial_{kl} N - N (\partial_k D \partial_l N + \partial_k N \partial_l D) + 2 D \partial_k N \partial_l N \right\} \right) \\ \times \sum_{k,l} \alpha_{kl} \left\{ D^2 \partial_{kl} N - N D \partial_{kl} D - D (\partial_k D \partial_l N + \partial_k N \partial_l D) + 2 N \partial_k D \partial_l D \right\}.$$

Le degré (en s) du dénominateur est au plus égale à $4 \deg D + 2 \deg N$, et le degré du numérateur à $\deg N + 2 \max\{\deg N, \deg D\}$. (Rappelons que nous considérons des degrés par rapport à s , tandis que les dérivées écrites avec le symbole ∂ sont effectuées par rapport aux paramètres.)

Le premier facteur dans le second terme du numérateur vient du dénominateur : ses coefficients peuvent être considéré comme connus. En utilisant cette astuce, on voit que les équations obtenues sont systématiquement *linéaires par rapport à α et β* .

En ce qui concerne la question de leur sur-ou sous-détermination, il ne semble pas évident d'y répondre en général. Cependant, l'exemple précédent montre une structure supplémentaire très particulière.

Exemple 6.1 (suite). La formule (6.26a) s'écrit :

$$\frac{\bar{\theta}_2}{\bar{\theta}_1} = \frac{y_1}{3} \quad (6.27a)$$

d'après (6.26f) on obtient :

$$\frac{\bar{\theta}_2}{\bar{\theta}_1} \frac{\alpha_{33}}{\bar{\theta}_1^2} = \frac{y_6}{x_3}$$

et donc

$$\frac{\alpha_{33}}{\bar{\theta}_1^2} = \frac{3y_6}{x_3 y_1} \quad (6.27b)$$

Ainsi, par (6.26e) :

$$\frac{\alpha_{23}}{\bar{\theta}_1^3} = \frac{x_2}{x_3} \frac{\bar{\theta}_2}{\bar{\theta}_1} \frac{\alpha_{33}}{\bar{\theta}_1^2} - \frac{y_5}{x_3} = \frac{x_2 y_6}{x_3^2} - \frac{y_5}{x_3} \quad (6.27c)$$

Les identités (6.26b) à (6.26d) apparaissent comme des contraintes, données par :

$$\begin{aligned} 3 \left(\frac{y_1}{3} \right)^2 - \left(\frac{x_2 y_6}{x_3^2} - \frac{y_5}{x_3} \right) x_0 &= y_2 \\ \left(\frac{y_1}{3} \right)^3 + \frac{y_6}{x_3} x_0 - \left(\frac{x_2 y_6}{x_3^2} - \frac{y_5}{x_3} \right) x_1 &= y_3 \\ \frac{x_1 y_6}{x_3} - x_2 \left(\frac{x_2 y_6}{x_3^2} - \frac{y_5}{x_3} \right) &= y_4 \end{aligned}$$

soit :

$$\frac{1}{3} x_3^2 y_1^2 + x_3 y_5 x_0 = x_3^2 y_2 + x_2 y_6 x_0 \quad (6.28a)$$

$$\frac{1}{27} x_3^2 y_1^3 + x_0 x_3 y_6 + x_1 x_3 y_5 = x_3^2 y_3 + x_1 x_2 y_6 \quad (6.28b)$$

$$x_1 x_3 y_6 + x_2 x_3 y_5 = x_3^2 y_4 + x_2^2 y_6. \quad (6.28c)$$

Les trois conditions de compatibilité (6.28), associés aux trois relations (6.27) sont équivalentes aux six formules en (6.26).

Les conditions de compatibilité peuvent être comparées à des relations de parité, leur violation (ou l'impossibilité de trouver x_r , $r = 0, \dots, 3$ et y_r , $r = 1, \dots, 6$ qui ajustent convenablement les données expérimentales) peut être le signe que 'quelque chose ne va pas'. Ici, elles sont intégrés à la première étape de la procédure d'identification : on cherche directement x_r , $r = 0, \dots, 3$ et y_r , $r = 1, \dots, 6$, qui correspondent le mieux aux données expérimentales parmi les valeurs qui remplissent les contraintes algébriques précédentes. Cependant, ces dernières sont non linéaires, et de manière générale, cela ne garantit pas que le problème d'optimisation correspondant est bien soluble (en particulier, il n'y aura pas en général la garantie d'absence de minima locaux).

Nous montrons sur la Figure 6.5 une comparaison entre le diagramme de Nyquist du modèle exact (le même que dans la Figure 6.3), en rouge, avec la fonction de type (6.24) (en bleu), où x_r , $r = 0, \dots, 3$ et y_r , $r = 1, \dots, 6$ sont choisis de tel façon qu'ils minimisent la somme des distances euclidiennes

$$\sum_{\omega \in \Omega} \|Z_{\text{mes}}(j\omega) - Z_{\text{red}}(j\omega)\|^2$$

tout en respectant les contraintes (6.28). Ω est l'ensemble des fréquences, fixé et discret, ici égal à $\{10^{-2}, 10^{-1.9}, 10^{-1.8}, \dots, 10^3\}$. Un très bon accord peut être obtenu.

En ce qui concerne maintenant les formules (6.25), avec l'introduction de la valeur de $\frac{\bar{\theta}_2}{\bar{\theta}_1}$ précédemment trouvée dans (6.27), elles peuvent être réécrites comme :

$$\frac{1}{\bar{\theta}_1} + \frac{\beta_{11}}{\bar{\theta}_1^3} = x_0 \quad (6.29a)$$

$$2 \left(\frac{y_1}{3} \frac{1}{\bar{\theta}_1} + \frac{\beta_{12}}{\bar{\theta}_1^3} \right) = x_1 \quad (6.29b)$$

$$\left(\frac{y_1}{3} \right)^2 \frac{1}{\bar{\theta}_1} + \frac{\beta_{22}}{\bar{\theta}_1^3} - 2 \frac{y_1}{3} \frac{\beta_{13}}{\bar{\theta}_1^2} + \frac{\beta_{23}}{\bar{\theta}_1^2} = x_2 \quad (6.29c)$$

$$\frac{y_1}{3} \frac{\beta_{33}}{\bar{\theta}_1} + \frac{y_1}{3} \frac{\beta_{23}}{\bar{\theta}_1^2} + x_3 = 0. \quad (6.29d)$$

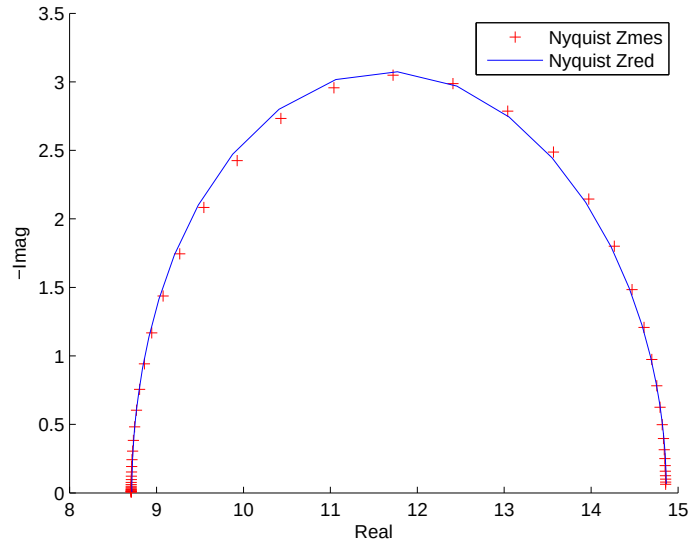


FIGURE 6.5 – Comparaison entre les données exactes et la courbe d'impédance identifiée

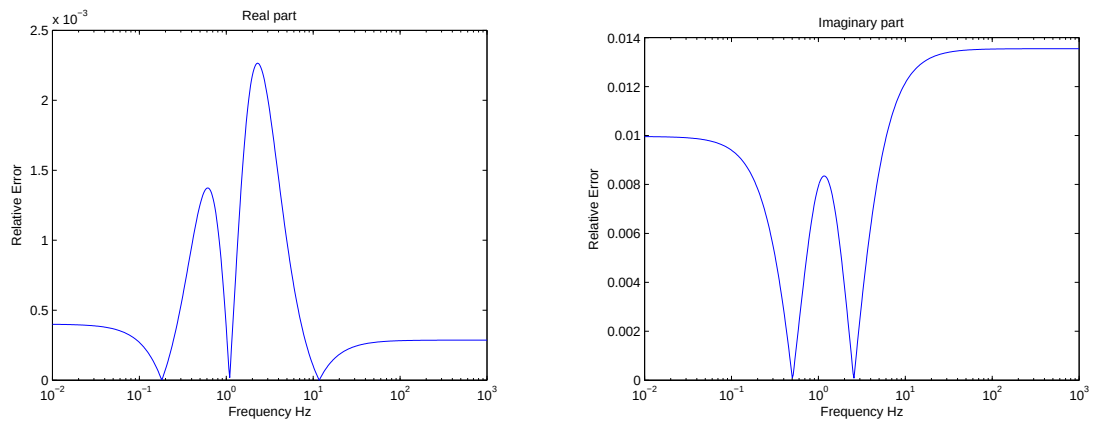


FIGURE 6.6 – Erreur relative entre Z_{mes} et Z_{red} en partie réelle et partie imaginaire

Les quatre quantités du côté gauche des identités précédentes sont ainsi identifiables, et l'on peut maintenant résumer la liste des huit quantités qui peuvent être identifiées à partir des mesures. Il s'agit de

$$\bar{\theta}_3, \quad \frac{\bar{\theta}_2}{\theta_1}, \quad \frac{\alpha_{33}}{\bar{\theta}_1^2}, \quad \frac{\alpha_{23}}{\bar{\theta}_1^3} \quad (6.30a)$$

$$\frac{1}{\theta_1} + \frac{\beta_{11}}{\bar{\theta}_1^3}, \quad \frac{\bar{\theta}_2}{\bar{\theta}_1^2} + \frac{\beta_{12}}{\bar{\theta}_1^3}, \quad (6.30b)$$

$$\frac{\bar{\theta}_2^2}{\bar{\theta}_1^3} + \frac{\beta_{22}}{\bar{\theta}_1^3} - 2\frac{\bar{\theta}_2\beta_{13}}{\bar{\theta}_1^3} + \frac{\beta_{23}}{\bar{\theta}_1^2}, \quad \frac{\beta_{33}}{\bar{\theta}_1} + \frac{\beta_{23}}{\bar{\theta}_1^2}. \quad (6.30c)$$

Application numérique : identification de l'exemple étudié. L'identification est réalisée comme présenté ci-dessus, sur les données déjà générées dans la section 6.1.3, voir la Figure 6.1. Comme on le voit, la comparaison entre les valeurs identifiées et les valeurs exactes est très bonne pour cinq grandeurs.

Quantité	Valeur exacte	Valeur identifiée	Erreur relative
$\bar{\theta}_3$	1.0757	1.0485	2.5%
$\frac{\bar{\theta}_2}{\theta_1}$	0.6611	0.6758	2.2%
$\frac{1}{\theta_1} + \frac{\beta_{11}}{\bar{\theta}_1^3}$	0.1105	0.1119	1.3%
$\frac{\bar{\theta}_2}{\bar{\theta}_1^2} + \frac{\beta_{12}}{\bar{\theta}_1^3}$	0.0724	0.0763	5.4%
$\frac{\bar{\theta}_2^2}{\bar{\theta}_1^3} + \frac{\beta_{22}}{\bar{\theta}_1^3} - 2\frac{\bar{\theta}_2\beta_{13}}{\bar{\theta}_1^3} + \frac{\beta_{23}}{\bar{\theta}_1^2}$	0.0495	0.0475	4.0%
$\frac{\beta_{33}}{\bar{\theta}_1} + \frac{\beta_{23}}{\bar{\theta}_1^2}$	0.0032	0.0039	22%

FIGURE 6.7 – Les valeurs calculées et exactes de six des huit quantités identifiables dans (6.30)

Cependant, concernant les paramètres $\bar{\theta}, \alpha, \beta$ eux-mêmes, il reste une sous-détermination, avec les quatre identités linéaires (6.29) reliant les sept inconnues

$$\frac{1}{\theta_1}, \quad \frac{\beta_{11}}{\bar{\theta}_1^3}, \quad \frac{\beta_{12}}{\bar{\theta}_1^3}, \quad \frac{\beta_{22}}{\bar{\theta}_1^3}, \quad \frac{\beta_{13}}{\bar{\theta}_1^2}, \quad \frac{\beta_{23}}{\bar{\theta}_1^2}, \quad \frac{\beta_{33}}{\bar{\theta}_1}.$$

Une fois que ces dernières sont fixées, tous les paramètres sont trouvés, en utilisant (6.27).

Si nécessaire, une possibilité pour lever la *sous-détermination*, est de minimiser une norme donnée (ou bien un semi-norme) de $\gamma \doteq \left(\frac{1}{\theta_1}, \frac{\beta_{11}}{\bar{\theta}_1^3}, \frac{\beta_{12}}{\bar{\theta}_1^3}, \frac{\beta_{22}}{\bar{\theta}_1^3}, \frac{\beta_{13}}{\bar{\theta}_1^2}, \frac{\beta_{23}}{\bar{\theta}_1^2}, \frac{\beta_{33}}{\bar{\theta}_1} \right)^\top$, soit $\gamma^\top M^\top M \gamma$, pour une matrice M (rectangulaire) donnée, sous les contraintes (6.29). En particulier, une possibilité serait de supposer que les termes non-diagonaux $\beta_{12}, \beta_{13}, \beta_{23}$ sont nuls, ce qui est raisonnable dans la mesure où les termes diagonaux sont généralement prédominants. En utilisant successivement (6.29b), (6.29a), (6.29c) et (6.29d), on obtient alors une solution

unique

$$\begin{aligned}\bar{\theta}_1 &= \frac{2y_1}{3x_1} \\ \beta_{11} &= \left(\frac{2y_1}{3x_1}\right)^3 x_0 - \left(\frac{2y_1}{3x_1}\right)^2 = \frac{8x_0y_1^3}{27x_1^3} - \frac{4y_1^2}{9x_1^2} \\ \beta_{22} &= \left(\frac{2y_1}{3x_1}\right)^3 x_2 - \left(\frac{2y_1^2}{9x_1}\right)^2 = \frac{8y_1^3x_2}{27x_1^3} - \frac{4y_1^4}{81x_1^2} \\ \beta_{33} &= -\frac{3}{y_1} \left(\frac{2y_1}{3x_1}\right) x_3 = -\frac{2x_3}{x_1}.\end{aligned}$$

Cela correspond au cas où $M = \begin{pmatrix} 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 \end{pmatrix}$. En général, le problème d'optimisation quadratique correspondant peut être résolu directement, comme suit.

Proposition 6.18. *Soit $p, q, r \in \mathbb{N}$ avec $p, q \leq r$, soit $A \in \mathbb{R}^{p \times r}$, $b \in \mathbb{R}^{p \times 1}$ et $M \in \mathbb{R}^{q \times r}$, avec A est de rang plein. Supposons que $\ker A \cap \ker M = \{0\}$. Alors*

$$\min \left\{ \frac{1}{2} \gamma^\top M^\top M \gamma : A\gamma = b \right\} = -\frac{1}{2} \begin{pmatrix} 0_{1 \times r} & b^\top \end{pmatrix} \begin{pmatrix} M^\top M & A^\top \\ A & 0_{p \times p} \end{pmatrix}^{-1} \begin{pmatrix} 0_{r \times 1} \\ b \end{pmatrix}. \quad (6.31)$$

6.3 Conclusion

Dans ce chapitre, nous avons étudié la modélisation réduite d'un réseau constitué d'un ensemble de sous-cellules approximativement identiques, couplées par des liaisons électriques série-parallèle. Le modèle réduit obtenu permet de fournir une meilleure approximation de la fonction d'impédance de l'ensemble du réseau, en introduisant des termes correctifs qui caractérisent la dispersion par rapport au comportement moyen. Nous avons étudié aussi le problème lié à l'identification des quantités décrivant le comportement moyen et la dispersion, pour les utiliser comme des données d'alerte pour la surveillance et le diagnostic.

Nous avons étudié ensuite le problème d'identification, plus en détail, dans le cas où les fonctions d'impédances individuelles sont des transferts du premier ordre.

Chapitre 7

Conclusions générales

La motivation initiale de cette thèse était l'utilisation de la spectroscopie d'impédance pour la surveillance et le diagnostic de la pile à combustible PEM en se basant sur des modèles physiques. Le point de départ de cette étude a consisté en la modélisation du comportement dynamique de différents phénomènes physiques et électrochimiques au sein de la pile. Dans un premier temps, nous avons détaillé le modèle multiéchelle d'A. Franco de la pile [37] et nous l'avons simplifié de façon à aboutir à un système d'équations aux dérivées partielles (EDPs) en une unique dimension spatiale. La semi-discrétisation en espace des EDPs issues du modèle de Franco nous a conduits à un modèle composé d'un système d'équations algébro-différentielles en 0D. Ce dernier, simple à implémenter dans un logiciel de calcul scientifique, permet de faire des simulations numériques relativement rapides avec une vingtaine de points de discrétisation.

En nous inspirant des travaux classiques sur l'analyse géométrique des réseaux de réactions électrochimiques [64, 1], nous avons obtenu un modèle d'évolution de réseaux de réactions électrochimiques compatible avec la thermodynamique. Nous avons présenté ensuite un modèle précis de double couche électrique avec la correction de Frumkin. Ce modèle donne lieu à un modèle de la pile entière écrit dans le cadre d'une réaction élémentaire à l'anode et à la cathode.

Nous avons ensuite analysé le comportement harmonique des modèles obtenus précédemment en utilisant la méthode de la balance harmonique. Nous avons calculé analytiquement l'impédance du modèle des réseaux de réactions électrochimiques, ainsi que le modèle de la pile. Les modèles précédents supposent que la pile est représentée par une cellule unique et homogène. Afin de permettre de décrire les hétérogénéités spatiales observées dans une cellule ou dans un stack, nous avons proposé finalement le modèle réduit d'un réseau de cellules représentées par leur impédance. Ce modèle approxime l'impédance globale du réseau par une "cellule moyenne", connectée à deux cellules "série" et "parallèle" représentatives d'écart par rapport à la moyenne. Nous avons ensuite étudié le problème lié à l'identification des quantités permettant de décrire le comportement moyen et la dispersion.

Plusieurs modèles ont été présentés dans cette étude, les relations entre ces différents modèles sont les suivantes :

Dans le Chapitre 3 nous avons réduit le modèle multi échelle d'A. Franco de la pile en un modèle 0D. Le modèle obtenu permet de tenir compte de tous les phénomènes physiques et électrochimiques du modèle de Franco et permet de réaliser des simulations numériques relativement rapides.

Dans la deuxième partie, nous avons proposé dans le Chapitre 4 un modèle générique d'évolution de réseau des réactions électrochimiques. Ce modèle qui s'écrit sous forme d'un

système dynamique est capable de fournir par rapport au modèle de Franco un modèle détaillé de la couche compacte. La couche diffuse est également incluse dans le modèle moyennant la correction de Frumkin. Cet aspect méthodologique de la modélisation rend le modèle obtenu modulable et permet de l'adapter à n'importe quel mécanisme réactionnel. Nous avons ensuite proposé un schéma électrique analogique du modèle obtenu.

Dans le Chapitre 5, nous avons présenté, moyennant la méthode de la balance harmonique, un modèle d'impédance issu du modèle obtenu dans le Chapitre 4. Ce modèle est ensuite enrichi en tenant compte de l'impédance de la couche de diffusion de gaz (GDL) pour laquelle nous fournissons le schéma électrique équivalent.

Enfin, la modélisation réduite d'un réseau de cellule que nous avons présenté dans le Chapitre 6 permet de décrire l'hétérogénéité dans une cellule ou dans un stack. Cette méthode a été testée sur un modèle simple d'impédance (du premier ordre). Au plus long terme, cette méthode sera appliquée sur le modèle de l'impédance obtenu dans le Chapitre 5 en vue de l'utiliser pour la surveillance et le diagnostic de la pile.

Parmi les perspectives qui peuvent être envisagées à la suite de cette étude, on peut citer les suivantes :

1/ Le modèle de réseau électrochimique développé est générique et il est utilisable pour n'importe quels mécanismes réactionnels. Il est donc possible de prendre en compte d'autres phénomènes électrochimiques qui peuvent avoir lieu sur les sites catalytiques de la pile par exemple l'empoisonnement par le CO ou par d'autre gaz. Par ailleurs, il est possible d'étendre le modèle obtenu afin de tenir compte de la gestion d'eau au cœur de la pile moyennant une variable d'état représentant la concentration eau en chaque point du segment joignant les deux couches compactes. En plus, le modèle obtenu est utilisable dans d'autres contextes : autre type de pile ou systèmes électrochimiques.

2/ L'expression de l'impédance de la pile obtenue analytiquement dépend de certains paramètres physiques. Il serait intéressant d'étudier l'identification de ces paramètres en utilisant des mesures expérimentales et d'étudier l'influence de ces paramètres sur la performance de la pile. Cela permettra ensuite de proposer des outils qui peuvent aider à la surveillance/diagnostic ainsi qu'à la conception de la pile.

3/ Une méthode de surveillance et de diagnostic à partir des mesures d'impédance a été proposée dans le Chapitre 6. Cette méthode a été testée sur un modèle simple d'impédance (du premier ordre) et elle a permis d'identifier un certain nombre de paramètres et de décrire l'hétérogénéité dans une cellule ou dans un stack. Il serait intéressant d'appliquer cette méthode au modèle d'impédance de la pile obtenue dans le Chapitre 5.

Troisième partie

Annexes

Annexe A

Couches de Stern et de diffusion et formules de Butler-Volmer

A.1 Introduction

Dans chaque électrode de la pile, des réactions d'oxydo-réduction ont lieu sur des sites catalytiques qui sont à la surface de contact entre un conducteur électronique, appelé ici la phase métallique, et un conducteur ionique, appelé électrolyte. Elles s'accompagnent du transfert de un ou plusieurs électrons vers la phase métallique, ce qui est à l'origine du courant électrique. Pour les piles, les réactions qui dominent sont $H_2 \rightarrow 2H^+ + 2e^-$ à l'anode (oxydation de H_2) et $O_2 + 4H^+ + 4e^- \rightarrow 2H_2O$ à la cathode (réduction de $O_2 + 4H^+$). Plus généralement, ces réactions d'oxydo-réduction sont de la forme :



L'étude cinétique a pour objet de déterminer les fréquences, disons v_O et v_R , des réactions d'oxydation $Red \rightarrow Ox + ne^-$ et de réduction $Red \leftarrow Ox + ne^-$. Pour cette étude, les réactions sont décomposées en séries d'étapes élémentaires au sens où un seul électron est transféré : $n = 1$. Une telle réaction implique une espèce réduite Red adsorbée sur les sites catalytiques ; une espèce oxydée Ox, mobile dans l'électrolyte ; un électron e^- , produit par la réaction d'oxydation ou consommé par celle de réduction ou objet d'un transfert entre le lieu de la réaction et la phase métallique.

Pour les piles PEM qui nous intéressent, un catalyseur Pt/C permet les interactions $H_2/Pt/C$ et $O_2/Pt/C$ et donc l'adsorption de H_2 à l'anode et O_2 à la cathode, comme on l'a vu au Chapitre 4. L'électrolyte est ici un système constitué d'une partie solide, le Nafion qui imprègne la membrane, et d'eau qui peut circuler dans le système. Le Nafion est un polymère qui se comporte comme un réseau de charges négatives fixes, des groupes sulfonates SO_3^- , avec lesquels les protons H^+ , mobiles dans l'eau, peuvent réagir pour former des groupes SO_3H pouvant se défaire sous l'effet du gradient de potentiel électrochimique de H^+ . Le comportement de ce système est alors assimilé à celui d'un conducteur de protons H^+ qui ne conduit pas les électrons. Il sera résumé ici par un simple coefficient de diffusion de protons et une concentration constante de charges négatives fixes.

A.2 Les couches de Stern et de diffusion

A.2.1 Fonction et géométrie de l'interface entre la phase métallique et l'électrolyte

On suppose que, au moins à petite échelle, la surface de la phase métallique est assez régulière pour être représentée par un plan sur lequel est adsorbé Red qui se trouve alors au

plus près de la phase métallique, dans un plan parallèle appelé plan de Helmholtz intérieur. Pendant les réactions chimiques, représentées par (A.1), Ox et les électrons sont supposés dans un autre plan, plus à l'intérieur de l'électrolyte, le plan extérieur de Helmholtz. On peut donc supposer que la réaction se passe entre ces deux plans, dans la couche de Stern, aussi appelée couche compacte, et d'épaisseur λ_S (voir la partie b) de la figure 1 tirée de [86].

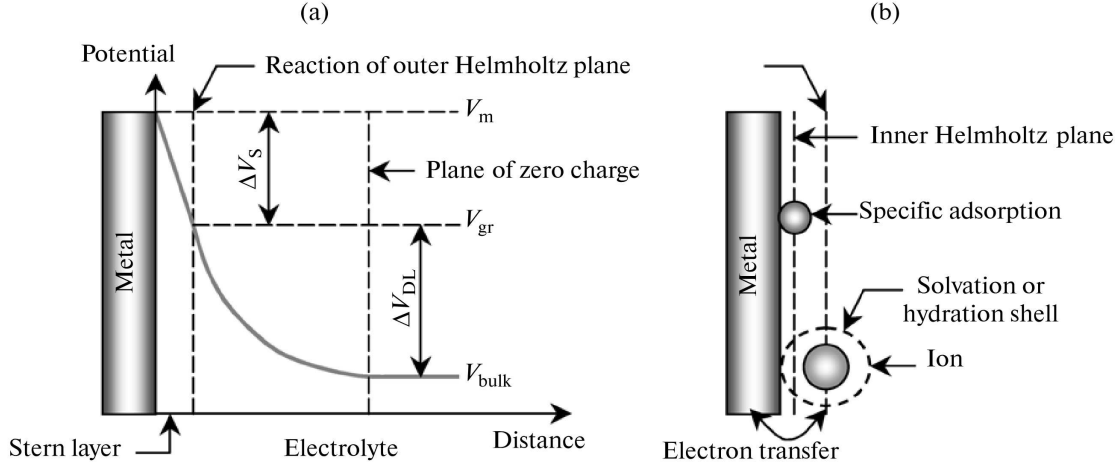


FIGURE A.1 – a) : Structure d'une double couche à l'interface entre une phase métallique et la membrane : couche compacte de Stern et couche de diffusion de Debye. b) : zoom sur la couche de Stern (illustration tirée de [86])

Les protons et les électrons sont produits ou consommés aux mêmes fréquences par les réactions d'oxydation ou de réduction. Cela correspond à des densités de courants protoniques et électroniques égales, disons à $J_O = nFv_O$ et $J_R = nFv_R$. Les protons étant mobiles dans l'électrolyte, vont pouvoir diffuser hors de la couche de Stern, dans la couche de diffusion jusqu'à se retrouver tous associés à un anion pour former une zone électroneutre dans la membrane. On note λ_D l'épaisseur de la couche diffuse et J_{F-} la densité de courant de protons quittant la couche de Stern, et J_{F+} celle se retrouvant dans la couche diffuse. On a $J_{F-} = J_O - J_R$ et J_{F+} sera discuté plus loin. Les électrons qui se trouvent dans la couche de Stern ont soit été produits par (A.1), soit résultent de transferts par effet tunnel avec la phase métallique, transferts, qui au total correspondent à une densité de courant électrique notée J qui est à un facteur près (surface active) le courant extérieur I , dit courant faradique, car résultant d'une réaction d'oxydo-réduction.

A une échelle plus large, cette géométrie plane se complique, mais on peut encore garder une vision simplifiée, même de la situation complexe de l'électrode poreuse, en imaginant que les plans de Helmholtz précédents deviennent les surfaces latérales de pores cylindriques de diamètre h_p par lesquels la partie liquide de l'électrolyte pénètre dans la partie solide.

Ce système double couche a essentiellement trois échelles. Si on note L la distance entre les électrodes, ce sont $L \gg h_p \gg \lambda_D$ (voir la figure A.2 tirée de [17]). Son comportement dépend de sa géométrie et en particulier des rapports $\varepsilon = \frac{\lambda_D}{L}$, de $\varepsilon_p = \frac{\lambda_D}{h_p}$ et de $\delta = \frac{\lambda_S}{\lambda_D}$. Dans tous les cas $\varepsilon \ll 1$. Lorsque $\varepsilon_p \ll 1$ (cas des macropores), on est dans la situation des électrolytes liquides considérée par les modèles de type Gouy-Chapman-Stern (GCS). Dans les autres cas on est en présence de micropores.

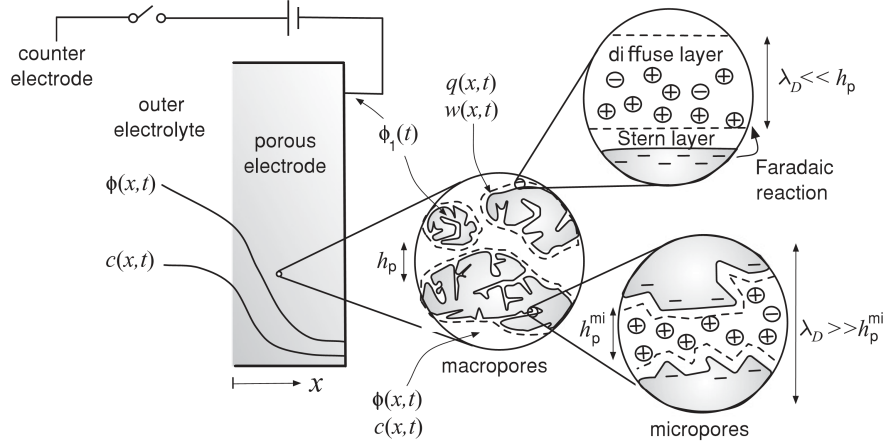


FIG. 1. Schematic of our porous electrode theory, which describes ion transport, diffuse charge, and Faradaic reactions across a hierarchy of three length scales: (i) left, the macroscopic continuum where the volume-averaged variables, such as the bulk concentration $c(x,t)$ and electrostatic potential $\phi(x,t)$, are defined; (ii) middle, the macropores in the pore-particle interphase with thin DLs (dashed lines) whose extent λ_D (the Debye screening length) is much smaller than the mean pore thickness h_p , and which are characterized by their mean charge density $q(x,t)$ and excess salt concentration $w(x,t)$ per area; and (iii) right, the nanoscale diffuse charge distribution, separated from the electron-conducting phase by dashed lines, by a molecular Stern layer (dashed lines) across which Faradaic electron-transfer reactions occur, and occurring either in thin DLs, in the macropores (upper right) or in charged micropores (lower right) of thickness $h_p^{mi} \ll \lambda_D$ with strongly overlapping DLs.

FIGURE A.2 – Les trois échelles (illustration tirée de [17]).

A.2.2 Le comportement électrique de la double couche

Notations. On commence par compléter les notations en suivant autant que possible celles de Bazant et al (voir par exemple [16]). On note, avec $i = a$ pour l'anode et $i = c$ pour la cathode :

$\varphi_{m,i}$: les potentiels des phases métalliques,

$\varphi_{rp,i}$: les potentiels des plans extérieurs de Helmholtz où on localise les réactions.

$\varphi_{M,i}$: les potentiels à l'interface entre une couche de diffusion et la zone où la membrane peut être supposée électroneutre.

On définit les différences de potentiel aux bornes des couches de Stern, de diffusion et de la membrane :

$$\Delta\varphi_{S,i} = \varphi_{m,i} - \varphi_{rp,i}, \quad \Delta\varphi_{D,i} = \varphi_{rp,i} - \varphi_{M,i}, \quad i = c, a; \quad \Delta\varphi_M = \varphi_{a,M} - \varphi_{c,M} \quad (\text{A.2})$$

Lorsqu'on n'aura pas besoin de préciser l'électrode concernée, on omettra l'indice $i = a, c$.

Finalement, la différence de potentiel aux bornes de la pile, U , est donnée par

$$U \doteq \varphi_{m,c} - \varphi_{m,a} = \Delta\varphi_{S,c} + \Delta\varphi_{D,c} - \Delta\varphi_M - \Delta\varphi_{D,a} - \Delta\varphi_{S,a} \quad (\text{A.3})$$

La pile est connectée au circuit extérieur par ses phases métalliques, le courant extérieur est I , le sens positif allant de la cathode vers l'anode, de sorte que pour une résistance de charge R_{ext} , on ait $U = R_{ext}I_M$, où I_M est le courant dans la membrane. En régime stabilisé $I = I_M$.

Avec ces conventions de signe, en fonctionnement normal, on a U , $-\Delta\varphi_{rp,a}$, $\Delta\varphi_{rp,c}$ strictement positifs et I , $-\Delta\varphi_{D,a}$, $\Delta\varphi_{D,c}$, $\Delta\varphi_M$, positifs ou nuls.

On note I_M le courant de protons traversant la membrane qui sera supposé suivre une simple loi d'Ohm avec une résistance R_M , on a $\Delta\varphi_M = R_M I_M$. I_M qui va de l'anode à la cathode sera positif. En régime stabilisé on aura $I = I_M$.

L'effet capacitif de la couche diffuse. Cette couche est un continuum dont l'intérieur contenant les charges et la frontière à travers laquelle passent des densités de courant sont bien définis. On peut appliquer la loi de conservation des charges. Notant J_M la densité de courant de protons quittant la couche diffuse pour se retrouver dans le reste de la membrane, on a :

$$\frac{\partial q_D}{\partial t} = J_F - J_M \quad (\text{A.4})$$

Pour obtenir une relation charge-tension afin de compléter le modèle de condensateur, il faut faire des hypothèses sur le comportement de cette couche.

Sous l'hypothèse du modèle de Gouy-Chapman, $\epsilon_p \ll 1$ et $\delta \ll 1$, on peut résoudre l'équation de Poisson-Boltzmann (voir les détails par exemple dans [2], page 547). La démarche est la suivante :

On suppose d'abord qu'une loi de Boltzmann régit la concentration de l'ion sur le plan de réaction, où on note c_i^0 sa concentration hors de la couche diffuse, en posant $\phi(x) = \varphi(x) - \varphi(0)$ ($x = 0$ correspondant à la sortie de la couche diffuse) :

$$c_i = c_{\infty,i} \exp\left(-\frac{z_i e \phi(x)}{k_B T}\right) \quad (\text{A.5})$$

où e est la charge élémentaire, k_B est la constante de Boltzmann, T est la température et z_i la valence de l'ion i .

Il reste alors à résoudre l'équation de Poisson qui s'écrit dans ce cas de plusieurs ions :

$$\frac{d^2 \phi}{dx^2} = -\frac{e}{\epsilon_D} \sum_i c_{\infty,i} z_i \exp\left(-\frac{z_i e \phi(x)}{k_B T}\right) \quad (\text{A.6})$$

avec $\epsilon_D = \epsilon \epsilon_0$, la permittivité diélectrique de la couche diffuse (supposée constante pour simplifier). La remarque de base est alors que $\frac{d^2 \phi}{dx^2} = \frac{1}{2} \frac{d}{d\phi} \left(\frac{d\phi}{dx}\right)^2$. On peut alors résoudre (A.6) qui devient une équation en le champ électrique considéré comme une fonction du potentiel $E(\phi)$ tel que $E(\phi(x)) = -\frac{d\phi}{dx}(x)$. On a

$$\frac{d}{d\phi} E(\phi)^2 = -\frac{2e}{\epsilon_D} \sum_i c_{\infty,i} z_i \exp\left(-\frac{z_i e \phi}{k_B T}\right)$$

En supposant que $E(0) = 0$ (i.e. $\phi(0) = \frac{d\phi}{dx}(0)$), cela s'intègre pour donner :

$$E(\phi)^2 = \frac{2k_B T}{\epsilon_D} \sum_i c_{inf,y,i} \left(\exp\left(-\frac{z_i e \phi}{k_B T}\right) - 1 \right) \quad (\text{A.7})$$

Le second membre est positif dès que $\phi \leq 0$.

On peut maintenant calculer la charge q_D de la couche diffuse en utilisant le théorème de Gauss qui la relie au flux de E à travers une surface Σ_D entourant la couche diffuse. Avec le modèle des pores cylindriques, on peut prendre $\Sigma_D = \Sigma_{rp} \cup \Sigma_{lat}$ constituée du cylindre Σ_{rp} qui remplace à cette échelle le plan de Helmholtz extérieur, et des surfaces latérales orthogonales à l'axe du cylindre Σ_{lat} , qui ferment Σ_D . Le théorème de Gauss s'écrit ainsi à l'aide de la permittivité diélectrique ϵ du milieu (non uniforme ici) : $q_D = \int_{\Sigma_D} \epsilon(s) E(s) \cdot dS$.

Du fait du comportement uniforme du pore suivant son axe, le champ électrique $E(s)$ est nul sur Σ_{lat} , d'où $q_D = \int_{\Sigma_{rp}} \epsilon(s) E(s) \cdot dS$.

Pour calculer cette intégrale, on a implicitement supposé que le déplacement électrique $\epsilon(s)E(s)$ était continu au voisinage des points Σ_{rp} , ce qui a permis de prendre sa trace sur cette frontière de la couche diffuse. Pour justifier cela, on fait donc l'hypothèse suivante :

Hypothèse A.1. Comportement du déplacement $\epsilon(s)E(s)$ au voisinage de la surface de réaction. On suppose que :

- i) Le déplacement électrique ϵE est continu à la traversée du plan de réaction.
- ii) A l'échelle intermédiaire, les couches de Stern et de diffusion constituent des pores cylindriques au comportement uniforme suivant leur axe.

Les intégrales sont constantes sur Σ_D et on peut donc intégrer à vue. Notant S_{rp} l'aire de Σ_{rp} , cette hypothèse permet de décrire $q_D = S_{rp}\epsilon_D E_D$. On peut calculer E_D avec (A.7) : $E_D = E(\Delta\varphi_D)$. On a donc :

$$q_D = S_{rp} \sqrt{2\epsilon_D k_B T \sum_i c_{\infty,i} \left(\exp\left(-\frac{z_i e \Delta\varphi_D}{k_B T}\right) - 1 \right)}, \quad \text{pour } \Delta\varphi_D \leq 0 \quad (\text{A.8})$$

On peut étendre cette formule à $\Delta\varphi_D$ de signe quelconque ainsi, en veillant à ce que la racine carrée soit définie en remplaçant $-\Delta\varphi_D$ par $|\Delta\varphi_D|$:

$$q_D = -\text{sgn}(\Delta\varphi_D) S_{rp} \sqrt{2\epsilon_D k_B T \sum_i c_{\infty,i} \left(\exp\left(\frac{z_i e |\Delta\varphi_D|}{k_B T}\right) - 1 \right)} \quad (\text{A.9})$$

Pour les faibles tensions, c'est à dire lorsque $|\Delta\varphi_D| \ll \frac{k_B T}{e}$, on a

$$\sum_i c_{\infty,i} \left(\exp\left(-\frac{z_i e \Delta\varphi_D}{k_B T}\right) - 1 \right) = \sum_i c_{\infty,i} \left(-\frac{z_i e \Delta\varphi_D}{k_B T} + \frac{1}{2} \left(\frac{z_i e \Delta\varphi_D}{k_B T} \right)^2 + \dots \right)$$

Mais $\sum_i c_{\infty,i} z_i = 0$ du fait de l'électroneutralité dans l'électrolyte, hors de la couche diffuse, d'où

$$\left(\frac{q_D}{S_{rp}} \right)^2 = 2\epsilon_D k_B T \frac{\sum_i c_{\infty,i} (z_i e)^2 + \dots}{2(k_B T)^2} \Delta\varphi_D^2 \approx \epsilon_D \frac{\sum_i c_{\infty,i} (z_i e)^2}{k_B T} \Delta\varphi_D^2$$

Il est alors naturel d'introduire la longueur de Debye qui est l'épaisseur λ_D de la couche diffuse :

$$\lambda_D = \sqrt{\frac{\epsilon_D k_B T}{\sum_i c_{\infty,i} (z_i e)^2}} \quad (\text{A.10})$$

Ce modèle linéarisé s'écrit ainsi comme le modèle d'un condensateur :

$$q_D = -C_{D0} \Delta\varphi_D, \quad \text{avec } C_{D0} = \epsilon_D \frac{S_{rp}}{\lambda_D} \quad (\text{A.11})$$

L'effet capacitif de la couche de Stern. La situation est différente ici du fait du caractère "compact" de cette couche qui par exemple dans le modèle de pore cylindrique, dans les conditions courantes $\delta \ll 1$ est une surface. Dans le modèle 1D (disons le long du rayon du pore), c'est un point. Il est donc ici difficile de donner directement du sens à la charge du point ou aux courants qui y entrent ou en sortent. Un passage à la limite de type couche mince est à faire.

L'hypothèse faite fréquemment pour résoudre ce problème ([10], [16]) est celle de la continuité du déplacement électrique ϵE . Nous faisons donc l'Hypothèse 1, ce qui assure que l'on a (A.9) et $q_D = S_{rp}\epsilon_S E_S$. On fait maintenant l'hypothèse supplémentaire suivante :

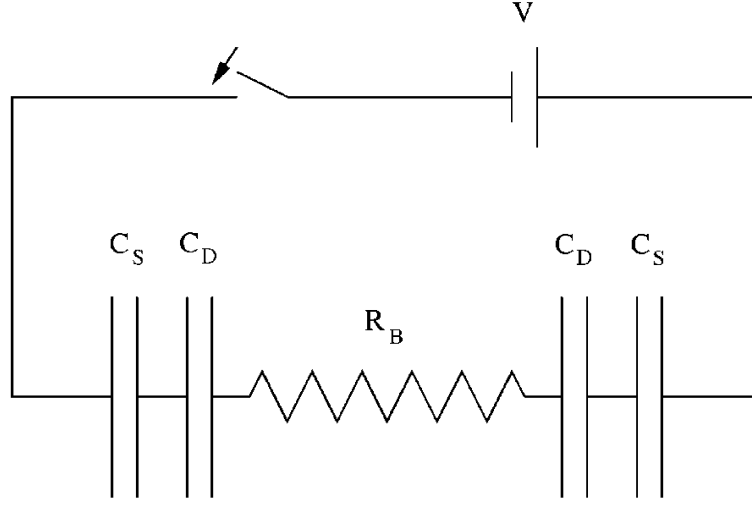


FIG. 5. Sketch of the equivalent RC circuit for the leading-order weakly nonlinear approximation: compact-layer and diffuse-layer capacitors in series with a bulk resistor. Although remarkably robust, the circuit approximation is violated by higher-order corrections, especially at large voltages.

FIGURE A.3 – Schéma de principe de la pile en non stationnaire (illustration tirée de [10]).

Hypothèse A.2. *Le champ électrique E est constant dans la couche de Stern.*

On a alors $E_S = -\frac{\Delta\varphi_S}{\lambda_S}$, d'où la :

Propriété A.3. *La charge dans la couche diffuse est proportionnelle à la différence de potentiel de la couche de Stern :*

$$q_D = -C_{S0}\Delta\varphi_S, \quad \text{avec } C_{S0} = \varepsilon_S \frac{S_{rp}}{\lambda_S} \quad (\text{A.12})$$

Remarques :

1/ Les formules (A.11) et (A.12) s'écrivent à l'échelle de l'électrode en remplaçant S_{rp} par $\gamma S_{EME} e_{CA}$ où S_{EME} est la surface géométrique de l'assemblage électrode-membrane, γ est la surface spécifique et e_{CA} l'épaisseur de la couche active.

2/ La formule (A.12) ressemble à la loi de comportement d'un condensateur du type (A.11), mais dans (A.12), si la différence de potentiel $\Delta\varphi_S$ est bien relative à la couche de Stern, la charge q_D est elle, celle de la couche diffuse : ce n'est donc pas un condensateur usuel. L'interprétation usuelle de cette situation en terme de circuit est que les deux couches constituent un système de deux condensateurs en série, cf par exemple [10] et la Figure A.3 qui en est extraite ou la discussion plus détaillée dans [9] qui montre les limites de ce point de vue dont une conséquence est de supposer que la couche de Stern reste électroneutre. En effet, avec ce montage en série, le courant faradique net se retrouve à l'extérieur (J), avant et après la traversée du plan de la réaction, et on a, à la place d'une équation de conservation comme (A.4), la relation :

$$J = J_{F-} = J_O - J_R = J_{F+} \quad (\text{A.13})$$

En pratique cette relation est fondamentale puisqu'elle permettra d'utiliser la relation de Butler-Volmer comme une fonctionnelle $J_F = \mathcal{R}(c_O, c_R, \Delta\varphi_S)$ définie indépendamment des autres éléments du système et utilisable simplement dans sa modélisation. Nous illustrerons cela dans la section suivante.

3/ On déduit de (A.11), (A.12) que pour les faibles potentiels, $q_D = -C_{D0}\Delta\varphi_D = -C_{S0}\Delta\varphi_S$, et par ailleurs, $\frac{C_{D0}}{C_{S0}} = \frac{\varepsilon_D}{\varepsilon_S} \frac{\lambda_S}{\lambda_D}$ qui est de l'ordre de grandeur de $\delta = \frac{\lambda_S}{\lambda_D}$, le facteur d'échelle géométrique entre les couches de Stern et diffuse. Il semble plus commode de prendre en compte dans ce facteur les déplacements électriques et de redéfinir δ par

$$\delta = \frac{C_{D0}}{C_{S0}} \quad (\text{A.14})$$

On a alors :

$$\Delta\varphi_S = \frac{\delta}{1+\delta}(\Delta\varphi_S + \Delta\varphi_D) \quad (\text{A.15})$$

Cas particulier d'un électrolyte ($z : z$). On considère donc le cas d'un anion et d'un cation de valence z et de concentration respectives $c_{\infty,-}$ et $c_{\infty,+}$ hors de la double couche. Les concentrations sur la couche de Stern deviennent :

$$c_+ = c_{\infty,+} \exp\left(-\frac{ze\Delta\varphi_D}{k_B T}\right), \quad c_- = c_{\infty,-} \exp\left(\alpha_D \frac{ze\Delta\varphi_D}{k_B T}\right) \quad (\text{A.16})$$

où on a introduit le facteur de symétrie pour pouvoir différentier les caractéristiques de la diffusion : $\alpha_D = 1$ correspond au cas d'un électrolyte symétrique et $\alpha_D = 0$ à celui pour lequel la concentration d'anion est constante. Cela revient à remplacer z_+ par z et z_- par $-\alpha_D z$ dans les calculs précédents qui sont donc toujours valables. On a maintenant, avec pour respecter l'électroneutralité, $c_{\infty,+} = c_{\infty,-} = c_\infty$:

$$q_D = -\text{sgn}(\Delta\varphi_D) S_{rp} \sqrt{2\varepsilon_D k_B c_\infty T \left(\exp\left(\frac{ze|\Delta\varphi_D|}{k_B T}\right) + \exp\left(\frac{-\alpha_D ze|\Delta\varphi_D|}{k_B T}\right) - 2 \right)} \quad (\text{A.17})$$

Dans le même temps, (A.10) devient :

$$\lambda_D = \sqrt{\frac{\varepsilon_D k_B T}{(1+\alpha_D)c_\infty (ze)^2}} \quad (\text{A.18})$$

Cas d'un électrolyte symétrique. On a $\alpha_D = 1$, d'où

$$\exp\left(\frac{ze|\Delta\varphi_D|}{k_B T}\right) + \exp\left(\frac{-\alpha_D ze|\Delta\varphi_D|}{k_B T}\right) - 2 = 2 \left(\cosh\left(\frac{ze|\Delta\varphi_D|}{k_B T}\right) - 1 \right) = 4 \sinh^2\left(\frac{ze|\Delta\varphi_D|}{2k_B T}\right)$$

et en utilisant l'impairité de la fonction $\sinh$:

$$q_D = -2S_{rp} \sqrt{2\varepsilon_D c_\infty T} \sinh\left(\frac{ze\Delta\varphi_D}{2k_B T}\right) = -\varepsilon_D S_{rp} \frac{2k_B T}{ze} \sqrt{\frac{2c_\infty (ze)^2}{\varepsilon_D k_B T}} \sinh\left(\frac{ze\Delta\varphi_D}{2k_B T}\right)$$

On a ici, en utilisant (A.18) :

$$q_D = -\varepsilon_D \frac{S_{rp}}{\lambda_D} \frac{2k_B T}{ze} \sinh\left(\frac{ze\Delta\varphi_D}{2k_B T}\right), \quad \text{avec } \lambda_D = \sqrt{\frac{\varepsilon_D k_B T}{2c_\infty (ze)^2}} \quad (\text{A.19})$$

C'est la formule classique de Chapman (voir [2], page 549) où c'est la charge par un unité de surface qui est calculée, ce qui fait que le facteur S_{rp} est absent . Voir également la formule [12] dans [16].

Cas d'un électrolyte asymétrique. Cette variante est mieux adaptée au cas des électrolytes solides, par exemple au sens où les anions sont fixes comme dans les piles PEM (voir [16]). On a $\alpha_D = 0$, d'où

$$\exp\left(\frac{ze|\Delta\varphi_D|}{k_B T}\right) + \exp\left(\frac{-\alpha_D ze|\Delta\varphi_D|}{k_B T}\right) - 2 = \exp\left(\frac{ze|\Delta\varphi_D|}{2k_B T}\right) - 1$$

et

$$q_D = -\text{sgn}(\Delta\varphi_D) S_{rp} \sqrt{2\varepsilon_D k_B c_\infty T \left(\exp\left(\frac{ze|\Delta\varphi_D|}{k_B T}\right) - 1 \right)}$$

Inversement, on a

$$q_D^2 = 2 \left(C_{D0} \frac{k_B T}{ze} \right)^2 \left(\exp\left(\frac{ze|\Delta\varphi_D|}{k_B T}\right) - 1 \right)$$

d'où

$$|\Delta\varphi_D| = \frac{k_B T}{ze} \ln \left(1 + \frac{1}{2} \left(\frac{ze q_D}{C_{D0} k_B T} \right)^2 \right)$$

combinant avec (A.12), cela conduit à ($\Delta\varphi_D$ et $\Delta\varphi_S$ ayant le même signe, d'après (A.17) et (A.12)) :

$$\Delta\varphi_D = \frac{k_B T}{ze} \ln \left(1 + \frac{1}{2} \left(\frac{ze \Delta\varphi_S}{\delta k_B T} \right)^2 \right) \text{sgn} \Delta\varphi_S$$

qui sera aussi utile sous la forme exponentielle :

$$\exp\left(\frac{ze \Delta\varphi_D}{k_B T} \text{sgn} \Delta\varphi_S\right) = 1 + \frac{1}{2} \left(\frac{ze \Delta\varphi_S}{\delta k_B T} \right)^2$$

c'est à dire :

$$\text{Pour } \Delta\varphi_S \geq 0, \exp\left(\frac{ze \Delta\varphi_D}{k_B T}\right) = 1 + \frac{1}{2} \left(\frac{ze \Delta\varphi_S}{\delta k_B T} \right)^2$$

$$\text{Pour } \Delta\varphi_S < 0, \exp\left(\frac{ze \Delta\varphi_D}{k_B T}\right) = \left(1 + \frac{1}{2} \left(\frac{ze \Delta\varphi_S}{\delta k_B T} \right)^2 \right)^{-1}.$$

Finalement on a pour le cas $\alpha_D = 0$

$$q_D = -\text{sgn}(\Delta\varphi_D) \varepsilon_D \frac{\sqrt{2} S_{rp} k_B T}{\lambda_D ze} \sqrt{\exp\left(\frac{ze|\Delta\varphi_D|}{k_B T}\right) - 1}, \quad \text{avec } \lambda_D = \sqrt{\frac{\varepsilon_D k_B T}{c_\infty (ze)^2}} \quad (\text{A.20})$$

formule à comparer avec la formule [13] de [16].

Propriété A.4. Pour un électrolyte ($z : z$) avec des anions fixes, la relation entre la charge q_D et la tension $\Delta\varphi_D$ dans la couche diffuse, est donnée par

$$q_D = -\text{sgn}(\Delta\varphi_D) \sqrt{2} C_{D0} \frac{RT}{zF} \sqrt{\exp\left(\frac{zF|\Delta\varphi_D|}{RT}\right) - 1} \quad (\text{A.21})$$

La relation entre $\Delta\varphi_D$ et $\Delta\varphi_S$ est :

$$\Delta\varphi_D = \frac{RT}{zF} \ln \left(1 + \frac{1}{2} \left(\frac{zF \Delta\varphi_S}{\delta RT} \right)^2 \right) \text{sgn} \Delta\varphi_S \quad (\text{A.22})$$

qui s'écrit encore :

$$\exp\left(-\frac{ze\Delta\varphi_D}{k_B T}\right) = k_D \left(\frac{ze\Delta\varphi_S}{\delta k_B T}\right), \text{ avec } k_D(\phi) = \begin{cases} 1 + \frac{1}{2}\phi^2 & \text{si } \phi \leq 0 \\ \left(1 + \frac{1}{2}\phi^2\right)^{-1} & \text{si } \phi > 0 \end{cases} \quad (\text{A.23})$$

$$\text{avec } C_{D0} = \varepsilon_D \frac{S_{rp}}{\lambda_D} \text{ et } \delta = \frac{C_{D0}}{C_{S0}}.$$

Remarque : dans le cas non symétrique que considère cette proposition, pour les faibles tensions ($\frac{zF|\Delta\varphi_D|}{RT} \ll 1$), (A.21) se comporte comme un condensateur non linéaire :

$$q_D \approx -\sqrt{2}C_{D0}\sqrt{\frac{RT}{zF|\Delta\varphi_D|}}\Delta\varphi_D \quad (\text{A.24})$$

à comparer par exemple avec le comportement linéaire du cas symétrique (anions et cations de même mobilité) de (A.11) : $q_D = -C_{D0}\Delta\varphi_D$.

L'expression (A.23) conduit à

$$\Delta\varphi_D := \Phi_D(\Delta\varphi_S) \approx \frac{zF}{2\delta^2 RT} \Delta\varphi_S |\Delta\varphi_S| \quad (\text{A.25})$$

Ici encore, le comportement est non linéaire (monotone croissant) et très dépendant de δ .

Pour δ grand, on a $k_D\left(\frac{\phi}{\delta}\right) \approx 1$: la couche de Stern domine dans le comportement total. On peut interpréter cela en disant que dans la série des deux condensateurs C_{S0} et C_{D0} , le plus petit des deux, C_{S0} impose son comportement.

A.2.3 Correction de Frumkin pour la formule de Butler-Volmer

On considère une formule de Butler-Volmer du type (A.45). On note $\alpha_O = z\beta_O$, $\alpha_R = z\beta_R$ et on suppose que $\beta_O + \beta_R = 1$:

$$\mathcal{R}(c_O, c_R, \Delta\varphi) = k_O c_R \exp\left(\frac{\alpha_O F \Delta\varphi}{RT}\right) - k_R c_O \exp\left(\frac{-\alpha_R F \Delta\varphi}{RT}\right) \quad (\text{A.26})$$

Cette formule était à l'origine appliquée à l'ensemble constitué par la couche de Stern et la couche diffuse, c'est à dire, avec les notations actuelles, que $\Delta\varphi = \Delta\varphi_S + \Delta\varphi_D$. La correction de Frumkin consiste à l'appliquer à la couche de Stern seule, donc à prendre $\Delta\varphi = \Delta\varphi_S$ puis à la corriger pour tenir compte de la couche diffuse.

Le signe du courant est choisi positif dans la direction $R \rightarrow O + e^-$, donc par exemple, il est positif à l'anode (les électrons sortent et le courant entre) où $\Delta\varphi_S$ est négatif, et il est négatif à la cathode où $\Delta\varphi_S$ est positif.

Correction de Frumkin de base. Elle consiste à définir

$$\tilde{\mathcal{R}}(c_R, \Delta\varphi_S) = \mathcal{R}(c_O, c_R, \Delta\varphi_S), \quad \text{avec } c_O = c_{\infty, O} \exp\left(-\frac{z_O e \Delta\varphi_D}{k_B T}\right) \quad (\text{A.27})$$

ce qui donne :

$$\tilde{\mathcal{R}}(c_R, \Delta\varphi_S) = k_O c_R \exp\left(\frac{\alpha_O F \Delta\varphi_S}{RT}\right) - k_R c_{\infty} \exp\left(-\frac{\alpha_R F \Delta\varphi_S + zF \Delta\varphi_D}{RT}\right) \quad (\text{A.28})$$

avec $\Delta\varphi_D = \Phi_D(\Delta\varphi_S)$ que l'on doit choisir suivant la nature de la couche diffuse.

Précisons les notations qui seront utilisées à l'anode et à la cathode.

On note $\Delta\varphi_{Di} = \Phi_D(\Delta\varphi_{Si})$ pour $i = a, c$.

- A l'anode de la pile :

$$\tilde{\mathcal{R}}_a(c_{Ra}, \Delta\varphi_{Sa}) = \mathcal{R}_a(c_O, c_{Ra}, \Delta\varphi_{Sa}), \quad \text{avec } c_O = c_\infty \exp\left(-\frac{zF\Delta\varphi_{Da}}{RT}\right)$$

c'est à dire :

$$\tilde{\mathcal{R}}_a(c_{Ra}, \Delta\varphi_{Sa}) = k_{Oa}c_{Ra} \exp\left(\frac{\alpha_{Oa}F\Delta\varphi_{Sa}}{RT}\right) - k_{Ra}c_\infty \exp\left(-\frac{\alpha_{Ra}F\Delta\varphi_{Sa} + zF\Delta\varphi_{Da}}{RT}\right) \quad (\text{A.29})$$

- A la cathode de la pile, de façon symétrique (on prend c_{Oc} comme entrée) :

$$\tilde{\mathcal{R}}_c(c_{Oc}, \Delta\varphi_{Sc}) = \mathcal{R}_c(c_{Oc}, c_R, \Delta\varphi_{Sc}), \quad \text{avec } c_R = c_\infty \exp\left(\frac{zF\Delta\varphi_{Dc}}{RT}\right)$$

c'est à dire,

$$\tilde{\mathcal{R}}_c(c_{Oc}, \Delta\varphi_{Sc}) = k_{Oc}c_\infty \exp\left(\frac{\alpha_{Oc}F\Delta\varphi_{Sc} + zF\Delta\varphi_{Dc}}{RT}\right) - k_{Rc}c_{Oc} \exp\left(-\frac{\alpha_{Rc}F\Delta\varphi_{Sa}}{RT}\right) \quad (\text{A.30})$$

Ces formules (A.29) et (A.30) avec correction de Frumkin sont identiques à la formule (9) de [16] (à la convention de signe des j près), et comme dans cet article il est facile de l'adapter pour tenir compte du taux d'occupation des sites actifs par Ra et Oc respectivement.

Correction de Frumkin pour l'électrolyte de type $(z : z)$ asymétrique. La relation $\Delta\varphi_D = \Phi_D(\Delta\varphi_S)$ est ici plus commode à utiliser sous la forme (A.23). On peut écrire, avec les notations correspondant à une anode (entrée c_R) :

$$\tilde{\mathcal{R}}(c_R, \Delta\varphi_S) = k_{OcR} \exp\left(\frac{\alpha_{Oc}F\Delta\varphi_S}{RT}\right) - k_Rk_D \left(\frac{ze\Delta\varphi_S}{\delta k_B T}\right) c_\infty \exp\left(-\frac{\alpha_R F \Delta\varphi_S}{RT}\right) \quad (\text{A.31})$$

avec

$$k_D(\phi) = \begin{cases} 1 + \frac{1}{2}\phi^2 & \text{si } \phi \leq 0 \\ \left(1 + \frac{1}{2}\phi^2\right)^{-1} & \text{si } \phi > 0 \end{cases}$$

Nous allons étudier quelques propriétés élémentaires de $\tilde{\mathcal{R}}(c_R, \Delta\varphi_S)$ utiles par la suite. Nous poserons $f = \frac{ze}{k_B T} = \frac{zF}{RT}$.

On a

$$\tilde{\mathcal{R}}(c_R, \Delta\varphi_S) = k_{OcR} \exp(\beta_O f \Delta\varphi_S) - k_Rk_D \left(\frac{f\Delta\varphi_S}{\delta}\right) c_\infty \exp(-\beta_R f \Delta\varphi_S) \quad (\text{A.32})$$

1/ *Comportement lorsque le courant est nul.*

Remarquons que la condition "courant nul" ne correspond pas nécessairement à un équilibre.

On a donc :

$$k_{OcR} \exp(\beta_O f \Delta\varphi_S) = k_Rk_D \left(\frac{f\Delta\varphi_S}{\delta}\right) c_\infty \exp(-\beta_R f \Delta\varphi_S)$$

d'où l'équation en $\Delta\varphi_{S0}$ à résoudre :

$$k_D \left(\frac{f\Delta\varphi_{S0}}{\delta}\right) \exp(-f\Delta\varphi_{S0}) = \frac{k_{OcR}}{k_R c_\infty}$$

En posant $\phi = \frac{f\Delta\varphi_{S0}}{\delta}$, il s'agit d'étudier l'équation

$$k_D(\phi) \exp(-\delta\phi) = \frac{k_{OCR}}{k_R c_\infty}$$

On vérifie facilement que la fonction $k_D(\phi) \exp(-\delta\phi)$ est continue, décroissante et va de $+\infty$ à 0 quand ϕ va de $-\infty$ à $+\infty$, en passant par 1 pour $\phi = 0$. Cette équation a donc une unique solution négative pour tout $\frac{k_{OCR}}{k_R c_\infty} > 1$. On définit alors la fonction L_δ de sorte que cette solution s'écrive $-\frac{1}{\delta}L_\delta\left(\frac{k_{OCR}}{k_R c_\infty}\right)$ en rappelant ainsi qu'elle dépend du paramètre δ . On a donc

$$\Delta\varphi_{S0} = -\frac{RT}{zF}L_\delta\left(\frac{k_{OCR}}{k_R c_\infty}\right) \quad (\text{A.33})$$

Remarque :

Pour δ très grand, la solution est proche de celle de $\exp(-\delta\phi) = \frac{k_{OCR}}{k_R c_\infty}$ et on retrouve l'usuelle formule de Nernst, $\Delta\varphi_{S0} = -\frac{RT}{zF} \ln\left(\frac{k_{OCR}}{k_R c_\infty}\right)$.

Pour δ très petit, la solution est proche de la solution négative de $k_D\left(\frac{f\Delta\varphi_{S0}}{\delta}\right) = \frac{k_{OCR}}{k_R c_\infty}$, c'est à dire de la solution négative de $1 + \frac{1}{2}\left(\frac{f\Delta\varphi_{S0}}{\delta}\right)^2 = \frac{k_{OCR}}{k_R c_\infty}$, c'est à dire que

$$\Delta\varphi_{S0} = -\frac{\delta RT}{zF} \sqrt{2\left(\frac{k_{OCR}}{k_R c_\infty} - 1\right)}$$

2/ *Comportement au voisinage du courant nul.*

$$\frac{\partial}{\partial\varphi}\tilde{\mathcal{R}}(c_R, \varphi) = k_{OCR}\beta_O f \exp(\beta_O f \varphi) + k_R c_\infty \left(k_D\left(\frac{f\varphi}{\delta}\right) \beta_R f - \frac{dk_D}{d\varphi}\left(\frac{f\varphi}{\delta}\right) \right) \exp(-\beta_R f \varphi)$$

soit encore,

$$\frac{\partial}{\partial\varphi}\tilde{\mathcal{R}}(c_R, \varphi) = k_R c_\infty \exp(\beta_O f \varphi) \left[\frac{k_{OCR}}{k_R c_\infty} \beta_O f + \left(k_D\left(\frac{f\varphi}{\delta}\right) \beta_R f - \frac{dk_D}{d\varphi}\left(\frac{f\varphi}{\delta}\right) \right) \exp(-f\varphi) \right]$$

et on voit facilement, d'après le raisonnement fait pour l'équilibre, que dans un voisinage de l'équilibre, $\frac{dk_D}{d\varphi} < 0$. On a donc $\frac{\partial}{\partial\varphi}\tilde{\mathcal{R}}(c_R, \varphi) > 0$. Il s'agit d'une propriété de stabilité : l'amplitude de la différence de potentiel diminue lorsque le courant augmente.

2/ *Comportement pour les petits courants.* On a montré un peu plus au point précédent : l'équation $\tilde{\mathcal{R}}(c_R, \Delta\varphi_S) = J_M$ qui a une solution pour $J_M = 0$, a encore une solution pour J_M suffisamment petit. Pour résumer :

Propriété A.5. *Pour un électrolyte ($z : z$) asymétrique, la prise en compte de la couche diffuse conduit à la modification suivante de la formule de Butler-Volmer :*

i) *Le courant s'écrit sous la forme*

$$\tilde{\mathcal{R}}(c_R, \Delta\varphi_S) = k_{OCR} \exp\left(\frac{\alpha_O F \Delta\varphi_S}{RT}\right) - k_R c_\infty \left(1 + \frac{1}{2} \left(\frac{F \Delta\varphi_S}{\delta RT} \right)^2 \right) \exp\left(-\frac{\alpha_R F \Delta\varphi_S}{RT}\right) \quad (\text{A.34})$$

ii) Lorsque le courant est nul, le potentiel $\Delta\varphi_S$ s'écrit :

$$\Delta\varphi_{S0} = -\frac{RT}{zF}L_\delta\left(\frac{k_{OC}c_R}{k_Rc_\infty}\right), \text{ où } L_\delta(\gamma) \text{ vérifie } \left(1 + \frac{1}{2}\left(\frac{L_\delta(\gamma)}{\delta}\right)^2\right) \exp(L_\delta(\gamma)) = \gamma \quad (\text{A.35})$$

en particulier,

$$L_\infty(\gamma) \approx \ln(\gamma), \text{ pour } \delta \gg 1, \quad L_0(\gamma) \approx \delta \sqrt{2\left(\frac{k_{OC}c_R}{k_Rc_\infty} - 1\right)}, \text{ pour } \delta \ll 1 \quad (\text{A.36})$$

iii) La relation courant - tension a la propriété de stabilité suivante :

$$\frac{\partial}{\partial\varphi}\tilde{\mathcal{R}}(c_R, \varphi) > 0 \text{ pour } \varphi \leq 0 \quad (\text{A.37})$$

iv) Pour J_M suffisamment petit, l'équation suivante a une solution unique :

$$\tilde{\mathcal{R}}(c_R, \Delta\varphi_S) = J_M \quad (\text{A.38})$$

A.3 Construction élémentaire des formules de Butler-Volmer

Nous nous intéressons à des réactions électrochimiques impliquant des transferts d'électrons à l'interface d'une phase métallique dans laquelle seuls les électrons sont mobiles, et d'une phase, l'électrolyte, dans laquelle les ions sont mobiles. Deux types de phénomènes physique peuvent contrôler la réaction :

- Le transfert, entre cette interface et la phase métallique, des électrons pris à un donneur lors d'une réaction d'oxydation ou acceptés par un ion lors d'une réaction de réduction ;
- La diffusion des ions dans l'électrolyte en présence d'un champ électrique.

Par ailleurs la réaction peut être couplée à d'autres réactions (électro-)chimiques ou d'adsorption, désorption.

Nous présentons dans cette sections les ingrédients de base des formules de Butler-Volmer qui représentent la relation courant-tension qui résume du point de vue électrique les phénomènes précédents.

Les modèles de base : les lois d'Arrhenius et de Tafel. Les modèles de cinétique utilisés dans ces situations combinent deux types de comportement décrits par ces lois :

- Pour une réaction chimique ordinaire, la loi d'Arrhenius décrit l'influence de la température sur la vitesse k_{chem} à l'aide d'une énergie d'activation, E_a et d'un facteur pré-exponentiel, A_{chem} :

$$k_{\text{chem}} = A_{\text{chem}} \exp\left(-\frac{E_a}{RT}\right) \quad (\text{A.39})$$

Pour une réaction élémentaire, $A \rightarrow B$, dans laquelle l'activité chimique de A est c_A , par exemple sa concentration, et son potentiel chimique, $\mu_A = \mu_A^0 + RT \ln c_A$, on a $E_a = -\mu_A$ et

$$k_{\text{chem}} = A_{\text{chem}} \exp\left(\frac{\mu_A^0}{RT}\right) c_A \quad (\text{A.40})$$

- Pour une réaction avec transfert d'électron unidirectionnel (oxydation ou réduction) de vitesse k_{elec} , la loi de Tafel décrit l'influence sur la densité de courant, $j = Fk_{\text{elec}}$, de la différence de potentiel $\Delta\varphi$ à l'interface à l'aide d'un coefficient α et d'un facteur pré-exponentiel $\frac{j_0}{F}$:

$$k_{\text{elec}} = \frac{j_0}{F} \exp\left(\frac{\alpha F \Delta\varphi}{RT}\right) \quad (\text{A.41})$$

Elle traduit directement les relations expérimentales observées qui sont de la forme

$$\Delta\varphi = a + b \ln j, \quad \text{avec } b = \frac{RT}{\alpha F}, \quad a = b_T \ln j_0 \quad (\text{A.42})$$

dans lesquelles α caractérise la pente de Tafel b_T .

Application à l'anode et à la cathode d'une pile : loi de Butler-Volmer. On précise les notations pour ces deux situations.

Pour les réactions d'oxydation du type $\text{Red} \rightarrow \text{Ox} + e^-$, qui dominant à l'anode, une augmentation de $\Delta\varphi$ va accélérer le transfert des électrons vers le métal (rappelons que $\Delta\varphi = \varphi_m - \varphi_{rp}$) et donc augmenter la fréquence de la réaction et le courant qui lui est proportionnel. Convenant que ce courant entrant dans l'électrode est positif, cela conduit à écrire (A.41) ainsi :

$$j = j_O \exp\left(\frac{\alpha_O F \Delta\varphi}{RT}\right) \quad (\text{A.43})$$

Pour les réactions de réduction, $\text{Red} \leftarrow \text{Ox} + e^-$, qui dominant à la cathode, une augmentation de $\Delta\varphi$ va freiner le transfert des électrons vers l'électrolyte et donc diminuer la fréquence de la réaction et le courant. Avec la convention précédente, ce courant sortant de l'électrode est compté négativement et (A.41) devient :

$$j = -j_R \exp\left(\frac{-\alpha_R F \Delta\varphi}{RT}\right) \quad (\text{A.44})$$

En pratique les facteurs j_O et j_R ne sont pas constants dès que les niveaux d'adsorption de Ox ou Red sur la surface active sont forts ou variables en temps. Il est alors nécessaire de prendre en compte les activités chimiques, par exemple sous une forme de type (A.40), la plus simple étant $j_O = k_O c_R$, $j_R = k_R c_O$ avec, suivant l'électrode, c_R qui peut être pris égal par exemple à θ_R la fraction de couverture de la surface active par Red et c_O égal à la la concentration de Ox dans l'électrolyte.

A l'anode comme à la cathode, les deux types de réaction on lieu simultanément et la densité de courant j résulte de la superposition des courants précédents :

$$j = k_O c_R \exp\left(\frac{\alpha_O F \Delta\varphi}{RT}\right) - k_R c_O \exp\left(\frac{-\alpha_R F \Delta\varphi}{RT}\right) \quad (\text{A.45})$$

ce qui est une relation de type Butler-Volmer.

En terme d'énergie, cela s'écrit, pour des facteurs pré-exponentiels k_O^0 , k_R^0 convenables :

$$j = k_O^0 \exp\left(\frac{\mu_R + \alpha_O F \Delta\varphi}{RT}\right) - k_R^0 \exp\left(\frac{\mu_O - \alpha_R F \Delta\varphi}{RT}\right) \quad (\text{A.46})$$

avec

$$\mu_k = \mu_k^0 + RT \ln c_k, \quad k = R, O \quad (\text{A.47})$$

Une variante classique de la loi de Butler-Volmer. Elle considère la situation plus générale du transfert d'électron bidirectionnel, les deux courants (oxydation et réduction) pouvant se compenser pour une valeur j_0 correspondant à l'équilibre $j = 0$ pour lequel $\Delta\varphi = \Delta\varphi_e$. La loi décrit l'influence du surpotentiel $\eta = \Delta\varphi - \Delta\varphi_e$ sur j à l'aide de j_0 et d'un coefficient de transfert de charge α lié à α_O et α_R . Par exemple, pour des raisons qui apparaîtront plus loin, on a souvent $\alpha_O + \alpha_R = 1$ et on peut alors prendre $\alpha = \alpha_R$. Une écriture courant est alors

$$j = j_0 \left[\exp\left(-\frac{\alpha F \eta}{RT}\right) - \exp\left(\frac{(1 - \alpha) F \eta}{RT}\right) \right] \quad (\text{A.48})$$

La densité de courant d'échange j_0 apparaît souvent sous la forme (A.39) :

$$j_0 = J_0 \exp\left(-\frac{E_a}{RT}\right) \quad (\text{A.49})$$

Des lois de Butler-Volmer de ce type sont identifiées avec succès dans de nombreuses situations expérimentales qui ne se réduisent pas à une réaction élémentaire. Dans deux situations extrêmes du champ électrique à l'interface, intéressantes en pratique, on peut simplifier (A.48) :

- Pour un champ faible, $|\eta| \ll \frac{RT}{F}$. On peut linéariser :

$$\eta \approx \frac{RT}{F j_0} j \quad (\text{A.50})$$

Si on note S_e la surface de l'électrode et $I = S_e j$ le courant, pour cette situation de faible courant, il s'agit d'une relation de type loi d'Ohm, qui fait apparaître la résistance faradique :

$$\eta \approx R_F I, \quad \text{avec } R_F = \frac{RT}{F S_e j_0} \quad (\text{A.51})$$

- Pour un champ fort, $|\eta| \gg \frac{RT}{F}$, il ne reste qu'une exponentielle, par exemple pour $\eta < 0$:

$$\eta \approx -\frac{RT}{\alpha F} \ln\left(\frac{j}{j_0}\right) \quad (\text{A.52})$$

C'est une loi de Tafel analogue à (A.41) qui apparaît alors comme une approximation lorsqu'un sens de transfert domine, conduisant à un courant plus fort. L'effet de α devient visible.

A.4 Formule de Butler-Volmer et théorie de Marcus

Pour modéliser la cinétique de réactions contrôlées par le transfert d'électrons, on peut utiliser le modèle d'énergie d'activation donné par la théorie de Marcus [60] qui lui a valu le Prix Nobel en 1992 et dont nous ne rappelons que les constructions de base qui suffiront pour obtenir une paramétrisation simple. Même si dans le cas qui nous intéresse, la chimie est contrôlée par le transfert des protons, plus lent que celui des électrons, cette théorie conduit à une formule de Butler-Volmer utilisable plus généralement en choisissant de façon convenable ses paramètres, en particulier les facteurs pré-exponentiels en fonction des densités d'espèces. On s'inspire des présentations de [19, 2] de la théorie de Marcus de base qui considère des réactions élémentaires dans lesquelles la transformation d'un réactif s'accompagne du transfert d'un seul électron. Elles sont donc de la forme



A.4.1 Les bases du modèle de Marcus.

Modélisation énergétique. Du point de vue mathématique, la modélisation peut être résumée ainsi : on introduit d'abord un espace métrique dans lequel le vecteur q des coordonnées d'un point est constitué des coordonnées de tous les atomes intervenant dans la réaction. On peut voir q comme une variable de déformation permettant de passer d'un point à un autre. La fréquence des réactions $X_R \rightarrow X_P$ et $X_R \leftarrow X_P$ entre un réactif X_R et un

produit X_P de coordonnées respectives q_R et q_P dépend alors de la géométrie des chemins reliant q_R et q_P définie par des fonctions d'énergie que l'on va décrire maintenant.

Les complexes X_R et X_P ont une propriété de stabilité au sens où leurs énergies de Gibbs de formation, ΔG^R , ΔG^P sont des minima locaux de fonctions d'énergie potentielle $V_R(q)$ et $V_P(q)$ définies dans l'espace des configurations q , ces minima étant atteints en q_R et q_P respectivement :

$$\Delta G^i := V_i(q_i) = \min_q V_i(q), \quad \text{pour } i = R, P. \quad (\text{A.54})$$

L'énergie de Gibbs de la réaction est donc donnée par

$$\Delta G^0 := \Delta G^P - \Delta G^R = V_P(q_P) - V_R(q_R) \quad (\text{A.55})$$

Pendant la réaction, la coordonnée q passe de la vallée de q_R à celle de q_P en franchissant éventuellement des cols dont les hauteurs seront des barrières de potentiel. La fréquence de la réaction sera limitée par la barrière de potentiel constituée par le col le plus haut, disons de hauteur $\Delta G^\ddagger$, rencontré à la position $q^\ddagger$ sur un chemin critique passant par les cols les moins hauts. Les caractéristiques $q^\ddagger$, $\Delta G^\ddagger$ étant celles d'un col sont solution de l'équation d'intersection de V_R et de V_P :

$$\Delta G^\ddagger := V_R(q^\ddagger) = V_P(q^\ddagger) \quad (\text{A.56})$$

La résolution de cette équation permet alors calculer les énergies d'activation $\Delta G^{R\ddagger}$ et $\Delta G^{P\ddagger}$ (le "double dag" $\ddagger$ désigne l'état activé) de X_R et X_P qui interviennent dans les fréquences k_f et k_b des changements d'état correspondant aux réactions $X_R \rightarrow X_P$ et $X_R \leftarrow X_P$:

$$\Delta G^{i\ddagger} = \Delta G^\ddagger - \Delta G^i, \quad \text{pour } i = R, P. \quad (\text{A.57})$$

Les fréquences de réaction. Les fréquences de transition sont calculées à partir de la fréquence élémentaire $\kappa = \frac{k_B T}{h}$ (k_B et h sont les constantes de Boltzmann et de Planck) :

$$k_f = \kappa \exp\left(-\frac{\Delta G^{R\ddagger}}{RT}\right), \quad k_b = \kappa \exp\left(-\frac{\Delta G^{P\ddagger}}{RT}\right) \quad (\text{A.58})$$

ce qui conduit à la fréquence de la réaction :

$$k = \kappa \exp\left(-\frac{\Delta G^{R\ddagger}}{RT}\right) - \kappa \exp\left(-\frac{\Delta G^{P\ddagger}}{RT}\right) \quad (\text{A.59})$$

Un calcul plus détaillé des énergies d'activation va permettre de préciser cette fréquence.

A.4.2 Calcul des énergies d'activation et du facteur de symétrie avec la théorie de Marcus

Approximations pour des calculs effectifs. Pour mener à bien ces calculs, il faut éviter les très grandes dimensions des espaces de configuration réels. Pour cela, on peut remplacer la coordonnée q initiale par une coordonnée locale définie le long d'un chemin critique déterminé en faisant des hypothèses simplificatrices sur les fonctions V_R et V_P choisies de façon à passer par les divers complexes participants à la réaction (réactifs, possibles réactifs activés, produits ...). En pratique, avec des approximations paraboliques de V_R et V_P au voisinage des minima, cette démarche conduit à la formule de Marcus (A.66) que l'on va rappeler et qui permet de représenter de nombreux résultats expérimentaux.

Energie d'activation. Dans la théorie de Marcus, q est scalaire et V_R et V_P sont de simples paraboles dépendant d'un paramètre de "raideur" k , et des positions q_R et q_P , ce qui donne :

$$V_R(q) = \frac{1}{2}k(q - q_R)^2 + \Delta G^R, \quad V_P(q) = \frac{1}{2}k(q - q_P)^2 + \Delta G^P \quad (\text{A.60})$$

On peut alors calculer $q^\ddagger$, solution de $V_R(q^\ddagger) = V_P(q^\ddagger)$, soit encore de

$$\frac{1}{2}k(q^\ddagger - q_R)^2 = \Delta G^0 + \frac{1}{2}k(q^\ddagger - q_P)^2 \quad (\text{A.61})$$

d'où, $k(q_P - q_R) \left(q^\ddagger - \frac{q_R + q_P}{2} \right) = \Delta G^0$, soit enfin

$$q^\ddagger = \frac{q_R + q_P}{2} + \frac{\Delta G^0}{k(q_P - q_R)} \quad (\text{A.62})$$

On peut calculer maintenant les énergies d'activation. On a $\Delta G^\ddagger = V_R(q^\ddagger) = V_P(q^\ddagger)$, soit encore :

$$\Delta G^{R\ddagger} := \Delta G^\ddagger - \Delta G^R = \frac{1}{2}k \left(\frac{q_P - q_R}{2} + \frac{\Delta G^0}{k(q_P - q_R)} \right)^2 \quad (\text{A.63})$$

$$\Delta G^{P\ddagger} := \Delta G^\ddagger - \Delta G^P = \frac{1}{2}k \left(\frac{q_R - q_P}{2} + \frac{\Delta G^0}{k(q_P - q_R)} \right)^2 \quad (\text{A.64})$$

On introduit la quantité suivante, interprétée par Marcus comme l'énergie de réorganisation de l'électrolyte avant et après le transfert d'électron. C'est cette réorganisation qui limite la vitesse de la réaction, le transfert d'électron lui-même étant très rapide :

$$\lambda = \frac{1}{2}k(q_P - q_R)^2 \quad (\text{A.65})$$

Il vient :

$$\Delta G^{R\ddagger} = \frac{\lambda}{4} \left(1 + \frac{\Delta G^0}{\lambda} \right)^2, \quad \Delta G^{P\ddagger} = \frac{\lambda}{4} \left(1 - \frac{\Delta G^0}{\lambda} \right)^2 \quad (\text{A.66})$$

Il apparaît ainsi que les énergies d'activation $\Delta G^{R\ddagger}$, $\Delta G^{P\ddagger}$, ne dépendent que de deux paramètres, ΔG^0 et λ .

Le facteur de symétrie β . Il mesure la dissymétrie de la barrière de potentiel associée à la réaction. On peut le définir ainsi

$$\beta := \frac{\partial \Delta G^{R\ddagger}}{\partial \Delta G^0} = \frac{1}{2} \left(1 + \frac{\Delta G^0}{\lambda} \right) \quad (\text{A.67})$$

On a en particulier pour l'autre pente :

$$-\frac{\partial \Delta G^{P\ddagger}}{\partial \Delta G^0} = \frac{1}{2} \left(1 - \frac{\Delta G^0}{\lambda} \right) = 1 - \beta \quad (\text{A.68})$$

Remarquons que $q^\ddagger - q_R = \frac{q_P - q_R}{2} + \frac{\Delta G^0}{k(q_P - q_R)} = \frac{q_P - q_R}{2} \left(1 + \frac{\Delta G^0}{k(q_P - q_R)^2} \right)$, d'où :

$$\beta = \frac{q^\ddagger - q_R}{q_P - q_R} \quad (\text{A.69})$$

On a aussi

$$\Delta G^{R\dagger} = \lambda\beta^2, \quad \Delta G^{P\dagger} = \lambda(1 - \beta)^2 \quad (\text{A.70})$$

Avec ce modèle la réaction est possible dès que $q_R \leq q^\ddagger \leq q_P$, c'est à dire $0 \leq \beta \leq 1$, ou encore $\left| \frac{\Delta G^0}{\lambda} \right| \leq 1$. Dans les situations usuelles $\frac{\Delta G^0}{\lambda}$ est petit devant 1 et β est proche de 0.5 qui correspond à la situation symétrique d'une barrière de potentiel au milieu de la zone de la réaction, $q^\ddagger = \frac{q_R + q_P}{2}$, les pentes étant égales d'un côté et de l'autre.

La figure A.4 illustre ce paysage énergétique avec quelques situations courantes.

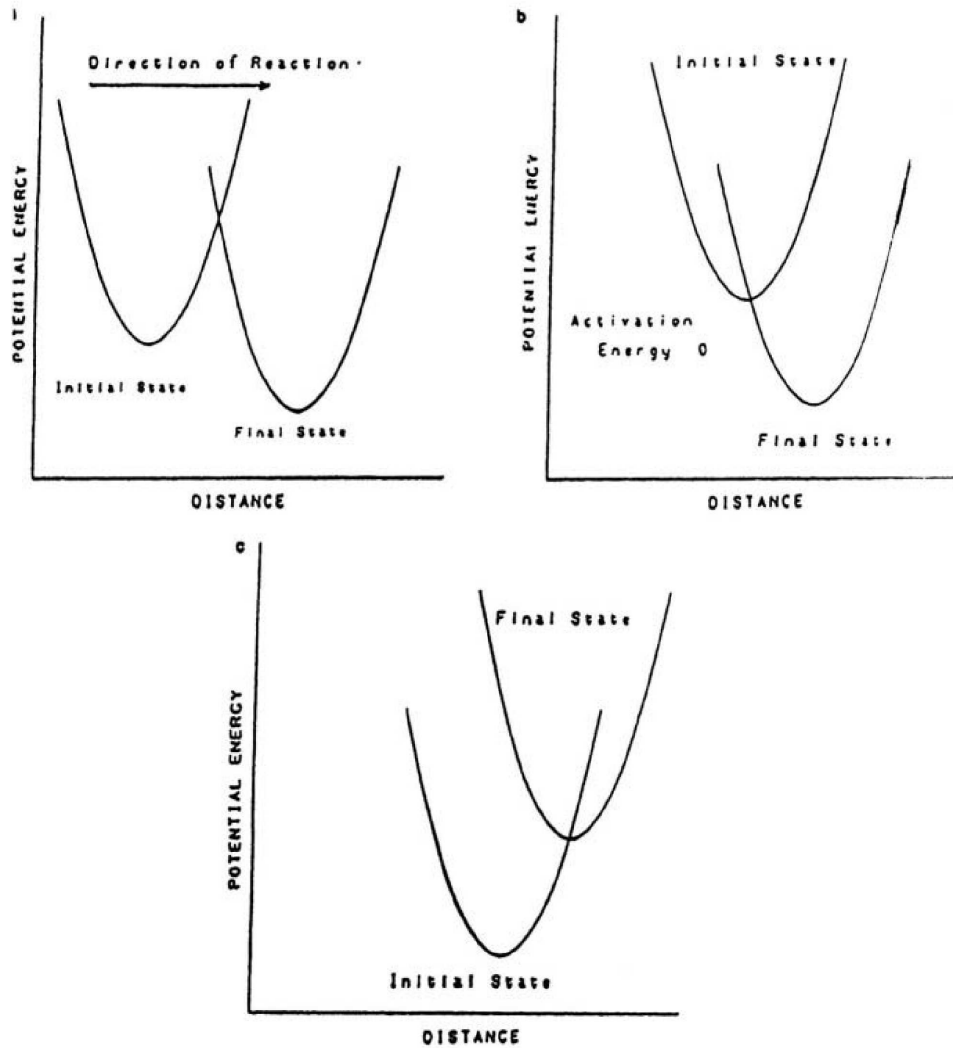


FIGURE A.4 – Trois exemples d'électrodes d'après Bockris [19] : a) Le cas typique ; b) Cas sans énergie d'activation ; c) Cas sans barrière d'activation

A.4.3 Formule de Butler-Volmer et coefficient de transfert de charge α

On note $\varphi(q_R) = \varphi_m$ le potentiel électrique dans le métal de l'électrode au contact de l'électrolyte, à la position q_R et $\varphi(q_P) = \varphi_m - \Delta\varphi$ celui dans l'électrolyte juste à l'extérieur du plan de Helmholtz qu'on situe en q_P . On suit ici les conventions de signe de l'introduction :

$\Delta\varphi$ représente suivant le cas $\Delta\varphi_{a,rp}$ ou $\Delta\varphi_{c,rp}$. Dans les représentations graphiques, q_R sera situé à gauche de q_P). Dans le calcul des énergies d'activation on doit tenir compte de l'énergie électrique nécessaire pour que les porteurs de charges franchissent la barrière de potentiel. Dans (A.66) c'est donc l'énergie de Gibbs électrochimique $\Delta\tilde{G}^0$ qui doit être utilisée. Pour la calculer on suppose que $\varphi(q)$ varie linéairement dans le domaine de la réaction, entre q_R et q_P (voir par exemple [19], Section 6.6). On a alors, d'après (A.69) :

$$\varphi(q^\ddagger) - \varphi(q_R) = -\beta\Delta\varphi, \quad \varphi(q_P) - \varphi(q^\ddagger) = -(1-\beta)\Delta\varphi \quad (\text{A.71})$$

Cas d'une réaction élémentaire. Considérons une réaction d'oxydation, $A \rightarrow A^+ + e^-$, correspondant, dans notre représentation de l'anode, à q allant de q_R , situé la gauche, côté métal, vers q_P , situé à droite, côté électrolyte. Lorsque $-\Delta\varphi > 0$, entre q_R et $q^\ddagger$, le champ électrique attire les électrons vers l'électrolyte, ce qui les freine pour rejoindre le métal et dans le même temps, il repousse les protons, ce qui les freine pour rejoindre l'électrolyte au delà de $q^\ddagger$: la barrière de potentiel est donc augmentée de $-\beta\Delta\varphi$ par cet effet électrique. Pour la réaction élémentaire de réduction, $A \leftarrow A^+ + e^-$, le potentiel croît de $-(1-\beta)\Delta\varphi$ entre $q^\ddagger$ et q_P , ce qui a encore comme effet de retenir protons et électrons dans la zone de q_P , favorisant cette fois la réaction : la barrière de potentiel est abaissée. On a donc à l'échelle des moles (la charge élémentaire étant remplacée par F) :

$$\Delta\tilde{G}^{R\ddagger} = \Delta G^{R\ddagger} - F\beta\Delta\varphi, \quad \Delta\tilde{G}^{P\ddagger} = \Delta G^{P\ddagger} + F(1-\beta)\Delta\varphi \quad (\text{A.72})$$

Finalement, pour une réaction élémentaire, la densité de courant électrique s'écrit sur le modèle de (A.59) :

$$j = F\kappa \left[\exp\left(-\frac{\Delta G^{R\ddagger} - F\beta\Delta\varphi}{RT}\right) - \exp\left(-\frac{\Delta G^{P\ddagger} + F(1-\beta)\Delta\varphi}{RT}\right) \right] \quad (\text{A.73})$$

A l'équilibre, on a $j = 0$, d'où un courant d'échange j_0 qui vérifie

$$j_0 = F\kappa \exp\left(-\frac{\Delta G^{R\ddagger} - F\beta\Delta\varphi_e}{RT}\right) = F\kappa \exp\left(-\frac{\Delta G^{P\ddagger} + F(1-\beta)\Delta\varphi_e}{RT}\right)$$

d'où

$$F\Delta\varphi_e = \Delta G^{R\ddagger} - \Delta G^{P\ddagger} = \Delta G^0$$

et finalement, $\Delta G^{R\ddagger} - F\beta\Delta\varphi_e = \Delta G^\ddagger - \Delta G^R - \beta(\Delta G^P - \Delta G^R)$, d'où

$$j_0 = F\kappa \exp\left(-\frac{\Delta G^\ddagger - (1-\beta)\Delta G^R - \beta\Delta G^P}{RT}\right) \quad (\text{A.74})$$

$$\Delta\varphi_e = \frac{\Delta G^P - \Delta G^R}{F} \quad (\text{A.75})$$

En introduisant le surpotentiel

$$\eta = \Delta\varphi - \Delta\varphi_e \quad (\text{A.76})$$

on obtient la relation de Butler-Volmer usuelle :

$$j = j_0 \left[\exp\left(\frac{\beta F \eta}{RT}\right) - \exp\left(-\frac{(1-\beta)F\eta}{RT}\right) \right] \quad (\text{A.77})$$

La relation courant-tension (A.77) dépend des paramètres spécifiques, j_0 et β . Avec la densité de courant d'échange à l'équilibre j_0 donnée par (A.74), elle est bien de la forme (A.48), (A.49).

Le facteur de symétrie β représente ici le coefficient de transfert de charge souvent noté α qui, suivant les conventions sur l'électrode de référence, est représenté par β ou $1-\beta$ (dans la présentation précédente nous avons fait le choix d'orientation correspondant à une anode avec le potentiel électrique de référence pris dans le métal de l'électrode).

A.4.4 Cas d'un système de réactions

Une adaptation de la formule (A.77) au cas de quelques systèmes de réactions élémentaires a été proposée par plusieurs auteurs, dont Bockris, voir [19], Section 7.6.

Le système considéré par cet Auteur est la décomposition en étapes élémentaires d'une réaction correspondant au transfert de n électrons :



La décomposition a la structure suivante :



Résistance faradique et coefficient de transfert de charge. Supposons qu'au voisinage d'un équilibre en temps la relation courant-surtension de chaque réaction soit de la forme

$$j_i = \sigma_i(\eta_i), \quad i = 1, n \quad (\text{A.80})$$

On cherche des informations, dans les mêmes conditions, sur la relation courant-surtension globale, c'est à dire sur la loi de type Butler-Volmer du système de réaction. La traduction électrique de la mise en série des étapes chimiques comme en (A.79) est la suivante :

i) Les courants produits par chacune des réactions sont égaux :

$$j_1 = \dots = j_n \quad (\text{A.81})$$

ii) La surtension totale η est la somme des surtensions de chaque réaction :

$$\eta = \eta_1 + \dots + \eta_n \quad (\text{A.82})$$

iii) Le courant j à l'interface est la somme des courants :

$$j = j_1 + \dots + j_n \quad (\text{A.83})$$

On voit que si les deux premières relations correspondent à la mise en série des dipôles électriques, il n'en est pas de même de la troisième qui n'est pas $j = j_1$ mais $j = nj_1$. Tout se passe comme si les transferts d'électrons avec la phase métal étaient en parallèle, d'où (A.83), mais synchronisés par les transitions $A_i \rightarrow A_{i+1}$, d'où (A.81).

Dans le cas d'un champ faible (courant faible), on peut utiliser l'approximation (A.51) de (A.77). On a pour chaque réaction, en supposant que la surface S_e est commune, $\eta_i \approx R_{Fi} S_e j_i$

avec $R_F = \frac{RT}{FS_e j_{0i}}$, pour $i = 1, \dots, n$.

On a alors la résistance faradique équivalente, en notant $I = S_e j$, $\eta \approx R_F I$ avec

$$R_F = \frac{\sum_1^n R_{Fi}}{n} = \frac{RT}{FS_e} \frac{1}{n} \sum_1^n \frac{1}{j_{0i}} \quad (\text{A.84})$$

Dans le cas d'un champ fort (courant fort), on peut utiliser l'approximation (A.52).

On a $\eta_i \approx -\frac{RT}{\beta_i F} \ln \left(\frac{j}{nj_{0i}} \right) = \frac{RT}{\beta_i F} (\ln(nj_{0i}) - \ln j)$, $i = 1, \dots, n$. D'où

$$\eta \approx \frac{RT}{F} \left(\sum_1^n \frac{1}{\beta_i} \ln(nj_{0i}) - \sum_1^n \frac{1}{\beta_i} \ln j \right)$$

D'où la formule de Tafel dans laquelle α est le coefficient de transfert de charge équivalent :

$$\eta \approx -\frac{RT}{\alpha F} \ln \left(\frac{j}{nj_0} \right), \quad \text{avec} \quad \frac{1}{\alpha} = \sum_1^n \frac{1}{\beta_i}, \quad \ln j_0 = \sum_1^n \frac{\alpha}{\beta_i} \ln j_{0i} \quad (\text{A.85})$$

Cas d'une étape limitante. Si on suppose maintenant que l'étape s est lente comparée aux autres, la conductivité pour cette étape est beaucoup plus petite que celle de toute autre étape, à savoir,

$$\sigma_s \ll \sigma_k, \quad k \neq s$$

on a alors dans ce cas :

$$\sum_{k=1}^n \frac{1}{\sigma_k} = \frac{1}{\sigma_1} + \frac{1}{\sigma_2} + \dots + \frac{1}{\sigma_n} \approx \frac{1}{\sigma_s} \quad (\text{A.86})$$

En utilisant (A.80) et (A.81), on obtient :

$$\eta_s \gg \eta_k, \quad k \neq s$$

Les surtensions des étapes ($k \neq s$) deviennent tous insignifiants par rapport aux η_s , d'où d'après (A.82) :

$$\eta \approx \eta_s. \quad (\text{A.87})$$

Par conséquent, l'étape s contrôle le courant total. En d'autres termes, à l'état stationnaire, les courants $j_k, k \neq s$ sont égaux au courant de l'étape déterminante j_s .

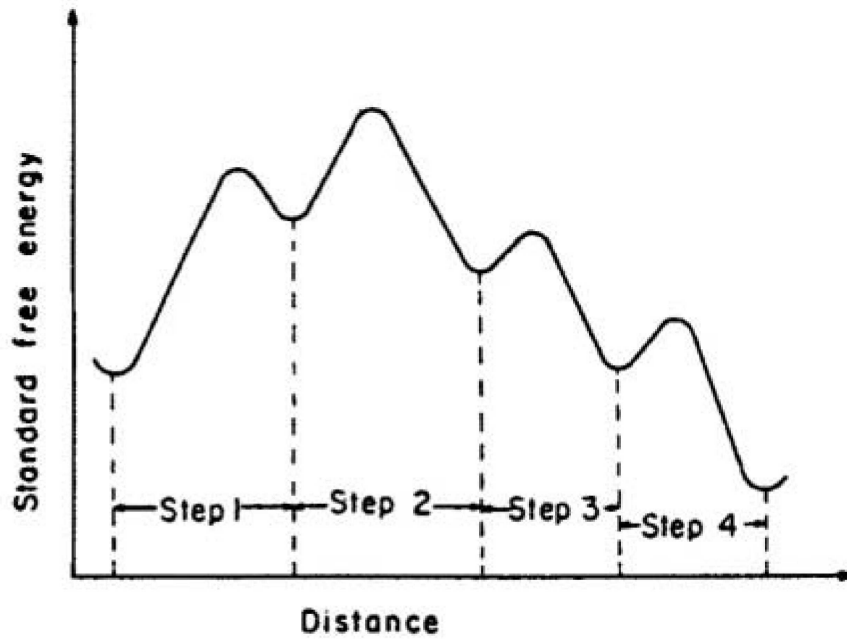


FIGURE A.5 – Barrière de potentiel d'une réaction à plusieurs étapes d'après Bockris [19]

Annexe B

Calcul de la surtension dans la couche compacte du modèle de Franco

L'expression de la surtension η introduite dans le Chapitre 3 est très largement inspiré des travaux d'A. Franco [37]. Dans cette annexe nous résumons la description et les calculs faites dans [37].

Une différence de potentiel électrique interfaciale entre l'électrolyte et le métal permet le développement de la réaction Redox sur la surface métallique avec une densité de charge électronique surfacique. Cette différence de potentiel interfaciale, est nommée par certains auteurs "surtension de Frumkin" [18], par d'autres "correction de Frumkin" [33]. D'après [37], la surtension dans la couche compacte η résulte de la superposition de deux phénomènes : l'effet capacitif dû à la présence de charges positives et négatives en regard ; et un phénomène dû à la nature dipolaire des molécules d'eau. Nous écrivons :

$$\eta(t) = \delta\phi(t) + \delta\phi_w(t) \quad (\text{B.1})$$

où $\delta\phi$ est la différence de potentiel de la couche compacte dû à la présence des charges de signes opposés. Elle s'écrit sous la forme :

$$\delta\phi(t) = -\frac{d}{\varepsilon_{cc}}\sigma(t)$$

où σ est la densité de charges électroniques dans la couche compacte et d est l'épaisseur d'une molécule d'eau. La quantité $\delta\phi_w$ représente la chute de potentiel, liée à la nature dipolaire des molécules d'eau. Le calcul de $\delta\phi_w$ est très largement inspiré des travaux d'A. Franco voir [39], ainsi que [37]. On a [39] :

$$\delta\phi_w(t) = \frac{\mu}{\varepsilon_{cd}} n_{\max} \left(\vec{\theta}_w(t) - \overleftarrow{\theta}_w(t) \right) \quad (\text{B.2})$$

où μ est le moment dipolaire d'une molécule d'eau, $\vec{\theta}_w$ (respectivement $\overleftarrow{\theta}_w$) est la fraction de recouvrement en dipôles orientés vers le métal (respectivement opposé au métal). ε_{cd} est la permittivité électrique dans la couche diffuse. D'après la loi d'action de masse, on a :

$$\vec{\theta}_w(t) = a_w \theta_s(t) \exp\left(-\frac{\overrightarrow{\Delta G^0}(t)}{RT}\right), \quad \overleftarrow{\theta}_w(t) = a_w \theta_s(t) \exp\left(-\frac{\overleftarrow{\Delta G^0}(t)}{RT}\right), \quad (\text{B.3})$$

où θ_s indique la fraction de sites catalytiques libres, a_w l'activité de l'eau dans la couche diffuse. ΔG^0 est l'enthalpie libre de la réaction d'adsorption dipolaire des molécules d'eau sur les sites catalytiques, elle est donnée par [37], [14] :

$$\overrightarrow{\Delta G^0}(t) = \frac{A\mu^2}{d^3}(\overrightarrow{\theta}_w(t) - \overleftarrow{\theta}_w(t)) + \frac{\mu}{\varepsilon_{cc}}\sigma(t) + \Delta G_{c,a}^0, \quad (\text{B.4a})$$

$$\overleftarrow{\Delta G^0}(t) = -\frac{A\mu^2}{d^3}(\overrightarrow{\theta}_w(t) - \overleftarrow{\theta}_w(t)) - \frac{\mu}{\varepsilon_{cc}}\sigma(t) + \Delta G_{c,a}^0 \quad (\text{B.4b})$$

$\Delta G_{c,a}^0$ est l'énergie chimique d'adsorption (supposée constante et indépendante de l'orientation dipolaire), $A = \frac{1.2}{2\pi\varepsilon_{cd}}$, ε_{cc} est la permittivité électrique dans la couche active. En introduisant les quantités $X(t)$ et K définies par :

$$X(t) = \frac{A\mu^2}{k_B T d^3}(\overrightarrow{\theta}_w(t) - \overleftarrow{\theta}_w(t)) + \frac{\mu}{k_B T \varepsilon_{cc}}\sigma(t), \quad K = 2 \exp\left(-\frac{\Delta G_{c,a}^0}{RT}\right), \quad (\text{B.5})$$

on en déduit d'après (B.3) et (B.4), les relations suivantes :

$$\overrightarrow{\theta}_w(t) = \frac{1}{2} K a_w \theta_s(t) \exp(X(t)), \quad \overleftarrow{\theta}_w(t) = \frac{1}{2} K a_w \theta_s(t) \exp(-X(t)) \quad (\text{B.6a})$$

$$\overrightarrow{\theta}_w(t) - \overleftarrow{\theta}_w(t) = -K a_w \theta_s(t) \sinh(X(t)) \quad (\text{B.6b})$$

où k_B est la constante de Boltzmann. En regroupant les équations (B.2) et (B.5) on obtient :

$$\begin{aligned} \delta\phi_w(t) &= \frac{\mu}{\varepsilon_{cd}} n_{\max} (\overrightarrow{\theta}_w(t) - \overleftarrow{\theta}_w(t)) \\ &= \frac{\mu}{\varepsilon_{cd}} n_{\max} \left(X(t) - \frac{\mu}{k_B T \varepsilon_{cc}} \sigma(t) \right) \frac{k_B T d^3}{A \mu^2} \\ &= \frac{d^3}{\varepsilon_{cd} A} n_{\max} \left(\frac{k_B T}{\mu} X(t) - \frac{\sigma(t)}{\varepsilon_{cc}} \right). \end{aligned} \quad (\text{B.7})$$

D'après (B.7), la formule (B.1) devient :

$$\eta(t) = - \left(1 + \frac{d^2 n_{\max}}{\varepsilon_{cd} A} \right) \frac{\sigma(t)}{\varepsilon_{cc}} d + \frac{k_B T d^3 n_{\max}}{\varepsilon_{cd} A \mu} X(t). \quad (\text{B.8})$$

Calcul de $X(t)$ Afin de calculer l'expression de X indépendamment de $\overrightarrow{\theta}_w$ et $\overleftarrow{\theta}_w$, on écrit :

$$\theta_s(t) = 1 - \theta^{int}(t) - \overrightarrow{\theta}_w(t) - \overleftarrow{\theta}_w(t)$$

où θ^{int} est la fraction molaires des espèces actifs intermédiaires de la réaction. On alors d'après (B.6a) :

$$\theta_s(t) = 1 - \theta^{int}(t) - K a_w \theta_s(t) \cosh(X(t)) \quad (\text{B.9})$$

d'où

$$\theta_s(t) = \frac{1 - \theta^{int}(t)}{1 + K a_w \cosh(X)}, \quad (\text{B.10})$$

En utilisant (B.6b), on a d'après (B.5) :

$$X(t) = \frac{\mu}{k_B T \varepsilon_{cc}} \sigma(t) - \frac{A\mu^2}{k_B T d^3} K a_w \theta_s(t) \sinh(X(t)). \quad (\text{B.11})$$

On a finalement d'après (B.11), (B.10) et (B.8) :

$$\eta(t) = - \left(1 + \frac{d^2 n_{\max}}{\varepsilon_{cd} A} \right) \frac{d}{\varepsilon_{cc}} \sigma(t) + \frac{k_B T d^3 n_{\max}}{\varepsilon_{cd} A \mu} X(t) \quad (\text{B.12a})$$

$$K a_w \theta_s(t) \sinh(X(t)) = \frac{d^3}{\varepsilon_{cc} A \mu} \sigma(t) - \frac{k_B T d^3}{A \mu^2} X(t) \quad (\text{B.12b})$$

$$\theta_s(t) = \frac{1 - \theta^{int}(t)}{1 + K a_w \cosh(X(t))}. \quad (\text{B.12c})$$

L'équation (B.12b) est transcendante en X , on montre qu'elle a une seule solution, qui dépend uniquement de σ et de θ^{int} .

Annexe C

Discrétisation des EDPs du modèle de Franco

On cherche ici à approximer l'équation (3.57) du modèle B par une semi-discrétisation, définies par les valeurs $\tilde{C}_N(\pm z_j, t)$, $j \in \{1, \dots, N+1\}$, telles que pour tout $j \in \{1, \dots, N+1\}$:

$$\begin{aligned} \frac{d}{dt} \int_{\pm z_{j-1}}^{\pm z_{j+1}} \tilde{C}_N(z, t) \varphi_j(z) dz = & \left(\alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) C_N(\pm z_{N+1}, t) \right) \varphi_j(\pm z_{N+1}) \\ & - D_+ \int_{\pm z_{j-1}}^{\pm z_{j+1}} \left(\frac{\partial \tilde{C}_N}{\partial z} + \frac{F}{RT} (\tilde{C}_N + C_{FIX}) \Phi_{\pm}(\tilde{C}_N, \mathcal{X}_{\pm}) \right) \frac{d\varphi_j}{dz} dz \end{aligned} \quad (\text{C.1a})$$

$$\Phi_{\pm}(\tilde{C}_{H+}, \mathcal{X}_{\pm})(z) = \mp \frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{\varepsilon_{cd}} \int_z^{\pm z_{N+1}} \tilde{C}_{H+}(z') dz' . \quad (\text{C.1b})$$

où

$$\tilde{C}_N(z, t) = \frac{(\pm z - z_{j-1}) \tilde{C}_N(\pm z_j, t) - (\pm z - z_j) \tilde{C}_N(\pm z_{j-1}, t)}{d_{j-1}}, \quad \pm z \in [z_{j-1}, z_j]. \quad (\text{C.2})$$

L'entier positif N en indice se réfère au nombre de points de discrétisation. Les fonctions test φ_j sont affines sur chaque intervalle $\pm[z_{j-1}, z_j]$ et $\varphi_j(\pm z_i) = \delta_{ij}$. En d'autres termes, pour tout $j = 1, \dots, N$ on a :

$$\varphi_j(z) = \begin{cases} \frac{\pm z - z_{j-1}}{d_{j-1}} & \text{si } \pm z \in [z_{j-1}, z_j], \\ \frac{z_{j+1} \mp z}{d_j} & \text{si } \pm z \in [z_j, z_{j+1}], \\ 0 & \text{si } \pm z \leq z_{j-1} \text{ ou } \pm z \geq z_{j+1}, \end{cases} \quad (\text{C.3a})$$

et

$$\varphi_{N+1}(z) = \begin{cases} \frac{\pm z - z_N}{d_N} & \text{si } \pm z \in [z_N, z_{N+1}], \\ 0 & \text{si } \pm z \leq z_N. \end{cases} \quad (\text{C.3b})$$

On rappelle que :

$$\mathcal{X}_-(t) \doteq \begin{pmatrix} C_{H_2}(-z_{N+1}, t) \\ \theta_{Hs}(t) \\ \sigma_-(t) \end{pmatrix}, \quad \mathcal{X}_+(t) \doteq \begin{pmatrix} C_{O_2}(z_{N+1}, t) \\ \theta_{OHs}(t) \\ \theta_{O_2Hs}(t) \\ \sigma_+(t) \end{pmatrix}$$

et

$$\begin{aligned} \alpha_-(t) &= -\gamma e_{CA} \left(k_{HEY} \theta_{s,-}(t) C_{H_2}(-z_{N+1}, t) + k_{VOL} \theta_{Hs}(t) \right) \exp \left(\frac{\eta_-(t)}{2RT/F} \right) \\ \alpha_+(t) &= \gamma e_{CA} \left(k_{-1} \theta_{O_2Hs}(t) \exp \left(\frac{(1-\alpha_1)\eta_+(t)}{RT/F} \right) + k_{-3} \theta_{s,+}(t) a_w \exp \left(\frac{(1-\alpha_3)\eta_+(t)}{RT/F} \right) \right) \\ \beta_-(t) &= \gamma e_{CA} \left(k_{-HEY} \theta_{Hs}(t) + k_{-VOL} \theta_{s,-}(t) \right) \exp \left(\frac{-\eta_-(t)}{2RT/F} \right) \\ \beta_+(t) &= -\gamma e_{CA} \left(k_1 \theta_{s,+}(t) C_{O_2}(z_{N+1}, t) \exp \left(\frac{-\alpha_1 \eta_+(t)}{RT/F} \right) + k_3 \theta_{OHs}(t) \exp \left(\frac{-\alpha_3 \eta_+(t)}{RT/F} \right) \right). \end{aligned}$$

Calcul du membre de gauche de (C.1a). On a, pour tout $j \in \{1, \dots, N+1\}$:

$$\begin{aligned} \int_{\pm z_{j-1}}^{\pm z_j} \tilde{C}_N(z) \varphi_j(z) dz &= \frac{1}{d_{j-1}^2} \int_{\pm z_{j-1}}^{\pm z_j} \left((\pm z - z_{j-1})^2 \tilde{C}_N(\pm z_j) - (\pm z - z_j)(\pm z - z_{j-1}) \tilde{C}_N(\pm z_{j-1}) \right) dz \\ &= \pm \frac{1}{d_{j-1}^2} \int_{z_{j-1}}^{z_j} \left((z - z_{j-1})^2 \tilde{C}_N(\pm z_j) - (z - z_j)(z - z_{j-1}) \tilde{C}_N(\pm z_{j-1}) \right) dz \\ &= \pm \frac{1}{d_{j-1}^2} \left(\frac{d_{j-1}^3}{3} \tilde{C}_N(\pm z_j) - \left(\frac{z_j^3 - z_{j-1}^3}{3} - \frac{z_j^2 - z_{j-1}^2}{2} (z_j + z_{j-1}) + z_j z_{j-1} (z_j - z_{j-1}) \right) \tilde{C}_N(\pm z_{j-1}) \right) \\ &= \pm \frac{d_{j-1}}{3} \tilde{C}_N(\pm z_j) \mp \frac{1}{6d_{j-1}} \left(2(z_j^2 + z_j z_{j-1} + z_{j-1}^2) - 3(z_j + z_{j-1})^2 + 6z_j z_{j-1} \right) \tilde{C}_N(\pm z_{j-1}) \\ &= \pm \frac{d_{j-1}}{6} \left(2\tilde{C}_N(\pm z_j) + \tilde{C}_N(\pm z_{j-1}) \right), \end{aligned} \tag{C.4a}$$

et de même pour tout $j \in \{1, \dots, N\}$

$$\int_{\pm z_j}^{\pm z_{j+1}} \tilde{C}_N(z) \varphi_j(z) dz = \pm \frac{d_j}{6} \left(2\tilde{C}_N(\pm z_j) + \tilde{C}_N(\pm z_{j+1}) \right). \tag{C.4b}$$

On en déduit que pour tout $j \in \{1, \dots, N\}$:

$$\begin{aligned} \frac{d}{dt} \int_{\pm z_{j-1}}^{\pm z_{j+1}} \tilde{C}_N(z, t) \varphi_j(z) dz &= \pm \left(\frac{d_{j-1} + d_j}{3} \frac{d\tilde{C}_N}{dt}(\pm z_j, t) + \frac{d_{j-1}}{6} \frac{d\tilde{C}_N}{dt}(\pm z_{j-1}, t) \right. \\ &\quad \left. + \frac{d_j}{6} \frac{d\tilde{C}_N}{dt}(\pm z_{j+1}, t) \right) \end{aligned} \tag{C.5a}$$

et

$$\frac{d}{dt} \int_{\pm z_N}^{\pm z_{N+1}} \tilde{C}_N(z, t) \varphi_{N+1}(z) dz = \pm \frac{d_N}{6} \left(2\frac{d\tilde{C}_N}{dt}(\pm z_{N+1}, t) + \frac{d\tilde{C}_N}{dt}(\pm z_N, t) \right). \tag{C.5b}$$

Calcul du second terme dans le membre de droite de (C.1a). On discrétise $\Phi_{\pm}(\tilde{C}_N, \mathcal{X}_{\pm})(z)$ de la manière suivante. Pour tout $j \in \{1, \dots, N+1\}$ et pour tout $\pm z \in [z_{j-1}, z_j]$, on a :

$$\begin{aligned}
 \Phi_{\pm}(\tilde{C}_N, \mathcal{X}_{\pm})(z) &= \mp \frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{\varepsilon_{cd}} \left[\sum_{i=j+1}^{N+1} \frac{1}{d_{i-1}} \int_{\pm z_{i-1}}^{\pm z_i} \left((\pm z' - z_{i-1}) \tilde{C}_N(\pm z_i) - (\pm z' - z_i) \tilde{C}_N(\pm z_{i-1}) \right) dz' \right. \\
 &\quad \left. + \frac{1}{d_{j-1}} \int_z^{\pm z_j} \left((\pm z' - z_{j-1}) \tilde{C}_N(\pm z_j) - (\pm z' - z_j) \tilde{C}_N(\pm z_{j-1}) \right) dz' \right] \\
 &= \mp \frac{\sigma_{\pm}}{\varepsilon_{cc}} \pm \frac{F}{\varepsilon_{cd}} \left[\sum_{i=j+1}^{N+1} \frac{1}{d_{i-1}} \int_{z_{i-1}}^{z_i} \left((z' - z_{i-1}) \tilde{C}_N(\pm z_i) - (z' - z_i) \tilde{C}_N(\pm z_{i-1}) \right) dz' \right. \\
 &\quad \left. + \frac{1}{d_{j-1}} \int_{\pm z}^{z_j} \left((z' - z_{j-1}) \tilde{C}_N(\pm z_j) - (z' - z_j) \tilde{C}_N(\pm z_{j-1}) \right) dz' \right] \\
 &= \mp \frac{\sigma_{\pm}}{\varepsilon_{cc}} \pm \frac{F}{\varepsilon_{cd}} \left[\sum_{i=j+1}^{N+1} d_{i-1} \frac{\tilde{C}_N(\pm z_i) + \tilde{C}_N(\pm z_{i-1})}{2} - \frac{(\pm z - z_{j-1})^2}{2d_{j-1}} \tilde{C}_N(\pm z_j) \right. \\
 &\quad \left. + \frac{(\pm z - z_j)^2}{2d_{j-1}} \tilde{C}_N(\pm z_{j-1}) + \frac{d_{j-1}}{2} \tilde{C}_N(\pm z_j) \right] \tag{C.6}
 \end{aligned}$$

Par suite de (C.3), (C.2) et (C.6), on a alors pour tout $j = \{1, \dots, N\}$:

$$\begin{aligned}
 &\int_{\pm z_j}^{\pm z_{j+1}} (\tilde{C}_N + C_{FIX}) \Phi_{\pm}(\tilde{C}_N, \mathcal{X}_{\pm}) \frac{d\varphi_j}{dz} dz \\
 &= \mp \frac{1}{d_j^2} \int_{\pm z_j}^{\pm z_{j+1}} \left[(\pm z - z_j) C_N(\pm z_{j+1}) - (\pm z - z_{j+1}) C_N(\pm z_j) \right] \times \\
 &\quad \left[\mp \frac{\sigma_{\pm}}{\varepsilon_{cc}} \pm \frac{F}{\varepsilon_{cd}} \left(\sum_{i=j+2}^{N+1} d_{i-1} \frac{\tilde{C}_N(\pm z_i) + \tilde{C}_N(\pm z_{i-1})}{2} - \frac{(\pm z - z_j)^2}{2d_j} \tilde{C}_N(\pm z_{j+1}) \right. \right. \\
 &\quad \left. \left. + \frac{(\pm z - z_{j+1})^2}{2d_j} \tilde{C}_N(\pm z_j) + \frac{d_j}{2} \tilde{C}_N(\pm z_{j+1}) \right) \right] dz \\
 &= \mp \frac{1}{d_j^2} \int_{z_j}^{z_{j+1}} \left[(z - z_j) C_N(\pm z_{j+1}) - (z - z_{j+1}) C_N(\pm z_j) \right] \times \\
 &\quad \left[- \frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{\varepsilon_{cd}} \left(\sum_{i=j+2}^{N+1} d_{i-1} \frac{\tilde{C}_N(\pm z_i) + \tilde{C}_N(\pm z_{i-1})}{2} - \frac{(z - z_j)^2}{2d_j} \tilde{C}_N(\pm z_{j+1}) \right. \right. \\
 &\quad \left. \left. + \frac{(z - z_{j+1})^2}{2d_j} \tilde{C}_N(\pm z_j) + \frac{d_j}{2} \tilde{C}_N(\pm z_{j+1}) \right) \right] dz. \tag{C.7}
 \end{aligned}$$

On calcule l'expression précédente en deux temps. D'une part, pour tout $j = \{1, \dots, N\}$:

$$\begin{aligned}
 & \int_{z_j}^{z_{j+1}} \left[(z - z_j)C_N(\pm z_{j+1}) - (z - z_{j+1})C_N(\pm z_j) \right] \times \\
 & \quad \left[-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{\varepsilon_{cd}} \left(\sum_{i=j+2}^{N+1} d_{i-1} \frac{\tilde{C}_N(\pm z_i) + \tilde{C}_N(\pm z_{i-1})}{2} + \frac{d_j}{2} \tilde{C}_N(\pm z_{j+1}) \right) \right] dz \\
 &= \frac{d_j^2}{2} (C_N(\pm z_j) + C_N(\pm z_{j+1})) \left[-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{2\varepsilon_{cd}} \left(\sum_{i=j+2}^{N+1} d_{i-1} (\tilde{C}_N(\pm z_i) + \tilde{C}_N(\pm z_{i-1})) + d_j \tilde{C}_N(\pm z_{j+1}) \right) \right] \\
 &= \frac{d_j^2}{2} (C_N(\pm z_j) + C_N(\pm z_{j+1})) \left[-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{2\varepsilon_{cd}} \left(\sum_{i=j+1}^N (d_i + d_{i-1}) \tilde{C}_N(\pm z_i) + d_N \tilde{C}_N(\pm z_{N+1}) \right) \right]. \tag{C.8}
 \end{aligned}$$

Pour obtenir la dernière égalité, on a utilisé le fait que :

$$\sum_{i=j+2}^{N+1} d_{i-1} (\tilde{C}_N(\pm z_i) + \tilde{C}_N(\pm z_{i-1})) + d_j \tilde{C}_N(\pm z_{j+1}) = \sum_{i=j+1}^N (d_i + d_{i-1}) \tilde{C}_N(\pm z_i) + d_N \tilde{C}_N(\pm z_{N+1}). \tag{C.9}$$

D'autre part, on vérifie¹ que pour tout $j = \{1, \dots, N\}$:

$$\begin{aligned}
 & \int_{z_j}^{z_{j+1}} \left((z - z_j)C_N(\pm z_{j+1}) - (z - z_{j+1})C_N(\pm z_j) \right) \left((z - z_j)^2 \tilde{C}_N(\pm z_{j+1}) - (z - z_{j+1})^2 \tilde{C}_N(\pm z_j) \right) dz \\
 &= \frac{d_j^4}{12} (C_N(\pm z_j) + 3C_N(\pm z_{j+1})) \tilde{C}_N(\pm z_{j+1}) - (C_N(\pm z_{j+1}) + 3C_N(\pm z_j)) \tilde{C}_N(\pm z_j). \tag{C.10}
 \end{aligned}$$

Pour tout $j = \{1, \dots, N\}$, l'équation (C.7) s'écrit donc, en utilisant (C.8) et (C.10), comme suit :

$$\begin{aligned}
 & \int_{\pm z_j}^{\pm z_{j+1}} (\tilde{C}_N + C_{FIX}) \Phi_{\pm}(\tilde{C}_N, \mathcal{X}_{\pm}) \frac{d\varphi_j}{dz} dz \\
 &= \mp \frac{1}{2} (C_N(\pm z_j) + C_N(\pm z_{j+1})) \left[-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{2\varepsilon_{cd}} \left(\sum_{i=j+1}^N (d_i + d_{i-1}) \tilde{C}_N(\pm z_i) + d_N \tilde{C}_N(\pm z_{N+1}) \right) \right] \\
 & \quad \pm \frac{d_j}{24 \varepsilon_{cd}} \left[(C_N(\pm z_j) + 3C_N(\pm z_{j+1})) \tilde{C}_N(\pm z_{j+1}) - (C_N(\pm z_{j+1}) + 3C_N(\pm z_j)) \tilde{C}_N(\pm z_j) \right] \tag{C.11a}
 \end{aligned}$$

1. On utilise le fait que :

$$\int_{z_j}^{z_{j+1}} (z - z_j)^3 dz = \frac{(z_{j+1} - z_j)^4}{4}$$

et

$$\int_{z_j}^{z_{j+1}} (z - z_{j+1})(z - z_j)^2 dz = \frac{1}{3} [(z - z_{j+1})(z - z_j)^3]_{z_j}^{z_{j+1}} - \frac{1}{3} \int_{z_j}^{z_{j+1}} (z - z_j)^3 dz = -\frac{1}{12} (z_{j+1} - z_j)^4$$

De la même manière, on a pour tout $j \in \{1, \dots, N+1\}$:

$$\begin{aligned}
 \int_{\pm z_{j-1}}^{\pm z_j} (\tilde{C}_N + C_{FIX}) \Phi_{\pm}(\tilde{C}_N, \mathcal{X}_{\pm}) \frac{d\varphi_j}{dz} dz &= - \int_{\pm z_{j-1}}^{\pm z_j} (\tilde{C}_N + C_{FIX}) \Phi_{\pm}(\tilde{C}_N, \mathcal{X}_{\pm}) \frac{d\varphi_{j-1}}{dz} dz \\
 &= \pm \frac{1}{2} (C_N(\pm z_{j-1}) + C_N(\pm z_j)) \left[-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{2\varepsilon_{cd}} \left(\sum_{i=j}^N (d_i + d_{i-1}) \tilde{C}_N(\pm z_i) + d_N \tilde{C}_N(\pm z_{N+1}) \right) \right] \\
 &\quad \mp \frac{d_{j-1}}{24} \frac{F}{\varepsilon_{cd}} \left[(C_N(\pm z_{j-1}) + 3C_N(\pm z_j)) \tilde{C}_N(\pm z_j) - (C_N(\pm z_j) + 3C_N(\pm z_{j-1})) \tilde{C}_N(\pm z_{j-1}) \right].
 \end{aligned} \tag{C.11b}$$

Nous calculons maintenant l'intégrale du produit des dérivées dans le second terme du membre de droite de (C.1a). Nous avons pour tout $j \in \{1, \dots, N\}$:

$$\begin{aligned}
 \int_{\pm z_{j-1}}^{\pm z_{j+1}} \frac{\partial \tilde{C}_N}{\partial z} \frac{d\varphi_j}{dz} dz &= \int_{\pm z_{j-1}}^{\pm z_j} \frac{\tilde{C}_N(\pm z_j) - \tilde{C}_N(\pm z_{j-1})}{d_{j-1}^2} dz - \int_{\pm z_j}^{\pm z_{j+1}} \frac{\tilde{C}_N(\pm z_{j+1}) - \tilde{C}_N(\pm z_j)}{d_j^2} dz \\
 &= \pm \int_{z_{j-1}}^{z_j} \frac{\tilde{C}_N(\pm z_j) - \tilde{C}_N(\pm z_{j-1})}{d_{j-1}^2} dz \mp \int_{z_j}^{z_{j+1}} \frac{\tilde{C}_N(\pm z_{j+1}) - \tilde{C}_N(\pm z_j)}{d_j^2} dz \\
 &= \mp \left(\frac{1}{d_{j-1}} \tilde{C}_N(\pm z_{j-1}) - \left(\frac{1}{d_{j-1}} + \frac{1}{d_j} \right) \tilde{C}_N(\pm z_j) + \frac{1}{d_j} \tilde{C}_N(\pm z_{j+1}) \right)
 \end{aligned} \tag{C.12a}$$

et

$$\int_{\pm z_N}^{\pm z_{N+1}} \frac{\partial \tilde{C}_N}{\partial z} \frac{d\varphi_{N+1}}{dz} dz = \pm \frac{\tilde{C}_N(\pm z_{N+1}) - \tilde{C}_N(\pm z_N)}{d_N}. \tag{C.12b}$$

En regroupant les équations (C.12) et (C.11), on obtient pour tout $j \in \{1, \dots, N\}$:

$$\begin{aligned}
 &\pm \int_{\pm z_{j-1}}^{\pm z_{j+1}} \left(\frac{\partial \tilde{C}_N}{\partial z} + \frac{F}{RT} (\tilde{C}_N + C_{FIX}) \Phi_{\pm}(\tilde{C}_N, \mathcal{X}_{\pm}) \right) \frac{d\varphi_j}{dz} dz \\
 &= -\frac{1}{d_{j-1}} \tilde{C}_N(\pm z_{j-1}) + \left(\frac{1}{d_{j-1}} + \frac{1}{d_j} \right) \tilde{C}_N(\pm z_j) - \frac{1}{d_j} \tilde{C}_N(\pm z_{j+1}) \\
 &\quad - \frac{F}{2RT} (C_N(\pm z_j) + C_N(\pm z_{j+1})) \left(-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{2\varepsilon_{cd}} \left(\sum_{i=j+1}^N (d_i + d_{i-1}) \tilde{C}_N(\pm z_i) + d_N \tilde{C}_N(\pm z_{N+1}) \right) \right) \\
 &\quad + \frac{F}{2RT} (C_N(\pm z_{j-1}) + C_N(\pm z_j)) \left(-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{2\varepsilon_{cd}} \left(\sum_{i=j}^N (d_i + d_{i-1}) \tilde{C}_N(\pm z_i) + d_N \tilde{C}_N(\pm z_{N+1}) \right) \right) \\
 &\quad + \frac{F^2 d_j}{24RT\varepsilon_{cd}} \left[(C_N(\pm z_j) + 3C_N(\pm z_{j+1})) \tilde{C}_N(\pm z_{j+1}) - (C_N(\pm z_{j+1}) + 3C_N(\pm z_j)) \tilde{C}_N(\pm z_j) \right] \\
 &\quad - \frac{F^2 d_{j-1}}{24RT\varepsilon_{cd}} \left[(C_N(\pm z_{j-1}) + 3C_N(\pm z_j)) \tilde{C}_N(\pm z_j) - (C_N(\pm z_j) + 3C_N(\pm z_{j-1})) \tilde{C}_N(\pm z_{j-1}) \right]
 \end{aligned}$$

$$\begin{aligned}
&= -\frac{1}{d_{j-1}}\tilde{C}_N(\pm z_{j-1}) + \left(\frac{1}{d_{j-1}} + \frac{1}{d_j}\right)\tilde{C}_N(\pm z_j) - \frac{1}{d_j}\tilde{C}_N(\pm z_{j+1}) \\
&+ \frac{F}{2RT}\left(C_N(\pm z_{j-1}) - C_N(\pm z_{j+1})\right)\left(-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{2\varepsilon_{cd}}\left(\sum_{i=j}^N(d_i + d_{i-1})\tilde{C}_N(\pm z_i) + d_N\tilde{C}_N(\pm z_{N+1})\right)\right) \\
&+ \frac{F^2}{24RT\varepsilon_{cd}}\left[d_{j-1}\left(C_N(\pm z_j) + 3C_N(\pm z_{j-1})\right)\tilde{C}_N(\pm z_{j-1}) + d_j\left(C_N(\pm z_j) + 3C_N(\pm z_{j+1})\right)\tilde{C}_N(\pm z_{j+1})\right. \\
&\quad \left.+ \left((5d_j + 6d_{j-1})C_N(\pm z_{j+1}) + 3(d_j + d_{j-1})C_N(\pm z_j) - d_{j-1}C_N(\pm z_{j-1})\right)\tilde{C}_N(\pm z_j)\right]
\end{aligned} \tag{C.13a}$$

et d'après (C.11b) et (C.12b), on a :

$$\begin{aligned}
&\pm \int_{\pm z_N}^{\pm z_{N+1}} \left(\frac{\partial \tilde{C}_N}{\partial z} + \frac{F}{RT}(\tilde{C}_N + C_{FIX})\Phi_{\pm}(\tilde{C}_N, \mathcal{X}_{\pm}) \right) \frac{d\varphi_{N+1}}{dz} dz = \frac{\tilde{C}_N(\pm z_{N+1}) - \tilde{C}_N(\pm z_N)}{d_N} \\
&- \frac{F}{2RT}\left(C_N(\pm z_N) + C_N(\pm z_{N+1})\right)\frac{\sigma_{\pm}}{\varepsilon_{cc}} - \frac{F^2 d_N}{24RT\varepsilon_{cd}}\left[\left(C_N(\pm z_N) + 3C_N(\pm z_{N+1})\right)\tilde{C}_N(\pm z_{N+1})\right. \\
&\quad \left.- \left(C_N(\pm z_{N+1}) + 3C_N(\pm z_N)\right)\tilde{C}_N(\pm z_N)\right].
\end{aligned} \tag{C.13b}$$

On introduit les vecteurs $\tilde{C}_{\pm}, C_{\pm} \in \mathbb{R}^{N+1}$, $u_k \in \mathbb{R}^{N+1}$, $v_k \in \mathbb{R}^{N+2}$, $k = 1, \dots, N+1$ et les matrices M_i , $i = 1, \dots, 9$, définies par :

$$C_{\pm}(t) \doteq \begin{pmatrix} C_N(\pm z_1, t) \\ \vdots \\ C_N(\pm z_{N+1}, t) \end{pmatrix}, \quad \tilde{C}_{\pm}(t) \doteq \begin{pmatrix} \tilde{C}_N(\pm z_1, t) \\ \vdots \\ \tilde{C}_N(\pm z_{N+1}, t) \end{pmatrix} = \begin{pmatrix} C_N(\pm z_1, t) - C_{FIX} \\ \vdots \\ C_N(\pm z_{N+1}, t) - C_{FIX} \end{pmatrix} \in \mathbb{R}^{N+1} \tag{C.14}$$

$$u_k = \begin{pmatrix} \delta_{k1} \\ \vdots \\ \delta_{k(N+1)} \end{pmatrix}, \quad v_k = \begin{pmatrix} \delta_{k1} \\ \vdots \\ \delta_{k(N+2)} \end{pmatrix}, \quad \delta_{ij} = \begin{cases} 1 & \text{si } i = j \\ 0 & \text{si } i \neq j \end{cases} \tag{C.15}$$

$$\begin{aligned}
 M_1 &\doteq -\frac{1}{6D_+} \begin{pmatrix} 2(d_0 + d_1) & d_1 & & & \\ & d_1 & 2(d_1 + d_2) & d_2 & \\ & & \ddots & \ddots & \ddots \\ & & & d_{N-1} & 2(d_{N-1} + d_N) & d_N \\ & & & & d_N & 2d_N \end{pmatrix} \\
 &= -\frac{1}{6D_+} \left[v_1 \left(2(d_0 + d_1)v_1 + d_1v_2 \right)^\top + \sum_{i=2}^N v_i \left(d_{i-1}v_{i-1} + 2(d_{i-1} + d_i)v_i + d_iv_{i+1} \right)^\top \right. \\
 &\quad \left. + v_{N+1}d_N(v_N + 2v_{N+1})^\top \right] \in \mathbb{R}^{(N+1) \times (N+1)} \quad (\text{C.16a})
 \end{aligned}$$

$$\begin{aligned}
 M_2 &\doteq \begin{pmatrix} \frac{1}{d_0} + \frac{1}{d_1} & -\frac{1}{d_1} & & & \\ -\frac{1}{d_1} & \frac{1}{d_1} + \frac{1}{d_2} & -\frac{1}{d_2} & & \\ & \ddots & \ddots & \ddots & \\ & & -\frac{1}{d_{N-1}} & \frac{1}{d_{N-1}} + \frac{1}{d_N} & -\frac{1}{d_N} \\ & & & -\frac{1}{d_N} & \frac{1}{d_N} \end{pmatrix} \\
 &= v_1 \left(\left(\frac{1}{d_0} + \frac{1}{d_1} \right) v_1 - \frac{1}{d_1} v_2 \right)^\top + \sum_{i=2}^N v_i \left(-\frac{1}{d_{i-1}} v_{i-1} + \left(\frac{1}{d_{i-1}} + \frac{1}{d_i} \right) v_i - \frac{1}{d_i} v_{i+1} \right)^\top \\
 &\quad + \frac{1}{d_N} v_{N+1} (-v_N + v_{N+1})^\top \in \mathbb{R}^{(N+1) \times (N+1)} \quad (\text{C.16b})
 \end{aligned}$$

$$\begin{aligned}
 M_3 &\doteq \frac{F}{2RT} \begin{pmatrix} 1 & 0 & -1 & & \\ & 1 & 0 & -1 & \\ & & \ddots & \ddots & \ddots \\ & & & 1 & 0 & -1 \\ & & & & -1 & -1 \end{pmatrix} \\
 &= \frac{F}{2RT} \left(\sum_{i=1}^N v_i (u_i - u_{i+2})^\top - v_{N+1} (u_{N+1} + u_{N+2})^\top \right) \in \mathbb{R}^{(N+1) \times (N+2)} \quad (\text{C.16c})
 \end{aligned}$$

$$\begin{aligned}
 M_4 &\doteq \frac{F}{2\varepsilon_{cd}} \begin{pmatrix} d_0 + d_1 & d_1 + d_2 & \dots & d_{N-1} + d_N & d_N \\ & d_1 + d_2 & \dots & d_{N-1} + d_N & d_N \\ & & \ddots & & \\ & & & d_{N-1} + d_N & d_N \\ & & & & 0 \end{pmatrix} \\
 &= \frac{F}{2\varepsilon_{cd}} \sum_{i=1}^N v_i \left(\sum_{j=i}^N (d_{j-1} + d_j) v_j + d_N v_{N+1} \right)^\top \in \mathbb{R}^{(N+1) \times (N+1)} \quad (\text{C.16d})
 \end{aligned}$$

$$M_5^+ = -\frac{1}{\varepsilon_{cc}} \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} \begin{pmatrix} 0 & 0 & 1 \end{pmatrix} \in \mathbb{R}^{(N+1) \times 3}, \quad M_5^- = -\frac{1}{\varepsilon_{cc}} \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} \begin{pmatrix} 0 & 0 & 0 & 1 \end{pmatrix} \in \mathbb{R}^{(N+1) \times 4} \quad (\text{C.16e})$$

$$\begin{aligned}
 M_6 &\doteq \frac{F^2}{24RT\varepsilon_{cd}} \begin{pmatrix} d_1 & 3d_1 & & & \\ & d_2 & 3d_2 & & \\ & & \ddots & \ddots & \\ & & & d_N & 3d_N \\ & & & & 0 \end{pmatrix} \\
 &= \frac{F^2}{24RT\varepsilon_{cd}} \sum_{i=1}^N v_i d_i (v_i + 3v_{i+1})^\top \in \mathbb{R}^{(N+1) \times (N+1)} \quad (\text{C.16f})
 \end{aligned}$$

$$\begin{aligned}
 M_8 &\doteq \frac{F^2}{24RT\varepsilon_{cd}} \begin{pmatrix} 0 & & & & \\ 5d_1 + 6d_2 & 3d_1 + 6d_2 & & & \\ & 5d_2 + 6d_3 & 3d_2 + 6d_3 & & \\ & & \ddots & \ddots & \\ & & & 5d_{N-1} + 6d_N & 3d_{N-1} + 6d_N \\ & & & & 3d_N & d_N \end{pmatrix} \\
 &= \frac{F^2}{24RT\varepsilon_{cd}} \sum_{i=2}^N v_i ((5d_{i-1} + 6d_i) v_{i-1} + (3d_{i-1} + 6d_i) v_i)^\top + v_{N+1} d_N (3v_N + v_{N+1})^\top \in \mathbb{R}^{(N+1) \times (N+1)} \quad (\text{C.16g})
 \end{aligned}$$

$$\begin{aligned}
 M_9 &\doteq \begin{pmatrix} 0_{1 \times N} & 0 \\ I_N & 0_{N \times 1} \end{pmatrix} \in \mathbb{R}^{(N+1) \times (N+1)}, \quad M_7 = M_9^\top \in \mathbb{R}^{(N+1) \times (N+1)}, \\
 M_{10} &\doteq \frac{F^2}{24RT\varepsilon_{cd}} \begin{pmatrix} d_0 & 3(d_1 + d_0) & d_1 & & & \\ & d_1 & 3(d_2 + d_1) & d_2 & & \\ & & \ddots & & \ddots & \ddots \\ & & & d_{N-1} & 3(d_N + d_{N-1}) & d_N \\ & & & & -d_N & -3d_N \end{pmatrix} \\
 &= \frac{F^2}{24RT\varepsilon_{cd}} \left[\sum_{i=1}^N v_i \left(d_{i-1}u_i + 3(d_i + d_{i-1})u_{i+1} + d_i u_{i+2} \right)^\top - v_{N+1} d_N (v_N + v_{N+1})^\top \right] \in \mathbb{R}^{(N+1) \times (N+2)}.
 \end{aligned} \tag{C.16h}$$

Le système des équations données par (C.1a) s'écrit alors sous la forme :

$$\frac{d\tilde{\mathcal{C}}_\pm(t)}{dt} = \mathcal{G}_\pm(\mathcal{X}_\pm(t), \mathcal{C}_\pm(t)) \tag{C.17}$$

où

$$\begin{aligned}
 \mathcal{G}_\pm(\mathcal{X}_\pm, \mathcal{C}_\pm) &= M_1^{-1} \left[M_2 \tilde{\mathcal{C}}_\pm + M_3 \begin{pmatrix} C_{FIX} \\ \mathcal{C}_\pm \end{pmatrix} \cdot \left(M_4 \tilde{\mathcal{C}}_\pm + M_5^\pm \mathcal{X}_\pm \right) + M_6 \mathcal{C}_\pm \cdot M_7 \tilde{\mathcal{C}}_\pm \right. \\
 &\quad \left. + M_8 \mathcal{C}_\pm \cdot M_9 \tilde{\mathcal{C}}_\pm + M_{10} \begin{pmatrix} C_{FIX} \\ \mathcal{C}_\pm \end{pmatrix} \cdot \tilde{\mathcal{C}}_\pm \pm \left(\alpha_\pm(\mathcal{X}_\pm) + \beta_\pm(\mathcal{X}_\pm) v_{N+1}^\top \mathcal{C}_\pm \right) v_{N+1} \right] \tag{C.18}
 \end{aligned}$$

où “ $\cdot$ ” représente le produit terme à terme de vecteurs.

Sortie du modèle. La tension aux bornes de chaque électrode de la pile s'écrit sous la forme :

$$U_\pm = \int_{\pm z_0}^{\pm z_{N+1}} \frac{\partial \phi}{\partial z} dz = \sum_{j=1}^{N+1} \int_{\pm z_{j-1}}^{\pm z_j} \Phi_\pm(\tilde{C}_N, \mathcal{X}_\pm)(z) dz. \tag{C.19}$$

En utilisant (C.6), l'équation précédente devient :

$$\begin{aligned}
 U_{\pm}(t) &= \sum_{j=1}^{N+1} \int_{\pm z_{j-1}}^{\pm z_j} \left[\mp \frac{\sigma_{\pm}}{\varepsilon_{cc}} \pm \frac{F}{\varepsilon_{cd}} \left(\sum_{i=j+1}^{N+1} d_{i-1} \frac{\tilde{C}_N(\pm z_i) + \tilde{C}_N(\pm z_{i-1})}{2} - \frac{(\pm z - z_{j-1})^2}{2d_{j-1}} \tilde{C}_N(\pm z_j) \right. \right. \\
 &\quad \left. \left. + \frac{(\pm z - z_j)^2}{2d_{j-1}} \tilde{C}_N(\pm z_{j-1}) + \frac{d_{j-1}}{2} \tilde{C}_N(\pm z_{j-1}) \right) \right] dz \\
 &= \sum_{j=1}^{N+1} \int_{z_{j-1}}^{z_j} \left[-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{\varepsilon_{cd}} \left(\sum_{i=j+1}^{N+1} d_{i-1} \frac{\tilde{C}_N(\pm z_i) + \tilde{C}_N(z_{i-1})}{2} - \frac{(z - z_{j-1})^2}{2d_{j-1}} \tilde{C}_N(\pm z_j) \right. \right. \\
 &\quad \left. \left. + \frac{(z - z_j)^2}{2d_{j-1}} \tilde{C}_N(\pm z_{j-1}) + \frac{d_{j-1}}{2} \tilde{C}_N(\pm z_{j-1}) \right) \right] dz \\
 &= \sum_{j=1}^{N+1} d_{j-1} \left[-\frac{\sigma_{\pm}}{\varepsilon_{cc}} + \frac{F}{2\varepsilon_{cd}} \left(\sum_{i=j+1}^{N+1} d_{i-1} (\tilde{C}_N(\pm z_i) + \tilde{C}_N(\pm z_{i-1})) - \frac{d_{j-1}}{3} \tilde{C}_N(\pm z_j) \right. \right. \\
 &\quad \left. \left. + \frac{4d_{j-1}}{3} \tilde{C}_N(\pm z_{j-1}) \right) \right]. \tag{C.20}
 \end{aligned}$$

En utilisant (C.9), on obtient

$$\begin{aligned}
 U_{\pm}(t) &= -\frac{\sigma_{\pm}}{\varepsilon_{cc}} L + \frac{F}{2\varepsilon_{cd}} \sum_{j=1}^{N+1} d_{j-1} \left[\sum_{i=j+1}^N (d_i + d_{i-1}) \tilde{C}_N(\pm z_i) + d_N \tilde{C}_N(\pm z_{N+1}) \right. \\
 &\quad \left. + \left(d_j - \frac{d_{j-1}}{3} \right) \tilde{C}_N(\pm z_j) + \frac{4}{3} d_{j-1} \tilde{C}_N(\pm z_{j-1}) \right] \\
 &= -\frac{\sigma_{\pm}}{\varepsilon_{cc}} L + \frac{F}{2\varepsilon_{cd}} \left[\sum_{i=2}^N \sum_{j=0}^{i-2} d_j (d_i + d_{i-1}) \tilde{C}_N(\pm z_i) + \frac{1}{3} \sum_{i=1}^N (4d_i^2 - 3d_i d_{i-1} + d_{i-1}^2) \tilde{C}_N(\pm z_i) \right. \\
 &\quad \left. + d_N \left(d_{N+1} - \frac{d_N}{3} + L \right) \tilde{C}_N(\pm z_{N+1}) \right] \\
 &= \left(-\frac{\sigma_{\pm}}{\varepsilon_{cc}} L + \frac{F}{2\varepsilon_{cd}} w^{\top} \right) \tilde{C}_{\pm}(t) \tag{C.21}
 \end{aligned}$$

où L , l'épaisseur de la couche diffuse, vaut $\sum_{i=1}^{N+1} d_{j-1}$, et

$$w \doteq \begin{pmatrix} d_0 \left(d_1 - \frac{d_0}{3} \right) + \frac{4}{3} d_1^2 \\ d_1 \left(d_2 - \frac{d_1}{3} \right) + d_0(d_2 + d_1) + \frac{4}{3} d_2^2 \\ \vdots \\ d_{N-1} \left(d_N - \frac{d_{N-1}}{3} \right) + (d_N + d_{N-1}) \sum_{j=0}^{N-2} d_j + \frac{4}{3} d_N^2 \\ d_N \left(d_{N+1} - \frac{d_N}{3} \right) + (d_{N+1} + d_N) \sum_{j=0}^{N-1} d_j + L d_N \end{pmatrix} \in \mathbb{R}^{N+1}. \tag{C.22}$$

Calcul de courant dans la membrane, $i_{H^+,M}(t)$. On a d'après (3.47a)

$$\frac{\partial C_{H^+}}{\partial t} = D_+ \frac{\partial}{\partial z} \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right), \quad \pm z \in]z_0, z_{N+1}[\tag{C.23}$$

avec

$$D_+ \frac{\partial}{\partial z} \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right) (\pm z_{N+1}, t) = \alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) v_{N+1}^{\top} \mathcal{C}_{\pm}(t) \quad (\text{C.24})$$

et

$$D_+ \frac{\partial}{\partial z} \left(\frac{\partial C_{H^+}}{\partial z} + \frac{F}{RT} C_{H^+} \phi_z \right) (\pm z_0, t) = \frac{1}{FS_{EME}} i_{H^+,M}(t) \quad (\text{C.25})$$

En intégrant (C.23) en z entre $\pm z_0$ et $\pm z_{N+1}$, on obtient en utilisant (C.24) et (C.25) :

$$\frac{d}{dt} \int_{\pm z_0}^{\pm z_{N+1}} C_{H^+}(z, t) dz = \alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) v_{N+1}^{\top} \mathcal{C}_{\pm}(t) - \frac{1}{FS_{EME}} i_{H^+,M}(t) \quad (\text{C.26})$$

D'où

$$i_{H^+,M}(t) = FS_{EME} \left(\alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) v_{N+1}^{\top} \mathcal{C}_{\pm}(t) - \frac{d}{dt} \int_{\pm z_0}^{\pm z_{N+1}} C_{H^+}(z, t) dz \right) \quad (\text{C.27})$$

On approxime $C_{H^+}(z, t)$ dans (C.27) par $\tilde{C}_N(z, t) + C_{FIX}$, donné par (3.58). On a :

$$\begin{aligned} \int_{\pm z_0}^{\pm z_{N+1}} \tilde{C}_N(z, t) dz &= \frac{d_0}{2} \tilde{C}(z_1, t) + \frac{d_1}{2} (\tilde{C}(z_1, t) + \tilde{C}(z_2, t)) + \dots + \frac{d_N}{2} (\tilde{C}(z_N, t) + \tilde{C}(z_{N+1}, t)) \\ &= \frac{d_0 + d_1}{2} \tilde{C}(z_1, t) + \frac{d_1 + d_2}{2} \tilde{C}(z_2, t) + \dots + \frac{d_{N-1} + d_N}{2} \tilde{C}(z_N, t) + \frac{d_N}{2} \tilde{C}(z_{N+1}, t) \\ &= \frac{1}{2} v^{\top} \tilde{\mathcal{C}}_{\pm}(t) \end{aligned}$$

où v est donné par

$$v \doteq \begin{pmatrix} d_0 + d_1 \\ d_1 + d_2 \\ \vdots \\ d_{N-1} + d_N \\ d_N \end{pmatrix} \in \mathbb{R}^{N+1}. \quad (\text{C.28})$$

En utilisant (C.17), on a alors :

$$i_{H^+,M}(t) = FS_{EME} \left(\alpha_{\pm}(\mathcal{X}_{\pm}(t)) + \beta_{\pm}(\mathcal{X}_{\pm}(t)) v_{N+1}^{\top} \mathcal{C}_{\pm}(t) - \frac{1}{2} v^{\top} \mathcal{G}_{\pm}(\mathcal{X}_{\pm}(t), v_{N+1}^{\top} \mathcal{C}_{\pm}(t)) \right) \quad (\text{C.29})$$

Potentiel électrique ϕ . On a pour tout $j \in \{1 \dots N+1\}$:

$$\phi(\pm z_j) - \phi(\pm z_{j-1}) = \int_{\pm z_{j-1}}^{\pm z_j} \phi_z(z) dz = \int_{\pm z_{j-1}}^{\pm z_j} \Phi_{\pm}(\tilde{C}_N)(z, t) dz. \quad (\text{C.30})$$

En intégrant $\Phi_{\pm}(\tilde{C}_N)(z, t)$, donnée par (C.6), en z entre $\pm z_{j-1}$ et $\pm z_j$, on obtient :

$$\begin{aligned} \phi(\pm z_j) - \phi(\pm z_{j-1}) &= d_{j-1} \left[\frac{\alpha_M}{d_0} \tilde{C}_N(\pm z_1, t) - \frac{F}{2\varepsilon_{cd}} \left(\sum_{i=1}^{j-1} d_{i-1} (\tilde{C}_N(\pm z_i, t) + \tilde{C}_N(\pm z_{i-1}, t)) \right. \right. \\ &\quad \left. \left. + \frac{d_{j-1}}{3} \tilde{C}_N(\pm z_j, t) + \frac{2d_{j-1}}{3} \tilde{C}_N(\pm z_{j-1}, t) \right) \right] \quad (\text{C.31}) \end{aligned}$$

En utilisant (C.9), l'équation précédente s'écrit sous la forme :

$$\begin{aligned} \phi(\pm z_j) - \phi(\pm z_{j-1}) &= \frac{d_{j-1}}{d_0} \alpha_M \tilde{C}_N(\pm z_1) \\ &\quad - \frac{F}{2\varepsilon_{cd}} \left[\sum_{i=1}^{j-1} (d_i + d_{i-1}) \tilde{C}_N(\pm z_i, t) + \frac{d_{j-1}}{3} \left(\tilde{C}_N(\pm z_j, t) - \tilde{C}_N(\pm z_{j-1}, t) \right) \right] \end{aligned} \quad (\text{C.32})$$

Annexe D

Outils généraux de surveillance et de diagnostic à base de modèles

Comme tous les procédés industriels, les piles à combustible nécessitent la mise en place d'une procédure de supervision. Cette procédure doit permettre de détecter rapidement les défaillances graves pour assurer la sécurité, et les dégradations lentes pour optimiser la maintenance. Deux grands types de méthodes de surveillance et de diagnostic des pannes peuvent être distingués : d'une part, les méthodes liées essentiellement à l'analyse des signaux acquis sans connaissance de modèle et de comportement. D'autre part, les méthodes fondées sur un modèle mathématique du système à surveiller, celles-ci sont basées sur des méthodes analytiques des domaines de l'automatique, du traitement du signal et la statistique. Dans la première section nous traitons uniquement les méthodes basées sur des modèles, et dans la deuxième section nous présentons les outils électriques spécifiques pour la surveillance et la diagnostic de la pile. Cette exposition s'inspire largement des thèses de C. Cussenot [31] et d'O. Perrin [68].

La surveillance et le diagnostic de pannes basés sur des modèles reposent, de manière générale, sur la comparaison entre les mesures fournies par les capteurs du système surveillé et les estimés de ces mesures fournies par le modèle du système surveillé. L'écart entre le système réel (sortie des capteurs) et le modèle physique (sortie du modèle) est appelé "résidu". Ces résidus reflètent les changements possibles dans le système ; leur valeur moyenne est proche de zéro en l'absence de changement et différente de zéro si un changement s'est produit. Le principe de la surveillance et du diagnostic de pannes se décompose en deux étapes : la génération et l'évaluation des résidus. Le résidu, une fois généré, est examiné en tant que signature de défauts, dont le principe est d'observer des valeurs non nulles, lors de l'apparition d'un défaut. À partir de là, il est nécessaire de définir de règles de décision (ou détecteur) basées sur ces résidus. En effet, on doit décider si l'écart au zéro détecté sur les résidus est dû à un défaut (nature et localisation du défaut) ou tout simplement au "bruit" (incertitudes de mesure, erreurs de modélisation...). La méthode locale apporte une réponse à l'évaluation des résidus dans un cadre stochastique.

D.1 Redondance physique et analytique

La redondance physique consiste à utiliser des capteurs et actionneurs multiples pour détecter et localiser les défauts. Si ces composantes identiques placées dans le même environnement émettent des signaux identiques, on considère que ces composantes sont à l'état sain, et dans le cas contraire on considère qu'une panne s'est produite dans au moins une des composantes. Cette approche est conceptuellement simple, mais est coûteuse à mettre en œuvre. La méthode de redondance analytique [40] permet de pallier à ce dernier inconvé-

nient et d'éviter la multiplication de capteurs. Cette méthode est caractérisée par des relations dynamiques entre les sorties des capteurs et les entrées des actionnaires. Elle consiste à faire intervenir des termes additifs ou multiplicatifs dans le modèle dynamique décrivant le fonctionnement sain du système surveillé. Dans le paragraphe suivant, nous présentons la méthode de surveillance basée sur des modèles entrée-sortie.

Le problème de la surveillance et de diagnostic de pannes additives des systèmes linéaires s'écrit de la façon suivante [25] :

$$\begin{cases} \dot{x}(t) &= Ax(t) + Bu(t) + R_1 f(t) \\ y(t) &= Cx(t) + Du(t) + R_2 f(t) \end{cases} \quad (\text{D.1})$$

où $x(t) \in \mathbb{R}^n$, $u(t) \in \mathbb{R}^r$ et $y(t) \in \mathbb{R}^m$ sont respectivement l'état, l'entrée et la sortie du système, sujet à des pannes additives de la forme $R_1 f(t)$ et $R_2 f(t)$. $f(t)$ est le vecteur de défaut (égale à zéro en l'absence de pannes). Les matrices A, B, C, D, R_1 et R_2 sont supposées connues.

Il existe deux formes de la redondance analytique : la redondance statique et la redondance temporelle.

Redondance statique : La redondance statique ou directe est caractérisée par l'utilisation de capteurs de types distincts mais redondants via la relation statique :

$$y(k) = Cx(k) + f(k)$$

où $x(k) \in \mathbb{R}^n$ représente les n variables d'états, $y(k) \in \mathbb{R}^m$ la sortie des m capteurs et $f(k) \in \mathbb{R}^m$ le vecteur de pannes avec $m > n$, *i.e.* les capteurs sont redondants (mais différents). C est une matrice $m \times n$ de rang colonne plein. On génère le résidu suivant appelé *vecteur de parité* [40] :

$$r(k) = Vy(k)$$

où la matrice $V \in \mathbb{R}^{(m-n) \times m}$ (de rang ligne plein) vérifie :

$$VC = 0$$

$$V^T V = I_m - C(C^T C)^{-1} C^T$$

et

$$V^T V = I_{m-n}$$

on a donc,

$$r(k) = Vf(k). \quad (\text{D.2})$$

L'équation (D.2) révèle que le résidu (le vecteur de parité) dépend seulement des erreurs liées aux défauts, indépendamment de $x(k)$ qui n'est pas directement mesurable. On détecte une panne ($f(k) \neq 0$), lorsque le résidu $r(k)$ est non nul. Et afin d'éviter de fausses alarmes, il est possible d'introduire un seuil plus grand que zéro.

Redondance temporelle : Le concept présenté dans le cas de la redondance statique a été généralisé pour les systèmes dynamiques. Cette technique est décrite dans [25]. On suppose que le système à étudier est donné par les équations d'états linéaire discret :

$$\begin{cases} x(k+1) &= Ax(k) + Bu(k) + R_1 f(k) \\ y(k) &= Cx(k) + Du(k) + R_2 f(k) \end{cases} \quad (\text{D.3})$$

Les relations de redondance sont mathématiquement spécifiées de la manière suivante : regroupons les équations (D.3) de l'instant $k - s + 1$ à l'instant k , on obtient :

$$Y(k) - \mathcal{T}_s(A, B, C, D)U(k) = \mathcal{O}_s x(k - s + 1) + \mathcal{T}_s(A, R_1, C, R_2)F(k)$$

avec :

$$Y(k) = \begin{pmatrix} y(k - s + 1) \\ \dots \\ y(k) \end{pmatrix}, \quad U(k) = \begin{pmatrix} u(k - s + 1) \\ \dots \\ u(k) \end{pmatrix}, \quad F(k) = \begin{pmatrix} f(k - s + 1) \\ \dots \\ f(k) \end{pmatrix}$$

et

$$\mathcal{T}_s(A, B, C, D) = \begin{pmatrix} D & 0 & \dots & 0 \\ CB & D & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ CA^{s-2}B & CA^{s-3}B & \dots & D \end{pmatrix}, \quad \mathcal{O}_s(C, A) = \begin{pmatrix} C \\ CA \\ \vdots \\ CA^{s-1} \end{pmatrix}.$$

où $U \in \mathbb{R}^{sr}$, $Y, F \in \mathbb{R}^{sm}$, $\mathcal{T}_s(A, B, C, D) \in \mathbb{R}^{sm \times sr}$ et $\mathcal{O}_s \in \mathbb{R}^{sm \times n}$.

On génère le résidu suivant :

$$r(k) = V(Y(k) - \mathcal{T}_s(A, B, C, D)U(k))$$

où $V \in \mathbb{R}^{p \times sm}$, (p est la dimension du vecteur résiduel) qui vérifie :

$$V\mathcal{O}_s(C, A) = 0, \quad V\mathcal{T}_s(A, R_1, C, R_2) \neq 0$$

on obtient donc :

$$r(k) = V\mathcal{T}_s(A, R_1, C, R_2)F(k). \quad (\text{D.4})$$

D'après la relation (D.4), le résidu est sensible uniquement aux pannes. Autrement dit, lorsque le vecteur des pannes $F(k)$ est égale à zéro le résidu $r(k)$ devient nul, et en pratique on détecte une panne par franchissement de seuil $r(k)^T r(k)$.

D.2 Approche à base d'observateur

Un observateur est un algorithme qui permet d'estimer l'état d'un système à partir des signaux d'entrée et de la sortie. L'idée de l'approche à base d'observateur est de reconstruire la sortie du système surveillé à partir des mesures à l'aide des observateurs ou des filtres de Kalman. L'erreur (pondérée) d'estimation de la sortie est utilisé comme résidu pour la détection et l'isolation des défauts [40].

Approche à base d'observateur pour des systèmes linéaires

L'observateur de Luenberger est généralement utilisé pour les systèmes linéaires, il s'écrit pour le système (D.1), en l'absence des défauts, sous la forme [25] :

$$\begin{cases} \dot{\hat{x}}(t) &= A\hat{x}(t) + Bu(t) + K(y(t) - \hat{y}(t)) \\ \hat{y}(t) &= C\hat{x}(t) + Du(t) \end{cases} \quad (\text{D.5})$$

où $\hat{x}(t)$ et $\hat{y}(t)$ sont respectivement les estimations de l'état $x(t)$ et de la sortie $y(t)$ du système (D.1). K est la matrice du gain. Le résidu du système (D.1) obtenu en utilisant l'observateur (D.5) est décrit par :

$$\begin{cases} \dot{\hat{x}}(t) &= (A - KC)\hat{x}(t) + (B - KD)u(t) + Ky(t) \\ \hat{y}(t) &= C\hat{x}(t) + Du(t) \\ r(t) &= Q(y(t) - \hat{y}(t)) \end{cases} \quad (\text{D.6})$$

où $Q \in \mathbb{R}^{p \times m}$ est une matrice de pondération. L'erreur de l'estimation d'état $e(t) = x(t) - \hat{x}(t)$, et le résidu $r(t)$ sont donc gouvernés par les équations suivantes :

$$\begin{cases} \dot{e}(t) &= (A - KC)e(t) + (R_1 - KR_2)f(t) \\ r(t) &= QCe(t) + QR_2f(t). \end{cases} \quad (\text{D.7})$$

La matrice de gain K doit être choisie de façon à ce que l'erreur sur l'état $e(t)$, en l'absence des pannes ($f(t) = 0$), converge asymptotiquement vers 0 :

$$\lim_{t \rightarrow \infty} e(t) = 0.$$

Si des pannes apparaissent soudainement ($f(t) \neq 0$), l'erreur d'estimation d'état $e(t)$ ainsi que le résidu $r(t)$ dévieront de zéro. Le résidu est donc sensible uniquement aux pannes qui sont généralement détectées par franchissement de seuil $r(t)^T r(t)$.

Approche à base d'observateur pour des systèmes non linéaires

Dans cette sous-section nous introduisons le problème de diagnostic des pannes dans un système dynamique non linéaire en utilisant un observateur non linéaire d'entrée inconnue [67]. Cet observateur sera utilisé pour générer des résidus qui détectent les écarts des paramètres physiques du modèle étudié par rapport aux valeurs nominales. Les résidus obtenus sont sensibles uniquement à un certain nombre des paramètres physiques, tout en étant robustes vis-à-vis d'incertitudes sur les autres paramètres du modèle.

Modèle non linéaire académique. On considère le modèle non linéaire à entrée affine :

$$\begin{cases} \dot{x}(t) &= A(x(t), \theta) + B(x(t), \theta)u(t) \\ y(t) &= C(x(t), \theta) \end{cases} \quad (\text{D.8})$$

où $x(t) \in \mathbb{R}^n$ est le vecteur d'état, $u(t) \in \mathbb{R}^m$ et $y(t) \in \mathbb{R}^p$ sont respectivement l'entrée et la sortie du système, $\theta \in \mathbb{R}^q$ est le vecteur des paramètres physiques du modèle. Ce modèle est défini par les applications non linéaires $A : \mathbb{R}^{n \times q} \rightarrow \mathbb{R}^n$, $B : \mathbb{R}^{n \times q} \rightarrow \mathbb{R}^{n \times m}$ et $C : \mathbb{R}^{n \times q} \rightarrow \mathbb{R}^p$. D'autre part, le vecteur des paramètres physiques θ est divisé en deux parties disjointes θ_1 et θ_2 , où θ_1 représente les paramètres associés à l'incertitude tandis que θ_2 reflète les défauts dans les composants et les capteurs du système. On note par $\theta_0 = (\theta_{01}, \theta_{02})$, où θ_{01} et θ_{02} sont les valeurs nominales (supposés connues) de θ_1 et θ_2 respectivement.

Développons en série de Taylor au premier ordre les équations du système (D.8), au voisinage des paramètres nominaux θ_{01} et θ_{02} , on obtient :

$$\left\{ \begin{array}{l} \dot{x}(t) = A(x, \theta_0) + B(x, \theta_0)u + \frac{\partial A(x, \theta_0)}{\partial \theta_1}(\theta_1 - \theta_{01}) + \frac{\partial A(x, \theta_0)}{\partial \theta_2}(\theta_2 - \theta_{02}) \\ \quad + \frac{\partial (B(x, \theta_0)u)}{\partial \theta_1}(\theta_1 - \theta_{01}) + \frac{\partial (B(x, \theta_0)u)}{\partial \theta_2}(\theta_2 - \theta_{02}) \\ y(t) = C(x, \theta_0) + \frac{\partial C(x, \theta_0)}{\partial \theta_1}(\theta_1 - \theta_{01}) + \frac{\partial C(x, \theta_0)}{\partial \theta_2}(\theta_2 - \theta_{02}). \end{array} \right. \quad (D.9)$$

On peut réécrire le système (D.9) sous la forme :

$$\left\{ \begin{array}{l} \dot{x}(t) = A(x, \theta_0) + B(x, \theta_0)u + E_1(x, u)d + K_1(x, u)f \\ y(t) = C(x, \theta_0) + E_2(x)d + K_2(x)f \end{array} \right. \quad (D.10)$$

où

$$\begin{aligned} E_1(x, u) &= \frac{\partial A(x, \theta_0)}{\partial \theta_1} + \frac{\partial (B(x, \theta_0)u)}{\partial \theta_1}, \\ K_1(x, u) &= \frac{\partial A(x, \theta_0)}{\partial \theta_2} + \frac{\partial (B(x, \theta_0)u)}{\partial \theta_2}, \\ E_2(x) &= \frac{\partial C(x, \theta_0)}{\partial \theta_1}, \quad K_2(x) = \frac{\partial C(x, \theta_0)}{\partial \theta_2}, \end{aligned}$$

et $d \doteq \theta_1 - \theta_{01}$ et $f \doteq \theta_2 - \theta_{02}$ sont respectivement les écarts inconnus de θ_1 et de θ_2 par rapport à leurs valeurs nominales θ_{01} et θ_{02} . La modification du paramètre par rapport à sa valeur nominale désigne, en général, un vieillissement ou bien des pannes susceptibles d'affecter le système.

Rejet de perturbation : Afin de découpler les perturbations associées au vecteur d et les défauts représentés par le vecteur f dans (D.10), on utilise une transformation non linéaire :

$$z = T(x, u) \quad (D.11)$$

La dérivation par rapport au temps de la transformation (D.11) conduit à :

$$\dot{z} = \frac{\partial T(x, u)}{\partial x} \dot{x} + \frac{\partial T(x, u)}{\partial u} \dot{u} \quad (D.12)$$

Substituons (D.10) dans (D.12), on obtient :

$$\left\{ \begin{array}{l} \dot{z}(t) = \frac{\partial T(x, u)}{\partial x} \left(A(x, \theta_0) + B(x, \theta_0)u + E_1(x, u)d + K_1(x, u)f \right) + \frac{\partial T(x, u)}{\partial u} \dot{u} \\ y(t) = C(x, \theta_0) + E_2(x)d + K_2(x)f \end{array} \right. \quad (D.13)$$

On peut donc éliminer les perturbations associées au vecteur d en introduisant la condition :

$$\frac{\partial T(x, u)}{\partial x} E_1(x, u) = 0, \quad \forall x, u \quad (\text{D.14})$$

Cette dernière condition implique que la dynamique du système suivant les coordonnées de z est complètement découplée des perturbations inconnues d . D'autre part, pour que la transformation (D.11) reflète les défauts f dans le modèle transformé il est nécessaire d'introduire la condition suivante [67] :

$$\text{rang} \left(\frac{\partial T(x, u)}{\partial x} K_1(x, u) \right) = \text{rang} (K_1(x, u)), \quad \forall x, u. \quad (\text{D.15})$$

Cette condition assure que les défauts qui affectent la dynamique du système d'origine (D.10) affectent aussi la dynamique donnée par (D.13) et qui est décrite suivant la transformation en z [67].

Construction d'un observateur. Il existe plusieurs approches pour construire un observateur de détection des défauts pour le système (D.13), ces approches sont détaillées dans [67]. Dans ce paragraphe nous présentons l'approche nommée *Linéarisation de l'erreur d'estimation*.

Dans le paragraphe précédent, nous avons vu que lorsqu'on applique la transformation $z = T(x, u)$ avec les deux conditions (D.14) et (D.15) sur le système (D.10), nous obtenons :

$$\begin{cases} \dot{z} = \frac{\partial T(x, u)}{\partial x} (A(x, \theta_0) + B(x, \theta_0)u) + \frac{\partial T(x, u)}{\partial u} \dot{u} + \frac{\partial T(x, u)}{\partial x} K_1(x, u)f, \\ y = C(x, \theta_0) + E_2(x)d + K_2(x)f, \end{cases} \quad (\text{D.16})$$

On suppose maintenant que le système (D.16) s'écrit sous la forme :

$$\begin{cases} \dot{z} = Fz + \Phi(\bar{y}, u, \dot{u}) + \frac{\partial T(x, u)}{\partial x} K_1(x, u)f, \\ \bar{y} = \rho(y), \end{cases} \quad (\text{D.17})$$

où les valeurs propres de la matrice F sont négatives, Φ est une application non linéaire. $\bar{y}$ est une transformation de la sortie y et elle vérifie les deux conditions :

$$\frac{\partial \rho(y)}{\partial d} = 0, \quad \text{rang} \left(\frac{\partial \rho(y)}{\partial f} \right) = \text{rang}(K_2(x)), \quad (\text{D.18})$$

Ces deux conditions impliquent que la sortie $\bar{y}$ est indépendante des perturbations d mais elle est sensible aux défauts qui sont liés à f . En plus, s'il existe une relation entre z et $\bar{y}$ telle que :

$$\Psi(T(x, u), \bar{y}, u) = 0 \text{ ssi } f = 0 \quad (\text{D.19})$$

le système (D.17) avec les conditions (D.18) et (D.19) nous permet de définir un observateur non linéaire qui s'écrit sous la forme :

$$\begin{cases} \dot{\hat{z}} = F\hat{z} + \Phi(\bar{y}, u, \dot{u}) \\ r = \Psi(\hat{z}, \bar{y}, u) \end{cases} \quad (\text{D.20})$$

L'erreur d'estimation $e = \hat{z} - z$ et le résidu r sont donc gouvernés par :

$$\begin{cases} \dot{e} &= Fe - \frac{\partial T(x, u)}{\partial x} K_1(x, u) f \\ r &= \Psi(T(x, u) + e, \bar{y}, u). \end{cases} \quad (\text{D.21})$$

En absence des écarts dans les paramètres physiques θ_2 par rapport aux valeurs nominales ($f = 0$), le résidu r tend vers zéro asymptotiquement et indépendamment des éventuelles perturbations des paramètres θ_1 .

D.3 Approche fréquentielle

La matrice de transfert du système (D.1) s'écrit sous la forme :

$$y(s) = G_u(s)u(s) + G_f(s)f(s) \quad (\text{D.22})$$

où

$$\begin{aligned} G_u(s) &= C(sI - A)^{-1}B + D \\ G_f(s) &= C(sI - A)^{-1}R_1 + R_2 \end{aligned}$$

La génération de résidus peut être également synthétisé dans le domaine fréquentiel via la factorisation de la matrice de transfert $G_u(s)$ [25]. Cette approche est basée sur le fait que si la matrice de fonctions de transfert $G_u(s)$ est rationnelle et propre¹, on peut la factoriser comme :

$$G_u(s) = \tilde{M}^{-1}(s)\tilde{N}(s)$$

où $\tilde{M}(s) \in \mathbb{R}^{m \times m}$ et $\tilde{N}(s) \in \mathbb{R}^{m \times r}$ sont deux matrices de transfert rationnelles et stables, données par :

$$\begin{aligned} \tilde{M}(s) &= -C(sI - A + KC)^{-1}L + I \\ \tilde{N}(s) &= C(sI - A + KC)^{-1}(B - KD) + D \end{aligned}$$

où la matrice de gains K assure la stabilité de la matrice $A - KC$. En utilisant cette factorisation de la matrice $G_u(s)$, on peut générer le résidu suivant [34] :

$$r(s) = Q(s) \left(\tilde{M}(s)y(s) - \tilde{N}(s)u(s) \right) \quad (\text{D.23})$$

où $Q \in \mathbb{R}^{p \times m}$ est une matrice de pondération. Remplaçons (D.22) dans (D.23), on obtient :

$$r(s) = Q(s)\tilde{M}(s)G_f(s)f(s).$$

Le résidu est nul en absence des pannes ($f(s) = 0$). La matrice Q , au sens physique, est considérée comme un *post-filtre* qui génère des degrés de liberté additionnels pour isoler les défauts ou bien générer des résidus robustes [41].

1. Une fonction de transfert est dite propre si le degré de son numérateur est inférieur à celui de son dénominateur.

D.4 Méthode locale pour la détection et la diagnostic de changement

La méthode locale [11, 95] consiste à caractériser le système surveillé par un modèle paramétré par un vecteur θ , dont la valeur nominale θ_0 a été préalablement identifiée en fonctionnement sain (sans défauts) sur des données réelles. La modification du paramètre par rapport à sa valeur nominale désigne, en général, un vieillissement ou bien des pannes susceptibles d'affecter le système. En général, le modèle dont on dispose s'écrit, après réduction, sous la forme d'un modèle d'état non linéaire :

$$\frac{dX(t, \theta)}{dt} = f(X(t), u(t), \theta) \quad (\text{D.24})$$

$$Y(t, \theta) = g(X(t), u(t), \theta) \quad (\text{D.25})$$

où $X(t, \theta)$ est le vecteur d'état, $u(t)$ et $Y(t, \theta)$ sont respectivement l'entrée et la sortie du modèle. f et g sont deux fonctions non linéaires, et θ est le vecteur des paramètres du modèle. Compte-tenu du fait que les mesures d'entrée et de sortie dont on dispose sont échantillonnées, l'estimation du paramètre θ est généralement réalisée par la minimisation d'un critère que nous notons :

$$J_N(\theta) = \frac{1}{N} \sum_{k=1}^N \|Z_k - Y(k, \theta)\|^2 \quad (\text{D.26})$$

où la norme $\|\cdot\|$ est typiquement euclidienne, Z_k et $Y(k, \theta)$ sont respectivement les mesures prises sur l'installation et la sortie du modèle mathématique pour un échantillon de taille N . On note par θ_0 le paramètre optimale qui minimise l'erreur (D.26), c.-à-d. celui qui annule la dérivée par rapport à θ de la quantité (D.26) donnée par :

$$\frac{\partial J_N(\theta)}{\partial \theta} = -\frac{2}{N} \sum_{k=1}^N \left(\frac{\partial Y(k, \theta)}{\partial \theta} \right)^T (Z_k - Y(k, \theta)). \quad (\text{D.27})$$

Le résidu primaire noté par $H(\theta, Z_k)$ est défini par :

$$H(\theta, Z_k) = \left(\frac{\partial Y(k, \theta)}{\partial \theta} \right)^T (Z_k - Y(k, \theta)). \quad (\text{D.28})$$

Par définition le *résidu secondaire* est la somme cumulée normalisée du résidu primaire, il est défini pour un échantillon de données de taille N par :

$$\xi_N(\theta) \doteq \frac{1}{\sqrt{N}} \sum_{k=1}^N H(\theta, Z_k)$$

Le résidu secondaire reflète ainsi les changements dans le vecteur paramètre θ . En effet, $\xi_N(\theta)$ est nulle si $\theta = \theta_0$ et non nulle si $\theta \neq \theta_0$. Notons que le choix de $H(\theta, Z_k)$ dépend de la distance choisie entre les données et le modèle, par exemple la distance quadratique.

Dans le paragraphe suivant, la surveillance et la détection des changements se fait à l'aide d'une statistique appropriée grâce à la méthode locale [11, 95]. Cette méthode est basée sur une approche statistique du traitement des signaux.

Validation du modèle

Une fois la signature nominale θ_0 acquise, et à partir d'un nouvel échantillon de données de taille N fixée, la détection de changement dans les paramètres θ est ramenée à un test d'hypothèses statistiques, qui s'écrit sous la forme :

$$\mathcal{H}_0 : \theta = \theta_0, \text{ pour } 0 \leq k \leq N$$

$$\mathcal{H}_1 : \theta = \theta_0 + \frac{\tilde{\theta}}{\sqrt{N}}, \text{ pour } 0 \leq k \leq N$$

où $\tilde{\theta} \neq 0$ est le changement inconnu.

Détection de changement

Le problème de détection de changement dans le paramètre θ revient à vérifier si avec les nouvelles mesures Z_k , la valeur de $\xi_N(\theta_0)$ est significativement éloignée de zéro ou non. D'après [11] et [95], le théorème de la limite centrale montre que $\xi_N(\theta_0)$ a le comportement asymptotique suivant :

$$\xi_N(\theta_0) \sim \begin{cases} \mathcal{N}(0, R(\theta_0)) & \text{sous } \mathcal{H}_0 \\ \mathcal{N}(M(\theta_0) \cdot \tilde{\theta}, R(\theta_0)) & \text{sous } \mathcal{H}_1 \end{cases}$$

où $M(\theta_0)$ est la matrice d'incidence définie par :

$$M(\theta_0) \doteq -\mathbb{E}_{\theta_0} \frac{\partial}{\partial \theta} H(\theta, Z_k) \Big|_{\theta=\theta_0}$$

et $\xi_N(\theta_0)$ est la matrice de covariance limite, définie par :

$$R(\theta_0) \doteq \lim_{N \rightarrow \infty} \mathbb{E}_{\theta_0} \xi_N(\theta_0) \xi_N^T(\theta_0)$$

où $\mathcal{N}$ est la loi normale ou gaussienne, et $\mathbb{E}_{\theta}$ est l'espérance lorsque la vraie valeur du paramètre est θ . Le problème de changement de paramètre θ se ramène donc à un problème de changement dans la moyenne $M(\theta_0)$ du processus gaussien.

La détection de changement repose sur la maximisation du logarithme du rapport de vraisemblance (*Likelihood ratio*) :

$$S_{\tilde{\theta}} = \max_{\tilde{\theta}} \sum_{k=1}^N \ln \frac{P_{\tilde{\theta}}(Z_k)}{P_0(Z_k)} \quad (\text{D.29})$$

où $P_{\tilde{\theta}}$ et P_0 sont respectivement les densités de probabilité des gaussiennes $\mathcal{N}(M(\theta_0) \cdot \tilde{\theta}, R(\theta_0))$ et $\mathcal{N}(0, R(\theta_0))$. Si $S_{\tilde{\theta}}$ est supérieure à un certain seuil, on décide qu'il y a eu changement dans la moyenne. On cherche donc à maximiser le rapport de vraisemblance (D.29), et on note :

$$\tilde{\theta}^* = \arg \max_{\tilde{\theta}} S_{\tilde{\theta}},$$

dont la valeur est obtenue en résolvant l'équation qui annule la dérivée de $S_{\tilde{\theta}}$ par rapport à $\tilde{\theta}$.

On définit donc le test, $t \doteq 2S_{\tilde{\theta}^*}$ qui est une variable aléatoire suivant une loi de χ^2 , de nombre de degrés de liberté égale à la dimension de $\tilde{\theta}$. Ce test est centré sous l'hypothèse $\mathcal{H}_0$ et a pour paramètre de non centralité sous $\mathcal{H}_1$:

$$\gamma = \tilde{\theta}^T M^T R^{-1} M \tilde{\theta}.$$

L'alarme sera déclenchée lorsque t dépassera un seuil fixé à l'avance, et il est possible de régler le seuil en fonction de la probabilité de la fausse alarme que l'on souhaite obtenir.

Diagnostic

Une fois que l'alarme est déclenchée, il reste à isoler son origine en identifiant les paramètres ayant subi de changement. On partitionne $\tilde{\theta}$ de la façon suivante :

$$\tilde{\theta} = \begin{pmatrix} \tilde{\theta}_a \\ \tilde{\theta}_b \end{pmatrix} \in \mathbb{R}^q,$$

où $\tilde{\theta}_a \in \mathbb{R}^{q_a}$, contient les composantes à surveiller et $\tilde{\theta}_b \in \mathbb{R}^{q_b}$, contient les paramètres de nuisance inconnus, dont on ignore leurs variations. La procédure de diagnostic consiste à tester un éventuel écart du vecteur paramètre d'intérêt $\tilde{\theta}_a$ par rapport à 0. À cette fin, deux types de test peuvent être utilisés : le test de sensibilité et le test de min-max [11, 8, 96].

Test de sensibilité : La méthode par sensibilité consiste à traiter uniquement les paramètres que l'on souhaite surveiller ($\tilde{\theta}_a$ dans notre cas) et à considérer les autres comme fixes. Cela revient à faire une projection du vecteur paramètre sur l'espace des paramètres que l'on désire surveiller [96].

Bibliographie

- [1] R. Aris. *Introduction to the Analysis of Chemical Reactors*. Prentice-Hall, 1st edition, 1965.
- [2] A. Bard and L. Faulkner. *Electrochemical Methods : Fundamentals and Applications*. John Wiley & Sons, 2nd edition, 2001.
- [3] E. Barsoukov and J. Macdonald. *Impedance Spectroscopy : Theory, Experiment, and Applications*. Wiley-Interscience, 2nd edition, 2005.
- [4] J.J. Baschuk and X. Li. Modelling of polymer electrolyte membrane fuel cells with variable degrees of water flooding. *Journal of Power Sources*, 86 :181–196, 2000.
- [5] J.J. Baschuk and X. Li. Carbon monoxide poisoning of proton exchange membrane fuel cells. *Int. J. Energy Res.*, 25 :695–713, 2001.
- [6] J.J. Baschuk and X. Li. Modelling CO poisoning and O₂ bleeding in a PEM fuel cell anode. *Int. J. Energy Res.*, 27 :1095–1116, 2003.
- [7] J.J. Baschuk, A.M. Rowe, and X. Li. Modelling and simulation of pem fuel cell with co poisoning. *Journal of energy resources technology*, 125(2) :94–100, 2003.
- [8] M. Basseville. Information criteria for residual generation and fault detection and isolation. *Automatica*, 33(5) :783–803, 1997.
- [9] M.Z. Bazant, M.S. Kilic, B.D. Storey, and A. Ajdari. Towards an understanding of induced-charge electrokinetics at large applied voltages in concentrated solutions. *Advances in Colloid and Interface Science*, 152(1-2) :48–88, 2009.
- [10] M.Z. Bazant, K. Thornton, and A. Ajdari. Diffuse-charge dynamics in electrochemical systems. *Phys. Rev. E*, 70 :021506, Aug 2004.
- [11] A. Benveniste, M. Basseville, and G. Moustakides. The asymptotic local approach to change detection and model validation. *Automatic Control, IEEE Transactions*, 32(7) :583–592, 1987.
- [12] D. Bernadi and M. Verbrugge. Mathematical model of a Gas Diffusion Electrode Bounded to a polymer electrolyte. *AIChE*, 37 :1151–1163, 1991.
- [13] D. Bernadi and M. Verbrugge. A mathematical model of the solid polymer electrolyte fuel cell. *Journal of the Electrochemical Society*, 139 :2477–2491, 1992.
- [14] J. Besson. *Precis de Thermodynamique et Cinétique Électrochimique*. Ellipses, 1984.
- [15] K.K. Bhatia and C-Y. Wang. Transient carbon monoxide poisoning of a polymer electrolyte fuel cell operating on diluted hydrogen feed. *J. Electrochimica Acta*, 49 :2333–2341, 2004.
- [16] P.M. Biesheuvel, A.A. Franco, and M.Z. Bazant. Diffuse charge effects in fuel cell membranes. *Journal of The Electrochemical Society*, 156(2) :B225–B233, 2009.
- [17] P.M. Biesheuvel, Y. Fu, and M.Z. Bazant. Diffuse charge and faradaic reactions in porous electrodes. *Phys. Rev. E*, 83 :061507, Jun 2011.

- [18] J. O'M. Bockris, M.A.V. Devanathan, and K. Müller. On the structure of charged interfaces. *Proceedings of the Royal Society of London. Series A, Mathematical and Physical Sciences*, 274(1356), 1963.
- [19] JOM Bockris, AKN Reddy, and ME Gamboa-Aldeco. *Modern Electrochemistry, 2A : Fundamentals of Electrodics*. Kluwer New York, 2002.
- [20] E. Bolinder. A note on the matrix representation of linear two-port networks. *Circuit Theory, IRE Transactions on*, 4(4) :337–339, december 1957.
- [21] E. Bolinder. Survey of some properties of linear networks. *Circuit Theory, IRE Transactions on*, 4(3) :70–78, sep 1957.
- [22] L. Boulon, K. Agbossou, D. Hissel, P. Sicard, A. Bouscayrol, and M.-C. Piéra. A macroscopic pem fuel cell model including water phenomena for vehicle simulation. *Renewable Energy*, (0), 2012.
- [23] A. Bouscayrol. *Formalisme de représentation et de commande appliqués aux systèmes électromécaniques multimachines multiconvertisseurs*. HDR, Université de sciences et technologies de Lille, Lille, 2003.
- [24] J-M. Le Canut, R. Abouatallah, and D. Harrington. Detection of Membrane Drying, Fuel Cell Flooding, and Anode Catalyst Poisoning on PEMFC Stacks by Electrochemical Impedance Spectroscopy. *Journal of The Electrochemical Society*, 153(5) :A857–A864, 2006.
- [25] J. Chen and R.J. Patton. *Robust Model-Based Fault Diagnosis For Dynamic Systems*. Kluwer, 1999.
- [26] E.-S.J. Chia, J.B. Benziger, and I.G. Kevrekidis. STR-PEM Fuel Cell as a Reactor Building Block. *AIChE Journal*, 52(11) :3903–3910, 2006.
- [27] D. Chrenko, M-C Péra, D. Hissel, and A. Bouscayrol. Inversion-based control of a proton exchange membrane fuel cell system using energetic macroscopic representation. *Journal of Fuel Cell Science and Technology*, 6(2) :024501, 2009.
- [28] D. Chrenko, M-C Péra, D. Hissel, and M. Geweke. Macroscopic modeling of a PEFC system based on equivalent circuits of fuel and oxidant supply. *ASME Journal of Fuel Cell Science and Technology*, 5(4-7), 2008.
- [29] D. Chrenko, M-C Péra, and Daniel Hissel. Fuel cell system modeling and control with energetic macroscopic representation. *IEEE ISIE, Vigo, Spain*, (4-7), 2007.
- [30] H.S. Chu, C.P. Wang, W.C. Liao, and W.M. Yan. Transient behavior of co poisoning of the anode catalyst layer of a pem fuel cell. *Journal of Power Sources*, 159(2) :1071 – 1077, 2006.
- [31] C. Cussenot. *Surveillance et diagnostic de la chaîne de pollution d'une automobile*. Phd thesis, Université de Rennes 1, 1998.
- [32] A. Damjanovic and V. Brusic. Electrode kinetics of oxygen reduction on oxide-free platinum electrodes. *Electrochimica Acta*, 12(6) :615 – 628, 1967.
- [33] J-P. Diard, B. Le Gorrec, and C. Montella. *Cinétique électrochimique*. Hermann, 1996.
- [34] X. Ding and P-M. Frank. Fault detection via factorization approach. *Systems & Control Letters*, 14(5) :431 – 436, 1990.
- [35] P. Érdi and J. Tóth. *Mathematical models of chemical reaction*. Prentice-Hall, 1st edition, 1989.
- [36] T. Escobet, D. Feroldi, S. de Lira, V. Puig, J. Quevedo, J. Riera, and M. Serra. Model-based fault diagnosis in pem fuel cell systems. *Journal of Power Sources*, 192(1) :216 – 223, 2009.

- [37] A. Franco. *Un modèle physique multiéchelle de la dynamique dans une pile à combustible à électrolyte polymère*. Thèse de doctorat, Université Claude Bernard, Lyon, 2005.
- [38] A. Franco, P. Schott, C. Jallut, and B. Maschke. A dynamic mechanistic model of an electrochemical interface. *Journal of the Electrochemical Society*, 153(6) :A1053, 2006.
- [39] A. Franco, P. Schott, C. Jallut, and B. Maschke. A Multi-Scale Dynamic Mechanistic Model for the Transient Analysis of PEFCs. *Fuel Cells*, 7 :99–117, 2007.
- [40] P. Frank. Fault Diagnosis in Dynamic Systems Using Analytical and Knowledge-based Redundancy. A Survey and Some New Results. *Automatica*, 26(3) :459–474, 1990.
- [41] P.M. Frank and X. Ding. Survey of robust residual generation and evaluation methods in observer-based fault detection systems. *Journal of Process Control*, 7(6) :403 – 424, 1997.
- [42] D.A. Frickey. Conversions between S, Z, Y, H, ABCD, and T parameters which are valid for complex source and load impedances. *Microwave Theory and Techniques, IEEE Transactions on*, 42(2) :205–211, feb 1994.
- [43] A. Gelb and W.E. Vander Velde. *Multiple-Input Describing Functions and Nonlinear System Design*. McGraw Hill, New-York, 2nd edition, 1968.
- [44] H. Gohr. Impedance modelling of porous electrodes. *Electrochemical Applications*, 1, 1997.
- [45] J. Golbert and D.R. Lewin. Model-based control of fuel cells : (1) regulatory control. *Journal of Power Sources*, 135(1-2) :135 – 151, 2004.
- [46] D.A. Harrington and B.E. Conway. Kinetic theory of the open-circuit potential decay method for evaluation of behaviour of adsorbed intermediates : Analysis for the case of the H_2 evolution reaction. *Journal of Electroanalytical Chemistry and Interfacial Electrochemistry*, 221(1-2) :1–21, 1987.
- [47] J-P. Hautier and J-P. Barre. The causal ordering graph - a tool for modelling and control law synthesis. *Studies in Informatics and Control Journal*, 13(4) :265 – 283, 2004.
- [48] J.P. Hautier and J. Faucher. Le graphe informationnel causal. *Bulletin de l'Union des Physiciens*, (90) :167 – 189, 1996.
- [49] D. Hissel, M-C Péra, A. Bouscayrol, and D.Chrenko. Représentation énergétique macroscopique d'une pile à combustible. *Revue internationale de génie électrique*, 11(4-5) :603 – 623, 2008.
- [50] D. Hissel, M.C Péra, and J.M Kauffmann. Diagnosis of automotive fuel cell power generators. *Journal of Power Sources*, 128(2) :239–246, 2004.
- [51] S. Jemei. *Modélisation d'une pile à combustible de type PEM par réseaux de neurones*. Thèse de doctorat, Université de technologie de Belfort-Montbéliard, Belfort, 2004.
- [52] A.M. Jhonson and J. Newman. . *Journal of the Electrochemical Society*, 118 :510, 1971.
- [53] A.Y. Karnik, A.G Stefanopoulou, and J.Sun. Water equilibria and management using a two-volume model of a polymer electrolyte fuel cell. *Journal of Power Sources*, 164 :590–605, 2007.
- [54] D. Kondepudi and I. Prigogine. *Modern Thermodynamics*. John Wiley and Sons, 1998.
- [55] A.A. Kulikovsky. The voltage-current curve of a polymer electrolyte fuel cell : exact and fitting equations. *Electrochemistry Communications*, 4(11) :845 – 852, 2002.
- [56] J. Larminie and A. Dicks. *Fuel Cell Systems Explained*. John-Wiley & Sons, 2nd edition, 2003.

- [57] Qingfeng Li, Ronghuan He, Jens Oluf Jensen, and Niels J. Bjerrum. Approaches and recent development of polymer electrolyte membranes for fuel cells operating above 100 Åřc. *Chemistry of Materials*, 15(26) :4896–4915, 2003.
- [58] S. Litster and G. McLean. Pem fuel cell electrodes. *Journal of Power Sources*, 130(1Ů2) :61 – 76, 2004.
- [59] M. Marchand. *Gestion de l’eau dans les piles à combustibles*. Thèse de doctorat, INPG, Grenoble, 1998.
- [60] R.A. Marcus. Electron transfer reactions in chemistry theory and experiment. *Journal of Electroanalytical Chemistry*, 438(1) :251–259, 1997.
- [61] T.T. Nwoga and J.W. Van Zee. Estimated Rate Constants for CO Poisoning of High-Performance Pt/Ru–C Anodes. *Journal of the Electrochemical Society*, 11 :B122–B126, 2008.
- [62] H.-F. Oetjen, V.M. Schmidt, U. Stimming, and F. Trila. Performance Data of a Proton Exchange Membrane Fuel Cell Using H_2/CO as Fuel Gas. *Journal of The Electrochemical Society*, 143(12) :3838–3842, 1996.
- [63] R. Onanena, L. Oukhellou, D. Candusso, F. Harel, D. Hissel, and P. Aknin. Fuel cells static and dynamic characterizations as tools for the estimation of their ageing time. *International Journal of Hydrogen Energy*, 36(2) :1730 – 1739, 2011.
- [64] G. Oster and A. Perelson. Chemical reaction networks. *Circuits and Systems, IEEE Transactions on*, 21(6) :709 – 721, nov 1974.
- [65] S.K. Park and S.Y. Choe. Dynamic modeling and analysis of a shell-and-tube type gas-to-gas membrane humidifier for pem fuel cell applications. *International Journal of Hydrogen Energy*, 33(9) :2273 – 2282, 2008.
- [66] P.R. Pathapati, X. Xue, and J. Tang. A new dynamic model for predicting transient phenomena in a pem fuel cell system. *Renewable Energy*, 30(1) :1 – 22, 2005.
- [67] R. Patton, P.M Frank, and R.N Clark. *Issues of Fault Diagnosis for Dynamic Systems*. Springer, 2000.
- [68] O. Perrin. *Modélisation et diagnostic de pannes dans des organes de véhicules automobiles à basse consommation*. Phd thesis, Université de Rennes 1, 2003.
- [69] S.P. Philipps and C. Ziegler. Computationally efficient modeling of the dynamic behavior of a portable pem fuel cell stack. *Journal of Power Sources*, 180(2) :309 –321, 2008.
- [70] B.V. Pillay and J. Newman. The influence of side reactions on the performance of electrochemical double-layer capacitors. *Journal of the Electrochemical Society*, 143 :1806–1814, 1996.
- [71] J.T. Pukrushpan, A.G. Stefanopoulou, and H. Peng. Modeling and control for pem fuel cell stack system. In *Proc. American Control Conference the 2002*, volume 4, pages 3117–3122, May 8-10 2002.
- [72] J.T. Pukruspan, A.G. Stefanopoulou, and H. Peng. Control-oriented model and analysis for automotive fuel cell systems. *Journal of Dynamic Systems*, 126 :pp. 14–25, 2004.
- [73] R. Saisset, G. Fontes, C. Turpin, and S. Astier. Bond graph model of a pem fuel cell. *Journal of Power Sources*, 156(1) :100 – 107, 2006.
- [74] W. Schmickler. *Interfacial Electrochemistry*. Oxford University Press, 1996.
- [75] I.A. Schneider, H. Kuhn, A. Wokaun, and G.G. Scherer. Fast Locally Resolved Electrochemical Impedance Spectroscopy in Polymer Electrolyte Fuel Cells. *Journal of The Electrochemical Society*, 152(10) :A2092, 2005.

- [76] I.A. Schneider, H. Kuhn, A. Wokaun, and G.G. Scherer. Study of Water Balance in a Polymer Electrolyte Fuel Cell by Locally Resolved Impedance Spectroscopy. *Journal of The Electrochemical Society*, 152(12) :A2383, 2005.
- [77] P. Schott and P. Baurens. Fuel cell operation characterization using simulation. *Journal of Power Sources*, 156(1) :85 – 91, 2006.
- [78] A.A. Shah, G.-S. Kim, W. Gervais, A. Young, K. Promislow, J. Li, and S. Ye. The effects of water and microstructure on the performance of polymer electrolyte fuel cells. *Journal of Power Sources*, 160(2) :1251 – 1268, 2006.
- [79] A.A. Shah, P.C. Sui, G.-S. Kim, and S. Ye. A transient pemfc model with CO poisoning and mitigation by O₂ bleeding and ru-containing catalyst. *Journal of Power Sources*, 166(1) :1 – 21, 2007.
- [80] Y. Shan, S. Choe, and S-H. Choi. Unsteady 2d pem fuel cell modeling for a stack emphasizing thermal effects. *Journal of Power Sources*, 165(1) :196 – 209, 2007.
- [81] T.E. Springer, T. Rockward, T.A. Zawodwinski, and S. Gottesfeld. Model for polymer electrolyte fuel cell operation on reformed feed. *Journal of the Electrochemical Society*, 148 :A11–A23, 2001.
- [82] T.E. Springer, T.A. Zawodwinski, and S. Gottesfeld. Polymer Electrolyte Fuel Cell Model. *Journal of the Electrochemical Society*, 138 :2334–2342, 1991.
- [83] V. Srinivasan and J. Weidner. Mathematical Modeling of Electrochemical Capacitors. *Journal of the Electrochemical Society*, 146 :1650–1658, 1999.
- [84] N. Steiner, D. Hissel, P. Moçotéguy, and D. Candusso. Non intrusive diagnosis of polymer electrolyte fuel cells by wavelet packet transform. *International Journal of Hydrogen Energy*, 36(1) :740 – 746, 2011.
- [85] A. Van der Schaft, S. Rao, and B. Jayawardhana. On the Mathematical Structure of Balanced Chemical Reaction Networks Governed by Mass Action Kinetics. October 2011.
- [86] M. van Soestbergen. Frumkin-Butler-Volmer theory and mass transfer in electrochemical cells. *Russian journal of electrochemistry*, 48(6) :570–579, 2012.
- [87] M.W. Verbrugge and P. Liu. Microstructural Analysis and Mathematical Modeling of Electric Double-Layer Supercapacitors. *Journal of the Electrochemical Society*, 152 :79–87, 2005.
- [88] C. Vidal and H. Lemarchand. *La réaction créatrice*. Hermann, 1988.
- [89] N Wagner and E Güllow. Change of electrochemical impedance spectra (EIS) with time during CO-poisoning of the Pt-anode in a membrane fuel cell. *Journal of Power Sources*, 127 :341–347, 2004.
- [90] Z. Wang, C-Y. Wang, and K.S. Chen. Two-phase flow and transport in the air cathode of proton exchange membrane fuel cells. *Journal of Power Sources*, 94(1) :40 – 50, 2001.
- [91] A.Z. Weber and J. Newman. Transport in Polymer-Electrolyte Membranes. Model Validation in a Simple Fuel-Cell Model. *Journal of the Electrochemical Society*, 151 :326–339, 2004.
- [92] O. Wing. *Classical circuit theory*. Springer, New-York, 2008.
- [93] N. Yousfi Steiner, D. Candusso, D. Hissel, and P. Moçoteguy. Model-based diagnosis for proton exchange membrane fuel cells. *Mathematics and Computers in Simulation*, 81(2) :158 – 170, 2010.
- [94] N. Yousfi Steiner, D. Hissel, Ph. Moçotéguy, and D. Candusso. Diagnosis of polymer electrolyte fuel cells failure modes (flooding & drying out) by neural networks modeling. *International Journal of Hydrogen Energy*, 36(4) :3067 – 3075, 2011.

- [95] Q. Zhang, M. Basseville, and A. Benveniste. Early warning of slight changes in systems. *Automatica*, 30(30) :95–113, 1994.
- [96] Q. Zhang, M. Basseville, and A. Benveniste. Fault detection and isolation in nonlinear dynamic systems : A combined input-output and local approach. *Automatica*, 34(11) :1359 – 1373, 1998.