

13^e année

La Revue 3E.I



Ressource publiée sur Culture Sciences de l'Ingénieur : <https://eduscol.education.fr/sti/si-ens-paris-saclay>



**Les
matériaux
électroactifs**

Publication trimestrielle du Cercle Thématique 13.01 de la SEE

ENSEIGNER L'ELECTROTECHNIQUE ET L'ÉLECTRONIQUE INDUSTRIELLE



**Société de l'Electricité, de l'Electronique
et des Technologies de l'Information
et de la Communication**

N° 47 - Décembre 2006

rts 15^e édition

EMBEDDED SYSTEMS

Le salon des solutions informatiques temps réel et des systèmes embarqués

OS TEMPS RÉELS ET EMBARQUÉS - OPEN SOURCE - WINDOWS - JAVA - INTERNET, RÉSEAUX - CARTES & SOLUTIONS MATÉRIELLES - WIRELESS



CNIT

Paris, La Défense

6, 7 et 8 mars 2007

100% EMBEDDED !

Aujourd'hui, l'embarqué fait partie intégrante de notre quotidien. Depuis 14 ans, le salon "rts **EMBEDDED SYSTEMS**" rend compte des évolutions technologiques du temps réel et de l'embarqué, et de l'état de l'art en la matière. **L'édition 2007** sera comme chaque année la vitrine de l'offre matérielle et logicielle, le rendez-vous de la profession, rassemblée pour proposer aux visiteurs son lot de nouveautés. Sous la houlette des meilleurs experts, conférences débats, présentations techniques et tables rondes nourriront la partie information de l'événement. **En 2007, "rts EMBEDDED SYSTEMS" fêtera ses 15 années d'existence !**

Pour exposer, visiter l'exposition ou s'inscrire aux conférences : www.birp.com/rts

Salon strictement réservé aux professionnels.

Organisation



97, rue du Cherche-Midi 75006 Paris - France - Tel : +33 (0)1 44 78 99 30 - Fax : +33 (0)1 44 78 99 49 - e-mail : rts@birp.fr



SOCIETE de l'ELECTRICITE, de l'ELECTRONIQUE et des TECHNOLOGIES de l'INFORMATION et de la COMMUNICATION.

17, rue Hamelin, PARIS 75 783 CEDEX 16
Tel : 01 56 90 37 00 Fax : 01 56 90 37 19
site web : www.see.asso.fr

La Revue 3EI
publication trimestrielle
du **Cercle Thématique 13-01**
de la SEE

SEE, association reconnue d'utilité publique par le décret du 7 décembre 1886
Siret 785 393 232 00026, APE 731 Z, n° d'identification FR 44 785 393 232

3EI : Enseigner l'Electrotechnique et l'Electronique industrielle

La Revue 3EI, Édition SEE, 17 rue Hamelin 75 783 PARIS CEDEX 16	Sommaire du n°47
Directeur de la publication Alain BRAVO Président de la SEE	Thème : les matériaux électroactifs (suite)
Rédacteur en Chef François BOUCHER	p.1 Sommaire
Adresser les propositions d'article à F. Boucher : revue3ei.art@voila.fr	p.3 Publications , Annonces
Communication Micheline BERTAUX communication@see.asso.fr	p.8 Matériaux magnétostrictifs : applications industrielles et exploitations pédagogiques Francisco ALVES Jean Baptiste DESMOULINS Christian OLLIER IUT CACHAN , SUPELEC
Publicité en Régie TRENDICE CONSEIL	P.18 Croissance et structure du niobate de lithium, matériau roi de l'opto-électronique Michel FERRIOL Université Paul Verlaine-Metz / Supélec
Philippe MINGORI 01 45 74 96 47	P.24 Cristaux photoniques intégrés sur niobate de lithium Nadege COURJAL, Richard FERRIERE, Maria-Pilar BERNAL Institut FEMTO-ST, Département d'Optique
Martine FERRON 01 45 74 96 48	P.31 Les biopiles Jean-Michel MONIER, Naoufel HADDOUR, Lorris NIARD, Timothy M. VOGEL, François BURET Ecole Centrale de LYON
Abonnement (4 numéros par an) déc. 2005, mars, juin, sept. 2006. tarifs TTC : <u>Individuel</u> : France et CEE.....35 € Pays hors CEE.....45 € <u>Collectivités</u> France et CEE.....50 € Pays hors CEE.....63 €	Hors Thème
Réalisation et impression Repro-Systèmes 23, rue de Verdun 77 181 Le Pin	p.40 Leds blanches et conversion DC-DC - 1 ^{ERE} PARTIE Frédéric NARCY Lycée Léonard de Vinci - Saint-Michel-sur-Orge
Routage et Expédition Départ Presse ZI les Richardets 93 966 Noisy le Grand	p.51 Etude d'une suspension magnétique active Pascale COSTA* Lycée Raspail Paris, François CARRERE Société S2M
Dépôt Légal : décembre 2006	p.63 Systèmes logiques séquentiels circuits câblés ou programmation directe (vhdl) ? Jérôme FAUCHER, Jérémie REGNIER, Marcel GRANDPIERRE ENSEIHT TOULOUSE
Commission Paritaire 1207 G 78028 ISSN 1252-770X	Sciences appliquées
	p.69 De la clepsydre à Michel Platini : notions fondamentales en mécanique des fluides F. LEMAIRE Lycée Colbert, LORIENT
	p.75 La perturbographie et l'analyse des réseaux Martin HANKER / LEM INSTRUMENTS & Denis KOBLER / QUALITROL
	Histoire de l'automatique
	p.81 A l'origine de l'automatique : Black, Nyquist, Bode et les Bell Laboratories Patrice REMAUD Jean-Claude TRIGEASSOU Ecole Supérieure d'Ingénieurs de Poitiers

Toute reproduction ou représentation intégrale ou partielle, par quelque procédé que ce soit, des pages publiées dans la présente édition, faite sans l'autorisation de l'éditeur est illicite et constitue une contrefaçon. Seules sont autorisées, d'une part, les reproductions strictement réservées à l'usage privé du copiste et non destinées à une utilisation collective et, d'autre part, les analyses et courtes citations justifiées par le caractère scientifique ou d'information de l'oeuvre dans laquelle elles sont incorporées.

Toutefois des copies peuvent être utilisées avec l'autorisation de l'éditeur. Celle-ci pourra être obtenue auprès du Centre Français du Droit de Copie, 20, rue des Grands Augustins, 75006 Paris, auquel la Revue 3EI a donné mandat pour la représenter auprès des utilisateurs. (loi du 11 mars 1957, art.40 et 41 et Code Pénal art. 425).

Le numéro de Septembre a eu un retard important ; nous vous prions de nous en excuser et nous espérons que son contenu riche et novateur compensera ce désagrément. De plus cet incident va induire un retard en cascade pour les prochains numéros, car le n°46 contenait le bulletin d'abonnement ou de réabonnement pour 2007 et **il nous faut souvent solliciter plusieurs fois certains abonnés étourdis**. L'article de Mr F. Lemaire de Lorient sur la mécanique des fluides, qui faisait partie du travail de l'académie de Rennes, est publié dans le n°47 et inaugure la rubrique Enseignement des Sciences Appliquées.

La partie hors thème de ce numéro est importante car beaucoup de communications étaient restés en attente. Nous remercions les auteurs de leur patience ! De plus plusieurs articles prévus sur le thème des **matériaux électroactifs** ont été réorientés : par exemple les articles sur le thème de la **thermoélectricité** seront regroupés dans un même numéro. Les panneaux photovoltaïques ayant été déjà bien étudiés dans de précédents numéros, nous prévoyons plutôt un article sur les matériaux utilisés pour la fabrication de ces panneaux photovoltaïques.

Nous continuons à raconter **l'Histoire de l'Automatique** avec les travaux de Messieurs Bode, Black et Nyquist que nous font découvrir avec bonheur Mr Remaud et Mr Trigeassou.

Les prochains thèmes restent pour le mois de Mars **le stockage de l'énergie** et pour le mois de Juin **les gisements d'économie d'énergie**.

Nous attendons avec impatience vos réponses à la consultation 2006 publiée dans le numéro de Septembre : vous pouvez envoyer votre réponse par courrier postal envoyé à la SEE 17 rue Hamelin Paris 75783 Cedex16 ou par courriel à l'adresse Revue3ei.enquete@voila.fr

Pour continuer à faire paraître cette revue qui contribue à la transmission des connaissances, nous avons toujours besoin de votre aide. Nous comptons sur vous pour saisir toutes les occasions pour faire connaître et faire vivre notre revue : réunions de jurys d'examen, rencontres entre enseignants pour l'application de nouveaux programmes, congrès, rencontres ou partenariat entre enseignants et industriels...

FAITES CONNAITRE LA REVUE !!! ..

Nous serons heureux de recevoir vos articles que vous aurez déposés dans la boîte aux lettres dont l'adresse e-mail est revue3ei.art@voila.fr. Vous pouvez également nous faire parvenir vos idées, vos réactions, vos suggestions à l'adresse e-mail revue3ei.cour@voila.fr

Bonne lecture.

Le Comité de Publication de la Revue3EI

La Revue 3EI

Comité de publication

Hamid BEN AHMED
Jean BONAL
François BOUCHER
Lucien DESCHAMPS
Jean FAUCHER
Gilles FELD
Jean-Philippe ILARY
Chérif LAROUCI
Marie Michèle LE BIHAN
Franck LE GALL
Sylvaine LELEU
Pascal LOOS
Claude OBERLIN
Oviglio SALA
Jean-François SERGENT
Jean-Claude VANNIER
Pierre VIDAL

Pour vos insertions
publicitaires, contacter :
TRENDICE CONSEIL

Philippe MINGORI
01 45 74 96 47

Martine FERRON
01 45 74 96 48

Abonnement à la Revue 3EI, année 2005-2006 :
Numéros : 43 (décembre 2005), 44 (mars), 45 (juin) et 46 (septembre 2006).

<u>Abonnement individuel :</u>		<u>Abonnement collectif souscrit par bon de commande (bibliothèque, CDI, laboratoire, entreprise, etc.)</u>	
France et Pays de la CEE :	35 €, TTC	France et Pays de la CEE :	50 €, TTC
Pays hors CEE :	45 €, TTC	Pays hors CEE :	63 €, TTC

Une seule adresse :

La Revue 3EI - SEE,

17, rue Hamelin, 75 783 PARIS Cedex 16

**pour nous contacter au sujet de votre abonnement écrivez à
revue3ei.cour@voila.fr**



Les convertisseurs de l'électronique de puissance. Vol. 1 : La conversion alternatif-continu (3° Ed.)

SÉGUIER Guy

Collection tec et doc Lavoisier

Résumé ;

Les convertisseurs de l'électronique de puissance, volume 1, est le premier ouvrage d'une série de cinq consacrée à l'étude approfondie des convertisseurs. Les quatre premiers traitent du fonctionnement et des caractéristiques des quatre grandes familles de convertisseurs tandis que le cinquième présente les procédés de commande. Le volume 1 est naturellement dédié aux montages qui assurent la transformation alternatif-continu, autrement dit aux redresseurs, car ce sont les plus utilisés et, croit-on, les plus simples. Les redresseurs classiques utilisent des diodes ou des thyristors, c'est-à-dire des semi-conducteurs non commandés ou commandés seulement à la fermeture. De ce fait, le fonctionnement et les caractéristiques sont très tributaires de la source alternative et du récepteur continu. Si la présentation des redresseurs peut être simple et rapide, cet ouvrage procède à une étude quantitative assez précise qui nécessite plus de calculs que celle des convertisseurs utilisant des semi-conducteurs totalement commandés. Il propose ainsi de nombreuses planches de caractéristiques qui permettent l'utilisation directe des résultats obtenus. L'évolution des composants semi-conducteurs et l'extension du domaine d'application de l'électronique de puissance accroissent le recours à des structures avec semi-conducteurs totalement commandés pour la transformation alternatif-continu. Cette troisième édition s'enrichit donc d'un long chapitre sur les redresseurs à modulation de largeur d'impulsions et sur les correcteurs de facteur de puissance. La bibliographie commentée en fin de volume est réactualisée et complétée pour tenir compte de cette addition.



Nouvelles technologies de l'énergie 1 : les énergies renouvelables (Traité EGEM, série Génie électrique)

SABONNADIÈRE Jean-Claude

Editions Hermès - Lavoisier

Resumé

Cet ouvrage consacré aux énergies renouvelables est le premier d'un ensemble de trois tomes destinés à faire le point sur l'état de l'art dans les nouvelles technologies de l'énergie. L'ouvrage débute sur un ensemble de plusieurs chapitres sur l'énergie solaire (thermique, photovoltaïque, thermodynamique) complété par un chapitre consacré aux technologies de l'éolien. Enfin une partie importante est consacrée aux nouvelles énergies de nature hydraulique où l'on trouve les énergies de la mer et la très petite hydraulité avec en particulier l'exploitation de l'énergie motrice des conduites d'eau potable et des eaux usées.

Sommaire

Avant-propos -J.-C. Sabonnadière. L'énergie électrique. Le nouveau paradigme. La production décentralisée. Les moyens de conversion de l'énergie. Solaire. Production photovoltaïque d'électricité -J.-C. Muller. Systèmes photovoltaïques couplés en réseau -S. Bacha, D. Chatroux. Chauffage solaire -Ch. Marvillet. Les centrales solaires thermodynamiques -A Ferrière. Éolien. Technologie des systèmes éoliens -R. Belhomme, D. Roye, N. Laverdure. Intégration des générateurs éoliens au réseau -R. Belhomme, D. Roye, N. Laverdure. Énergies hydraulique et marines. Systèmes de conversion des ressources énergétiques marines -B. Multon, A. Clément, M. Ruellan, J. Seigneurbieux, H. Ben Ahmed. La petite hydroélectricité -R. Chenal, A. Choulot, V. Denis, N. Tissot. Index



RÉSÉLEC

Réseau National de Ressources en Électrotechnique

NOUVEAUTÉS

nouveaux référentiels

habilitation électrique

productions

contributions

événements

concours

communiquer

forum

liens

partenaires industriels

présentation de Résélec

contacts

rechercher

ministère
éducation
nationale
enseignement
supérieur
recherche



Nos missions 2006 / 2007 :

Production de ressources pédagogiques innovantes en génie électrique

- ♦ pour l'accompagnement de la rénovation
 - du BTS Électrotechnique,
 - du référentiel "Habilitation électrique",
- ♦ et le développement de la formation en économie d'énergie, éco conception, énergie renouvelables.

Missions antérieures :

Élaboration de documents pédagogiques pour

- ♦ le CAP PROELEC,
- ♦ le Bac Pro ELEEC,
- ♦ le BEP des métiers de l'électrotechnique,
- ♦ l'habilitation électrique.

www.iufmrese.cict.fr

Contacts :

Responsable pédagogique : Pascal MAUSSION
pmaussio@cict.fr
Chargé de la communication : Michel LEFÈVRE
mlefevre@cict.fr

Réseau National de Ressources en Électrotechnique
IUFM Midi-Pyrénées Site de Rangueil 118 route de Narbonne
31078 TOULOUSE CEDEX 4
Tél.: 05 62 25 21 85 ~ Télécopie: 05 62 25 21 58



TEXTE DE PRESENTATION RESELEC REVUE 3EI DECEMBRE 2006

Le ministère de l'Éducation nationale a créé en 1991 le Réseau National de Ressources en Électrotechnique à l'initiative de l'Inspection Générale. Il s'intègre dans le dispositif de formation continue géré par la Direction générale de l'enseignement scolaire (D.G.E.S.C.O.).

La mission du Réseau National de Ressources en Electrotechnique est de découvrir, de soutenir la production et de diffuser des ressources pédagogiques innovantes au service des enseignants de la filière génie électrique. Il s'agit principalement de promouvoir les nouvelles orientations liées aux rénovations de diplômes, mais en se basant sur les pratiques innovantes acquises par les enseignants dans leurs classes.

Outre les développements de TP spécifiques à la filière, le Réseau a publié des documents multimédia testés devant élèves en direction des filières : Génie Électrique, MI, MAI, Énergétique, Électronique, dans le domaine de la prévention des risques d'origine électrique depuis 1995.

Objectif de la mission 2006/2007

Les axes prioritaires suivants ont été retenus sur cet exercice par l'Inspection Générale :

- accompagnement de la rénovation du BTS d'électrotechnique,
- accompagnement de la rénovation du CAP PROELEC et du Bac Pro ELEEC,
- approche des énergies renouvelables et applications dans le cadre actuel des diplômes.

Les productions du Réseau National de Ressources en Electrotechnique

Les productions du RéSéLec émanent de propositions d'enseignants de la filière génie électrique, parfois impulsées par des Inspecteurs, en rapport avec les axes retenus. Elles ont été validées par les corps d'Inspection (IPR ou IEN) avant sélection. La répartition des moyens est effectuée sous l'autorité de l'Inspection Générale.

Les publications sélectionnées pour l'exercice 2006-2007 traitent notamment des sujets suivants :

- BTS électrotechnique : l'activité de projet,
- Mise en place de maquettes numériques (modeleur volumique SW) sous forme réelle et schématique des systèmes de base des ateliers de BTSET (axe x, axe z, levage...),
- Production décentralisée d'énergie électrique à partir d'un générateur photovoltaïque industriel. Dans ce cas, les systèmes objets d'étude ne sauraient être dissociés de l'étude de consommateurs d'énergie associés. Il s'agit de promouvoir une efficacité énergétique durable prenant sens, dans le cycle : génération, conversion, tarification/gestion de l'énergie produite et consommateurs en fonction des usages (chauffage, eau chaude sanitaire, transport personnel ou éclairage....).

Tête de Réseau TOULOUSE

			
Responsable pédagogique: Pascal MAUSSION , Maître de conférences, IUFM de Midi-Pyrénées	Animateur chargé de la communication : Michel LEFÈVRE ; IUFM de Midi-Pyrénées	Coordonnateur des publications lycée professionnel : Henri MARTEL-HÉBRARD Lycée Professionnel Renée Bonnet, Toulouse,	Coordonnateur des publications lycée technologique : Gérard OLLÉ Lycée Technologique Déodat de Séverac, Toulouse



Journées 2007 de la section Électrotechnique du club EEA



« *Energie et développement durable* »

Antenne de Bretagne de l'ENS de Cachan

14 et 15 mars 2007

Pour traiter cette problématique d'actualité, des spécialistes renommés interviendront dans un large domaine de compétences. Les expériences pédagogiques de collègues, déjà sensibles à cette problématique, feront l'objet d'une session posters et démonstrations. A la fin du colloque, la table ronde s'est fixée pour objectif d'accélérer la dissémination des problématiques énergétiques et environnementales au sein des enseignements des domaines du Génie Electrique et de la Physique Appliquée.

Programme

■ Session 1 « Problématiques générales »

Thomas GUERET (NegaWatt)
Pierre MATARASSO (CIRED, CNRS)
Bernard CHABOT (ADEME)

■ Session 2 « Eco-pédagogie » (Posters expositions)

Présentation des posters « Eco-pédagogie »
par Christophe ESPANET (UFC, L2ES)

■ Session 3 « Economies d'énergie et analyses sur cycle de vie »

Yves LOERINCIK (EPF Lausanne)
Maxime TROCME (GTM-Construction)
Benoît LEBOT (PNUD-FEM)

■ Session 4 « L'électricité et le développement durable

Georges ZISSIS (UPS, CPAT)
François COSTA (ENS Cachan, SATIE)
Laurent GONTHIER (ST Microelectronics)
Joseph BERETTA (PSA)
Jean-Claude MULLER (INESS, CNRS)

■ Session 5 « Scénarios du futur ? »

Vincent MAZAURIC (Schneider Electric)
MAIZI-MENARD (Ecole des Mines de Paris)
Pierre MATARASSO (CIRED, CNRS)



Appel à communications (posters ou démonstrations) *Eco-pédagogie*

Merci d'adresser vos propositions à Christophe ESPANET : Christophe.Espanet@univ-fcomte.fr

Informations : www.bretagne.ens-cachan.fr/jeea2007,

Email secrétariat (Solène DUMAST) : dumast@bretagne.ens-cachan.fr

Electronique de puissance : du régime continu au régime impulsif

JEAN BONAL

Expert énergie et conversion de l'énergie – Association ECRIN

Résumé des 10ème entretiens « Physique Industrie : Electronique de puissance » qui ont eu lieu le Jeudi 19 Octobre 2006 Salon Mesurexpo Paris

Pour leur dixième édition, les entretiens « Physique-Industrie » ont été consacrés à l'électronique de puissance.

L'électronique de puissance « en régime continu » est maintenant bien connue, elle concerne notamment les domaines de la conversion, du stockage, et du transport de l'énergie électrique [convertisseurs de tension (AC/DC, DC/DC, DC/AC) – convertisseurs de fréquence, variateurs de vitesse, FACTS, chargeurs de batteries...]. Les progrès récurrents des performances des semi-conducteurs ont permis d'améliorer de manière continue l'efficacité des systèmes de conversion, les moyens de commande et de réglage des réseaux de transport de l'électricité ; et c'est aussi grâce aux performances des systèmes électroniques de puissance que les énergies dites renouvelables ont pu se développer et devenir une alternative crédible pour compléter le panel énergétique.

L'électronique de puissance en régime impulsif est moins connue, pourtant elle permet d'accéder à des effets physiques rapides ou transitoires induits par interactions électroniques, magnétiques de grande puissance crête, de l'ordre de quelques gigawatts, dans des durées allant de quelques picosecondes à des microsecondes. Les générateurs d'impulsions capables de provoquer ces effets sont rares, coûteux et difficiles à réaliser, les conditions expérimentales sont inusuelles, peu explorées et donc peu employées. Il en découle que les effets induits mécaniques, électrochimiques, radioactifs et biologiques ... d'un grand intérêt potentiel sont moins exploités que ceux obtenus en optique impulsif.

La conférence placée sous la présidence de Mr Yves Farge, membre de l'Académie des Technologies, a réuni 130 participants et a permis de mettre en évidence :

- l'évolution de la manière dont on appréhende l'utilisation de l'énergie électrique à bord des vecteurs civils et militaires (avions, véhicules terrestres, bateaux...)

- le fait que confrontés aux exigences d'efficacité énergétique, les semi-conducteurs de puissance évoluent par saut plutôt que par facteur d'échelle. Ils sont en effet placés au centre d'une problématique multi dimension (nature du semi-conducteur, technologie utilisée, application au sein du système) de sorte que, si pour les circuits intégrés et les mémoires le référentiel est la loi de Moore, pour les composants de puissance le référentiel est le système dans lequel ils sont insérés.

- l'importance des systèmes de conversion électrique dans le fonctionnement d'un outil de recherche et de développement tel qu'Iter.

- la nécessité de posséder des équipements de réglage et de commande des réseaux de transport électrique appropriés pour répondre aux exigences induites par la dérégulation des marchés et l'émergence d'une foule de centres de production d'énergie électrique d'origine renouvelable, éolien, solaire, qui se caractérisent par des niveaux de puissance unitaire relativement faible, des livraisons fluctuant au gré des conditions climatiques avec des durées de fonctionnement aléatoires ou encore fortement imprévisibles.

- les contraintes afférentes à la définition des systèmes impulsifs :

- tension : du kV à quelques MV
- courant : du kA à quelques MA
- impédances : de la centaine de mΩ à quelques dizaines d'ohm
- temps de montée : de la centaine de ps à quelques μs et même quelques ms
- durée : de la centaine de ps à quelques μs et même quelques ms
- fréquence : monocoup à quelques Hz

- de présenter quelques machines d'électronique de puissance impulsives.

Les actes de la manifestation sont disponibles sur le site de la Société Française de Physique (SPF) en suivant : conférences 2006, EPI 10^{ème}, présentations.

<http://sfp.in2p3.fr/expo/>

Matériaux magnétostrictifs : applications industrielles et exploitations pédagogiques.

FRANCISCO ALVES* (Professeur des Universités) Francisco.alves@iut-cachan.u-psud.fr, francisco.alves@lqep.supelec.fr

JEAN-BAPTISTE DESMOULINS ** (Professeur Agrégé) desmouli@physique.ens-cachan.fr

CHRISTIAN OLLIER** (Assistant Ingénieur) ollier@physique.ens-cachan.fr

* Département GEI1, IUT de Cachan, 9 Avenue de la Division Leclerc 94234 Cachan

Laboratoire LGEP UMR 8507, Membre de SPEE Labs, Supélec, Univ Paris Sud, Univ Paris et Marie Curie, 11 Rue Joliot-Curie, Plateau du Moulon, 91192 Gif sur Yvette.

** Département de Physique, ENS Cachan

61, Av. Pdt. Wilson 94235 Cachan.

Résumé : *Le phénomène de magnétostriction est connu depuis le milieu du 19^{ème} siècle. Il a été mis en évidence par le physicien J.P.Joule, qui en appliquant un champ d'excitation magnétique de plus en plus intense à un barreau de fer doux a montré que ce dernier s'allongeait puis se contractait. Quelques années plus tard, E. Villari montra que les propriétés magnétiques d'un matériau ferromagnétique étaient modifiées lorsque ce dernier était soumis à des contraintes mécaniques. La magnétostriction, observable dans les matériaux ferrimagnétiques et ferromagnétiques, peut être une source de nuisance dans certains cas, mais elle est également utilisée dans différents types d'applications comme la génération d'ondes acoustiques pour les sonars, la réalisation d'actionneurs linéaires ou de capteurs...*

Dans cet article, nous allons commencer par présenter quelques points importants concernant le phénomène de magnétostriction, notamment les ordres de grandeurs et quelques exemples d'applications. Nous présenterons ensuite quelques exemples de techniques de mesure permettant de déterminer les allongements relatifs caractéristiques du phénomène.

1. Présentation du phénomène de magnétostriction

1.1. Origine du phénomène, cas de la magnétostriction linéaire.

Lorsqu'un champ d'excitation magnétique est appliqué à un matériau comportant des atomes ayant des caractéristiques ferro- ou ferri-magnétiques, les interactions entre les atomes de la maille cristalline vont être modifiées, ce qui va provoquer une modification des distances inter-atomiques dans cette dernière. A une échelle plus macroscopique, il faudra tenir compte de la structure en domaines, puisque l'organisation magnétique change d'un domaine à l'autre. Globalement, pour l'échantillon complet, on va observer une déformation, qui sera soit une dilatation, soit une contraction, suivant le matériau et le niveau de champ appliqué. Si l'observation est réalisée dans la direction d'aimantation, on parle d'effet Joule longitudinal.

La magnétostriction va se manifester de différentes manières suivant les échantillons étudiés et les conditions

auxquelles sont soumis ces derniers. On peut notamment citer :

- L'effet Joule transversal : Cet effet découle de l'effet longitudinal puisque l'allongement s'accompagne d'une contraction transversale. Les effets Joule longitudinal et transversal n'introduisent pas de variation de volume si le matériau est isotrope.

- L'effet Wiedemann : On observe la torsion d'un barreau cylindrique aimanté et parcouru par un courant suivant son axe, et réciproquement.

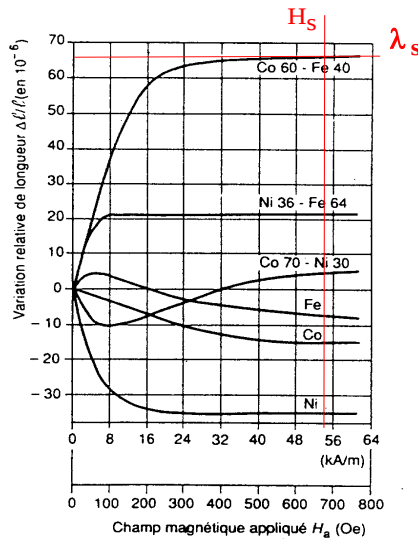
- L'effet Guillemin : On observe la flexion d'un barreau lorsqu'il est soumis à un champ longitudinal.

- L'effet de variation du module d'Young, ou effet ΔE : sous l'action d'une contrainte, la déformation due à l'aimantation s'ajoute à la déformation élastique, ce qui entraîne une modification des constantes élastiques, donc du module d'Young.

a/ Evolution avec le champ d'excitation appliqué

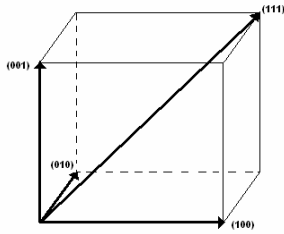
L'allongement relatif, appelé magnétostriction, sera noté λ . Sur la figure suivante, pour différents matériaux et alliages, on observe des allongements relatifs en fonction

du champ. Très souvent, on donne le coefficient de magnétostriction à saturation, puisque l'allongement dépend de l'état magnétique. On le note λ_s .



b/ Magnétostriction et anisotropie

La magnétostriction est souvent anisotrope, notamment dans les structures cristallines. Par exemple, dans le cas des systèmes à structure cristallographique cubique, on peut décrire les variations de dimensions à l'aide de deux constantes, λ_{100} et λ_{111} associées aux directions cristallographiques correspondantes. Sur la figure suivante, on a représenté les différentes directions particulières du cristal.



Si on imagine un alignement de tous les moments atomiques dans une direction dont les cosinus directeurs par rapport à l'axe cristallographique [100] sont α_i et si on s'intéresse à l'allongement dans une direction de cosinus directeurs β_i , on trouve que l'allongement est de la forme:

$$\frac{\Delta l}{l} = \frac{3}{2} \lambda_{100} \left(\alpha_1^2 \beta_1^2 + \alpha_2^2 \beta_2^2 + \alpha_3^2 \beta_3^2 - \frac{1}{3} \right) + 3 \lambda_{111} (\alpha_1 \alpha_2 \beta_1 \beta_2 + \alpha_2 \alpha_3 \beta_2 \beta_3 + \alpha_1 \alpha_3 \beta_1 \beta_3)$$

Par conséquent, l'allongement mesuré vaut λ_{100} dans la direction [100] ($\beta_1=1, \beta_2=0, \beta_3=0$) et vaut $-\lambda_{100}/2$ dans la direction perpendiculaire. Dans le cas simple de matériaux isotropes, par exemple, les matériaux amorphes ou nanocristallins (grains répartis aléatoirement dans une phase amorphe), on a $\lambda_{100}=\lambda_{111}=\lambda$. L'allongement s'écrit alors:

$$\frac{\Delta l}{l} = \lambda(\theta) = \frac{3}{2} \lambda \cdot (\cos^2 \theta - \frac{1}{3})$$

si λ est la magnétostriction obtenue lorsque la direction d'observation de l'allongement est parallèle à la direction

d'aimantation et θ l'angle entre la direction d'observation de l'allongement et la direction d'aimantation.

1.2. Exemples de matériaux magnétostrictifs et ordres de grandeur.

a/ Ordres de grandeur

Suivant les matériaux employés, on peut obtenir des ordres de grandeur très variés pour le coefficient de magnétostriction à saturation. Le tableau donné page suivante donne quelques exemples. Le coefficient k_{33} représente l'aptitude d'un matériau à convertir son énergie magnétique en énergie mécanique et inversement.

b/ Cas particulier du Terfenol-D.

Le Terfenol-D (abréviation de **T**erbium **F**e **N**aval **O**rdnance **L**aboratory **D**ysprosium) est utilisable à température ambiante. Sa caractéristique reliant excitation magnétique et magnétostriction dépend de la contrainte. Pour obtenir une réponse la plus linéaire possible, on a intérêt à le polariser en excitation et en contrainte. Comme la plupart des matériaux magnétiques métalliques, le Terfenol est sujet aux courants de Foucault, ce qui limite la plage de fréquence sur laquelle on peut le solliciter et conduit à des pertes qui dégradent la qualité du transfert de puissance. Par ailleurs, le matériau est très cassant, ce qui le rend délicat à laminier et donc interdit son emploi dans des structures feuilletées.

1.3. Exemples de manifestations de la magnétostriction.

Le phénomène de magnétostriction peut être source de problèmes dans les structures électrotechniques en provoquant l'apparition de vibrations, sources de bruit et éventuellement d'usure mécanique. Cependant, le phénomène peut être à l'origine de nombreuses applications, notamment lorsqu'il correspond à des allongements relatifs de quelques 1000 ppm comme pour le Terfenol.

a/ Aspects négatifs

La magnétostriction va se manifester dans tous les circuits magnétiques. Si ces derniers sont soumis à des champs variables, ils seront sièges de vibrations qui peuvent être sources de bruit. Ce bruit peut être simplement désagréable, mais peut également poser des problèmes mécaniques. Par ailleurs, les différentes contraintes mécaniques auxquelles sont soumis les matériaux peuvent dégrader leurs propriétés magnétiques par l'intermédiaire du couplage auquel est associée la magnétostriction. On peut par exemple citer les colonnes de transformateurs soumises à leur propre poids, les contraintes de découpe des tôles, ou l'effet des contraintes lors de la réalisation de tores enroulés... Cette modification des propriétés magnétiques peut être responsable d'une partie des pertes « fer » des circuits magnétiques.

matériau		Induction à saturation B_s (T)	λ_s (10^{-6} ou ppm)	T_c (K)	k_{33}
Fe			-9	1040	0,3
Co			-62		0,3
Ni			-35	630	0,3
Amorphe base fer	$Fe_{66}Co_{18}B_{15}Si_1$	1,8	35		0,71
	$Fe_{78}B_{13}Si_9$	1,56	27		
Amorphe base cobalt	$Co_{66}Fe_4B_{12}Si_{16}Mo_2$	0,55	0,2		
Nanocristallin	$Fe_{73,5}Si_{13,5}B_9Cu_1Nb_3$	1,35	2,3		
Cristallins	Tôle FeSi3,5%	1,9	7,8		0,26
	$Fe_{49}Co_{49}V_2$	2,35	70		
Ferrite AVX H2	NiZn	0,3	-10		
Terfenol-D	$Tb_{0,3}Dy_{0,7}Fe_{1,9}$		1600-2400	650	0,75
	$Tb_{0,6}Dy_{0,4}$ (77K)		6300	210	
	$Tb_{1-x}Dy_xZn$		5000	200	

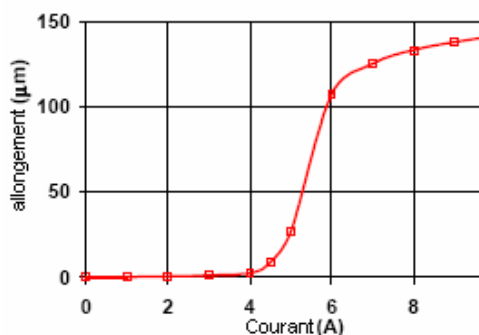
b/ Actionneurs magnétostrictifs

Différentes entreprises proposent désormais des actionneurs utilisant des matériaux à magnétostriction géante, soit à température ambiante (Terfenol-D), soit à température cryogénique (TbDyZn) suivant le type d'applications recherchées. On trouve principalement deux types d'actionneurs.

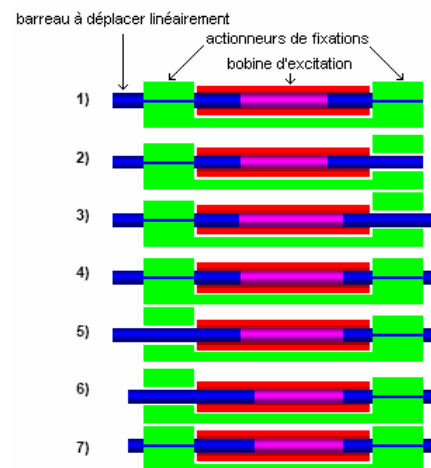
Les actionneurs dits « linéaires », sont réalisés à partir d'un barreau de matériau magnétostrictif entouré par une bobine dans laquelle on fera passer un courant qui permettra de contrôler l'allongement. Quelques actionneurs de ce type sont présentés sur la photographie suivante [Ener].



Ils permettent d'appliquer une force voisine de 1000N sur une course allant jusqu'à une centaine de μm à 77K et moins pour des températures plus élevées. Des actionneurs permettant de créer des efforts plus importants peuvent être réalisés en augmentant la taille du barreau magnétostrictif. On peut alors atteindre des efforts d'une vingtaine de kN. Dans ces systèmes, l'allongement en fonction du courant n'est pas linéaire mais présente plutôt l'allure suivante [Ener].



Il existe aussi des actionneurs linéaires pas à pas construits à partir des actionneurs linéaires qui viennent d'être présentés. Ils sont réalisés à partir de trois actionneurs linéaires. L'un, construit autour d'un barreau magnétostrictif, destiné à réaliser l'effet de translation recherché, et les deux autres permettant de fixer, ou de laisser libre chacune des extrémités de l'échantillon précédent, la fixation étant assurée lorsque l'actionneur correspondant n'est pas excité. Le fonctionnement du système, appelé « inchworm », peut être décrit par la séquence suivante [Ener].



Lors de l'étape 1, le système est au repos. A l'étape 2, on active l'actionneur de fixation à droite qui relâche la fixation du barreau. Lors de l'étape 3, l'actionneur linéaire de translation est excité, ce qui provoque l'allongement du barreau vers la droite. On coupe l'excitation de l'actionneur de droite lors de l'étape 4, ce qui entraîne la fixation du barreau de ce côté. Lors de l'étape 5, on excite alors l'actionneur de fixation de gauche qui relâche la pression de ce côté. On coupe ensuite le bobinage longitudinal, ce qui provoque une contraction du barreau lors de l'étape 6. A l'étape 7, l'actionneur de fixation de gauche cesse d'être alimenté et le barreau est à nouveau fixé de ce côté.

Globalement, avec ces 7 étapes, on s'est déplacé d'un pas vers la droite. Si on recommence dans ce sens, on pourra se déplacer d'un pas supplémentaire dans cette direction. Si on procède dans l'ordre inverse en débloquent à gauche puis à droite, le barreau se déplacera en sens inverse.

L'intérêt de ce système est de permettre un allongement dans les deux directions du barreau central et de maintenir un état d'allongement en l'absence d'alimentation, lorsque les deux extrémités du barreau restent fixées. Par ailleurs, on peut atteindre des déplacements importants (plusieurs mm), avec une bonne résolution (jusqu'au nm) tout en tenant des charges mécaniques importantes (2500N). La vitesse maximale de déplacement va dépendre du pas choisi.

Un actionneur de ce type est présenté sur la photographie suivante :



Ce type d'actionneur est destiné à être utilisé sur le NGST (New Generation Space Telescope) destiné à remplacer le satellite Hubble [Ener]. Il doit permettre de réaliser des miroirs dont l'état de surface peut être contrôlé par une matrice d'actionneurs commandant une fraction du miroir, afin de corriger les aberrations et les effets thermiques. Pour ce type d'application, on peut utiliser un matériau fonctionnant dans des conditions cryogéniques.



Les actionneurs magnétostrictifs de ce type sont également utilisés dans les cavités SRF (Supraconducting radiofrequency) pour compenser les déformations dues aux effets thermiques et aux limites de tolérances, afin d'obtenir des faisceaux d'électrons homogènes. Pour cela, l'actionneur va allonger ou contracter la cavité.

c/ Génération d'ondes acoustiques : les sonars

Durant la Seconde Guerre Mondiale, les premiers sonars furent réalisés à partir de transducteurs magnétostrictifs [Pier] en utilisant des matériaux à base de nickel. A partir des années 50, ils furent remplacés par des céramiques piézoélectriques. Avec l'apparition des matériaux à magnétostriction géante à température ambiante, il a été à nouveau possible d'envisager leur emploi dans la réalisation des sonars, car ces matériaux permettent d'obtenir des échos à de plus grandes distances qu'avec les sonars de précédente génération.

En excitant un matériau à magnétostriction géante soumis à une contrainte mécanique, à sa fréquence de résonance, on peut émettre une puissance acoustique très importante (jusqu'à 10kW). Ce type de dispositif a été utilisé pour les sonars et permet d'obtenir une meilleure densité de puissance qu'avec les systèmes piézoélectriques équivalents. La figure suivante présente un prototype de transducteur basse fréquence et forte puissance mis au point pour la marine française (projet R&D Cedrat/eramer, [Cedr]).



L'émission d'ondes acoustiques est aussi utilisée pour du nettoyage par ultrasons et même pour transmettre des informations [Cedr].

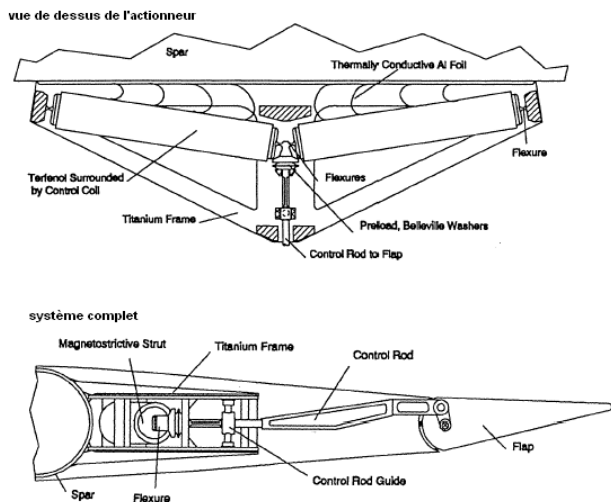
On rencontre désormais des systèmes à base de Terfenol-D destinés à émettre des ondes sonores dans des applications plus ludiques [Newl]. Le système est fixé sur un support (table, vitre...) qui va servir de caisse de résonance pour le signal sonore à diffuser. L'emploi de ce type de système peut également être envisagé pour réduire l'effet des bruits extérieurs dans une pièce via un système actif de bruit ou pour nettoyer des vêtements.

d/ Contrôle actif de vibrations

Les transducteurs magnétostrictifs permettent d'envisager de faire du contrôle actif de vibrations. Par exemple, l'une des applications est la réduction des vibrations dans les pales d'hélicoptères [Fenn], qui peut se faire sur un plus grand nombre d'harmoniques avec ces transducteurs qu'avec ceux utilisés jusqu'à présent. Les figures suivantes présentent un actionneur magnétostrictif construit autour de barreaux de Terfenol excités par des bobinages ainsi que leur intégration dans une palle du rotor.

Le résultat attendu est une réduction de 90% des vibrations. Pour ce type d'applications, la solution d'une transduction magnétostrictive est préférable aux autres solutions existantes [Fenn]. Les actionneurs hydrauliques sont moins intéressants en raison du problème posé par le stockage du fluide, la complexité mécanique de la

réalisation et les coûts de maintenance. Les actionneurs piézoélectriques ne sont pas satisfaisants, en raison d'une densité d'énergie de l'actionneur trop faible (15 fois moins pour du PZT que pour du Terfenol). Par ailleurs, ils nécessitent de travailler à haute tension, ce qui est inacceptable ici.



e/ Capteurs

Les capteurs magnétostrictifs sont utilisés essentiellement comme capteurs de position sans contact. Ils exploitent la caractéristique magnétique et la déformation micro-élastique d'un élément magnétostrictif pour détecter la position exacte d'un aimant. Grâce à l'absence de contact entre les parties mobiles et l'élément sensible, les transducteurs magnétostrictifs garantissent le maintien des performances d'origine du produit tout au long de son cycle de vie. Le capteur est entièrement étanche, ce qui le rend compatible avec des conditions d'utilisation difficiles (par exemple, présence de saletés ou immersion).

Ce type système est donc constitué d'un barreau magnétostrictif qui va jouer le rôle de guide d'ondes et d'un aimant mobile dont on cherche à déterminer la position. Le principe utilisé par les capteurs de position repose sur deux effets magnétostrictifs : l'effet *Wiedemann* et l'effet *Villari*.

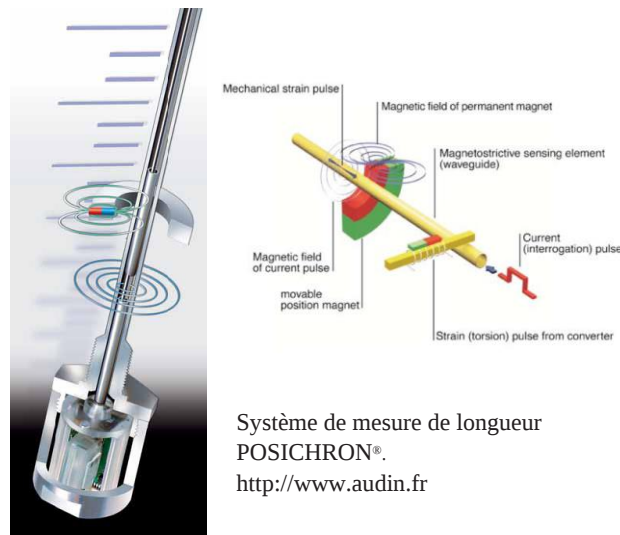
Pour créer l'effet *Wiedemann*, une impulsion de courant est envoyée dans l'axe longitudinal du barreau magnétostrictif. Cette impulsion génère un champ magnétique orthoradial autour du barreau qui « suit » l'impulsion, à une vitesse proche de la vitesse de la lumière. Dès que le champ magnétique orthoradial arrive au voisinage de la zone où règne le champ magnétique de l'aimant du curseur de position, une onde élastique torsionnelle est créée par magnétostriction. Cette onde se propage alors à une vitesse sonique dans le barreau magnétostrictif.

L'extrémité du barreau contient un système de détection qui repère l'arrivée de l'onde élastique. La détection repose sur l'effet magnéto-élastique *Villari*. Afin de connaître la distance qui sépare le détecteur de l'aimant

mobile, l'électronique intégrée permet de récupérer une image du décalage temporel entre l'impulsion de courant électrique initiale et l'impulsion de tension générée par l'effet *Villari* dans la bobine du détecteur (principe du « temps de vol »). L'écart temporel est alors converti en signal analogique ou numérique selon des méthodes traditionnelles de traitement du signal (0-10 VDC ou 4-20 mA).

L'aimant curseur est soit guidé le long du profilé par une liaison mécanique avec la partie mobile du système via une rotule, soit libre de mouvement sans guidage. Les courses varient de 50 mm à 5000 mm avec une linéarité de 0,02%, une répétabilité de 0,001% et une fréquence de mesure de 16 kHz.

Avec un degré de protection IP 65, ces capteurs fonctionnent dans une plage de température comprise entre -40°C et + 75°C.



Système de mesure de longueur
POSICHRON®.
<http://www.audin.fr>

Détails du fonctionnement d'un capteur de position magnétostrictif. Source [Aud]

f/ Bilan

Les actionneurs magnétostrictifs fonctionnent avec des tensions moindres que les actionneurs piézoélectriques et ils permettent d'obtenir des efforts plus importants mais leur emploi reste limité à quelques gammes d'applications très particulières (émission d'ondes basse fréquence de forte puissance, positionnement dans des conditions cryogéniques, transducteur développant des efforts importants...). Cependant, la diffusion plus simple du Terfenol-D à un coût plus faible, que lorsque son emploi était restreint à des applications militaires, permet d'envisager un nombre d'applications croissant.

Les capteurs de position magnétostrictifs trouvent parfaitement leur place dans le domaine industriel, nombre de fabricants fournissent désormais ce type de produit.

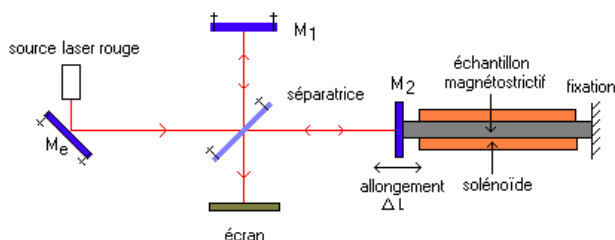
2. Exemples de mesure de magnétostriction sur différents types d'échantillons.

Les déformations relatives à la magnétostriction sont habituellement si faibles qu'on ne peut les mesurer qu'en les amplifiant à l'aide de moyens mécaniques, électriques ou optiques. Lors des premières expérimentations, la déformation de longues tiges soumises à un champ magnétique était observée par une méthode optique : un rayon lumineux était dévié par la rotation d'un miroir associé à un levier mécanique qui amplifiait l'allongement de l'échantillon. La déviation du spot restait très faible, c'est pourquoi des méthodes interférométriques ont été introduites en 1883 par Lochner. Plus tard, des méthodes électriques ont été développées avec le dilatomètre capacitif décrit par Whiddington en 1920, et la méthode des jauges de déformation développée par Goldman en 1947. Des méthodes de plus en plus précises et sensibles ont été employées dans un passé récent, et il est maintenant possible de mesurer directement de très faibles déformations magnétostrictives sur des fils fins (50 μm de diamètre), des rubans (typiquement 20 μm d'épaisseur), et même des films fins (100 nm d'épaisseur) par la méthode SAMR (Small Angle Magnetization Rotation, développée par Narita [Nari]).

2.1. Mesure sur un barreau par interférométrie.

a/ Structure générale du dispositif.

Pour observer les faibles allongements (qq μm) mis en jeux dans l'expérience proposée, nous avons choisi une structure d'interféromètre de Michelson. Comme source lumineuse, on utilise un laser rouge He-Ne (longueur d'onde λ de l'ordre de 625 nm). Il faut noter que dans ce type d'expérience, lorsque l'on cherche à mesurer l'allongement du trajet optique sur un bras de l'interféromètre, la résolution augmente lorsque la longueur d'onde de la source diminue... Cette dernière vaut en effet $\lambda/2$. Globalement, le système présente la structure suivante :



On cherche à faire en sorte que M_2 et M_1 soient perpendiculaires afin d'observer des anneaux. Une fois les anneaux obtenus, on va augmenter le courant dans le solénoïde afin d'augmenter le champ H . Avec l'évolution de H , le barreau à étudier va changer de longueur (se dilater ou se contracter suivant les cas et les matériaux). L'une des extrémités du barreau étant fixée, l'autre, celle

où est fixé le miroir, va se déplacer. Entre deux maxima d'intensité successifs au centre des anneaux, le miroir M_2 s'est déplacé de $\lambda/2$.

Le solénoïde comporte 6000 spires afin d'atteindre des champs intenses sans avoir à appliquer des courants trop importants. Expérimentalement, on trouve que

$$H = 11710 \cdot I \quad (H \text{ en A/m et } I \text{ en A})$$

L'échantillon est placé au centre du solénoïde afin de rester dans la zone où le champ d'excitation reste homogène. Le solénoïde fait 50 cm de long et l'échantillon testé 44 cm.

Avec un tel circuit et une fois l'échantillon placé au centre, si on applique un échelon de tension, le temps de réponse à 63% du courant vaut 10 ms environ. On peut donc supposer que le courant met 50 ms environ pour atteindre le régime permanent.

b/ Réalisation pratique de la mesure.

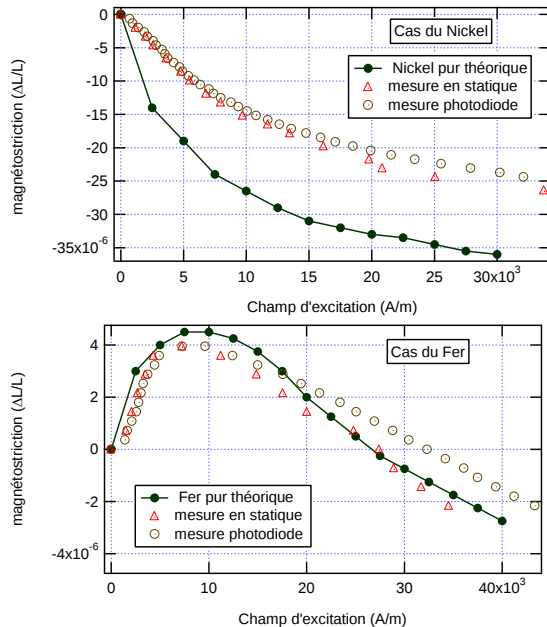
On fait défiler un nombre d'anneaux fixé et on mesure le courant qu'il a fallu appliquer pour obtenir ce résultat. Expérimentalement, on procède de deux façons à chaque fois. La première consiste à compter les anneaux qui défilent à l'œil. La seconde consiste à récupérer simultanément le courant dans le solénoïde et la tension image du courant délivré par une photodiode placée au centre des anneaux et à en déduire la valeur du courant appliquée pour chaque anneau qui a défilé.

Tant que le courant est assez faible, les deux façons de procéder sont équivalentes. Mais dès lors que le courant dans le solénoïde dépasse l'ampère, il faudra tenir compte de la dilatation thermique qui va contribuer à une variation de longueur de l'échantillon qui n'a rien à voir avec la magnétostriction. Dans notre cas, il semble que l'échantillon se dilate essentiellement radialement, ce qui conduit à une contraction longitudinale. En effet, expérimentalement, on observe que les anneaux défilent dans le même sens que quand on augmente le champ pour de fortes valeurs, avec le « nickel » et avec le « fer ».

Avec la première approche, il faut faire en sorte, dès que l'on observe des anneaux qui défilent en l'absence de variation de champ, de laisser le système prendre une température stable. Pour cela, on attend que les anneaux s'arrêtent de défiler. On peut alors augmenter à nouveau le champ et compter les anneaux qui défilent. Très vite, la variation de température va poser problème et il faudra à nouveau attendre que la température se stabilise. On recommence cette démarche jusqu'à la valeur de champ souhaitée. Le principal défaut de cette méthode, c'est que la température met plusieurs minutes à se stabiliser. Elle est assez fastidieuse.

Avec la seconde approche, en une dizaine de secondes, on a fait défiler tous les anneaux souhaités, et il ne reste qu'à dépouiller le fichier de points acquis pour en déduire la courbe souhaitée. Tant que le courant dans le solénoïde ne prend pas de valeur trop élevée, on peut supposer que la dilatation provoquée par l'élévation de température n'a pas trop d'effet sur le nombre d'anneaux qui ont défilé. Ceci étant, il faut s'attendre à avoir de plus en plus d'erreur due à la dilatation quand on va augmenter le champ.

Expérimentalement, nous avons testé le dispositif avec un barreau de fer Armco et un barreau de nickel. Pour le barreau de « nickel », on trouve bien une contraction de l'échantillon. En revanche, la contraction relative est beaucoup plus faible que ce qui est attendu pour le nickel pur. Le barreau employé n'est donc pas constitué de nickel pur.



Pour le barreau de fer Armco, on observe un allongement de l'échantillon pour les faibles valeurs de champ et une contraction pour les valeurs plus importantes, conformément à ce qui est attendu pour le fer pur. Le maximum d'allongement a bien lieu pour la valeur de champ attendue. De même, avec la mesure à l'œil, on constate que le barreau revient à la longueur initiale pour la bonne valeur de champ.

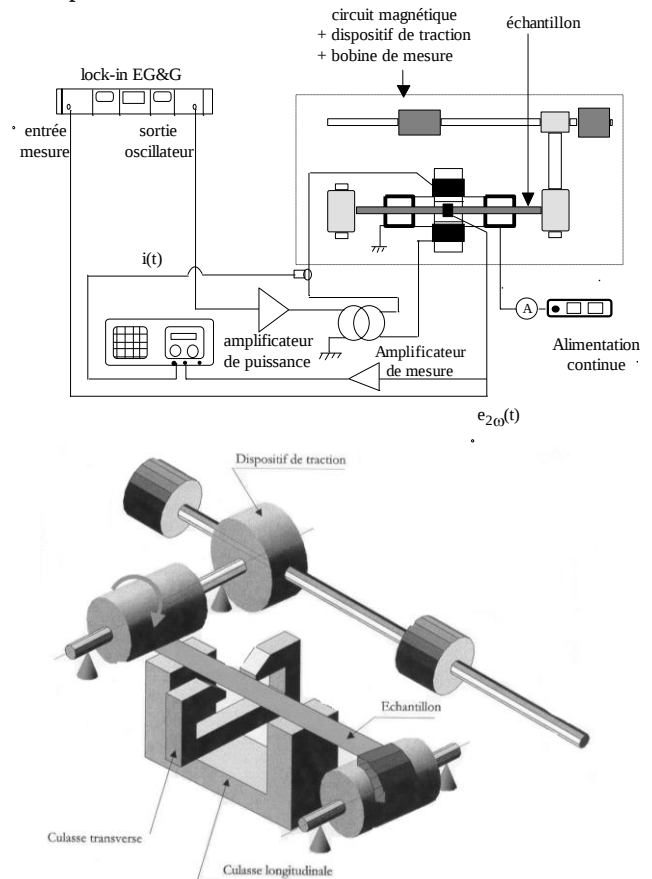
Pour les deux matériaux, quand le courant est important, on constate que la longueur obtenue est toujours plus grande en travaillant avec la photodiode. A l'œil, l'effet de contraction lié à la dilatation radiale semble plus important, alors qu'à la photodiode, cet effet semble pouvoir être négligé compte tenu de la rapidité de la prise de mesure devant les constantes de temps thermiques. Cette hypothèse risque d'être de plus en plus difficile à accepter lorsque le courant va augmenter. Il faudra réaliser la montée de courant de plus en plus vite pour ne pas laisser aux effets thermiques le temps de se manifester.

Avec un système mécaniquement plus perfectionné, on pourrait travailler avec un courant à 50 Hz et relever la réponse d'une photodiode en fonction du courant dans le solénoïde. Ainsi, la mesure s'effectuerait sur un quart de période du secteur, soit 5 ms, ce qui ne laisserait pas le temps aux phénomènes thermiques de poser problème. Par ailleurs, en attendant assez longtemps, le système, thermiquement sensible à la valeur efficace du courant, finirait par atteindre une température fixée, pour un courant maximum donné et prendrait donc une longueur moyenne fixée.

2.2. Mesure sur des rubans : technique S.A.M.R.

a/ Structure générale du dispositif

Ce dispositif comporte un circuit magnétique longitudinal (culasse simple et échantillon) dimensionné pour pouvoir saturer l'échantillon en champ continu. Il se compose également d'une culasse transversale qui permet de soumettre l'échantillon à une composante de champ transverse alternative (de pulsation ω). Le ruban est positionné dans un dispositif de traction. Une bobine de mesure entoure l'échantillon au centre du dispositif. Elle recueille une tension liée à la variation du flux longitudinal de l'échantillon $[Alv1][Nari]\{MCHou\}$. Le dispositif complet avec les alimentations et les appareils de visualisation et de mesure est représenté sur la figure suivante. On a également détaillé le circuit magnétique et le dispositif de traction.



b/ Bilan énergétique

Dans le dispositif S.A.M.R., le champ extérieur est formé par la somme d'une composante longitudinale continue et d'une composante transversale alternative. Ces composantes de champ sont respectivement nommées $H_{//}$ et $H_{\perp}$. La contrainte appliquée au ruban est appelée σ et M_s représente l'aimantation à saturation du matériau considéré.

L'énergie interne du ruban est la somme de l'énergie d'échange, de l'énergie magnétostatique et de l'énergie magnéto-élastique (on peut négliger l'énergie magnéto-

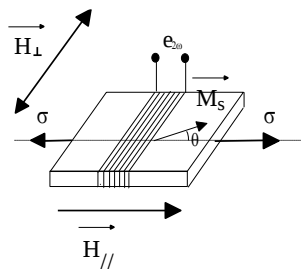
cristalline pour un ruban amorphe ou nanocristallin). A partir de la figure suivante, on peut développer les différentes énergies (cf. Annexe).

En comptant positives les énergies fournies au système, on trouve que

$$E = [E_{ech}] + \left[\frac{3}{2} \lambda_S \cdot \sigma \cdot \cos^2 \theta \right] + [\mu_0 \cdot M_S \cdot H_{//} \cdot \cos \theta + \mu_0 \cdot M_S \cdot H_{\perp} \cdot \sin \theta] - \frac{\mu_0 \cdot M_S^2}{2} \cdot (N_{//} \cdot \cos^2 \theta + N_{\perp} \cdot \sin^2 \theta)$$

Avec le dispositif S.A.M.R., la composante transversale de champ a une amplitude faible par rapport à la composante longitudinale. Le matériau étant considéré comme isotrope, on peut supposer que le champ et l'aimantation ont des directions quasiment identiques. Par conséquent, on peut supposer que θ a une valeur très faible, qui justifie un développement limité au second ordre des termes en θ . Cela nous conduit à :

$$E = E_{ech} + \frac{3}{2} \lambda_S \cdot \sigma \cdot (1 - \theta^2) + \mu_0 \cdot M_S \cdot H_{//} \cdot (1 - \frac{\theta^2}{2}) + \mu_0 \cdot M_S \cdot H_{\perp} \cdot \theta - \frac{\mu_0 \cdot M_S^2}{2} \cdot (N_{//} \cdot (1 - \theta^2) + N_{\perp} \cdot \theta^2)$$



c/ Force électromotrice induite dans la bobine de flux

Aux bornes de la bobine de flux placée sur l'échantillon, on récupère la tension $e_{2\omega}$. Pour une bobine de N spires et un échantillon de section S, cette dernière a pour valeur :

$$e_{2\omega} = -\mu_0 \cdot N \cdot M_S \cdot S \cdot \sin \theta \cdot \frac{d\theta}{dt} \approx -\mu_0 \cdot N \cdot M_S \cdot S \cdot \theta \cdot \frac{d\theta}{dt}$$

La valeur de θ à considérer est celle qui correspond à un minimum d'énergie interne soit :

$$\frac{dE}{d\theta} = 0$$

Ce qui entraîne :

$$\theta = \frac{\mu_0 \cdot M_S \cdot H_{\perp}}{\mu_0 \cdot M_S \cdot H_{//} + 3 \lambda_S \cdot \sigma + \mu_0 \cdot M_S^2 \cdot (N_{\perp} - N_{//})}$$

Sachant que la composante transversale du champ est de la forme :

$$H_{\perp} = H_{\perp \max} \cdot \sin(\omega t)$$

On en déduit que :

$$e_{2\omega} = \frac{-\mu_0 \cdot N \cdot M_S \cdot S \cdot \omega}{2} \cdot \frac{(H_{\perp \max})^2}{(H_{//} + \frac{3 \lambda_S \cdot \sigma}{\mu_0 \cdot M_S} + H_d)^2} \cdot \sin(2\omega t)$$

avec $H_d = (N_{\perp} - N_{//}) \cdot M_S$

On remarque que la tension recueillie a pour pulsation 2ω alors que la composante transversale de champ, à l'origine de la variation de flux, a pour pulsation ω .

d/ Principe de la mesure du coefficient de magnétostriction à saturation

Il faut d'abord déterminer expérimentalement la relation entre le champ et le courant dans la direction longitudinale. A titre d'exemple, dans le dispositif employé, on a $H_{//} = 350 \cdot I_{//}$, ce qui est suffisant pour saturer les échantillons, sachant que l'on peut imposer un courant de 10A [Alv1]{MCHou}. Pour obtenir ce coefficient à partir du dispositif S.A.M.R., on procède de la façon suivante :

- On aimante à saturation l'échantillon dans cette direction et on mesure le courant continu $I_{//}$ (on prend en général $I_{//} = 10A$).

- On impose un champ magnétique transversal alternatif (fréquence de 1kHz environ) dont la valeur efficace est 350 A/m.

- On mesure la tension $e_{2\omega}$ aux bornes de la bobine de flux (on vérifie que le signal reste sinusoïdal de pulsation 2ω en cours de mesure) ; on note cette valeur, que l'on cherchera à conserver dans la suite.

- On soumet le ruban à une variation de contrainte $\Delta\sigma$. Du fait de l'énergie magnétoélastique apportée au ruban, la tension $e_{2\omega}$ varie.

- On agit alors sur $I_{//}$ afin de modifier le champ $H_{//}$ pour ramener $e_{2\omega}$ à sa valeur initiale. Sachant que dans l'expression de $e_{2\omega}$, seuls $H_{//}$ et le terme dépendant de σ ont été modifiés, cela signifie que l'on a :

$$\Delta H_{//} + \frac{3 \lambda_S \cdot \Delta\sigma}{\mu_0 \cdot M_S} = 0$$

De cette expression, on déduit le coefficient de magnétostriction à saturation, soit :

$$\lambda_S = -\frac{\mu_0 \cdot M_S}{3} \cdot \frac{\Delta H_{//}}{\Delta\sigma}$$

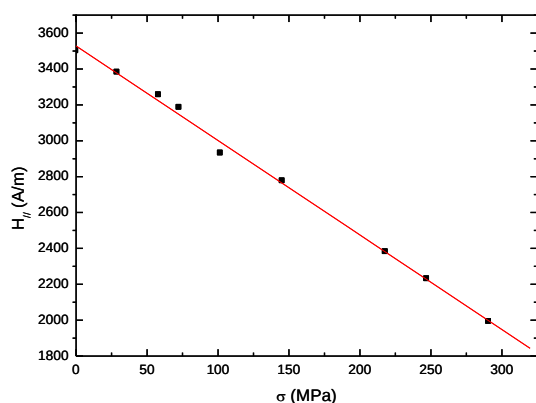
On constate que si le coefficient de magnétostriction à saturation est positif, une augmentation de la contrainte appliquée impose une diminution de champ $H_{//}$ pour maintenir $e_{2\omega}$ constante.

e/ Réalisation pratique de la mesure et limites du dispositif

Pour mesurer précisément le coefficient de magnétostriction à saturation, on installe l'échantillon dans le dispositif de traction sous faible contrainte (quelques MPa). Puis, on augmente la contrainte par pas réguliers en ajustant progressivement $I_{//}$ pour maintenir $e_{2\omega}$ constant. Durant toute la manipulation, on vérifie que $e_{2\omega}$ reste bien sinusoïdale de pulsation 2ω . On trace alors $H_{//}$ en fonction de la contrainte. On obtient une droite dont la pente nous indique le coefficient de magnétostriction à saturation.

A titre d'exemple, la figure suivante présente la relation entre $H_{//}$ et σ pour un matériau nanocristallin. De cette courbe, on peut déduire que la magnétostriction à saturation λ_s vaut 2.1 ± 0.2 ppm sachant que l'induction à saturation dans ce matériau vaut 1,25 T. Ce dispositif permet de mesurer des coefficients supérieurs à 10^{-7} (en deçà, une mesure précise nécessite une valeur $\Delta H_{//}$ imposant des contraintes trop importantes pour le banc expérimental utilisé).

La valeur maximale de λ_s mesurable est limitée par la sensibilité de mesure sur $e_{2\omega}$. Cette f.e.m. peut prendre de très faibles valeurs lorsque l'énergie magnéto-élastique apportée au ruban est trop importante. C'est pourquoi on a recours à une mesure par détection synchrone. L'oscillateur interne de ce dernier permet de générer le champ transverse de fréquence f par l'intermédiaire d'un amplificateur de puissance. On récupère alors très précisément la valeur efficace de $e_{2\omega}$ (fréquence $2f$) grâce à la fonction $2f$ de l'appareil. On peut alors travailler avec des matériaux dont la magnétostriction impose des f.e.m. $e_{2\omega}$ de l'ordre de quelques centaines de μV .



2.3. Mesure de magnétostriction dans des couches minces.

On considère la déflexion d'une poutre constituée d'un substrat (non magnétique) et d'une couche magnétique mince. On suppose que la poutre est sous l'influence d'une contrainte isotrope due à la dilation thermique et au champ d'anisotropie uniaxial (magnétostriction de Joule ou effet Joule longitudinal). Une première formule a été proposée en 1976 par Klockholm donnant la relation entre la déflexion et le coefficient de magnétostriction de la couche, lequel dépend de différents paramètres tels que l'épaisseur de la couche ou le module d'Young. Cependant, les valeurs prédites étaient le double de celles observées, une correction a donc été apportée en prenant en compte le caractère anisotrope de la magnétostriction de Joule [Trem].

$$\lambda_s = \frac{2}{9} \frac{(\Delta_{//} - \Delta_{\perp})}{L^2} \frac{E_s}{E_c} \frac{h_s^2}{h_c} \frac{1 + \nu_c}{1 + \nu_s}$$

Avec:

- $\Delta_{//}$: déflexion mesurée dans la direction du champ magnétique appliqué.
- $\Delta_{\perp}$: déflexion mesurée dans la direction perpendiculaire.
- λ_s : coefficient de magnétostriction à saturation.
- E_s et E_c : respectivement modules d'Young du substrat et de la couche.
- h_s et h_c : respectivement épaisseurs du substrat et de la couche.
- ν_s et ν_c : respectivement coefficients de Poisson du substrat et de la couche.
- L : longueur de la poutre.

Cette équation est valide pour des couches magnétostrictives isotropes suffisamment épaisses pour négliger les effets de surface ($h_c > 50$ nm), déposées sur des substrats plus épais encore ($h_s > 100 h_c$).

Annexe : Energies dans un ruban autre que l'énergie magnéto-élastique.• **Energie d'échange :**

L'énergie d'échange est une énergie résultant des interactions entre atomes. Pour deux atomes isolés (i) et (j), l'expression de l'énergie d'échange est donnée par :

$$E_{ech} = -2J_{ech,ij} \cdot (\vec{S}_i \cdot \vec{S}_j)$$

où $J_{ech,ij}$ est l'intégrale d'échange, S_i et S_j les spins totaux des atomes (i) et (j). Le signe de l'intégrale d'échange définit le magnétisme du matériau. Dans le cas du ferromagnétisme, $J_{ech,ij}$ est positif, ce qui impose à S_i et S_j d'être alignés dans le même sens. Cela permet de minimiser l'énergie d'échange, afin de tendre vers la structure magnétique la plus stable. Au niveau macroscopique, si on ne modifie pas la nature des interactions atomiques, l'énergie d'échange est constante. Elle n'interviendra donc pas dans les variations de l'énergie interne du matériau considéré.

• **Energie magnétostatique :**

Quand on aimante un matériau ferromagnétique, l'énergie magnétostatique dans l'échantillon à la saturation est donnée par :

$$E_{ms} = \mu_0 \int_0^{M_s} \left(\vec{H}_{ext} - \vec{H}_d \right) \cdot d\vec{M} = \mu_0 \int_0^{M_s} \left(\vec{H}_{ext} - N_d \cdot \vec{M} \right) \cdot d\vec{M}$$

où $\vec{H}_{ext}$ représente le champ extérieur appliqué, $\vec{H}_d$ le champ démagnétisant, $\vec{M}$ l'aimantation et N_d est appelé facteur démagnétisant. Le champ démagnétisant représente la réaction du matériau à l'application d'un champ extérieur. Dans le cas d'un échantillon allongé, celui-ci s'aimante plus facilement le long de son grand axe que dans une direction perpendiculaire. Cela provient de la valeur du champ démagnétisant, qui est différente suivant les directions. Pour un ruban, le champ démagnétisant peut être exprimé par la formule suivante :

$$\vec{H}_d = N_{//} \cdot \vec{M}_{//} + N_{\perp} \cdot \vec{M}_{\perp}$$

où $N_{//}$ et $N_{\perp}$ représentent respectivement les facteurs démagnétisants dans les directions longitudinales et transverses, $\vec{M}_{//}$ et $\vec{M}_{\perp}$ respectivement les composantes longitudinales et transverses de l'aimantation.

On en déduit que :

$$E_{ms} = \mu_0 \cdot \vec{H}_{ext} \cdot \vec{M}_s - \mu_0 \int_0^{M_s} \left(N_{//} \cdot \vec{M}_{//} + N_{\perp} \cdot \vec{M}_{\perp} \right) \cdot d\vec{M}$$

Soit :

$$E_{ms} = \mu_0 \cdot \vec{H}_{ext} \cdot \vec{M}_s - \mu_0 \int_0^{M_s} \left(N_{//} \cdot \vec{M}_{//} \cdot d\vec{M}_{//} + N_{\perp} \cdot \vec{M}_{\perp} \cdot d\vec{M}_{\perp} \right)$$

D'où finalement :

$$E_{ms} = \mu_0 \cdot \vec{H}_{ext} \cdot \vec{M}_s - \mu_0 \cdot \frac{M_s^2}{2} \cdot (N_{//} \cdot \cos^2 \theta + N_{\perp} \cdot \sin^2 \theta)$$

Références bibliographiques**Aspects généraux**

[Bri] *Magnétisme et matériaux magnétiques pour l'électrotechnique*, P. Brissonneau, 1997, collection Lavoisier-Hermès.

[Noz] *Ferromagnétisme*, J.P. Nozières, Techniques de l'Ingénieur, E 1730.

[Har] *Effets et matériaux magnétostrictifs*, P. Hartemann, Techniques de l'ingénieur, E 1880.

[Pier] *The changing shape of magnetostriction*, A.R. Piercy, University of Brighton, New Materials.

[Trem] *Magnetostriction and internal stresses in thin films: the cantilever method revisited*, Du Trémolet de Lacheisserie E. and Peuzin J.C., J.M.M.M. 136 (1994) 189-196.

Applications

[Fenn] *Terfenol-D driven flaps for helicopter vibration reduction*, R.C. Fenn & al. Smart Mater. Struct. 5 (1996) 49-57.

SAMR

[Nari] *Measurement of saturation magnetostriction of a thin amorphous ribbon by means of small-angle magnetization rotation*, K. Narita, J. Yamasaki and H. Fukunaga, I.E.E.E. Trans. Mag., vol. MAG-16, NO. 2 March 1980.

{MCHou} *Contribution à la mesure du coefficient de magnétostriction de matériaux d'épaisseur micrométrique*, P. Houée, Mémoire CNAM, 1996.

[Alv1] *New design of small-angle magnetization rotation device: evaluation of saturation magnetostriction in wide thin ribbons*, F. Alves, P. Houée, M. Lécivain, F. Mazaleyrat, J.A.P. 81 (8), 15 April 1997.

Sites web

[Etre] <http://www.etrema-usa.com/>

[Ener] <http://www.energeninc.com/> et *Compact magnetostrictive actuators and linear motors*, C.H. Joshi, Actuator 2000, Bremen Germany, June 2000.

[Cedr] <http://www.cedrat.com/> et *Magnetostrictive actuators compared to piezoelectric actuators*, F. Claeysen, N. Lhermet, T. Maillard, ASSET 2002.

[Newl] <http://www.newlandsscientific.plc.uk> et <http://www.soundbug.biz>

[Aud] *Système de mesure de longueur POSICHRON®*. <http://www.audin.fr>

Croissance et structure du niobate de lithium, matériau roi de l'opto-électronique

Michel FERRIOL, Laboratoire Matériaux Optiques, Photonique et Systèmes, UMR-CNRS 7132, Université Paul Verlaine-Metz / Supélec, 2 rue Edouard Belin, 57070 – Metz.

mferriol@univ-metz.fr

Résumé : La phase niobate de lithium existe dans un domaine de composition assez étendu et les cristaux les plus faciles à produire correspondent à une composition d'environ 48,6 mol % en Li_2O pour laquelle la fusion est congruente. Les cristaux obtenus sont donc déficitaires en lithium, ce qui génère l'apparition de défauts structuraux intrinsèques qui influencent fortement leurs propriétés. De par sa structure, le matériau purement stoechiométrique ne peut pas présenter ce type de défauts et possède ainsi des propriétés supérieures. Mais sa croissance est délicate et nécessite des techniques particulières. Le but de cet article est de présenter les différentes méthodes de croissance utilisées actuellement en fonction de la composition des cristaux et des applications recherchées ainsi que les principales caractéristiques structurales de la phase niobate de lithium..

1 Introduction

Plus de 40 000 tonnes de cristaux de niobate de lithium sont produites chaque année. C'est un matériau ferroélectrique doué d'une combinaison unique de propriétés piézo-électriques, non-linéaires et électro-optiques qui lui ont conféré le succès qu'on lui connaît. La majorité des cristaux est destinée à la production de dispositifs SAW (Surface Acoustic Waves) utilisés surtout dans les téléphones mobiles et les récepteurs radio des téléviseurs. Cependant, les applications en optique non-linéaire ne cessent de se développer essentiellement dans les télécommunications. Le niobate de lithium a été également un des premiers matériaux à montrer de la photoréfractivité, phénomène mis à profit dans le traitement des images holographiques et le stockage de l'information. Ses propriétés assez exceptionnelles, sa disponibilité et son usage très répandu l'ont fait surnommé « le silicium de l'optique non-linéaire ».

L'ensemble des propriétés de ce matériau résulte pour une bonne part de ses caractéristiques structurales et dépend de la qualité des cristaux produits qui est également dépendante de la technique de croissance utilisée. À côté des défauts bien connus des pratiquants de la croissance cristalline : inclusions, bulles, fractures, dislocations..., la qualité s'entend aussi par la présence de défauts intrinsèques (lacunes, atomes en position anormale) et extrinsèques (impuretés chimiques) dont l'influence sur le comportement du matériau est

particulièrement significative dans le cas de LiNbO_3 (seuil d'endommagement laser notamment). Il nous semble donc opportun de présenter au lecteur quelques repères sur la structure des cristaux de niobate et sur leurs conditions de croissance, facteurs clés de leurs performances.

2 Croissance cristalline

La première description du columbate de lithium LiCbO_3 est due à W.H. Zachariassen en 1895 qui remarqua tout de suite sa biréfringence et son indice de réfraction élevé. C'est en 1949 que le columbium (Cb) devint le niobium (Nb) en référence à Niobé, fille du roi Tantale afin de marquer les profondes analogies entre les deux métaux. Cette appellation est due à J.P. Remeika qui fut le premier auteur à signaler, à la même époque, la croissance de monocristaux de LiNbO_3 par une technique de flux et de nucléation spontanée. Les propriétés ferro-électriques des cristaux de niobate de lithium furent mises en évidence pour la première fois à cette occasion.

En 1958, A. Reisman et F. Holtzberg publièrent la première étude expérimentale du diagramme de phases du système binaire $\text{Li}_2\text{O}-\text{Nb}_2\text{O}_5$ qui montra que le niobate de lithium était à fusion congruente (solide et liquide de même composition) à 1253 °C [1].

DIAGRAMME DE PHASES

Le diagramme de phases est une représentation graphique donnant les limites d'existence et de co-existence des différentes phases en équilibre d'un système en fonction des conditions de température, pression et composition. Les équilibres sont régis par la règle des phases de Gibbs : $v = c + 2 - j$ (où v est la variance du système (nombre de degrés de liberté), c est le nombre de constituants indépendants et j est le nombre de phases). Le nombre 2 correspond aux deux variables intensives : pression, température. Dans des conditions isobares (ce qui est le cas en général des diagrammes d'équilibre liquide-solide) : $v = c + 1 - j$. Ainsi, pour un système binaire (2 constituants) à la pression atmosphérique, la représentation se fera en 2 dimensions, pour un système ternaire (3 constituants), la représentation se fera en 3 dimensions, etc...

En 1964, A.A. Ballmann obtint de gros cristaux de LiNbO_3 à partir de sa phase liquide par la technique de Czochralski. Dans la foulée, K. Nassau et coll. étudièrent les propriétés importantes du niobate de lithium et développèrent plusieurs techniques pour produire des cristaux mono-domaine [2, 3]. En 1968, furent publiés des travaux relatifs à l'influence d'un écart à la stoechiométrie sur la température de Curie et la biréfringence de cristaux obtenus en faisant varier en plus ou en moins la composition en Li_2O du bain de tirage [4].

Les résultats obtenus, suggérant l'existence d'un domaine de solution solide au voisinage de LiNbO_3 (en contradiction avec le diagramme de Reisman et Holtzberg), nécessitèrent un réexamen de celui-ci. P. Lerner et coll. [5] ont ainsi montré qu'il y avait un étroit domaine de solution solide au voisinage de la stoechiométrie et également que la congruence n'était obtenue que pour une composition déficitaire en lithium entraînant nécessairement l'apparition de défauts intrinsèques pour les cristaux obtenus à partir de cette composition (figure 1). La détermination exacte du point de congruence a fait l'objet de plusieurs études entre 1968 et 1995 qui situent ce point entre 48,35 et 48,6 mol% en Li_2O , la valeur obtenue dépendant de la pureté des réactifs de départ, des conditions de croissance, de la méthode de mesure et aussi de l'orientation du cristal. Les cristaux obtenus à partir de compositions du bain différentes de la composition congruente montrent toujours une évolution de composition le long de l'axe de croissance, conformément à celle du solidus de la phase niobate de lithium. Parmi les propriétés dépendant de la composition et donc, de la teneur en défauts intrinsèques, citons l'indice de réfraction extraordinaire, la température de Curie, la transparence dans l'UV, le

seuil d'endommagement laser, la vitesse de propagation des ondes acoustiques.

TECHNIQUE DE CZOCHRALSKI

La méthode dite « par tirage » mise au point par Czochralski (1918) fait intervenir un germe placé au contact de la matière à cristalliser fondue dans un creuset en platine ou iridium. Le germe est ensuite remonté tirant avec lui du liquide qui refroidit et cristallise. Pendant le tirage, creuset et germe tournent en sens inverse. Cette technique s'adresse aux composés à fusion congruente.

Pour les tirages de composés à fusion non-congruente nécessitant l'utilisation d'un flux (c'est le cas du niobate stoechiométrique), le principe reste sensiblement identique, mais liquide et solide cristallisant n'ayant pas la même composition, la température du liquide n'est pas constante (son évolution est donnée par le diagramme de phases). Le chauffage du creuset doit donc être asservi de façon à maintenir l'équilibre liquide-solide à l'interface de cristallisation.

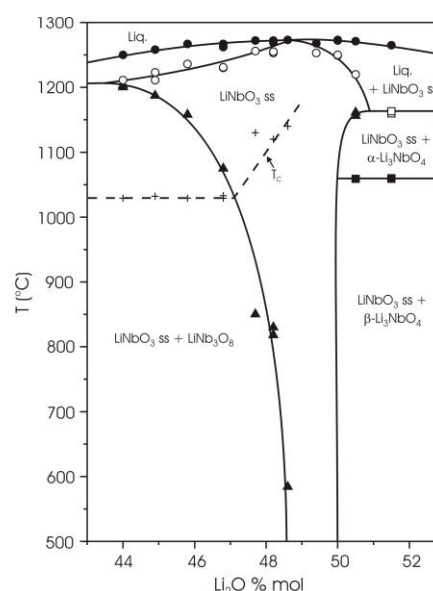


Figure 1 : Diagramme de phases du système $\text{Li}_2\text{O}-\text{Nb}_2\text{O}_5$ au voisinage de LiNbO_3 (A. Dakki, M. Ferriol, M.T. Cohen-Adad, *Eur. J. Solid State Inorg. Chem.*, 33, 19-31 (1996))

Dans ces conditions, des cristaux exactement stoechiométriques (50 mol% en Li_2O) ne présenteront aucun défaut intrinsèque venant grever les performances du matériau. Mais, comme le montre le diagramme de phases de la figure 1, le niobate strictement stoechiométrique n'est pas à fusion congruente et donc, sa croissance est difficile. Plusieurs méthodes ont été envisagées pour préparer de tels cristaux. Dès 1971, J.R. Carruthers et coll. revendiquent la croissance de niobate stoechiométrique à partir d'un bain titrant 58 mol% en Li_2O , mais apparemment de piètre qualité. En 1984, I. Földvári et coll. ont obtenu des cristaux à partir de bains de concentration supérieure à 50 mol% en Li_2O qui montraient toujours une variation de composition le long de l'axe de croissance, nuisible à leurs propriétés. Pour pallier à cet inconvénient, K. Kitamura et coll. ont mis au point la méthode du double creuset à chargement continu dans le but de compenser les pertes en Li_2O par vaporisation et d'assurer une composition constante du liquide et donc, du solide cristallisant (1992). Ces auteurs obtinrent effectivement des cristaux de composition uniforme très proche de la stoechiométrie.

D'autres essais de croissance utilisant l'oxyde de vanadium V_2O_5 comme flux ont été effectués entre 1980 et 1996 avec pour principal inconvénient la solubilisation d'une partie du vanadium (0,2 %) dans la matrice cristalline donnant une légère couleur verte aux cristaux de niobate. Mais les développements les plus récents de la préparation de cristaux massifs de niobate de lithium stoechiométrique trouvent leur origine dans une communication de E.S. Vartanyan faite en 1985 et revendiquant l'élaboration de cristaux stoechiométriques à partir d'un bain de composition congruente contenant de l'oxyde de potassium K_2O dont l'auteur montra l'insolubilité dans le réseau cristallin de LiNbO_3 . Ces résultats furent confirmés par G. Malovichko et coll. (1992) qui obtinrent des cristaux de composition comprise entre 49,95 et 50,00 mol% en Li_2O à partir d'un bain de composition congruente additionné de 6 % de K_2O . En 1997, K. Polgár et coll. ont élaboré des cristaux strictement stoechiométriques par une méthode de flux (« top seeded solution growth ») et ont confirmé l'insolubilité du potassium. Très récemment, l'étude du système ternaire $\text{Li}_2\text{O}-\text{Nb}_2\text{O}_5-\text{K}_2\text{O}$ [6] a permis à un groupe franco-hongrois de définir précisément les conditions de croissance à réaliser pour obtenir des cristaux de niobate purement stoechiométriques à partir d'une solution de Li_2O , Nb_2O_5 et K_2O [7].

Parallèlement à la croissance de cristaux massifs, l'élaboration de fibres monocristallines s'est développée depuis le début des années 1980 en raison du potentiel offert par la forme fibrée en termes d'applications. Deux techniques sont principalement utilisées pour obtenir des cristaux de haute qualité (sans dislocations) avec des diamètres allant de 10 μm à 1 mm :

1/ la méthode dite L.H.P.G. (« Laser Heated Pedestal Growth ») dont l'essor est dû au professeur R. Feigelson à l'Université de Stanford. Dans cette méthode (sans creuset), l'extrémité d'un barreau source fritté et constitué d'une céramique de composition égale à celle du cristal à élaborer, est fondue au moyen d'un laser à CO_2 , focalisé par un dispositif optique adapté. Le cristal est obtenu par tirage d'un germe immergé dans la zone fondue maintenue en équilibre par les forces de tension superficielle (figure 2a).

2/ la méthode dite μ -PD (« micro-pulling down ») développée au Japon par le professeur T. Fukuda. Dans ce cas, la fibre monocristalline est obtenue par cristallisation d'un liquide de composition adéquate s'écoulant d'un creuset (généralement chauffé par résistivité) par un capillaire de diamètre adapté à celui du cristal souhaité (figure 2b).

L'une des caractéristiques de ces deux techniques est l'existence de forts gradients thermiques axiaux (quelques centaines de degrés par millimètre) autorisant ainsi des vitesses de tirage élevées (jusqu'à 20 cm.h^{-1}) et des mécanismes de croissance basés sur un effet de trempe. Dans ces conditions, les coefficients de distribution (rapport des concentrations dans le solide et le liquide) effectifs des différentes espèces chimiques sont pratiquement égaux à l'unité assurant une composition homogène du cristal le long de l'axe de tirage et égale à celle du liquide de départ. Cela permet d'envisager le tirage de cristaux de niobate très enrichis en lithium, voire même stoechiométriques. Des essais de tirage ont ainsi été réalisés par L.H.P.G. et ont conduit à des cristaux quasi-stoechiométriques de composition comprise entre 49,50 et 49,85 mol% en Li_2O selon le degré d'enrichissement en lithium des barreaux source. L'écart à 50 %, même avec une teneur en Li_2O des barreaux source égale à 58 mol%, ne peut être qu'attribué à la vaporisation de l'oxyde de lithium en raison d'un rapport surface/volume de la zone fondue très important et peu favorable dans ce cas.

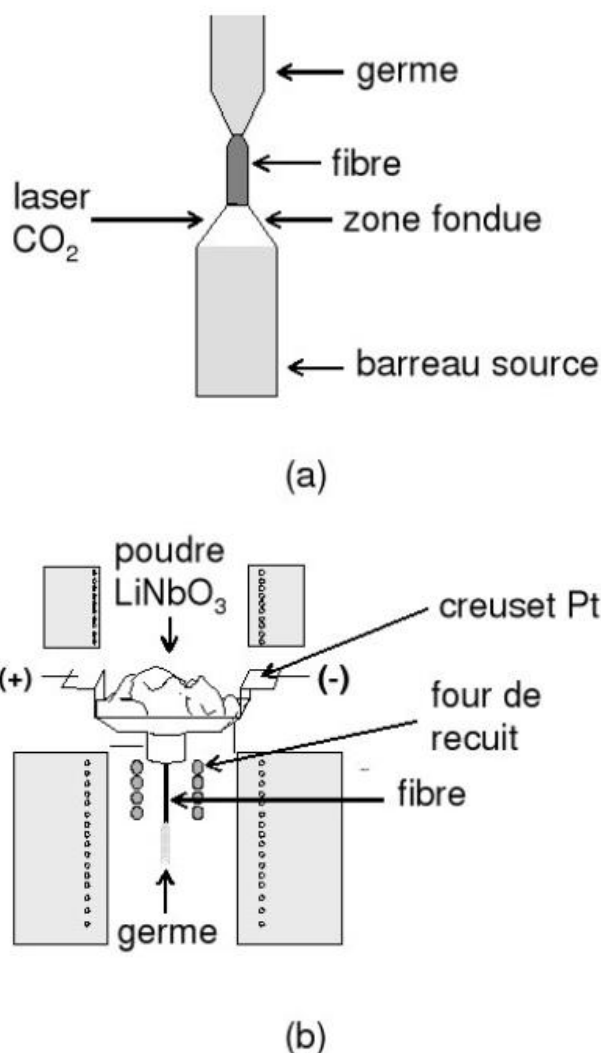


Figure 2 : Schémas des techniques de croissance de fibres de LiNbO_3 ((a) : L.H.P.G., (b) : μ -PD)

Les gradients thermiques axiaux mis en jeu lors de la croissance des fibres génèrent également un champ électrique de l'ordre de 1 V.cm^{-1} largement supérieur à celui utilisé pour le poling des cristaux massifs obtenus par Czochralski. Dans ces conditions, contrairement à une organisation aléatoire des domaines ferroélectriques observée sur des cristaux Czochralski orientés selon l'axe cristallographique, les fibres tirées suivant cette orientation présentent une structure bi-domaine, la frontière des deux domaines étant parallèle à la face c du cristal et correspondant à l'axe de la fibre. Cette structure particulière pourrait être mise à profit dans la fabrication de Q-switches pour lasers à fibres ou de convertisseurs de modes.

Une autre technique peut être exploitée pour obtenir des fibres cristallines de niobate stoechiométrique : la méthode V.T.E. (« Vapour Transport Equilibration »). Il s'agit de faire diffuser à température élevée de l'oxyde de lithium, à partir d'une source riche en lithium (en général, un mélange $\text{LiNbO}_3/\text{Li}_3\text{NbO}_4$ sous forme de

poudre), dans le cristal congruent enrobé par cette poudre. Basé sur les phénomènes de diffusion, le traitement est long et ne peut donc être appliqué qu'à des échantillons de faible épaisseur si l'on veut obtenir une composition homogène en lithium en un temps raisonnable.

3 Structure cristallographique

À la température ambiante, la structure du niobate de lithium est une cousine assez éloignée de la structure pérovskite bien connue. Elle dérive en fait de la structure ilménite et est de symétrie rhomboédrique (groupe d'espace $R3c$). Elle a été résolue par S.C. Abrahams et coll. au milieu des années 1960, à la fois par diffraction des rayons X et diffraction de neutrons. À cette époque, les auteurs n'avaient pas connaissance de l'existence d'une solution solide au voisinage de la stoechiométrie et il a fallu attendre 1984 pour que Abrahams et Marsh déterminent les paramètres cristallins exacts du niobate stoechiométrique et du niobate congruent [8]. Comme nous l'avons déjà indiqué, le solide cristallin est en général déficitaire en lithium. Ce déficit se manifeste par des défauts répartis aléatoirement sans perturbation des relations de symétrie du réseau. Ainsi, la structure cristalline peut toujours être traitée comme si le cristal était stoechiométrique.

Décrite dans le système hexagonal, la structure (figure 3) est constituée d'un empilement, le long de l'axe cristallographique (qui est également l'axe polaire), d'octaèdres déformés d'anions oxydes formant des plans équidistants. Les cations prennent place dans ces octaèdres selon la séquence : $\text{Li}^+ - \text{Nb}^{5+} - \text{site inoccupé} - \text{Li}^+ - \text{Nb}^{5+} - \text{site inoccupé}$, etc.... Dans la phase ferro-électrique, les cations lithium et niobium occupent une position décalée par rapport au centre des octaèdres, cause de la polarisation observée. Les paramètres de maille sont les suivants [8] :

- composé à fusion congruente :

$$a=5,15052 \text{ \AA}, c=13,86496 \text{ \AA} (\rho=4,648 \text{ g.cm}^{-3})$$

- composé stoechiométrique :

$$a=5,14739 \text{ \AA}, c=13,85614 \text{ \AA} (\rho=4,635 \text{ g.cm}^{-3})$$

La maille élémentaire hexagonale contient six formules LiNbO_3 .

Plusieurs études ont montré que dans le composé à fusion congruente, le taux d'occupation des sites niobium est de 95,3 %, celui des sites lithium est de 94,1 % en Li^+ et de 5,9 % en Nb^{5+} . Ainsi, chaque cation Li^+ manquant est remplacé par un cation Nb^{5+} . La charge excédentaire est compensée par la création de lacunes.

Nombre de travaux ont été consacrés à cette structure lacunaire, tant il est vrai qu'elle influence fortement les propriétés du matériau. Trois possibilités ont été proposées :

- lacunes d'oxygène
- lacunes de niobium
- lacunes de lithium

Les modèles à lacunes de niobium et de lithium semblent les plus satisfaisants, mais ces différents modèles ne devraient être considérés que comme des cas limites, la réalité se situant vraisemblablement entre les trois.

Quoi qu'il en soit, l'influence des défauts intrinsèques est bien illustrée, par exemple, par la valeur du champ coercitif qui passe de 240 kV.cm⁻¹ pour le composé à fusion congruente à 80 kV.cm⁻¹ pour le composé stoechiométrique. Le problème de la structure se complique lorsque l'on introduit des dopants (ions de terres rares, par exemple) dans la matrice cristalline. Ces ions peuvent occuper quatre sites différents : trois sites octaédriques (correspondant aux sites occupés par Li⁺, Nb⁵⁺ et aux sites inoccupés) et un site tétraédrique (figure 4). Les ions dopants n'ont donc pas tous le même environnement et ne possèdent donc pas exactement les mêmes niveaux d'énergie. Ils constituent ainsi des centres optiques différents et de nombreux travaux ont été consacrés à l'étude de la structure multisite du niobate de lithium pour des ions dopants comme : Nd³⁺, Er³⁺, Cr³⁺,...

4 Conclusion

Bien que plusieurs centaines de publications lui soient consacrées chaque année, l'intérêt pour le niobate de lithium ne faiblit pas en raison de la combinaison unique de propriétés qu'il manifeste et de la sensibilité de celles-ci à la nature des défauts qu'il présente. Les conditions d'élaboration (pureté de la matière première de départ, taux de vaporisation du lithium, par exemple) et la méthode de croissance ont aussi une influence directe sur la structure de défauts du matériau et la qualité des cristaux. Il est donc nécessaire de bien les maîtriser si l'on veut obtenir des cristaux aux propriétés bien définies et les plus reproductibles possibles.

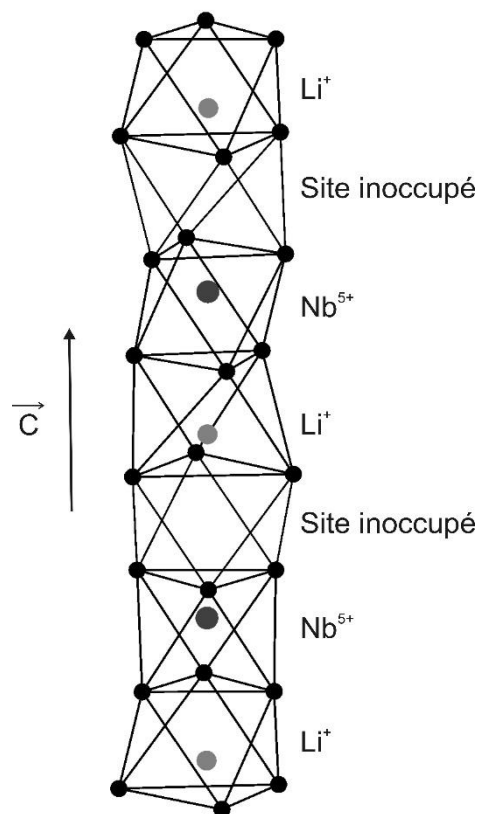
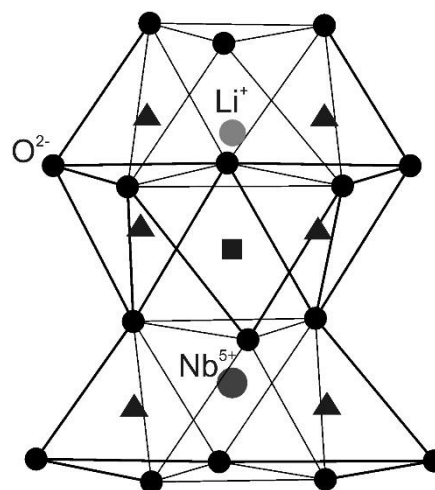


Figure 3 : Structure du niobate de lithium



- ▲ : sites tétraédriques libres
 ■ : site octaédrique libre

Figure 4 : Sites du niobate de lithium

Bibliographie

- [1] A. Reisman et F. Holtzberg, J. Amer. Chem. Soc., 80, 6503 (1958)
- [2] K. Nassau, H.J. Levinstein et G.M. Loiacono, Appl. Phys. Lett., 6, 228 (1965)
- [3] K. Nassau, H.J. Levinstein et G.M. Loiacono, J. Phys. Chem. Solids, 27, 983 et 989 (1966)
- [4] J.G. Bergman, A. Ashkin, A.A. Ballman, J.M. Dziedzic, H.J. Levinstein et R.G. Smith, Appl. Phys. Lett., 12, 92 (1968)
- [5] P. Lerner, C. Legras et J.P. Dumas, J. Cryst. Growth, 3-4, 231 (1968)
- [6] M. Cochez, M. Ferriol, L. Pöppl, K. Polgár et Á. Péter, J. All. Comp., 386, 238 (2005)
- [7] Á. Péter, K. Polgár, M. Ferriol, L. Pöppl, I. Földvári, M. Cochez et Z.S. Szaller, J. All. Comp., 386, 246 (2005)
- [8] S.C. Abrahams et P. Marsh, Acta Cryst. B, 42, 61 (1986)

Biographie

Michel Ferriol est enseignant-chercheur au Département Chimie de l'I.U.T. de Metz. Tout d'abord, Chargé de Recherches au C.N.R.S. à l'Université Claude Bernard Lyon I depuis 1984, il a été recruté comme Professeur à l'Université Paul Verlaine-Metz en 1998. Il effectue ses travaux de recherche au laboratoire Matériaux Optiques, Photonique et Systèmes (U.M.R. C.N.R.S.-U.P.V.M.-Supélec 7132). Il est spécialiste de la thermodynamique des équilibres entre phases et de son application à l'élaboration de matériaux à pureté contrôlée. Au sein de l'opération « Guides et microstructures sur niobate de lithium », ses préoccupations scientifiques sont orientées vers l'élaboration et la caractérisation de fibres monocristallines de niobate de lithium présentant différentes fonctions optiques.

Cristaux photoniques intégrés sur niobate de lithium

Nadège COURJAL, Richard FERRIERE, Maria-Pilar BERNAL

Institut FEMTO-ST, Département d'Optique P.M. Duffieux UMR 6174,
16 rte de Gray, 25030 Besançon Cedex, France
nadege.bodin@univ-fcomte.fr

Résumé : La fabrication de cristaux photoniques LiNbO₃ est très attractive pour l'intégration de circuits optiques de taille micrométrique commandables par effet acousto-optique, non-linéaire ou électro-optique. Pour atteindre cette perspective, une démonstration expérimentale de faisabilité par caractérisation optique est impérative. Une étude théorique préalable nous permet de montrer qu'un arrangement triangulaire de trous peut assurer un contrôle efficace de la transmission, par variation de l'indice de réfraction. Nous décrivons ensuite comment une gravure par faisceau ionique focalisé peut mener à la performance d'une matrice de 21*19 trous larges de 213nm et profonds de 1,5mm. La caractérisation optique de cette matrice est la première effectuée sur un cristal photonique LiNbO₃ : elle confirme la faisabilité de ces structures en exhibant un creux de -12dB dans le spectre en transmission. Finalement, nous proposons des solutions vers la mise en œuvre de méthodes de fabrication collectives par holographie.

1 Introduction

Depuis leur première apparition vers la fin des années 1980, les cristaux photoniques ont suscité de nombreuses recherches, motivées par la perspective de miniaturiser les circuits optiques. Un cristal photonique se présente sous la forme d'une alternance périodique de motifs diélectriques ou métalliques. Lorsque la période est inférieure à la longueur d'onde, la lumière qui traverse la structure est soumise à des interférences : l'énergie lumineuse est ainsi quantifiée en structure de bandes, similaire à celle des semi-conducteur en électronique. Si le motif périodique est bien choisi, il existe une gamme de longueurs d'onde pour laquelle la lumière ne peut pas être transmise. Le contrôle actif de cette bande interdite photonique offre des perspectives très attrayantes pour la fabrication de composants optiques actifs ultra-compacts. Différents matériaux ont déjà été exploités pour fabriquer par exemple des lasers de taille micrométrique [1], ou pour montrer la commande d'une bande interdite par effet thermo-optique, nonlinéaire acousto-optique ou électro-optique [2,3]... Parmi les candidats potentiels pour la fabrication de cristaux photoniques commandables, le niobate de lithium présente un intérêt particulier en raison de ses forts coefficients non linéaires, électro-optique, piézoélectrique et acousto-optique qui ouvrent la voie à des composants à commande multiple. En outre, ce matériau a déjà largement montré son intérêt dans les domaines de la détection ou des télécommunications. Malgré tout l'attrait que peut présenter le niobate de

lithium pour la commande des structures à bande interdites photoniques (BIP), les publications relatives à ce sujet restent rares. Ce phénomène s'explique par la difficulté rencontrée pour graver le matériau à des dimensions nanométriques avec un facteur de forme supérieur à 1. A ce jour, seules deux équipes sont parvenues à graver des trous de quelques centaines de nanomètres avec une profondeur de quelques microns dans le matériau [4,5]. Ces deux équipes sont celle de Yin et al. [4] et la notre [5], qui exploitent toutes deux la même méthode : la gravure directe par faisceau ionique focalisé (FIB en anglais : Focused Ion Beam). Dans ce papier, nous montrons comment cette technique peut être mise à profit pour la fabrication et la caractérisation de cristaux photoniques 2D. Nous montrons finalement comment une fabrication collective peut être envisagée, comme alternative au FIB.

Etude théorique : les cristaux photoniques LiNbO₃ commandables

Notre objectif est de contrôler la transmission lumineuse au travers d'un circuit optique intégré de taille micrométrique. Pour réaliser cet objectif, nous nous appuyons sur une structure périodique bidimensionnelle (2D) qui est constituée par une répétition périodique de trous ou de plots dans un arrangement carré ou triangulaire (cf vue schématique des fig.1a et 2a). Dans des conditions de maille bien choisies, une bande interdite pourra apparaître dans le spectre de transmission.

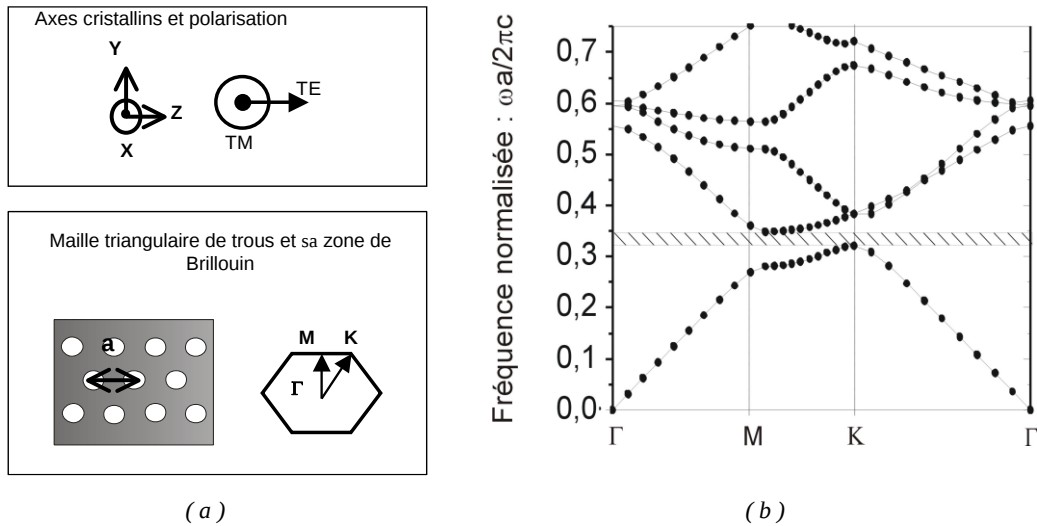


Figure 1 : Structure autorisant une bande interdite photonique pour une propagation TE dans un cristal LiNbO3 en coupe X.
a) Schéma de la structure photonique autorisant la BIP : maille triangulaire de trou et sa zone de Brillouin. En haut : axes et polarisation associés.
b) Diagramme de bande pour la maille triangulaire de trou. (Rapport rayon/période=0.25)

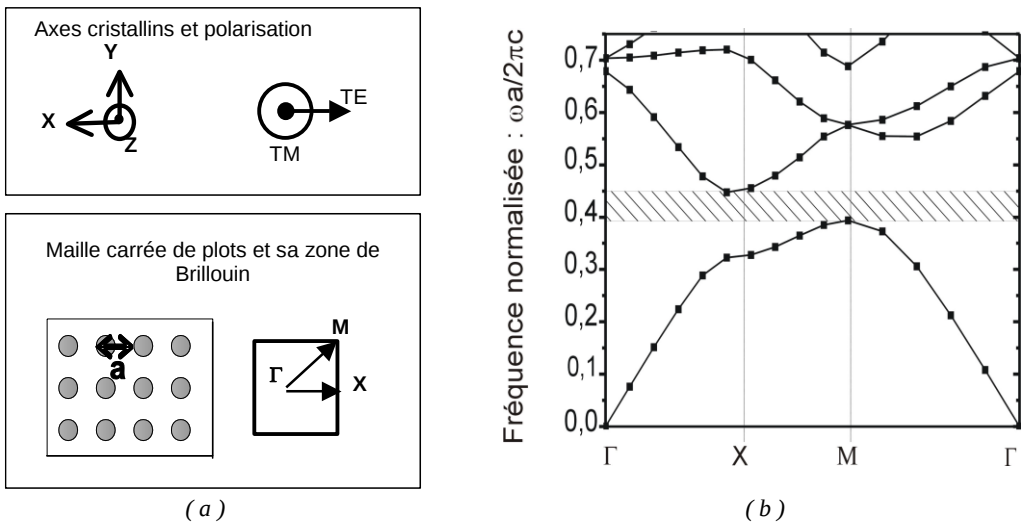


Figure 2 : Structure autorisant une bande interdite photonique pour une propagation TM dans un cristal LiNbO3 en coupe Z.
a) Schéma de la structure photonique autorisant la BIP : maille carrée de plots et sa zone de Brillouin. En haut : axes et polarisation associés.
b) Diagramme de bande pour la maille carrée de plots. (Rapport rayon/période=0.25, $n_e=2.143$)

Pour contrôler la bande interdite, nous nous proposons de modifier l'indice de réfraction du cristal. L'indice de réfraction est en effet un des paramètres régisseurs de la position spectrale de la bande interdite. Pour modifier l'indice par effet électro-optique, la lumière doit être polarisée suivant l'axe extraordinaire du cristal (axe Z), car cette direction de polarisation est celle qui subit la plus forte interaction entre champ optique et champ électrique pour les cristaux LiNbO3. Cette contrainte de polarisation ne permet de sélectionner que deux types de propagation possibles. La première configuration repose sur l'utilisation d'un cristal LiNbO3 en coupe X. Dans les cristaux en coupe X, l'axe Z du cristal est parallèle au plan du substrat : le champ électrique doit donc se

propager avec une polarisation parallèle au plan. Ce type de polarisation est nommé TE (Transverse Electrique). La deuxième configuration possible utilise un cristal LiNbO3 en coupe Z : l'axe Z du cristal est dans ce cas perpendiculaire au plan du substrat. Pour bénéficier d'une interaction électro-optique forte, on s'arrangera donc pour polariser le champ électrique perpendiculairement au plan du substrat. On dit alors que la polarisation est TM (Transverse Magnétique) car c'est le champ magnétique qui est parallèle au plan du substrat. Les figures 1a et 2a illustrent les axes cristallins associés pour les deux conditions de propagation sélectionnées.

Tout d'abord, nous avons déterminé quelles sont les configurations susceptibles d'assurer une bande interdite photonique (BIP). Le calcul de la structure de bande a été effectué par la méthode des ondes planes à l'aide du logiciel Bandsolve. Il nous a permis d'isoler les deux seuls arrangements périodiques qui présentent une BIP dans leur structure de bande : la maille triangulaire de trous lorsque la polarisation est TE, et la maille carrée de plots lorsque la polarisation est TM. La représentation schématique de ces deux mailles avec les principales directions de propagation sont indiquées en figures 1a et 2a. Les diagrammes de bande associés sont représentés en figures 1b et 2b. Ils illustrent la dispersion des différents modes de propagation : l'ordonnée correspond à la fréquence normalisée $\omega a/2\pi c$, où ω est la pulsation de l'onde optique, a la période du cristal photonique et c la vitesse de la lumière dans le vide. L'abscisse du diagramme de bande représente le vecteur d'onde pour différentes directions de propagation. La bande interdite est indiquée par la partie hachée sur les figures. Il existe un rayon minimal des motifs en deçà duquel la bande interdite disparaît : ce rayon minimal, déterminé par la méthode des ondes planes, vaut respectivement 0.17 pour la maille triangulaire et 0.15 pour la maille carrée de plots.

Direction	Largeur de la BIP δ (nm)	Minimum de transmission $T_{\min}$ (%)	Longueur d'onde au centre de la BIP	Décalage $\Delta\lambda$ de la BIP pour $\Delta n=1$
GM	316	3.7e-7	1372	696 nm
GK	233	3.4e-2	1211	600 nm

Tableau 1 : Position théorique de la BIP et transmission minimale au centre de la BIP, pour les directions de propagation GK et GM du cristal photonique arrangé en maille triangulaire (polarisation TE). La colonne de droite indique le décalage $\Delta\lambda$ en nm de la BIP lorsque la variation d'indice Δn du substrat est égale à 1. Les calculs sont effectués par ftd 2D pour 21 rangées de trous, de rayon $0.25*a$, et de période $a=425\text{nm}$.

La deuxième étape de calcul a consisté à déterminer quelle direction de propagation est la plus sensible à la variation d'indice. Pour des raisons pratiques, nous nous sommes principalement focalisés sur la maille triangulaire de trous. Nous avons choisi de travailler avec un rayon égal à $0.25*a$. Ce choix est motivé par des raisons technologiques : si le facteur de remplissage est trop élevé, les murs entre trous risquent de s'effondrer. Par ailleurs, la période a est prise arbitrairement comme égale à 425nm pour les calculs :

le spectre en transmission pour une autre période pourra toujours être déduit par un simple produit en croix. Nous avons évalué par la méthode FDTD2D (méthode des différences finies bidimensionnelle) quelle est la direction de propagation la plus favorable à un contrôle de la transmission par variation de l'indice de réfraction. Les résultats sont exposés dans le tableau 1 : la direction GM est plus sensible à une variation d'indice que la direction GK. En effet, nous prévoyons un décalage de 7nm de la bande interdite si la variation d'indice est de 0.01, ce qui correspond à 1nm de plus que le décalage obtenu pour la direction GK de propagation avec la même variation d'indice. Pour tous les calculs, nous avons choisi 21 rangées de trous, car ce nombre correspond au minimum de rangées permettant d'atteindre une transmission minimale au centre de la BIP.

2 Fabrication et caractérisation :

Nous avons ensuite fabriqué les structures photoniques à partir des prévisions théoriques ci-dessus. Pour que le bord droit du gap (creux) coïncide avec 1550nm lorsque la direction de propagation est GM, nous définissons les motifs avec une période de 425nm, et un diamètre de 123nm. Le substrat est un cristal LiNbO3 en coupe X, de 300mm d'épaisseur. Pour confiner la lumière verticalement, nous procédons à la fabrication d'un guide optique classique par échange protonique suivi d'un recuit (APE : annealed proton exchange) : les détails de fabrication sont précisés en référence [5]. Les guides ainsi réalisés présentent un gradient de l'indice extraordinaire de 2.138 dans le substrat à 2.143 en surface pour $l=1.55\text{mm}$. Le guide est ainsi monomode à 1.55mm, et enterré à 1.4mm de la surface, avec une largeur à mi-hauteur de 4mm. L'enterrement du guide est contraignant, car les trous devront atteindre une profondeur supérieure 1.4mm pour une largeur de 213nm, ce qui implique un facteur de forme supérieur à 7. Il faut néanmoins remarquer que ces conditions de guidage représentent une amélioration en comparaison des guides classiquement réalisés sur niobate de lithium, qui peuvent être enterrés jusqu'à 4mm de profondeur par rapport à la surface : les conditions de fabrication ont été étudiées ici de façon à rapprocher les guides de la surface tout en préservant les propriétés électro-optique du matériau, ainsi que le caractère monomode du guide.

La structure photonique est ensuite gravée dans le guide optique à l'aide d'un faisceau ionique focalisé (FIB) [5]. Cette étape consiste à bombarder localement le substrat par des ions Ga^+ . La gravure est assistée par ordinateur, avec une résolution de 70nm. Par un choix approprié des vitesses d'accélération et temps de gravure, nous avons réussi à définir des trous larges de

quelques centaines de nanomètres avec une profondeur de 2 μm . La figure 3 montre une des matrices usinées par FIB sur un guide optique (visualisation au microscope à balayage électronique (MEB)).

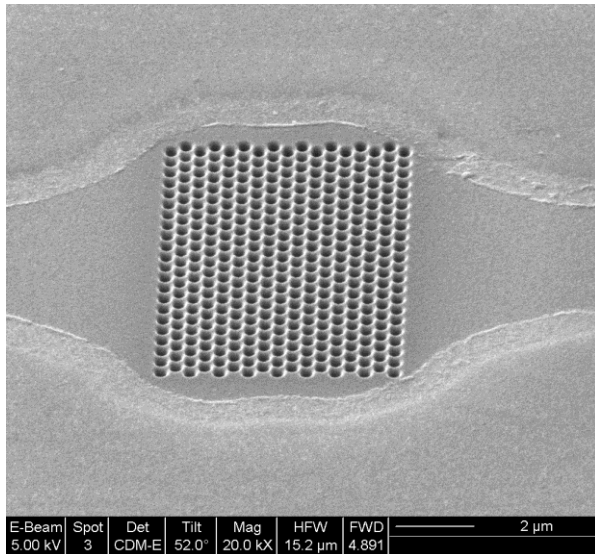


Figure 3 : Vue inclinée au microscope à balayage électronique d'une matrice de trous triangulaire usinée par FIB avec une période de 426nm. (tilt de 52°C°). Les larges traits clairs correspondent au bord du guide optique, qui est élargi au centre pour accueillir la structure photonique.

La matrice correspond à 21 x 19 rangées de trous en arrangement triangulaire. Le temps de gravure de cette matrice est de 20minutes et la profondeur des trous est estimée à 1.5 μm . Ces résultats de gravure, qui peuvent sembler triviaux en comparaison des résultats obtenus sur d'autres matériaux, sont pourtant remarquables car le niobate de lithium est très difficile à usiner aux dimensions nanométriques avec un facteur de forme supérieur à 1.

Le dispositif expérimental pour la caractérisation en transmission des structures photoniques est schématisé en figure 4.

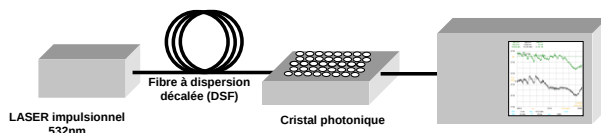


Figure 4 : Dispositif expérimental de mesure du spectre de transmission.

Le spectre en transmission est obtenu en injectant une lumière à large bande spectrale (supercontinuum) dans le guide optique sur lequel est gravée la structure photonique [6]. Le supercontinuum est généré par un microchip laser subnanoseconde émettant à 532 nm (puissance moyenne de sortie : 30 mW, largeur à mi-hauteur : 0.4ns). La lumière en sortie du guide optique micro-usiné est collectée par une fibre et envoyée vers un analyseur de spectre optique. Il est à remarquer

qu'aucun contrôle de polarisation n'est nécessaire car le guide optique APE est polarisant : il sélectionne la polarisation TE.

Les résultats expérimentaux sont représentés en figure 5. Le spectre en transmission au travers du guide avec structure photonique (traits continus) y est comparé au spectre de la transmission au travers d'un guide optique standard réalisé dans les mêmes conditions que le précédent et sur le même substrat, mais sans structure photonique gravée (traits pointillés).

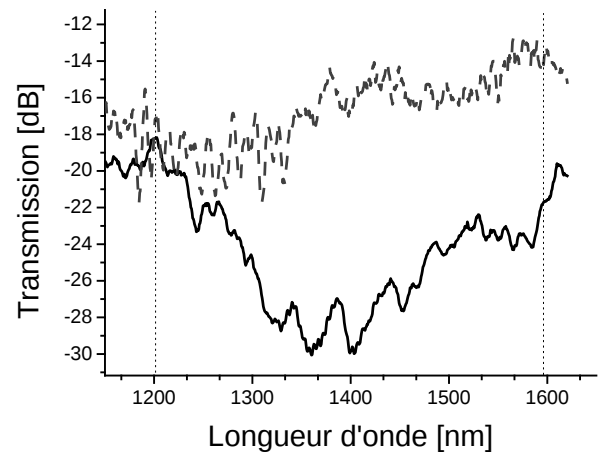


Figure 5 : Spectre en transmission expérimental au travers de deux guides optiques réalisés par échange protonique suivi d'un recuit sur le même substrat. Trait noir continu : spectre au travers d'un guide APE avec structure photonique. Trait pointillé noir : spectre au travers d'un guide APE non usiné.

Le taux d'extinction de la BIP est évalué à -12dB. Ce taux d'extinction est difficile à évaluer précisément, en raison du bruit observé dans la transmission, résultant principalement du caractère multimode de la fibre et du guide optique sur cette gamme de longueur d'ondes. Le creux en transmission s'étale de 1300nm à 1600nm, ce qui correspond bien aux prédictions théoriques qui peuvent être déduites du tableau 1.

Le spectre en transmission présente des bords rugueux qui seront difficilement exploitables pour une commande efficace de la BIP. Son intérêt est néanmoins non négligeable, car c'est la première caractérisation optique d'une structure BIP sur niobate de lithium. Ce résultat confirme la possibilité de fabriquer à terme des structures optiques ultra-compactes commandables.

Pour viser la production collective, il faut cependant s'affranchir de l'utilisation du FIB, qui requiert des temps de gravure rédhibitoires. Le paragraphe suivant est destiné à montrer des solutions alternatives de définition des motifs submicroniques.

3 Fabrication collective par holographie

L'inscription de structures photoniques dans des guides d'onde nécessite des moyens de fabrication collective dans une perspective de production industrielle à bas coût. Compte tenu de la dimension des motifs, les procédés de photolithographie classiques avec usage d'une résine photosensible s'avèrent difficiles à employer. En effet, les dimensions des motifs de l'ordre de $0.3\ \mu\text{m}$ et l'épaisseur de la résine photosensible ($0.5\ \mu\text{m}$) s'approchent de la valeur de la longueur d'onde utilisée par l'aligneur de masque pour l'insolation. Il en résulte que la technique classique de masquage par ombroscopie fait apparaître des phénomènes de diffraction à distance finie très difficiles à modéliser mais qui cependant doivent être pris en compte lors de la conception du masque. Pour pallier ces difficultés, nous avons développé une technique d'insolation des résines basée sur un principe holographique.

L'idée de base pour la réalisation du motif périodique du cristal photonique est fondée sur l'utilisation d'un système de franges d'interférences tridimensionnel dont le pas puisse être ajusté de manière très précise et dont la structure soit conservée dans l'épaisseur de la résine photosensible en évitant les phénomènes de diffraction parasites. L'objectif est de permettre la fabrication d'un masque sur la surface (on chip) du guide d'onde optique qui permette de procéder ensuite directement à la gravure du niobate de lithium [7].

Le dispositif étudié offre une grande souplesse au niveau des réglages et de la polyvalence d'utilisation tout en demeurant simple, précis et stable au niveau des réglages. Le principe du dispositif est représenté en figure 6. Il est basé sur la formation de l'image d'une source ponctuelle S à travers un interféromètre triangulaire Bs - $M1$ - $M2$. Après division d'amplitude par le cube séparateur Bs , la lumière parcourt les bras de l'interféromètre dans les deux sens opposés et en suivant des trajets de longueurs rigoureusement égales. La différence de phase entre les deux bras de l'interféromètre est nulle quelle que soit la position respective des miroirs $M1$ et $M2$. Cette propriété assure une quasi-invulnérabilité du dispositif par rapport aux vibrations extérieures qui seraient susceptibles de brouiller le champ de franges d'interférences et ceci sans qu'aucun asservissement sur le contrôle de phase ne soit nécessaire.

Lorsque le miroir $M2$ est déplacé linéairement de D la source S est dédoublée en deux sources jumelles, parfaitement synchrones, et translatées latéralement et symétriquement par rapport à l'axe optique (sources d'Young). Les deux sources sont placées dans le plan focal objet de la lentille $L2$ de longueur focale f_2 et le phénomène d'interférence est enregistré dans le plan focal objet. Le pas i des franges d'interférence en fonction du déplacement D est donné par la relation :

$$i(\Delta) = \lambda / \left(2 \sin \left(\arctan \left(\frac{\Delta \sqrt{2}}{2 f_2} \right) \right) \right)$$

Le miroir $M2$ est monté sur une platine de translation linéaire de précision $0.1\ \mu\text{m}$ ce qui permet d'atteindre une précision théorique de $0.5\ \text{nm}$ sur le pas des franges d'interférence. Le pas minimum des franges d'interférences est déterminé par les dimensions des composants optiques. En outre la visibilité du phénomène d'interférence résultant de la superposition de deux ondes est égale à 1.

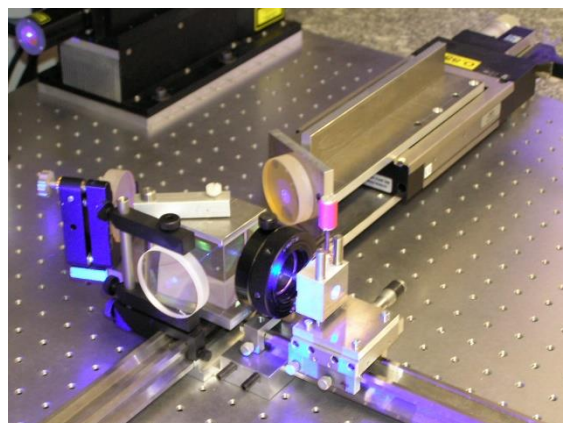
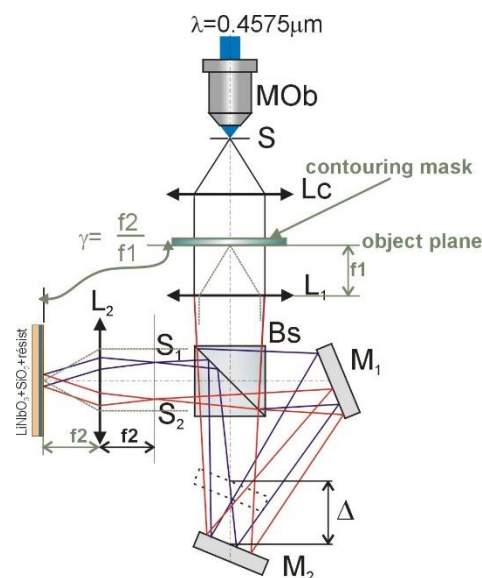


Figure 6 a : Vue schématique et photographique du dispositif holographique

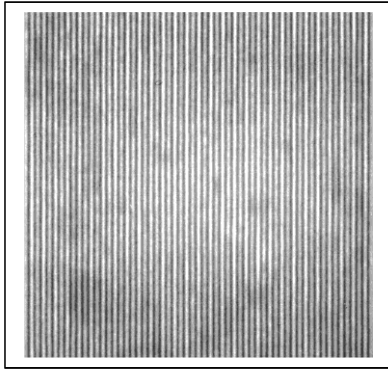


Figure 6 b : Résultat visualisé au microscope (le pas des motifs est de $1.2\ \mu\text{m}$)

Les franges d'interférences obtenues à la sortie du dispositif holographique constituent le motif élémentaire utilisé pour la microstructuration. Leur pas peut être ajusté avec une précision manométrique par déplacement linéaire du miroir M2

Une autre particularité intéressante du dispositif tient au fait que le plan focal objet de la lentille L1 est conjugué avec le plan focal image de la lentille L2. De ce fait la fonction de transparence d'une pupille placée dans ce plan se trouve être multipliée par la fonction de frange d'interférence formée dans le plan de sortie. Il devient alors très aisé d'affecter une fonction de contourage au masque on-chip pour inscrire par exemple la géométrie de la zone guidante. Enfin la dernière particularité du dispositif vient du fait que la différence de marche entre les deux bras de l'interféromètre est nulle de construction. Le dispositif peut donc fonctionner aussi bien avec une source lumineuse à émission continue qu'avec une source impulsionnelle.

Le dispositif holographique peut donc être employé pour l'enregistrement en double exposition des motifs de microstructuration 2D utilisés pour la réalisation de cristaux photoniques à maille carrée. Dans ce cas le substrat enduit de résine photosensible est exposé une première fois par un réseau de franges linéaires puis une seconde fois, après rotation mécanique de 90° . Les zones où la résine a été insolée deux fois disparaissent lors du développement tandis que les zones non insolées subsistent et forment des plots qui protègent la couche mince de métal (Cr, Ni) déposée pour former le masque. Après gravure du métal, des plots métalliques subsistent à la surface du niobate et peuvent servir à nouveau de masque pour la gravure du niobate de lithium par une méthode collective telle que la RIE (gravure par plasma ionique réactif).

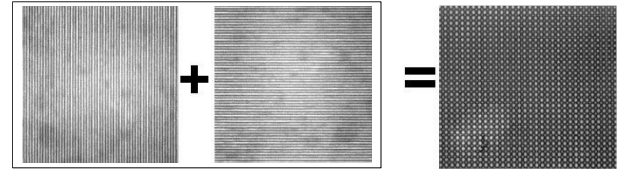


Figure 7 : Vue au microscope des résines après développement. L'enregistrement successif de deux champs de franges orthogonales permet, après développement de la résine, d'obtenir un masque de plots pour la réalisation d'un cristal photonique à maille carrée. La période est ici de $1.2\ \mu\text{m}$

4 Conclusion

Les propriétés électro-optiques, acousto-optiques ou non linéaires du niobate de lithium justifient l'intérêt que l'on peut porter au nano-usinage du matériau pour la fabrication de circuits micrométriques. Une première évaluation de structures de bandes nous a permis de sélectionner une configuration permettant d'assurer un contrôle de la transmission par décalage d'une bande interdite photonique. Cette configuration, qui consiste en un arrangement triangulaire de trous avec un diamètre de 213nm et une période de 426nm , a été fabriquée par faisceau ionique focalisée, avec une profondeur de $1,5\text{mm}$. Ce résultat représente une performance sur un substrat LiNbO_3 , car les facteurs de forme supérieurs à 1 sont très difficiles à réaliser sur ce matériau quasi-inerte chimiquement. La caractérisation de cette structure constitue la première démonstration optique de la faisabilité de structures photoniques dans le niobate de lithium. Elle confirme les simulations numériques qui évaluaient un creux dans la transmission pour la direction de propagation GM de 1200nm à 1550nm . Par ailleurs, nous avons montré que dans la perspective de fabrication collective de circuits ultra-compacts sur niobate de lithium, il est possible de définir les structures par technique holographique.

Références

- [1] S. Mahnkopf, M. Kamp, A. Forchel, F. Lelarge, G.-H. Duan and R. März., Mode anti-crossing and carrier transport effects in tunable photonic crystal coupled-cavity lasers, *Optics Communications*, 239, 187-191 (2004)
- [2] S.M. Weiss, M. Haurylau, P.M. Fauchet, Tunable photonic bandgap structures for optical interconnects, *Optical Materials*, 27, 740-744 (2005)
- [3] M. Schmidt, M. Eich, U. Huebner, R. Boucher, Electro-optically tunable photonic crystals, *Applied Physics Letters*, 87, 121110 (2005)
- [4] S. Yin, Fabrication of high-aspect-ratio submicron-to-nanometer range microstructures in LiNbO₃ for the next generation of integrated optoelectronic devices by focused ion beams (FIB), *Microwave and optical technology letters* 22, 396-398 (1999)
- [5] F. Lacour, N. Courjal, M.P. Bernal, M. Spajer, A. Sabac, C. Bainier, Nanostructuring Lithium Niobate substrates by focused ion beam milling, *Optical Materials*, 27, 1421-1423 (2005).
- [6] A. Mussot, T. Sylvestre, L. Provino, and H. Maillote, *Opt. Lett.*, 28, 1820 (2003).
- [7] R. Ferriere, B.E. Benkelfat, Novel holographic setup to realize on-chip photolithographic mask for Bragg grating inscription" *Optics Communication*, 206/4-6, 275-280, (2002)

Biographies :

Nadège COURJAL, née à Rennes en 1974, est agrégée de sciences physiques depuis 1997 et maître de Conférence au département d'optique P.M. Duffieux du laboratoire FEMTO-ST de Besançon depuis septembre 2004. Elle a effectué un D.E.A. spécialisé en optoélectronique à l'Institut National Polytechnique de Grenoble. Ces travaux de DEA à l'IMEP (Grenoble), puis de thèse au laboratoire FEMTO-ST, ont porté respectivement sur la fabrication de modulateurs non linéaire ou électro-optiques intégrés sur verre avec polymères photochromique ou sur niobate de lithium. Elle travaille maintenant avec Maria-Pilar Bernal dans les domaines de la nanophotonique sur niobate de lithium.

Richard FERRIERE, né en 1947 à Besançon, a reçu le diplôme de docteur en physique de l'Université de Franche Comté en 1973. Il intègre le CNRS comme chercheur au Laboratoire d'Optique P.M. Duffieux en 1978 et obtient le grade de docteur ès sciences physiques en 1982. Ses travaux de recherche ont porté sur l'interaction acousto-optique et l'imagerie ultrasonore, le traitement optique de l'information en lumière polychromatique et la conception de composants optiques holographiques blazés, les télécommunications optiques par modulation de cohérence, le calcul systolique optique. Ses thèmes de recherche actuels concernent les filtres de fréquence actifs intégrés sur niobate de lithium pour les télécommunications et la télémétrie optique.

Née à Zaragoza (Espagne) en 1970, **Maria-Pilar BERNAL** a fait son Ph. D. au IBM Almaden Research Center, San Jose, Californie où elle a travaillé de 1994 à 1998 dans le domaine du stockage holographique de données en trois dimensions. De 1998 à 2001 elle a effectué un stage postdoctoral à l'Ecole Polytechnique Fédérale de Lausanne (Suisse) et ses travaux de recherche ont porté sur la microscopie optique en champ proche et plus spécifiquement sur l'amélioration des sondes locales. Elle a rejoint le Département d'Optique P.M. Duffieux de FEMTO-ST en 2003 comme Chargée de Recherche CNRS. Ses recherches actuelles concernent la nanophotonique appliquée au niobate de lithium.

Les biopiles

Jean-Michel MONIER, Naoufel HADDOUR, Lorris NIARD, Timothy M. VOGEL, François BURET

Laboratoire Ampère UMR CNRS 5005, Equipe Microsystèmes et Microbiologie

Ecole Centrale de Lyon, 36 avenue Guy de Collongue, 69134 Ecully CEDEX.

Résumé : Les piles à combustible « microbiennes » (MFCs : « microbial fuel cells ») apportent de nouvelles opportunités en termes de production d'énergie à partir de composants biodégradables. Une MFC convertit directement en électricité une partie de l'énergie disponible dans un substrat biodégradable. C'est une voie intéressante, bien qu'encore un peu futuriste, d'élimination de déchets à moindre coût énergétique. Cet article donne un aperçu de l'état de l'art dans ce domaine.

1. Introduction

De manière très schématique, les piles mettent en œuvre une réaction d'oxydoréduction en récupérant une partie de l'énergie, autre que la chaleur, dégagée par cette réaction. Le métabolisme des bactéries et la dégradation de composants complexes, nécessaire à leur multiplication et diverses activités, se réalisent également par une chaîne de réactions d'oxydoréduction. En principe et sans préjuger des performances que l'on pourra obtenir, on peut donc espérer récupérer une partie de l'énergie produite par une ou plusieurs de ces réactions d'oxydoréduction pour produire de l'énergie électrique. Cette énergie électrique peut être récupérée dans des réacteurs appelés biopiles (MFC : *microbial fuel cells*) constitués, de manière très schématique, d'une culture de bactéries et de deux électrodes. Cet article présente une synthèse des connaissances actuelles dans ce domaine ainsi que les applications potentielles de cette technologie émergente dont l'optimisation et la mise en place à l'échelle industrielle nécessiteront une approche pluridisciplinaire.

2. Pile à hydrogène

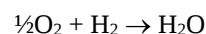
L'approche un peu simpliste évoquée dans l'introduction laisserait à penser que tous les couples redox permettent la réalisation de piles. En fait, si on se limite au domaine des piles à combustibles (PAC), on ne peut que constater que seules les piles à hydrogène/oxygène, et dans une moindre mesure au méthanol/oxygène¹, ont atteint un degré de maturité significatif.

Les quelques rappels qui suivent concernent la pile à hydrogène et permettront de mettre en regard quelques

aspects fondamentaux et technologiques des PAC et des biopiles.

2.1. Principe

La pile à hydrogène met en œuvre la réaction d'oxydoréduction bien connue :



L'énergie dégagée par cette réaction est donnée par la variation d'enthalpie :

$$\Delta H = \Delta G + T \cdot \Delta S \quad (1)$$

où ΔS est la variation de l'entropie molaire de la réaction, T est la température de la réaction et le terme $T\Delta S$ correspond à la partie de l'énergie transformée en chaleur.

ΔG est la variation de l'énergie libre de Gibbs et correspond à la partie qui est transformée en énergie électrique dans le cas d'une pile (combustion contrôlée) et en énergie mécanique dans le cas de la combustion détonante. On peut déduire que le rendement énergétique théorique d'une pile à hydrogène correspond alors au rapport de la variation de l'enthalpie libre (ΔG) sur la variation de l'enthalpie de la réaction (ΔH).

$$\eta = \Delta G / \Delta H$$

Dans les conditions standards de pression et de température :

$\Delta H^\circ = -285,8 \text{ KJ.mole}^{-1}$ et $\Delta G^\circ = -237,1 \text{ KJ.mole}^{-1}$ d'où un rendement attendu dans ces mêmes conditions d'environ 80%.

Le fonctionnement d'une pile à hydrogène est souvent décrit comme un processus inverse de l'électrolyse de l'eau. En fait, il s'agit d'une combustion contrôlée d'hydrogène et d'oxygène. Cette réaction s'opère au sein d'une structure essentiellement

¹ On parle ici des piles où le méthanol est oxydé directement dans la pile et non pas utilisé dans un

réformeur amont pour produire de l'hydrogène utilisé ensuite dans une PAC classique.

composée de deux électrodes, l'anode et la cathode, séparées par un électrolyte (Fig. 1).

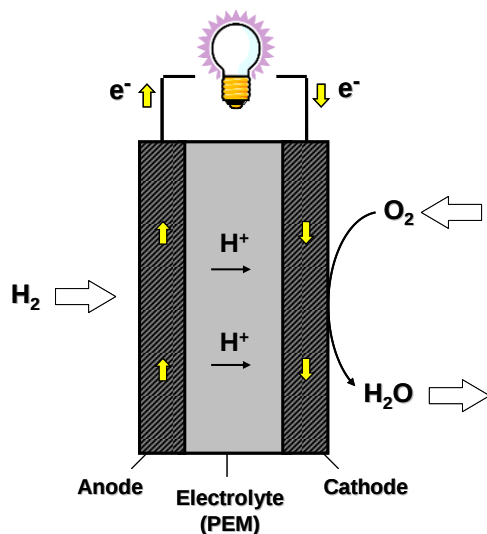
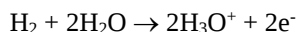


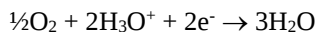
Figure 1 : Principe de la pile à hydrogène

La chambre anodique est alimentée en combustible (hydrogène) et la chambre cathodique en comburant (O_2).

A l'anode l'hydrogène est oxydé et décomposé en protons et électrons selon la réaction :



Les protons traversent ensuite l'électrolyte sous l'action du champ électrique et les électrons parcourent le circuit externe de la pile pour se retrouver à la cathode et réduire l'oxygène en eau selon la réaction :



La réaction peut donc être contrôlée par l'intermédiaire de l'intensité circulant dans le circuit électrique extérieur. L'électrolyte permet d'une part la migration des protons d'une électrode à l'autre tout en empêchant la diffusion de l'oxygène vers la chambre anodique et celle de l'hydrogène vers la chambre cathodique. Le combustible et le comburant ne doivent évidemment pas entrer en contact direct dans la pile à hydrogène.

2.2. Tension

La tension d'une pile à combustible peut être obtenue à partir de l'équation suivante :

$$W_{\text{élec}} = n.F.E = -\Delta G$$

où :

n est le nombre d'électron participant à la réaction (ici $n = 2$)

F la constante de Faraday ($9,64853 \cdot 10^4$ C/mole)

E la force électromotrice (fem) de la pile qui est donc fonction de la température et des pressions de dioxygène et dihydrogène².

Dans les conditions standard, la fem à vide est de 1,26 volts.

2.3. Tension réelle de la pile et problème technologique.

La tension en charge (et même la force électromotrice d'une pile à hydrogène) est inférieure à la valeur théorique annoncée précédemment. Cette différence entre la tension réelle et la tension théorique est due à des pertes³ dans le processus électrochimique.

Ces pertes apparaissent au niveau des électrodes et de l'électrolyte. Une part importante des recherches et développement dans le domaine des piles s'attache à diminuer ces pertes.

Les pertes ou chute de tension aux électrodes ont principalement deux causes :

- Chute de tension d'activation

Cette chute de tension d'activation est provoquée par la cinétique de la réaction électrochimique à la surface de l'électrode. En effet, pour que la réaction électrochimique puisse avoir lieu, les réactifs doivent franchir une barrière d'activation. Pour abaisser cette barrière d'activation et augmenter la vitesse de la réaction, il faut utiliser un catalyseur au niveau de l'interface électrode-électrolyte où se déroule chaque demi-réaction d'oxydoréduction. On est ainsi amené à utiliser des électrodes en platine (ou recouverte de platine). Le phénomène d'adsorption sur le platine baisse l'énergie d'activation de la dissociation de l'hydrogène et dans une moindre mesure de l'oxygène. Cette utilisation obligatoire du platine est un handicap bien connu des PAC.

- Chute de tension de concentration

Au voisinage des électrodes, il y a consommation des réactifs (O_2 à la cathode et H_2 à l'anode) ainsi que la formation d'eau. La faible diffusion des gaz à travers les électrodes poreuses et des protons dans l'électrolyte entraîne l'apparition d'un gradient de concentration qui traduit l'incapacité du système à maintenir la concentration initiale des réactifs. Cette diminution de l'activité des réactifs au voisinage des électrodes se traduit par une diminution de la tension.

- Chute ohmique

Cette chute de tension est principalement due à la résistance que rencontre le flux de protons dans l'électrolyte. Il est clair que cette chute de tension peut

² L'enthalpie libre est fonction de la température et des pressions respectives de l'oxygène et de l'hydrogène.

³ Toute perte venant amputer la variation d'enthalpie libre se traduit par une diminution de la tension et de l'énergie électrique générée.

être réduite en diminuant l'épaisseur de l'électrolyte et en améliorant sa conductivité ionique. La technologie la plus courante de pile à hydrogène (PEMFC) utilise une membrane polymère⁴ humidifiée de 50 µm d'épaisseur. Cette membrane est un point sensible de la technologie des piles en introduisant des pertes significatives et en limitant les températures de fonctionnement.

Compte tenu de tous ces facteurs, le rendement d'une pile réelle se situe plutôt aux alentours de 60% avec une fem de 0,8 volt.

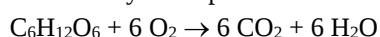
3. Biopile

3.1. Les acteurs et leur métabolisme

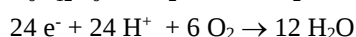
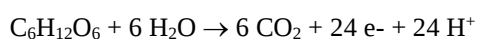
Les bactéries sont des microorganismes ubiquistes ne dépassant pas les quelques microns en taille et sont cependant les organismes les plus nombreux sur notre planète puisque leur nombre total est estimé à 10^{30} . Bien qu'elles soient souvent perçues comme des organismes nuisibles (pathogènes, multi-résistantes aux antibiotiques, sources d'infections nosocomiales), les bactéries, qui ont une vie en dehors des hôpitaux, jouent un rôle essentiel dans l'équilibre écologique de notre planète. Leur exceptionnelle capacité d'adaptation leur a permis de coloniser des environnements aussi divers et extrêmes que les fonds océaniques, la stratosphère, et plus proche de nous, les sols, sédiments, plantes, animaux et milieux aquatiques. Prises dans leur ensemble, les bactéries utilisent une vaste gamme de substrats organiques⁵ et inorganiques (sucres, acides aminés, stérols, alcools, ammonium, ...) comme source d'énergie. La diversité bactérienne est telle que l'on considère qu'aucune molécule organique ne peut être considérée comme strictement non biodégradable.

La diversité des activités métaboliques qui sont associées à des réactions d'oxydoréduction libérant des protons et électrons peut être exploitées dans les biopiles afin de générer de l'électricité. Le glucose par exemple peut être dégradé de plusieurs manières :

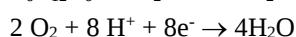
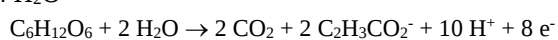
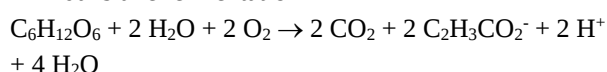
- dans un cycle respiratoire



avec les demi réactions



- dans une fermentation



Une autre particularité des bactéries, exploitée dans les biopiles, est que leur multiplication et leur activité métabolique s'effectuent principalement lorsqu'elles sont attachées à la surface d'un support où elles forment un film (la plaque dentaire étant un exemple de biofilm bactérien).

3.2. Principe de la biopile

La structure de principe d'une biopile n'est pas fondamentalement différente de celle d'une pile à hydrogène. Comme cette dernière, la biopile est composée de deux chambres contenant chacune une électrode (l'anode ou la cathode), séparées par un électrolyte, le plus souvent une membrane échangeuse d'ions (Fig. 2). Comme dans la pile classique, la chambre cathodique est alimentée en O_2 (ou air), les protons traversent l'électrolyte et les électrons parcourent le circuit externe de la pile (de l'anode vers la cathode) permettant ainsi la réduction de l'oxygène à la cathode.

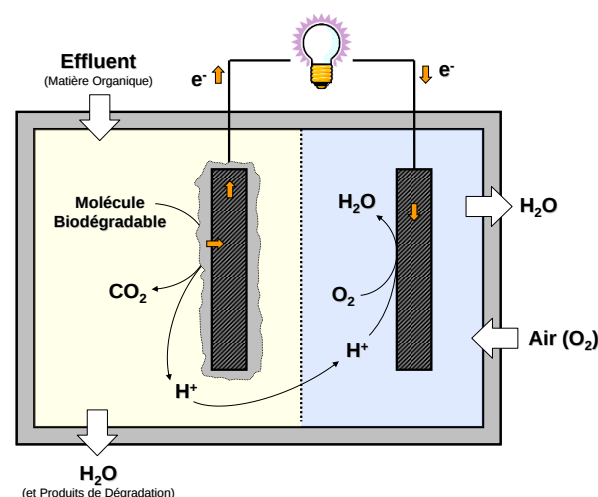


Figure 2 : Représentation schématique d'une biopile

La différence essentielle se situe au niveau de la chambre anodique. En effet c'est dans celle-ci que s'opère l'activité bactérienne qui permet, au final, la libération de protons et d'électrons. Cette libération est la conséquence des réactions d'oxydoréduction mises en œuvre par les bactéries pour extraire l'énergie qui leur est nécessaire à partir de différents substrats.

Comme nous l'avons mentionné dans un paragraphe précédent la dégradation de ce(s) substrat(s) (catabolisme bactérien) s'effectue à la surface de l'anode où se multiplie les bactéries.

⁴ nafion ® de DuPont ®

⁵ le terme substrat doit ici être compris dans le sens de nutriment (support de la vie bactérienne)

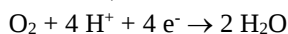
Alors qu'en terme de combustible, les PAC classiques sont limitées à l'hydrogène et au méthanol qui doivent même être « purifiés » avant injection, toute molécule organique sera un combustible potentiel pouvant être convertie en électricité à partir du moment où les bactéries pourront la dégrader.

3.3. Génération de courant dans les biopiles

Au même titre que les piles à combustibles classiques, la production d'électricité peut être évaluée à partir de la variation de l'énergie libre de Gibbs mise en jeu dans la réaction d'oxydoréduction. On peut également, de la même manière déterminer la fem.

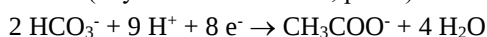
On peut considérer le cas de la dégradation de l'acétate qui est un produit intermédiaire apparaissant, par exemple dans le cycle respiratoire d'une bactérie. Cette dégradation se traduit par les demi réactions :

Cathode (réduction du dioxygène, pH=7) :



($E_0 = 805 \text{ mV}$)

Anode (oxydation de l'acétate, pH=7) :



($E_0 = -296 \text{ mV}$)

Ce qui donne une fem théorique d'environ 1,1 volts

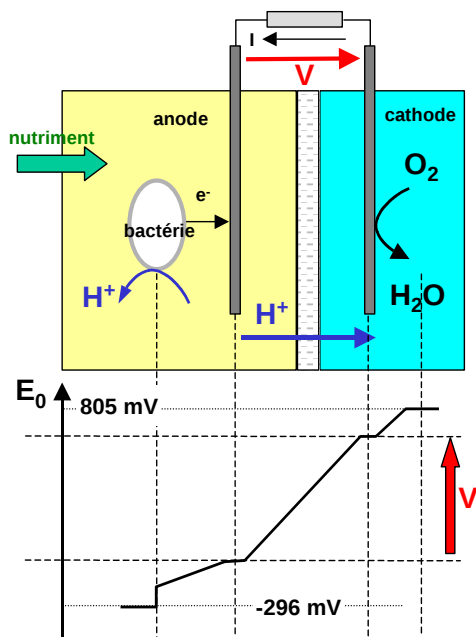


Figure 3 : Représentations schématiques d'une biopile et sources de perte de tension au sein de la biopile.

Cependant cette valeur n'est que théorique puisque de nombreuses pertes de tension peuvent survenir (Fig. 3). Outre celles décrites pour les piles à hydrogène, des pertes additionnelles peuvent survenir au sein des biopiles dues au métabolisme bactérien. En effet, si le potentiel de l'anode est trop bas, le métabolisme de la bactérie peut basculer de la respiration à la fermentation du substrat qui lui fournira alors plus d'énergie. D'une

manière générale, l'ordre de grandeur des tensions obtenues varie entre 0,6 V et 0,8 V.

4. Etat de l'art

Depuis le début du siècle dernier et les travaux de M.C. Potter (University of Durham, Stockton, UK), on sait que les bactéries sont capables de produire de l'électricité. Ce phénomène est resté pendant de nombreuses années à l'état de curiosité de laboratoire, les puissances générées étant a priori considérées comme trop faibles pour être facilement valorisées. Cependant, depuis le début des années 2000, un nombre croissant de laboratoires à travers le monde s'intéressent aux biopiles en cherchant à obtenir les meilleures performances possibles.

D'un côté on trouve des équipes de biologistes dont la préoccupation principale concerne moins la génération d'électricité que l'identification des phénomènes mis en jeu dans les biopiles. Ceci leur apporte, du point de vue fondamental, une meilleure connaissance des réactions associées au métabolisme des bactéries. D'un autre côté, certaines équipes de recherche ont une approche plus pragmatique et s'intéressent aux applications potentielles qui sont fort nombreuses dans le domaine du traitement des effluents. La reprise des travaux de recherche sur les biopiles a d'ailleurs été initiée par des chercheurs travaillant dans cette thématique (H.B. Kim; Korea Institute of Science and Technology, Séoul).

Nous aborderons dans cette partie la description des différents types de réacteurs développés au sein des laboratoires leaders dans la recherche sur les biopiles ainsi que leurs caractéristiques techniques et performances électriques.

4.1. Les différents types de biopiles

Il existe principalement deux types de biopiles : celles qui sont à chambres doubles séparées par une membrane et celles qui mettent en œuvre une seule chambre avec une cathode à air.

Le type le plus couramment utilisé est celui de la figure 2, à deux chambres, qui nous a servi à présenter le principe du fonctionnement des biopiles. La figure 4 présente une réalisation particulière d'une pile de ce type. Deux bouteilles sont reliées entre elles par un tube contenant la membrane de séparation. Cette membrane joue un rôle identique à celui qu'elle remplit dans la pile à hydrogène. Elle doit permettre le passage des protons mais pas celui du substrat ou des bactéries. Il est en effet impératif de limiter la croissance bactérienne dans la chambre cathodique afin que la réaction d'oxydation ne s'y effectue pas. Le plus souvent c'est une membrane échangeuse de cations (Nafion, Ultrex) qui est utilisée.

Ce type de biopiles ne génère que peu de puissance électrique et n'est le plus souvent utilisé que pour tester des paramètres particuliers : biologiques (substrat, espèce bactérienne) ou électrochimiques (matériau utilisé pour l'électrode).



Figure 4 : Exemple de réalisation de biopile à deux chambres (www.microbialfuelcell.org).

Le deuxième type de pile supprime la chambre cathodique en mettant en œuvre une cathode à air. L'oxygène pénètre dans la biopile à travers cette cathode poreuse au contact de laquelle se déroule la réduction des protons. Ce modèle ingénieux (Fig. 5) permet de supprimer la membrane échangeuse d'ions qui est sensible aux substrats complexes (risque de colmatage). Ce procédé a permis d'accroître de manière significative les puissances générées en diminuant les pertes par conduction due à la membrane. Ce modèle a été développé initialement aux Etats-Unis dans le laboratoire de B. Logan à Penn State University.

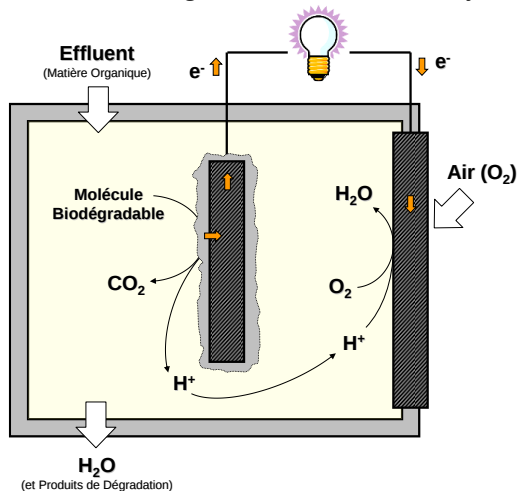


Figure 5 : Représentation schématique d'une biopile composée d'une seule chambre et d'une cathode directement exposée à l'air.

Dans le modèle le plus simplifié de ce système, l'anode et la cathode sont placées aux deux extrémités d'un tube, l'anode étant entièrement plongée dans la chambre évidemment sans contact avec l'air. Ce principe a été depuis décliné sous de multiples variantes.

Différentes architectures de biopiles sont présentées sur la figure 6.

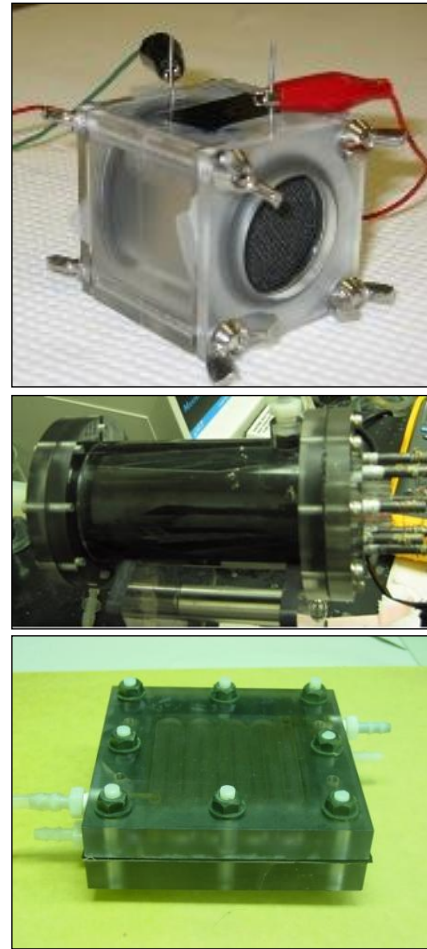


Figure 6 : Exemples de réalisation biopiles à une seule chambre (source : www.microbialfuelcell.org).

- (haut) cathode à air placée à l'extrémité de la chambre.
- (centre) cathode en forme de tube au centre de la chambre.
- (bas) cathode plaquée directement contre la membrane échangeuse de cations et l'anode, design se rapprochant de celui d'une pile à combustible classique.

Il est également possible de générer de l'électricité en plaçant les électrodes directement dans l'environnement dans des conditions reproduisant naturellement celles rencontrées dans les biopiles (ie cathode exposée à l'oxygène et croissance des bactéries à l'anode en présence de matière organique). C'est le cas par exemple des fonds marins, où l'anode est placée dans les sédiments marins et la cathode maintenue au dessus dans l'eau de mer contenant de l'oxygène (Fig. 7). Cette approche a été appliquée par les équipes de L. Tender (US Naval Research Lab, Washington DC) et D. Lovley (University of Massachusetts, Amherst) afin d'alimenter en électricité et de rendre autonome des appareils de mesure placés dans les fonds marins.

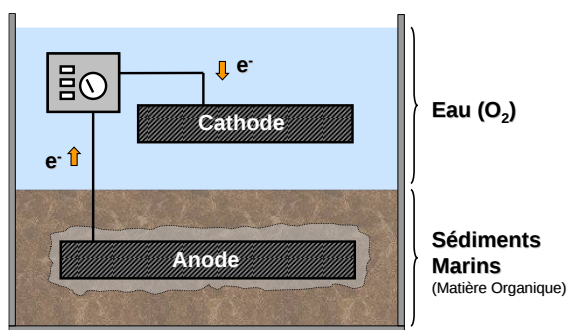


Figure 7 : Représentation schématique d'une biopile utilisant les conditions de l'environnement naturel des fonds marins pour générer de l'électricité.

4.2. Performances

Les puissances générées par les différents types de biopiles n'ont cessé d'augmenter depuis la reprise des recherches dans ce domaine à la fin des années 90 et suivent une progression exponentielle (Fig. 8).

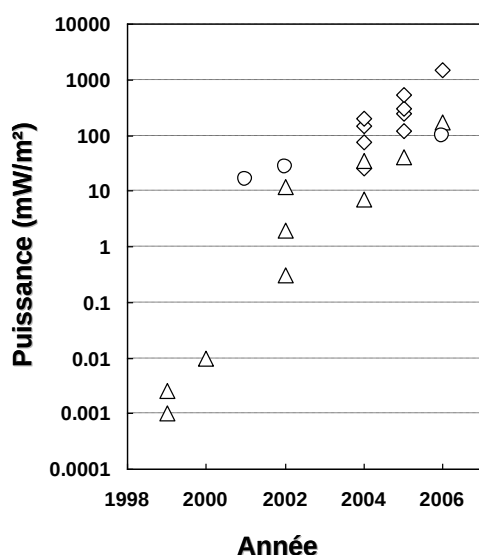


Figure 8 : Puissances électriques générées par les trois types de biopiles les plus couramment utilisées (d'après Logan et Regan, 2006b).

- (triangles) : chambre double.
- (losanges) : chambre simple, cathode à air.
- (cercles) : sédiments.

On caractérise les performances par la puissance par unité de surface active d'anode (surface où les bactéries se développent et les électrons sont transférés). Alors que les puissances générées par les biopiles placées dans les sédiments semblent plafonner autour de 100 mW/m², l'utilisation des cathodes à air a permis d'augmenter les puissances générées d'un facteur 10 permettant d'atteindre plusieurs W/m². Bien que ce niveau de performance puisse encore paraître dérisoire, il n'est pas déraisonnable de penser qu'il existe encore une marge de progrès conséquente.

5. Développements technologiques

L'optimisation des réacteurs tant sur le plan biologique que physique devrait permettre d'ici à quelques années de proposer une technologie économiquement viable. Une approche intégrée et pluridisciplinaire alliant génie électrique, électrochimie et microbiologie est indispensable afin d'aller au delà des progrès déjà réalisés. Les axes de recherche concernent essentiellement les électrodes, les communautés bactériennes susceptibles de produire de l'électricité et la structure des réacteurs.

5.1. Matériaux d'électrodes

5.1.1. Anode

Les matériaux utilisés pour l'anode doivent être conducteurs, chimiquement stables et évidemment biocompatibles. Alors que l'on peut envisager l'utilisation d'acier inoxydable, le cuivre est exclu du fait de la toxicité qu'il présente généralement vis-à-vis des bactéries. Actuellement, les matériaux les plus couramment employés sont à base de carbone. L'anode peut se présenter sous forme de plaques, granules ou bâtonnets en graphite mais également sous forme de fibres ou papier carbone.

Puisque l'activité bactérienne s'effectue à la surface de l'anode, un ratio surface/volume élevé doit être privilégié. Les physico-chimistes travaillant dans le domaine de la catalyse savent développer des matériaux présentant de très grandes surfaces d'échange par unité de volume. Dans le cas des biopiles il faut veiller à maintenir une porosité suffisante permettant d'assurer la circulation et donc la disponibilité pour les bactéries du substrat organique à dégrader. L'addition de molécules électrocatalytiques (poly-alinilins/Pt), bien que pouvant être coûteuse, permet d'augmenter les courants générés en assurant l'oxydation directe de certaines molécules (Fig. 9 - D).

5.1.2. Cathode

Les matériaux utilisés à ce jour pour la cathode sont également à base de carbone. Au niveau de la cathode, l'oxygène est l'accepteur idéal d'électrons pour une biopile du fait de son potentiel oxydatif élevé, de sa disponibilité et de sa gratuité (air). De plus sa réduction ne génère aucun déchet puisque la réaction aboutie à la formation d'eau comme dans la PAC classique. Cependant afin d'augmenter le faible taux de réduction de l'oxygène, l'addition de catalyseur tel que le platine, reste nécessaire. Il faut noter que la stabilité des catalyseurs sur le long terme n'a pas été précisément abordée dans les biopiles et qu'il y a là une grosse incertitude technologique. Afin de diminuer les coûts des biopiles, les recherches actuelles s'orientent vers le

développement de « biocathodes » où la réduction de l'oxygène serait assurée par des bactéries, formant comme à l'anode, un film à la surface de la cathode.

5.2. Microbiologie

Une connaissance plus approfondie des espèces bactériennes responsables de la production d'électricité permettrait d'optimiser la nature des matériaux ainsi que l'architecture des électrodes. Alors que différentes communautés de bactéries issues d'environnements extrêmement diversifiés sont capables de produire de l'électricité, il apparaît qu'un certain nombre d'espèces bactériennes (*Geobacter*, *Shewanella*...) jouent un rôle prépondérant dans la production d'électricité au sein des biopiles. L'utilisation de pressions de sélection environnementale favorisant de telles espèces pourrait permettre d'accroître les puissances générées. La caractérisation des modes de transferts des électrons à l'anode (direct, par l'intermédiaire de « navettes » excrétées par les bactéries, l'utilisation de nanotubes...) illustre cependant la complexité des modes de transferts et d'échanges d'électrons possibles entre bactéries et entre bactéries et anode (Fig. 8).

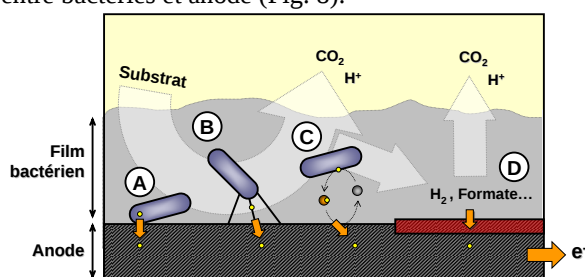


Figure 8 : Modes de transferts d'électrons identifiés chez les bactéries.

- (A) direct,
- (B) par le biais de « nanotubes » produits par les bactéries,
- (C) par le biais de médiateurs chimiques excrétés par les bactéries et jouant le rôle de navette entre la cellule bactérienne et l'anode.
- (D) l'addition d'électro-catalyseurs à la surface de l'anode permet de catalyser chimiquement la dégradation de molécules simples (hydrogène, formate...) et le transfert d'électrons supplémentaires directement à l'anode.

5.3. Utilisation de médiateurs chimiques

Comme nous l'avons signalé précédemment l'utilisation d'accepteurs d'électrons au niveau de la chambre cathodique peut permettre d'améliorer les performances des biopiles. L'optimisation de la

technologie peut également se faire par l'addition de molécules chimiques (médiateurs) jouant le rôle de navettes dans le transfert d'électrons entre bactéries et anode. Jusqu'à ce jour l'identité de tels médiateurs reste obscure et leur utilisation a été plutôt empirique. Il faut signaler en outre que l'instabilité chimique des médiateurs redox est une limitation majeure à leur utilisation. Une stratégie envisagée est la sélection, dans les biopiles, de bactéries produisant naturellement des médiateurs possédant, si possible, un faible potentiel d'oxydoréduction.

5.4. Architecture des systèmes

Une des limitations principales de l'utilisation des biopiles à grande échelle reste le passage du réacteur de laboratoire au réacteur de taille industrielle. La plupart des réacteurs développés à ce jour ne peuvent être « agrandis » pour traiter des mètres cubes d'effluent et générer de l'électricité. Il faudrait d'une part encore augmenter le taux de conversion de l'énergie contenue dans les molécules organiques en énergie électrique et d'autre part simplifier les designs des réacteurs et matériaux utilisés pour les électrodes afin de réduire leur coût. L'utilisation de stacks comme ceux utilisés dans les piles à hydrogène est envisagée afin d'accroître les puissances.

5.5. Mesures des caractéristiques des biopiles

Afin de comprendre les mécanismes mis en place pour la production d'électricité dans les biopiles, toute une batterie d'outils de biologie moléculaire combinés à des mesures de paramètres électrochimiques et électriques sont mis en place. Les analyses microbiologiques ont pour but d'identifier les espèces bactériennes « électroactives », de déterminer leur nombre, leur activité et leur distribution spatiale sur l'anode afin d'optimiser l'architecture des électrodes.

Des mesures de paramètres électriques sont effectuées en continu à divers instants du cycle biologique afin d'évaluer leur potentiel de production.

Enfin des mesures telles que l'efficacité de la charge électrique, la demande en oxygène ou la disparition du carbone organique total sont effectuées afin de déterminer l'efficacité de la dégradation de la matière organique.

Les performances des biopiles (conversion de l'énergie chimique fournie en énergie électrique récupérée) peuvent varier de 2 à 50 %, selon les types de réacteurs.

6. Applications

6.1. Le traitement d'effluents

Dans le domaine des applications, le traitement des effluents vient évidemment à l'esprit. Un nombre croissant de travaux tend à prouver que de nombreux effluents d'origine domestique, agricole ou encore industrielle constituent des combustibles potentiels pour les biopiles. Ces effluents contiennent en fortes quantités non seulement des déchets organiques solubles mais également les bactéries pouvant les dégrader. C'est un moyen élégant de traitement des déchets puisqu'en éliminant ceux-ci on peut récupérer une partie de leur contenu énergétique.

Aux Etats-Unis, près de 25 milliards de dollars sont dépensés chaque année pour le traitement des eaux usées avec un coût énergétique correspondant à 4% de l'électricité produite au plan national. On estime que les eaux usées contiennent potentiellement 10 fois plus d'énergie, sous forme de substrat organique, que celle qui est nécessaire à leur traitement. Dans ce contexte, le développement d'un système de traitement basé sur la technologie des biopiles offre des perspectives économiques intéressantes. B.E. Logan (Penn State University, USA) et son équipe sont des pionniers dans ce domaine. Ils visent à développer des biopiles à échelle industrielle qui viendraient suppléer les systèmes de traitement actuels, extrêmement coûteux en énergie. Dans un de leur article, ils estiment à 300 kW la puissance disponible à partir d'une production de 7500 kg de déchets organiques par jour (en assumant un rendement énergétique de 30%). D'autres laboratoires de recherche travaillent également sur la génération d'électricité à partir d'effluents plus spécifiques d'origines industrielles ou agricoles.

6.2. Appareils de mesure autonomes

La mesure en continu de paramètres environnementaux dans des sites difficilement accessibles peut être facilitée par l'utilisation de biopiles utilisant les substrats organiques et bactéries naturellement présentes dans l'environnement ; les biopiles permettant alors d'assurer l'alimentation en énergie des systèmes de mesures et de transmission. C'est le cas des sédiments marins, que nous avons abordé précédemment, ou des systèmes autonomes placés au fond des mers collectant des données qui sont ensuite transmises vers la surface. Cette technologie est a priori applicable à tous types d'environnements aquatiques (rivières, marais, océans...) puisqu'on y est assuré de trouver des bactéries et des substrats adaptés.

6.3. Production d'hydrogène

Une modification des réacteurs décrits ci-dessus peut permettre à partir des mêmes substrats organiques de produire non plus de l'électricité mais de l'hydrogène. Il suffit pour cela d'appliquer à l'aide d'un générateur externe une tension relativement faible de l'ordre de 0,25 V entre les électrodes⁶ pour pouvoir générer de l'hydrogène au niveau de la chambre cathodique. Cette transformation est énergétiquement intéressante puisque le contenu énergétique de l'hydrogène produit est plus important que l'énergie fournie sous forme d'électricité.

6.4. Autres applications

Un certain nombre d'applications se profile à un horizon plus ou moins lointain. On estime, par exemple, que d'ici quelques années, il sera possible de faire évoluer les biopiles afin de générer de l'électricité directement à partir de la biomasse : résidus végétaux issus de l'agriculture ou de l'industrie forestière.

Une autre application potentielle est d'aider à la dépollution de sols en convertissant les polluants solubles en forme insolubles. Dans ce cas les bactéries sont utilisées non plus comme donneuses d'électrons mais jouent au contraire le rôle d'accepteurs d'électrons afin d'assurer la précipitation des composés solubles au niveau de la cathode. Cette technologie est déjà appliquée notamment à la précipitation de l'uranium ou encore la conversion de nitrate en nitrite.

Ces deux exemples ne constituent évidemment pas une liste exhaustive.

7. Conclusions

Les biopiles sont des systèmes qui suscitent un intérêt évident quel que soit le public auquel on s'adresse. Il y a en effet un côté un peu magique, auquel personne n'échappe, dans le fait de pouvoir produire de l'énergie électrique, qui est l'énergie noble par excellence, à partir de déchets dont on ne sait pas trop quoi faire. Il n'en reste pas moins que les enjeux industriels et scientifiques sont présents et font sortir les biopiles de la simple curiosité de laboratoire.

Les enjeux industriels et sociétal apparaissent clairement même si il reste évidemment beaucoup de chemin à parcourir avant une mise en œuvre à grande échelle.

Au niveau scientifique, les biopiles constituent un challenge passionnant pour les chercheurs issus de communautés très différentes. Biologistes, physico-chimistes, électriciens doivent s'associer afin de relever le défi que représente le développement de ces systèmes.

⁶ Anode à la polarité positive et cathode à la polarité négative

8. Références:

- Buckley M, Wall J** (2006) Microbial Energy Conversion: Report from the American Society of Microbiology (24 pages), Washington, DC, USA (www.asm.org).
- Liu H, Grot S, Logan BE** (2005) Electrochemically assisted microbial production of hydrogen from acetate. *Environmental Science and Technology* **39**:4317-4320.
- Logan BE and Regan JM** (2006a) Microbial fuel cells: challenges and applications. *Environmental Science and Technology* **40**: 5172-5180.
- Logan BE and Regan JM** (2006b) Electricity-producing bacterial communities in microbial fuel cells. *Trends in Microbiology* **14**: 512-518.
- Logan BE et al.**, (2006) Microbial Fuel Cells: Methodology and Technology. *Environmental Science and Technology* **40**: 5181-5192.
- Lovley DR** (2006) Bug juice: harvesting electricity with microorganisms. *Nature Reviews* **4**: 497-508.
- Potter MC** (1912) Electrical effects accompanying the decomposition of organic compounds. *Proceedings of the Royal Society* **84**: 290-276.
- Rabaey K, Verstraete W** (2005) Microbial fuel cells: novel biotechnology for energy generation. *Trends in Biotechnology* **23**: 291-298.
- www.microbialfuelcell.org** : Site dédié aux principaux laboratoires impliqués dans la recherche sur les biopiles.

Leds blanches et conversion DC-DC - 1^{ERE} PARTIE

Frédéric NARCY

PROFESSEUR DE PHYSIQUE APPLIQUÉE

Lycée Léonard de Vinci - 91240 Saint-Michel-sur-Orge

Frederic.Narcy@ac-versailles.fr

Résumé : Grâce à une progression constante de leur puissance unitaire et de leur efficacité lumineuse, les diodes électroluminescentes blanches voient leur utilisation se généraliser. L'étude de ces composants, souvent alimentés par un hacheur série ou parallèle, permet d'aborder de nombreux thèmes liés à l'éclairage : photométrie, économies d'énergie, développement durable... Après une présentation générale des leds, ce premier article décrit des expériences simples dont le but est de faire comprendre les limites du circuit classique alimentation+résistance+led et de justifier le recours à l'électronique de puissance.

1. LES DIODES ÉLECTROLUMINESCENTES

1.1. D'abord, la couleur...

Les diodes électroluminescentes de couleur font partie de notre quotidien depuis plus de trente ans. D'abord utilisées comme voyants, elles ont aussi constitué les premiers affichages numériques de petite taille (premières montres électroniques et premières calculatrices) avant l'arrivée des afficheurs à cristaux liquides. Puis c'est dans le domaine de la signalisation routière et ferroviaire qu'elles ont permis de véritables économies d'énergie. Cette application montre bien les avantages des leds pour la signalisation :

- Les leds ont une meilleure efficacité lumineuse que les lampes à incandescence, elles consomment donc moins d'énergie pour une même quantité de lumière produite.

- Pratiquement toute la lumière d'une led est émise dans la bonne direction. En revanche, une lampe à incandescence nécessite un réflecteur qui ne récupère pas toute la lumière, ainsi qu'un filtre coloré qui l'atténue.

- Un feu de signalisation à leds présente une meilleure visibilité, notamment en cas de forte luminosité ambiante.

- La durée de vie des leds est très supérieure à celle des autres sources lumineuses. Le nombre de 100 000 heures est souvent avancé. Cette donnée doit cependant être nuancée car le vieillissement d'une led est accompagné d'une diminution de son flux lumineux. Avec une durée de vie de l'ordre de 1000 heures seulement, une lampe à incandescence doit être remplacée régulièrement, ce qui multiplie les interventions.

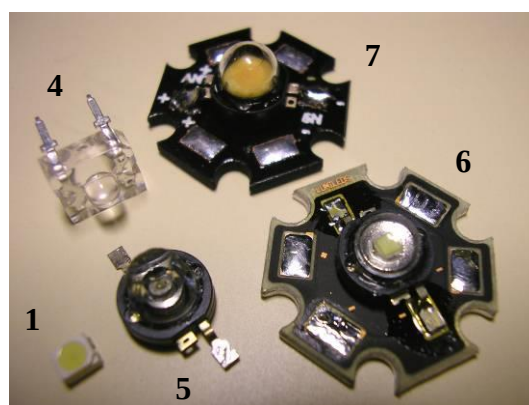
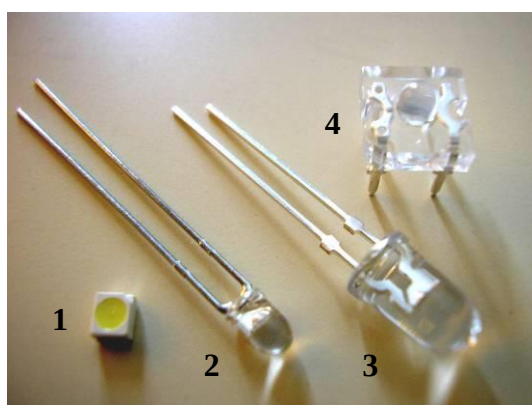


Figure 1 : différents types de leds blanches

- 1 : Led CMS (composant monté en surface) - 300 mcd 2 : Led ronde Ø 3 mm 3 : Led ronde Ø 5 mm - 20 mA - 9200 mcd
4 : Led carrée - 30 mA - angle d'émission 20° - coordonnées colorimétriques $x = 0,33$ $y = 0,34$
5 : Emetteur Luxeon - 1 W - 350 mA - émission latérale (pour rétroéclairage ou éclairage indirect)
6 : StarLuxeon - 1 W - 350 mA - angle d'émission 160° - température de couleur 5500 K - indice de rendu des couleurs 70 - 25 lm (c'est un émetteur comme le précédent mais soudé sur une platine facilitant les connexions et le refroidissement)
7 : Led de puissance - 3 W - 700 mA - blanc chaud - température de couleur 3300 K - indice de rendu des couleurs 90 - 45 lm



Ainsi, un feu tricolore à leds consomme moins d'énergie, et coûte moins cher en maintenance. On cite souvent le cas de la ville de Grenoble qui a équipé de leds tous ses feux tricolores. Cet investissement a été amorti en trois ans et permet depuis une économie annuelle de 55000 €.

Une augmentation significative de la puissance unitaire à la fin des années 1990 va donner un nouvel essor à l'utilisation des leds. Actuellement, les puissances les plus courantes sont 1 W, 3 W et 5 W. Le principal fabricant, la firme Lumileds, a donné à ses leds de puissance le nom de Luxeon (diodes 5 et 6 de la figure 1). Lumileds est une "joint venture" issue de Agilent Technologies et de Philips Lighting.

Ces composants se prêtent à des applications qui exigent des flux lumineux plus importants comme la signalisation marine et aéroportuaire ou l'éclairage décoratif de monuments et de bâtiments.

Contrairement aux autres sources, on peut régler le flux lumineux d'une diode électroluminescente en agissant sur son intensité, tout en gardant une bonne efficacité lumineuse et avec peu de variations du spectre émis. Cette propriété est mise à profit pour réaliser la synthèse additive des couleurs en combinant dans un même luminaire des diodes de puissance rouges, vertes et bleues. On peut ainsi, grâce à une programmation, moduler la teinte d'un éclairage de manière progressive. Dans le cas d'un bâtiment possédant plusieurs zones éclairées indépendantes, chaque étage d'une tour par exemple, on peut créer de superbes animations.

Les écrans vidéo géants d'extérieur utilisent aussi des leds. Le plus grand d'entre eux (en 2005), situé dans un stade à Atlanta, comporte 1080×1200 pixels de quatre leds et consomme 400 kW.

En choisissant correctement le nombre de leds de chacune des couleurs primaires, on peut obtenir une lumière blanche d'excellente qualité dont la principale application est le rétroéclairage d'écrans vidéo à cristaux liquides.

1.2. ...Puis le blanc

Les premières leds bleues n'apparurent qu'en 1990. Ce n'est qu'après qu'il a été possible de fabriquer des leds émettant de la lumière blanche. Mais une diode blanche n'est pas composée, comme on pourrait le croire, de trois diodes, une rouge, une verte et une bleue, intégrées dans un même boîtier. Une led blanche est en fait une led bleue dont la partie émettrice est recouverte d'une ou plusieurs substances fluorescentes, appelées chromophores, qui absorbent une partie de la lumière bleue et réémettent de la lumière de longueur d'onde plus élevée, donc moins énergétique. C'est un peu le principe d'un tube fluorescent, sauf que l'on part d'une lumière

bleue dont une partie n'est pas convertie et non d'un rayonnement ultraviolet qui doit être entièrement converti en lumière.

Comme pour les tubes fluorescents, on trouve différentes qualités de blanc.

Les diodes blanches les plus courantes ne comportent qu'une substance fluorescente qui émet dans le jaune. Sur le spectre de la figure 2, on distingue parfaitement le pic étroit dans le bleu de la puce émettrice et le pic large du chromophore. La lumière résultant du bleu non absorbé et du jaune est d'un blanc bleuté assez froid. En observant une diode de ce type hors alimentation, on peut apercevoir la couleur jaune du chromophore. Le faisceau lumineux, observé sur un écran blanc, est plutôt bleu au milieu et jaune sur les bords car le chromophore diffuse la lumière jaune dans toutes les directions.

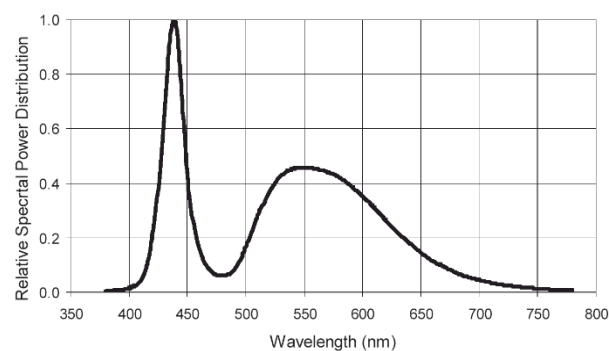


Figure 2 : Spectre d'une led blanche "ordinaire"
(document Lumileds)

D'autres LED blanches émettent un blanc qualifié de "chaud", plus proche de la lumière d'une ampoule à incandescence et qui donne un meilleur rendu des couleurs. Grâce à un mélange de chromophores, le spectre est plus étalé et comporte plus de rouge (figure 3). Sur certaines de ces diodes on peut observer une couche de chromophores de couleur orangée.

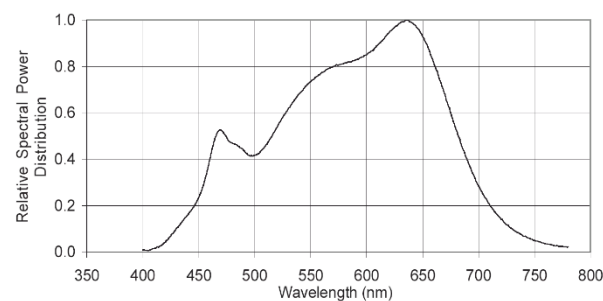


Figure 3 : Spectre d'une led "blanc chaud"
(document Lumileds)

D'autres méthodes permettent de réaliser des leds blanches mais elles ne semblent pas encore utilisées pour des fabrications en série.

1.3. Les paramètres caractéristiques d'une led blanche

Ces différents paramètres peuvent apparaître sur les fiches techniques des constructeurs et les catalogues des distributeurs :

- **L'intensité** est toujours indiquée : 20 mA, 50 mA, 350 mA, 700 mA, 1400 mA sont les valeurs les plus courantes.

- **La puissance**, exprimée en watts, est donnée en général pour les LED de puissance, entre 1 W et 5 W.

- **La tension**. Ce paramètre pouvant varier énormément d'une diode à l'autre d'un même lot, le constructeur donne une valeur typique, une valeur minimale et une valeur maximale. La tension est généralement comprise entre 3 V et 4 V.

- **Le flux lumineux**, exprimé en lumens (lm), mesure la quantité de lumière émise. Cette grandeur dépend de la sensibilité de l'œil humain aux différentes longueurs d'onde. Par exemple, pour une même énergie émise, le flux lumineux est plus élevé pour une lumière jaune que pour une lumière bleue. Le flux lumineux typique d'un Luxeon blanc de 5 W est de 120 lm.

- **L'intensité lumineuse** dépend de la direction d'émission, elle est exprimée en millicandelas (mcd). On donne soit l'intensité lumineuse maximale, soit une courbe en fonction de l'angle entre la direction considérée et l'axe optique.

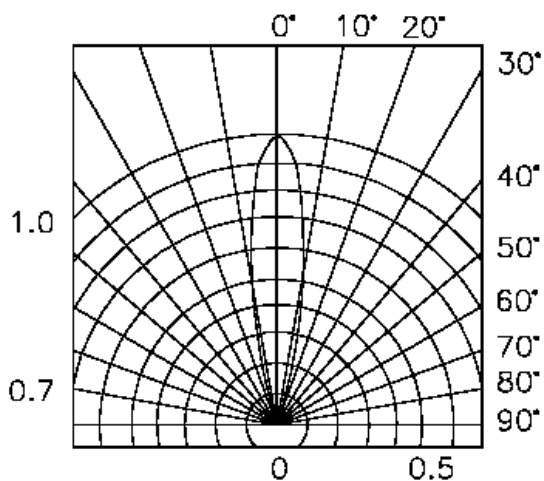


Figure 4 : Courbe d'intensité lumineuse relative de la diode 4 de la figure 1 (document Kingbright)

Si on calcule l'intégrale de l'intensité lumineuse sur toutes les directions d'émission, avec comme variable l'angle solide, on obtient le flux lumineux.

On trouve actuellement des leds blanches de diamètre 5 mm dont l'intensité lumineuse maximale est égale à 25000 mcd. Ce paramètre, comme beaucoup d'autres, ne cesse de progresser.

- **L'angle d'émission** donne une idée de la "largeur" du faisceau. On trouve deux définitions. D'après la première, c'est l'angle au sommet du cône qui contient 90 % du flux lumineux, le sommet de ce cône étant la puce émettrice de la diode (diode 6 de la figure 1). La deuxième définition est utilisée quand l'intensité lumineuse décroît avec l'angle depuis l'axe optique : c'est deux fois l'angle qui correspond à la moitié de l'intensité lumineuse maximale (diode 4 de la figure 1 et figure 4). Les intensités lumineuses les plus élevées sont obtenues grâce à un angle d'émission faible. Les faisceaux lumineux des diodes de puissance étant assez larges, il existe des optiques à rajouter sur ces composants pour obtenir des faisceaux plus étroits.

- **La température de couleur**, exprimée en kelvins (K), est la température à laquelle le rayonnement du corps noir présente le spectre le plus proche de celui de la diode. Cette définition est donc approximative car les deux spectres ne coïncident pas. Elle vaut généralement 5500 K pour le blanc des diodes blanches ordinaires et 3300 K pour le blanc chaud. Pour les leds blanches, cette grandeur est donnée à la place de la longueur d'onde centrale des leds de couleur. On notera que l'association communément faite entre couleur et température ne correspond pas au rayonnement du corps noir puisque l'on qualifie de chaude une couleur qui correspond à une température plus faible.

- **Les coordonnées colorimétriques x et y** selon la Commission Internationale de l'Eclairage permettent de quantifier la nature du blanc émis. Le blanc parfait pour l'œil humain est caractérisé par $x = 0,333$ et $y = 0,333$. Ces données sont beaucoup plus précises que la température de couleur. On trouvera le diagramme correspondant à ce système de repérage des couleurs dans l'article de Pierre ALBOU sur l'éclairage automobile dans le numéro 33 de La Revue 3EI.

- **L'indice de rendu des couleurs (IRC)** mesure la capacité d'une source lumineuse à restituer les couleurs des objets éclairés. La valeur maximale de 100 est obtenue avec un spectre comportant toutes les longueurs d'onde, ce qui est pratiquement le cas de la lumière solaire. Les lampes à incandescence ont un IRC supérieur à 90. Pour les leds blanches ordinaires, il est de l'ordre de 70 et pour les leds "blanc chaud" de l'ordre de 90. Dans les locaux scolaires et les bureaux, une valeur minimale de 80 est exigée.



1.4. Le refroidissement des leds de puissance

Une lampe à incandescence ne convertit en lumière que 5 % de l'énergie absorbée, 83 % de cette énergie est rayonnée sous forme d'infrarouges et les 12 % restants sont transformés en chaleur. Une led convertit 15 % de l'énergie absorbée en lumière mais tout le reste est perdu sous forme de chaleur qu'il faut évacuer par conduction et convection. Le refroidissement est donc un problème important à prendre en compte lors de la conception d'un luminaire utilisant des leds de puissance. Il existe des dissipateurs spécialement dessinés pour ces diodes.

1.5. L'efficacité lumineuse

Le paramètre utilisé pour comparer des sources de lumière sur le plan des économies d'énergie est l'efficacité lumineuse exprimée en lumens par watt (lm/W). C'est le rapport entre le flux lumineux produit et la puissance électrique absorbée.

source lumineuse	efficacité lumineuse (lm/W)
lampe à pétrole	0,03
première lampe d'Edison	2
lampe à incandescence	5 - 20
lampe halogène	10 - 40
led blanche	15 - 50
lampe fluocompacte	30 - 70
lampe à induction	60 - 80
tube fluorescent	50 - 100
lampe à vapeur de sodium basse pression	100 - 200

Figure 5 : efficacité lumineuse de différentes sources

Dans le tableau ci-dessus, on donne pour chaque type de source une plage de valeurs assez large. En effet, pour une technologie donnée, l'efficacité lumineuse dépend de plusieurs facteurs :

- elle peut dépendre de la puissance unitaire par effet d'échelle, comme c'est le cas du rendement des machines électriques,
- les progrès réalisés ont doté les modèles les plus récents d'un meilleur rendement,
- une enveloppe diffusante ou légèrement colorée fait chuter l'efficacité lumineuse.
- il peut exister différentes sortes d'alimentation pour un même type de source, par exemple les ballasts électromagnétiques ou électroniques pour les tubes.

On retiendra que les leds blanches sont loin d'être les sources lumineuses les plus économes en énergie. Cependant, elles présentent les meilleures perspectives de progression : en laboratoire, on a déjà dépassé les 60 lm/W. Pour les produits de série, on attend 75 lm/W en 2007 et 150 lm/W dans 10 ou 20 ans.

On rappelle la limite supérieure théorique de l'efficacité lumineuse : une source qui convertirait toute l'énergie électrique en énergie lumineuse à la longueur d'onde de 555 nm atteindrait 683 lm/W.

Grâce à leur faisceau directionnel, les leds peuvent se montrer plus économes qu'il n'y paraît pour certaines applications. En effet, une source de meilleure efficacité lumineuse mais qui émet dans toutes les directions nécessite un réflecteur qui ne renvoie jamais toute la lumière produite dans la bonne direction, ce qui fait chuter l'efficacité lumineuse du luminaire.

1.6. Les applications des leds blanches

C'est surtout pour l'éclairage que l'on utilise les leds blanches. Aujourd'hui, si elles ont montré une très nette supériorité pour certaines applications, en revanche elles sont encore loin d'accéder au statut de source de lumière universelle. Les diodes électroluminescentes ne présentent pas autant d'avantages décisifs pour l'éclairage que pour la signalisation. C'est surtout leur efficacité lumineuse moyenne et leur prix élevé qui limitent leur utilisation, du moins pour l'instant.

Les leds blanches s'imposent quand on ne peut pas utiliser les sources lumineuses plus économes faute de réseau électrique. On distinguera deux cas :

- l'éclairage portable par torches et lampes frontales pour les activités d'extérieur (loisirs, travaux, secours).
- l'éclairage des habitations sans distribution électrique, ce qui concerne 1/3 de l'humanité, surtout en zone rurale dans les pays en voie de développement.



figure 6 : un luxeon de 1 W éclaire une habitation en Inde (photo LUTW)

Enfin les leds blanches sont utilisées dans tous les domaines pour leur facilité d'intégration et pour l'uniformité de leur faisceau lumineux dans une zone étendue.

1.6.1. L'éclairage portable

Dans une lampe de poche, les leds remplacent les ampoules à incandescence beaucoup plus gourmandes. On gagne ainsi en autonomie et cela d'autant plus qu'une led continue à éclairer jusqu'à l'usure complète des piles, alors qu'une ampoule à incandescence n'émet plus que dans l'infrarouge en dessous d'une certaine tension d'alimentation. Les spéléologues qui ont besoin d'un éclairage sûr ont adopté les leds blanches très rapidement. Au début, des montages artisanaux permettaient de garder la même lampe en remplaçant l'ampoule par une association de diodes munies d'une résistance. Maintenant, on trouve des lampes à led, équipées ou non d'un convertisseur, chez tous les grands distributeurs d'articles de sport.

Les leds ont considérablement fait évoluer ce domaine de l'éclairage car leur faible consommation autorise d'autres sources d'énergie que des piles ou des accumulateurs rechargés par le secteur, notamment l'énergie humaine. On trouve une grande variété de "chaînes d'alimentation" :

- panneau solaire + accumulateur,
- manivelle + génératrice + accumulateur,
- aimant + bobine + redresseur + supercapacité.

Le dernier type de lampe doit être secoué pour recharger le condensateur. La durée d'éclairage est égale à plusieurs fois la durée de recharge.



figure 7 : une lampe de poche à leds et à manivelle

La durée de vie des leds est telle qu'il n'est pas nécessaire de prévoir leur remplacement, ce qui simplifie la fabrication de la lampe.

1.6.2. L'éclairage des habitations en absence de réseau électrique

Dans ce deuxième cas, les leds remplacent une flamme obtenue par combustion d'un hydrocarbure ou d'huile. Ce mode d'éclairage archaïque ne présente qu'un

avantage : le coût initial du luminaire est faible, souvent la lampe à huile ou à pétrole est fabriquée par l'utilisateur lui-même. Par contre, le prix du combustible représente une charge élevée et ce mode d'éclairage est dangereux en raison de sa toxicité et des risques d'incendie et de brûlures.

Pour l'éclairage à leds c'est exactement l'inverse : un investissement très élevé au départ puis aucune dépense pendant des années. A long terme, cela revient beaucoup moins cher qu'un éclairage à combustible !

On retrouve là un problème récurrent des pays en voie de développement : de nombreux objets techniques permettent des économies d'énergie ou facilitent certains travaux mais restent inaccessibles à la majorité à cause de l'investissement initial. La distribution locale ne se développe donc pas. Pour le moment, le rôle des organisations caritatives et des associations locales est encore essentiel pour développer ces technologies par le biais de dons ou du micro-crédit.

La fondation canadienne Light Up The World (LUTW) créée par le professeur Dave Irvine-Halliday de l'université de Calgary se consacre entièrement à la diffusion de l'éclairage à leds (figures 6 et 8). C'est au Népal en 2000 que commence cette aventure électrique et humanitaire : plusieurs villages sont "électrifiés" grâce à des panneaux solaires et à de petites génératrices hydrauliques ou à main. Les luminaires sont alors composés de 6 ou 9 diodes de diamètre 5 mm. Aujourd'hui, LUTW intervient dans de nombreux pays et propose un ensemble composé d'un panneau solaire de 5 W, d'une batterie 12 V - 7 Ah et de deux lampes comportant chacune un luxon de 1 W alimenté par un convertisseur (figures 6 et 8).



figure 8 : une installation électrique complète
(photo LUTW)

La puissance électrique mise en oeuvre peut paraître dérisoire à un consommateur occidental énergivore, mais ce dernier doit savoir que l'on peut lire à la lueur d'une led blanche de 0,07 W !



1.6.3. Facilité d'intégration

Très souvent, le recours à des leds blanches se justifie par leur petite taille qui permet de réaliser des éclairages discrets. Il s'agit alors le plus souvent de produits de haute technologie.

Il y a d'abord le rétroéclairage des écrans couleur à cristaux liquides comme ceux des téléphones portables, des agendas électroniques ou des appareils photo numériques où les leds entrent en concurrence avec les mini tubes fluorescents. L'utilisation des leds va aussi se développer dans le domaine des écrans vidéo de grandes dimensions.

On trouvera bientôt des projecteurs vidéo de petite taille dont la source lumineuse sera une led de puissance, beaucoup moins fragile que les lampes actuelles.

Dans un luminaire d'intérieur, comme une lampe de bureau, une source de lumière conventionnelle occupe beaucoup de place avec son réflecteur. Avec des leds, sans réflecteur et que l'on peut répartir, il est possible d'adopter des lignes très épurées.

En architecture, les diodes électroluminescentes sont souvent utilisées comme éclairage ponctuel ou décoratif ou encore pour la signalisation. On les incorpore aux nez de marche, dans les planchers, ou en corniche. Des rubans flexibles de diodes montées en surface permettent de réaliser des lignes lumineuses.

Les leds partent également à la conquête de l'automobile. Après les feux stop et les clignotants, puis l'éclairage intérieur, c'est au tour des feux de position à l'avant, mais seulement sur des prototypes ou des modèles de série de très haute gamme (Audi A8 L 6.0). Pour les feux de route, il faudra encore attendre.

On pourra bientôt apprécier la lumière du luxeon blanc des liseuses du futur TGV-Est ainsi que celle du plafond lumineux du futur Boeing 787.

1.7. Intérêt des leds pour l'enseignement

D'abord, les leds blanches présentent un attrait indéniable pour les élèves. En effet, le fonctionnement d'une led s'observe directement, sans appareil de mesure. Comme un moteur, un haut-parleur ou une résistance de chauffage, une led "fait quelque chose" de perceptible. A cela s'ajoute la nouveauté : l'aspect d'une led de puissance en fonctionnement est différent de celui des sources lumineuses habituelles.

En raison de leur faible puissance comparée à celle d'autres récepteurs, les leds sont faciles à mettre en œuvre. Le niveau de tension ne pose pas de problème de sécurité et le coût reste raisonnable à l'échelle d'un laboratoire d'enseignement. Pour ces raisons, les leds blanches se prêtent particulièrement bien à la réalisation de mini projets. La différence entre le montage de

laboratoire et l'application sur le terrain ne porte que sur le nombre de leds utilisées.

Les leds permettent d'illustrer des notions électriques essentielles. Il existe de très nombreuses façons de les alimenter, notamment avec des convertisseurs sous forme de circuits intégrés qui ne nécessitent que quelques composants supplémentaires.

Au-delà de l'électricité et de l'électronique de puissance, les leds constituent des exemples intéressants dans les domaines de la photométrie, de la colorimétrie, de l'optique et aussi des échanges thermiques pour ce qui est des leds de puissance.

1.8. Un domaine en pleine évolution

Cette présentation des leds, loin d'être exhaustive, n'a pour but que de montrer la richesse d'un domaine de la technologie appelé à progresser encore longtemps. De nombreuses données chiffrées concernant les leds blanches évoluent à un rythme rapide, c'est le cas par exemple de l'efficacité lumineuse et de la puissance maximale unitaire. Aussi, les données de cet article devront être actualisées avant d'être présentées à des élèves ou à des étudiants. On pourra consulter les sites internet indiqués en fin d'article.

2. TRAVAUX PRATIQUES D'ÉLECTRICITÉ AVEC DES LEDS BLANCHES

2.1. Objectifs

Les expériences décrites dans cette partie sont d'une grande simplicité. Elles illustrent les propriétés de base de l'électricité et permettent de justifier, dans le cas de l'éclairage par leds blanches, le recours à l'électronique de puissance. Aussi, elles trouveront leur place en début de cycle de formation. Ces expériences sont réalisées en classe de première STI ainsi qu'en première année de BTS électrotechnique, avec évidemment dans les deux cas une présentation, des questions et une exploitation adaptées au niveau des élèves ou des étudiants.

Cette étude a pour objectifs de mettre en évidence :

- le comportement électrique d'une led,
- les problèmes liés à l'association de plusieurs leds,
- les insuffisances du circuit source+résistance+led parfaitement adapté à la signalisation de faible puissance mais inadapté à l'éclairage par leds blanches ou à la signalisation de puissance.

2.2. Relevé de la caractéristique U(I)

On commencera par faire relever la caractéristique $I(U)$ d'une led blanche. Certes, cette expérience est plutôt classique et peut paraître ennuyeuse mais on pourra en tirer de nombreuses informations. Le montage

étant très simple, on pourra demander d'en faire le schéma. On pourra aussi faire calculer la valeur de la résistance de protection, ainsi que sa puissance, connaissant les données de constructeur, par exemple 20 mA et 3,6 V pour une led de diamètre 5 mm ou 350 mA et 3,42 V pour un luxeon de 1 W.

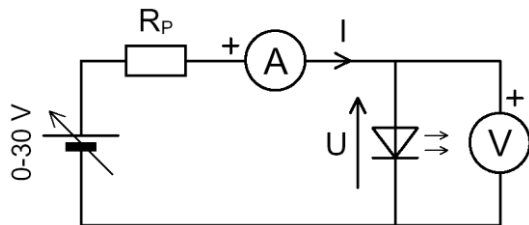


Figure 9 : Montage pour le relevé de la caractéristique $I(U)$ d'une diode électroluminescente

En fait, on donnera à chaque binôme deux led du même modèle et du même lot et ils relèveront les deux caractéristiques dans le même repère jusqu'à atteindre l'intensité nominale. On peut aussi indiquer la valeur de la résistance de protection mais pas la valeur de la tension aux bornes de la diode.

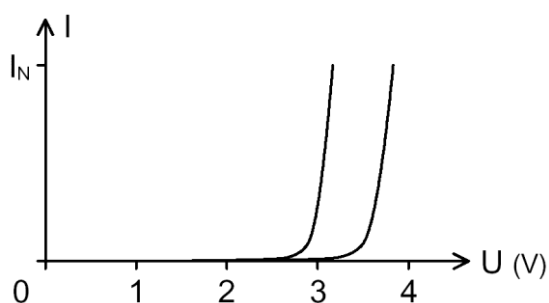


Figure 10 : Les caractéristiques $I(U)$ de deux diodes électroluminescentes d'un même lot

D'abord on constate qu'une led blanche a une caractéristique de même forme que celle d'une diode de redressement, mais avec une tension à l'état passant beaucoup plus élevée : entre 3 V et 4 V.

Ensuite, on est frappé par le fait que les deux diodes, en principe identiques, présentent des tensions assez différentes pour une même intensité. On connaît ce phénomène avec la tolérance des résistances, mais ici, les écarts entre tensions à intensité nominale peuvent largement dépasser les 5 ou 10 %. Cette dispersion des valeurs de la tension aura des conséquences notables pour l'association des leds.

2.3. Exploitation de la caractéristique

On pourra faire vérifier au passage à l'aide de documents constructeur que les leds blanches ont bien une tension identique à celle des leds bleues puisque vues du côté électrique ce sont effectivement des leds bleues (voir le paragraphe 1.2.). Les leds rouges ont une tension de l'ordre de 2 V seulement.

On pourra demander aux élèves de préciser si une led blanche se comporte plutôt comme une résistance, ou plutôt comme une source de tension ou plutôt comme une source de courant, sans oublier de justifier la réponse. La compréhension de ce point est primordiale pour aborder les convertisseurs spécifiques aux leds blanches.

On devra aussi poser le problème de l'alimentation d'une led blanche avec une ou plusieurs piles et une résistance de protection, pour la fabrication d'une lampe de poche par exemple. Avec une ou deux piles de 1,5 V ou avec un ou deux accumulateurs de 1,2 V cela n'est pas possible car on n'atteint pas une tension suffisante. On met ici en évidence un **premier problème** qui sera résolu par la suite grâce à l'électronique de puissance.

Avec trois piles de 1,5 V ou une pile de 4,5 V il est possible d'alimenter une led blanche. Seulement, là il faut se placer dans des conditions de production en série : on ne relève pas la caractéristique de chaque diode mise en oeuvre. Pour calculer la valeur de la résistance de protection, il faut donc se fier à la valeur typique U_T de la tension indiquée par le constructeur. Les diodes étant différentes, elles ne seront pas toutes traversées par la même intensité, certaines subissant une surintensité. On peut d'ailleurs faire une prédétermination de l'intensité pour les deux diodes étudiées : il suffit de tracer la caractéristique $I(U)$ de l'ensemble pile+résistance dans le même repère que celles des diodes.

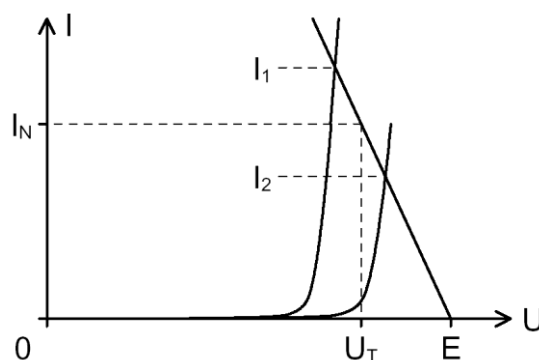
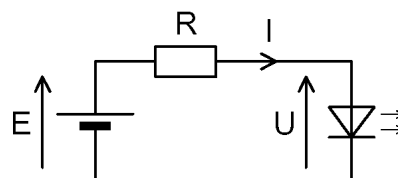


Figure 11 : Alimentation d'une led par une pile de 4,5 V ou trois de 1,5 V - Prédétermination de l'intensité

Aucun calcul n'est nécessaire pour ce tracé car on sait que cette caractéristique passe par les points $(E; 0)$ et $(U_T; I_N)$ sans qu'il y ait besoin de déterminer R.



Les intensités dans deux diodes différentes peuvent donc présenter un écart important avec l'intensité nominale et cela dans les deux sens. Ce phénomène est d'autant plus marqué que E est proche de U_T . L'incertitude sur l'intensité est encore accentuée par le fait que la tension délivrée par la pile diminue quand celle-ci s'use.

Ce montage classique ne donne donc pas de bons résultats si la tension des piles est trop proche de la tension de la led. Voilà un **deuxième problème** à signaler.

Il y a donc beaucoup de travail à effectuer sur ces caractéristiques. Une exploitation complète serait longue et fastidieuse. On pourra donc se contenter de l'essentiel dans un premier temps quitte à revenir sur les autres points plus tard, sous forme de travaux dirigés par exemple.

2.4. Préparation de l'étude expérimentale

En fait, pour que l'étude expérimentale précédente soit significative, il faut éviter de se retrouver avec deux diodes ayant des caractéristiques trop proches. Cela suppose que l'on parte d'un lot de leds suffisamment important, que l'on mesure la tension à intensité nominale de chacune puis que l'on choisisse les paires de leds. On fera en sorte que les deux leds aient la plus grande différence de tension possible, les groupes de manipulation ayant des écarts semblables.

Ce sujet peut faire l'objet d'une étude statistique avec les étudiants : après le relevé des caractéristiques complètes des deux premières diodes, on leur demande de mesurer simplement la tension U à 20 mA de quelques diodes supplémentaires. En mettant en commun les résultats de tous les élèves, on peut traiter rapidement un lot important.

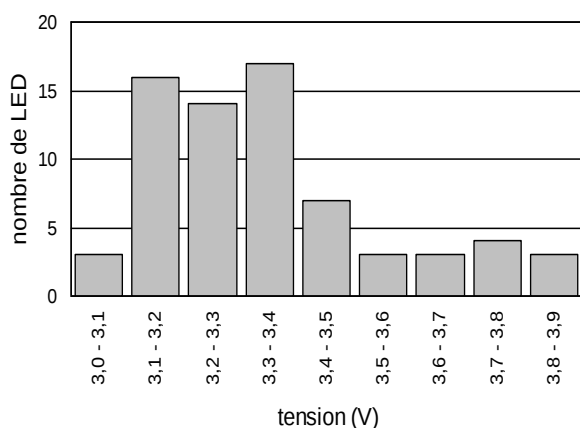


Figure 12 : Histogramme de répartition de la tension à 20 mA dans un lot de 70 diodes

La tension inverse maximale que peut supporter une diode électroluminescente n'est que de 5 V et dépasser cette valeur ne pardonne pas ! Pour protéger une led blanche contre les inversions de polarité accidentelles, il est fortement recommandé de lui souder une diode de redressement en anti-parallèle.

Le faisceau lumineux d'une led blanche est très intense et peut présenter un danger pour les yeux. Si certains distributeurs mettent en garde contre le risque de brûlure de la rétine, en revanche les constructeurs n'abordent pas ce sujet. Il est vrai qu'il existe d'autres sources de lumière très intenses dans notre environnement et que nous sommes habitués à protéger instinctivement notre vue de ces sources. On ne négligera pas pour autant ce problème lors de la préparation d'une séance de travaux pratiques. Les leds de puissance, au moins, seront montées de façon à diriger leur faisceau vers le côté ou vers le bas pour que les manipulateurs puissent regarder leur plan de travail sans danger.

2.5. Associations de leds

En raison d'une puissance unitaire assez faible, ou pour obtenir un éclairage réparti, il est souvent nécessaire d'alimenter plusieurs leds simultanément.

A cause de la dispersion des valeurs de la tension, l'association en parallèle donne de mauvais résultats : les leds de tensions différentes ne brillent pas autant. Par contre, toutes les leds d'un même lot brillent du même éclat quand elles sont associées en série.

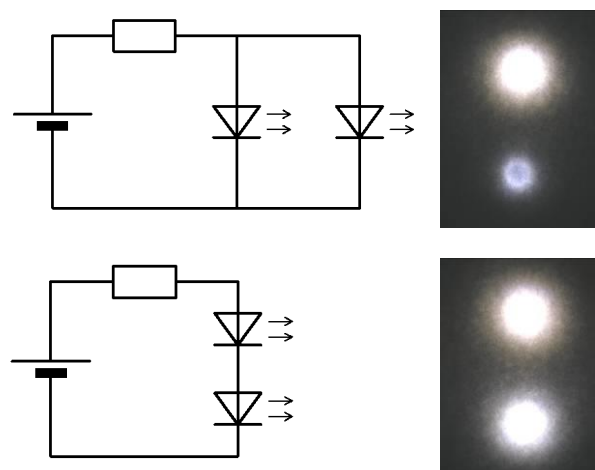


Figure 13 : Deux leds blanches en parallèle puis en série. Les faisceaux lumineux sont observés à travers une feuille de papier plaquée sur les boîtiers des leds.

Ainsi, une même intensité donne une même quantité de lumière pour des leds d'un même modèle, même si leurs tensions sont différentes. Sachant cela, les caractéristiques $I(U)$ permettent de comprendre ce qui se passe lors d'une association en parallèle.



Pour aborder cette question avec les étudiants, on peut leur demander de chercher expérimentalement la réponse à la question suivante :

"Pour que deux leds brillent autant, faut-il qu'elles soient soumises à la même tension ou qu'elles soient traversées par la même intensité ?"

Ils devront penser à réaliser les deux montages précédents, parallèle et série. C'est l'occasion d'utiliser une des propriétés de base de l'électrocinétique : "en série des dipôles sont traversés par le même courant, en parallèle ils sont soumis à la même tension".

Pour la réalisation pratique, on effectue le montage avec les deux leds déjà étudiées et la même résistance de protection. Ainsi les intensités seront plus faibles et la led de plus faible tension ne subira pas de surintensité lors du montage en parallèle.

Il n'est pas possible de comparer directement les intensités lumineuses des diodes sans être ébloui. Pour atténuer les faisceaux lumineux, on posera une feuille de papier ordinaire directement sur les boîtiers (voir figure 12).

Les leds devant être associées en série pour briller de la même façon, la tension nécessaire pour les alimenter devient vite assez élevée et le "premier problème" évoqué plus haut prend encore plus d'importance.

Par exemple, dans le cas du rétroéclairage d'un écran LCD, il est impératif que toutes les leds brillent avec la même intensité pour que l'éclairage de l'écran soit uniforme. Il faut alors les brancher en série, ce qui nécessite une tension assez élevée (14,4 V pour six leds). Si l'appareil est portable, il ne dispose que de deux ou trois piles ou accus, ce qui est largement insuffisant. On pourra alors indiquer aux élèves qu'un circuit appelé "hacheur élévateur" permet de résoudre ce problème et sera étudié ultérieurement.



Figure 14 : Trois LED blanches soudées sur le culot de l'ampoule à incandescence d'une lampe de poche.

On trouve parfois des associations de leds en parallèle. C'est par exemple le cas des assemblages permettant de remplacer l'ampoule à incandescence d'une lampe de poche classique alimentée en 4,5 V : par

manque de place, une seule résistance est mise en série avec les leds en parallèle (figure 14). Avec cette configuration, il est indispensable de choisir des leds de tensions suffisamment proches, ce qui complique sérieusement la mise en oeuvre.

2.6. Etude du rendement

Il s'agit maintenant de montrer que le circuit classique source+résistance+led présente un mauvais rendement quand la tension d'alimentation est très supérieure à celle de la led ou des leds en série.

On propose aux élèves de se placer dans le cas où l'alimentation est une batterie de 12 V, comme c'est le cas des systèmes d'éclairage de l'association LUTW (paragraphe 1.2.6.).

On demande de calculer le rendement pour une, pour deux et enfin pour trois leds. Cela n'est pas évident pour tous les étudiants : à quoi correspondent ici le " P_U " et le " P_A " de la célèbre formule ? On pourra orienter leur réflexion en les invitant à prendre du recul et à considérer l'objet "appareil d'éclairage" dans son environnement. P_U est évidemment la puissance consommée par la ou les diodes puisque le circuit sert à éclairer. P_A correspond à ce que l'on paye ou à ce que l'on investit pour que le circuit fonctionne, c'est donc la puissance fournie par la batterie. Si le rendement est mauvais, il faudra un panneau solaire plus grand pour recharger la batterie.

Pour une led avec les notations de la figure 11 :

$$\rho = \frac{P_U}{P_A} = \frac{U \cdot I}{E \cdot I} = \frac{U}{E}$$

Pour N leds de tension U : $\rho = \frac{N \cdot U}{E}$

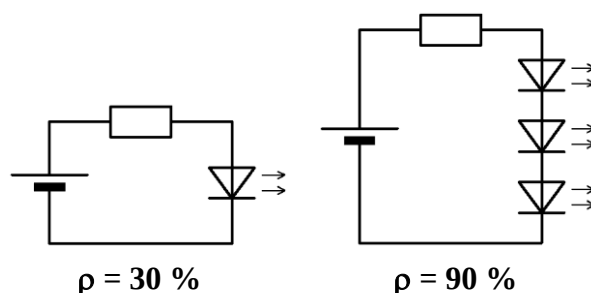


Figure 15 : Rendement du circuit source+résistance+led
source : $E = 12 \text{ V}$ - led : $U = 3,6 \text{ V}$

En fait, pour une valeur de l'intensité imposée, les pertes sont proportionnelles à la tension aux bornes de la résistance. Plus il y a de leds, plus cette tension est faible et meilleur est le rendement, le nombre N de leds étant limité par la tension de la source.



On met là en évidence un **troisième problème** : le circuit utilisant une résistance de protection présente un très mauvais rendement quand la tension de la source est très supérieure à celle de la ou des leds. Avec le développement des leds de puissance, c'est justement cette situation qui devient plus courante. En effet, une led de puissance présente une meilleure efficacité lumineuse qu'un ensemble de leds, elle occupe beaucoup moins de place et la mise en oeuvre est plus simple : une seule suffit.

Cette fois-ci, c'est le hacheur série qui permettra d'obtenir un rendement acceptable, comme dans les luminaires de la figure 8.

On peut aussi demander de déterminer le nombre maximal de leds en série que l'on peut alimenter avec la source choisie, puis de comparer le rendement pour ce nombre maximal de leds au rendement pour une seule led.

Comme application de l'étude précédente, on peut demander de justifier les nombres de leds utilisés dans les premiers luminaires de l'association LUTW au Népal (figures 16 et 17) puis d'en proposer un schéma de câblage.



Figure 16 : Un luminaire à 6 leds pour une batterie de 12 V.



Figure 17 : Un luminaire à 9 leds pour une batterie de 12 V.

Les élèves comprennent très vite le principe et regroupent bien les leds par branches de trois, mais beaucoup proposent un schéma avec une résistance de protection unique en série avec toutes les branches.

Avec cette disposition, la dispersion des tensions des leds peut entraîner des différences entre les intensités des différentes branches, comme pour des leds directement en parallèle. De plus, un mauvais contact dans une branche entraîne une surintensité dans les autres branches. La meilleure solution, pour une source de 12V, est donc de constituer des branches de trois leds, chaque branche ayant sa propre résistance de protection.

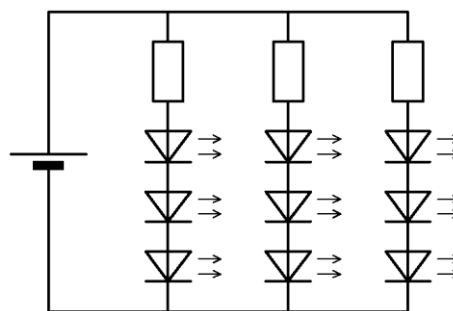


Figure 18 : Schéma d'un luminaire à 9 leds pour une batterie de 12 V, avec une résistance de protection par branche.

Il est intéressant de faire dimensionner la résistance de protection d'une diode de puissance, de 1 W au moins, pour une tension d'alimentation de 12 V : très souvent, les élèves oublient de déterminer la puissance qu'elle dissipe. Une résistance de 1/4 W ne tient pas longtemps et une résistance de puissance maximale suffisante s'échauffe de manière significative.

On peut également faire dimensionner une résistance variable qui en série avec la led de puissance et la résistance de protection permettra de faire varier la luminosité, en imposant par exemple un facteur de 10 entre les intensités maximale et minimale. Le résultat de cette étude est intéressant : la résistance variable trouvée est beaucoup plus chère que la led, prend énormément de place et pèse un bon poids. On arrive à une véritable aberration technologique !

On pourra faire le parallèle avec les gradateurs de lumière que les étudiants connaissent bien. Ils savent que la commande actionne un potentiomètre mais certains pensent que ce potentiomètre est en série avec la lampe. En fait, cela exigerait un rhéostat proprement monstrueux.

Cette dernière étude amène à un constat : pour les diodes électroluminescentes, le recours à l'électronique de puissance est déjà indispensable pour une puissance de 1 W !



2.7. Idées fausses sur les leds

Ce premier contact avec les leds blanches peut être l'occasion de rectifier certaines idées fausses sur les diodes électroluminescentes en général.

Certains élèves ou étudiants pensent que la lumière est émise par un filament. En observant de près une led ayant un boîtier translucide, on peut effectivement distinguer un petit fil. En fait, il s'agit de la connexion entre l'anode et la cathode qui porte la puce émettrice. On pourra faire observer de près une led traversée par un courant très faible pour éviter l'éblouissement : on voit bien que la lumière provient d'une surface et non d'un filament.

Les premières leds rouges puis vertes avaient toutes un boîtier coloré et c'est encore souvent le cas aujourd'hui. De-là à penser que c'est ce boîtier qui donne la couleur à la lumière, il n'y a qu'un pas. On pourra donc montrer aux élèves des diodes de différentes couleurs ayant un boîtier translucide. Il est alors évident qu'il ne s'agit pas d'une lumière blanche filtrée.

2.8. Conclusion de l'étude électrique

Pour exploiter au mieux l'étude de convertisseurs électroniques spécifiques aux leds, il est important que les élèves aient bien compris les points qui suivent.

D'abord, le comportement électrique d'une led (partie 2.2.) :

→ Une led traversée par un courant se comporte quasiment comme **une source de tension**. Le convertisseur qui l'alimente doit donc se comporter comme une source de courant.

Ensuite, la conséquence la plus importante de la dispersion de la tension à intensité nominale dans un lot de leds (partie 2.5.) :

→ Pour que plusieurs leds brillent autant les unes que les autres, il faut les **brancher en série**.

Et enfin, les limites du circuit classique constitué d'une source et d'une résistance de protection en série avec la led (parties 2.3. et 2.6.) :

→ En raison d'une tension nominale assez élevée et de la nécessité d'associer les leds en série, **la tension d'alimentation disponible est souvent insuffisante** (premier problème).

→ Quand la tension d'alimentation est légèrement supérieure à la tension requise, c'est la dispersion des valeurs de la tension des leds et les fluctuations de la tension d'alimentation qui **entraînent de grandes variations de l'intensité** (deuxième problème).

→ Quand la tension d'alimentation est largement supérieure à la tension totale des leds, **le rendement est mauvais** (troisième problème).

Ainsi, on met en évidence des problèmes que seule l'électronique de puissance permet de résoudre. Les convertisseurs électroniques qui seront étudiés dans les prochains articles permettent d'obtenir un rendement excellent, mais en plus ils garantissent une valeur précise de l'intensité grâce à des fonctions de régulation.

Led ou del ?

Certains lecteurs auront peut-être déploré dans cet article l'utilisation du mot "led" venant de l'anglais au lieu du mot "del" pour "diode électroluminescente" qui vient du français. C'est tout simplement que le mot "led" est celui que l'on rencontre le plus souvent dans les documents en français. A noter que "del" et "led" figurent tous les deux dans Le Petit Larousse 2006.

Références

- Sites généralistes consacrés aux leds :

<http://www.ledsmagazine.com>

<http://www.led-fr.net>

- Constructeurs :

<http://www.lumileds.com>

<http://www.LaminaCeramics.com>

<http://www.zled.com>

- Distributeurs :

<http://www.luxeonstar.com>

<http://www.ledsupply.com>

<http://www.conrad.fr>

<http://www.radiospares.fr>

<http://fr.farnell.com>

<http://www.mouser.com>

<http://www.superbrightleds.com>

<http://www.alscomposants.com>

<http://www.totalstyles.fr>

<http://www.led1.de>

- Lampes à leds :

<http://www.lampe-led-sans-pile.com>

<http://www.lampesdepoeche.com>

<http://www.howell-lighting.com.cn>

- Association Light Up The World :

<http://www.lutw.org>

Merci à Hervé Ricard, professeur de Physique Appliquée au lycée Clément Ader d'Athis-Mons, pour ses relectures de cet article.



Etude d'une suspension magnétique active

Pascale COSTA*, François CARRERE**

* Lycée Raspail Paris 14

** Société S2M, BP 2282, 2 rue des champs, 27950 St-Marcel

Cette étude d'une suspension magnétique active pour turbocompresseur a fait l'objet du **sujet d'automatique et d'informatique industrielle de l'agrégation de Génie Electrique session 2005**. Cet article reprend l'approche de ce sujet et correspond à une suite de questions-réponses.

Après avoir décrit le système étudié et précisé les avantages d'une suspension magnétique, les caractéristiques d'une suspension en terme de raideur seront définies. La constitution de la suspension magnétique sera ensuite abordée : l'électroaimant, l'électronique de commande, l'amplificateur de puissance et les détecteurs.

Nous évoquerons ensuite la commande multi-variable dite de « translation basculement ». L'influence des modes flexibles et des effets gyroscopiques sera ensuite décrite. Nous terminerons enfin par la description d'une méthode de contrôle anti-vibratoire.

1. Présentation du système étudié

Le système étudié, annexe 1, est un turbocompresseur industriel de 110 kW destiné à une station d'épuration des eaux. Il permet la ré-oxygénation de l'air ; le débit d'air souhaité est de 3000m³/h avec une pression de 1 à 2 bars. L'air doit être dépourvu de toute trace d'huile de lubrification.

Lorsqu'elles sont chargées du choix d'un turbocompresseur, les sociétés d'ingénierie fixent les critères suivants :

- coût initial, efficacité thermodynamique,
- volume / poids,
- rendement,
- fiabilité / disponibilité,
- absence d'huile (maintenance, risque d'accident),
- propreté de l'air
- durée de vie de l'équipement (30 ans et plus),
- démarrage rapide (disponibilité immédiate).

Si l'efficacité thermodynamique ne dépend pas du palier (assurant le guidage en rotation), mais de la conception du constructeur ; le Palier Magnétique Actif (PMA) optimise tous les autres critères. En effet, il réduit et allège le turbocompresseur grâce à la suppression de son installation de lubrification (habituellement 4 à 5 fois plus grande que le turbocompresseur lui-même). Il permet en effet d'optimiser les dimensions par des technologies basées sur la haute vitesse. Ainsi, même en tenant compte du coût du PMA, le coût initial de l'équipement diminue. L'absence de toute pièce mécanique soumise à l'usure lui procure une fiabilité et une disponibilité inégalées. La flexibilité du turbocompresseur est accrue car le PMA permet d'éviter la mise en température nécessaire pour atteindre un fonctionnement nominal à pleine vitesse.

La figure 1 présente la solution conventionnelle sur paliers à huile où 25% de la puissance du moteur est perdue dans les paliers et la boîte de vitesse.

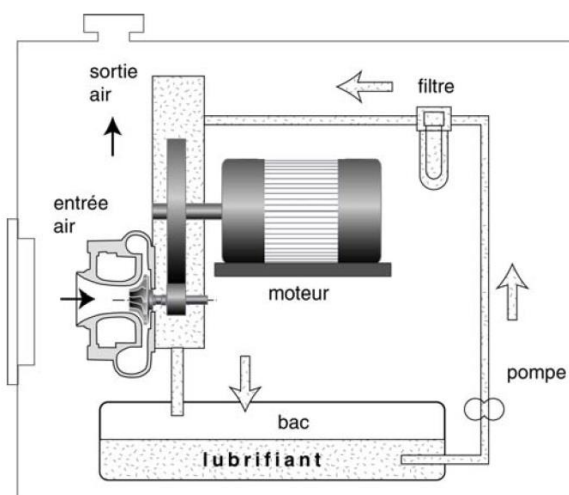


Figure 1 : Compresseur avec boîte de vitesse et circuit de lubrification (engrenages, paliers)

Dans le cas d'un entraînement direct sur paliers magnétiques (figure 2), il y a seulement 8% de pertes.

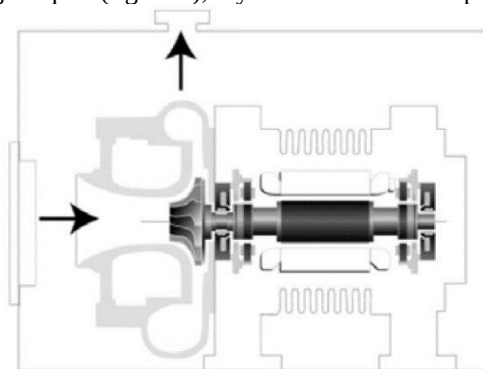


Figure 2 : Compresseur avec entraînement direct sur PMA

1.1. Qu'est ce qu'une suspension magnétique ?

Une suspension magnétique est un système permettant de contrôler 5 degrés de liberté d'une pièce par interaction magnétique, la motorisation pilotant le degré de liberté restant. Dans un turbocompresseur, quatre degrés de liberté contrôlent la position radiale et un degré de liberté contrôle la position axiale. Une telle suspension nécessite au moins deux paliers radiaux et une butée axiale solidaires du rotor.

1.2. Cinq degrés de liberté à contrôler

La suspension magnétique doit assurer à tout instant l'équilibre du rotor. Cinq asservissements sont par conséquent nécessaires. Une suspension magnétique (figure 3) est composée de deux paliers radiaux (chacun étant constitué de deux asservissements de position) et d'une butée axiale (comportant un asservissement de position).

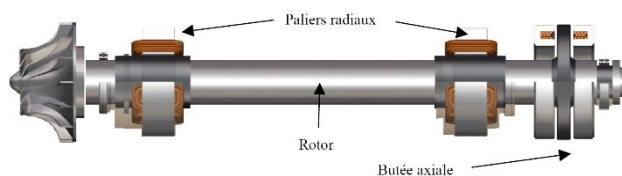


Figure 3 : Suspension magnétique

La figure 4 définit les cinq axes d'asservissement usuellement utilisés. Les axes V13 (représentés sur la figure 4 par V_1V_3) et W13 (respectivement V24 et W24) définissent le plan 1-3 (respectivement 2-4) du palier 1 (respectivement palier 2). Les axes V13 et V24 d'une part, W13 et W24 d'autre part, définissent respectivement les plans V et W.

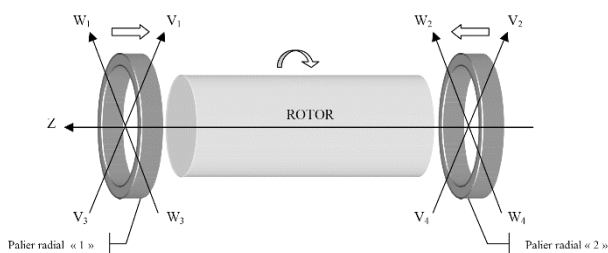


Figure 4 : Convention d'axes

1.3. Palier radial

Le schéma de la figure 5 représente un palier radial. La position du rotor est mesurée au moyen de capteurs qui appréhendent en permanence les déplacements par rapport à la position de référence.

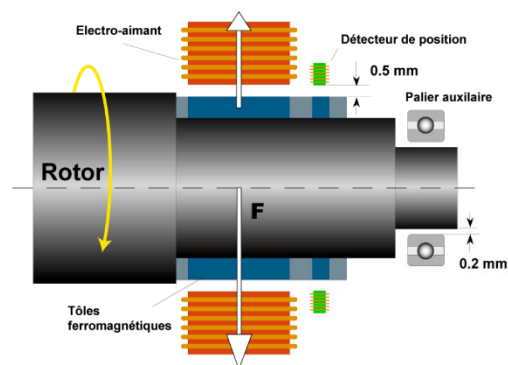


Figure 5 : Palier radial

La présence de deux électro-aimants (figure 6) respectant la symétrie centrale permet d'exercer une force $\vec{F}_{\text{palier}} = \vec{F} + \vec{F}'$ sur le rotor, ici suivant l'axe V13. Un mécanisme similaire permet de créer une force suivant W13.

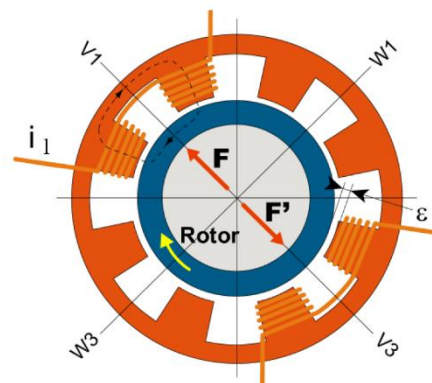


Figure 6 : Centrage du rotor suivant V13

1.4. Butée axiale

Un palier axial (figure 7) est fondé sur le même principe que le palier radial, le rotor étant constitué par un disque dans un plan perpendiculaire à l'axe de rotation et en face duquel se trouvent les électro-aimants. Ce palier sert de butée axiale. Le détecteur de position est souvent situé à l'extrémité de l'arbre. La conception de cette chaîne d'asservissement est en général moins complexe, car les problèmes de flexions, de résonances et de largeur d'entrefer sont peu critiques suivant la direction axiale Z.

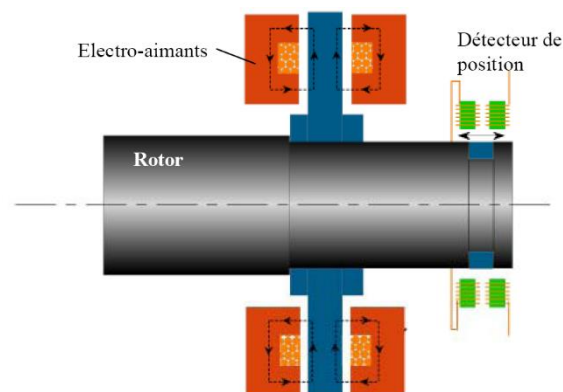


Figure 7 : Palier axial



2. Caractéristiques générales d'une suspension

L'ensemble du système "palier + rotor en suspension" présente différents modes de vibrations avec et sans déformations. Nous nous intéressons, dans un premier temps pour simplifier l'étude, aux modes rigides (encore appelés modes de paliers) pour lesquels le rotor ne se déforme pas. Ce paragraphe a pour but de rappeler les caractéristiques d'une suspension et de définir ainsi les "propriétés" des paliers afin de comprendre les avantages apportés par une suspension magnétique par rapport à une suspension classique.

L'étude concerne uniquement les mouvements radiaux du rotor sur un axe noté x et représentant indifféremment les axes V13, V24, W13 ou W24.

2.1. Cas d'une suspension classique

Les modes rigides peuvent s'analyser en utilisant l'approche mécanique classique d'un système en suspension. Les perturbations s'appliquant sur la suspension sont représentées par la force $\vec{F}_{\text{ext}}$. Cette force extérieure perturbatrice modélise les effets de balourds, d'efforts, ou toute autre force liée à l'utilisation de la suspension.

Nous allons, tout d'abord, étudier les modèles classiques d'une suspension. L'origine des axes ($x = 0$) est obtenu pour la position de repos de la suspension.

La qualité de la suspension dépendra du comportement fréquentiel de sa raideur, on note $K(p) = \frac{F_{\text{ext}}(p)}{X(p)}$ la fonction de transfert de la raideur.

• Système masse-ressort

On suppose, dans un premier temps, que la suspension idéalisée sur un seul axe x , se comporte comme un système masse-ressort (figure 8).

On note :

- k la constante de raideur du ressort,
- M la masse équivalente sur un axe du rotor à sustenter.

La masse M peut se déplacer sur l'axe x sans frottement.

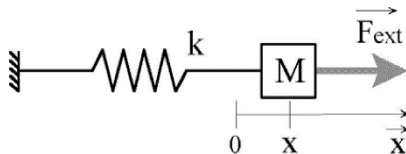


Figure 8 : Système masse-ressort

En appliquant la relation fondamentale de la dynamique suivant l'axe x , on obtient alors :

$$M \frac{d^2 x}{dt^2} = F_{\text{ext}} - kx$$

et une raideur définie par :

$$K_{\text{mr}}(p) = k \left(1 + \left(\frac{p}{\omega_0} \right)^2 \right) \text{ avec } \omega_0 = \sqrt{\frac{k}{M}}.$$

La raideur est alors nulle à la pulsation ω_0 (pulsation de résonance). En théorie, lors de la moindre perturbation (F_{ext}), la suspension peut entrer en résonance. En pratique, il existe toujours un faible amortissement dû aux pertes qui empêche la suspension d'entrer en résonance d'elle-même.

• Système masse-ressort avec amortissement

Pour passer cette pulsation propre, il est nécessaire de créer un amortissement au niveau du palier. L'étude se fait alors selon le système modélisé figure 9. On garde les mêmes notations que précédemment, avec en plus :

- a le coefficient de frottements visqueux de la suspension suivant l'axe x .

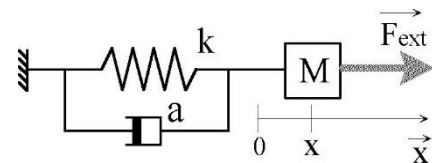


Figure 9 : Système masse-ressort-amortisseur

En reprenant la même démarche que précédemment, on obtient :

$$M \frac{d^2 x}{dt^2} = F_{\text{ext}} - kx - a \frac{dx}{dt}$$

La fonction de transfert du système H étant l'inverse de la raideur K , on en déduit :

$$H_{\text{mra}}(p) = \frac{1/k}{1 + 2\xi \frac{p}{\omega_0} + \left(\frac{p}{\omega_0} \right)^2} \text{ avec } \xi = \frac{a}{2\sqrt{kM}}.$$

Si la force perturbatrice est un échelon d'amplitude F_0 , $x(t)$ peut donc évoluer selon la figure 10.

On définit l'amortissement critique par $a_c = 2\sqrt{kM}$.

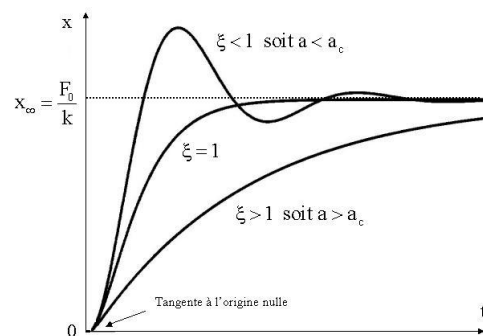


Figure 10 : Evolution de $x(t)$

• Système masse-ressort amorti avec contrôle intégral

La suspension précédente est limitée dans les basses fréquences où la raideur est constante et égale à k . Pour amener le rotor à sa position de référence, il est nécessaire de créer une action intégrale (intégration de la variable x qui s'oppose au mouvement). Ceci est possible en utilisant, comme nous le verrons par la suite, un palier magnétique actif (PMA).

L'action générée par le palier doit donc être équivalente à celle d'un système masse-ressort amorti avec contrôle intégral. L'étude se fait alors selon le système modélisé figure 11. On garde les mêmes notations que précédemment, avec en plus :

- b le gain de l'intégrateur.

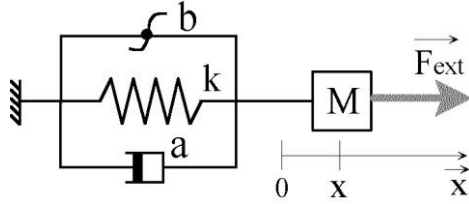


Figure 11 : Système masse-ressort amorti avec contrôle intégral

On a alors :

$$M \frac{d^2 x}{dt^2} = F_{\text{ext}} - kx - a \frac{dx}{dt} - b \int_{-\infty}^t x(u) du$$

soit $K_{\text{mrai}}(p) = k + ap + \frac{b}{p} + Mp^2$.

On montre alors facilement que cette suspension ne possède plus d'erreur statique :

$$x_{\infty} = \lim_{p \rightarrow 0} pX(p) = \lim_{p \rightarrow 0} p \frac{F_{\text{ext}}(p)}{K_{\text{mrai}}(p)} = \lim_{p \rightarrow 0} p \frac{F_0/p}{K_{\text{mrai}}(p)} = 0.$$

2.2 Cas d'une suspension magnétique

Dans une suspension magnétique active, les caractéristiques introduites précédemment (k , a et b) proviennent des performances de l'asservissement, et sont donc paramétrables. On considère dans cette partie que les quatre mouvements radiaux et le mouvement axial sont indépendants, on a donc ici une stratégie de commande monovariable sur les cinq axes, dite axe par axe.

Le principe d'un palier magnétique radial est alors simple. Une variation de position du rotor entraîne une variation de l'entrefer. Cet écart est comparé à une référence de position, le signal d'erreur est ensuite envoyé au correcteur. La correction peut se faire de façon analogique ou numérique. Le correcteur calcule la consigne nécessaire pour ramener le rotor dans sa position de référence. Cette consigne est finalement amplifiée de façon à alimenter les électro-aimants. Ces derniers vont créer les forces de rappel nécessaires au guidage du rotor.

Le schéma bloc d'un palier radial, figure 12, suivant un axe (nommé ici x), est donc le suivant :

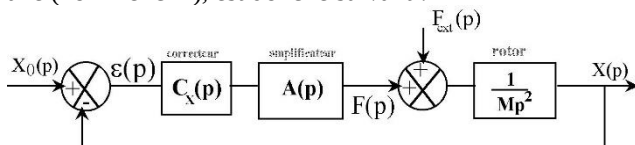


Figure 12 : Schéma bloc d'un palier radial (contrôle axe par axe)

Il y a deux entrées sur le système étudié : l'une est la consigne de position (elle est fixe et nulle) et l'autre est

la force extérieure perturbatrice qui modélise les effets de balourd, ou toute autre force liée à l'utilisation de la suspension.

L'amplificateur et le capteur de position seront étudiés dans la partie suivante. L'étude sera ainsi simplifiée en prenant arbitrairement : $A(p) = 1$.

On s'intéresse au comportement du système lorsque celui-ci fonctionne en régulation puisqu'il faut maintenir l'entrefer constant quelques soient les perturbations.

•Système non perturbé

On étudie, dans un premier temps, le système non perturbé soit $F_{\text{ext}}(p) = 0$.

La fonction de transfert en boucle ouverte, notée $H_{\text{BOX1}}(p)$ est donnée par :

$$H_{\text{BOX1}}(p) = C_x(p)A(p) \frac{1}{Mp^2} D(p) = \frac{C_x(p)}{Mp^2}$$

La suspension se comporte alors comme un double intégrateur en boucle ouverte. Le système oscillera à la pulsation $\sqrt{\frac{k_0}{M}}$. Pour stabiliser ce système, on introduit

un premier correcteur à avance de phase :

$$C_x(p) = C_{x1}(p) = k_0 \frac{1 + \alpha \tau_1 p}{1 + \tau_1 p}.$$

Les différents paramètres de ce correcteur fixent les performances de la suspension. La valeur de α_0 fixe le maximum de phase, celle de τ_1 permet d'obtenir en boucle ouverte un maximum de phase pour la fréquence propre f_0 de la suspension. Enfin, k_0 règle la marge de phase.

•Système perturbé

Si on tient compte de l'influence des perturbations, l'erreur statique est non nulle. Il est donc nécessaire d'introduire un second correcteur proportionnel et intégral en amont de la perturbation :

$$C_{x2}(p) = 1 + \frac{1}{\tau_2 p} = \frac{1 + \tau_2 p}{\tau_2 p}$$

Un simple correcteur intégrateur serait impossible, le système serait instable puisqu'il possède intrinsèquement deux intégrations.

La constante τ_2 est choisi de façon à ne pas trop diminuer la marge de phase obtenue avec le correcteur $C_{x1}(p)$.

•Synthèse

La suspension, ainsi corrigée, possède maintenant toutes les caractéristiques du système masse-ressort-amorti avec contrôle intégral.

La fonction de transfert de la raideur $K_f(p)$ vaut en effet :

$$K_f(p) = \frac{F_{\text{ext}}(p)}{X(p)} = Mp^2 + C_x(p)$$

avec $C_x(p) = k_0 \frac{1 + \alpha \tau_1}{1 + \tau_1} \frac{1 + \tau_2 p}{\tau_2 p}$.

Dans l'exemple étudié, on souhaite avoir les performances suivantes $f_0 = 83$ Hz, une raideur



minimale de $1,3 \text{ N}/\mu\text{m}$, une marge de phase de 30° et la masse équivalente sur un axe du rotor est de $12,3 \text{ Kg}$. Ces données numériques permettent de calculer le correcteur, on obtient ainsi la réponse en fréquence de la raideur donnée figure 13.

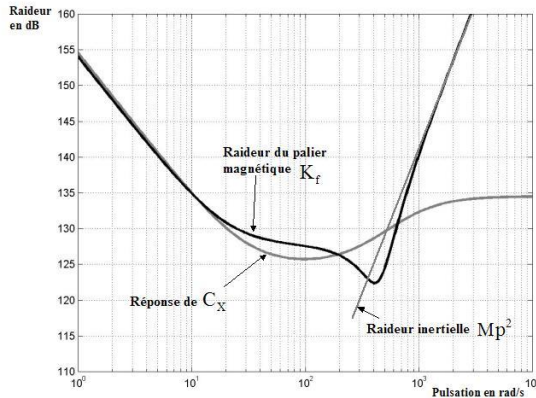


Figure 13 : Réponse en fréquence de la raideur

3. Réalisation d'un palier magnétique actif

La figure 14 donne une vue schématique du système de contrôle dans un plan avec les principaux éléments de la chaîne permettant l'asservissement en position du rotor dans ses paliers. Nous étudierons dans cette partie l'électro-aimant, l'électronique de commande, l'amplificateur de courant ainsi que le détecteur de position.

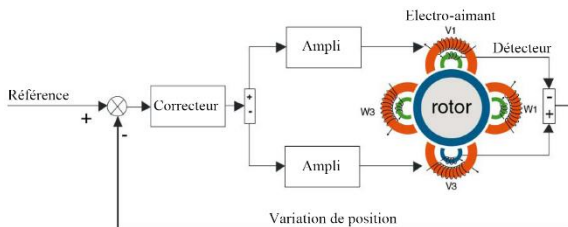


Figure 14 : Vue schématique du système de contrôle dans le plan V

3.1. Electro-aimant

L'électro-aimant (figure 15) est constitué d'une culasse en matériau ferro-magnétique autour de laquelle est bobiné un enroulement de N spires. Le rotor est maintenu en lévitation en commandant le courant i qui circule dans l'enroulement. Les perméabilités magnétiques relatives des tôles magnétiques de la culasse de l'électro-aimant et du rotor sont considérées de valeur infinie. La section du tube de flux magnétique reste partout égale à la section d'un pôle (flux de fuites négligées).

L'entrefer e est variable puisqu'il dépend de la position du rotor.

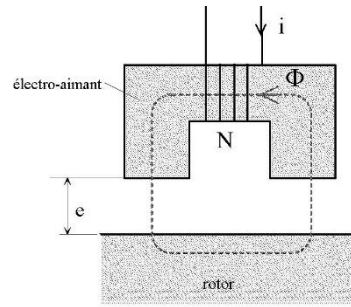


Figure 15 : Electro-aimant

La force d'un électro-aimant est donnée par :

$$F = \frac{B^2 S}{2\mu_0}$$

avec B l'induction magnétique dans l'entrefer (noté e), S la surface active utile et μ_0 la perméabilité magnétique du milieu constituant l'entrefer.

On démontre, en appliquant le théorème d'Ampère, que :

$$B = \frac{\mu_0 N i}{2e}$$

La force produite s'exprime alors par :

$$F = \lambda \left(\frac{i}{e} \right)^2 \text{ avec } \lambda = \frac{\mu_0 S N^2}{8}$$

La force F fournie par l'électro-aimant est donc uniquement une force d'attraction dépendant du carré du courant. Le sens du courant est par conséquent indifférent. Pour contrôler un axe, il est nécessaire d'avoir 2 électro-aimants situés de part et d'autre du rotor. Pour contrôler le plan, il en faudra donc 4.

Si on considère de petites variations autour d'un point de fonctionnement F_0 obtenu pour un courant $i = I_0$ et un entrefer e_0 , on obtient :

$$F = F_0 + dF \text{ avec } dF = -2 \frac{F_0}{e_0} de + 2 \frac{F_0}{i_0} di$$

Le terme $-2 \frac{F_0}{e_0}$ est appelé raideur négative ou raideur

instable du palier magnétique. En effet, ce terme est homogène à une raideur d'un point de vue mécanique (N/m) et est négatif.

Lorsque la raideur d'un palier est positive, le palier peut être comparé à un ressort car il s'oppose au déplacement. Il est stable.

Par contre, lorsque la raideur est négative, le palier favorise l'écartement par rapport à sa position d'équilibre. Si e augmente, la force appliquée par l'électro-aimant diminue, le rotor s'écarte davantage. On parle alors d'instabilité en boucle ouverte. Il est donc nécessaire de tenir compte de cette raideur négative dans l'asservissement.

3.2. Electronique de commande

Le but de l'électronique de commande utilisée est de réaliser un transfert linéaire entre le signal de sortie du régulateur et la force résultante sur un axe. Plusieurs

méthodes peuvent être envisagées pour commander les électro-aimants.

On peut commander les électro-aimants en classe B : un seul électro-aimant est alors alimenté. Cette commande présente un certain nombre d'inconvénients : utilisation de linéariseurs permettant de réaliser un transfert linéaire entre la commande et la force résultante, mauvais comportement dynamique, force maximale réduite...

Il est préférable de commander simultanément les électro-aimants (classe A). Chaque électro-aimant est modulé en opposition de phase autour d'une valeur de repos. Le schéma synoptique de cette commande est donnée figure 16. On suppose dans ce paragraphe que l'amplificateur de courant se comporte comme un gain pur k_a .

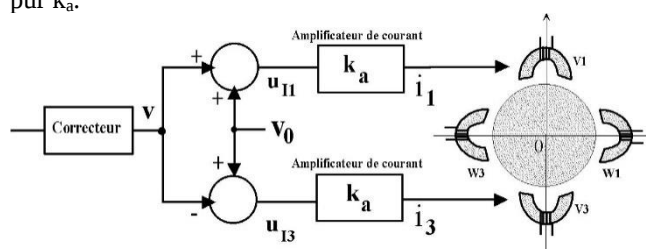


Figure 16 : Electronique de commande en classe A

Lorsque le rotor est centré, l'entrefer est constant et de valeur e_0 . La sortie v du correcteur est alors nulle ; un courant i_1 (resp. i_3) valant $k_a v_0$ produit une force F_1 (resp. F_2) tel que $F = F_1 - F_3 = 0$.

Si on suppose que le rotor effectue un déplacement x dans le sens positif de l'axe V13. La force résultante F (orientée suivant V13) appliquée au rotor s'exprime alors par :

$$F = F_1 - F_3 = \lambda \left(\frac{i_1}{e_1} \right)^2 - \lambda \left(\frac{i_3}{e_3} \right)^2 = \lambda \left[\left(\frac{i_1}{e_0 - x} \right)^2 - \left(\frac{i_3}{e_0 + x} \right)^2 \right]$$

$$F = \lambda \left[\left(\frac{k_a(v + v_0)}{e_0 - x} \right)^2 - \left(\frac{k_a(v - v_0)}{e_0 + x} \right)^2 \right]$$

si $x \ll e_0$, on a :

$$F \approx \lambda \left[\left(\frac{k_a(v + v_0)}{e_0} \right)^2 - \left(\frac{k_a(v - v_0)}{e_0} \right)^2 \right]$$

$$\text{soit } F \approx 4\lambda \left(\frac{k_a}{e_0} \right)^2 v_0 v.$$

Pour des déplacements tels que $x \ll e_0$, la force résultante suivant l'axe V est alors proportionnelle à v (signal de sortie du correcteur).

On a donc $A(p) = 4\lambda v_0 \left(\frac{k_a}{e_0} \right)^2$, la commande devient ainsi linéaire.

Chaque amplificateur de puissance devra donc fournir d'une part la charge statique F_0 et d'autre part une force instantanée F_p permettant de compenser la force extérieure perturbatrice F_{ext} . La dynamique de l'amplificateur de puissance étant limité en tension et en courant, au delà d'une certaine fréquence, l'amplificateur de puissance saturera et on ne pourra plus moduler la force dans sa totalité.

3.3. Amplificateur de courant

L'amplificateur de courant, mentionné figure 16, est obtenu par un convertisseur à découpage asservi en courant.

Ce convertisseur (figure 17) est réversible en tension afin de pouvoir imposer des variations rapides tant à la croissance qu'à la décroissance du courant. On suppose une stratégie de commande où la fréquence de découpage est fixe.

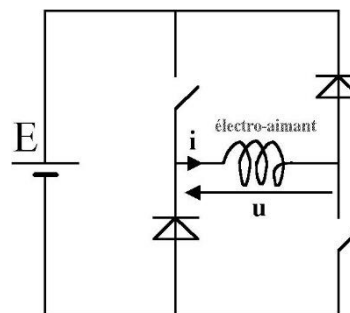


Figure 17 : Convertisseur

3.4. Détecteur de position

Le capteur est un élément essentiel de l'asservissement. Il doit fournir une image de l'entrefer sans engendrer de contact entre la partie fixe et la partie mobile de la suspension. Dans un palier radial, la position du rotor est déterminée au moyen de quatre détecteurs inductifs montés en pont et utilisant le principe du pont de Wheatstone (figure 18) à courant porteur alternatif U_d . Deux détecteurs permettent de mesurer le déplacement suivant l'axe V (mesure en U_V), les deux autres suivant l'axe W (mesure en U_W).

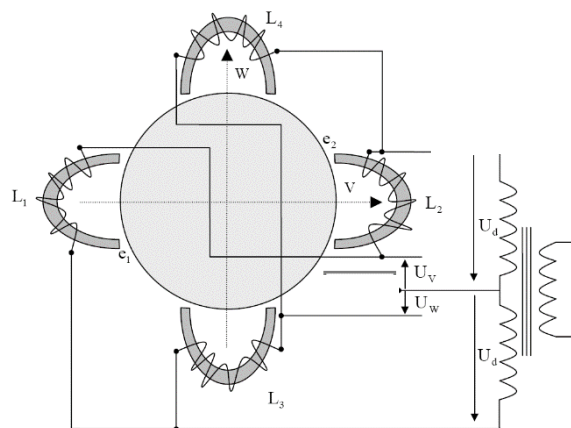


Figure 18 : Détecteurs de position montés en pont de Wheatstone



Chaque détecteur est un capteur à inductance variable, l'entrefer variant suivant la position du rotor. On considère les capteurs purement inductifs d'impédance équivalente L_i :

$$L_i = \frac{\mu_0 S_d N^2}{2e_i}$$

avec S_d la section droite de l'entrefer, e_i l'épaisseur de l'entrefer et N le nombre de spires.

A l'équilibre, le rotor est centré ($e_i = e_0$), on a alors $L_1 = L_2 = L_3 = L_4$ soit $u_v = u_w = 0$ V.

On montre que $\frac{U_v(p)}{U_d(p)} = \frac{L_2 - L_1}{L_1 + L_2}$.

Si le rotor effectue un déplacement x dans le sens positif de l'axe V ($e_1 = e_0 + x$ et $e_2 = e_0 - x$), la relation précédente devient :

$$\frac{U_v(p)}{U_d(p)} = \frac{x}{e_0}$$

Le module de U_v est proportionnel au déplacement relatif x , on élimine ensuite la partie alternative par détection synchrone.

Il est nécessaire pour assurer la stabilité de la suspension de tenir compte des différentes fonctions de transfert de ces éléments.

4. Commande multi-variables dite de « translation – basculement »

En réalité, l'approche utilisée pour stabiliser la suspension n'est pas mono-variable. Il peut y avoir un intérêt particulier à coupler les axes entre eux afin de simplifier ou rendre plus efficace la commande au regard des lois fondamentales qui régissent un solide en mouvement. On considère deux mouvements combinés pour décrire les mouvements radiaux. Il s'agit du mode de translation et du mode de basculement autour du centre de gravité de l'arbre (figure 19) dans les plan V et W.

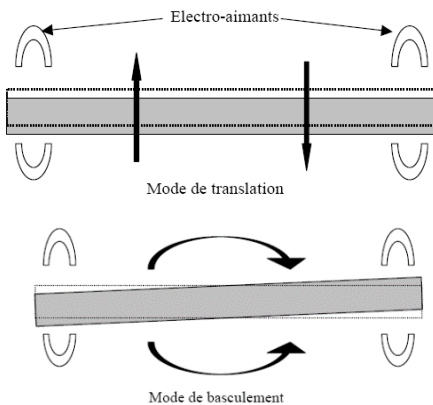


Figure 19 : Modes rigides

Cette décomposition est associée à une commande de type "translation-bascullement" qui fait intervenir des boucles de régulation en translation et en basculement

dans chaque plan V et W. Une cinquième boucle de régulation, non étudiée ici, est dédiée à la régulation axiale.

4.1. Principe de la commande

On suppose une translation verticale x du rotor dans le sens positif du plan V et un basculement d'angle θ autour de son centre de gravité G (figure 20).

On rappelle que le détecteur de position, présenté au paragraphe B.4 mesure un déplacement relatif dans l'entrefer. On note d_1 et d_2 les déplacements relatifs mesurés par les détecteurs par rapport à la position d'équilibre du rotor respectivement suivant les axes V13 et V24.

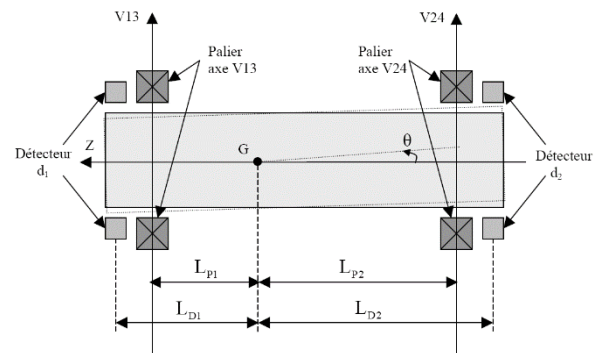


Figure 20 : Conventions utilisées dans le plan V

Dans le cas d'une translation x positive du rotor et d'un basculement d'angle θ positif autour du centre de gravité G, on obtient :

$$d_1 = x - \tan \theta \cdot L_{D1} \quad \text{et} \quad d_2 = x + \tan \theta \cdot L_{D2}$$

Le débattement maximal autorisé au niveau des détecteurs étant égal à 20% de l'entrefer, l'approximation $\tan \theta \approx \theta$ est justifiée. On a donc :

$$\begin{cases} d_1 = x - \theta L_{D1} \\ d_2 = x + \theta L_{D2} \end{cases}$$

Le principe de l'asservissement plan par plan utilisé est illustré figure 21.

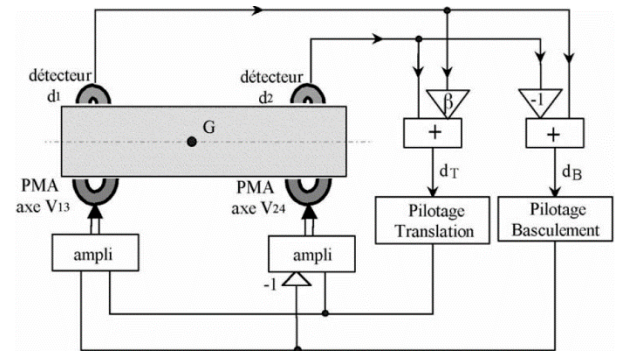


Figure 21 : Synoptique de l'asservissement en translation-bascullement dans le plan V

Les entrées d_T et d_B des pilotages en translation et basculement valent :

$$\begin{cases} d_T = \beta d_1 + d_2 = x(\beta + 1) + \theta(L_{D2} - \beta L_{D1}) \\ d_B = d_1 - d_2 = -\theta(L_{D1} + L_{D2}) \end{cases}$$

Si $\beta = \frac{L_{D2}}{L_{D1}}$ alors $\begin{cases} d_T = x(\beta + 1) \\ d_B = -\theta(L_{D1} + L_{D2}) \end{cases}$, d_T et d_B sont

alors proportionnelles à x et θ .

Lors d'une translation x dans le sens positif du plan V, d_T étant positif, les paliers V13 et V24 sont activés positivement. Le rotor "redescend". De même, lors d'un basculement d'angle θ positif, d_B est négatif. Le PMA suivant l'axe V13 est activé positivement, par contre le PMA suivant l'axe V24 est activé négativement. Le rotor pivote dans le sens inverse au basculement θ . Les actions ainsi engendrées sur les paliers des axes V13 et V24 permettent de corriger le déplacement ; la commande se fait donc bien plan par plan.

4.2. Modèle d'état

On peut modéliser sous forme d'état le système rigide en le décomposant sous forme modale.

Le rotor est modélisé par une tige de diamètre négligeable, de centre de gravité G et de masse totale M' . On note J le moment d'inertie axiale ou inertie transverse du rotor.

Les actions de la pesanteur sont supposées nulles, cette annulation étant due à l'utilisation d'un effet intégral dans la régulation.

On utilise, pour le modèle d'état les notations suivantes :

- x : variation verticale de la position du centre de gravité,
- θ : angle de basculement du rotor autour du centre de gravité,
- d_1, d_2 : variations de position au niveau des détecteurs 1 et 2,
- $\vec{F}_1, \vec{F}_2$: forces exercées au niveau des paliers 1 et 2 (orientées suivant V13 et V24).

En étudiant le mouvement de translation et le mouvement de basculement, on obtient :

$$M' \ddot{x} = F_1 + F_2$$

$$J \ddot{\theta} = -F_1 L_{P1} + F_2 L_{P2}$$

On note $[U]$ et $[Y]$ les vecteurs d'entrée et de sortie du système et $[X]$ le vecteur d'état, soit :

$$[U] = \begin{bmatrix} F_1 \\ F_2 \end{bmatrix}, [Y] = \begin{bmatrix} d_1 \\ d_2 \end{bmatrix} \text{ et } [X] = \begin{bmatrix} x \\ \dot{x} \\ \theta \\ \dot{\theta} \end{bmatrix}.$$

D'après les questions précédentes, on a :

$$\begin{cases} \ddot{x} = \frac{F_1}{M'} + \frac{F_2}{M'} \\ \ddot{\theta} = -F_1 \frac{L_{P1}}{J} + F_2 \frac{L_{P2}}{J} \end{cases} \text{ et } \begin{cases} d_1 = x - \theta L_{D1} \\ d_2 = x + \theta L_{D2} \end{cases}.$$

La représentation d'état du rotor rigide peut donc se mettre sous la forme classique :

$$[\dot{X}] = [A][X] + [B][U]$$

$$[Y] = [C][X] + [D][U]$$

$$\text{avec } [A] = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{bmatrix}, [B] = \begin{bmatrix} 0 & 0 \\ b_1 & b_1 \\ 0 & 0 \\ b_2 & b_3 \end{bmatrix},$$

$$[C] = \begin{bmatrix} 1 & 0 & c_1 & 0 \\ 1 & 0 & c_2 & 0 \end{bmatrix}, [D] = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \text{ et } b_1 = 1/M',$$

$$b_2 = -L_{P1}/J, b_3 = L_{P2}/J, c_1 = -L_{D1} \text{ et } c_2 = L_{D2}.$$

Pour piloter les sorties d_1 et d_2 , on peut envisager une commande par retour d'état. La loi de commande s'écrit : $[U] = [K][X]$.

5. Influences des modes flexibles et des effets gyroscopiques

Le système du système "paliers + rotor en suspension" comporte, comme nous l'avons vu dans les parties précédentes, des modes rigides. Il présente également des modes flexibles dits mode d'arbre, ainsi nommés parce qu'ils représentent la flexion de l'arbre.

La figure 22 donne la représentation du second mode flexible d'un rotor.

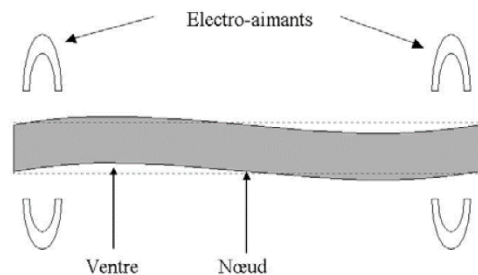


Figure 22 : Second mode flexible

5.1. Modélisation des modes flexibles

Pour cette étude, on reprend la démarche utilisée au paragraphe 2 pour les modes rigides, rappelée figure 23.

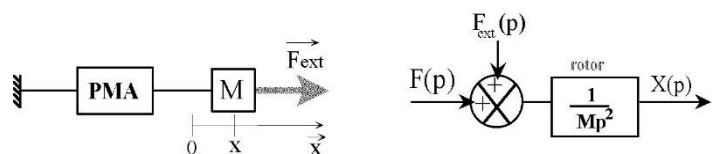


Figure 23 : Modélisation sur un axe des modes rigides

Lorsque la fréquence augmente, l'arbre peut être considéré comme une somme de masses reliées entre elles par une raideur.



L'étude du mode 1 se ramène au cas simple où la masse M est décomposée en deux masses γM et $(1-\gamma)M$ reliées par une raideur k' , comme le montre la figure 23.

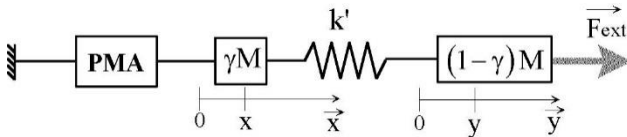


Figure 23 : Contraintes dans le mobile en suspension (mode 1 non amorti)

En appliquant les relations fondamentales de la dynamique à la masse $(1-\gamma)M$ suivant l'axe y et à la masse γM suivant l'axe x , on obtient :

$$(1-\gamma)M \frac{d^2 y}{dt^2} = F_{\text{ext}} - k'(y-x)$$

$$\gamma M \frac{d^2 x}{dt^2} = F + k'(y-x).$$

Le schéma fonctionnel de l'ensemble du système "palier + rotor en suspension" peut alors être représenté par la figure 24.

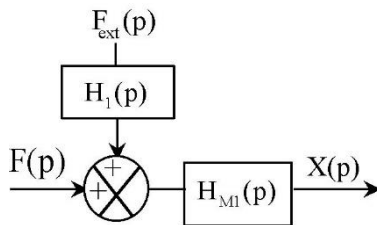


Figure 24 : Schéma bloc du rotor + palier (mode flexible 1)

Après calculs, on montre que $H_{M1}(p)$ peut se mettre sous la forme :

$$H_{M1}(p) = \frac{1}{Mp^2} + \frac{k_{M1}}{1 + \left(\frac{p}{\omega_{M1}}\right)^2}$$

avec $k_{M1} = \frac{(1-\gamma)^2}{k'}$ et $\omega_{M1} = \sqrt{\frac{k'}{\gamma(1-\gamma)M}}.$

Les modes flexibles sont modélisés à partir d'un logiciel à éléments finis (FEM MADYN). Ce logiciel calcule les déformations de l'arbre en différents nœuds. La figure 25 représente le rotor du turbocompresseur ainsi que les déformations associées pour les modes 1 et 2.

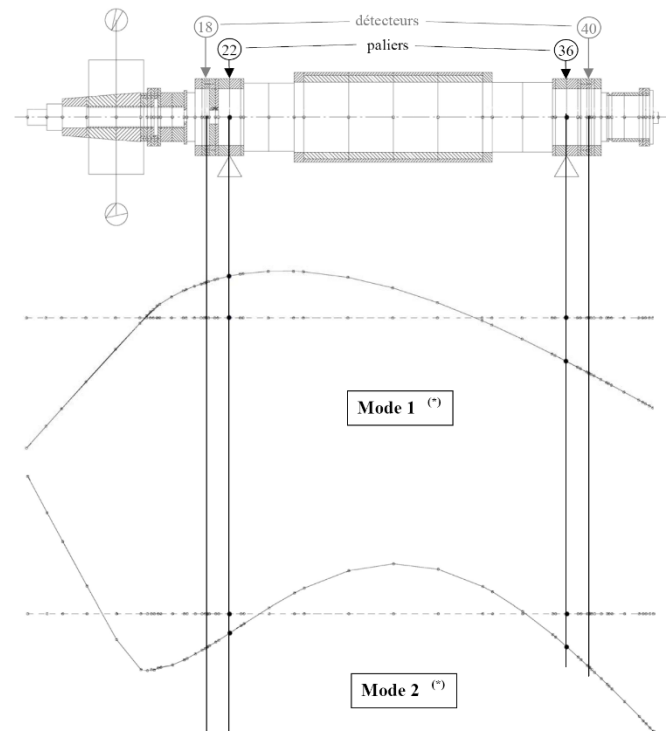
Après réduction nodale, on obtient les différentes fonctions de transfert du rotor suspendu. Ces modes flexibles induisent des résonances qui déstabilisent la suspension.

Pour rendre l'amortissement stable, il faudrait :

- soit amortir le système en considérant que la raideur entre les deux masses possède un coefficient d'amortissement ; si cet amortissement ne suffit pas, il faut créer un trou de gain au niveau des résonances. Ceci est possible en ajoutant un filtre double T au correcteur. Le problème de cette solution est qu'elle fait également diminuer la phase autour de la pulsation propre ω_0 .
- soit filtrer beaucoup plus vite après la première fréquence propre ou au contraire allonger le plus

possible la bande de phase pour englober les différentes pulsations propres.

Ces différentes méthodes sont difficiles à mettre en place et ne sont pas satisfaisantes. Les réglages sont peu précis car les valeurs des fréquences propres, lorsque le rotor tourne, sont modifiées par les effets gyroscopiques. Pour passer ces vitesses critiques on utilisera une stratégie fonction de la géométrie du rotor en mouvement : en effet, l'utilisation de la commande multi-variables de translation basculement permet de découpler les propriétés de commandabilité. Ainsi, le mode 1 sera principalement contrôlé en basculement, alors que le mode 2 le sera en translation.



(*) seule l'échelle des abscisses est conservée.

Figure 25 : Déformation de l'arbre pour les modes 1 et 2

5.2. Effets gyroscopiques

Dans le cas d'une machine tournante, les effets gyroscopiques se présentent comme un couplage entre les plans V et W du rotor, qui apparaît lors de la mise en rotation du système.

Ces effets gyroscopiques sont à l'origine d'une évolution des modes de la machine (modes rigides et flexibles) lorsque sa vitesse évolue. Cette évolution des modes est connue sous le nom de diagramme de Campbell (figure 26).

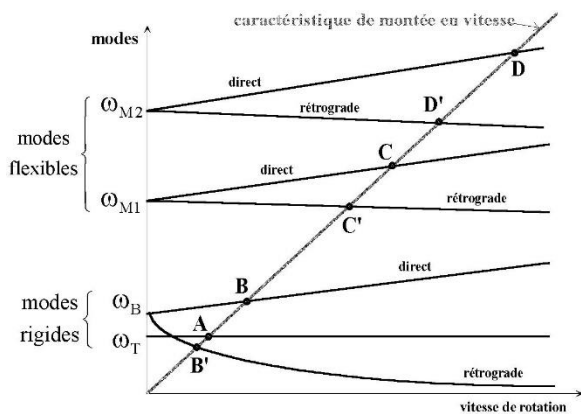


Figure 26 : Diagramme de Campbell

Si l'asservissement est de type translation-basculement, les deux premiers modes rigides correspondent au mode de translation ω_T et au mode de basculement ω_B . Le mode rigide de translation ω_T est généralement peu affecté par les effets gyroscopiques donc par la vitesse. Dans le cas particulier du turbocompresseur étudié, on a $\omega_T = (\omega_B)_{\Omega=0} = \omega_0$.

Par contre, le mode de basculement ω_B ainsi que les modes de flexion ω_{M1} et ω_{M2} se dédoublent lorsque la vitesse est non nulle. Lorsqu'un mode se dédouble, le mode résultant dans le sens de rotation est appelé mode de nutation ou mode direct, le second mode est appelé mode de précession ou rétrograde.

Durant une montée en vitesse, les modes de résonance correspondant aux différents modes propres du rotor sont croisés et peuvent être potentiellement excités (diagramme de Campbell). Les points A, B, C et D, qui correspondent à des modes directs, peuvent être excités par le balourd (car tournant dans le même sens de rotation). Les points B', C' et D' ne sont pas excités par le balourd, mais peuvent l'être par des perturbations rétrogrades (effet de gaz autour des roues ou contact rotor stator par exemple).

En général, pour les turbomachines, on ne se contente pas d'examiner les excitations liées au fondamental de la vitesse mais aussi à des multiples de cette valeur (nombre de pales des roues...).

5.3. Contrôle anti-vibratoire

Les machines tournantes subissent des perturbations harmoniques de la fréquence de rotation (fondamental, harmoniques supérieures ou inférieurs). Ces forces extérieures harmoniques dépendent de la roue du turbocompresseur, des défauts de symétrie spatiale du système (défauts d'usinage et de montage du rotor, des paliers et des capteurs)... Les différents modes de résonance introduits précédemment sont ainsi excités entraînant l'instabilité de la suspension ou des vibrations au niveau du bâti. Les industriels cherchent donc à éliminer ces vibrations synchrones.

La perturbation principale est due au balourd qui provient d'un décalage entre l'axe d'inertie G et l'axe des paliers S (figure 27) provoquant des forces perturbatrices importantes. Le but de l'asservissement étant de tourner autour de l'axe géométrique S des paliers, ceux-ci devront supporter les forces dynamiques du balourd.

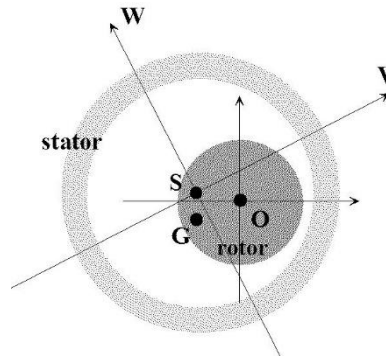


Figure 27 : Position des axes (rotor très déséquilibré)

Dans le cas de paliers magnétiques actifs, il est possible de déplacer l'axe de rotation à travers l'entrefer. Si on déplace l'axe de rotation sur l'axe d'inertie du rotor ($S \approx G$), on fait disparaître les forces provoquées par le balourd. Cette méthode permet de s'affranchir de l'influence du balourd au delà des vitesses critiques, il s'agit de l'ABS (Active Balancing System).

On montre que le schéma fonctionnel de la suspension obtenue figure 12 peut être remplacé pour cette étude par celui de la figure 28. En effet, il est possible de modéliser le balourd comme une perturbation additive sur la position d'amplitude ($OG = b$) et de phase ϕ constante. Ceci correspond au fait que le vecteur $\overline{OG}$ est fixe dans un repère lié au rotor.

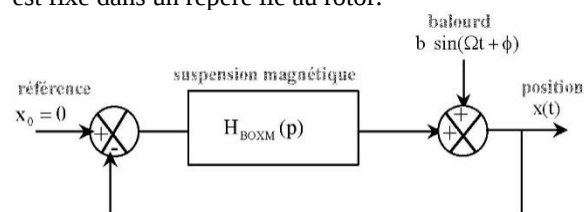


Figure 28 : Schéma fonctionnel (balourd perturbant la position)

Si on injecte au système décrit figure 28 un signal $s(t)$ telle que :

où b , $\hat{\Omega}$ et ϕ sont les estimés de b , Ω et ϕ , on s'affranchit du balourd. En effet, si le balourd est parfaitement connu, on a alors un minimum de variation de commande donc un minimum de vibrations.

Le signal $s(t)$ est obtenu par détection synchrone (figure 29). On suppose qu'une boucle à verrouillage de phase génère à partir d'une valeur estimée de la vitesse de rotation deux signaux $\sin(\hat{\Omega}t)$ et $\cos(\hat{\Omega}t)$.

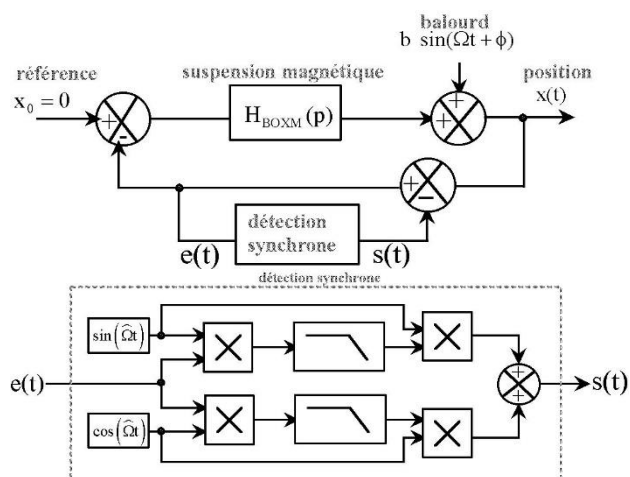


Figure 29 : Schéma de principe de la compensation de balourd par ABS

Une mesure de montée en vitesse est donnée figure 30. La courbe 1 représente le signal de position pour l'axe V13. Les courbes 2 et 3 représentent respectivement les signaux de commande pour les axes V24 et V13.

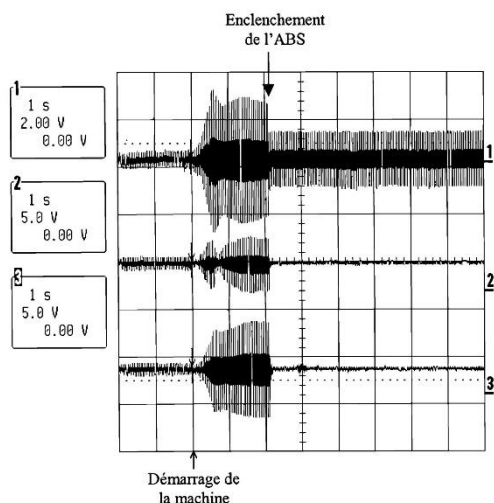


Figure 30 : Montée en vitesse avec déclenchement de l'ABS

Ces courbes représentent la réponse du système à un balourd quasi-sinusoidal dont la fréquence augmente progressivement avec la montée en vitesse.

On voit que l'amplitude des différentes courbes passe par un maximum lors du franchissement du mode rigide. Dès que le contrôle automatique du balourd (ABS) est enclenché, l'enveloppe de la position reste à une valeur constante alors que la commande synchrone est annulée. Le rotor tourne autour de son axe d'inertie naturel, en ne consommant presque plus d'énergie (courbes 2 et 3).

Conclusion

Ce sujet a permis de mettre en évidence les constituants et les propriétés d'une suspension magnétique active. L'approche a souvent été simplifiée, le lecteur pourra se référer à la bibliographie pour compléter cette étude.

La technologie des paliers magnétiques est de plus en plus accessible et devient concurrentielle dans des applications élargies. Elle commence à sortir de ses domaines de prédilections qui ont fait son succès : turbocompresseurs, pompes turbomoléculaires, électrobroches, volants d'inertie...

La société S2M prévoit le développement d'une machine dédiée à l'enseignement. Ce système permettrait une sensibilisation à cette technologie qui englobe tous les domaines des sciences industrielles : la mécanique, l'électromagnétisme, l'électronique de puissance, l'électronique numérique matérielle et logicielle et la commande.

Bibliographie utilisée pour la réalisation de ce sujet :

H. HABERMANN, Paliers magnétiques, Techniques de l'ingénieurs, B 5345, Déc. 1984.

V. TAMISIER, Machines tournantes sur paliers magnétiques actifs : modélisation et contrôle anti-vibratoire numérique, Thèse de doctorat Paris XI, Février 2003

E. MASLEN, Magnetic Bearings, Cours de l'Université de Virginie, Juin 2000.

S2M, Rapports internes.

J. DELAMARE, Suspensions magnétiques partiellement actives, Thèse de Doctorat INPG, Janvier 1994.

www.s2m.fr

Annales des épreuves de l'agrégation :

www.iufmrese.cict.fr/concours/2005/AgExt/AgExt2005.shtml

www.education.gouv.fr/siac/siac2/jury/2005/detail/agreg_ext_get.htm

Annexe

Turbocompresseur sur paliers magnétiques



TURBO COMPRESSEUR 110KW

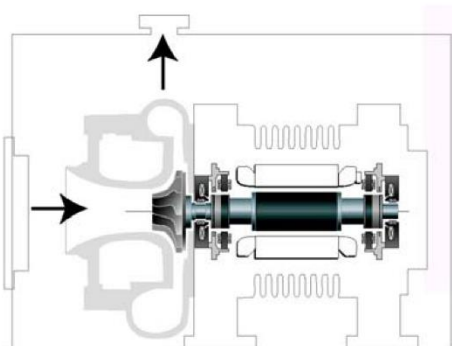


Fig. 1 Le compresseur centrifuge



Fig. 2 La roue

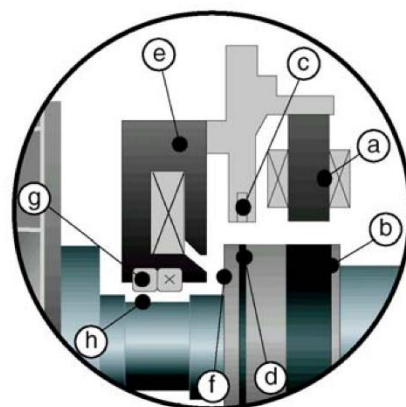


Fig. 3 Détail des paliers

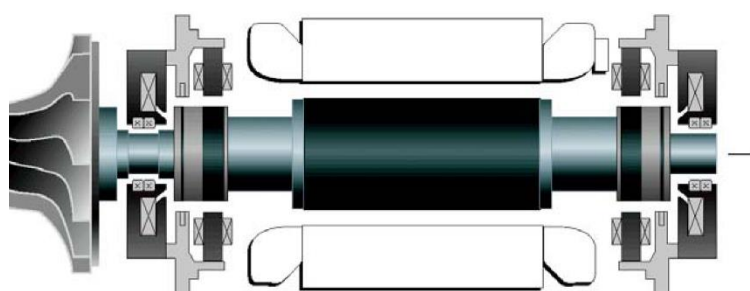
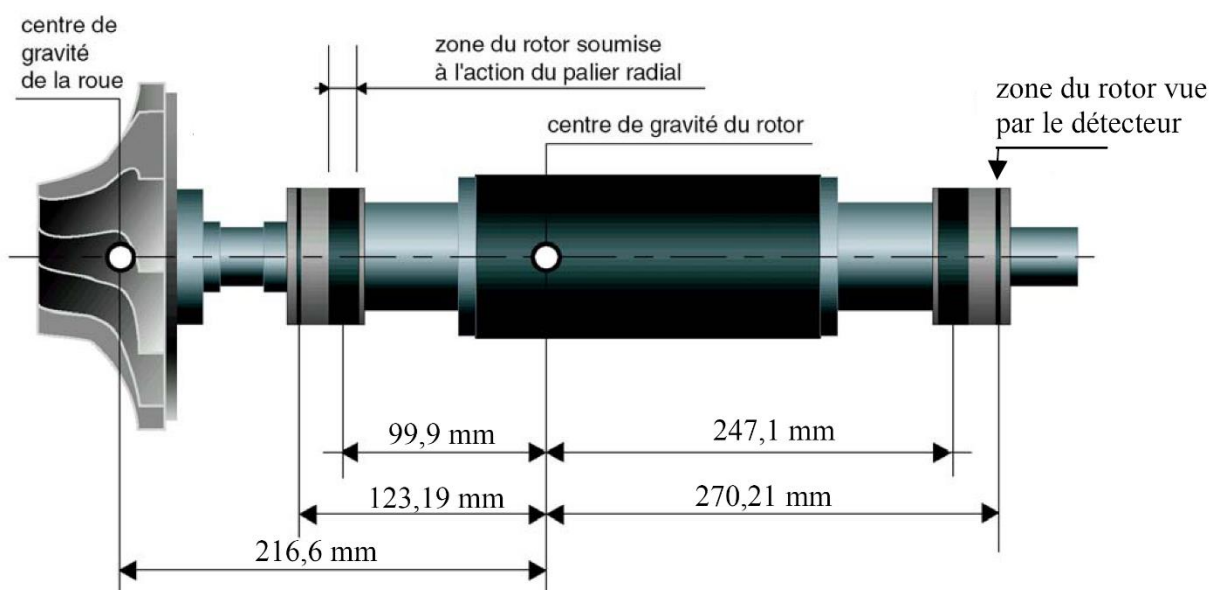


Fig. 3 LEGENDES

- a) Palier radial : électro-aimant (stator)
- b) Palier radial : tôles fretées sur le rotor
- c) Détecteur combiné radial/axial (stator)
- d) Zone du rotor vue par le détecteur
- e) Palier axial : électro-aimant (stator)
- f) Zone du rotor soumise à l'action du palier axial
- g) Roulements auxiliaires
- h) Zone d'appui des roulements aux. en cas de besoin

Fig. 3 Rotor + stator (paliers & moteur)





Systèmes logiques séquentiels

Optimisation théorique et circuits câblés ou programmation directe (vhdl) ?

Jérôme FAUCHER, Jérémie REGNIER, Marcel GRANDPIERRE

Ecole Nationale Supérieure d'Electrotechnique, d'Informatique, d'Hydraulique et de télécommunications

Département de formation Génie Electrique-Automatique

2 Rue Camichel 31071 Toulouse Cedex 7

Remerciements : Olivier Durrieu et René Dirat

Résumé : L'objectif du présent article est de présenter une séquence pédagogique consacrée à la modélisation, l'analyse et la mise en œuvre de systèmes logiques séquentiels.

Prenant appui sur l'aspect ludique d'une séquence d'ouverture d'un coffre fort, les étudiants doivent mener tout d'abord une approche théorique traditionnelle (Huffman). L'exemple proposé permet de mettre en évidence les concepts de formalisation (diagrammes d'état), les notions de minimisation (fusionnement) et les problèmes de codage (cours critiques). Le tout conduit à une mise en œuvre "câblée" au moyen de portes logiques élémentaires (FPGA).

Une approche plus globale est ensuite proposée par implantation directe de la machine à états par programmation VHDL et implantation toujours sur FPGA.

L'accent est ensuite mis sur la comparaison des deux approches. Dans un premier temps, l'implantation VHDL semble largement privilégiée. Mais les étudiants sont amenés à réfléchir en intégrant à la fois les temps de développement, les coûts de mise en œuvre effective mais aussi l'environnement matériel et logiciel nécessaires. Les conclusions sont alors souvent moins évidentes !

1. Introduction

Cet article présente une séquence pédagogique de travaux pratiques dédiée à l'acquisition et à l'approfondissement des concepts de la logique séquentielle. Partant des méthodes théoriques traditionnelles de formalisation et de synthèse, elle se propose de conduire les étudiants à une réflexion objective sur différentes méthodes d'implantation.

Les cours de logique séquentielle s'appuient encore aujourd'hui sur les notions incontournables de représentations d'état, mais les points de vue divergent lorsqu'il est question de mise en œuvre. Les méthodes de minimisation et de codage optimaux (synthèse d'Huffman par exemple) restent bien évidemment d'actualité dans la mesure où elles conduisent à l'obtention de solutions réduisant le nombre d'opérateurs logiques. Cependant, avec l'avènement des composants numériques programmables de type FPGA, l'utilisation de méthodes de réalisation programmées (VHDL,...) ont pris au cours de cette dernière décennie un essor important. L'approche d'implantation est beaucoup plus directe, et ne nécessite plus forcément les efforts d'optimisation et de synthèse des méthodes plus traditionnelles.

L'objectif de la manipulation consiste donc, à travers une application prenant le prétexte de l'ouverture d'un coffre fort, d'implanter, après en avoir fait la synthèse, différentes solutions afin de les analyser et de les comparer en terme de temps de développement et de comportement fonctionnel.

2. Présentation de la manipulation

À partir d'un cahier des charges simple et "ludique" (séquence d'ouverture d'un coffre fort), le premier travail consiste à déterminer le graphe d'état de la machine correspondante. Les huit états de cette machine peuvent être réduits à quatre par fusionnement (Moore ou Mealy). Des précautions doivent ensuite être prises pour le codage de la table obtenue pour éviter l'apparition de courses critiques. Pour des raisons pratiques, l'implantation est réalisée en utilisant un FPGA, mais en mettant en œuvre uniquement des opérateurs logiques élémentaires (portes OU, ET,...). Le FPGA utilisé est un EP1K100QC208-2 de chez Altera et l'environnement de programmation utilisé est Quartus II.

Les différents cas de courses critiques mis en évidence de façon théorique peuvent ensuite être

visualisés tant à partir d'affichages sur une platine spécifique qu'en utilisant un analyseur logique sur PC.

En repartant du graphe d'état initial, la seconde étape consiste à programmer une nouvelle solution directement en VHDL et à l'implanter sur une autre zone du même FPGA.

Une comparaison simpliste des deux méthodes conduit invariablement les étudiants à privilégier l'approche VHDL. Mais une analyse plus approfondie amène à intégrer dans le jugement certes des notions de coût de développement, mais aussi la surface de silicium nécessaire, et de ramener le tout au contexte dans lequel se situe le problème : problème isolé pour lequel le coût d'étude s'avère prohibitif pour les solutions optimales théoriques où circuits à grande diffusion qui peuvent tout à fait justifier la mise en œuvre de minimisations.

2. Présentation de la maquette

La présente partie est dédiée à la description de la plateforme pédagogique permettant d'illustrer la problématique considérée, c'est-à-dire l'étude et la réalisation d'un système séquentiel : une serrure électrique à combinaison. La plate forme de travail est composée d'un coffre fort, d'une carte de commande comportant un FPGA, d'une interface utilisateur et d'un micro-ordinateur de développement.

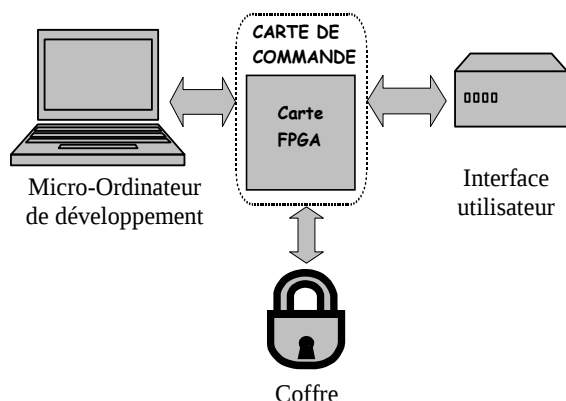


Figure 1 : Architecture de la plate forme de travail

L'objectif est l'ouverture d'un coffre fort commandé par un moteur électrique. Deux configurations physiques sont alors possibles : l'ouverture et la fermeture du coffre. Une seule sortie combinatoire pour les décrire est ainsi nécessaire.



Figure 2 : Coffre fort

Le coffre fort est alimenté au moyen de relais situés sur la carte de commande et commandés par le FPGA. Rappelons que deux moyens de programmation sont proposés : la réalisation du schéma logique d'une part et la programmation par langage VHDL d'autre part. Le cycle de conception de ce type d'application est présenté figure 3. Le concepteur agit uniquement sur la première phase (conception) et le reste du cycle est automatisé. A l'issue de celui-ci, il est possible d'observer les détails du taux d'occupation du FPGA permettant ainsi de comparer les tailles en terme d'implantation des différentes structures de programmation utilisées.

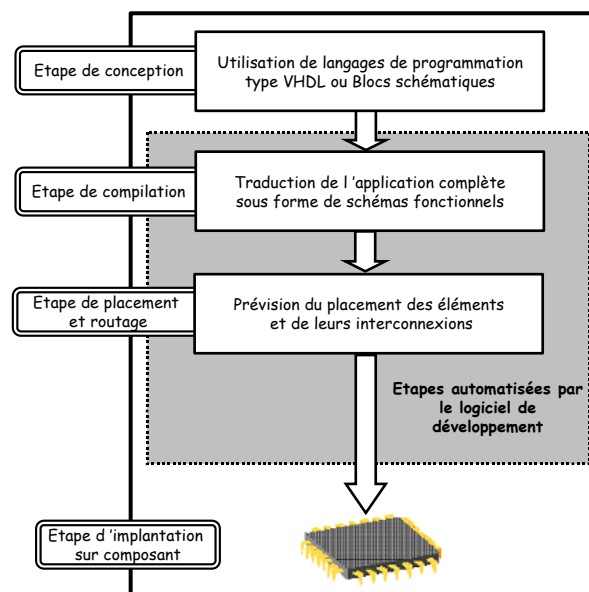


Figure 3 : Cycle de conception d'une application

La carte de commande permet également de remonter les informations binaires des différentes grandeurs utilisées vers le micro-ordinateur pour permettre leur observation au moyen d'un analyseur logique.

Enfin, l'interface utilisateur, située sur la même platine que la carte de commande (figure 4), permet à l'étudiant d'agir sur le système et de visualiser l'état de certaines variables.



Figure 4 : Platine de commande et d'interface

La combinaison d'ouverture du coffre passe par l'action sur deux boutons bistables. Ceux-ci pourront être actionnés à travers la platine et leur état est visualisé au moyen de leds. De plus, des interrupteurs permettront de sélectionner le programme actif, c'est-à-dire le câblage issu de la programmation VHDL ou du schéma logique.

3. Cahier des charges et aspects théoriques

L'objectif est donc d'aboutir à l'ouverture du coffre au moyen d'une combinaison. Cette combinaison d'ouverture est temporelle et passe par l'action sur deux boutons bistables notés A et B. La variable binaire de sortie S est associée à l'ouverture du coffre lorsqu'elle est à l'état haut. Le cahier des charges, définissant la combinaison, est construit de manière à mettre en évidence le phénomène de course critique et de pouvoir mettre en œuvre la synthèse d'Huffman. De plus, il est réalisé en vu de limiter le nombre d'états afin de conserver un système relativement simple et exploitable par les étudiants au niveau pédagogique.

Le graphe de fluence, établi à partir de ce cahier des charges, est représenté sur la figure 5 et décrit le fonctionnement désiré du système :

3. Programmation par schéma câblé

La présente partie illustre la programmation par logique câblée à l'issue de l'étude par la méthode d'Huffman.

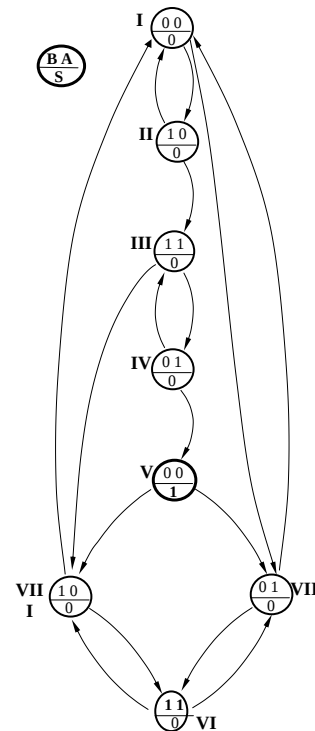


Figure 5 : Graphe de fluence

Cette méthode, qui permet d'aboutir à une solution de programmation d'une machine séquentielle par logique câblée, met en œuvre des processus de minimisation et de codages optimaux. Son principal intérêt réside dans la solution optimale qu'elle fournit en terme de simplicité des fonctions utilisées. Cependant, la durée de développement de l'étude est souvent longue et augmente de façon exponentielle avec le nombre d'état à traiter.

La méthode d'Huffman ne sera pas complètement détaillée ici. Seuls les points permettant d'illustrer l'intérêt pédagogique de la démarche seront présentés.

A partir du graphe de fluence présenté figure 5, la méthode d'Huffman, au moyen de la table des transitions primitives et de la fusion d'état (non présentés ici), conduit au graphe de fusionnement de la figure 6.

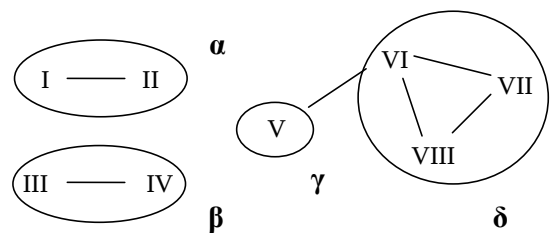


Figure 6 : Graphe de fusionnement



Il est à noter que même si le choix des états fusionnés conduit à une machine de Moore, les étudiants ont toute latitude pour choisir de réaliser une machine de Mealy (même si par la suite l'expression de la sortie devient plus compliquée).

Le choix effectué sur la figure 6 conduit au graphe d'adjacence des états présenté figure 7.

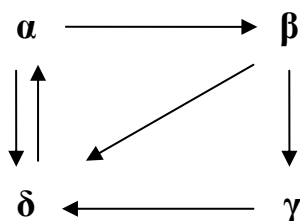


Figure 7 : Graphe d'adjacence des états

Il apparaît donc ici un problème d'adjacence entre β et δ qui peut être source de course critique (le cahier des charges a été énoncé de façon à ce qu'il y ait toujours une course critique quelque soit le choix des états fusionnés).

Le codage de ces 4 états fusionnés nécessite l'utilisation de deux variables d'excitations secondaires, X et Y , et conduit à la configuration de la machine séquentielle représentée figure 8.

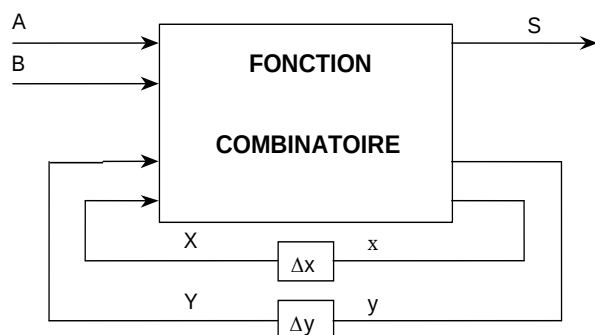


Figure 8 : Configuration de la machine séquentielle

Dans un premier temps, les étudiants doivent réaliser la programmation par schéma logique du système sans résoudre la course critique. Le logiciel *Quartus II* permet de réaliser directement le schéma à partir d'une bibliothèque de composant, comme illustré figure 9.

Après compilation et transfert du programme sur le FPGA, ce phénomène de course critique est mis en évidence au moyen de l'analyseur logique permettant de tracer les chronogrammes des entrées et des excitations secondaires. En effet, la carte d'interface de la platine permet de modifier les délais relatifs de transition des variables d'excitations secondaires.

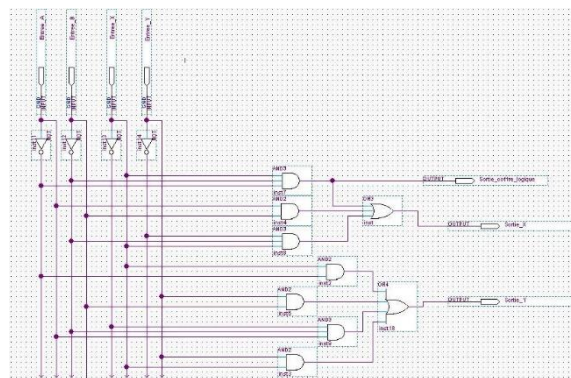


Figure 9 : Câblage sous Quartus II

La figure 10 montre les deux cas de figures possibles, c'est-à-dire, lorsque l'évolution de x est plus rapide, on observe le phénomène de course critique, alors que lorsque l'évolution de y est plus rapide, le cahier des charges est respecté.

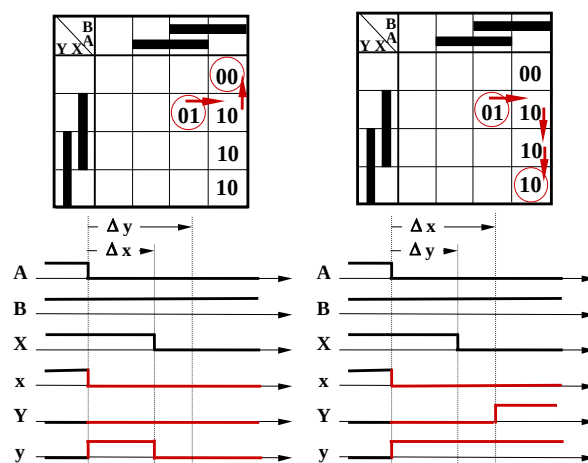


Figure 10 : Mise en évidence de la course critique

Dans un second temps, il est demandé aux étudiants de résoudre par la programmation le phénomène de course critique et de valider expérimentalement leur solution.

Ce premier exercice leur aura donc permis de mettre en œuvre la méthode d'Huffman sur un processus réel et d'être confrontés aux difficultés liées à l'utilisation de celle-ci. Il est évident que l'utilisation d'un FPGA avec tous les éléments connexes devant lui être associés pour l'application considérée est dans ce cas peu judicieuse. Cependant, cette solution présente l'avantage d'utiliser une plateforme de programmation unique pour réaliser la comparaison avec l'approche par programmation en langage VHDL.



4. Programmation en VHDL

La programmation en VHDL est réalisée directement à partir du graphe de fluence de la figure 5. Le choix de ce langage de programmation est motivé par l'avantage qu'il a d'être indépendant du type de FPGA utilisé, ce qui n'est pas le cas de certains langages de programmation constructeurs, comme le AHDL chez Altera. De plus, le VHDL s'impose de plus en plus comme le langage de référence pour la conception d'applications câblées. Les étudiants ayant peu de connaissances concernant ce langage, il leur est fourni une ossature de programme comportant trois processus différents qu'ils doivent compléter. Le processus 1 consiste à décrire la structure du graphe de fluence en indiquant le nombre d'état et les conditions de transition entre états. Le processus 2 consiste à décrire les valeurs à affecter à la sortie pour les différents états. Enfin, le troisième processus consiste à mettre à jour le graphe d'état en fonction des évolutions liées au processus 1.

La programmation en langage VHDL est assez intuitive et s'appréhende très vite, de telle sorte que les étudiants arrivent vite à la conclusion que cette approche est largement préférable à la méthode précédente. Les analyses approfondies proposées dans la suite de cette étude permettent de modérer cette première impression en sensibilisant les étudiants sur les différents critères à intégrer pour choisir un type d'implantation particulier.

5. Comparaison fonctionnelle

L'idée de cette partie est d'illustrer le fonctionnement du système dans le cas de la réduction d'état par programmation en langage VHDL. Pour ce faire, à partir du choix des états fusionnés, le graphe d'états fusionnés, figure 11, va être réalisé puis programmé en VHDL.

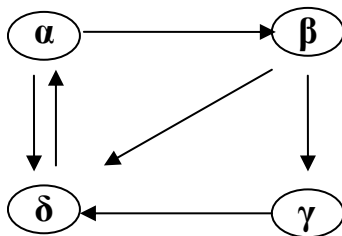


Figure 11 : Graphe d'états fusionnés

Il est évident que pour cette partie de l'exercice, l'utilisation d'une machine de Moore facilite considérablement la programmation, le processus 2 décrivant directement les sorties à partir des états fusionnés.

Cette étude supplémentaire présente deux intérêts. En premier lieu, utilisant le même graphe réduit que pour la programmation par logique câblée, il est intéressant de souligner la disparition du phénomène de course critique, notion qui est parfois difficile à faire passer, la différence de la nature des deux approches posant quelquefois des difficultés d'appréhension.

De plus, la première phase d'optimisation de la méthode d'Huffman est réutilisée et va permettre d'optimiser le nombre d'états et donc le codage en VHDL. Cette considération conduit alors à la partie suivante visant à comparer les deux approches en terme de coût global d'implantation, l'information sur le taux d'occupation du FPGA étant communiquée par le logiciel de compilation pour chacune des méthodes de programmation.

6. Comparaison du coût d'implantation

L'objectif est ici de comparer les deux approches de programmation préalablement utilisées en terme de coût d'implantation. Le logiciel Quartus II permet de d'obtenir le taux d'occupation du FPGA pour chacune des solutions testées. Il apparaît alors que la solution VHDL nécessite l'utilisation d'environ 7 fois plus de cellules logiques que la solution obtenue avec la synthèse d'Huffman (voir figure 7).

Hierarchy							
Entity	Logic Cells	LC Registers	Memory Bits	Pins	LUT-Only LCs	Register-Only LCs	LUT/Register LCs
ACE1K1K: EP1K100Q208-2							
TP_COFFRE	1331 (0)	453	0	24	878 (0)	419 (0)	34 (0)
Synthese_logique_cour...	3 (3)	0	0	0	3 (3)	0 (0)	0 (0)
gestion_variabes:inst4	458 (0)	136	0	0	322 (0)	128 (0)	8 (0)
Synthese_VHDL:inst5	20 (20)	8	0	0	12 (12)	0 (0)	8 (8)
Gestion_coffre:inst16	850 (0)	309	0	0	541 (0)	291 (0)	18 (0)

Figure 12 : Taux d'occupation du FPGA pour les solutions VHDL et synthèse d'Huffman

Dans ces conditions, si le taux d'occupation du FPGA est une contrainte importante dans la conception du système, la solution VHDL, certes plus simple à concevoir, n'est pas idéale.

Parmi les autres critères pouvant influencer sur le choix d'une des deux solutions, figurent également le temps de développement et le coût de production. Les phases de développement réalisées précédemment ont montré que la méthode d'Huffman nécessite un coût d'étude plus important que l'approche par VHDL mais conduit à une solution minimale en terme de silicium utilisé. De plus, même si pour notre manipulation, un FPGA a été utilisé pour des raisons pratiques pour implanter la solution par synthèse d'Huffman, une implantation à l'aide de composants discrets est envisageable, ce qui réduit



d'autant plus le coût de l'application. La synthèse en VHDL conduit inévitablement à disposer d'un composant FPGA ainsi que de son environnement de développement, ce qui représente un investissement financier plus important que de simples portes logiques qui restent très bon marché.

L'intérêt est alors d'amener les étudiants à mettre en perspectives ces considérations et ne pas considérer une approche optimale en absolu mais de rapporter le processus de conception / réalisation à la problématique considérée : amortissement du coût de l'étude d'une part et du coût en silicium d'autre part par rapport au nombre d'exemplaires réalisés.

7. Conclusion

Au travers d'une plate forme expérimentale, les étudiants ont pu mettre en œuvre deux méthodes de programmation conduisant à la réalisation d'une machine séquentielle. Les intérêts pédagogiques de cette manipulation sont multiples.

Tout d'abord, l'étudiant peut mettre en œuvre deux méthodes de synthèse différentes pour répondre au cahier des charges, montrant ainsi, d'un point de vue méthodologique, que plusieurs approches peuvent être menées pour traiter un seul et même problème.

Ensuite, les étudiants sont mis en contact avec la technologie FPGA. Ils sont ainsi familiarisés avec les composants numériques programmables, qui sont devenus de nos jours des éléments incontournables dans le domaine de l'informatique industrielle. L'introduction de ces technologies dès le début de leur formation nous semble indispensable.

Enfin, au-delà des aspects méthodologiques et technologiques, cette manipulation permet de sensibiliser les étudiants, futurs ingénieurs, aux contraintes et critères industriels (coût de développement, coût d'implantation, coût matériel, ...) qui doivent être pris en considération pour étudier et proposer une solution à un cahier des charges. Ces aspects, en général peu abordés dans les enseignements, doivent faire parties des réflexions à intégrer au cours du développement d'une application.

De la clepsydre à Michel Platini : quelques notions fondamentales en mécanique des fluides

François LEMAIRE

Professeur de physique appliquée
Lycée Colbert, Lorient

Contact : francois.lemaire1@ac-rennes.fr

Résumé : *La mécanique des fluides est l'un des sujets qui apparaît dans les nouveaux programmes de sciences appliquées du BTS électrotechnique.*

La mécanique des fluides peut paraître parfois un peu indigeste, il suffit d'ouvrir le « mécanique des fluides » de Landau et Lifschitz aux éditions Mir pour s'en convaincre. Il m'a semblé intéressant de rechercher une autre présentation qui puisse être un peu plus attractive pour nos élèves, tout en leur apportant quelques repères historiques qui sont pour moi des éléments fondamentaux d'une culture scientifique et technique.

1. Introduction

Cette présentation est organisée en quatre parties :

- 1- Une question sur des phénomènes simples destinée à éveiller la curiosité des auditeurs.
- 2- Une approche historique de l'explication des phénomènes, permettant d'appréhender la démarche scientifique et de montrer l'intérêt d'une étude théorique.
- 3- Une formulation des lois physiques qui sous-tendent l'explication, et l'utilisation de ces lois dans des exemples simples.
- 4- Une ou plusieurs applications technologiques usuelles dont la connaissance doit faire partie du bagage scientifique et technique de nos élèves.

Il y a quatre items dans le programme, il m'a semblé plus logique de faire suivre le théorème de Bernouilli, du débit, et les pertes de charge, de la viscosité.

Il ne s'agit pas d'un cours complet et terminé, je laisse aux lecteurs le soin d'adapter les données que j'ai collectées. J'ai juste voulu apporter quelques informations historiques sur le sujet. Pour consulter un cours finalisé, je vous conseille d'aller voir les sites internet mentionnés en références bibliographiques [1 à 3].

2. Le débit

Question : Est-il possible de mesurer le temps à l'aide de l'écoulement d'un liquide ?

Approche historique : Un récipient rempli d'eau laisse échapper le liquide par un petit trou à sa base. Cet objet est une horloge à eau dénommée aussi **clepsydre** du grec klepsudra qui signifie "dérobe l'eau" ou "voleuse d'eau".

Les clepsydres seront utilisées très longtemps, elles ne deviendront obsolètes qu'avec le développement des horloges et des montres de précision au XVIII^e siècle. Elles seront employées dans deux situations : soit pour déterminer une durée précise, comme le temps de parole pour les procès dans la Grèce antique, soit pour mesurer l'écoulement du temps quand les cadrans solaires ne pouvaient être utilisés.

La plus ancienne réalisation connue date du pharaon Amenophis III (1408-1372 av J.C.), elle serait l'oeuvre ou appartenait à un dénommé Amenemhat. Elle a été retrouvée à Karnak en Egypte (c f. figure 1).



Figure 1 : Dessin de la clepsydre d'Amenemhat

Cette horloge à eau est un vase tronconique en albâtre. Sa hauteur est de 95 cm et son diamètre de 48,5 cm. La surface extérieure de cette clepsydre est décorée de figurines et de textes, non représentés sur le dessin. Ceux-ci représentent certaines planètes et constellations et dressent une liste des esprits protecteurs pour chacun des dix jours de la semaine de l'Égypte ancienne. Sur la surface intérieure sont gravées douze colonnes de onze petits points servant de repères. L'eau s'échappait du vase par un petit trou percé dans le bas du vase. Pour connaître l'heure, il suffisait de regarder le niveau d'eau dans le récipient et de repérer le point le plus proche.

Dès le cinquième siècle avant-Jésus-Christ les tribunaux grecs utiliseront des clepsydes pour limiter le temps de parole des orateurs, celles retrouvées contiennent 6,5 litres d'eau ce qui correspond à une durée d'écoulement de l'eau d'environ 6 minutes. Les orateurs disposaient de plusieurs remplissages en fonction de l'importance du procès. Les tribunaux romains utiliseront le même système un peu plus tard.

On peut citer comme autre exemple la clepsydre offerte à Charlemagne par le calife de Bagdad Haroun al rachid en 807, considérée comme un présent de très grande valeur.

Plus près de nous Galilée en 1610, utilisa une clepsydre pour les mesures de temps lors de ses études expérimentales sur la chute des corps : *"quant à la mesure du temps, nous la fîmes à l'aide d'un grand seau plein d'eau d'où sortait, par un fin tuyau soudé sur le fond, un mince filet d'eau reçu dans un petit verre durant tout le temps de la descente de la boule. Les quantités d'eau recueillies étaient pesées chaque fois sur une balance très exacte donnant par la différence et proportion de leurs poids la différence et proportion des temps"* [4].

Sa précision est limitée par le phénomène suivant : Le débit dépend de la vitesse d'écoulement du liquide, lui-même dépendant de la pression exercée au niveau de l'orifice, liée elle à la hauteur de l'eau dans le récipient, du diamètre de l'orifice, et de la viscosité du liquide. Diverses améliorations seront apportées au cours des siècles pour obtenir une horloge mesurant des durées égales.

Il faudra attendre 1640 pour que, le père Marin Mersenne (1588-1648), René Descartes (1596-1650) et Evangelista Torricelli (1608-1647) commencent à formaliser les lois de ce que Daniel Bernoulli (1700-1782) nommera en 1738 hydrodynamique. C'est le nom de son ouvrage "hydrodynamica" publié à Strasbourg en 1738.

Marin Mersenne montrera le premier que la vitesse du jet d'eau est proportionnelle à la racine carrée de la hauteur d'eau dans le récipient. Il faudra plusieurs années d'expérimentations pour parvenir à ce résultat; La loi dite de Torricelli sera en fait établie sous sa forme actuelle: $v = \sqrt{2gh}$ par Bernoulli.

Tous ces scientifiques constatent aussi que l'eau ne peut remonter jusqu'à sa hauteur de départ, ce qui s'explique par une perte d'énergie et conduira à la notion de viscosité.

Les "fontainiers" qui s'occupaient de l'eau des villes savaient que les siphons utilisés pour franchir des vallées ne permettaient pas de remonter à une hauteur supérieure à 18 brasses soit approximativement 10m. Si

le phénomène est parfaitement compris aujourd'hui (la pression exercée par la colonne d'eau ne peut dépasser la pression atmosphérique $P_{atm} = \rho gh$), le problème est toujours là. Une pompe aspirante ne peut aspirer l'eau sur une telle hauteur. On doit alors utiliser des pompes refoulantes, éventuellement immergées.

Beaucoup d'expérimentations vont avoir lieu sur ce sujet mais une colonne d'eau de 10 m dans un tube n'est pas très aisée à manipuler¹. Torricelli a alors l'idée d'utiliser le mercure de densité 13,6 donc tel que la hauteur d'aspiration maximale soit 10,3/13,6 soit 0,76m. Le principe du baromètre à mercure venait d'être trouvé.

Le site de Versailles, mal alimenté en eau et le goût immodéré de Louis XIV pour les spectacles aquatiques vont entraîner un développement considérable de la recherche en hydraulique. L'académie des sciences sera d'ailleurs chargée officiellement par Louvois de ces recherches. Il en résultera la publication de nombreux ouvrages tel le "Traité du mouvement des eaux et des autres corps fluides" publié à Paris en 1686 par Edme Mariotte (1620-1684).

Lois physiques (définitions de base) :

Fluide : Dénomination englobant l'état fluide et l'état gazeux. L'état fluide est caractérisé par l'absence de forme propre et la possibilité de diffusion mutuelle de plusieurs fluides les uns à travers les autres. La principale différence entre liquide et gaz est la compressibilité. Les liquides sont incompressibles et ont un volume propre, les gaz sont compressibles et n'ont pas de volume propre. Un fluide est dit parfait si l'on peut négliger sa viscosité, grandeur qui sera abordée plus loin.

Débit massique : dm étant la masse de fluide qui s'écoule à travers une section S durant l'intervalle de temps dt , le débit massique, appelé aussi débit masse s'écrit :

$$q_m = \frac{dm}{dt} \quad (1)$$

Débit volumique : dV étant le volume de fluide qui s'écoule à travers une section S durant l'intervalle de temps dt le débit volumique, appelé aussi débit volume s'écrit :

$$q_v = \frac{dV}{dt} \quad (2)$$

Masse volumique : c'est le rapport de la masse au volume : $\rho = \frac{dm}{dV}$, cette grandeur est constante pour les liquides qui sont incompressibles et peu dilatables. Elle dépend de la température et de la pression pour les gaz.

¹ A ce sujet, il existe un curieux baromètre à eau dans un escalier du palais de la découverte à Paris.

Pression : Une force $\vec{F}$ de grandeur F , appliquée normalement et uniformément sur une surface S exerce sur celle-ci une pression $P = \frac{F}{S}$. L'unité S.I. de pression est le Pascal (symbole Pa), mais c'est certainement la grandeur physique qui s'exprime dans le plus d'unités selon le domaine d'emploi (cf. tableau ci-dessous).

domaine	unités
Météorologie	1mbar = 100 Pas = 1 hectopascal (hPas)
Hautes pressions : (chimie, plongée sous-marine,...)	1 bar = 105 Pa 1 atm = 1.033105 Pas
Très basses pressions : (pompes à vide,...)	Le millimètre de mercure, autrefois appelé Torr : 1 mm Hg = 1 torr = 133,32 Pa à 0°C
Le gonflage des pneus :	le PSI c'est à dire Pound-force by Square Inch d'origine américaine 1 PSI = 6895 10 ⁻² bar = 6895 Pa

Tableau 1

Applications technologiques :

Le débit est l'un des éléments de choix d'une pompe, c'est en général le premier paramètre à prendre en compte.

En général on trouve plutôt des pompes à faibles débits, quelques m³ /h, mais ayant de grandes hauteurs d'aspiration ou de refoulement, telles les pompes centrifuges ou les pompes volumétriques. Il existe aussi des pompes à hélices ou rotatives axiales qui permettent un très grand débit jusqu'à plusieurs milliers de m³ /h mais avec une très faible hauteur d'aspiration ou de refoulement.

Rappelons encore une fois qu'une pompe à eau ne peut aspirer au-delà de la hauteur théorique de 10,3 m, ce qui limite pratiquement l'aspiration à 7 ou 8 m.

3. Le théorème de Bernoulli

Question :

Le 18 novembre 1981 au parc des princes à Paris, l'équipe de France de football rencontrait l'équipe des Pays-Bas pour un match de qualification à la coupe du monde. A la cinquante-deuxième minute du match, Michel Platini ouvre le score sur un coup franc direct, situé à 18 m à gauche, à hauteur du premier poteau. Un tir croisé, brossé à mi-hauteur, contournant le mur adverse. Comment le ballon a-t-il pu avoir une trajectoire courbe lui permettant de passer derrière le mur des défenseurs ?

L'explication est l'effet Magnus que l'on peut démontrer en utilisant le théorème de Bernoulli.

Approche historique :

Daniel Bernoulli publie en 1738 à Strasbourg l'ouvrage " Hydrodynamica, siva de Viribus et Motibus Fluidorum

commentarii ", rédigée bien sûr en latin qui est encore à cette époque la langue universelle que personne ne parle, mais que tous les savants de tous les pays savent lire. Il y énonce le principe de conservation de l'énergie et y expose un traité sur le phénomène des marées. L'académie royale des sciences le cite comme "professeur de médecine en l'université de Bâle", il fût aussi connu comme mathématicien, et il semble qu'on le confonde parfois avec son père Jean Bernoulli (1667-1768) mathématicien très célèbre en son temps qui a été le professeur de Leonhard Euler (1707-1783). On doit à Euler un théorème de dynamique des fluides d'un usage plus général que celui de Bernoulli.

Heinrich Gustav Magnus (1802-1870) a été le premier à étudier l'effet qui porte son nom et qui permet d'expliquer les effets de balle dans les sports. Cet effet permet aussi d'expliquer le fonctionnement du boomerang que les aborigènes d'Australie utilisent depuis des milliers d'années.

Explication : Quand la vitesse d'un fluide augmente, sa pression diminue. lorsqu'une balle en rotation se déplace dans l'air, elle va par frottement modifier la vitesse du courant d'air autour d'elle. L'effet sera dissymétrique : d'un côté de la balle la vitesse de rotation s'ajoute à la vitesse de l'air, qui se trouve accéléré et sa pression diminuée, tandis que de l'autre côté la vitesse de rotation se soustrait à celle de l'air, qui se trouve freiné et sa pression augmentée. On aura donc une différence de pression et la balle va se déplacer du côté où la pression est la plus faible. Tout l'art du footballeur est dans la manière de communiquer la vitesse de rotation au ballon.

Dans les années 1920, un ingénieur allemand, Anton Flettner, connu aussi pour avoir construit des hélicoptères durant la seconde guerre mondiale, a mis en application sur un navire, un système de propulsion basé sur l'effet Magnus (un cylindre en rotation dans un courant d'air développe une force perpendiculaire à ce courant d'air). Sur un voilier démâté de 52 m, le BUCKAU, il fit installer deux cylindres verticaux de 15 m de hauteur et 2,75 m de diamètre. Les cylindres étaient entraînés en rotation par un petit moteur. L'expérience fut concluante puisque le navire se déplaça plus vite qu'à la voile.

En 1926, Flettner entreprit la traversée Hambourg-New York via Les Canaries, avec son navire alors renommé BADEN-BADEN. Parti le 2 avril, il arriva le 9 mai. La Hambourg Amerika Line fit construire alors un navire à rotors. Long de 92 m et doté de 3 rotors (h = 17 m, Ø = 4 m, vitesse de rotation = 150 tr/min), il emportait 3 000 t de marchandises. Après six années d'utilisation, les cylindres furent démontés; Le navire fut équipé d'une propulsion classique. Le bateau Alcyone de la fondation Cousteau est le seul navire actuel qui fonctionne avec une turbo voile utilisant l'effet Magnus



Figure 2 : La turbo-voile de l'Alcyon photo de l'université de Rhode island [5]

Lois physiques : Le Théorème de Bernoulli

Hypothèses :

On considère l'écoulement permanent d'un fluide parfait incompressible. Fluide parfait signifie que les forces de frottement (viscosité) sont négligées. Il n'y a pas de machine hydraulique, les seules forces extérieures sont la pesanteur et les forces pressantes.

En appliquant le théorème de l'énergie cinétique à ce fluide, on obtient :

$$\rho \frac{v^2}{2} + \rho gz + p = \text{Constante} \quad (3)$$

Cette équation traduit la conservation de l'énergie. Elle est écrite ici en terme de pression : p étant la pression statique, $\rho \frac{v^2}{2}$ étant la pression cinétique, ρgz étant la

pression de pesanteur. On peut aussi l'écrire en terme de hauteur, et c'est l'usage le plus répandu en hydraulique :

$$\frac{v^2}{2g} + z + \frac{p}{\rho g} = \text{Constante} = h \quad (4)$$

où h est la hauteur totale, $\frac{v^2}{2g}$ est la hauteur cinétique,

$\frac{p}{\rho g}$ est la hauteur de pression, z est la cote.

Applications technologiques :

Cette équation est une des plus importantes de la mécanique des fluides.

Beaucoup de phénomènes peuvent être démontrés ou expliqués grâce à elle. Un exemple parmi tant d'autres est le tube de Pitot, ce qui va me permettre de revenir dans le domaine historique.

Un tube de Pitot est un appareil de mesure de la vitesse des fluides très utilisé en aéronautique. Il doit son nom au physicien français Henri Pitot (1695-1771).

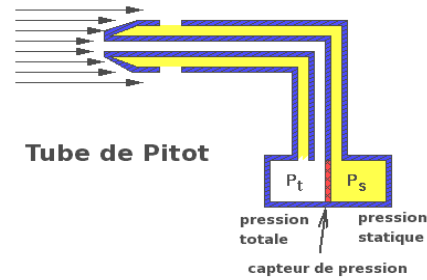


Figure 3 : Tube de Pitot (image issue de [6])

Sa constitution usuelle est la suivante : Deux tubes coudés concentriques, l'un placé perpendiculairement à l'écoulement, a une vitesse relative v égale à la vitesse du fluide et une pression statique p_s égale à la pression ambiante, l'autre placé dans le sens de l'écoulement a une vitesse relative nulle et une pression égale à la pression totale p_t somme de la pression statique et de la pression dynamique. Un capteur de pression ou un simple manomètre mesure alors la différence des deux pressions, celle-ci étant égale à la pression dynamique. Le théorème de Bernoulli montre qu'elle est proportionnelle au carré de la vitesse.

$$\frac{v^2}{2} + \frac{p_s}{\rho} + \frac{p_t}{\rho}, \text{ d'où } v^2 = 2 \frac{(p_t - p_s)}{\rho} \quad (5)$$

Henri Pitot (3 mai 1695 – 27 décembre 1771) était mathématicien et physicien, membre de l'académie des sciences. Il s'intéressa aux problèmes de fluides comme celui de l'écoulement de l'eau dans les rivières, et il a imaginé sa machine pour mesurer la vitesse des cours d'eau. Il est surprenant que cette invention du dix-huitième siècle soit toujours utilisée de nos jours (il y a encore un tube de Pitot à bord de l'airbus A 380).

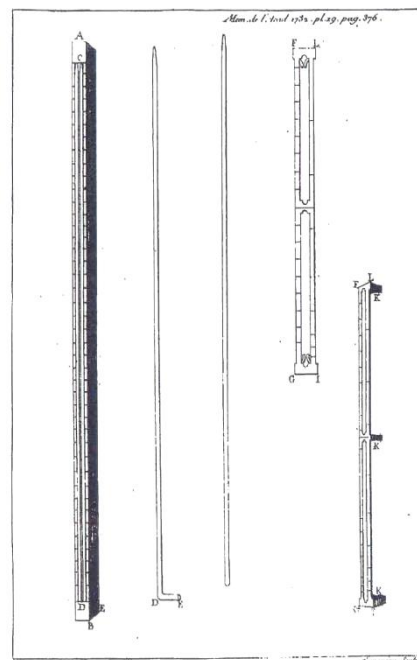


Figure 4 : (image issue de [7])

Il avait proposé sa machine à l'académie royale des sciences le 12 novembre 1732 sous le titre : *"Description d'une machine pour mesurer la vitesse des eaux courantes et le sillage des vaisseaux"*, par M. Pitot (texte complet sur [8]).

Voici la description qu'il fit de sa machine :

"AB est une tringle de bois taillée en forme de prisme triangulaire; sur le milieu d'une des trois faces de cette tringle on a creusé une rainure capable de loger deux tuyaux de verre blanc, l'un de ces tuyaux est courbé à angle droit en D et le bout DE passe dans un trou fait à la tringle. La face CD, dans laquelle les tuyaux sont logés est divisées en pieds et pouces. FGIL est une règle mobile de cuivre refendue dans le milieu sur presque toute sa longueur de la quantité de la somme des diamètres des tuyaux, en sorte qu'elle ne couvre les tuyaux qu'à ses extrémités et un peu à son milieu. Un des côtés de cette règle est divisé en pieds et pouces pour les hauteurs des chutes d'eau, et l'autre côté en pieds et pouces de vitesse de l'eau relative aux hauteurs."

Il faut remarquer que la présentation de la "machine de Pitot" précède de 6 ans la publication de l'ouvrage où Bernoulli exprime ce qui deviendra le théorème de Bernoulli. Pourtant la présentation usuelle des ouvrages de mécanique des fluides fait que le "tube de Pitot" est souvent donné comme le premier exercice d'application du théorème de Bernoulli. Il est toujours intéressant de faire constater à nos élèves qu'historiquement beaucoup d'applications ont précédées les théories dont elles découlent dans les manuels de physique.

Mais laissons Pitot nous expliquer le principe de sa machine :

"Pour savoir maintenant la quantité de vitesse des eaux courantes relative à leur ascension dans le tuyau recourbé de la machine, il faut se rappeler le principe fondamental de presque toute la théorie du mouvement des eaux, qui est, que les vitesses des eaux sont en raison sousdoublée de la hauteur de chute. Mr Varignon a eu le premier la gloire de démontrer ce principe... mais les élévations ou ascensions de l' eau dans notre tube étant égales aux chutes, il s'ensuit que les vitesses des courants seront en raison sousdoublée des élévations de l'eau, ou comme les carrés des vitesses"

Je n'avais pas cité Pierre Varignon (1654-1722), mathématicien et physicien français, membre de l'académie royale des sciences, mais il faut reconnaître qu'il est un peu oublié et rarement cité dans l'histoire de la mécanique des fluides.

4. La Viscosité – les pertes de charge

Question:

Comment savoir si un oeuf est cuit ou cru ?

A priori rien dans l'aspect ne distingue un oeuf cru d'un oeuf dur, tout le monde connaît le petit truc qui permet de les distinguer :

On fait tourner l'œuf sur lui-même, l'œuf cru s'arrête rapidement de tourner tandis que l'œuf dur continu sa rotation. On peut aussi faire le test suivant, on fait tourner rapidement l'œuf sur lui-même, on l'arrête brusquement et on le relâche, s'il recommence à tourner, il est cru.

L'explication étant l'existence de forces de frottements entre la coquille et le jaune d'œuf qui dissipent rapidement l'énergie cinétique que l'on a communiqué à l'œuf lors de sa mise en rotation ou permettent un stockage de cette énergie cinétique lors d'un arrêt brutal.

La viscosité (du latin viscum) désigne la capacité d'un fluide à s'écouler. En langage courant, on utilise aussi le terme de *fluidité*.

Approche historique :

Les scientifiques (Bernoulli, Mersenne, Pitot, etc) du XVIIIe siècle qui sont les fondateurs de la mécanique des fluides, ont bien conscience des frottements exercés par les fluides en mouvement sur les parois des récipients. Voici un extrait du texte de Pitot précédemment cité :

" D'un autre côté on oppose à toutes ces raisons la quantité de frottements des eaux contre le fond ou le lit et les bords des fleuves: il est vrai, comme je crois l'avoir démontré dans un mémoire en 1730 (réflexions sur le mouvement des eaux académie royale des sciences 1730), que cette quantité de frottements est prodigieuse, et il est heureux qu'elle le soit, car sans les frottements, les fleuves et les rivières ne seraient pas navigables...ainsi sans les frottements presque toutes les eaux courantes seraient des torrents affreux dont on ne tirerait aucun avantage".

La formulation mathématique de la viscosité attendra Jean-Louis-Marie Poiseuille (22 avril 1797, Paris - 26 décembre 1869, Paris) médecin à qui l'on doit différents mémoires sur le cœur et la circulation du sang dans les vaisseaux (l'hémodynamique), il établira en 1844 - dans son ouvrage "Le mouvement des liquides dans les tubes de petits diamètres" les lois de l'écoulement laminaire des fluides visqueux dans les tuyaux cylindriques et où il sera le premier à faire intervenir un coefficient de viscosité pour les liquides.

Lois physiques :

Définitions :

Soit un fluide qui s'écoule dans la direction Oy, le module v de sa vitesse en un point quelconque n'étant fonction que de la coordonnée x . La force $\vec{f}$ qui s'exerce sur une surface S parallèle au mouvement du fluide, est donnée par :

$$\vec{f} = \eta \frac{\delta \vec{v}}{\delta x} . S \quad (6)$$

Le coefficient η est appelé viscosité dynamique. Son unité dans le S.I. est le pascal-seconde (Pa.s) autrefois appelé Poiseuille Pl.

Le coefficient $\nu = \frac{\eta}{\rho}$ où ρ est la masse volumique du fluide est appelé viscosité cinématique son unité S.I. est : m².s⁻¹

La viscosité des liquides diminue beaucoup lorsque la température augmente.

Applications technologiques : pertes de charge :

Il me semble naturel de considérer la notion de pertes de charge comme une application de la notion de viscosité, car c'est dans ce domaine que l'on aura le plus de calculs d'application à faire effectuer à nos élèves.

Du fait de la viscosité, la pression d'un fluide réel diminue tout au long d'une canalisation, ce qui nécessite parfois d'ajouter des pompes intermédiaires.

Cette diminution de la pression correspond à une dissipation d'énergie que l'on appelle pertes de charge quand on l'exprime en mètre de colonne de liquide.

On distingue deux types de pertes de charge :

- Des pertes de charge systématiques dues aux frottements du liquide sur les parois des tuyaux, et qui existent aussi dans des tuyaux lisses.
- Des pertes de charge accidentelles liées aux accidents de parcours du fluide, coude, rétrécissement, vanne, etc.

Un petit tour du côté de l'histoire s'impose ne serait-ce que pour citer l'incontournable Reynolds.

Henry Darcy (1803-1858) a proposé une variante de l'équation de Gaspard de Prony (1755-1839) permettant de calculer la perte de charge due aux frottements dans une conduite modifiée à son tour par Julius Weissbach en 1845

On doit à Osborne Reynolds (1842-1912), en 1883, la caractérisation par un nombre sans dimension, le nombre de Reynolds, de la nature d'un écoulement : laminaire, transitoire ou turbulent. Il peut se comprendre comme le rapport des forces d'inertie aux forces de frottements (viscosité). Aux faibles valeurs de ce nombre, inférieure à 2000, les forces de viscosité l'emporte et l'écoulement est laminaire, aux fortes

valeurs, supérieures à 3000, les forces d'inertie prédominent et l'écoulement est turbulent. Entre les deux, on a un régime intermédiaire.

Pour terminer cette présentation, il me semble indispensable de citer Henri Navier (1785-1835) et de George Stokes (1819-1903) dont les noms sont associés aux "équations de Navier-Stokes" qui décrivent le mouvement des fluides. Ce sont des équations aux dérivées partielles non-linéaires que je me garderai bien d'écrire et qui traduisent la conservation de la masse, la conservation de la quantité de mouvement et la conservation de l'énergie.

Cette présentation est terminée, j'espère qu'elle plaira à mes élèves. Si des lecteurs attentifs trouvent des erreurs dans le texte, merci de me les signaler, je serais ravi de les corriger. La plus grande partie des renseignements provient de recherches sur internet. J'ai essayé de ne mettre que des éléments que j'ai pu vérifier. Mais les livres anciens sont loin d'être tous numérisés.

5. Bibliographie

- [1] www.ac-nancy-metz.fr/enseign/physique/PHYS/Term/Mecaflu/mecanique_des_fluides.htm (pour un cours niveau terminale sciences et techniques de laboratoire)
- [2] www.ac-nancy-metz.fr/enseign/physique/PHYS/Bts-Cira/mecaflu/mecaflu.htm (pour un cours niveau BTS contrôle industriel régulation automatique)
- [3] www.grc.nasa.gov/WWW/K-12/airplane/short.html (pour ceux qui comprennent l'anglais et qui aiment bien les applets java le site de la NASA sur l'aérodynamique est très intéressant)
- [4] "Discorsi" paru à Leyde en 1638 (traduction du musée de Florence).
- [5] www.gso.uri.edu
- [6] wikipedia.org
- [7] (image: Académie des sciences numérisée par la bibliothèque nationale de France gallica.bnf.fr)
- [8] http://www.academie-sciences.fr/archives/doc_anciens/hmvol3529_pdf/p363_376_vol3529m.pdf

La perturbographie et l'analyse des réseaux

Martin HANKER / LEM INSTRUMENTS & Denis KOBLER / QUALITROL

Résumé : la nécessité de surveiller les réseaux s'impose, l'oscilloperturbographie est une des réponses à ce besoin.

Pourquoi et comment assurer le fonctionnement des réseaux ?

«Le marché de l'électricité s'ouvre à la concurrence» ! De nombreux articles de presse évoquent cette ouverture et en précisent les modalités sans réellement en évoquer les enjeux. Les termes de libéralisation, dérégulation ou déréglementation sont employés, sans que soient apportées quelques nuances.

L'électricité n'est pas un bien comme les autres

Les directives européennes sur la dérégulation du marché de l'énergie concernent autant le gaz que l'électricité. Cependant, le débat s'est focalisé sur l'électricité, notamment du fait de son caractère plus englobant.

L'électricité, vitale et peu substituable dans la plupart de ses usages, ne se stocke pas de façon économiquement viable. Elle doit donc être acheminée depuis son lieu de production jusqu'à ses points de consommation, aussi éloignés fussent-ils, au travers de réseaux.

Son caractère stratégique, la maîtrise technologique et l'importance des enjeux économiques ont longtemps justifiés le concept de monopole ; mais la pression de l'économie de marché a fini par l'emporter, avec certaines nuances/

En effet, de plus en plus de libéraux souhaitent une distinction entre ce qui relève du réseau (les infrastructures, coûteuses, qui doivent être centralisées) et le service à forte valeur ajoutée (la commercialisation, qui peut être assurée par diverses entreprises). Tous ces éléments, dans un contexte général de dénigrement de l'action publique (jugée moins efficace) et de valorisation de la concurrence depuis la fin des années 70 (néolibéralisme), sont en faveur de la déréglementation.

Néanmoins, la libéralisation ne sera jamais totale. Elle se fera progressivement par introduction de mécanismes concurrentiels respectant ses caractéristiques de réseau et sous le contrôle d'instance de régulation. C'est pour cela que nous parlons de dérégulation du marché, vue comme une modification des règles de fonctionnement du marché, et non de déréglementation (les règlements persisteront).

L'interconnexion des réseaux

Les interconnexions régionales, nationales ou internationales des réseaux électriques sont généralisées à l'échelle mondiale. Cette tendance se justifie aisément par les avantages importants qui en découlent :

- Amélioration de la fiabilité générale des réseaux
 - Aide mutuelle en cas de défaillance
 - Sécuriser l'alimentation en électricité
- Favoriser la compétition
 - Favoriser les échanges transfrontaliers
 - Transporter (vendre) l'électricité la moins chère
- Economiques
 - Partage des réserves de génération particulièrement en cas de demandes décalées (France – GB)
 - Utilisation optimale de grande centrales, plus économiques
 - Maintien de réserve de sécurité conservatrices
 - Exportation d'électricité d'un pays en surcapacité vers un pays en sous-capacité, que ce soit temporaire ou permanent (France – Italie)
 - Possibilité d'implanter les nouvelles centrales en des lieux plus favorables
- Ecologiques
 - Migrer vers une filière électrique moins polluantes (hydroélectricité)
 - Installation de points de génération (éoliens) en des localisations favorables

Ces avantages indéniables ont cependant leurs revers. Ils induisent, en effet, une complexité accrue. En outre, et c'est plus préoccupant, les investissements dans les réseaux de transmissions évoluent nettement moins vite que la demande. A titre d'exemple, au cours des dix dernières années la demande aux Etats-Unis a augmenté de 35% alors que les investissements dans les réseaux de transport n'ont progressés que de 18%. Les projections pour les dix prochaines années vont dans le même sens : une demande en hausse de 20%, mais moins de 6% en transmission. Et par voie de conséquence les systèmes opèrent et opèreront de plus en plus près de leurs limites.

Les black-out : conséquences de la déréglementation ?

Pour les partisans de la régulation étatique, les nombreuses coupures d'électricité, en Californie, à l'est des Etats-Unis, en Grande-Bretagne, dans les pays scandinaves ou en Italie, sont le résultat de la logique concurrentielle. C'est probablement en partie exact, même si dans un certain nombre de cas l'origine de la panne est liée à des causes simples comme la chute d'un arbre. Il n'en reste pas moins vrai que les conséquences ont des effets économiques dévastateurs, à cause de l'effet domino rendu possible ... par les interconnexions

Contrôle et surveillance des réseaux

La surveillance des réseaux est assurée en temps réel (en ligne) par les systèmes SCADA. Ceci n'empêche pas des défaillances de se produire. Celles-ci sont causées par des phénomènes complexes, généralement transitoires, dont l'analyse, indispensable, ne peut se réaliser qu'après leur survenance, sur base d'enregistrements de l'ensemble des signaux (courants, tensions, position des protections, ...). On parle alors d'analyse hors ligne ou post-mortem.

Ces analyses se basent sur des appareils indépendants, type « boîte noire » présentant les caractéristiques suivantes

• avantages:

- performances sophistiquées
- fiabilité
- vision de l'ensemble du réseau
- déclenchement sur des grandeurs non-électriques

• inconvénients

- coût

Il est également possible de se baser sur des enregistrements provenant des protections :

• inconvénients

- un non-fonctionnement d'une protection n'engendrera pas d'enregistrement (juge et partie)
- Vision limitée du réseau
- Performances limitées

Le perturbographe

Le perturbographe (ex oscilloperturbographe) est un équipement qui permet l'enregistrement et la visualisation des perturbations (défauts) affectant les réseaux électriques. Il offre à l'électricien une image des grandeurs électriques analogiques (variables en amplitude) et logiques (états tout ou rien).

Souvent considéré comme la "boîte noire" du réseau par analogie avec celle des avions, le perturbographe a montré son utilité bien au-delà de cette description restrictive. Outre son utilisation pour le diagnostic d'incident "post mortem", il se révèle en effet un outil important pour assurer la maîtrise opérationnelle et la pérennité de l'équipement électrique haute tension.

Le schéma de la Figure 1 décrit les éléments fondamentaux de son architecture, dont les caractéristiques principales sont décrites dans l'article (immunité, flexibilité, puissance de calcul, capacité de mémoire, communication, ...).

La fonctionnalité perturbographique évolue en raison de plusieurs facteurs: augmentation de l'intégration fonctionnelle, élaboration de systèmes experts, construction d'une base de connaissance et obligations contractuelles suite à la dérégulation.

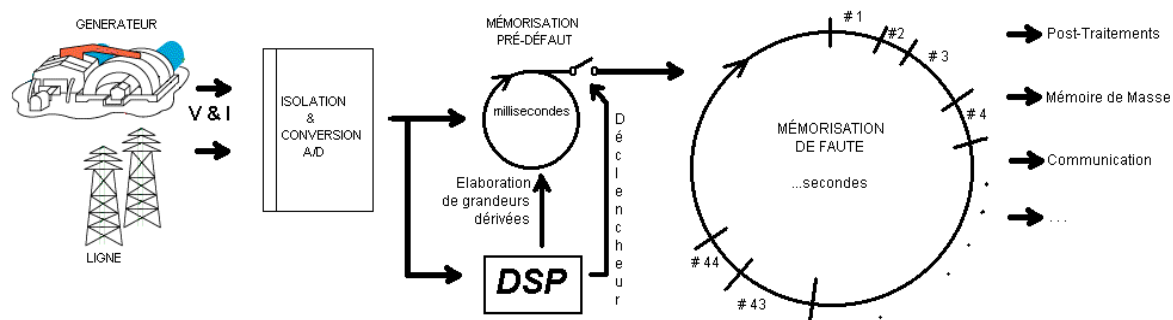


figure 1 ; schéma-bloc du perturbographe

Applications du perturbographe

Les informations capturées par le perturbographe sont exploitées à différentes fins et de différentes manières:

⇒ Analyse “a posteriori”

Résolution d'un défaut complexe pour laquelle les autres équipement de contrôle ne suffisent pas (Identification et résolution du défaut: 1ère analyse).
Vérification de la performances de divers équipements pendant l'incident (2ème analyse).
Localisation du défaut en ligne.
Compilation de statistiques de performance du réseau et/ou de maintenance des protections.

⇒ Utilisation opérationnelle

Mesures de contrôle régulières.
Enregistrement en continu

⇒ Etude préventive

Utilisation d'enregistrement réels pour la validation de nouveaux modèles de protections.
Validation des modèles mathématiques de simulation du réseau.

Considérations technologiques

Le perturbographe étant installé au sein même des sous-stations, il est lui-même soumis à toutes les perturbations d'ordre électromagnétique conduites et/ou induites qui peuvent être importantes, surtout en cas d'incident. Par analogie avec la “boîte noire” de l'avion, c'est essentiellement dans les conditions environnementales les plus critiques que l'on attend du perturbographe qu'il fonctionne normalement. Il doit donc être caractérisé par une excellente immunité aux perturbations.

Des normes (CEI 6100-4-x) décrivent de manière spécifique les niveaux de perturbation auxquelles les perturbographes peuvent être soumis ainsi que les comportements attendus.

Le perturbographe (voir le schéma bloc en Figure 1) se caractérise aussi par son nombre de canaux d'entrées. Ainsi, selon sa taille, il pourra capturer les signaux relatifs à une ligne (typiquement 3 tensions et 3 ou 4 courants, plus les contacts des protections associées) ou à plusieurs lignes dans un même poste à haute tension, permettant ainsi la comparaison de signaux instantanés. Dans les centrales électriques, il pourra enregistrer des dizaines de grandeurs relatives au fonctionnement d'un groupe générateur.

La perturbographie moderne permet en outre l'enregistrement de grandeurs dérivées des grandeurs physiques d'entrées, telles que: les puissances active et réactive déduites des tensions et courants d'entrée, la tension homopolaire à partir des tensions des trois phases, la fréquence du réseau, etc.

Suivant le type de perturbation à capturer, la fréquence d'échantillonnage (la résolution en temps) sera adaptée. Les appareils les plus puissants permettent la capture de plusieurs types de perturbations à des vitesses différentes au sein d'un même équipement.

Types de perturbation	Courts-circuits	Oscillations de puissance fréquence	Harmoniques
Composantes fréquentielles du phénomène	...1000 Hz	...15 Hz	...2500 Hz
Durée du phénomène	≤ 3 s	≤ 10 mn	...x cycles
Fréquence d'échantillonnage	≥ 3000 Hz	30 ≤ 200 Hz	≥ 5000 Hz

Les déclencheurs sont fondamentaux dans la description de la fonction de perturbographie. En effet, ce sont ces organes qui permettent la reconnaissance automatique des phénomènes perturbateurs. Les plus simples reconnaîtront une sous- ou surtension, un sur-courant ou le changement d'état d'une entrée logique. Des déclencheurs plus élaborés font intervenir la capacité de traitement du perturbographe pour calculer des grandeurs dérivées auxquelles un critère sera appliqué : tension/courant homopolaire, direct ou inverse, taux de variation de la puissance active ou réactive, oscillation lente de la puissance, ...

Les perturbographes numériques modernes sont équipés de processeurs de traitement de signal (Digital Signal Processors : DSP) qui permettent des traitements complexes des grandeurs d'entrée en temps réel.

Les enregistrements ainsi capturés sont mémorisés et mis à la disposition des utilisateurs locaux ou à distance. Un nombre plus ou moins élevé d'enregistrements sera conservé dans l'appareil selon la capacité de la mémoire. Cette mémoire est temporaire ou permanente en fonction de la technologie utilisée (semi-conducteurs, disques, etc.). Elle se chiffre généralement en secondes d'enregistrement pour l'ensemble des canaux plutôt qu'en Mega-octets.

Une fois acquise cette assurance que les défauts peuvent être capturés, l'accent sera mis essentiellement sur la manipulation et le traitement de ces informations.

Les données mesurées sont généralement conservées à proximité de leur lieu d'acquisition, qui peut être très distant de l'utilisateur (parfois multiple). Des procédés de communication rapides et sûrs sont mis à la disposition par les sociétés d'électricité afin de transmettre les enregistrements dans des temps réduits, que ce soit de manière automatique ou à la demande.

Il est toutefois important de noter que certains standards de communication existants, dédiés au contrôle/commande (ModBus™, DNP3.0™, IEC870-5-90, ...), ne sont pas adaptés au transfert de perturbographies de grande taille (de quelques centaines de Kilo-octets à plusieurs Mega-octets) tels que celles générées par les perturbographes. Ceci a poussé au développement de la norme CEI 61850 qui définit les systèmes de communication dans les postes et les centrales.

D'une part, un équipement performant générera un grand nombre d'enregistrements :

- de phénomènes jusque là ignorés ,
- de phénomènes résultant de charges et variations importantes imposées au réseau suite à la dérégulation des sociétés d'électricité.

D'autre part, la dérégulation augmente la pression sur le personnel des départements de protections et d'analyse (rationalisation, concentration sur les activités d'exploitation).

Loin de requérir une communication instantanée, l'exploitant souhaite quand même dans certains cas (voir les applications ci-dessus) avoir accès aux enregistrements dans un délai restreint. L'information doit en outre éventuellement être accessible à plusieurs utilisateurs.

La figure 2 illustre quelques visualisations d'un même incident qui diffèrent selon que l'utilisateur est intéressé par le transitoire rapide (valeurs de court-circuit, etc.), par le profil de qualité de l'électricité, par les phaseurs ou par la séquence des événements.

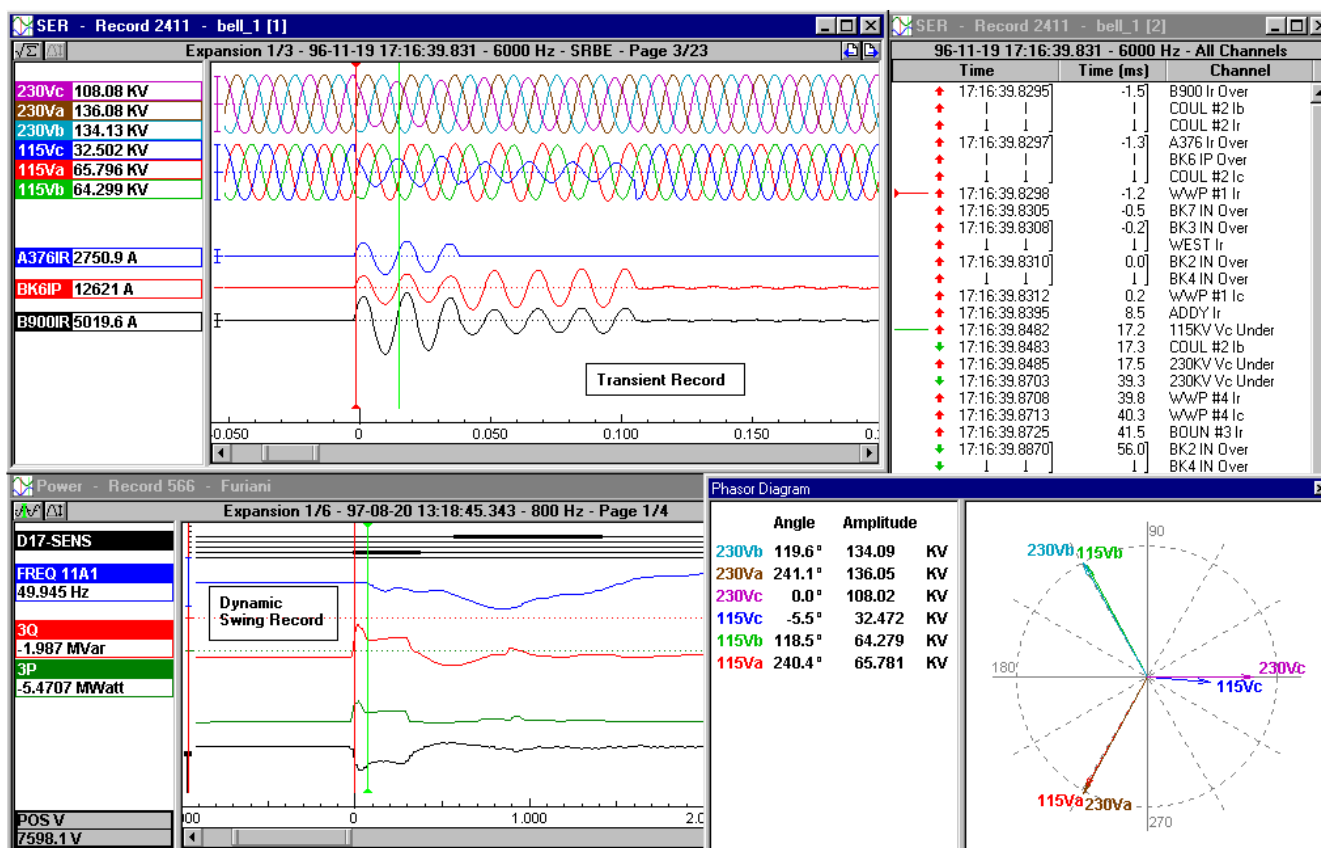


figure 2: Visualisation de données mesurées.

Vu le nombre et la variété croissants des enregistrements, ils peuvent éventuellement être transmis selon un ordre de priorité (importance en fonction du type d'incident, par exemple) afin de faciliter la résolution de problèmes urgents, par opposition à la collecte d'informations à des fins statistiques. La charge imposée au système de communication et la rapidité d'analyse en seront améliorés.

Tendances

Intégration :

Tandis que certaines fonctions ont tendance à être de plus en plus intégrées au sein d'un même équipement (par ex. plusieurs algorithmes de protection fonctionnant dans le même relais), l'intégration de la perturbographie au sein des équipements de protection soulève une question plus philosophique que technique : « Est-il souhaitable que l'équipement de protection fournisse lui-même toutes les informations nécessaires à l'évaluation de sa performance ? ».

Il va de soi que la fiabilité - tant matérielle que logicielle - et les caractéristiques d'acquisition des divers éléments de contrôle vont déterminer jusqu'où la chaîne d'acquisition peut éventuellement être intégrée. A l'heure actuelle, les caractéristiques de la mesure par les relais de protection numériques (vitesse, bande passante, synchronisme, précision, ...) ne permettent pas de franchir ce pas sans abandonner certaines des caractéristiques fondamentales de la perturbographie moderne.

Système expert (pour l'aide opérationnelle):

Dans la mesure ou des moyens de transmission rapides sont mis en place et que l'information relative à la structure physique du réseau est aussi disponible, l'ordinateur recevant les enregistrements peut réaliser une analyse automatique.

Cette analyse peut avoir pour objectif :

- # l'identification précise du type de défaut rencontré
- # la localisation du défaut sur la ligne
- # l'analyse de la performance du système
- # la compilation de données à des fins de maintenance statistique

Cette automatisation aura bien sûr pour effet de permettre aux analystes et ingénieurs de protection de se concentrer sur les défauts les plus complexes et/ou sur l'analyse plus approfondie permettant éventuellement de déceler une dérive de performance avant que ne survienne une défaillance d'équipement (fondement de la maintenance préventive).

Source de connaissance

Élément fondamental à la bonne gestion du réseau, le niveau de maîtrise des équipements sera fort dépendante de la capacité des ingénieurs à acquérir une bonne compréhension des phénomènes qui l'affectent. A cet effet, le perturbographe offre une vision privilégiée des phénomènes électriques rapides et lents. L'analyse qualitative et quantitative des enregistrements perturbographiques permet un accroissement des connaissances qui se traduit par l'ajustement plus rapide et plus précis des équipements de protection. Il faut aussi souligner l'intérêt pédagogique de l'analyse d'incidents pour les jeunes ingénieurs, qui pourront mieux maîtriser leur fonction.

En outre, les enregistrements de défauts réels saisis par les perturbographes peuvent être utilisés pour évaluer la performance de nouvelles protections (ou logiciels), avant même que ces systèmes ne soient installés dans le réseau. Cette pratique contribue à réduire les cycles de développement, d'implémentation et d'évaluation.

Certains laboratoires procèdent de manière systématique à l'essai d'un nouveau relais de protection en le soumettant à un ensemble de défauts propres au réseau où cette protection sera installée.

Obligations contractuelles :

La dérégulation dans les réseaux électriques pourrait avoir un impact sur les types d'équipements installés aux différents points stratégiques du réseau. Le National Electricity Reliability Council (NERC) aux USA a imposé aux nouveaux opérateurs libéralisés d'enregistrer systématiquement les perturbations et d'en effectuer leur compilation statistique faute de quoi des pénalités financières seront appliquées. La transparence du marché est à ce prix. Dans un marché libéralisé, il sera important de pouvoir prouver l'origine de perturbations affectant la fourniture d'électricité, qui ont parfois des conséquences contractuelles très importantes.

Conclusions

La perturbographie évoluera sous l'influence de facteurs de deux natures :

l'évolution du marché de l'électricité
les progrès technologiques.

Là où par le passé, l'installation d'un perturbographe était justifiée par le souci de performance technique d'un opérateur monopolistique, la libéralisation des marchés pourra associer à l'information fournie par les perturbographes une valeur qui se traduira éventuellement en pénalités financières pour les

opérateurs peu consciencieux. Il n'est plus rare de voir un gros consommateur d'électricité négocier avec son (ou ses ?) fournisseur(s) en des termes incluant l'impact des perturbations réciproques. Un dialogue similaire n'est pas impensable entre producteurs, transporteurs et distributeurs d'électricité.

Les besoins en équipements d'enregistrement et de contrôle seront d'autant plus importants que les conditions de marché seront tendues et que les équipements (lignes, transfos, ...) seront exploités au plus près de leur limites.

L'utilisation intensive des processeurs de signaux (DSP) permet l'élaboration de grandeurs dérivées en continu à partir d'un nombre limité de signaux de base (U, I). Ces grandeurs dérivées, parfois très élaborées (dérivée de la fréquence, composante inverse, ...) permettent de détecter des phénomènes autres que les simples courts-circuits et chutes de tension.

La numérisation de la mesure à des fins multiples (facturation profilée, protection, perturbographie, ...) pourra éventuellement fort affecter le paysage des équipements présents dans les sous-stations. En effet, pour autant que l'organe de mesure fournisse une précision, une bande passante et une fiabilité satisfaisant à toutes les fonctionnalités, il ne serait pas étonnant de voir apparaître cette combinaison de fonctions de mesure, avec redondance bien sûr, dans des équipements (ou programmes) de fabricants différents de ceux d'aujourd'hui.

On pourrait ainsi assister à un déplacement de l'expertise des fabricants d'équipement de la mesure vers la détection et le traitement des signaux. Dans un environnement où, pour des raisons diverses, il y a une tendance à la pénurie en personnel disposant de l'expertise nécessaire à l'exploitation des données, ce traitement des signaux pourrait fournir une assistance utile en vue d'augmenter l'efficacité de la gestion de ces données.

Il serait toutefois prudent de conserver une indépendance fonctionnelle de la "boîte noire" perturbographique, dernier rempart face à de possibles défaillances d'un système par ailleurs très intégré. Cela ne signifie pas que la perturbographie doive rester dans le domaine du traitement à posteriori. De plus en plus, l'amélioration des performances en communication aidant, le perturbographe pourra fournir des données en temps réel pour assister la gestion du réseau. Il doit cependant conserver sa spécificité d'acquisition de données pour la supervision du processus de production et distribution de l'énergie électrique et donc rester distinct des automatismes (protections, etc.) qu'il contribue à superviser.



perturbographe ; face avant et face arrière



A l'origine de l'automatique : Black, Nyquist, Bode et les Bell Laboratories

Patrice REMAUD, Jean-Claude TRIGEASSOU

Laboratoire d'Automatique et d'Informatique Industrielle de Poitiers
Ecole Supérieure d'Ingénieurs de Poitiers

Résumé : Cet article présente l'apport des électroniciens à l'automatique au début du XX^e siècle en relatant les travaux effectués par Harold Black (1898-1983), Harry Nyquist (1889-1976) et Heindrik Bode (1905-1982) au sein des Bell Laboratories, travaux qui ont abouti aux concepts de l'automatique actuelle.

Mots clés — Histoire des sciences, histoire des techniques, histoire de l'automatique.

I. INTRODUCTION

Black, Nyquist et Bode sont des noms familiers aux automaticiens qui les imaginent facilement en pionniers des asservissements. La réalité historique est sensiblement différente, mais tout aussi passionnante et enrichissante. Ces trois ingénieurs en électronique et télécommunications, possédant de solides compétences en mathématiques, ont travaillé au sein des Bell Laboratories. Ils avaient pour mission de développer le téléphone transcontinental aux USA, afin de joindre les côtes est et ouest, avec un maximum de canaux par ligne. La solution adoptée consistait à utiliser plusieurs porteuses modulées en amplitude, véhiculées par la même ligne. Compte tenu de l'atténuation due à la distance, des amplificateurs ou répéteurs étaient nécessaires. Ces amplificateurs, basés sur l'utilisation de lampes triode, introduisaient une inacceptable distorsion d'intermodulation à cause de leur non linéarité.

C'est Black qui a réinventé le principe de la boucle fermée pour linéariser avec succès ces amplificateurs. On lui doit la contre réaction et sa formalisation sous forme de boîte noire, telle que nous l'utilisons aujourd'hui. Nyquist, pour sa part, a formalisé le problème de la stabilité de ces amplificateurs grâce à l'approche fréquentielle : on lui doit son célèbre critère qui permet de prévoir la nature des régimes transitoires de la boucle fermée à partir de l'examen d'une courbe de réponse en fréquence obtenue expérimentalement. Enfin, Bode a repris les travaux de Black et surtout ceux de Nyquist afin de synthétiser des amplificateurs correspondant au cahier des charges imposé par les problèmes de téléphonie. On lui doit ses marges de gain et de phase, ainsi que la relation asymptotique entre atténuation du module et variation de la phase, à l'origine des diagrammes qui portent son nom.

Cet exposé a pour objectif de présenter les travaux de ces trois grands noms de l'automatique, en les replaçant dans la problématique de la téléphonie à grande distance, et en les situant dans la structure de recherche des Laboratoires Bell.

II. LA REVOLUTION DES TELECOMMUNICATIONS

Le téléphone, mais aussi le télégraphe et la radiodiffusion, furent inventés à la fin du XIX^e siècle en liaison avec le développement des applications de l'électricité. Dans un article de *L'illustration*, le Français Charles Bourseul proposait dès 1854 un mécanisme permettant de faire varier l'intensité du courant électrique en fonction de la voix. Le téléphone ne fut effectivement créé qu'une vingtaine d'années plus tard par Graham Bell (1847-1922) et Elisa Gray (1835-1901) qui déposèrent leur brevet de système téléphonique le même jour de 1876 aux Etats Unis. Cependant l'antériorité du brevet en fut reconnue à Graham Bell.

L'entreprise créée par Bell en 1877 aux Etats Unis, la Bell Telephone Company, devint en 1899 l'American Telephone and Telegraph Company (AT&T). Elle domina le marché du téléphone aux Etats Unis à partir des années 1910 en s'appuyant sur une structure de fonctionnement appelée depuis le Bell System. Cette structure était composée de la Western Electric pour la fabrication du matériel de télécommunications et d'un laboratoire de recherche, les Bell Laboratories, pour la recherche et le développement, inspirée en cela par l'exemple de la General Electric d'Edison (1847-1931).

Les Bell Laboratories, créés en tant que département indépendant au sein du Bell System en 1925, deviendront un des plus grands centres de recherche aux Etats Unis. Ce laboratoire jouera un rôle déterminant dans la domination d'AT&T sur l'industrie du téléphone aux Etats Unis en raison de la pertinence de sa production scientifique. Les Bell Laboratories emploieront à partir des années 1930 Harold Black, Harry Nyquist et Heindrik Bode (ce dernier devenant après la deuxième guerre mondiale le vice-président des Bell Laboratories).

III. LA RETROACTION

La rétroaction positive ou « regeneration », c'est à dire l'action d'ajouter une fraction de la grandeur de sortie d'un système à sa variable d'entrée dans le but d'en augmenter

la valeur, fut utilisée pour la génération des signaux de haute fréquence (par oscillation du dispositif) utilisés comme porteuses en modulation d'amplitude, grâce à la triode. L'invention de ce tube électronique (ou « audion ») par Lee de Forest en 1906, en ajoutant une électrode à la diode inventée par Fleming en 1904, a constitué une révolution dans le domaine des télécommunications. L'« audion » dans sa fonction amplificatrice avait un gain faible. Pour l'augmenter, il fallait en associer plusieurs en cascade, c'est à dire réaliser de coûteux amplificateurs multi-étages. La rétroaction positive, à condition de se placer en deçà de l'accrochage (ou instabilité) permet d'obtenir une amplification théoriquement infinie, à condition bien sûr d'accepter les distorsions inhérentes à cette approche [13].

Des brevets sur l'utilisation de la rétroaction positive avec une seule triode pour une amplification élevée seront déposés par plusieurs chercheurs dans différents pays au cours des années 1910, dont Edwin Armstrong [1] aux Etats Unis.

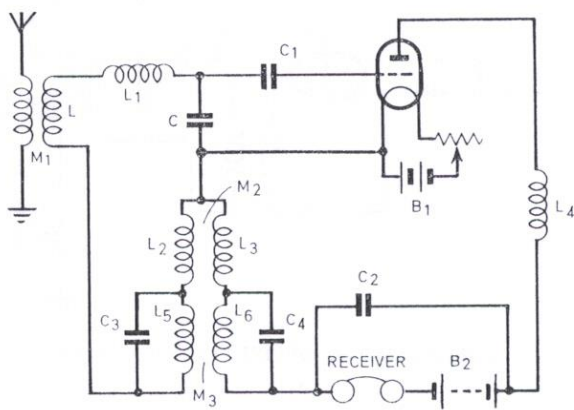


Figure 1 : Circuit de rétroaction positive d'Armstrong

Ainsi, historiquement, la rétroaction positive précéda la rétroaction négative ou contre-réaction. Aujourd'hui, la rétroaction positive n'est plus qu'une curiosité alors que la contre-réaction est le concept fondamental de toutes les applications de l'automatique.

IV. LA RETROACTION NEGATIVE: HAROLD BLACK

Le développement des lignes de téléphone sur de très grandes distances se heurtait à de nombreux problèmes et, entre autres, à celui de l'atténuation du signal sur la ligne. La première solution fut de placer sur la ligne, à intervalles réguliers, des associations récepteur-microphone qui jouaient le rôle d'amplificateurs ou « répéteurs » électromécaniques. Par la suite, ces dispositifs électromécaniques furent remplacés par de véritables amplificateurs électroniques à lampes triode.

En 1915, AT&T fit l'expérience d'une ligne téléphonique de longue distance entre New York et San Francisco destinée à montrer la capacité de l'entreprise à réaliser des liaisons téléphoniques sur des distances transcontinentales. L'atténuation du signal était bien sûr très grande et nécessitait l'utilisation d'un grand nombre

d'amplificateurs. D'autre part, afin de transmettre simultanément plusieurs messages sur la même ligne, les ingénieurs d'AT&T utilisaient un principe de modulation d'amplitude avec différentes porteuses. Le fonctionnement non linéaire des triodes, accentué par le nombre de répéteurs, introduisait une distorsion d'intermodulation totalement préjudiciable à la qualité des transmissions téléphoniques. Ce problème essentiel avait été parfaitement identifié et la solution résidait dans une amplification véritablement linéaire.

Les triodes présentaient naturellement des non-linéarités et une forte dispersion de leurs caractéristiques lors de leur fabrication, accentuée par des variations en fonction de la température. Harold Black, jeune ingénieur électricien embauché en 1921 aux Bell Laboratories, travailla d'abord sur l'amélioration des performances en boucle ouverte des amplificateurs téléphoniques à triode pour le compte de la Western Electric Company (département de production d'AT&T) puis aux Bell Telephone Laboratories d'A T & T. Il s'intéressa donc à la fabrication des triodes afin de limiter leur non linéarité. Ces perfectionnements, quoique bénéfiques à la technologie électronique, ne permirent pas bien sûr d'aboutir à l'objectif initial, à savoir la réduction de l'intermodulation.

Black imagina alors une autre solution, basée sur l'analyse du phénomène de distorsion. La sortie de l'amplificateur comportait le signal d'entrée, correctement amplifié de la valeur μ , plus un terme dû à la distorsion.

En atténuant ce signal de sortie de $\frac{1}{\mu}$ et en le comparant à

l'entrée, il reconstituait par différence le terme de distorsion, qui une nouvelle fois amplifié de la valeur μ , pouvait être retranché de la sortie pour éliminer la distorsion. Black a donc ainsi inventé la correction de type feed forward. Comme il l'a relaté, ce principe lui permit de réaliser un prototype satisfaisant à l'objectif de linéarisation. Néanmoins, cette structure très sophistiquée restait tributaire des non stationnarités (il lui fallait constamment régler le courant d'alimentation des filaments de triode) et ne pouvait être utilisé en pratique sans l'aide d'un opérateur. Après de nombreux essais, Harold Black acquit la conviction que cette approche constituait une impasse ; il chercha alors une autre solution.

Comme il l'a relaté bien plus tard [5], c'est le 6 août 1927, à bord du bateau lui permettant de traverser l'Hudson à New York pour se rendre à son travail, qu'il découvre la solution de son problème sous la forme d'une rétroaction négative. Bien qu'il soit illusoire de vouloir reconstituer le processus intellectuel qui lui permit d'aboutir à la rétroaction négative (ou negative feedback), on peut raisonnablement penser que l'étape de correction feedforward ne soit pas étrangère au murissement de la solution par feedback.

On présente souvent Harold Black comme l'inventeur de la rétroaction négative : en fait cette technique était déjà bien connue dans l'industrie, à titre d'exemple la régulation de vitesse des machines à vapeur et des turbines hydrauliques. De même, il semble que le principe de la contre-réaction ait été utilisé en France sur des



amplificateurs électroniques à la même époque [12]. Le réel mérite de Harold Black est d'avoir formulé le problème de la rétroaction négative, ou contre-réaction, en des termes très généraux, facilement réutilisables avec d'autres technologies, et d'avoir ainsi été à l'origine de la généralisation de ce concept à d'autres domaines techniques.

En particulier, il a montré que la rétroaction négative permet de réduire la non-linéarité d'un organe amplificateur en augmentant le gain de boucle de ce dispositif. Dans sa publication de 1934, nettement postérieure à sa découverte, parue dans *The Bell System Technical Journal* et intitulée *Stabilised Feedback Amplifiers* [4], il propose la formalisation de son principe sous la forme d'un schéma-bloc et d'une mise en équation de l'amplificateur à rétroaction négative.

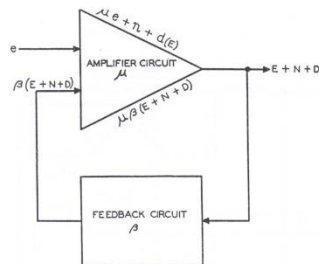


Fig. 1—Amplifier system with feedback.

e —Signal input voltage.
 μ —Propagation of amplifier circuit.
 μe —Signal output voltage without feedback.
 n —Noise output voltage without feedback.
 $d(E)$ —Distortion output voltage without feedback.
 β —Propagation of feedback circuit.
 E —Signal output voltage with feedback.
 N —Noise output voltage with feedback.
 D —Distortion output voltage with feedback.

The output voltage with feedback is $E + N + D$ and is the sum of $\mu e + n + d(E)$, the value without feedback plus $\mu\beta[E + N + D]$ due to feedback.

$$E + N + D = \mu e + n + d(E) + \mu\beta[E + N + D]$$

$$[E + N + D](1 - \mu\beta) = \mu e + n + d(E)$$

$$E + N + D = \frac{\mu e}{1 - \mu\beta} + \frac{n}{1 - \mu\beta} + \frac{d(E)}{1 - \mu\beta}$$

If $|\mu\beta| \gg 1$, $E \approx -\frac{e}{\beta}$. Under this condition the amplification is independent of μ but does depend upon β . Consequently the over-all characteristic will be controlled by the feedback circuit which may include equalizers or other corrective networks.

Figure 2 : Extrait de l'article 'Stabilised Feedback Amplifier'

Le prototype de l'amplificateur à rétroaction négative de Black fut testé avec succès en décembre 1927 et, en 1932, Black et ses collègues construisaient des amplificateurs à rétroaction négative présentant d'excellentes performances. (Figure3)

Mais l'incompréhension fut grande devant ce succès car le principe de la rétroaction négative, réellement novateur, restait difficile à accepter. En effet, le principe de rétroaction était déjà largement utilisé dans les oscillateurs pour amorcer et entretenir des oscillations, et dans les dispositifs à rétroaction positive pour augmenter l'amplification, avec des risques bien connus d'instabilité. Aussi, le concept même de rétroaction, utilisable pour améliorer les performances, sans risque d'instabilité, choquait les esprits. A preuve, la reconnaissance de la demande de brevet déposée par Black pour son amplificateur prit plus de neuf années. L'office des brevets américains affirma que la sortie de l'amplificateur ne pouvait pas être connectée à l'entrée sans que cela provoque une instabilité. En Angleterre, l'office des

brevets, selon les termes employés par Black, traita son amplificateur de la même manière que les machines à mouvement perpétuel !

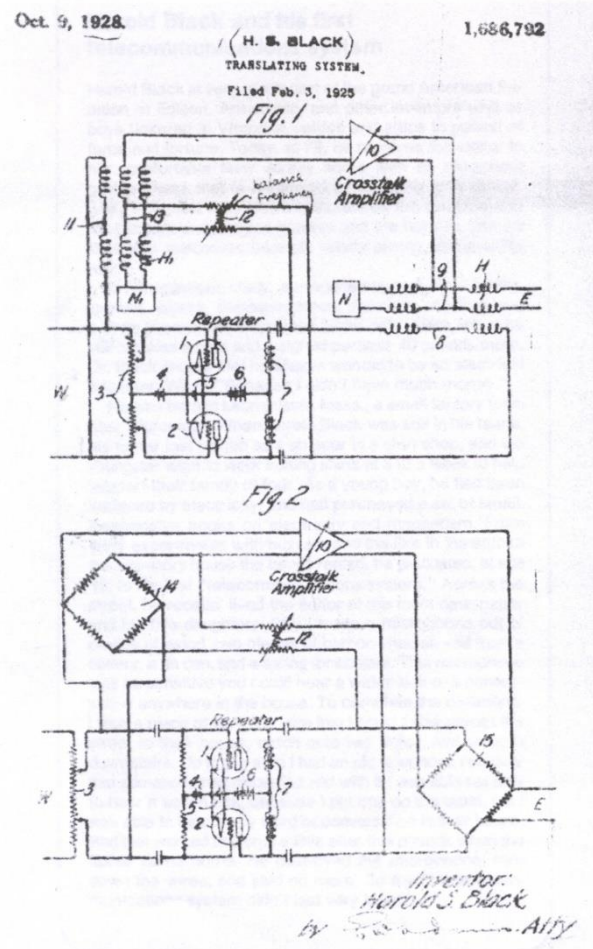


Figure 3 : Extrait de l'article 'Inventing the negative feedback amplifier'

V. LE CRITERE DE STABILITE DE HARRY NYQUIST

L'amplificateur à rétroaction négative construit par Black avait tout de même tendance à « chanter », expression des ingénieurs du téléphone signifiant que l'amplificateur devenait instable en se transformant en oscillateur. Avant 1932, la seule méthode d'étude de la stabilité disponible pour les ingénieurs électriciens était celle basée sur les équations différentielles et le critère de Routh. Son application aux amplificateurs à lampes triode aurait nécessité l'écriture d'un système de quelques dizaines d'équations différentielles, ce qui n'était pas envisageable.

C'est ainsi qu'Harry Nyquist, ingénieur possédant une solide culture mathématique, travaillant aux Bell Laboratories, fut sollicité pour résoudre ce problème. En 1932, il fit paraître dans *The Bell System Technical Journal* sa célèbre condition de stabilité dans un article intitulé *Regeneration theory* [11]. La publication débutait par une définition de la stabilité :

« The circuit will be said to be stable when an impressed small disturbance, which itself dies out, results in a response which dies out. It will be said to be instable when such a disturbance results in a response which goes on indefinitely, either staying at a relatively small value or increasing until it is limited by non-linearity of the amplifier. »

Il commença par cette définition dans le but d'éviter les confusions avec les notions utilisées dans les circuits oscillateurs. En effet, dans ces circuits, le terme de stabilité est employé pour préciser la constance de la fréquence de l'oscillateur. Les caractéristiques non-linéaires sont utilisées dans le principe même des oscillateurs pour limiter l'amplitude des oscillations. Il précisa même qu'il considérait « In the present discussion this difficulty will be avoided by the use of a strictly linear amplifier, which implies an amplifier of unlimited power carrying capacity. ».

La lecture de l'article de Nyquist s'avère difficile de nos jours. Tout d'abord, parce que nous nous attendons à y trouver une démonstration basée sur le théorème de Cauchy, telle que nous l'enseignons habituellement. Nyquist s'appuie bien évidemment sur la théorie des fonctions de la variable complexe, mais pas au sens où nous l'entendons. Ce malentendu, essentiellement actuel, résulte de l'interprétation qui prévalait à l'époque sur la naissance et l'entretien des oscillations. Ce que rappelle Nyquist dans un long préambule. Il considère un dispositif bouclé, constitué d'un amplificateur linéaire A et d'un réseau de transfert J . Un raisonnement habituel consistait à envisager un signal sinusoïdal F introduit dans la boucle et qui y circulait indéfiniment.

Après n tours de boucle, ce signal avait pour expression :

$$S_n = F(1 + AJ + A^2 J^2 + \dots + A^n J^n) e^{j\omega t}. \text{ Cette série de } n$$

termes ayant pour somme $F \left(\frac{1 - A^{n+1} J^{n+1}}{1 - AJ} \right)$ conduisait à

s'interroger sur sa convergence en fonction du module de AJ . La condition $AJ = 1$ correspondant à celle des oscillateurs, une solution apparemment évidente consistait à tester la stabilité grâce à la transmittance en boucle ouverte AJ . La somme précédente diverge si $|AJ| > 1$; cependant, des expériences paradoxales montraient que le système bouclé pouvait néanmoins rester stable. Nyquist a donc réussi à correctement formuler le problème de stabilité de la boucle en termes de convergence ou de divergence du régime transitoire $s(t)$ causé par une perturbation initiale de type impulsionnel avec

$$s(t) = \lim_{n \rightarrow \infty} \int_{-\infty}^{+\infty} S_n(z) e^{jzt} dz \text{ lorsque } n \rightarrow \infty, \text{ et non en}$$

régime sinusoïdal permanent, comme cela était envisagé habituellement. Si cette démonstration nous semble aujourd'hui curieuse et entourée d'un luxe de justifications mathématiques, c'est parce qu'elle avait pour objectif de convaincre la communauté de l'époque du bien fondé de son approche et de l'impasse de l'analyse habituelle. La suite de la démonstration s'appuie sur la définition d'un

contour d'exclusion dans le demi-plan complexe de droite et utilise le calcul des résidus.

Nyquist formula alors en ces termes une règle de stabilité dorénavant bien connue :

« Rule: Plot plus and minus the imaginary part of $AJ(\omega)$ against the real part for all frequencies from 0 to ∞ . If the point $1+i.0$ lies completely outside this curve the system is stable; if not it is instable. » .

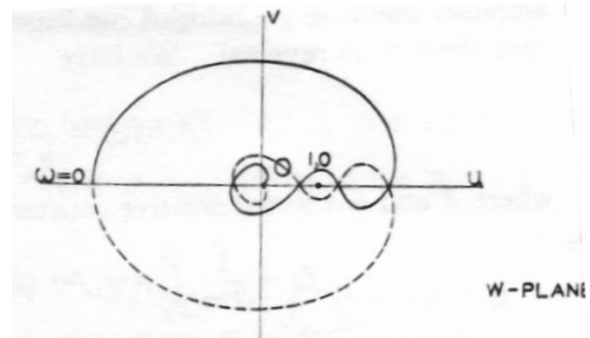


Figure 4 : Extrait de l'article Regeneration theory

L'immense intérêt pratique de la solution proposée par Nyquist est qu'elle s'exprime en fonction d'une grandeur, le gain complexe, qui est directement mesurable. Ainsi, l'application du critère de stabilité de Nyquist ne nécessite plus la connaissance d'un modèle du système sous forme d'équations différentielles. Par ailleurs, la forme du lieu de Nyquist donne un aperçu immédiat de la modification qu'il faut apporter au gain complexe en boucle ouverte pour améliorer les performances en stabilité du système.

Le caractère hermétique de l'article de Nyquist pourrait laisser penser, à tort, que ces travaux mathématiques étaient dépourvus de liens avec des préoccupations concrètes. Heureusement, un article de ses collègues des Laboratoires Bell, paru en 1934 [9], nous en apporte un démenti formel. Après une discussion rapide sur la technique de Nyquist et sur ses liens avec celle de Routh, les auteurs abordent le problème de la vérification expérimentale du critère, et plus particulièrement du caractère paradoxal de la stabilité conditionnelle. Le dispositif considéré est un amplificateur avec contre réaction où la transmittance en boucle ouverte nécessite une analyse harmonique précise de 0.5 kHz à 1.2 MHz ! Deux techniques étaient utilisées pour effectuer le tracé du lieu de Nyquist : la première, qualifiée d'approximative, était basée sur une visualisation sur écran cathodique (!) ; la seconde, précise, s'appuyait sur une technique de transposition de fréquence, pour éviter d'effectuer des mesures de module et de phase à fréquence élevée. Cet article nous prouve que la technique proposée par Nyquist est devenue immédiatement une réelle méthode d'analyse, puis de synthèse, au sein des Laboratoires Bell.

Cette technique d'analyse de la stabilité d'un système fut bien accueillie dans le domaine de l'amplification électronique, malgré son formalisme mathématique. Ainsi, en France, dès 1934, un article de Le Corbeiller relate les travaux de Nyquist dans les Annales des PTT [8].



Cependant, sa diffusion dans les autres domaines de l'automatique fut beaucoup plus lente, car contraire aux habitudes des praticiens. En France, à partir de 1945, de nombreux cours, des conférences, furent organisés à l'intention d'ingénieurs de toutes les disciplines pour leur faire découvrir les asservissements et les nouvelles méthodes d'analyse et de synthèse basées sur l'approche fréquentielle.

Ce concept a dû constituer une révolution pour beaucoup de thermiciens, de mécaniciens, voire d'électriciens. Remarquons qu'à l'inverse, l'approche fréquentielle, et la notion de boîte noire (et de système) se sont imposées naturellement aux électroniciens.

En effet, les oscillateurs ont précédé les montages à contre réaction, qui eux-mêmes avaient une fâcheuse tendance à se transformer en oscillateurs. Les électroniciens étant continuellement confrontés aux oscillations et aux signaux sinusoïdaux, ceci les a conduits à raisonner en termes d'approche fréquentielle. C'est au contraire la notion de régime transitoire qui pouvait leur paraître étrangère, à la différence des mécaniciens travaillant sur le réglage des régulateurs de vitesse. Il peut donc apparaître comme une conséquence logique qu'un électronicien ait proposé un critère basé sur une approche fréquentielle, mais en utilisant cette approche pour prévoir indirectement l'apparition d'un régime transitoire, stable ou instable.

VI. LES TRAVAUX DE HEINDRIK BODE

A la suite des travaux de Black et de Nyquist, il devint clair que la clef pour réaliser un système à rétroaction stable avec le meilleur amortissement consistait dans une modification judicieuse des caractéristiques d'amplitude et de phase de la fonction de transfert de la boucle ouverte au voisinage du point critique. Les études furent entreprises, particulièrement au sein du Bell Telephone Laboratories, sur la réalisation d'un amplificateur à rétroaction idéal qui présenterait une coupure rapide en amplitude avec une phase associée très faible. L'analyse de la relation entre l'amplitude et la phase d'une fonction de transfert en boucle ouverte devint donc à l'époque un enjeu essentiel. Ce sont les travaux de Heinrich Bode qui définirent clairement les relations entre module et phase et démontrèrent que la quête de l'amplificateur idéal était illusoire.

Les recherches de Heinrich Bode sur les circuits de rétroaction commencèrent en 1934. Dans la publication *Relations between attenuation and phase in feedback amplifier design* [6] de 1940 dans *The Bell System Technical Journal*, Bode montra, en s'appuyant sur la théorie des fonctions analytiques, qu'il existait une relation entre l'amplitude et la phase pour une fréquence donnée. Il introduisit pour cela le concept de circuit à minimum de phase.

Sa première relation $\int_{-\infty}^{+\infty} B \cdot du = \frac{\pi}{2} [A_{\infty} - A_0]$ pose que la

surface totale sous la courbe de phase tracée en fonction de la fréquence dépend seulement de la différence entre les

amplitudes aux fréquences nulle et infinie et non de l'allure de la courbe amplitude-fréquence entre ces limites.

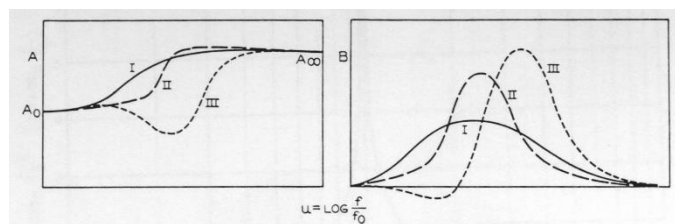


Figure 5 : Extrait de l'article 'Relations between attenuation and phase in feedback amplifier design'

Sa deuxième relation $B(f_c) = \frac{1}{\pi} \cdot \int_{-\infty}^{+\infty} \frac{dA}{du} \cdot \log \coth \frac{|u|}{2} \cdot du$

montre que la phase $B(f_c)$ est proportionnelle à la dérivée de l'amplitude (A représente le logarithme de l'amplitude) exprimée en fonction d'une échelle logarithmique de fréquence. La phase, pour une fréquence donnée, dépend aussi de la dérivée à toutes les fréquences (intégrale sur toute les fréquences), la fonction cotangente ayant pour objectif de réaliser une pondération fréquentielle.

La conclusion est que la dérivée à toutes les fréquences participe à la phase mais que la dérivée autour du point $f = f_c$ est prépondérante. Si on suppose $A = k \cdot u$, avec une pente $6k$ dB par octave, la phase associée est alors constante et a pour valeur $\frac{k \cdot \pi}{2}$ radians, résultat qui est devenu pour nous une évidence.

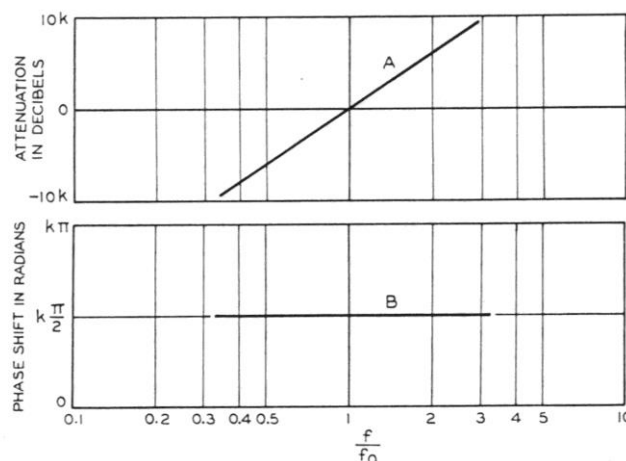


Figure 6 : Extrait de l'article 'Relations between attenuation and phase in feedback amplifier design'

Bode développa les techniques de construction asymptotique pour tracer les diagrammes utilisés dans la conception des systèmes à rétroaction. Il introduisit pour cela le logarithme du gain complexe (qui permet de séparer le gain de la phase dans l'écriture complexe $\ln[H(\omega)] = \ln[\rho(\omega)] + j\theta(\omega)$), et l'échelle logarithmique pour la fréquence. C'est la raison pour laquelle les courbes du gain en dB et de la phase en fonction du logarithme de la fréquence sont connues dorénavant sous le nom de courbes de Bode.

Bode relia son travail à celui de Nyquist. Il fut responsable de la rotation du diagramme de Nyquist de 180° ; le point critique apparaissant dans la théorie de la stabilité de Nyquist devint donc le point $-1 + j.0$. Il affirma que la rotation du diagramme permettait d'ignorer l'inversion de phase due aux amplificateurs à triode.

Il introduisit les marges de phase et de gain ; à ce titre, il peut être considéré comme l'inventeur de la robustesse. Ainsi, dans son article de 1940, il décrit une technique de synthèse du circuit de rétroaction permettant de satisfaire à un gabarit fréquentiel original (module décroissant à phase constante) afin de se prémunir contre les variations de gain dues à la disparité des caractéristiques des lampes triode et de leurs variations dans le temps.

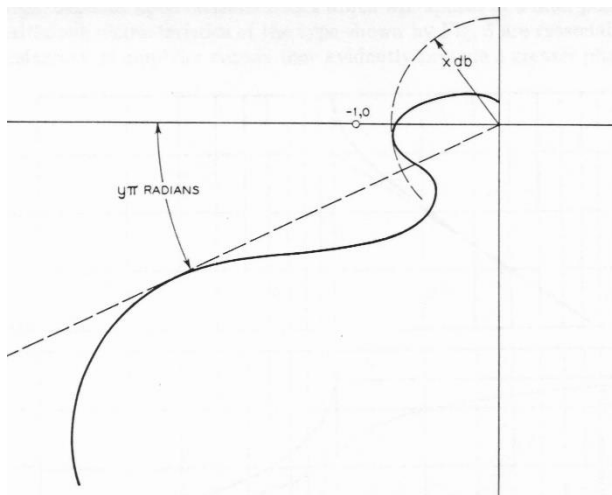


Figure 7 : Extrait de l'article 'Relations between attenuation and phase in feedback amplifier design'

Le travail de Bode fut exposé sous une forme développée dans son ouvrage de 1945 *Network analysis and feedback amplifier design* [7] qui permit la diffusion des concepts élaborés au sein des Bell Laboratories.

VII. CONCLUSION

En conclusion, Black, Nyquist et Bode, bien que de formations différentes, étaient à l'époque de véritables électroniciens. Ils n'en ont pas moins inventé les concepts de l'automatique et posé les bases de ses prolongements avancés en robustesse.

On leur doit la formalisation des propriétés de la boucle fermée, de la stabilité et surtout une méthode d'analyse par

approche fréquentielle, naturelle en électronique, mais révolutionnaire dans les autres domaines : mécanique, thermique, aéronautique, etc.

Leurs recherches en ce début du XX^e siècle ont stimulé l'apparition d'une nouvelle science, l'automatique, sa généralisation à tous les domaines de la technologie, et la création d'une nouvelle communauté d'ingénieurs et de chercheurs spécialisés en automatique.

VIII. REFERENCES

- [1] Erwin Armstrong, *Operating features of the audion : explanation of its action as an amplifier, as a detector of high-frequency oscillations and as a « valve »*, Electrical World (New York), 64, p. 1149, 12th December 1914
- [2] Stuart Bennett, *A history of control engineering 1800-1930*, Peter Peregrinus Ltd, London, UK, 1979
- [3] Stuart Bennett, *A history of control engineering 1930-1959*, Peter Peregrinus Ltd, London, UK, 1993
- [4] H. S.Black, *Stabilised feedback amplifiers*, The Bell System Technical Journal, January, 1934
- [5] H. S.Black, *Inventing the negative feedback amplifiers*, IEEE Spectrum, 14, 1977
- [6] H. W.Bode, *Relations between attenuation and phase in feedback amplifier design*, 1940
- [7] H. W.Bode, *Network analysis and feedback amplifier design*, D. Van Nostrand New-York, 1945
- [8] Le Corbeiller, *Etude de la stabilité d'un réseau à réaction, d'après H. Nyquist*, Annales des PTT, 1934, XI
- [9] E. Peterson, J.G. Kreer, L.A. Ware, *Regeneration theory and experiment*, The Bell System Technical Journal, 13, 1934
- [10] Harry Nyquist, *Certain factors affecting telegraph speed*, The Bell System Technical Journal, 3, 1924
- [11] Harry Nyquist, *Regeneration theory*, The Bell System Technical Journal, 11, 1932
- [12] A. Pagès, *Description d'un amplificateur basse fréquence à grande sélection*, Onde Electrique, Juin 1926, pp 276-283
- [13] D.G. Tucker, *The history of positive feedback : the oscillating audion, the regenerative receiver, and other applications up to around 1923*, The Radio and Electronic Engineer, Vol 42, N°2, February 1972

Conformément à la Loi Informatique et Libertés du 06/01/1978, vous disposez d'un droit d'accès et de rectification aux informations qui vous concernent. Contacter le Service Abonnements de la SEF

- Décembre 2006 (n° 47)
Les matériaux électro actifs (suite)
- Mars 2007 (n° 48)
Stockage de l'énergie électrique
- Juin 2007 (n° 49)
Gisements d'économies d'énergies
- Septembre 2007 (n° 50)
Pratiques pédagogiques et réalités industrielles

- **Support de cours**
Outil didactique pour les filières préparant à l'enseignement technique
- **Support documentaire**
Trame indispensable à la formation continue des Hommes de terrain



Les sommaires des derniers numéros
sont disponibles sur le site www.see.asso.fr

A retourner à la SEE - La Revue 3 E.I., 17 rue de l'Amiral Hamelin - 75783 Paris cedex 16 - France

Tarifs d'abonnements	France 1 an (4 numéros)	Etranger 1 an (4 numéros)	Prix de vente au numéro (Tarif France)
Tarif membre SEE *	32 €	42 €	* 1 ex. 11 €, 2 ex. 22 €, 3 ex. 27 €, 4 ex. 36 € <i>Liste complète des N° déjà parus disponible sur demande auprès du Service Adhésions : adhesions@see.asso.fr (Tél. : 01 56 90 37 09 - Fax : 01 56 90 37 19)</i>
Plein tarif (non membre SEE)	35 €	45 €	
Tarif collectif membre SEE *	47 €	60 €	Je règle la somme de : € Indiquer votre N° de membre (le cas échéant) : par : ☐ chèque à l'ordre de la SEE ☐ Carte bancaire (Visa, Eurocard, American Express)
Plein tarif collectif (non membre SEE) Bibliothèques, CDI, Laboratoires, Entreprises)	50 €	63 €	

**Pour devenir membre SEE, appeler le 01 56 90 37 09 ou adresser un courriel à : adhesion@see.asso.fr*

Oui, je m'abonne à la revue 3 E.I. pour 4 n^{os} (47 à 50 inclus)

Nom et prénom (ou raison sociale) :

Service/ département :

Activité (facultatif):

Adresse :

Code postal Ville :

Pays : e-mail souhaitable:

Conformément à la Loi Informatique et Libertés du 06/01/1978, vous disposez d'un droit d'accès et de rectification aux informations qui vous concernent. Contacter le Service Abonnements de la SEF

Je règle la somme de : €
Indiquer votre N° de membre (le cas échéant) :
par : ☐ chèque à l'ordre de la SEE
☐ Carte bancaire (Visa, Eurocard, American Express)
N° Carte | | | | | | | | | | | | | | Date de validité | | | | |
Date de souscription | | | | | |

Date, signature et cachet :

☐ Je joins le bon de commande administratif numéro :
et je désire recevoir une facture au nom de mon employeur pour paiement à réception

Raison sociale et adresse :

.....

Code postal Ville :

Pays : VAT :

MACHINE TO MACHINE²

ème / nd
édition

Salon des Solutions **MtoM** solutions Show

Exposition et conférences
Trade show and conferences

6, 7 & 8 mars/March 2007
CNIT Paris La Défense

Exposer, éditer son badge, s'informer sur le Salon
Exhibit, obtain a badge, learn more about the show
www.birp.com/mtom

APPLICATIONS

- Gestion de flotte • Gestion de la chaîne d'approvisionnement • Sécurité
- Domotique, bâtiment intelligent • Gestion routière • Télépaiement • Santé • ...

- Fleet management • Supply chain management • Security • Building intelligence
- Traffic flow • Retailing • Health • ...

OUTILS ET SERVICES / TOOLS AND SERVICES

- WiFi • IPv6 • Bluetooth • Wireless
- RFID • Edge • 3G • ADSL • GPRS
- GSM • Zigbee • ...

GOLD SPONSOR



97, rue du Cherche Midi - 75006 PARIS
Tél : 33 (0)1 44 78 99 30 Fax : 33 (0)1 44 78 99 49
mtom@birp.fr

Strictement réservé aux professionnels / Strictly reserved for professionals

RF & HYPER

33^e ÉDITION

EUROPE 2007

LE SALON DES RADIOFRÉQUENCES, DES HYPERFRÉQUENCES,
DU WIRELESS, DE LA FIBRE OPTIQUE ET DE LEURS APPLICATIONS

27, 28 & 29 MARS 2007

CNIT - PARIS LA DÉFENSE

AU CŒUR DES MARCHÉS ÉMERGENTS

Venez découvrir sur RF & Hyper Europe les dernières évolutions technologiques présentées par plus de 150 exposants experts des radiofréquences, hyperfréquences, et de la fibre optique.

Vous trouverez entre autres les dernières nouveautés en composants actifs et passifs, modules, systèmes, logiciels de simulation ou de conception, instrumentation et test pour les applications dans les télécommunications (Wifi, Wimax, 3G, Bluetooth, UWB, ZigBee, ...), les liaisons satellites, l'avionique, le militaire, la sécurité, et les nouveaux développements dans les systèmes RFID.

ÉVÉNEMENTS 2007

Les ANTENNES en vedette sur le salon !

2 journées de conférences
incontournables sur la CEM.

Après le succès remporté en 2006 :
nouvelle édition du séminaire
RF TECHNOLOGIES AND PACKAGING.

Pour exposer, demander son badge,
s'inscrire aux conférences :

www.RFHyper.com

EXPOSIUM

1, rue du Parc - 92593 Levallois-Perret Cedex, France

Tél.: +33 (0)1 49 68 51 00 - Fax : +33 (0)1 49 68 54 18 - E-mail : RFHyper@exposium.fr

RF & Hyper Europe, un salon
organisé par EXPOSIUM

www.exposium.fr

