

Chapitre C-V

Eléments d'électrocinétique.

Joël SORNETTE met ce cours à votre disposition selon les termes de la licence Creative Commons :

- Pas d'utilisation commerciale.
- Pas de modification, pas de coupure, pas d'intégration à un autre travail.
- Pas de communication à autrui sans citer son nom, ni en suggérant son autorisation.

Retrouvez l'intégralité du cours sur le site joelsornette.fr

RÉSUMÉ :

On commence par justifier la loi d'Ohm par deux modèles de conduction puis on en explore les conditions de validité. On en déduit la loi d'Ohm intégrale et l'on précise la définition des forces ou mieux tensions électromotrices. On présente quelques exemples de conducteurs non ohmiques.

On s'intéresse ensuite aux aspects thermodynamiques : effet Joule, effets électro-chimiques et thermo-électriques.

On laisse de côté l'étude pratique des réseaux électriques mais on justifie théoriquement tous les théorèmes utiles à celle-ci.

On montre comment les notations complexes et la transformation de Laplace permettent de traiter respectivement les régimes sinusoïdaux et les régimes transitoires avec le même formalisme que les régimes permanents.

Table des matières

C-V	Eléments d'électrocinétique.	1
1	Loi d'Ohm locale.	5
1.a	Un premier modèle de conduction.	5
1.b	Le modèle de Drude.	6
1.c	Densité de courant.	7
1.d	Synthèse : la loi d'Ohm locale.	8
1.e	Ordres de grandeur.	9
1.f	Conditions de validité.	10
1.g	Influence de la température. Conducteurs et semi-conducteurs. . .	11
2	Loi d'Ohm intégrale.	14
2.a	Champ électrostatique et champ électromoteur.	14
2.b	Résistance et tension électromotrice. Loi d'Ohm intégrale.	15
2.c	Premières indications sur l'induction magnétique.	18
2.d	Calcul des résistances mortes.	19
3	Les complications.	22
3.a	L'effet Hall.	22
3.b	Magnéto-résistance.	23
3.c	Diode à vide.	25
3.d	Diodes à semi-conducteur.	27
4	Aspects thermodynamiques.	29
4.a	Puissance fournie un dipôle. Effet Joule.	29
4.b	Piles électrochimiques.	31

4.c Effets Peltier et Thomson. Thermocouples.	32
5 Réseaux électriques en courant continu.	35
5.a Branches, nœuds, mailles.	35
5.b Modélisation d'une branche. Conventions de signe et conventions graphiques.	35
5.c Modélisation des générateurs.	36
5.d Les outils de base : loi des nœuds, loi des mailles, théorème de Millmann	36
5.e Association de résistances : série, parallèle, triangle-étoile.	38
5.f Résolution matricielle par les courants de mailles.	39
5.g Théorème de superposition.	41
5.h Théorèmes de Thévenin et de Norton. Exemples et applications.	42
5.i Cas des dipôles commandés.	47
5.j Théorème de compensation.	47
5.k Présentation en termes de quadripôle.	48
6 Réseaux électriques en régime variable.	50
6.a Indications sur les régimes quasi-stationnaires.	50
6.b Les nouveaux personnages : condensateurs et bobines.	50
6.c Régime sinusoïdal. Notation complexe. Impédances.	51
6.d Puissance en régime sinusoïdal.	51
6.e Régime transitoire. Transformation de Laplace.	53

1 Loi d'Ohm locale.

1.a Un premier modèle de conduction.

Considérons une charge libre, c'est à dire une particule chargée à l'échelle atomique ou moléculaire, en pratique un électron ou un ion, de masse m , de charge q , parmi beaucoup d'autres de son espèce, en mouvement relatif à vitesse relative $\vec{v}$ au sein d'un milieu quelconque en mouvement ou non. On se place dans le référentiel lié au milieu et l'on note $\vec{F}$ la force qu'elle subirait au point M où elle se trouve si elle avait une vitesse nulle, y compris la force d'inertie d'entraînement si le milieu est en mouvement, et notons de plus $\vec{E}(M, t) = \frac{\vec{F}}{q}$, grandeur homogène à un champ électrique, dépendant a priori de l'espace et du temps. Du fait de sa vitesse, la force exercée sur la charge diffère a priori, notons-la donc $q \vec{E}(M) + \vec{f}(\vec{v})$.

Cette force supplémentaire globale contient éventuellement la force d'inertie complémentaire ou force de CORIOLIS $-2m \vec{\Omega} \wedge \vec{v}$ ou une force magnétique $q \vec{v} \wedge \vec{B}$; mais la suite montrera que les vitesses en jeu rendent ces deux forces négligeables. Quoi d'autre? L'expérience montre qu'il y a des phénomènes dissipatifs. Dans une première approche, plutôt que de chercher à en comprendre la nature, on peut se contenter d'introduire un terme phénoménologique, c'est-à-dire une loi de force dont on ne cherche pas à justifier la validité mais que l'on choisit parce qu'elle conduit notoirement aux mêmes effets que ceux que l'on cherche à modéliser. En l'occurrence, on propose ici une force de type frottement fluide notée $\vec{f} = -\lambda \vec{v}$. L'équation du mouvement d'une charge libre est donc, successivement :

$$m \frac{d\vec{v}}{dt} = q \vec{E}(M, t) - \lambda \vec{v}$$

$$m \frac{d\vec{v}}{dt} + \lambda \vec{v} = q \vec{E}(M, t)$$

Ne nous emballons pas, ce n'est pas une équation en $\vec{v}$, mais en $\vec{r} = \overrightarrow{OM}$ avec $\vec{v} = \frac{d\vec{r}}{dt}$ et, si vous avez déjà lu le chapitre B-XIII de mécanique des fluides, le premier membre est d'approche lagrangienne axée sur la particule et le second d'approche eulérienne axée sur la notion de champ, ce qui ne simplifie pas les choses. Dans la pratique sur une échelle de temps assez courte pour que, d'une part, $\vec{E}$ n'ait pas le temps de varier de façon notable et, d'autre part, le déplacement soit assez court pour que les variations de $\vec{E}$ avec la position soient elle aussi négligeables, on peut considérer $\vec{E}$ comme constant. L'équation se réécrit ainsi :

$$m \frac{d\vec{v}}{dt} + \lambda \vec{v} = q \vec{E}$$

Et sa solution est classiquement, en notant $\tau = \frac{m}{\lambda}$:

$$\vec{v}(t) = \frac{q}{\lambda} \vec{E} + \left(\vec{v}(0) - \frac{q}{\lambda} \vec{E} \right) e^{-\frac{t}{\tau}}$$

Au delà d'un temps de l'ordre de 7τ (c'est cette durée qui doit être très courte pour valider notre démarche, nous en discuterons un peu plus loin) et quelle que soit la vitesse initiale, la vitesse limite est atteinte et l'on considérera que c'est la vitesse de la charge libre qui est donc $\vec{v}_{\text{lim}} = \frac{q}{\lambda} \vec{E}$ que l'on note $\vec{v} = \mu \vec{E}$ en escamotant l'indice « limite » et où $\mu = \frac{q}{\lambda}$ s'appelle *mobilité de la charge libre*.

Si le vecteur $\vec{E}$ varie avec le temps mais avec un temps caractéristique très grand devant 7τ , on peut considérer que cette relation est vraie à tout instant (là encore, nous en discuterons plus loin).

Une remarque pour finir : si l'on tient compte des forces en $-2m \vec{\Omega} \wedge \vec{v}$ ou $q \vec{v} \wedge \vec{B}$; il sera plus dur mais tout à fait possible de montrer l'existence d'une vitesse limite. Lorsque celle-ci est atteinte, alors $\frac{d\vec{v}}{dt} = \vec{0}$, on aura encore $\vec{v} = \mu \vec{E}$ en intégrant à $\vec{E}$ les deux nouvelles forces mais il faudra lire cette relation comme une équation en $\vec{v}$ à résoudre. On en donnera un exemple un peu plus loin (voir la magnéto-résistance au paragraphe 3.b p. 23).

1.b Le modèle de Drude.

Un modèle plus fin fut proposé en 1900 par Paul DRUDE¹. La force phénoménologique est supprimée et l'on suppose que la charge mobile subit de façon aléatoire des chocs avec les atomes du milieu².

Plaçons à un instant t arbitraire et choisissons une charge libre d'indice i . Appelons t_{i0} l'instant, antérieur à t , où elle a subi le dernier choc et $\vec{v}_{i0}$ sa vitesse juste après ce choc. Par intégration entre t_{i0} et t de $m \frac{d\vec{v}}{dt} = q \vec{E}$ avec les mêmes hypothèses permettant de considérer $\vec{E}$ comme uniforme et stationnaire, on arrive pour sa vitesse $\vec{v}_i(t)$ à l'instant t à :

$$\vec{v}_i(t) = \vec{v}_{i0} + \frac{q}{m} \vec{E} \cdot (t - t_{i0})$$

Effectuons la moyenne $\langle \vec{v} \rangle$ de la vitesse à l'instant t sur un grand nombre de particules du résultat. La moyenne étant une opération linéaire, on a :

$$\langle \vec{v} \rangle = \langle \vec{v}_i(t) \rangle = \langle \vec{v}_{i0} \rangle + \frac{q}{m} \vec{E} \cdot \langle t - t_{i0} \rangle$$

Le terme $\langle \vec{v}_{i0} \rangle$ est la moyenne des vitesses juste après un choc qui renvoie les charges dans toutes les directions possibles de façon aléatoire ; on peut raisonnablement considérer la moyenne comme nulle ou en tous cas négligeable.

Le facteur $\langle t - t_{i0} \rangle$ est la moyenne de l'« âge » de la charge depuis son dernier choc (âge remis à zéro à chaque nouveau choc ; on peut penser qu'il s'établit un régime permanent

1. physicien allemand, on prononce donc Droudeux avec un accent tonique sur le « ou ».

2. dans la version originale ; dans les métaux, on parle plutôt maintenant de choc avec les *phonons*, en gros les modes propres de vibration du cristal métallique.

entre les charges, le lent vieillissement de presque toutes étant compensé par le brusque retour à l'état de bébés des vieillards après un choc, ce qui confère à cet âge moyen une valeur indépendante du temps que l'on note τ . Dès lors, on peut affirmer que :

$$\langle \vec{v} \rangle = \frac{q}{m} \vec{E} \tau$$

On retrouve la notion de mobilité μ avec ici $\mu = \frac{q\tau}{m}$ et l'on identifie les deux modèles en prenant comme coefficient de la force de frottement fluide $\lambda = \frac{m}{\tau}$.

1.c Densité de courant.

Imaginons une surface élémentaire de vecteur surface³ $\vec{d_2S}$ fixe par rapport au milieu où se meuvent les charges libres. On se propose de calculer le nombre $\delta_3 N$ de charges libres qui la traversent pendant une durée élémentaire dt puis la charge⁴ $\delta_3 Q$ correspondante.

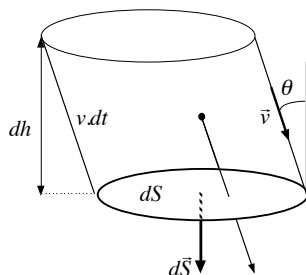


FIGURE 1 – Densité de courant.

La figure 1 p. 7 localise les charges à l'instant initial ; elles sont dans un cylindre de base d_2S et de génératrice, pas forcément orthogonale, $\vec{v} dt$ où la vitesse des charges de la forme $\mu \vec{E}$ est pratiquement une constante si d_2S et dt sont suffisamment petits pour que $\vec{E}$ puisse être considéré comme homogène dans le cylindre et stationnaire pendant la durée dt .

Le volume de ce cylindre est, en notant d_2S et v les modules (les normes) de $\vec{d_2S}$ et $\vec{v}$ et θ l'angle entre ces deux derniers vecteurs,

$$d_3\mathcal{V} = d_2S dh = d_2S v dt \cos \theta = \vec{v} \cdot \vec{d_2S} dt$$

Si l'on appelle n la *densité particulière* de charges libres de charge q , c'est à dire leur nombre par unité de volume, on a successivement :

$$\delta_3 N = n d_3\mathcal{V} = n \vec{v} \cdot \vec{d_2S} dt$$

3. par définition orthogonal à la surface dans un sens arbitraire qui sert de sens positif pour les charges qui le traversent et dont le module (la norme) soit l'aire de la surface.

4. La notation $\delta_2 S$ pour indiquer que c'est d'ordre deux (largeur élémentaire et longueur élémentaire) et puisque dt est d'ordre un, $\delta_3 N$ est d'ordre trois d'où la notation.

$$\delta_3 Q = q \delta_3 N = n q \vec{v} \cdot \overrightarrow{d_2 S} dt$$

On appelle *intensité élémentaire* $d_2 I$ traversant la surface $d_2 S$, orientée par son vecteur surface le rapport de la charge $\delta_3 Q$ qui la traverse à la durée dt de cette traversée ; on a donc :

$$d_2 I = \frac{\delta_3 Q}{dt} = n q \vec{v} \cdot \overrightarrow{d_2 S}$$

Dans le cas d'une surface quelconque, on la découpe en surfaces élémentaires et la charge qui la traverse par unité de temps s'obtient par addition des contributions des surfaces élémentaires ; l'*intensité* est donc :

$$I = \iint d_2 I = \iint n q \vec{v} \cdot \overrightarrow{d_2 S}$$

On a sous-entendu jusqu'ici qu'il n'y avait qu'un seul type de charges libres ; s'il y en a plusieurs, on procède par sommation sur les différents types, repérés ici par un indice i , soit

$$I = \iint d_2 I = \iint \left(\sum_i n_i q_i \vec{v}_i \right) \cdot \overrightarrow{d_2 S}$$

On appelle *densité de courant* (abrégé de densité surfacique de courant électrique) le vecteur, traditionnellement noté $\vec{j}$ défini par $\vec{j} = \sum_i n_i q_i \vec{v}_i$ est dont le flux à travers une surface est l'intensité qui la traverse puisque $I = \iint \vec{j} \cdot \overrightarrow{d_2 S}$.

1.d Synthèse : la loi d'Ohm locale.

Le raisonnement qui introduit le vecteur densité de courant est récurrent en physique et se retrouve dans maint domaine (thermodynamique, mécanique des fluides, etc) et n'est donc pas une loi d'électrocinétique. Par contre la proportionnalité $\vec{v} = \mu \vec{E}$ (ou encore $\vec{v}_i = \mu_i \vec{E}$ s'il y a plusieurs types⁵ de porteurs) en est une. La synthèse avec la définition $\vec{j} = \sum_i n_i q_i \vec{v}_i$ conduit à :

$$\vec{j} = \sum_i n_i q_i \vec{v}_i = \sum_i \left(n_i q_i \mu_i \vec{E} \right) = \left(\sum_i n_i q_i \mu_i \right) \vec{E}$$

Cette formule prouve la proportionnalité entre la densité de courant et le vecteur $\vec{E}$ (patience, nous creuserons plus loin sa signification au paragraphe 2.a p. 14, pour l'instant c'est inutile⁶ donc prématuré.). La constante de proportionnalité est appelée *conductivité* du milieu et est traditionnellement notée soit γ , soit σ .

5. A ce stade de l'exposé ($\vec{E}$ n'est pas encore assimilé au champ électrique), rien ne prouve que $\vec{E}$ ait le même signification pour toutes les espèces, mais c'est fréquent ; nous y reviendrons plus loin

6. Aphorisme personnel : ce qui est inutile est potentiellement nuisible.

La loi d'OHM locale est donc $\vec{j} = \gamma \vec{E}$ ou $\vec{j} = \sigma \vec{E}$ avec γ ou $\sigma = \sum_i n_i q_i \mu_i$.

Remarque : la mesure de la conductivité d'une solution ionique donne accès aux densités particulières n_i qui ne sont autres dans ce contexte que les concentrations. Nos amis chimistes sont friands de ce genre de mesure des concentrations par un procédé physique qui ne perturbe pas, comme le ferait un dosage qui consomme l'un des protagonistes, l'état d'équilibre d'une réaction réversible, ni la vitesse d'une réaction lente.

1.e Ordres de grandeur.

• Le cuivre, exemple de métal.

Le cuivre a une masse molaire $M = 64 \cdot 10^{-3} \text{ kg} \cdot \text{mol}^{-1}$ et une masse volumique $\rho = 9,0 \cdot 10^3 \text{ kg} \cdot \text{m}^{-3}$; en utilisant le nombre d'Avogadro $\mathcal{N}_A = 6,0 \cdot 10^{23} \text{ mol}^{-1}$, on en déduit sa densité particulière n qui est aussi celle des électrons libres car un atome de cuivre donne un ion Cu^+ dans le réseau cristallin et un électron libre, soit :

$$n = \mathcal{N}_A \frac{\rho}{M} = 8,5 \cdot 10^{28} \sim 10^{29} \text{ m}^{-3}$$

La conductivité est donnée par $\sigma = n q \mu$ avec, dans le modèle de DRUDE, le plus fin des deux, or $\mu = \frac{q\tau}{m}$ d'où $\sigma = \frac{n q^2 \tau}{m}$. L'expérience donne une conductivité $\sigma = 6,0 \cdot 10^7 \text{ S m}^{-1}$; avec ici $q = -e = -1,6 \cdot 10^{-19} \text{ C}$ et $m = 9,0 \cdot 10^{-31} \text{ kg}$, on en tire la valeur du paramètre τ , soit :

$$\tau = \frac{m \sigma}{n q^2} = 2,5 \cdot 10^{-14} \text{ s}$$

qui doit être lu comme l'ordre de grandeur du temps d'établissement d'un régime permanent; suffisamment court pour considérer que ce régime s'installe de façon instantanée.

On en déduit aussi une valeur de la mobilité de l'électron libre :

$$\mu = \frac{q \tau}{m} = \frac{\sigma}{n q} = 4,4 \cdot 10^{-3} \text{ m}^2 \text{ s}^{-1} \text{ V}^{-1}$$

Electricité de France considère qu'un fil de section $1,5 \text{ mm}^2$ peut supporter 16 ampères et un fil de $2,5 \text{ mm}^2$ peut supporter 20 ampères; retenons que la densité de courant est limité à quelque chose comme $j_{\text{max.}} = 10 \text{ A} \cdot \text{mm}^{-2} = 1 \cdot 10^7 \text{ A} \cdot \text{m}^{-2}$. Avec $j = n q v$, sans se soucier du sens qui fixe le sens du courant, on en déduit avec les données précédentes, la vitesse maximale admissible des électrons :

$$v_{\text{max.}} = \frac{j_{\text{max.}}}{n q} \sim 10^{-3} \text{ m} \cdot \text{s}^{-1} = 1 \text{ mm} \cdot \text{s}^{-1} \sim 4 \text{ m} \cdot \text{h}^{-1}$$

très faible à l'échelle macroscopique.

Remarque 1 : on peut raisonnablement penser que la distance parcourue entre deux chocs de l'ordre est de l'ordre de $v \tau$ soit au maximum de l'ordre de 10^{-17} m , du dix-millionième de la distance interatomique; ce qui n'est guère satisfaisant pour un modèle

basé sur les chocs entre électrons et ions. Le modèle de DRUDE est donc naïf ; comme il donne cependant de bons résultats, il contient une part de vérité. Celle-ci est quantique : les phonons, modes de vibration de l'édifice cristallin des ions ont une énergie quantifiée et τ représente le temps moyen écoulé depuis le dernier échange d'un quantum d'énergie entre l'électron libre et un phonon.

Remarque 2 : la faible valeur de $v_{\max}$ permet de négliger les forces de CORIOLIS (où de plus $\Omega \sim 10^{-4} \text{ rad} \cdot \text{s}^{-1}$) et même les forces en $q \vec{v} \wedge \vec{B}$ car $\vec{v} \wedge \vec{B}$, avec un champ intense de un telsa est homogène à un champ de $10^{-3} \text{ V} \cdot \text{m}^{-1}$ à comparer avec le champ donné par la loi d'OHM $E = \frac{j}{\sigma} \sim \frac{10^7}{6 \cdot 10^7} \sim 150 \cdot 10^{-3} \text{ V} \cdot \text{m}^{-1}$. L'influence des phénomènes magnétique sera étudiée plus en détail dans le chapitre C-VII sur l'induction.

• Les solutions aqueuses.

Par rapport à la conduction métallique, il y a plusieurs différences. La première est qualitative : l'existence d'au moins deux types de charges libres, une négative au moins et une positive au moins. On aura donc addition des contributions de chaque type d'ion.

Une deuxième est quantitative : la masse des porteurs, celle de leur(s) noyau(x) est de l'ordre de quelques dizaines de fois celle d'un nucléon, elle-même près de deux mille fois celle d'un électron et la mobilité en $\frac{1}{m}$ s'en trouve diminuée d'un facteur de l'ordre de 10^5 à 10^6 . L'expérience est en totale contradiction, on trouve par exemple pour l'ion Na^+ 42 000 fois plus lourd, une mobilité de $52 \text{ m}^2 \text{s}^{-1} \text{V}^{-1}$ soit 12 000 plus grande.

C'est dû à une autre différence : le mécanisme d'interaction avec le milieu est plutôt un mouvement avec frottement fluide des ions dans le solvant avec des temps τ beaucoup plus élevés. Le premier modèle est plus adapté.

Une dernière est elle aussi quantitative, une solution décimolaire, déjà très concentrée, donne des densités particulières en 10^{23} m^{-3} donc un million de fois plus petite que pour un métal.

Bref, c'est la même chose, sauf que tout est différent, si je puis dire. Pour donner une idée au lecteur, une solution décimolaire de chlorure de sodium est une centaine de fois moins conductrice que le cuivre ; l'écart n'est pas considérable.

1.f Conditions de validité.

Quelque soit le modèle choisi, on a supposé que $\vec{E}$ était constant pendant un temps de l'ordre de τ ; dans le cas fréquent d'un régime périodique de période T et de fréquence f , il faudra donc $T \gg \tau$ soit encore $f \ll \frac{1}{\tau}$. Si l'on veut présenter le résultat avec la longueur d'onde $\lambda = cT$ qu'aurait dans le vide une onde électromagnétique de même fréquence, la condition est $\lambda \gg c\tau$.

Dans le cas le plus courant d'une conduction métallique avec, pour l'exemple du cuivre, $\tau = 2,5 \cdot 10^{-14} \text{ s}$ et en prenant un facteur 100 pour traduire « très grand devant », on

devra avoir, selon la présentation :

$$T > 2,5 \cdot 10^{-12} \text{ s} = 2,5 \text{ ps}$$

$$f < 4 \cdot 10^{11} \text{ Hz} = 400 \text{ GHz}$$

$$\lambda > 7,5 \cdot 10^{-4} \text{ m} = 750 \text{ }\mu\text{m}$$

ce qui correspond aux fréquences dites industrielles, aux fréquences dite hertziennes et au lointain infra-rouge.

Remarque : une notion importante traitée dans le chapitre C-VIII, celui sur les équations de MAXWELL définit l'*approximation des régimes quasi-permanents* et le limite aux fréquences inférieures à environ 1 Mhz. On voit donc que la loi d'OHM est valable en régime quasi-permanent.

1.g Influence de la température. Conducteurs et semi-conducteurs.

• Influence de la température.

Dans les différents facteurs qui apparaissent dans $\sigma = \sum_i \frac{n_i q_i^2 \tau_i}{m_i}$, les charges q_i et les masses m_i sont des constantes. Les temps de relaxation τ_i , caractérisant l'interaction des charges avec le milieu, dépendent de la température T , essentiellement parce que le milieu se dilate légèrement quand la température croît. Pour la même raison, les densités particulières varient avec T . Sans autre complication, comme dans un métal, la conductivité varie lentement avec la température.

Comme c'est souvent le cas quand les variations sont faibles, on trouve en bonne approximation une variation affine de la forme :

$$\sigma(T) = \sigma(0) (1 + aT)$$

en prenant les températures en degrés CELSIUS et avec, pour les métaux des coefficients thermiques de l'ordre de $5 \cdot 10^{-3}$.

• Conducteurs et semi-conducteurs.

Commençons par exposer les grandes lignes du modèle par bandes, assez pour comprendre sans passer par la mécanique quantique. Dans un solide cristallin, les énergies de tous les électrons sont quantifiées en *niveaux d'énergie*, deux électrons ne peuvent être sur le même niveau⁷. Au zéro absolu, les électrons remplissent les niveaux les plus bas possible ; à plus haute température, certains électrons peuvent quitter un niveau inférieur pour un niveau supérieur avec une probabilité qui relève de la thermodynamique statistique. Les

7. En fait c'est deux maximum par niveau mais avec des spins opposés, mais ça ne change pas grand chose à la compréhension.

énergies des niveaux prennent des valeurs très serrées dans certains intervalles appelées *bandes d'énergie* séparés par des intervalles sans niveaux d'énergie et appelées *bandes interdites*. Au zéro absolu, avec les règles quantiques de construction des bandes, on s'attend à ce que toutes les bandes les plus basses soient totalement remplies (un électron sur chacun de ses niveaux) et les autres vides ; la plus haute des bandes remplies s'appelle la *bande de valence* et la plus basses des bandes vides s'appelle la *bande de conduction* ; la largeur en énergie de la bande interdite entre la bande de valence et celle de conduction s'appelle le *gap* (pas d'équivalent français officiel).

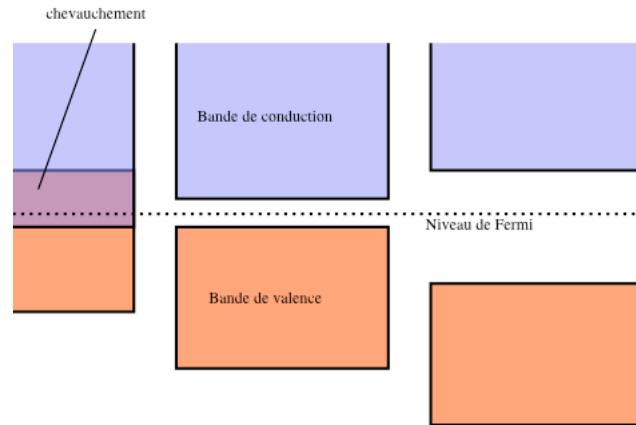


FIGURE 2 – Théorie des bandes.

Dans le cas d'un *isolant*, le gap ΔE est grand, à température ambiante T , devant kT (où k est la constante de BOLTZMANN) et la probabilité qu'un électron s'excite, en $e^{-\frac{\Delta E}{kT}}$, est négligeable ; les électrons sont coincés sur leurs niveaux sans pouvoir en changer, ils ne bougent donc pas. C'est la situation du schéma de droite de la figure 2 p. 12 issue de Wikipedia. Il y figure un *niveau de FERMİ* (voir chapitre E-IX traitant des particules indiscernables) mais peu importe ici.

Dans le cas d'un métal, le haut de la bande de valence et la bas de la bande de conduction se chevauchent. Il n'y a plus de gap et les niveaux d'énergie sont si serrés que même à très basse température les électrons peuvent s'exciter à l'intérieur de cette double bande ; changer de niveau, c'est changer aussi de localisation, les électrons peuvent bouger, ils sont libres et assurent la conduction. C'est la situation évoquée sur le schéma à gauche de la même figure.

Le schéma intermédiaire est celui des *semi-conducteurs*. Le gap est assez petit pour que quelques électrons passent par agitation thermique dans la bande de conduction ; ceux-ci, peu nombreux dans une bande riche en niveaux très serrés, peuvent y changer de niveau à volonté et assurer la conduction, comme dans un métal. Toutefois la densité particulière est beaucoup plus faible que dans un métal ; les élus sont rares. Mais ce n'est pas tout ! Dans la bande de valence, quelques niveaux sont désormais vides, on les appelle des *trous*.

La figure 3 p. 13 issue de Wikipedia schématise la formation d'une *paire électron-trou*.

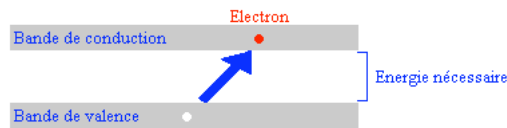


FIGURE 3 – Paire électron-trou.

Dans la bande de valence un électron sur un niveau proche de celui d'un trou peut s'y déplacer, comblant le trou mais en créant un autre ; le processus peut se réitérer, les électrons peuvent se déplacer et assure la conduction. On peut aussi considérer que c'est le trou, considéré comme une charge positive qui se déplace. La figure 4 p. 13 schématise le déplacement réel d'un électron et le déplacement virtuel du trou qui lui est opposé.

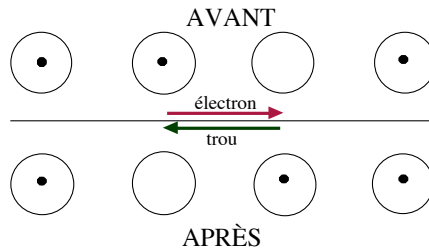


FIGURE 4 – Déplacement d'un trou.

La conséquence thermique est que le nombre de porteurs de charges varie en $e^{-\frac{\Delta E}{kT}}$ et que la conductivité d'un *semi-conducteur intrinsèque* (c'est à dire sans *dopage* comme ci-dessous) varie de la même façon et varie donc beaucoup plus rapidement avec la température qu'un métal. Un usage courant de cette propriété est la fabrication de thermistances capables de mesurer la température du milieu ambiant. Par contre, un semi-conducteur intrinsèque possède à température ordinaire peu de porteurs et sa conductivité est très faible.

• Semi-conducteurs dopés.

Pour augmenter et maîtriser la conductivité des semi-conducteurs, on procède par dopage, c'est à dire l'introduction d'impuretés dans le réseau cristallin. Au point de vue énergétique, cela conduit soit à la création de niveaux occupés juste au-dessous de la bande de conduction (dopage négatif ou N), soit à la création de niveaux vides juste au-dessus de la bande de valence (dopage positif ou P) ; dans le premier cas les électrons de ces niveaux migrent tous dans la bande de conduction, ce qui revient à y injecter un nombre déterminé d'électrons, dans le second les trous que sont ces niveaux migrent tous dans la bande de valence, ce qui revient à y injecter un nombre déterminé de trous.

La loi de GULDBERG et WAAGE (voir le chapitre E-VII de thermodynamique sur les potentiels chimiques) s'applique à la formation d'un trou (noté O) et d'un électron (noté $\bullet$) à partir d'un niveau occupé (noté Θ) dans la bande de valence. De la même façon que l'équilibre $H_2O \rightleftharpoons H^+ + OH^-$ a une constante d'équilibre $[H^+][OH^-] = K_e(T)$, l'équilibre $\Theta \rightleftharpoons O + \bullet$ a pour constante d'équilibre $[O][\bullet] = K(T)$. La conséquence est que le dopage N est analogue à l'acidification d'un milieu par apport d'ions H^+ (analogues aux électrons) et l'on en déduit que les trous (analogues des ions OH^-) sont ultra-minoritaires, ainsi donc que les électrons dus à la dissociation électron-trou et donc que les porteurs se réduisent à ceux apportés par le dopage. On raisonne de façon symétrique pour le dopage P.

• Supraconductivité.

Un phénomène longtemps inexpliqué est qu'au-dessous d'une température critique la conductance des matériaux s'annule brusquement. La découverte de cet effet fut tardive (1911 par Kamerlingh ONNES) car les températures critiques des métaux sont très basses (quelques kelvins) et accessibles uniquement dans l'hydrogène ou l'hélium liquide et sa compréhension encore plus (1957) : les électrons s'accouplent en *paires de COOPER* qui sont des bosons et subissent une condensation de BOSE-EINSTEIN (voir le chapitre E-IX sur la thermodynamique des particules indiscernables). Je n'en dirai pas plus⁸. On conçoit maintenant des alliages donc la température critique avoisine les cent kelvins, supérieure à celle de l'azote liquide moins coûteux à obtenir que l'hydrogène ou l'hélium liquide.

2 Loi d'Ohm intégrale.

2.a Champ électrostatique et champ électromoteur.

Dans tout ce qui précède, $\vec{E}$ n'est pas un champ électrique mais une grandeur homogène à un champ électrique, le rapport de la force exercée sur une charge libre à sa charge, grandeur donc très floue et c'est volontaire : la proportionnalité entre vitesse et force ne relève pas de l'électromagnétisme et l'on trouve de tels comportements en mécanique des fluides.

Il est temps désormais de détailler les forces possibles pour faire le lien avec l'électromagnétisme. Sur une charge libre de charge q et de masse m peuvent s'exercer :

- la force de LORENTZ $q(\vec{E} + \vec{v}_{\text{abs.}} \wedge \vec{B})$ où le champ électrique est lié au potentiel scalaire électrique V et au potentiel-vecteur magnétique $\vec{A}$ par $\vec{E} = -\text{grad } V - \frac{\partial \vec{A}}{\partial t}$ et où $\vec{v}_{\text{abs.}}$ est somme de la vitesse $\vec{V}_e$ du milieu (vitesse d'entraînement) et de la vitesse $\vec{v}$ de la charge libre par rapport au milieu (cf supra)
- des forces gravitationnelles $m \vec{g}$ et des forces d'inertie d'entraînement $-m \vec{a}_{\text{entr.}}$ et de CORIOLIS $-2m \vec{\Omega} \wedge \vec{v}_{\text{ref.}}$ si le référentiel choisi l'exige. Dans la pratique, même

8. faute d'avoir étudié cette théorie !

pour un ion, beaucoup plus lourd qu'un électron, on a $m \sim 10^{-25}$ kg d'où, avec $g \sim 10^1 \text{ m} \cdot \text{s}^{-2}$ et $q \sim 10^{-19} \text{ C}$, l'on a $\frac{mg}{q} \sim 10^{-5} \text{ V m}^{-1}$, totalement négligeable devant les champs électriques habituels et il en est de même pour les forces d'inertie. On peut donc négliger ces forces sauf à faire de la provocation.

- aucune force du style⁹ forces de Van der Waals ou forces de frottement de COULOMB qui ne sont que le bilan des forces précédentes sur un système à nombreuses particules.
- des forces fictives traduisant la moyenne de forces intenses mais brèves lors de chocs. Macroscopiquement, on en rend compte par des forces de frottement fluides en $-\lambda \vec{v}$ dans un milieu homogène
- des forces en $-Cte \overrightarrow{\text{grad}} n$ traduisant les phénomènes diffusifs dans un milieu hétérogène où la densité particulaire n des porteurs est inhomogène comme dans une pile électrochimique ou un électrolyte dans une couche mince autour des électrodes. Dans le cas où il y a plusieurs types de porteurs, les constantes de proportionnalité de ces deux types de forces possibles peuvent varier d'un type à l'autre.

Pour un type de charge libre d'indice i , et en négligeant (cf supra) les forces liées à la masse, chaque charge est soumise à une force totale, dont le dernier terme (on note κ la constante) n'existe pas toujours :

$$\vec{F}_i = q_i \left(-\overrightarrow{\text{grad}} V - \frac{\partial \vec{A}}{\partial t} + \vec{V}_e \wedge \vec{B} + \vec{v}_i \wedge \vec{B} \right) - \lambda_i \vec{v}_i - \kappa_i \overrightarrow{\text{grad}} n_i$$

On a défini plus haut le champ $\vec{E}$ par $q_i \vec{E} = \vec{F}_i$ quand la vitesse $\vec{v}_i$ est nulle, soit :

$$\vec{E} = -\overrightarrow{\text{grad}} V - \frac{\partial \vec{A}}{\partial t} + \vec{V}_e \wedge \vec{B} - \frac{\kappa_i}{q_i} \overrightarrow{\text{grad}} n_i$$

Le dernier terme apporte une complication, qui n'apparaît que dans les piles électrochimiques ou électrolytes. Pour alléger l'exposé, nous nous placerons en dehors de ce cas, que nous traiterons toutefois en remarque dans le prochain paragraphe. Hormis ce dernier terme provisoirement écarté, il est d'usage de décomposer formellement $\vec{E}$ en deux termes, le *champ électrostatique* $\vec{E}_s = -\overrightarrow{\text{grad}} V$ et le *champ électromoteur* $\vec{E}_m = -\frac{\partial \vec{A}}{\partial t} + \vec{V}_e \wedge \vec{B}$. Cette décomposition soulève toutefois un problème évoqué un peu plus loin au paragraphe 2.c p. 18.

2.b Résistance et tension électromotrice. Loi d'Ohm intégrale.

• Etude sans électrochimie.

On appelle ligne de courant une courbe en tout point parallèle au vecteur densité de courant $\vec{j}$ et tube de courant l'ensemble de lignes de courant passant par tous les points d'une courbe fermée.

9. genre forces de VdW, dirait ma fille.

On appellera *élément de courant filiforme* une portion de tube de courant dont la longueur $d\ell$, parallèle à $\vec{j}$ par définition des lignes de courant, soit suffisamment courte pour que les portions de lignes de champ qui le composent soient quasiment rectilignes et dont les dimensions latérales soient négligeables devant la longueur ; on appellera dS l'aire d'une section du tube orthogonale à $\vec{j}$ et $d\vec{S}$ son vecteur surface. Dans ces conditions $\vec{j}$ sera considéré comme uniforme dans le tube et l'on notera $d\vec{\ell}$ la longueur vectorielle de l'élément de courant orienté dans le sens de $\vec{j}$. On notera M et M' les deux extrémités dans le sens où $\overrightarrow{MM'} = d\vec{\ell}$. L'intensité dans cet élément est $dI = j dS$ où j est le module (la norme) de $\vec{j}$. La figure 5 p. 16 résume tout cela ; $\vec{u}$ y désigne le vecteur unitaire de la direction commune à $\vec{j}$ et $d\vec{\ell}$.

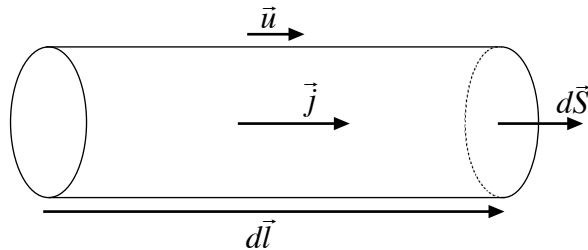


FIGURE 5 – Élément de courant.

La loi d'OHM locale et la définition des champs électrostatique et électromoteur conduisent à la relation :

$$\vec{j} = \sigma \vec{E} = \sigma (-\overrightarrow{\text{grad}} V + \vec{E}_m)$$

Multiplions scalairement par $d\vec{\ell}$ sans oublier la propriété classique du gradient, à savoir $\overrightarrow{\text{grad}} V \cdot d\vec{\ell} = dV = V(M') - V(M)$, on arrive à :

$$\vec{j} \cdot d\vec{\ell} = j d\ell = \sigma (-dV + \vec{E}_m \cdot d\vec{\ell}) = \sigma [V(M) - V(M') + \vec{E}_m \cdot d\vec{\ell}]$$

En introduisant l'intensité $dI = j dS$ on peut réécrire ainsi cette relation :

$$V(M) - V(M') = \frac{d\ell}{\sigma dS} dI - \vec{E}_m \cdot d\vec{\ell}$$

Le facteur $r = \frac{d\ell}{\sigma dS}$ est appelé *résistance* de l'élément de courant et le terme $de = \vec{E}_m \cdot d\vec{\ell}$ est appelé *tension électromotrice* ou *force électromotrice*, usuellement abrégée en f.e.m.¹⁰. Remarque : on appelle $g = \frac{1}{r}$ la conductance de l'élément. On peut donc écrire :

$$V(M) - V(M') = r dI - de \quad \text{ou} \quad dI = g [V(M) - V(M') + de]$$

10. Je vendrais mon âme au diable pour un calembour calamiteux et ne puis résister à celui-ci : l'électrocinétique c'est l'ohm et la f.e.m.

Ce raisonnement reste valable pour un conducteur métallique cylindrique homogène de longueur ℓ et d'extrémité A et B , de section d'aire S et de diamètre petit devant ℓ , pourvu que l'on puisse considérer $\vec{j}$ comme uniforme (c'est vrai en courant continu et basse fréquence, nous l'admettons ici ; ce sera démontré dans un chapitre ultérieur). On aboutit alors, en notant I_{AB} l'intensité du courant, comptée positivement de A vers B , à la *loi d'OHM intégrale* et avec la notation traditionnelle V_A pour $V(A)$ et analogues, à ces deux formulations équivalentes :

$$V_A - V_B = R I_{AB} - E_{AB} \quad \text{ou} \quad I_{AB} = G [(V_A - V_B) + E_{AB}]$$

où, en notant $\rho = \frac{1}{\sigma}$ la *résistivité* du milieu

$$R = \frac{\ell}{\sigma S} = \frac{\rho \ell}{S} \quad G = \frac{S}{\rho \ell} = \frac{\sigma S}{\ell} \quad E_{AB} = \int_A^B \vec{E}_m \cdot d\vec{\ell}$$

C'est la loi d'OHM sous sa forme historique et classique qui est une forme intégrale.

On reviendra un peu plus en détail sur le passage de l'élément de courant au cylindre et à bien d'autres formes.

• Etude avec électrochimie.

Comment gérer le problème des forces fictives gérant les phénomènes diffusifs dont la constante dépend du type de porteur ? Se poser la question c'est y répondre. Pour chaque type de porteur, le vecteur $\vec{E}$ qui devient logiquement $\vec{E}_i$ s'écrit $\vec{E}_i = -\overrightarrow{\text{grad}} V + \vec{E}_{m,i}$ avec :

$$\vec{E}_{m,i} = -\frac{\partial \vec{A}}{\partial t} + \vec{V}_e \wedge \vec{B} - \frac{\kappa_i}{q_i} \overrightarrow{\text{grad}} n_i$$

On pourra calculer la densité de courant par sommation des contributions de chaque espèce, soit

$$\vec{j} = \sum_i \vec{j}_i = \sum_i \sigma_i \vec{E}_i = \sum_i \sigma_i (-\overrightarrow{\text{grad}} V + \vec{E}_{m,i})$$

Notons $\sigma = \sum_i \sigma_i$, détaillons chaque $\vec{E}_{m,i}$ en un terme commun d'origine magnétique $\vec{E}_{m,0} = -\frac{\partial \vec{A}}{\partial t} + \vec{V}_e \wedge \vec{B}$ et un terme individuel diffusif $\vec{E}_{d,i} = -\frac{\kappa_i}{q_i} \overrightarrow{\text{grad}} n_i$; on arrive ainsi à :

$$\vec{j} = \sigma (-\overrightarrow{\text{grad}} V + \vec{E}_{m,0}) + \sum_i \sigma_i \vec{E}_{d,i}$$

Pour un élément de courant filiforme, multiplions comme plus haut par $d\vec{\ell}$, divisons par $\sigma = \sum_i \sigma_i$ et introduisons l'intensité ; on arrive (on ne détaille plus les calculs) à :

$$\frac{d\ell}{\sigma dS} dI = r dI = V(M) - V(M') + de$$

où l'on a :

$$de = \overrightarrow{E_{m,0}} \cdot \overrightarrow{d\ell} + \frac{\sum_i \sigma_i \overrightarrow{E_{d,i}} \cdot \overrightarrow{d\ell}}{\sum_i \sigma_i}$$

Donc rien d'insurmontable à gérer, mais comme on ne trouve pas ou peu cette explication dans la littérature, il fallait bien que je la donnasse ici.

Cela dit, le calcul des f.e.m. dans ce contexte électrochimique s'effectue avec beaucoup plus de pertinence par une approche thermochimique que vous trouverez un peu plus loin au paragraphe 4.b p. 31 car je ne peux rien vous refuser.

2.c Premières indications sur l'induction magnétique.

Les effets électrocinétiques des termes d'origine magnétique des f.e.m. ont des conséquences suffisamment délicates, variées et riches d'applications technologiques pour mériter d'être traités dans le chapitre C-VII qui leur sera consacré sous le titre historique d'induction magnétique.

En outre, le chapitre C-VIII sur les équations de MAXWELL va soulever un problème. L'interaction entre charges en mouvement est gérée par les champs électrique et magnétique, ou plutôt le champ électromagnétique ; le potentiel scalaire électrique V et le potentiel-vecteur magnétique $\overrightarrow{A}$ sont des outils pratiques qui se sur-ajoutent à ces champs et en sont en quelque sorte une description. Le problème est que pour un champ électromagnétique donné, il y a une infinité de couples de potentiels $(V, \overrightarrow{A})$ qui le décrivent ; on en choisit un, a priori arbitrairement, en pratique un qui simplifie les calculs par ce que l'on appelle une *condition de jauge*.

Pour un même élément de courant, si l'on change de jauge, les valeurs de $V(M) - V(M')$ et de la f.e.m. élémentaire $de = \overrightarrow{E_m} \cdot \overrightarrow{d\ell}$ changent toutes les deux, mais bien sûr la loi d'OHM reste valable avec ces nouvelles valeurs et, bien entendu, la même résistance r et le même courant dI qui n'ont rien à voir avec la condition de jauge. La notion de différence de potentiel et de f.e.m. a donc fondamentalement une part d'arbitraire.

Que me dis-tu, ô mon lecteur ? Qu'expérimentalement, un voltmètre, un vieux, à aiguille, donne la même indication quand l'expérimentateur décide dans sa tête de changer de jauge. C'est vrai, tu as raison mais on t'as roulé : un voltmètre, ça n'existe pas, il mesure l'intensité qui le traverse et son cadran effectue de façon statique la multiplication par sa résistance, il mesure donc $(V_A - V_B) - e_{AB}$ et ce n'est que quand on est sûr que la f.e.m. est nulle qu'il est un voltmètre. Pour les voltmètres électroniques, la même explication se cache dans les chaînes de rétroaction des amplificateurs opérationnels.

2.d Calcul des résistances mortes.

• Association en série.

Considérons deux éléments successifs $\overrightarrow{MM'}$ et $\overrightarrow{M'M''}$ d'un même tube de courant, donc parcourus par la même intensité dI . Si l'on écrit la loi d'OHM pour chacun des éléments, en adaptant la notation à notre propos et que l'on additionne, on a successivement à :

$$\begin{aligned} V_M - V_{M'} &= r_{MM'} dI - de_{MM'} \\ V_{M'} - V_{M''} &= r_{M'M''} dI - de_{M'M''} \\ V_M - V_{M''} &= (r_{MM'} + r_{M'M''}) dI - (de_{MM'} + de_{M'M''}) \end{aligned}$$

ce qui prouve à la simple lecture qu'il y a additivité d'une part des résistances, d'autre part des tensions électromotrices.

On généralise aisément à un tube de champ de faible section, éventuellement variable donc notée $dS(M)$, de longueur quelconque de A à B ; sa résistance et sa f.e.m. se calculent par intégration des contributions d'éléments infiniment courts, soit, en supposant le milieu homogène donc la conductivité constante

$$\begin{aligned} R_{AB} &= \int r_{MM'} = \int_A^B \frac{1}{\sigma dS(M)} d\ell = \frac{1}{\sigma} \int_A^B \frac{1}{dS(M)} d\ell \\ e_{AB} &= \int_A^B de_{MM'} = \int_A^B \vec{E}_m(M) \cdot d\vec{\ell} \end{aligned}$$

• Association en parallèle.

Imaginons deux éléments de courant juxtaposés MM' et NN' avec N quasiment confondu avec M de sorte que $V_M \approx V_N$ et $\vec{E}(M) \approx \vec{E}(N)$ ainsi que N' quasiment confondu avec M' de sorte que $V_{M'} \approx V_{N'}$. Leur réunion forme un nouvel élément parcouru par la somme des intensités que nous noterons ici $dI_{\text{tot.}} = I_M + dI_N$. On écrit cette fois la loi d'OHM sous son autre forme et pour les deux éléments, avant d'additionner, soit, en adaptant les notations et en confondant M et N , M' et N' :

$$\begin{aligned} dI_M &= g_{MM'} [(V_M - V_{M'}) - \vec{E}_m(M) \cdot \overrightarrow{MM'}] \\ dI_N &= g_{NN'} [(V_M - V_{M'}) - \vec{E}_m(M) \cdot \overrightarrow{MM'}] \\ dI_{\text{tot.}} &= (g_{MM'} + g_{NN'}) [(V_M - V_{M'}) - \vec{E}_m(M) \cdot \overrightarrow{MM'}] \end{aligned}$$

cette fois il y a additivité des conductances et identité des f.e.m. mais attention, si l'on procède par intégration (une intégrale double sur la surface traversée) sur de nombreux tubes dont les origines soient sur une même surface équipotentielle de valeur noté V_M et

les extrémités sur une autre équipotentielle à $V_{M'}$ (on notera $U = V_M - V_{M'}$), on ne pourra plus supposer le champ électromoteur comme constant d'un tube élémentaire à l'autre et l'on arrivera, en passant les détails de calculs qui sont aisés, à :

$$I_{\text{tot.}} = \iint g_{MM'} (U + \vec{E}_m(M) \cdot \overrightarrow{MM'}) = U \iint \frac{\sigma}{d\ell(M)} dS + \iint \frac{\sigma}{d\ell(M)} \vec{E}_m(M) \cdot \overrightarrow{d\ell(M)} dS$$

pour la conductivité, il y a toujours additivité, mais pour la f.e.m., rien de simple ni d'exploitable. En pratique, on ne peut utiliser cette approche que dans le cas où l'on est sûr qu'il n'y a pas de champ électromoteur ; on parle alors de *résistance morte*. Dans les autres cas, s'il s'agit de mise en parallèle de deux éléments (quelques-uns peu nombreux) filiformes, on utilisera l'approche par les réseaux électrocinétiques et s'il s'agit d'une masse métallique quelconque, on doit gérer la situation au cas par cas (voir au paragraphe 3.b p. 23 l'exemple de l'effet de magnétorésistance).

• Calcul d'une résistance morte non filiforme.

Parler de la résistance d'un bloc métallique sans autre précision n'a aucun sens car ce bloc peut, selon les circonstances être parcouru par des courants dans tous les sens. Il faut donc préciser le branchement, c'est-à-dire par où y entre le courant et par où il en sort, ou encore donner les deux surfaces équipotentielles d'entrée et de sortie du courant. Encore faut-il aussi connaître la géométrie des lignes de courant, que l'on ne peut pas imposer. Pour cela, on est aidé par trois remarques : la première est qu'en courant continu ou basse fréquence, la densité de courant a une divergence nulle¹¹, soit $\text{div } \vec{j} = 0$, la deuxième est que la loi d'OHM locale dans un milieu homogène ($\sigma = Cte$) permet d'en déduire que le champ électrique a lui aussi une divergence nulle $\text{div } \vec{E} = \frac{1}{\sigma} \text{div } \vec{j} = 0$ et enfin la troisième est que pour une résistance morte, il n'y a pas de champ électromoteur d'où, en particulier, $\vec{E} = -\text{grad } V$, d'où $\vec{0} = \text{div } \vec{E} = -\text{div grad } V = -\Delta V$. Dans une résistance morte, le potentiel a une divergence nulle et l'on connaît deux équipotentielles, exactement comme dans l'espace vide entre deux conducteurs à l'équilibre (voir chapitre C-II) ; on sait que le problème ne peut être résolu explicitement que dans des cas de haute symétrie, sinon par des algorithmes numériques gourmands et mis en œuvre sur ordinateur.

Prenons l'exemple d'un cylindre creux de hauteur h de rayon intérieur r_1 et de rayon extérieur r_2 petit devant h avec la surface intérieure au potentiel V_1 et la surface extérieure au potentiel V_2 inférieur à V_1 (on notera $U = V_1 - V_2$). La symétrie de révolution, aux effets de bords près donne, en coordonnées cylindriques, un potentiel $V(r)$ et un champ radial $E(r) \vec{e}_r$ et donc une densité de courant radiale $\vec{j} = \sigma E(r) \vec{e}_r$. Un élément de tube de courant élémentaire, radial donc, défini par sa « hauteur » élémentaire dz , sa « largeur » angulaire $d\theta$ et de longueur dr dans le sens du courant entre les rayons r et $r + dr$; sa surface est classiquement $r d\theta dz$, sa résistance est $\frac{d\ell}{\sigma dS}$ soit ici $\frac{dr}{\sigma r d\theta dz}$ et sa conductance en est l'inverse. On pourra s'aider de la figure 7 p. 24 qui étudie un cas plus complexe.

11. Ce sera démontré dans le chapitre C-VIII sur les équations de MAXWELL

Il faudra associer ces éléments en série de $r = r_1$ à $r = r_2$ et en parallèle de $z = 0$ à $z = h$ et de $\theta = -\pi$ à $\theta = \pi$; on fait cela en deux temps, peu importe l'ordre.

La mise en parallèle pour r donne, par addition, une conductance :

$$g_{r,r+dr} = \int_{z=0}^{z=h} \int_{\theta=-\pi}^{\theta=\pi} \frac{\sigma r d\theta dz}{dr} = \frac{2\pi\sigma h r}{dr}$$

que l'on s'est refusé à noter dg car c'est un infiniment grand ; par contre, son inverse est une résistance élémentaire dont la mise en série donne la résistance de ce dispositif, soit :

$$R = \int_{r=r_1}^{r=r_2} \frac{dr}{2\pi\sigma h r} = \frac{\ln\left(\frac{r_1}{r_2}\right)}{2\pi\sigma h}$$

Remarque 1 : Dans le cas d'un condensateur, près des bords, côté $z = 0$ et $z = h$, les lignes de champ cessent d'être radiales et font des incursions respectivement du côté $z < 0$ et $z > h$ et le calcul théorique n'est valable que si $h \gg r_2$. Dans le cas d'une résistance, les lignes de courant sont confinées dans le métal et ne peuvent s'échapper et la formule est valable même sans cette condition.

Remarque 2 : c'est par cette méthode que l'on montre, très rapidement que la résistance d'un cylindre de longueur ℓ et de section S est $R = \frac{\ell}{\sigma S}$

Remarque 3 : Si l'on compare un calcul de résistance et un calcul de capacité de même géométrie, le flux du champ électrique à travers une section élémentaire de vecteur surface $d\vec{S}$ sur une équipotentielle donne $E dS = \frac{1}{\sigma} j dS = \frac{1}{\sigma} dI$ dans le premier cas et, grâce au théorème de GAUSS, $E dS = \frac{1}{\varepsilon_0} dQ$ où dQ est la charge sur une armature au « pied » du tube de champ dans le second. On calcule par intégration sur dS respectivement :

$$GU = I = \iint dI = \sigma \iint E dS \quad \text{et} \quad CU = Q = \iint dQ = \varepsilon_0 \iint E dS$$

A différence de potentiel égale, donc à champ électrique égal en tous points, le rapport membre à membre donne la relation :

$$\frac{G}{C} = \frac{\sigma}{\varepsilon_0}$$

et si l'on connaît un résultat, on connaît l'autre. Pour la résistance étudiée ci-dessus, on a trouvé $G = \frac{1}{R} = \frac{2\pi\sigma h}{\ln\left(\frac{r_1}{r_2}\right)}$, on en déduit qu'un condensateur cylindrique a une capacité

$$C = \frac{2\pi\varepsilon_0 h}{\ln\left(\frac{r_1}{r_2}\right)}$$

3 Les complications.

3.a L'effet Hall.

Il s'agit ici d'étudier un premier effet d'un champ magnétique sur un courant électrique, connu sous le nom de *effet HALL*. Un conducteur parallélépipédique de longueur ℓ dans le sens attendu du courant est placé dans un champ magnétique $\vec{B}$ orthogonal à ses faces les plus larges de largeur b et sa hauteur dans le sens du champ est a . On a tracé la figure 6 p. 22 dans le cas où les porteurs de charge ont une charge q négative ; le courant y arrive par la droite et ressort par la gauche si bien que les porteurs se dirigent vers la droite.

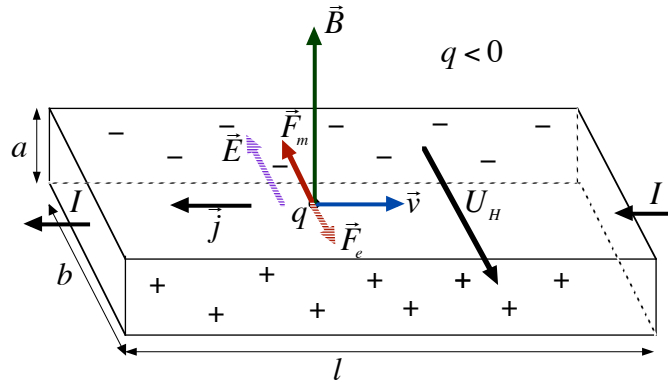


FIGURE 6 – Effet Hall.

Avec une vitesse $\vec{v}$ orientée vers la droite (en bleue), pas forcément parallèle à l'arête de longueur ℓ , le champ magnétique $\vec{B}$ (en vert foncé) exerce sur les porteurs une force $q \vec{v} \wedge \vec{B}$ (en marron) pointe vers l'arrière de la figure et les charges dévient dans cette direction (ce qui explique que $\vec{v}$ n'est pas dans la direction intuitive) mais, confinées dans le conducteur, elles ne peuvent aller plus loin que la face arrière et s'y accumulent ; de même, les charges qui se déplacent près de la face avant s'en éloignent et ne peuvent pas être remplacées car il n'y a pas de charges en avant de la face avant qui se charge donc positivement. Ces deux faces chargées se comportent comme les deux armatures d'un condensateur et créent un champ électrique supplémentaire $\vec{E}_H$ (en violet hachuré) dirigé vers l'arrière et exerçant sur les charges négatives une force $q \vec{E}_H$ (en marron hachuré) vers l'avant. Le phénomène se poursuit jusqu'à ce que la force électrique compense la force magnétique ; alors et seulement alors la vitesse des charges est parallèle à l'arête de longueur ℓ .

On a alors $q(\vec{E}_H + \vec{v} \wedge \vec{B}) = \vec{0}$ d'où, en module (en norme), $E_H = v B$. Expérimentalement, on peut mesurer d'une part l'intensité $I = j S = (n q v) (a b)$ et d'autre part la différence de potentiel entre les deux faces chargées $U_H = \int dV = - \int \vec{E}_H \cdot d\vec{l} = E_H b$, d'où, en valeur absolue, la valeur de la différence de potentiel latérale de cet effet HALL

(on passe les calculs élémentaires) :

$$U_H = \frac{1}{n q a} I B$$

où $\frac{1}{n q a}$ est appelée *constante de HALL* du matériau.

Remarque 1 : si avec l'intensité dans le même sens, les charges sont positives, celles-ci vont désormais vers la gauche donc $q \vec{v}$ n'a pas changé de sens, la force magnétique non plus mais ce sont cette fois des charges positives (les porteurs) qui s'accumulent sur la face arrière, le champ électrique change de sens et la différence de potentiel de signe. L'effet HALL permet donc de connaître le signe des porteurs ; on a pu ainsi constater que pour un semi-conducteur dopé positivement, ce sont bien des trous positifs qui se déplacent.

Remarque 2 : l'effet est d'autant plus marqué que a est petit. Dans les sondes à effet HALL, on choisit donc a le plus petit possible.

Remarque 3 : avec a , I et B connu et avec $q = \pm e$ (électrons ou trous), on peut mesurer n , ce qui permet, par exemple, de vérifier que dans le cuivre, il y a un seul électron libre par atome de cuivre.

Remarque 4 : l'utilisation la plus courante de l'effet HALL est la mesure des champs magnétiques, rôle dans lequel est à aujourd'hui la quasi-exclusivité.

Remarque 5 : on explique souvent ¹² que la force de LAPLACE (en $I d\vec{l} \wedge \vec{B}$) est le bilan macroscopique des forces de LORENTZ sur les porteurs. Ce n'est pas vrai, les porteurs sont soumis à des forces nulles ; par contre le réseau cristallin, de vitesse nulle, de charge opposée aux porteurs est soumis à l'action du champ électrique créé par l'effet HALL et c'est en fait l'origine de la force de LAPLACE.

3.b Magnéto-résistance.

Reprenons l'étude de la résistance cylindre prise plus haut en exemple et placée cette fois dans un champ magnétique uniforme et stationnaire, parallèle à l'axe, assez intense pour que ses effets soient décelables, c'est à dire pour que la force en $q \vec{v} \wedge \vec{B}$ négligée en début de chapitre ne le soit plus. Comme pour l'effet HALL les charges sont déviées mais cette fois de façon orthoradiale donc sans paroi pour les bloquer ; la densité de courant aura une composante radiale j_r et une composante orthoradiale j_θ ; enfin le champ électrostatique $\vec{E}_S = -\vec{\text{grad}} V$ est toujours radial par symétrie. La figure 7 p. 24 résume la situation en projection orthogonale à l'axe avec $\vec{B}$ vers l'arrière de votre écran. On y a dessiné un élément de courant, grossi pour la lisibilité.

Ici, il faut reprendre à zéro (mais plus vite) le problème de conduction car nous avons éludé ce cas qui pose problème. Dans le modèle à frottement fluide et sans complications supplémentaires (le conducteur est fixe, donc pas de force d'inertie ni de force en $\vec{V}_e \wedge \vec{B}$

12. et je l'ai souvent fait moi-même, y compris dans le chapitre C-III de ce cours.

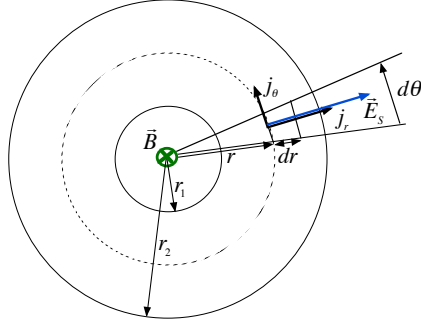


FIGURE 7 – Magnéto-résistance.

et sans effet de diffusion, dans un cadre de magnétostatique donc $\frac{d\vec{A}}{dt} = \vec{0}$), on a comme équation de mouvement d'une charge :

$$m \frac{d\vec{v}}{dt} = q (-\vec{\text{grad}} V + \vec{v} \wedge \vec{B}) - \lambda \vec{v}$$

Escamotons ici la démonstration¹³ que $\vec{v}$ atteint une vitesse limite. Lorsque celle-ci est atteinte, on a $\frac{d\vec{v}}{dt} = \vec{0}$ et, en reprenant la notation de la mobilité $\mu = \frac{q}{\lambda}$, on arrive à :

$$\vec{v} = \mu (-\vec{\text{grad}} V + \vec{v} \wedge \vec{B})$$

La relation entre densité de courant et vitesse des porteurs reste, avec un seul type de porteurs (pour alléger l'exposé) $\vec{j} = n q \vec{v}$ que l'on reporte dans la relation précédente pour arriver, en notant la conductivité $\sigma = n q \mu$, à :

$$\vec{j} = \sigma (-\vec{\text{grad}} V + \frac{1}{n q} \vec{j} \wedge \vec{B})$$

Avec $\vec{j} = j_r \vec{e}_r + j_\theta \vec{e}_\theta$, $\vec{B} = B \vec{e}_z$ (avec, dans le cas de la figure, $B < 0$) et enfin $\vec{\text{grad}} V = \frac{dV}{dr} \vec{e}_r$, on en déduit en projection sur les axes :

$$\begin{cases} j_r = -\sigma \frac{dV}{dr} + \frac{\sigma B}{n q} j_\theta \\ j_\theta = -\frac{\sigma B}{n q} j_r \end{cases}$$

d'où l'on tire, en éliminant j_θ , la relation :

$$\left[1 + \left(\frac{\sigma B}{n q} \right)^2 \right] j_r = -\sigma \frac{dV}{dr}$$

13. On projette sur l'axe radial et l'axe orthoradial, on fait apparaître une matrice dont le polynôme caractéristique a deux racines complexes conjugués à partie réelle négative, ce qui suffit à conclure.

Le flux de $\vec{j}$ à travers un cylindre de hauteur h et de rayon r est l'intensité I , indépendante de r en régime permanent, d'où $I = 2\pi r h j_r$; on en déduit :

$$\left[1 + \left(\frac{\sigma B}{n q}\right)^2\right] I = -2\pi h \sigma r \frac{dV}{dr}$$

qui prouve que $-r \frac{dV}{dr}$ est une constante notée K , d'où $V(r) = -K \ln r + Cte$ et puisque $V(r_1) = V_1$ et $V(r_2) = V_2$, on a $U = V_1 - V_2 = K \ln\left(\frac{r_2}{r_1}\right)$ et l'on arrive finalement à :

$$\left[1 + \left(\frac{\sigma B}{n q}\right)^2\right] I = 2\pi h \sigma \frac{U}{\ln\left(\frac{r_2}{r_1}\right)}$$

On en déduit la résistance par $R = \frac{U}{I}$ et l'on y reconnaît la résistance calculée plus haut en l'absence de champ magnétique, et notée ici R_0 , soit

$$R = R_0 \left[1 + \left(\frac{\sigma B}{n q}\right)^2\right]$$

La présence du champ magnétique dans ce contexte augmente la résistance d'un facteur égal, en reportant la relation $\sigma = n q \mu$, à $[1 + (\mu B)^2]$.

L'effet est quadratique donc imperceptible si μB est petit. Pour le cuivre, nous avons vu que $\mu = 4,4 \cdot 10^{-3} \text{ m}^2 \text{ s}^{-1} \text{ V}^{-1}$ donc pour un champ déjà intense de 0,25 tesla, $(\mu B)^2 = 10^{-6}$ et l'effet prévu théoriquement n'est pas décelable. Il ne peut l'être que dans un milieu où les porteurs ont une grande mobilité et c'est le cas pour les électrons (mais pas les trous) de certains semi-conducteurs exceptionnels comme l'antimoniure d'indium InSb avec $\mu = 7,7 \text{ m}^2 \text{ s}^{-1} \text{ V}^{-1}$

3.c Diode à vide.

On modélise une diode à vide par un tube dans lequel on a fait le vide dans lequel on a placé un condensateur plan dont les armatures ont une surface S et sont distantes de a ; on choisit les axes de sorte que l'une soit dans le plan $x = 0$ et l'autre $x = a$. L'armature $x = 0$ est chauffée par effet Joule grâce à un filament annexe à une température suffisante pour que certains électrons de charge $-e$ et de masse m acquièrent assez d'énergie pour s'échapper au métal (énergie d'extraction) et en sortent avec une vitesse négligeable. Ils sont alors attirés vers l'autre armature par l'action d'un champ électrique dirigé de $x = a$ vers $x = 0$ de la forme $\vec{E} = E(x) \vec{e}_x$ avec $E(x)$ négatif, aux effets de bords près, engendré par une différence de potentiel entre les armatures ; on prendra l'origine des potentiels en $x = 0$, de sorte que l'armature en $x = a$ soit au potentiel U positif et l'on notera $V(x)$ le potentiel entre les armatures. Entre celles-ci, il y a des électrons qui se déplacent à une vitesse $\vec{v}(x) = v(x) \vec{e}_x$, d'où l'existence d'une densité volumique de charges $\rho(x)$ négative et

d'une densité de courant $\vec{j} = \rho(x) \vec{v}(x)$. Le lecteur aura remarqué que nous nous sommes tacitement placés en régime permanent établi. Dans la pratique, il y a toutes sortes de géométries ; la plus courante est celle de deux cylindres coaxiaux, celui du centre étant creux et chauffé par une résistance chauffante qui le traverse ; parfois même, la résistance chauffante fait office de cylindre intérieur. On trouve aussi la géométrie étudiée ici et qui allège un peu le calcul. Je n'ai pas cru utile de tracer une figure.

Appliquons la loi de la dynamique à un électron, mais en gérant l'approche eulérienne par la dérivée particulaire (voir le chapitre B-XIII de mécanique des fluides). Le poids d'un électron est toujours négligeable devant les forces électriques, d'où :

$$m \frac{D\vec{v}}{Dt} = m \left(\frac{\partial \vec{v}}{\partial t} + (\vec{v} \cdot \text{grad}) \vec{v} \right) = -e \vec{E}$$

On est en régime permanent, $\vec{v} = v(x) \vec{e}_x$ et $\vec{E}$ dérive du potentiel V . On arrive donc à :

$$m v(x) \frac{dv}{dx} = e \frac{dV}{dx}$$

Par intégration, on en déduit, compte tenu qu'en $x = 0$ on a $v(0) = 0$ et $V(0) = 0$, que :

$$\frac{1}{2} m v(x)^2 = e V(x)$$

que l'on peut bien sûr trouver directement en appliquant le théorème de l'énergie mécanique. Toutefois, la première approche permet, en appliquant sa conclusion en $x = 0$, de montrer que $\frac{dV}{dx}$ s'y annule car $v(0) = 0$, ce qui simplifie la suite.

L'équation de POISSON (voir le chapitre C-VIII sur les équations de MAXWELL) donne ici :

$$\Delta V = \frac{d^2 V}{dx^2} = -\frac{\rho(x)}{\varepsilon_0}$$

Enfin le flux de $\vec{j} = \rho \vec{v}$ à travers un plan parallèle aux armatures, dont seule une surface S est traversé par les électrons, est l'intensité, négative¹⁴ car les électrons négatifs vont dans le sens positif, notée $-I$ dans ces conditions. On a donc :

$$\rho(x) v(x) S = -I$$

où I est indépendant de x en régime permanent, c'est valable aussi pour la diode à vide avec le raisonnement classique¹⁵.

14. Historiquement, on a choisi un sens positif avant de connaître la nature des porteurs. c'était une chance sur deux... et on a perdu. Cette erreur-là, c'est la faute à Volta, quel manque de pot, c'est la faute à Rousseau.

15. Si $I(x_1)$ était différent de $I(x_2)$, la charge entre x_1 et x_2 varierait et le régime ne serait pas permanent.

On en déduit une équation différentielle vérifiée par $V(x)$:

$$\frac{d^2V}{dx^2} = -\frac{\rho(x)}{\varepsilon_0} = \frac{I}{\varepsilon_0 S v(x)} = \frac{I}{\varepsilon_0 S} \sqrt{\frac{m}{2e}} \sqrt{\frac{1}{V(x)}}$$

On multiplie les deux membres par $\frac{dV}{dx}$ ce qui permet une première intégration où le constante est déterminée grâce à $V(0) = 0$ et $\frac{dV}{dx}$ nul en $x = 0$ (voir remarque ci-dessus) ; on arrive ainsi à :

$$\frac{1}{2} \left(\frac{dV}{dx} \right)^2 = \frac{I}{\varepsilon_0 S} \sqrt{\frac{m}{2e}} 2 \sqrt{V(x)}$$

On extrait la racine (dans ce contexte $V(x)$ croît de gauche à droite car $V(a) = U > 0$) et l'on sépare les variables :

$$\frac{dV}{V^{\frac{1}{4}}} = 2 \left(\frac{I}{\varepsilon_0 S} \right)^{\frac{1}{2}} \left(\frac{m}{2e} \right)^{\frac{1}{4}} dx$$

d'où par intégration et avec $V(0) = 0$:

$$\frac{4}{3} V(x)^{\frac{3}{4}} = 2 \left(\frac{I}{\varepsilon_0 S} \right)^{\frac{1}{2}} \left(\frac{m}{2e} \right)^{\frac{1}{4}} x$$

et enfin avec $V(a) = U$, on arrive à la relation qui lie U et I , qui est la *caractéristique* de la diode, soit ici :

$$U^{\frac{3}{4}} = \left(\frac{3a}{2} \right) \left(\frac{I}{\varepsilon_0 S} \right)^{\frac{1}{2}} \left(\frac{m}{2e} \right)^{\frac{1}{4}}$$

dont on retiendra surtout que c'est de la forme $U^{\frac{3}{4}} = Cte I^{\frac{1}{2}}$ ou mieux de la forme $U = Cte I^{\frac{2}{3}}$. La diode à vide est donc fondamentalement non linéaire.

La diode à vide a été remplacée à peu près partout par la diode à semi-conducteurs (et les triodes¹⁶ par les transistors), car elle a deux défauts majeurs : elle est encombrante et à l'allumage, il faut attendre une bonne minute pour que l'armature concernée soit suffisamment chaude pour émettre des électrons. Toutefois, elle a une inertie très faible car il s'agit d'électrons de masse infime se déplaçant dans le vide ; dans les fonctions amplificatrice de la triode, le temps de réponse à une impulsion est dès lors extrêmement faible et rien ne vaut un amplificateur à lampes pour amplifier sa guitare électrique.

3.d Diodes à semi-conducteur.

On se contente ici de l'explication minimale pour la diode la plus simple, dite *diode à jonction*. De plus amples explications, faisant intervenir la thermodynamique et des effets

16. On intercale une grille entre les armatures à un potentiel répulsif dont les variations modulent le courant de la diode ; à partir de là, on arrive aisément à amplifier les variations du potentiel de grille.

tunnel, d'avalanche et ZENER seront données dans le chapitre C-VI consacré aux bases de l'électronique.

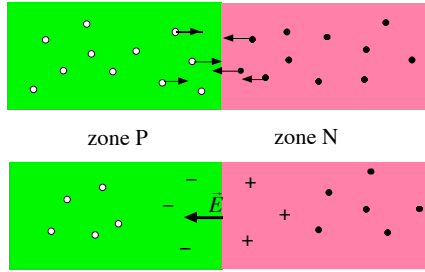


FIGURE 8 – Diode à jonction.

La figure 8 p. 28 schématise une diode à jonction qui résulte de la mise en contact d'un semi-conducteur dopé P (à gauche en vert), avec des trous majoritaires (les ronds blancs) et des électrons libres ultra minoritaires, non représentés et d'un semi-conducteur dopé N (à droite en rose), avec des électrons libres majoritaires (les ronds noirs) et des trous ultra minoritaires, non représentés. L'inhomogénéité ainsi créée provoque la diffusion sur une faible épaisseur des électrons libres de la zone N vers la zone P (et aussi des trous de la zone P vers la zone N ; on laisse le soin au lecteur de vérifier que ça produit le même effet) où ils se recombinaient avec les trous (voir l'analogie acido-basique développée au sein du paragraphe 1.g débutant p. 11 ; ici la diffusion est comme un apport d'acide dans un milieu basique et il y a neutralisation) avec un double effet, le premier que les porteurs de charges disparaissent, le second que le milieu devient électriquement négatif ; de même il y a la diffusion sur une faible épaisseur des trous de la zone P vers la zone N où ils se recombinaient avec les électrons libres (un apport de base dans un milieu acide) avec un double effet, le premier que les porteurs de charges disparaissent, le second que le milieu devient électriquement positif. Cette zone, créée par diffusion est peu épaisse comme tout phénomène diffusif, pratiquement sans porteurs donc quasiment isolante et chargée en deux demi-zones de charge opposée, ce qui génère un champ électrique dirigé vers la zone P et une différence de potentiel (une *barrière de potentiel*) avec $U_{NP} = V_N - V_P > 0$.

On observe toutefois deux courants opposés qui s'annulent en régime permanent lorsque la jonction n'est branchée à aucun générateur, celui des charges ultra-minoritaires des zones N et P situées aux bords de la zone chargée et qui se font happer par le champ dans le bon sens pour elles, courant de N vers P et celui des charges majoritaires des zones N et P dont l'énergie est suffisamment grande pour franchir la barrière de potentiel malgré le champ dans le mauvais sens pour elles et dont la proportion est proportionnelle à $e^{-\frac{eU_{NP}}{kT}}$ d'après les lois thermodynamiques, courant de P vers N.

Si l'on soumet la diode à une différence de potentiel U selon son sens, on favorise ou on défavorise la diffusion, on augmente ou on diminue l'épaisseur de la zone non conductrice, on augmente ou diminue la barrière de potentiel, on diminue ou on augmente le second

courant, le premier restant constant et l'on arrive à justifier des caractéristiques en :

$$I = I_0 \left(e^{\frac{eU}{kT}} - 1 \right)$$

ce qui est un autre exemple de conducteur non linéaire.

4 Aspects thermodynamiques.

4.a Puissance fournie un dipôle. Effet Joule.

Comme on a appris très jeune que $P = UI$, on croit que ça va être simple. Oh non !

L'énergie $\mathcal{E}$ d'un système sera ici la somme de l'énergie interne et de l'énergie potentielle liée au au champ électrostatique $-\overrightarrow{\text{grad}} V$; on fait ce choix pour se mettre en cohérence avec les modèles de conduction qui traitent à part la force due à ce champ (cf champ électrostatique et champ électromoteur).

Considérons, en régime permanent ou non, une portion de circuit filiforme, ou dipôle électrocinétique, d'extrémités A et B parcouru par un courant I de A vers B sous une différence de potentiel $U = V_A - V_B$, et recevant de l'extérieur une puissance P . Attention AB n'est pas un système mais un volume de contrôle puisqu'il y entre des charges d'un côté et qu'il en sort autant de l'autre. Soit le système formé, à l'instant t , du dipôle, d'énergie $\mathcal{E}_{AB}(t)$, et de la charge $\delta Q_e = I dt$, d'énergie $\delta \mathcal{E}_e$, qui y entre entre t et $t + dt$ et donc formé, à l'instant $t + dt$, du dipôle, d'énergie $\mathcal{E}_{AB}(t + dt)$, et de la charge $\delta Q_s = I dt = \delta Q_e$, d'énergie $\delta \mathcal{E}_s$, qui en sort entre t et $t + dt$. C'est la variation d'énergie de ce système qui est égale, par définition de la puissance, à $P dt$, donc :

$$P dt = (\mathcal{E}_{AB}(t + dt) + \delta \mathcal{E}_s) - (\mathcal{E}_{AB}(t) + \delta \mathcal{E}_e) = d\mathcal{E}_{AB} + \delta \mathcal{E}_s - \delta \mathcal{E}_e$$

Les énergies entrante et sortante comportent un terme cinétique qui ne varie pas car à intensité égale, les porteurs de charges ont la même vitesse et d'un terme potentiel produit de la charge par le potentiel électrique, d'où : $\delta \mathcal{E}_s - \delta \mathcal{E}_e = I dt V_B - I dt V_A = I(V_B - V_A) dt = -UI dt$.

Dans le cas linéaire où $U = RI - E$ où R est la résistance et E la force électromotrice, on a :

$$-UI = EI - RI^2$$

D'autre part, la puissance P reçue est somme d'une puissance thermique P_{th} et d'une puissance mécanique due aux forces d'interaction avec l'extérieur autres que celles dues au champ électrostatique (on en a déjà tenu compte via l'énergie potentielle), en pratique la puissance due à l'induction magnétique ; en effet les autres forces sont négligeables (celles liées à la masse) ou intérieures (les forces fictives de diffusion), nous la noterons P_{ind} .

Enfin, la variation $d\mathcal{E}_{AB}$ d'énergie du volume de contrôle se traduit soit par un changement de température qu'on pourra mettre sous la forme $K d\theta$ où θ est la température

et K , la capacité thermique¹⁷, soit par une réaction chimique que l'on pourra mettre sous la forme $\Delta U_\chi d\xi$ où ξ est l'avancement chimique (en moles) et ΔU_χ l'énergie interne de la réaction, soit les deux.

A ce stade, le bilan est, après division par dt :

$$P_{\text{th.}} + P_{\text{ind.}} = K \frac{d\theta}{dt} + \delta U_\chi \frac{d\xi}{dt} + E I - R I^2$$

$E I$ est la puissance apportée par les phénomènes électromoteurs (induction magnétique, effets électrochimiques). Bien qu'il soit rare que même dipôle soit siège d'un phénomène d'induction et d'un phénomène électrochimique, on peut considérer que E est somme de deux termes $E_{\text{ind.}}$ et E_χ dus à ces deux effets. Réécrivons donc la relation en profitant pour déplacer quelques termes (on verra pourquoi juste après), soit

$$P_{\text{ind.}} - \delta U_\chi \frac{d\xi}{dt} + P_{\text{th.}} = K \frac{d\theta}{dt} + E_{\text{ind.}} I + E_\chi I - R I^2$$

Comme $P_{\text{ind.}}$ et $E_{\text{ind.}} I$ sont liés aux mêmes effets d'induction, il est logique de les identifier et comme $-\delta U_\chi \frac{d\xi}{dt}$ et $E_\chi I$ sont liés aux mêmes effets chimiques, il est logique de les identifier. Après simplification et réécriture, il reste

$$K \frac{d\theta}{dt} = R I^2 + P_{\text{th.}}$$

que l'on interprète ainsi : l'échauffement du volume de contrôle est dû à une puissance thermique $R I^2$, appelée l'*effet JOULE* et une puissance thermique reçue du milieu extérieur, en général négative. En régime permanent ($\theta = Cte$), l'effet JOULE est donc intégralement restitué sous forme thermique au milieu extérieur.

Bien sûr, on ne va pas refaire à chaque fois ce raisonnement. Retenons que les phénomènes électromoteurs, liés à l'induction ou à l'électrochimie fournissent au dipôle une puissance $E I$ et qu'il en perd $R I^2$ par effet JOULE. Mais on présente souvent les choses à l'envers ; on dit :

- le courant apporte une puissance $U I$
- les phénomènes électromoteurs apportent la puissance $E I$
- l'effet JOULE dissipe la puissance $R I^2$ dont une partie sert éventuellement à chauffer le dipôle et le reste le milieu extérieur
- le bilan est nul d'où $U + E = R I$ soit $U = R I - E$

Comme vision des choses c'est très sommaire et pas très raisonnable. Comme résumé, c'est parfait.

Dans la formulation :

$$P dt = P_{\text{ind.}} dt + P_{\text{th.}} dt = d\mathcal{E}_{AB} + \delta\mathcal{E}_s - \delta\mathcal{E}_e = d\mathcal{E}_{AB} - U I t$$

17. Si le volume n'est pas constant mais la pression, on remplacera dans $\mathcal{E}$, l'énergie interne par l'enthalpie, soit ici la capacité thermique à volume constant par la capacité thermique à pression constante.

cela revient à faire basculer le dernier terme dans le premier membre et à le considérer comme un vrai apport énergétique baptisé travail électrique, associé à une puissance électrique, soit $P_{\text{elec.}} = U I$ et

$$P_{\text{elec.}} dt + P_{\text{ind.}} dt + P_{\text{th.}} dt = d\mathcal{E}_{AB}$$

On pourra avoir une approche plus fine de ces choses par la thermodynamique des milieux non homogènes¹⁸.

Remarque 1 : pour une résistance morte où $U = R I$, la puissance de l'effet JOULE peut s'écrire sous l'une de ces trois formes $P = U I = R I^2 = \frac{U^2}{R}$ ou encore ces trois autres, en introduisant la conductance $G = 1/R$, $P = U I = G U^2 = \frac{I^2}{G}$.

Remarque 2 : pour une résistance morte donnée, si U augmente, I et P aussi, l'échauffement du conducteur et la température d'équilibre aussi. Or la résistance dépend de la température, surtout ce c'est un semi-conducteur, ceci est donc source de non-linéarité.

4.b Piles électrochimiques.

Il est prudent de relire son cours de thermodynamique.

Revenons par une autre voie, qui enrichira le débat, sur l'énergétique de l'électrochimie.

Soit une pile électrochimique qui débite un courant I ou une pile rechargeable en cours de recharge ou un électrolyseur. Le courant sort par une borne B , donc les électrons y entrent et provoquent la réaction chimique d'oxydo-réduction de schéma $e^- + Ox_B \rightarrow Red_B$ accompagnée de variations d'enthalpie ΔH_B et d'entropie ΔS_B et le courant entre par une borne A donc en sortent des électrons qui sont générés au niveau de l'électrode par une réaction chimique dont le schéma est $Red_A \rightarrow Ox_A + e^-$ accompagnée de variations d'enthalpie $-\Delta H_A$ et d'entropie $-\Delta S_A$ avec une notation négative car la réaction se fait en sens inverse.

Considérons l'évolution de ce système ouvert parcouru par une intensité I pendant un temps t tel qu'il soit entré en A et sorti en B une charge q égale à un faraday, c'est à dire, au signe près la charge d'une mole d'électron, soit $q = \mathcal{N}_A e = \mathcal{F}$. En raisonnant comme au paragraphe précédent en terme de système ouvert, la variation d'enthalpie du système est :

$$\Delta H = \Delta H_B - \Delta H_A - U I t$$

Or cette variation est somme du travail reçu, hormis celui de pression (que l'on supposera constante), donc nul en l'absence de phénomènes d'induction, et de la chaleur reçue Q , d'où en tenant compte de la loi d'OHM :

$$Q = \Delta H_B - \Delta H_A + E I t - R I^2 t$$

18. voir chapitre E-XI.

dans laquelle le ΔH pour A ou B est différence de l'enthalpie à pression imposée entre l'état chimique final à la température finale et l'état chimique initial à la température initiale.

Si l'on fait tendre I vers 0 et t vers l'infini de sorte que $q = I t$ reste égal à $q = \mathcal{F}$, alors $R I^2 t = R I q$ tend vers zéro et $U = V_A - V_B = -E$. Dans ce cas, plus d'effet Joule, plus d'échauffement, donc équilibre thermique avec l'extérieur de température T et réversibilité de la transformation ; les deux principes de la thermodynamique¹⁹ puis la classique combinaison linéaire donnent, en introduisant l'enthalpie libre :

$$\begin{aligned} Q &= \Delta H_B - \Delta H_A + E \mathcal{F} \\ \frac{Q}{T} &= \Delta S_B - \Delta S_A \\ 0 &= (\Delta H_B - T \Delta S_B) - (\Delta H_A - T \Delta S_A) + E \mathcal{F} \\ E &= (V_B - V_A) = -\frac{\Delta G_A}{\mathcal{F}} + \frac{\Delta G_B}{\mathcal{F}} \end{aligned}$$

Où cette fois les ΔG pour A et B correspondent à des états chimiques initial et final à pression et température imposée, qui correspondent aux notations des chimistes.

La f.e.m. et la différence de potentiel apparaissent comme différence de deux termes que les chimistes appellent potentiels d'électrodes, $\Pi_B = -\frac{\Delta G_B}{\mathcal{F}}$ et $\Pi_A = -\frac{\Delta G_A}{\mathcal{F}}$ de sorte que mnémotechniquement $V_A - V_B = \Pi_A - \Pi_B$.

Remarque : si le ΔG correspond à $n e^- + Ox \rightarrow Red$, on aura bien sûr $\Pi = -\frac{\Delta G}{n \mathcal{F}}$

Remarque : dans un électrolyseur, les deux électrodes sont identiques et il se comporte symétriquement, à savoir que sa f.e.m. de valeur absolue E est toujours dans le sens inverse du courant qui le traverse. Si I est positif, sa caractéristique est donc $U = R I + E$, ce qui suppose $U > E$ et si I est négatif, sa caractéristique est $U = R I - E$, ce qui suppose $U < -E$. Si $-E < U < E$, on vérifie expérimentalement la seule possibilité logique, c'est-à-dire $I = 0$. Un électrolyseur est un dipôle non-linéaire dont la caractéristique est formée de trois zones linéaires.

4.c Effets Peltier et Thomson. Thermocouples.

Ici aussi, il est prudent de relire son cours de thermodynamique et en particulier de comparer l'approche de ce paragraphe avec celle de la partie du chapitre E-XI qui traite des mêmes effets.

Considérons un circuit filiforme en boucle formé de deux portions réalisées en deux métaux différents notés A et B et soudées en deux points numérotés 1 et 2 ; une des soudures, entre un point A_2 du métal A et un point B_2 du métal B est à la température T_2

¹⁹. L'entropie des charges entrante et sortante ne dépendent pas du potentiel électrique et sont donc identiques.

et l'autre, entre un point A_1 du métal A et un point B_1 du métal B est à la température T_1 . La figure 9 p. 33 (pour l'instant, faire semblant de ne pas voir ce qui est en bleu) schématise le dispositif et précise le sens positif arbitraire du courant I , le métal A est en rouge et le B en vert.

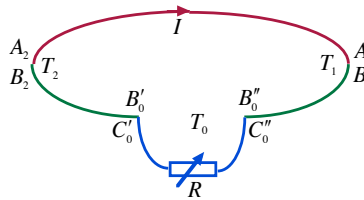


FIGURE 9 – Thermocouple.

La boucle est doublement dissymétrique, d'une part par les températures T_1 et T_2 et d'autre part par les métaux A et B , ce qui fait que les deux sens de parcours de la boucle ne sont pas équivalents et l'on constate expérimentalement le passage d'un courant I non nul. Son origine est simple : la non-uniformité de la température entraîne une diffusion thermique au sein des métaux dont le vecteur essentiel est l'ensemble des charges libres du métal qui sont aussi des porteurs de charges ; il y a donc forcément un couplage entre les propriétés thermiques et les propriétés électriques de cette boucle, appelée *thermocouple*.

La compréhension profonde du phénomène relève de la thermodynamique des états hors d'équilibre (voir note 18 p.31) et nous nous contenterons ici de constatations expérimentales commentées au crible de la thermodynamique classique.

On constate deux effets thermo-électriques. Le premier a lieu au niveau des soudures où l'on constate électriquement une discontinuité de potentiel dépendant de la température que nous noterons respectivement $V_{A_2} - V_{B_2} = V(T_2)$ et $V_{A_1} - V_{B_1} = V(T_1)$ et un échange thermique avec l'extérieur. Les deux phénomènes sont forcément liés. Electriquement, la loi d'OHM donne $-V(T_2) = R_{B_2A_2} I - E_{B_2A_2}$, or $R_{A_2B_2}$ est forcément nul car les deux points sont confondus, d'où $E_{B_2A_2} = V(T_2)$. De même $E_{B_1A_1} = V(T_1)$. Au niveau de la soudure à la température T_2 , cette force électromotrice apporte une puissance $V(T_2) I$ qui ne peut être accumulée car la soudure ponctuelle a une capacité thermique nulle, elle est donc intégralement restituée sous forme de chaleur au milieu extérieur ; il en est de même, en remplaçant I par $-I$ (courant de A vers B et non plus de B vers A) au niveau de l'autre soudure. Cet effet thermo-électrique est l'*effet PELTIER*.

L'autre effet est dû au gradient de température dans les conducteurs métalliques. Soit un élément infiniment petit du conducteur de métal B ; le courant y entre par un point de température T et en sort par un point de température $T + dT$. L'effet thermo-électrique, de la même façon que ci-dessus, introduit une force électromotrice dE_A et d'une puissance $I dE_A$ qui est restituée, en régime permanent dans lequel nous nous plaçons, sous forme thermique au milieu extérieur et qui s'ajoute ici à l'effet JOULE élémentaire. L'expérience s'accorde avec une formule en $dE_A = h_A(T) dT$. Il en va de même pour l'autre métal. Cet effet thermo-électrique est l'*effet THOMSON*.

La f.e.m. totale comptée positivement dans le sens de T est la somme de tous ces effets soit :

$$E(T_1, T_2) = V(T_2) + \int_{T_2}^{T_1} h_A(T) dT - V(T_1) + \int_{T_1}^{T_2} h_B(T) dT = V(T_2) - V(T_1) + \int_{T_1}^{T_2} (h_B - h_A) dT$$

et la puissance thermique fournie à l'extérieur est, compte tenu de l'effet JOULE est :

$$P_{th.} = E I + R I^2$$

Pour utiliser à plein le second principe, plaçons-nous dans un cadre réversible en faisant tendre l'effet JOULE vers zéro en augmentant la résistance pour diminuer l'intensité de sorte que $R I^2$ de degré deux devienne négligeable. A cet effet, on intercale une résistance dans un troisième métal dans un thermostat à T_0 (voir figure 9 p. 33) que sorte qu'il n'y ait pas d'effet THOMSON (température uniforme) et que les deux effets PELTIER aux deux soudures supplémentaires se compensent (même température).

A chacun des échanges thermiques de puissance en $P = V(T) I$ ou $dP = dE I$ pendant une durée τ correspond une chaleur algébriquement reçue en $Q = -P \tau$ ou $\delta Q = -dP \tau$ et pour l'expression du second principe des termes en $\frac{Q}{T}$ ou $\frac{\delta Q}{T}$. La variation d'entropie est nulle puisqu'on se place en régime permanent et le second principe s'exprime avec une égalité car on s'est placé dans les conditions de réversibilité, d'où, après simplification :

$$\frac{\Delta S}{-I \tau} = 0 = \frac{V(T_2)}{T_2} - \frac{V(T_1)}{T_1} + \int_{T_1}^{T_2} \frac{(h_B - h_A)}{T} dT$$

La seule mesure qui soit aisée est celle de $E(T_1, T_2)$. Comme elle dépend de deux températures, on fixe T_1 à une valeur de référence (glace fondante par exemple) et l'on mesure à différentes températures T_2 la f.e.m. et l'on obtient ainsi une formule expérimentale donnant, en renommant les températures de façon adaptée au point de vue, $e(T) = E(T, T_{réf.})$. Regroupons sa définition et l'expression du second principe :

$$\begin{cases} e(T) = V(T) - V(T_{réf.}) + \int_{T_{réf.}}^T [h_B(\tilde{T}) - h_A(\tilde{T})] d\tilde{T} \\ 0 = \frac{V(T)}{T} - \frac{V(T_{réf.})}{T_{réf.}} + \int_{T_{réf.}}^T \frac{[h_B(\tilde{T}) - h_A(\tilde{T})]}{\tilde{T}} d\tilde{T} \end{cases}$$

Dérivons ces deux expressions par rapport à T ; on rappelle que $e(T)$ donc ses dérivées successives sont des données expérimentales accessibles. On arrive à :

$$\begin{cases} \frac{de}{dT} = \frac{dV}{dT} + [h_B(T) - h_A(T)] \\ 0 = \frac{1}{T} \frac{dV}{dT} - \frac{1}{T^2} V + \frac{[h_B(T) - h_A(T)]}{T} \end{cases}$$

Si l'on multiplie la seconde relation par T et qu'on la soustrait à la première, on arrive à : $\frac{de}{dT} = \frac{1}{T} V(T)$ d'où :

$$V(T) = T \frac{de}{dT}$$

Si l'on reporte dans la première relation, on tire $\frac{de}{dT} = \frac{de}{dT} + T \frac{d^2e}{dT^2} + [h_B(T) - h_A(T)]$, d'où :

$$h_B(T) - h_A(T) = -T \frac{d^2e}{dT^2}$$

qui ne permet que de calculer les différences de fonctions h , soit encore les fonction h à une constante commune près.

Sans autres outils thermodynamiques, on ne peut en dire plus mais c'est déjà bien.

5 Réseaux électriques en courant continu.

L'objectif est ici modeste : donner, en les justifiant, les outils nécessaires à l'étude d'un réseau électrique.

5.a Branches, nœuds, mailles.

Nous ne chercherons pas ici une définition rigoureuse qui nécessiterait le recours à la théorie de graphes.

Une branche est un dipôle électrocinétique ou une succession de dipôles en série ; elle est caractérisée par le courant qui le traverse et la différence de potentiel entre ses extrémités.

Un nœud est un point où se raccordent par l'une de leurs extrémités au moins trois dipôles (donc trois branches).

Une maille est une succession de branches se raccordant par leurs extrémités et telle que la dernière se raccorde à la première.

5.b Modélisation d'une branche. Conventions de signe et conventions graphiques.

Soit une branche d'extrémités A et B . On note I_{AB} le courant algébrique qui la traverse quand le sens positif choisi va de A à B et I_{BA} quand le sens positif va de B à A . On note $U_{AB} = V_A - V_B$ et $U_{BA} = V_B - V_A$. La force ou tension électromotrice a été définie par $E_{AB} = \int_A^B \vec{E}_m \cdot d\vec{l}$ et donc par symétrie $E_{BA} = \int_B^A \vec{E}_m \cdot d\vec{l} = -E_{AB}$.

On a vu plus haut que l'on peut écrire $U_{AB} = RI_{AB} - E_{AB}$ (ou symétriquement $U_{BA} = RI_{BA} - E_{BA}$) ; cette écriture est connue sous le nom de *convention récepteur*. Graphiquement, tout revient à considérer la branche comme une résistance morte de valeur R (terme RI_{AB}) et d'un *générateur de tension parfait* sans résistance (terme $-E_{AB}$) mis en série. On dessine les choses comme sur la figure 10 p. 36, où le sens des flèches fait référence aux vecteurs des espaces affines (formellement $V_A - V_B$ est un vecteur d'origine B et d'extrémité A) et à l'addition des vecteurs.

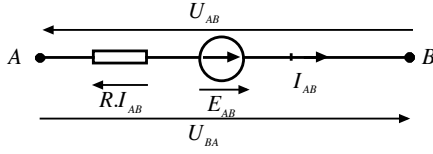


FIGURE 10 – Conventions récepteur et générateur.

La *convention générateur* écrit : $U_{BA} = E_{AB} - R I_{AB}$. Tout cela se fait naturellement avec la gestion visuelle des flèches de la figure.

5.c Modélisation des générateurs.

On vient de voir qu'une branche avec une f.e.m. non nulle E peut être considérée comme un générateur de tension parfait de même f.e.m. aux bornes duquel la différence de potentiel est E quel que soit le courant qui le traverse mis en série avec une résistance R , de sorte qu'en convention générateur on puisse écrire $U_{AB} = R I_{AB} - E_{AB}$.

Si l'on réécrit ainsi les choses : $I_{AB} = \frac{E_{AB}}{R} - \frac{U_{AB}}{R}$, une autre interprétation se dessine : celle d'un *générateur de courant parfait* délivrant un *courant caractéristique* $\eta = \frac{E_{AB}}{R}$ quelle que soit la différence de potentiel à ses bornes dont une partie est dérivée à travers une résistance en parallèle de conductance $G = \frac{1}{R}$ de sorte que le courant résiduel soit $I_{AB} = \eta - G U_{AB}$.

Le tout est résumé sur la figure 11 p. 36

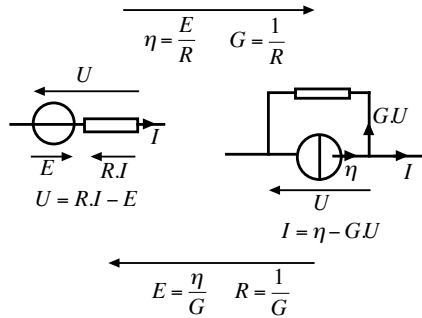


FIGURE 11 – Générateur parfait de tension, d'intensité.

5.d Les outils de base : loi des nœuds, loi des mailles, théorème de Millmann

La loi des nœuds et la loi des mailles sont connues sous le nom de lois de KIRSCHHOFF ; le théorème de MILLMANN est un corollaire de la première.

• **Loi des nœuds.**

Soit un nœud auquel aboutissent n branches numérotées de 1 à n ; notons I_k l'intensité dans la branche k , comptée positivement en direction du nœud. Si la somme de ces intensités était non nulle, il y aurait accumulation de charges au niveau du nœud, ce qui est exclu car nous sommes en régime permanent (courant continu). Donc

$$\sum_1^n I_k = 0$$

• **Loi des mailles.**

Soit une maille formée de la succession de n branches limitées par n nœuds ; la branche k étant limitée par les nœuds k et $k+1$ de potentiels V_k et V_{k+1} , sauf celle de rang n entre les nœuds n et 1 (ce qui revient à dire que $V_{n+1} = V_1$). La somme des $U_k = V_k - V_{k+1}$ va être nulle par télescopage additif, comme disent nos amis mathématiciens. Donc

$$\sum_1^n U_k = 0$$

• **Théorème de Millmann.**

Appliquons la loi des nœuds dans le cas où les n branches sont des résistances mortes de conductivité G_k entre les nœuds k de potentiels V_k et le nœud qui leur est commun dont nous noterons V_0 potentiel. Alors, nous avons successivement :

$$0 = \sum_1^n I_k = \sum_1^n G_k U_k = \sum_1^n G_k (V_k - V_0)$$

$$\sum_1^n G_k V_k = \left(\sum_1^n G_k \right) V_0$$

$$V_0 = \frac{\sum_1^n G_k V_k}{\sum_1^n G_k}$$

L'intérêt de ce théorème est de résoudre un problème de réseau en ne passant que par les potentiels inconnus. Pour une branche avec f.e.m. , il est astucieux de la considérer comme deux branches en série, l'une étant un générateur de tension parfait et l'autre une résistance morte et de considérer leur point commun comme un nœud. Ce théorème est surtout puissant en électronique où la masse est un nœud de potentiel connu et commun à beaucoup de branches ; dans le découpage précédent, on prend soin de brancher le générateur à la masse de sorte que le nœud fictif ainsi créé ait lui aussi un potentiel connu.

5.e Association de résistances : série, parallèle, triangle-étoile.

Pour alléger l'analyse d'un réseau, on peut tenter de le simplifier en remplaçant plusieurs branches par une seule. En pratique, ce n'est possible que si les branches ne comportent que des résistances mortes, c'est à dire pas de tensions électromotrices.

Si une branche comporte plusieurs résistances en série, les résistances s'ajoutent (cf supra).

Si plusieurs branches sont en parallèle, c'est-à-dire ont mêmes extrémités, on peut les remplacer par une branche unique dont la conductance est somme des conductances (cf supra).

Si trois branches forment un triangle, on peut les remplacer par une étoile, ce qui diminue d'une unité le nombre de mailles du réseau. Observons la figure 12 p. 38 ; à gauche un montage en triangle (ou en Δ ou en Π selon la façon de le dessiner) et à droite en étoile (ou en Y ou en T). Il y a trois courants qui entrent en A , B et C de somme nulle (une généralisation de la loi des nœuds, cf supra) ; on en choisit deux indépendants, I_A et I_B et le troisième est $I_A + I_B$ sortant. Il y a trois différences de potentiel de somme nulle (cf supra, loi des mailles) si on les prend dans le même sens giratoire ; on en choisit deux indépendants U_{AC} et U_{BC} (dans ce sens pour la symétrie) et le troisième s'exprime en fonction des deux autres (voir figure).

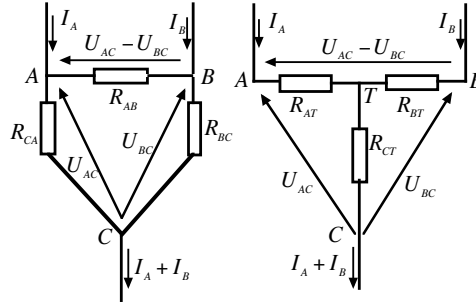


FIGURE 12 – Equivalence triangle-étoile.

On réalise l'équivalence triangle-étoile ou Δ -Y ou Π -T, si quand on choisit les mêmes intensités, on obtient les mêmes différences de potentiel (ou l'inverse).

Dans le montage en triangle, la loi des nœuds appliquée en A et B donne, en passant par les conductances :

$$\begin{cases} I_A = G_{CA} U_{AC} + G_{AB} (U_{AC} - U_{BC}) = (G_{CA} + G_{AB}) U_{AC} - G_{AB} U_{BC} \\ I_B = G_{BC} U_{BC} - G_{AB} (U_{AC} - U_{BC}) = -G_{AB} U_{AC} + (G_{BC} + G_{AB}) U_{BC} \end{cases}$$

Soit de façon matricielle :

$$\begin{pmatrix} I_A \\ I_B \end{pmatrix} = \begin{pmatrix} G_{CA} + G_{AB} & -G_{AB} \\ -G_{AB} & G_{BC} + G_{AB} \end{pmatrix} \begin{pmatrix} U_{AC} \\ U_{BC} \end{pmatrix}$$

Dans le montage en étoile, par simple addition, on tire, en passant par les résistances :

$$\begin{cases} U_{AC} = R_{AT} I_A + R_{CT} (I_A + I_B) = (R_{AT} + R_{CT}) I_A + R_{CT} I_B \\ U_{BC} = R_{BT} I_B + R_{CT} (I_A + I_B) = R_{CT} I_A + (R_{BT} + R_{CT}) I_B \end{cases}$$

Soit de façon matricielle :

$$\begin{pmatrix} U_{AC} \\ U_{BC} \end{pmatrix} = \begin{pmatrix} R_{AT} + R_{CT} & R_{CT} \\ R_{CT} & R_{BT} + R_{CT} \end{pmatrix} \begin{pmatrix} I_A \\ I_B \end{pmatrix}$$

Par simple comparaison des deux écritures matricielles, on en déduit que les deux matrices sont inverses l'une de l'autre. Pour simplifier un réseau, il faut remplacer le triangle par l'étoile donc calculons l'inverse²⁰ de la première matrice, soit :

$$\frac{1}{G_{AB} G_{BC} + G_{BC} G_{CA} + G_{CA} G_{AB}} \begin{pmatrix} G_{BC} + G_{AB} & G_{AB} \\ G_{AB} & G_{CA} + G_{AB} \end{pmatrix}$$

L'identification avec le seconde donne les expressions de R_{CT} , $R_{AT} + R_{CT}$ et $R_{BT} + R_{CT}$, d'où l'on tire celle de R_{AT} et R_{BT} , d'où en fonction des conductances puis des résistances (on n'explicite pas le passage, aisé, de l'un à l'autre) :

$$R_{AT} = \frac{G_{BC}}{G_{AB} G_{BC} + G_{BC} G_{CA} + G_{CA} G_{AB}} = \frac{R_{CA} R_{AB}}{R_{AB} + R_{BC} + R_{CA}}$$

et deux autres expressions obtenues par permutation circulaire sur A , B et C .

On donne parfois à cette relation le nom de *théorème de KENNELLY*.

5.f Résolution matricielle par les courants de mailles.

La méthode que nous allons exposer n'est pas forcément la plus adroite pour résoudre un problème particulier mais elle a l'avantage de permettre une démonstration aisée de théorèmes intéressants.

Elle consiste à diviser le réseau en un certain nombre de mailles linéairement indépendantes et d'affecter à chacune un *courant de maille*, étant entendu que si une branche est commune à deux mailles ou plus, le courant qui la parcourt est somme algébrique des courants des mailles auxquelles elles appartient. Expliquons cela par un exemple simple : un réseau à deux mailles, celui de la figure 13 p. 40.

La maille de gauche y est traversée dans le sens horaire par un courant I_1 et celle de gauche par un courant I_2 dans le sens horaire ; la branche commune est donc traversée de

20. Il est bon de savoir par cœur que l'inverse de $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ est $\frac{1}{a d - b c} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$

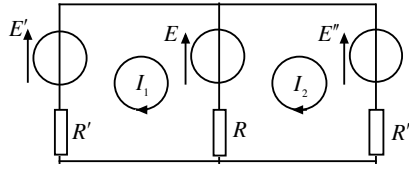


FIGURE 13 – Courant de maille.

haut en bas par $I_2 - I_1$. Si l'on applique la loi des mailles à chacune des deux en tournant dans le sens horaire, on a :

$$\begin{cases} R' I_1 - E' + E + R (I_1 - I_2) = 0 \\ R (I_2 - I_1) - E + E'' + R'' I_2 = 0 \end{cases}$$

soit en réorganisant

$$\begin{cases} (R' + R) I_1 - R I_2 = E' - E \\ -R I_1 + (R + R'') I_2 = E - E'' \end{cases}$$

On peut réécrire cela sous forme matricielle, soit en généralisant à n mailles :

$$\begin{pmatrix} R_{11} & R_{12} & \cdots \\ R_{21} & R_{22} & \cdots \\ \vdots & \vdots & \ddots \end{pmatrix} \begin{pmatrix} I_1 \\ I_2 \\ \vdots \end{pmatrix} = \begin{pmatrix} E_1 \\ E_2 \\ \vdots \end{pmatrix}$$

que l'on peut noter $(R)(I) = (E)$.

La matrice qui apparaît, on peut l'appeler *matrice des résistances* est forcément carrée, au vu de sa construction (une ligne par maille et une colonne par courant de maille) et c'est du reste ce qui justifie sa pertinence.

Détaillons les différents termes qui apparaissent, en laissant le soin au lecteur de vérifier mes assertions (il s'en convaincra en dix secondes mais il me faudrait trois pages) :

- les termes diagonaux de la matrice : R_{11} , par exemple, est la somme des résistances des branches qui constituent la maille numéro 1
- les termes diagonaux de la matrice : $R_{12} = R_{21}$, par exemple, est la résistance de la branche commune aux mailles numéros 1 et 2, comptée positivement ou négativement selon que les courants de ces mailles la traversent dans le même sens ou dans le sens contraire. Au passage remarquons que la matrice est symétrique.
- les termes de la matrice colonne du second membre : E_1 par exemple est la somme des tensions électromotrices des branches qui constituent la maille numéro 1, affectées individuellement d'un signe positif ou négatif selon que sa valeur donnée dans l'énoncé du problème correspond au sens du courant de maille ou non.

La matrice des résistances est inversible. En effet, si elle ne l'était pas, le problème $(R)(I) = (0)$ aurait d'autres solutions que $(I) = 0$, donc il y aurait au moins un courant non nul, donc un effet JOULE, sans aucun générateur, ce qui est énergétiquement impossible. La matrice est inversible, donc le problème est résolu.

Remarque : on peut se demander comment construire un ensemble de mailles linéairement indépendantes. Je propose bêtement l'algorithme suivant : on commence par une première maille arbitraire²¹ puis tant qu'il se trouve au moins une branche n'appartenant à aucune des mailles déjà choisies, on définit une nouvelle maille contenant au moins une de ces branches « vierges ».

5.g Théorème de superposition.

Dans l'exemple précédent associé à la figure 13 p. 40, les courants de mailles étaient solutions de $(R)(I) = (E)$, soit sous forme développée :

$$\begin{pmatrix} R' + R & -R \\ -R & R + R'' \end{pmatrix} \begin{pmatrix} I_1 \\ I_2 \end{pmatrix} = \begin{pmatrix} E' - E \\ E - E'' \end{pmatrix} = E \begin{pmatrix} -1 \\ 1 \end{pmatrix} + E' \begin{pmatrix} 1 \\ 0 \end{pmatrix} + E'' \begin{pmatrix} 0 \\ -1 \end{pmatrix}$$

que l'on peut généraliser sous la forme

$$(R)(I) = \sum_k E_k (U_k)$$

où E_k est la f.e.m. du générateur repéré par l'indice k et les (U_k) des vecteurs colonnes dont le terme de rang n est égal à 0 si le générateur k n'appartient pas à la maille k , +1 s'il lui appartient et est branché dans le sens du courant de maille I_k et -1 s'il est branché dans le sens inverse.

Il en résulte que :

$$(I) = \sum_k E_k (R)^{-1} (U_k)$$

Ce qui est le *théorème de superposition* : Les courants de mailles (et donc tous les courants et toutes les différences de potentiels), sont combinaisons linéaires, avec les f.e.m. comme coefficients, des courants que l'on obtiendrait en supprimant tous les générateurs sauf un, remplacé par un générateur de f.e.m. unité. En version plus souple : somme des courants que l'on obtiendrait en supprimant tous les générateurs sauf un. Bien sûr, il y a dans la somme un terme par générateur.

21. mais rien n'empêche, bien au contraire, qu'elle soit habilement choisie.

5.h Théorèmes de Thévenin et de Norton. Exemples et applications.

• La situation. Le lemme.

Plutôt qu'une démonstration aride avec plein d'indices partout, nous allons donner le principe de la démonstration sur un exemple simple, illustré par la figure 14 p. 42 (le schéma du haut).

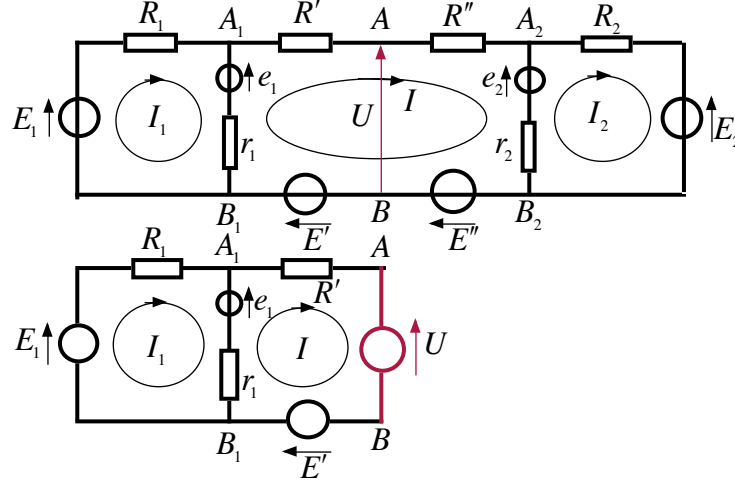


FIGURE 14 – Réseau type.

On peut « lire » le schéma comme deux réseaux, ici réduits à une maille, réunis par deux branches A_1A_2 et B_1B_2 et qui seraient disjoints sans ces branches. Nous appellerons I le courant dans la maille qui contient ces deux branches qui, par nature ne sont communes à aucune autre maille. Chacune de deux branches A_1A_2 et B_1B_2 sont possiblement divisées en deux sous-branches de points milieux A et B et nous appellerons U la différence de potentiel $U = V_A - V_B$. Le reste des notations se lira sur la figure.

Nous nous proposons de démontrer le lemme suivant : si l'on remplace (cf schéma du bas) tout ce qui est à droite des points A et B sur la figure (soit l'un des deux réseaux connectés par les deux branches particulières ainsi que les deux demi-branches du côté du réseau supprimé) par un générateur parfait de tension électromotrice U , le courant I est inchangé.

Supposons résolu le problème posé par la figure du haut, les courants de maille sont connus, d'où les courants dans chaque branche puis, par la loi d'OHM, la différence de potentiel de chaque branche et plus généralement toute différence de potentiel, en particulier U . La loi des mailles pour la maille centrale, de façon brute, est :

$$(V_B - V_{B_1}) + (V_{B_1} - V_{A_1}) + (V_{A_1} - V_A) + (V_A - V_{A_2}) + (V_{A_2} - V_{B_2}) + (V_{B_2} - V_B) = 0$$

puis

$$(V_B - V_{B_1}) + (V_{B_1} - V_{A_1}) + (V_{A_1} - V_A) + (V_A - V_B) = 0$$
$$-E' + r_1(I - I_1) - e_1 + R'I + U = 0$$

qui n'est rien d'autre que la loi des mailles pour la maille de courant I dans le réseau de la figure du bas. La solution de la figure du haut est donc solution de la figure du bas.

• Théorème de Norton.

Résolvons notre problème équivalent par superposition, de façon souple, d'une situation où la f.e.m. U est seule annulée et d'une autre où ce sont toutes les autres f.e.m. qui sont annulées mais pas U .

Si U est annulé, A et B sont court-circuités et le courant I est alors noté I_{cc} (*courant de court-circuit*) ; il se calcule par résolution du réseau (figure du bas) et est donc un cas d'espèce.

Si toutes les autres f.e.m. sont annulés, le réseau à gauche de A et B n'est formé que de résistances mortes dont la conductance équivalente est notée G_e , là encore un cas d'espèce. Le courant I est alors (attention au signe, voir la figure) $-G_e U$.

Par superposition le courant est donc $I = I_{cc} - G_e U$ et l'on y reconnaît la caractéristique d'un générateur modélisé par un générateur de courant parfait de courant caractéristique I_{cc} en parallèle avec une conductance G_e et pour lequel $I = I_{cc} - G_e U$.

Ceci constitue le *théorème de NORTON*. Quand un réseau complexe est formellement scindé en deux sous-réseaux disjoints se connectant en deux points, on peut remplacer l'un des deux par un générateur de courant parfait, dont le courant caractéristique est le courant qui traverserait un fil de court-circuit remplaçant l'autre sous-réseau, mis en parallèle avec sa conductance équivalente, tous générateurs éteints, entre les deux points de connexion.

Remarque 1 : le lecteur a bien sûr deviné que l'on peut procéder de même avec le second sous-réseau.

Remarque 2 : l'interprétation de ce théorème et du suivant est que, quelque compliqué qu'il soit, le comportement d'un réseau peut toujours se ramener à quelque chose de très simple, pourvu toutefois qu'il soit constitué uniquement de dipôles linéaires.

Remarque 3 : si ce n'est pas le cas, on cherchera à isoler le ou les dipôles non linéaires dans un sous-réseau simple, car la division d'un réseau en deux sous-réseaux doit se faire de façon pertinente et non au petit bonheur la chance.

• Théorème de Thévenin.

On a montré plus haut, en adaptant les notations, qu'un générateur de courant parfait, de courant caractéristique I_{cc} , en parallèle avec une conductance G_e est équivalent à un

générateur de tension parfait, de f.e.m. $U_v = \frac{I_{cc}}{G_e}$ en série avec une résistance $R_e = \frac{1}{G_e}$. Le remplacement de G_e par R_e n'est qu'une présentation variante d'une même réalité physique, par contre, on doit donner du sens au remplacement de I_{cc} par U_v . Le lien entre U et I est cette fois mis sous la forme $U = U_v - R_e I$, donc on a $U = U_v$ si $I = 0$ que l'on peut obtenir en ne branchant rien entre A et B ; on obtient alors ce que l'on appelle la *tension à vide* U_v .

On formule ainsi le *théorème de THÉVENIN*. Quand un réseau complexe est formellement scindé en deux sous-réseaux disjoints se connectant en deux points, on peut remplacer l'un des deux par un générateur de tension parfait, dont la f.e.m. est la différence de potentiel entre les deux points de connexion entre lesquels on ne branche rien après suppression l'autre sous-réseau, mis en série avec sa résistance équivalente, tous générateurs éteints, entre les deux points de connexion.

Les remarques figurant après le théorème de NORTON restent valables.

- **Un exemple d'application : le pont de Wheatstone.**

La figure 15 p. 44 est le schéma classique du montage appelé *pont de WHEATSTONE*; R_A représente un galvanomètre (ampèremètre très sensible), qui, à l'équilibre de son aiguille, est équivalent à une résistance morte et d'autre part E et R_G sont la représentation d'un générateur. Les notations sont expliquées par la figure. Il ne s'agit pas ici de trouver la condition d'équilibre du pont ($I = 0$) mais de calculer I pour illustrer ce qui précède.

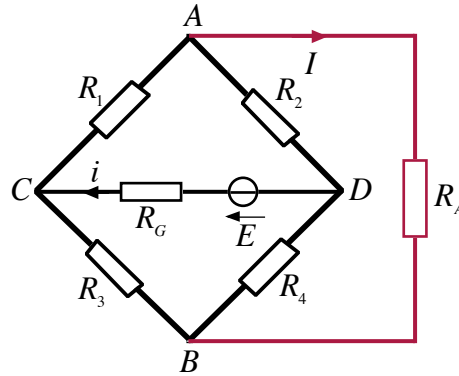


FIGURE 15 – Pont de Wheatstone.

On peut remplacer la partie dessinée en noir (tout le réseau sauf la branche du galvanomètre) par un générateur de tension imparfait débitant dans R_A ; pour calculer sa f.e.m. E_e , on enlève la branche du galvanomètre et l'on calcule $V_A - V_B$. Dans ce contexte, le générateur de f.e.m. débite dans R_G en série avec un montage parallèle comprenant d'un côté $R_1 + R_2$ et de l'autre $R_3 + R_4$, d'où un courant i dont l'expression est, en notant

$R' // R''$ la résistance équivalente à R' et R'' en parallèle :

$$i = \frac{E}{R_G + (R_1 + R_2) // (R_3 + R_4)}$$

d'où l'expression de $V_C - V_D$, calculée avec la résistance équivalente du montage parallèle :

$$U_{CD} = V_C - V_D = (R_1 + R_2) // (R_3 + R_4) i = E \frac{(R_1 + R_2) // (R_3 + R_4)}{R_G + (R_1 + R_2) // (R_3 + R_4)}$$

d'où après des calculs simples non explicités :

$$U_{CD} = V_C - V_D = E \frac{(R_1 + R_2) (R_3 + R_4)}{R_G (R_1 + R_2 + R_3 + R_4) + (R_1 + R_2) (R_3 + R_4)}$$

Toujours dans l'hypothèse $I = 0$, en appelant I_A le courant qui traverse R_1 et R_2 , on a $U_{CD} = V_C - V_D = (R_1 + R_2) I_A$ et $V_A - V_D = R_2 I_A$, d'où, par division membre à membre, $V_A - V_D = \frac{R_2}{R_1 + R_2} U_{CD}$. De même on a $V_B - V_D = \frac{R_4}{R_3 + R_4} U_{CD}$, d'où successivement :

$$E_e = V_A - V_B = \left(\frac{R_2}{R_1 + R_2} - \frac{R_4}{R_3 + R_4} \right) U_{CD}$$

$$E_e = \frac{R_2 (R_3 + R_4) - R_4 (R_1 + R_2)}{(R_1 + R_2) (R_3 + R_4)} U_{CD}$$

$$E_e = \frac{R_2 R_3 - R_4 R_1}{(R_1 + R_2) (R_3 + R_4)} U_{CD}$$

$$E_e = E \frac{R_2 R_3 - R_4 R_1}{R_G (R_1 + R_2 + R_3 + R_4) + (R_1 + R_2) (R_3 + R_4)}$$

Pour calculer la résistance équivalente, on supprime toujours la branche du galvanomètre et le générateur parfait de f.e.m. E et l'on considère le dipôle obtenu de bornes A et B . Une méthode possible est de remplacer le triangle ACD en une étoile de centre O dont les résistances sont (voir ci-dessus) $R_{AO} = \frac{R_1 R_2}{R_G + R_1 + R_2}$, $R_{CO} = \frac{R_1 R_G}{R_G + R_1 + R_2}$ et $R_{DO} = \frac{R_G R_2}{R_G + R_1 + R_2}$ puis de considérer qu'entre A et B , on trouve R_{AO} en série avec un groupement parallèle de $R_{CO} + R_3$ et $R_{DO} + R_4$. On laisse au lecteur le soin de terminer le calcul. Pour en revenir au montage global, il est équivalent à un générateur de f.e.m. E_e et de résistance R_e débitant dans la résistance R_A , d'où $I = \frac{E_e}{R_e + R_A}$.

Est-ce le moyen le plus rapide de calculer I ? Sûrement pas. Les théorèmes de THÉVENIN et NORTON n'ont pas ce but ; ce sont des théorèmes conceptuels : un sous-réseau, si compliqué soit-il, est équivalent à un générateur unique. Voyons cela dans un autre exemple.

- **Un autre exemple d'application : point de fonctionnement.**

Si l'on place un unique dipôle non-linéaire dans un réseau, on peut remplacer tout le reste du réseau par un générateur de THÉVENIN de tension électromotrice E , de résistance R et donc de caractéristique $U = E - RI$. Si l'on trace sur un même graphe cette caractéristique et celle du dipôle non-linéaire, leur intersection, appelée *point de fonctionnement*, fournit à la fois la différence de potentiel U_F aux bornes du dipôle et l'intensité I_F qui le traverse.

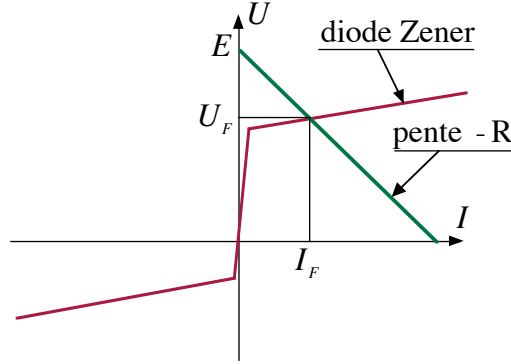


FIGURE 16 – Point de fonctionnement d'une diode Zener.

La figure 16 p. 46 est tracé dans le cas d'une diode ZENER montée en inverse, dont le caractéristique réelle²² a été remplacée par une succession de segments rectilignes ; cette approximation est classique. Sur cet exemple, on voit aisément que si E varie à R égal, la valeur de U_F varie peu (surtout que le graphe exagère la pente de la caractéristique de la diode, dans un souci de lisibilité), ce qui montre le rôle régulateur de tension de ce type de diode.

- **Adaptation d'impédance.**

Dans un réseau, composé de deux sous-réseaux dont l'un sans générateur, on peut remplacer le premier par un générateur de THEVENIN de f.e.m. E et de résistance R et le second par une résistance morte r . Le courant qui en résulte est $I = \frac{E}{R+r}$, la tension aux bornes de r est $U = rI$ et la puissance électrique reçue par r est :

$$\mathcal{P} = UI = rI^2 = \frac{rE^2}{(R+r)^2}$$

Si E et R sont constants, alors si $r = 0$ ou $r = \infty$, on a $\mathcal{P} = 0$; entre les deux, il doit y

22. On se contente ici de l'admettre comme un résultat expérimental.

avoir un maximum. On le recherche ainsi²³ :

$$0 = \frac{d\mathcal{P}}{dr} = \frac{E^2}{(R+r)^2} - \frac{2rE^2}{(R+r)^3} = E^2 \frac{(R+r) - 2r}{(R+r)^3} = E^2 \frac{R-r}{(R+r)^3}$$

Le maximum est obtenu pour $r = R$. Quand on branche un sous-réseau purement résistif à un sous-réseau générateur, on veillera toujours, pour avoir un bon transfert d'énergie, à ce que les deux sous-réseaux aient des résistances équivalentes sinon égales au moins voisines. Il s'agit de l'*adaptation d'impédance*²⁴.

5.i Cas des dipôles commandés.

Certains composants électroniques peuvent être présentés comme des générateurs dont la tension électromotrice E est fonction affine d'un courant d'une branche du réseau, elle-même somme ou différence de deux courants de maille (parfois plus). Dans une loi de maille, correspondant à une ligne de la présentation matricielle du style $\sum_i R_{ij} I_j = E_i$, si par exemple la f.e.m. d'une branche de la maille s'écrit $e = e_0 + \beta (I_1 - I_2)$, cette loi sera réécrite tout naturellement $\sum_i R'_{ij} I_j = E'_i$ où E'_i est E_i dans lequel on a remplacé e par e_0 , où $R'_{i1} = R_{i1} - \beta$, $R'_{i2} = R_{i2} + \beta$ et tous les autres inchangés, simple conséquence du passage des termes en β du second au premier membre. On procède bien sûr de même pour toutes les mailles qui contiennent la branche où est placé le dipôle commandé. On est alors ramené au cas général.

Dans le cas où le dipôle est commandé par une tension, c'est à dire où $e = e_0 + \beta U$ avec U une différence de potentiel d'une branche, on se ramène aisément au cas précédent grâce à la caractéristique $U = f(I)$ de la branche.

5.j Théorème de compensation.

Imaginons que l'on ait résolu un problème de réseau par n'importe quelle méthode ; sa solution est aussi celle de la méthode matricielle des courants de maille de la forme $(R)(I) = (E)$. Imaginons maintenant que la résistance R de la branche commune aux mailles 1 et 2 varie de ΔR , ce qui modifie les termes R_{11} , R_{22} , $R_{12} = R_{21}$ de la matrice ; on doit résoudre un nouveau problème de la forme $(R')(I') = (E)$, ce qui oblige a priori à reprendre la résolution à zéro. C'est frustrant.

Si ΔR est petit devant R , on se doute que tous les I'_k seront proches de I_k . La loi des mailles appliquée à la maille 1 et en supposant pour fixer les idées que les courants I_1 et I_2 y aient leurs sens positifs opposés, contient dans le nouveau problème et au titre de cette branche le terme :

$$U = (R + \Delta R)(I'_1 - I'_2) = R(I'_1 - I'_2) + \Delta R(I_1 - I_2) + \Delta R[(I'_1 - I_1) - (I'_2 - I_2)]$$

23. Il est plus simple de considérer $\frac{r}{(R+r)^2}$ comme produit de r par $\frac{1}{(R+r)^2}$

24. ici de résistance, mais ça se généralise dans toutes sortes de contextes et le terme « impédance » est plus général.

où l'on peut négliger le dernier terme qui est un correctif d'ordre 2 ; on arrive ainsi à :

$$U \approx R(I'_1 - I'_2) + \Delta R(I_1 - I_2)$$

Pour la maille 2, ce terme figurerait avec le signe opposé

Par analogie avec $U = R(I'_1 - I'_2) - E$, cela revient formellement à remplacer la résistance additionnelle par un générateur de f.e.m. $e_1 = -\Delta R(I_1 - I_2)$, connue puisque le premier problème est réputé résolu. Dans la branche 2, ici orientée dans l'autre sens, on a la f.e.m. fictive additionnelle opposée, soit $e_2 = -e_1$. Ainsi, on est amené formellement, avec une matrice inchangée, à $(R)(I') = (E) + (e)$ où (e) a pour termes e_1, e_2 puis des zéros. Le théorème de superposition permet d'ajouter à la solution connue la solution du problème $(R)(I') = (e)$, par toute méthode que l'on veut. En général, on gagne ainsi beaucoup de temps.

Le *théorème de compensation* que je vous laisse formuler à votre guise affirme et détaille le remplacement d'une résistance additionnelle par un générateur additionnel.

5.k Présentation en termes de quadripôle.

Considérons, en généralisant la démarche utilisée plus haut, un réseau scindé en trois sous-réseaux, celui « de gauche » réuni par deux branches à celui « du milieu » et celui-ci par deux autres branches à celui « de droite ». Le lemme démontré ci-dessus permet de remplacer le réseau de gauche par un générateur de tension parfait de f.e.m. égale à la différence de potentiel entre les deux points de jonction, appelée tension d'entrée U_E et celui de droite par un générateur de f.e.m. égal à la tension de sortie U_S . Les courants des mailles passant par les points de jonction sont les courants d'entrée I_E et de sortie I_S . Le tout est schématisé dans la figure 17 p. 48 où le réseau du milieu est représenté par une « boîte noire ».



FIGURE 17 – Quadripôle.

Dans le cas où le réseau du milieu ne comporte pas de générateurs (ou uniquement des générateurs commandés), la résolution matricielle permet de calculer les intensités en fonctions des deux f.e.m. fictives pour arriver à une solution de la forme :

$$\begin{cases} I_E = G_{11} U_E + G_{12} (-U_S) = G_{11} U_E - G_{12} U_S \\ I_S = G_{21} U_E + G_{22} (-U_S) = G_{21} U_E - G_{22} U_S \end{cases}$$

où l'on a noté $-U_S$ car I_S traverse U_s à rebrousse-poil et où $G_{12} = G_{21}$ car l'inverse d'une matrice symétrique est symétrique, mais cela importe peu ici ²⁵.

Il est plus intéressant de lier le couple (U_S, I_S) supposé inconnu au couple (U_E, I_E) supposé connu ; il s'agit d'une présentation en terme de *quadripôle électrocinétique* (l'usage hésite entre quadripôle et quadrupôle).

La première relation donne directement :

$$U_S = \frac{G_{11}}{G_{12}} U_E - \frac{1}{G_{12}} I_E$$

que l'on reporte dans la seconde pour obtenir :

$$I_S = \left(G_{21} - \frac{G_{11} G_{22}}{G_{12}} \right) U_E + \frac{G_{22}}{G_{12}} I_E$$

soit sous forme matricielle :

$$\begin{pmatrix} U_S \\ I_S \end{pmatrix} = \begin{pmatrix} \frac{G_{11}}{G_{12}} & -\frac{1}{G_{12}} \\ \frac{G_{12} G_{21} - G_{11} G_{22}}{G_{12}} & \frac{G_{22}}{G_{12}} \end{pmatrix} \begin{pmatrix} U_E \\ I_E \end{pmatrix}$$

ou sous forme condensée $(S) = (T) (E)$ où (T) s'appelle *matrice de transfert*.

Cette présentation est d'une remarquable efficacité lorsque l'on peut analyser un réseau complexe comme une succession de quadripôles, l'entrée de l'un étant la sortie du précédent. Avec $(S_k) = (T_k) (E_k) = (T_k) (S_{k-1})$ pour tout k , l'on tire :

$$(S_n) = (T_n) (T_{n-1}) \cdots (T_2) (T_1) (E_1)$$

c'est à dire que les matrices de transfert se multiplient, en sens rétrograde comme il se doit.

Remarque : s'il y a des générateurs dans le quadripôle, le théorème de superposition permet de se convaincre qu'au lieu d'avoir une relation linéaire entre entrée et sortie, on a cette fois une relation affine, c'est-à-dire avec des constantes en plus, ce qui est tout à fait gérable.

25. Cette symétrie permet quand même de montrer que la matrice de transfert définie juste en-dessous a un déterminant égal à 1, donc des valeurs propres de produit égal à 1, donc, dans le cas où elles sont imaginaires, complexes conjuguées. Dans le cas d'une cascade de quadripôles identiques, on pourra en déduire bien des choses intéressantes, surtout en régime sinusoïdal. Voir développements en physique vibratoire (chapitres D-II et D-IV).

6 Réseaux électriques en régime variable.

6.a Indications sur les régimes quasi-stationnaires.

Nous nous plaçons ici dans le cadre de l'approximation des régimes quasi-stationnaires²⁶. Les seules choses qu'il nous faille savoir ici sont d'une part que, dans un laboratoire, cette approximation est valable grosso modo pour des phénomènes périodiques de fréquence inférieure à un mégahertz et d'autre part que dans ce cas, l'intensité d'un courant est, à un instant donné, la même en tout point d'une même branche de réseau, exactement comme en régime stationnaire. On rappelle (cf paragraphe 1.f p. 10) que dans ce cadre la loi d'OHM est toujours valable.

6.b Les nouveaux personnages : condensateurs et bobines.

• Condensateurs.

En électrostatique, un condensateur est formé de deux conducteurs séparés par le vide ou un isolant dont deux faces en regard à faible distance l'une de l'autre portent des charges opposées et dont les autres faces portent des charges négligeables. Si l'on appelle A et B les deux conducteurs, q la charge du conducteur A , celle-ci est proportionnelle à la différence de potentiel $U = V_A - V_B$ et la constante de proportionnalité, notée C , est la *capacité* du condensateur.

Dans un usage normal²⁷, cette relation reste valide en régime variable.

Si le condensateur est placé dans une branche dont le courant I est compté positivement de A vers B , le courant véhicule pendant une durée dt une charge $I dt$; dans la portion de branche en amont de A , cette charge $I dt$ vient s'ajouter à q et dans la portion de branche en aval de B , elle se soustrait à $-q$. On a donc $dq = I dt$ soit $\frac{dq}{dt} = I$; or on a aussi $q = C U$. La règle du jeu pour un condensateur est donc :

$$I = C \frac{dU}{dt}$$

Remarque : A géométrie égale, on passe de la capacité C_0 d'un condensateur à vide à la capacité C d'un condensateur à isolant en multipliant C_0 par la *permittivité* ε_r de l'isolant. Malheureusement, en régime variable, ce n'est pas aussi simple. Dans le cas sinusoïdal, ε_r varie avec la fréquence et est en général complexe. Bien entendu, les condensateurs de bonne qualité sont réalisés avec un isolant pour lequel cela a des conséquences marginales.

26. voir chapitre C-VIII relatif aux équations de Maxwell.

27. Le théorème de Gauss reste valable mais le champ n'est plus le gradient du potentiel. Si l'on plaçait le condensateur dans un champ magnétique rapidement variable, il faudrait se creuser la tête.

- **Bobines.**

La théorie en sera faite dans le chapitre C-VII sur l'induction. Admettons ici, comme résultat expérimental que pour une bobine parcourue par un courant I sous une différence de potentiel U , il existe un *coefficient d'auto-induction* (ou *inductance*), noté L et tel que :

$$U = L \frac{dI}{dt}$$

Remarque : le cas des bobines en interaction par un *coefficient de mutuelle induction* ne sera traité que dans le chapitre C-VII sur l'induction.

6.c Régime sinusoïdal. Notation complexe. Impédances.

S'il y a des générateurs de tension sinusoïdaux dans un circuit ou réseau électrique, même de fréquences différentes, le théorème de superposition permettra de les traiter un à un.

En régime sinusoïdal de pulsation ω , on utilisera bien sûr les notations complexes²⁸. A toute fonction $f(t) = f_m \cos(\omega t + \varphi)$, on associe la fonction $f^*(t) = f_m e^{j(\omega t + \varphi)}$ ou le complexe $f^* = f_m e^{j\varphi}$ selon l'humeur ; on rappelant que le physicien appelle $\pm j$ les racines de -1 et non $\pm i$ comme les mathématiciens. L'intérêt est qu'à $\frac{df}{dt}$ correspond $j\omega f^*$.

Aux trois dipôles de base dont les caractéristiques sont $U = RI$ (résistance), $U = L \frac{dI}{dt}$ (bobine) et $I = C \frac{dU}{dt}$ (condensateur), les notations complexes associent des caractéristiques linéaires $U^* = RI^*$, $U^* = jL\omega I^*$ et $I^* = jC\omega U^*$ soit $U^* = \frac{1}{jC\omega} I^*$.

On généralise la notion de résistance (inverse de la conductance) par la notion d'*impédance* (inverse de l'*admittance*), traditionnellement notée Z (inverse A). Les impédances d'un résistance, d'une bobine et d'un condensateur sont respectivement R , $jL\omega$ et $\frac{1}{jC\omega}$.

Puisque nous sommes formellement revenus à une écriture linéaire, tout ce qui a été dit sur l'étude des réseaux reste valable. Nous ne développerons par ici mais en mécanique vibratoire et ondulatoire pour exploiter les nombreuses analogies.

6.d Puissance en régime sinusoïdal.

- **Puissance moyenne.**

Par contre la méthode des notations complexes est par nature incompatible avec le produit, en particulier pour toute approche énergétique.

Rappelons le résultat classique en le redémontrant. Soit deux fonctions sinusoïdales de même pulsation $f(t) = f_m \cos(\omega t + \varphi)$ et $g(t) = g_m \cos(\omega t + \psi)$, soit en notation complexe

28. On suppose la notion acquise. Sinon, voir le chapitre D-I du cours de physique vibratoire et ondulatoire.

$f^*(t) = f_m e^{j(\omega t + \varphi)}$ et $g^*(t) = f_m e^{j(\omega t + \psi)}$ ou $f^* = f_m e^{j\varphi}$ et $g^* = g_m e^{j\psi}$. Linéarisons leur produit et remarquons au passage qu'aucun des termes obtenus est sinusoïdal de pulsation ω :

$$f(t) g(t) = f_m g_m \cos(\omega t + \varphi) \cos(\omega t + \psi)$$

$$f(t) g(t) = \frac{1}{2} f_m g_m \cos(\varphi - \psi) + \frac{1}{2} f_m g_m \cos(2\omega t + \varphi + \psi)$$

Ce qui est intéressant, c'est la moyenne temporelle (sur une période ou sur une durée très grande vis-à-vis de la période, ça revient au même) qui est :

$$\langle f g \rangle = \frac{1}{2} f_m g_m \cos(\varphi - \psi)$$

• Mise en place d'un formalisme.

Pour tenter de faire apparaître les notations complexes dans ce résultat, on peut remarquer que, quelle que soit celle des deux définitions choisies et en notant le conjugué d'un complexe par une barre au-dessus, on a :

$$f^* \bar{g}^* = f_m g_m e^{j(\varphi - \psi)} \quad \text{et} \quad \bar{f}^* g^* = f_m g_m e^{j(-\varphi + \psi)}$$

d'où, en faisant intervenir la partie réelle (symbole $\Re$) :

$$\langle f g \rangle = \frac{1}{2} \Re(f^* \bar{g}^*) = \frac{1}{2} \Re(\bar{f}^* g^*)$$

En particulier, la puissance moyenne fournie dans un dipôle par une intensité $I(t)$ sous une différence $U(t)$ est :

$$\langle P \rangle = \langle U I \rangle = \frac{1}{2} \Re(U^* \bar{I}^*) = \frac{1}{2} \Re(\bar{U}^* I^*)$$

Remarque 1 : il est courant d'escamoter le facteur $\frac{1}{2}$ en définissant pour chaque fonction f une valeur efficace f_e liée à sa valeur maximale f_m par $f_e = \frac{f_m}{\sqrt{2}}$. Comme cette définition ne se généralise pas à d'autres formes de signaux périodiques (triangles, créneaux, etc.), je n'en suis pas fanatique.

Remarque 2 : On définit parfois la puissance complexe par $P^* = \frac{1}{2} U^* \bar{I}^*$ (ou encore $P^* = \frac{1}{2} \bar{U}^* I^*$, peu importe, pourvu que l'on se mette d'accord) et l'on appelle alors puissance réactive sa partie imaginaire. Comme ça n'a pas de signification physique et aucune application théorique mais uniquement des applications pratiques comme le théorème de BOUCHEROT que je me contente de citer, je n'abstiens de tout développement.

- **Adaptation d'impédance en régime sinusoïdal.**

Revenons sur un problème étudié un peu plus haut (en fin de paragraphe 5.h commençant p. 42) en régime permanent. Un générateur de f.e.m. sinusoïdale $E(t)$ et d'impédance Z débite dans une impédance z ; en notation complexe, l'intensité est $I^* = \frac{E^*}{Z+z}$ et la différence de potentiel aux bornes de l'impédance z est $U^* = z I^*$. La puissance moyenne apportée à z est :

$$\langle P \rangle = \frac{1}{2} \Re(U^* \bar{I}^*) = \frac{1}{2} \Re \left(z \frac{E^*}{Z+z} \frac{\bar{E}^*}{\bar{Z}+\bar{z}} \right) = \frac{1}{2} \Re \left(z \frac{|E^*|^2}{|Z+z|^2} \right) = \frac{1}{2} \Re(z) \frac{|E^*|^2}{|Z+z|^2}$$

Notons $|E^*| = E_m = \sqrt{2} E_e$, $Z = R + j B$ et $z = r + j b$, on a alors :

$$\langle P \rangle = \frac{r E_e^2}{(R+r)^2 + (B+b)^2}$$

Reprenons la problématique de trouver l'impédance z qui rende $\langle P \rangle$ maximum pour E_e et Z donnés. Procédons en deux temps. On commence par fixer r et chercher la valeur de b qui optimise $\langle P \rangle$, c'est-à-dire ici qui minimise le dénominateur ; pas besoin de dériver pour constater que cette valeur est $b = -B$ et qu'alors $\langle P \rangle = \frac{r E_e^2}{(R+r)^2}$; on retrouve alors formellement la situation du régime permanent ce qui dispense de refaire les calculs : le maximum est obtenu pour $r = R$, donc $z = r + j b = R - j B = \bar{Z}$.

La charge est adaptée au générateur si elle est le conjugué de l'impédance du générateur. Par exemple et en bonne approximation, les générateurs sinusoïdaux ont un comportement de bobine et une impédance en $R + j L \omega$; l'adaptation d'impédance imposera, si la charge est une résistance r de lui mettre en série un condensateur pour faire passer son impédance à $r + \frac{1}{j C \omega}$ en choisissant C de sorte que l'on ait $L C \omega^2 = 1$.

On remarque sur ce dernier exemple que l'adaptation d'impédance ne peut se faire que pour une fréquence donnée.

6.e Régime transitoire. Transformation de Laplace.

- **Le contexte.**

On s'intéresse souvent à ce type de situation relative aux circuits ou réseaux électrocinétiques : avant un instant initial noté $t = 0$, les tensions électromotrices sont nulles depuis si longtemps que les courants et les charges des condensateurs sont nulles. Au delà de $t = 0$ les tensions électromotrices sont des fonctions non nulles du temps. On se contente trop souvent d'aborder ce genre de situation par le biais d'équations différentielles alors que l'utilisation de la transformation de LAPLACE permet d'utiliser à plein les formalismes mis en place en régime permanent ou sinusoïdal. Voyons cela !

- **La transformation de Laplace.**

Soit une fonction f du temps, nulle jusque $t = 0$, bornée après $t = 0$. A tout réel p positif ou nul, voire à tout complexe p de partie réelle positive ou nulle, on associe à f , sa transformée de LAPLACE notée²⁹ F définie par :

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

Cette transformation est manifestement linéaire.

Par exemple la fonction échelon h , dite aussi fonction de HEAVISIDE, nulle jusque $t = 0$ et égale à l'unité au delà a pour transformée :

$$H(p) = \int_0^{\infty} e^{-pt} dt = \left[-\frac{1}{p} e^{-pt} \right]_0^{\infty} = \frac{1}{p}$$

- **Transformée d'une dérivée.**

Soit f une fonction et $g = f'$ sa dérivée ; la transformée de LAPLACE de sa dérivée va être évaluée grâce à une intégration par parties, soit puisque $f(0) = 0$:

$$\begin{aligned} G(p) &= \int_0^{\infty} f'(t) e^{-pt} dt = \int_0^{\infty} e^{-pt} df = [e^{-pt} f]_0^{\infty} - \int_0^{\infty} f(t) d(e^{-pt}) = \dots \\ &\dots = -f(0) + p \int_0^{\infty} f(t) e^{-pt} dt = p F(p) \end{aligned}$$

- **Application à l'électrocinétique.**

Pour les trois composants de base, résistance, bobine et condensateurs, le lien entre intensité et différence de potentiel est respectivement $u(t) = R i(t)$, $u(t) = L \frac{di}{dt}$ et enfin $i(t) = C \frac{du}{dt}$. On en déduit qu'après une transformation de LAPLACE, on a respectivement $U(p) = R I(p)$, $U(p) = L p I(p)$ et $I(p) = C p U(p)$ ou encore $U(p) = \frac{1}{C p} I(p)$, ce qui permet de se ramener à une formation linéaire introduisant des *impédances opérationnelles* dont l'expression est respectivement R , $L p$ et $\frac{1}{C p}$. Tout ce qui a été dit sur l'étude des réseaux reste donc valable. On remarque que le remplacement de p par $j\omega$ redonne les impédances du régime sinusoïdal.

Par exemple, si l'on branche à $t = 0$ un générateur de f.e.m. constante E , soit avec la fonction échelon h , $E(t) = E h(t)$, un circuit R-L-C, on aura, après transformation :

$$E H(p) = \frac{E}{p} = \left(R + L p + \frac{1}{C p} \right) I(p)$$

29. Nous noterons dans ce paragraphe les fonctions du temps avec une minuscule et leurs transformées de Laplace avec la majuscule associée.

d'où

$$I(p) = \frac{C E}{1 + R p + L p^2}$$

Et puis... c'est là que c'est plus difficile qu'en régime sinusoïdal. Il faut en déduire l'intensité $i(t)$ à partir de sa transformée. A moins de maîtriser la théorie des fonctions d'une variable complexe, on procède, un peu comme pour l'intégration où l'on cherche à reconnaître la dérivée d'une fonction connue, en cherchant à retrouver la transformée d'une fonction connue.

• **Transformées des fonctions habituellement rencontrées.**

Soit une fonction f définie par $f(t) = e^{-at} e^{-j\omega t} = e^{-(a+j\omega)t}$ où a est positif ou nul, soit en notant $\alpha = a + j\omega$, complexe de partie réelle positive ou nulle, $f(t) = e^{-\alpha t}$, on a alors :

$$F(p) = \int_0^\infty e^{-(p+\alpha)t} dt = \dots = \frac{1}{p + \alpha}$$

On remarque au passage que la fonction échelon est décrite par $\alpha = 0$.

Soit une fonction f_n définie par $f_n(t) = t^n e^{-\alpha t}$ où n est un entier strictement positif et α un complexe de partie réelle positive ou nulle, on a alors en intégrant par parties :

$$\begin{aligned} F_n(p) &= \int_0^\infty t^n e^{-(p+\alpha)t} dt = \dots \\ &= \left[-\frac{1}{p + \alpha} t^n e^{-(p+\alpha)t} \right]_0^\infty + \frac{n}{p + \alpha} \int_0^\infty t^{n-1} e^{-(p+\alpha)t} dt = \frac{n}{p + \alpha} F_{n-1}(p) \end{aligned}$$

d'où par récurrence et puisque la première étude traitait tacitement le cas $n = 0$:

$$F_n(p) = \frac{n!}{(p + \alpha)^{n+1}}$$

• **Résolution pratique d'un problème.**

Par analogie avec $I(p) = \frac{C E}{1 + R p + L p^2}$ et avec des générateurs dont la f.e.m. est l'une des fonctions ci-dessus (dans la pratique des constantes ou des sinusoïdes), on conçoit que l'on trouve $I(p)$ sous forme d'une fraction rationnelle de p . Si on la décompose en éléments simples sur le corps des complexes, on arrive à une combinaison linéaire de fractions en $\frac{1}{(p+\alpha)^m}$. En comparant avec le résultat des transformées de fonctions usuelle, on en déduit que $i(t)$ est la combinaison linéaire de mêmes coefficients des fonctions $\frac{1}{(m-1)!} t^{m-1} e^{-\alpha t}$.

Bien sûr, encore faut-il que la décomposition en éléments simples soit faisable. Mais si elle ne l'est pas, la résolution par équation différentielle linéaire ne le sera pas non plus ; on sera bloqué dans les deux cas par la recherche de racines d'un même polynôme de degré trop élevé.