

Cours de Génie Électrique

G. CHAGNON



Cours de Génie Electrique

G. CHAGNON

Table des matières

Introduction	11
1 Quelques mathématiques...	12
1.1 Généralités sur les signaux	12
1.1.1 Introduction	12
1.1.2 Les classes de signaux	12
1.1.2.1 Temps continu et temps discret	12
1.1.2.2 Valeurs continues et valeurs discrètes	13
1.1.2.3 Période, fréquence	14
1.1.3 Energie, puissance	14
1.1.3.1 Définitions	14
1.1.3.2 Remarques	14
1.2 La Transformée de Fourier	15
1.2.1 Généralités	15
1.2.1.1 Introduction	15
1.2.1.2 Définitions	15
1.2.2 Propriétés	16
1.2.2.1 Linéarité	16
1.2.2.2 Décalage en temps/fréquence	16
1.2.2.3 Dérivation	17
1.2.2.4 Dilatation en temps/fréquence	17
1.2.2.5 Conjugaison complexe	18
1.2.2.6 Convolution	18
1.2.3 Représentation de Fourier des signaux d'énergie infinie	19
1.2.3.1 Impulsion de Dirac	19
1.2.3.2 Spectre des signaux périodiques	21
1.2.3.3 Cas particulier : peigne de Dirac	22
1.3 Notion de filtre linéaire	24
1.3.1 Linéarité	24
1.3.2 Invariance	24
1.3.3 Fonction de transfert	25
2 Généralités	27
2.1 Le circuit électrique	27
2.1.1 Circuits électriques	27
2.1.2 Courant, tension, puissance	27
2.1.2.1 Courant électrique	27
2.1.2.2 Différence de potentiel	28
2.1.2.3 Energie, puissance	29
2.1.2.4 Conventions générateur/récepteur	29
2.1.3 Lois de Kirchhoff	30
2.1.3.1 Loi des nœuds	30
2.1.3.2 Loi des mailles	30
2.2 Dipôles électriques	31
2.2.1 Le résistor	31
2.2.1.1 L'effet résistif	31
2.2.1.2 Loi d'Ohm	31

2.2.1.3	Aspect énergétique	32
2.2.1.4	Associations de résistors	32
2.2.2	La bobine	33
2.2.2.1	Les effets inductif et auto-inductif	33
2.2.2.2	Caractéristique tension/courant d'une bobine	33
2.2.2.3	Aspect énergétique	34
2.2.3	Le condensateur	34
2.2.3.1	L'effet capacitif	34
2.2.3.2	Caractéristique tension/courant d'un condensateur	34
2.2.3.3	Aspect énergétique	34
2.3	Régime sinusoïdal, ou <i>harmonique</i>	35
2.3.1	Définitions	35
2.3.2	Puissance en régime sinusoïdal	35
2.3.2.1	Puissance en régime périodique	35
2.3.2.2	Puissance instantanée en régime sinusoïdal	35
2.3.2.3	Puissance moyenne en régime sinusoïdal	36
2.3.3	Représentation complexe d'un signal harmonique	36
2.3.4	Impédances	37
2.3.4.1	Rappel : caractéristiques tension/courant	37
2.3.4.2	Impédance complexe	37
2.3.4.3	Associations d'impédances	38
2.4	Spectre et fonction de transfert	38
2.4.1	Spectre d'un signal	38
2.4.1.1	Introduction	38
2.4.1.2	Signaux multipériodiques et apériodiques	39
2.4.2	Fonction de transfert	40
3	Du semi-conducteur aux transistors	42
3.1	Les semi-conducteurs	42
3.1.1	Semi-conducteurs intrinsèques	42
3.1.1.1	Réseau cristallin	42
3.1.1.2	Définitions	42
3.1.1.3	Exemples	43
3.1.2	Semi-conducteurs extrinsèques de type <i>n</i>	43
3.1.2.1	Réseau cristallin	43
3.1.2.2	Définitions	43
3.1.2.3	Modèle	43
3.1.3	Semi-conducteurs extrinsèques de type <i>p</i>	44
3.1.3.1	Réseau cristallin	44
3.1.3.2	Définition	44
3.1.3.3	Modèle	44
3.2	La jonction PN	45
3.2.1	Introduction	45
3.2.2	Description	45
3.2.3	Définitions	46
3.2.4	Barrière de potentiel	46
3.2.5	Caractéristique électrique	47
3.2.5.1	Description	47
3.2.5.2	Définitions	48
3.2.5.3	Caractéristique et définitions	48
3.3	Le transistor bipolaire	49
3.3.1	Généralités	49
3.3.1.1	Introduction	49
3.3.1.2	Définitions	49
3.3.1.3	Hypothèse	50
3.3.1.4	Transistor au repos	50
3.3.2	Modes de fonctionnement du transistor	51
3.3.2.1	Définitions	51

3.3.2.2	Blocage	51
3.3.2.3	Fonctionnement normal inverse	52
3.3.2.4	Fonctionnement normal inverse	52
3.3.2.5	Saturation	52
3.4	Le transistor MOS	53
3.4.1	Introduction	53
3.4.2	Définitions et principe de fonctionnement	53
4	Systèmes analogiques	55
4.1	Représentation quadripolaire	55
4.1.1	Introduction	55
4.1.2	Matrice de transfert	55
4.1.3	Exemple	56
4.1.4	Impédances d'entrée/sortie	56
4.2	Contreréaction	58
4.2.1	Généralités	58
4.2.1.1	Introduction	58
4.2.1.2	Conventions	59
4.2.1.3	Un exemple d'intérêt du bouclage	59
4.2.2	Un peu de vocabulaire...	60
4.2.2.1	Les signaux	60
4.2.2.2	Les « branches » de la boucle	60
4.2.2.3	Les gains	60
4.2.3	Influence d'une perturbation	61
4.2.4	Exemples de systèmes à contreréaction	61
4.2.4.1	Exemple détaillé : une file de voitures sur l'autoroute	61
4.2.4.2	Autres exemples	62
4.3	Diagramme de Bode ; Gabarit	62
4.3.1	Diagramme de Bode	62
4.3.1.1	Définition	62
4.3.1.2	Exemple	63
4.3.1.3	Les types de filtres	64
4.3.2	Gabarit	65
4.4	Bruit dans les composants	66
4.4.1	Densité spectrale de puissance	66
4.4.2	Les types de bruit	67
4.4.2.1	Bruit thermique	67
4.4.2.2	Bruit de grenaille	68
4.4.2.3	Bruit en $1/f$	68
4.4.2.4	Bruit en créneaux	69
4.4.3	Bruit dans un dipôle	69
4.4.3.1	Température équivalente de bruit	69
4.4.3.2	Rapport de bruit	70
4.4.4	Facteur de bruit	70
4.4.4.1	Définition	70
4.4.4.2	Température de bruit	70
4.4.4.3	Facteur de bruit d'un quadripôle passif	71
4.4.4.4	Théorème de Friiss	71
4.5	Parasites radioélectriques	73
4.5.1	Les sources de parasites	73
4.5.2	Classification des parasites...	73
4.5.2.1	... par leur propagation	74
4.5.2.2	... par leurs effets	74
4.5.3	Les parades	74

5	Systèmes numériques	76
5.1	Introduction	76
5.1.1	Généralités	76
5.1.2	Représentation logique	76
5.1.3	Familles de portes logiques	77
5.2	Logique combinatoire	77
5.2.1	Les opérateurs de base	77
5.2.1.1	Les opérateurs simples	77
5.2.1.2	Propriétés	78
5.2.1.3	Les opérateurs « intermédiaires »	79
5.2.2	Table de Karnaugh	80
5.2.2.1	Principe	80
5.2.2.2	Code binaire réfléchi	80
5.2.2.3	Exemple	80
5.2.3	Quelques fonctions plus évoluées de la logique combinatoire	81
5.2.3.1	Codage, décodage, transcodage	81
5.2.3.2	Multiplexage, démultiplexage	82
5.2.4	Fonctions arithmétiques	83
5.2.4.1	Fonctions logiques	83
5.2.4.2	Fonctions arithmétiques	83
5.2.5	Mémoire morte	84
5.2.6	Le PAL et le PLA	85
5.2.6.1	Le PAL	85
5.2.6.2	Le PLA	85
5.3	Logique séquentielle	86
5.3.1	Généralités	86
5.3.1.1	Le caractère séquentiel	87
5.3.1.2	Systèmes synchrones et asynchrones	87
5.3.1.3	Exemple : bascule RS asynchrone	87
5.3.2	Fonctions importantes de la logique séquentielle	88
5.3.2.1	Bascules simples	88
5.3.2.2	Bascules à fonctionnement en deux temps	90
5.3.2.3	Registres (ensembles de bascules)	91
5.3.3	Synthèse des systèmes séquentiels synchrones	93
5.3.3.1	Registres de bascules	93
5.3.3.2	Compteur programmable	93
5.3.3.3	Unité centrale de contrôle et de traitement (CPU) : microprocesseur	94
5.4	Numérisation de l'information	95
5.4.1	Le théorème de Shannon	95
5.4.1.1	Nécessité de l'échantillonnage	95
5.4.1.2	Exemple : échantillonnage d'une sinusoïde	95
5.4.1.3	Cas général	96
5.4.2	Les échantillonneurs	97
5.4.3	Convertisseur analogique/numérique (CAN)	98
5.4.3.1	Généralités	98
5.4.3.2	Les caractéristiques d'un CAN	98
5.4.3.3	Quelques CAN	98
5.4.4	Convertisseur numérique/analogique (CNA)	100
5.4.4.1	Généralités	100
5.4.4.2	Un exemple de CNA	100
5.4.4.3	Applications des CNA	101
6	Transmission de l'information	102
6.1	Généralités	102
6.1.1	Quelques dates	102
6.1.2	Nécessité d'un conditionnement de l'information	102
6.1.3	Transports simultanés des informations	103
6.1.4	Introduction sur les modulations	103

6.2	Emission d'informations	104
6.2.1	Modulation d'amplitude	104
6.2.1.1	Introduction	104
6.2.1.2	Modulation à porteuse conservée	104
6.2.1.3	Modulation à porteuse supprimée	106
6.2.2	Modulations angulaires	106
6.2.2.1	Introduction	106
6.2.2.2	Aspect temporel	107
6.2.2.3	Aspect fréquentiel de la modulation de fréquence	108
6.3	Réception d'informations	108
6.3.1	Démodulation d'amplitude	109
6.3.1.1	Démodulation incohérente	109
6.3.1.2	Détection synchrone	110
6.3.2	Démodulation angulaire	110
7	Notions d'électrotechnique	112
7.1	Le transformateur monophasé	112
7.1.1	Description, principe	112
7.1.1.1	Nécessité du transformateur	112
7.1.1.2	Principe du transformateur statique	112
7.1.2	Les équations du transformateur	113
7.1.2.1	Conventions algébriques	113
7.1.2.2	Détermination des forces électromotrices induites	114
7.1.2.3	Le transformateur parfait	114
7.2	Systèmes triphasés	115
7.2.1	Définition et classification	115
7.2.1.1	Définition d'un système polyphasé	115
7.2.1.2	Systèmes direct, inverse et homopolaire	116
7.2.1.3	Propriétés des systèmes triphasés équilibrés	116
7.2.2	Associations étoile et triangle	117
7.2.2.1	Position du problème	117
7.2.2.2	Association étoile	118
7.2.2.3	Association triangle	118
7.2.2.4	Bilan	118
7.2.3	Grandeurs de phase et grandeurs de ligne	119
7.2.3.1	Définitions	119
7.2.3.2	Relations dans le montage étoile	119
7.2.3.3	Relations dans le montage triangle	120
7.2.3.4	Bilan	120
7.3	Les machines électriques	121
7.3.1	Généralités	121
7.3.1.1	Mouvement d'un conducteur dans un champ d'induction magnétique uniforme	121
7.3.1.2	Le théorème de Ferraris	122
7.3.2	La machine à courant continu (MCC)	122
7.3.2.1	Principe de la machine	122
7.3.2.2	Réalisation	123
7.3.2.3	Modèle	123
7.3.2.4	Excitation parallèle, excitation série	124
7.3.3	La machine synchrone	125
7.3.4	La machine asynchrone	125
7.4	Conversion d'énergie	126
7.4.1	Introduction	126
7.4.2	Les interrupteurs	127
7.4.2.1	Principe de fonctionnement	127
7.4.2.2	Les types d'interrupteurs	128
7.4.3	Le redressement	128
7.4.3.1	Montages à diodes	128
7.4.3.2	Montage à thyristors	130

7.4.4	L'ondulation	131
7.4.4.1	Généralités	131
7.4.4.2	Exemple d'onduleur	131
A	Table de transformées de Fourier usuelles	133
A.1	Définitions	133
A.2	Table	134
B	Quelques théorèmes généraux de l'électricité	135
B.1	Diviseur de tension, diviseur de courant	135
B.1.1	Diviseur de tension	135
B.1.2	Diviseur de courant	136
B.2	Théorème de Millman	136
B.3	Théorèmes de Thévenin et Norton	137
B.3.1	Théorème de Thévenin	137
B.3.2	Théorème de Norton	138
B.3.3	Relation entre les deux théorèmes	138
C	L'Amplificateur Opérationnel (AO)	139
C.1	L'AO idéal en fonctionnement linéaire	139
C.1.1	Représentation	139
C.1.2	Caractéristiques	139
C.1.3	Exemple : montage amplificateur	140
C.2	L'AO non idéal en fonctionnement linéaire	140
C.2.1	Représentation	140
C.2.2	Caractéristiques	140
C.2.3	Exemples : montage amplificateur	141
C.2.3.1	Gain non infini	141
C.2.3.2	Impédance d'entrée non infinie	141
C.2.3.3	Réponse en fréquence imparfaite	142
C.3	L'AO en fonctionnement non linéaire	142
D	Lignes de transmission	143
D.1	Lignes sans perte	143
D.1.1	Quelques types de lignes	143
D.1.2	Equation de propagation	143
D.1.3	Résolution de l'équation	144
D.2	Interface entre deux lignes	144
D.2.1	Coefficients de réflexion/transmission	144
D.2.2	Cas particuliers	145
D.3	Ligne avec pertes	146
D.3.1	Equation de propagation	146
D.3.2	Résolution de l'équation	146
E	Rappels sur les nombres complexes	147
E.1	Introduction	147
E.2	Représentations algébrique et polaire	147
E.2.1	Représentation algébrique	147
E.2.1.1	Vocabulaire	147
E.2.1.2	Règles de calcul	148
E.2.1.3	Conjugaison	148
E.2.2	Représentation polaire	148
E.2.2.1	Interprétation géométrique	148
E.2.2.2	Représentation polaire	149
E.2.2.3	Règles de calcul et conjugaison	149
E.3	Tables récapitulatives	150
E.3.1	Quelques nombres complexes remarquables	150
E.3.2	Règles de calcul et propriétés	150

F Liste d'abréviations usuelles en électricité

151

Index

153

Introduction

Ce cours a pour but de présenter rapidement le plus large éventail possible des connaissances de base en électronique (analogique et numérique), électrotechnique, traitement et transport du signal.

- Le premier chapitre, à la lecture facultative, introduit la notion de transformée de Fourier et en établit les propriétés mathématiques ;
- Le deuxième chapitre aborde les notions de base des circuits électriques, et présente une approche plus « empirique » des définitions du chapitre précédent ;
- le chapitre suivant expose rapidement les principes de fonctionnement des semi-conducteurs, et présente succinctement transistors bipolaire et MOS ;
- Le quatrième regroupe sous le titre « Systèmes analogiques » des champs aussi divers que les notions de filtrage, de bruit dans les composants, de contre-réaction, etc. ;
- Le chapitre suivant aborde les « systèmes numériques » : circuits de logique combinatoire ou séquentielle et quelques contraintes techniques liées au traitement numérique de l'information ;
- Le sixième chapitre expose brièvement quelques modes de transport de l'information ;
- Le dernier introduit quelques concepts-clefs de l'électrotechnique et de l'électronique de puissance : transformateur, systèmes polyphasés, machines électriques et conversion d'énergie ;

On trouvera en fin de polycopié quelques annexes et un index.

Chapitre 1

Quelques mathématiques...

1.1 Généralités sur les signaux

1.1.1 Introduction

Le concept de signal est extrêmement vaste :

- le relevé en fonction du temps de l'actionnement d'un interrupteur ;
- une émission radiophonique ou télévisée ;
- une photographie...

... sont autant de signaux différents.

Un signal y dépend d'une variable x , sous la forme générale^{1.1} :

$$y = \mathcal{S}(x)$$

avec $y \in \mathbb{C}^m$ et $x \in \mathbb{C}^n$

On se limitera, sauf mention contraire, au cas où $m = 1$ et $n = 1$. Le cas le plus courant est celui où x est en fait le temps t . Nous considérerons donc à l'avenir que les signaux que nous allons étudier sont des fonctions de t .

1.1.2 Les classes de signaux

Les signaux peuvent être classés en diverses catégories :

1.1.2.1 Temps continu et temps discret

- Dans le premier cas, le signal x est une fonction continue du temps t .
Exemple :

1.1. Rappel : $\mathbb{N}$ désigne l'ensemble des *entiers naturels* (0, 1, 2, ... 33, etc.), $\mathbb{Z}$ l'ensemble des *entiers relatifs* (-10, -4, 0, 1, etc.), $\mathbb{Q}$ l'ensemble des *nombres rationnels* (tous les nombres qui peuvent s'écrire sous la forme d'une fraction), $\mathbb{R}$ l'ensemble des *nombres réels* (tous les nombres rationnels, plus des nombres comme $\pi, \sqrt{2}, e$, etc.), et $\mathbb{C}$ l'ensemble des *nombres complexes*.

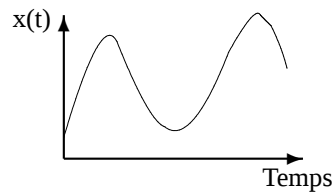


FIG. 1.1 – Signal à temps continu

On notera souvent un tel signal sous la forme $x(t)$, par exemple.

- Dans le deuxième, x n'est défini qu'en un ensemble dénombrable de points.

Exemple :

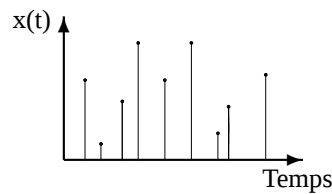


FIG. 1.2 – Signal à temps discret

On notera souvent un tel signal sous la forme $x(n)$, par exemple. Ces points sont souvent répartis à des intervalles de temps réguliers.

1.1.2.2 Valeurs continues et valeurs discrètes

- Dans le premier cas, le signal x peut prendre toutes les valeurs possibles dans un ensemble de définition donné (exemple $] -\infty; +2[$ ou $\mathbb{C}$). Un tel signal est également appelé *analogique* en référence à l'électronique.
- Dans le deuxième, le signal x ne peut prendre qu'un ensemble dénombrable de valeurs. Un tel signal est également appelé *numérique* en référence à l'électronique.

Exemple :

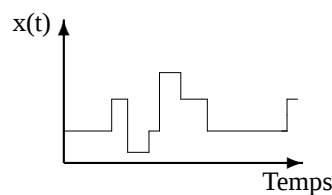


FIG. 1.3 – Signal à valeurs discrètes

Notez que les quatre combinaisons sont possibles : les figures 1.1, 1.2 et 1.3 donnent ainsi respectivement un exemple de signal analogique à temps continu, de signal analogique à temps discret, et de signal numérique à temps continu.

On se limitera dans la suite du chapitre aux signaux analogiques à temps continu. On peut passer d'un signal analogique à un signal numérique par *échantillonnage* : se reporter notamment au chapitre 5.4 pour plus de détails.

1.1.2.3 Période, fréquence

On parle également de **signaux périodiques** : un signal x est dit périodique de période T , ou par anglicisme T -périodique, si pour tout instant t_0 , $x(t_0 + T) = x(t_0)$: le signal se répète, identique à lui-même, au bout d'un intervalle de temps T .

On définit alors sa *fréquence* f ^{1.2} par

$$f = \frac{1}{T}$$

Une fréquence est l'inverse d'un temps, et s'exprime en Hertz (Hz).

1.1.3 Energie, puissance

1.1.3.1 Définitions

- **Energie** : soit un signal $x(t)$ à temps continu, tel que $\int_{-\infty}^{+\infty} |x(t)|^2 dt$ existe et converge. Alors le signal est dit à *énergie finie* et la valeur de cette intégrale est appelée *énergie* du signal x ^{1.3} :

$$E_x \triangleq \int_{-\infty}^{+\infty} |x(t)|^2 dt \quad (1.1)$$

- **Puissance** : pour le même type de signaux, on définit également la *puissance*, notée P_x , par :

$$P_x \triangleq \lim_{\theta \rightarrow +\infty} \frac{1}{2\theta} \int_{-\theta}^{+\theta} |x(t)|^2 dt \quad (1.2)$$

1.1.3.2 Remarques

1. Pour un signal périodique, l'intégrale 1.1 ne converge pas. On peut néanmoins définir la puissance d'un signal x T -périodique par :

$$P_x \triangleq \frac{1}{T} \int_{(T)} |x(t)|^2 dt$$

2. Il existe des signaux ni périodiques, ni d'énergie finie, pour lesquels la puissance ne peut être définie, comme par exemple la « rampe » $x(t) = t$.
3. Il s'agit là de définitions mathématiques. En pratique, un signal mesuré ne l'est *jamais* sur un intervalle de temps infini. Par exemple, on peut commencer à visualiser un signal à un instant qu'on prendra comme origine des temps, et dans ce cas on arrêtera son examen au bout d'un temps T_{obs} . Comme on ne *sait* pas ce que ce signal était avant qu'on ne l'observe, ni ce qu'il deviendra après, il serait présomptueux d'utiliser les bornes $-\infty$ et $+\infty$ dans l'intégrale 1.1, et on se limitera donc à l'écrire sous la forme $\int_0^{T_{obs}} |x(t)|^2 dt$. Remarquons d'ailleurs que cette dernière intégrale converge toujours.

Ce qu'il faut retenir

- Les signaux peuvent être à valeurs discrètes ou continues ; à temps discret ou continu ;

1.2. Parfois notée ν (à prononcer « nu »).

1.3. Le symbole $\triangleq$ désigne une définition.

- La *période* d'un signal est l'intervalle de temps au bout duquel il se répète identique à lui-même ; sa *fréquence* est l'inverse de la période ;
- L'*énergie* d'un signal x à temps continu vaut :

$$E_x \triangleq \int_{-\infty}^{+\infty} |x(t)|^2 dt$$

1.2 La Transformée de Fourier

1.2.1 Généralités

1.2.1.1 Introduction

Cet outil fut introduit pour la première fois par le physicien français Joseph Fourier, pour ses travaux sur la conduction de la chaleur au XIX^e siècle. Depuis lors, il a longuement été développé, et des extensions en ont été proposées.

Il existe plusieurs sortes de Transformées de Fourier, chacune adaptée aux classes de signaux qu'elle analyse, ou au type de signal qu'elle génère. On dénombre ainsi :

- une transformée continue pour les signaux à temps continu : la Transformée de Fourier à proprement parler ;
- une transformée continue pour les signaux à temps discret : la Transformée de Fourier à temps discret ;
- une transformée discrète pour les signaux périodiques à temps continu : le développement en série de Fourier, ou Transformée de Fourier au sens des distributions ;
- une transformée discrète pour les signaux à temps discret : la Transformée de Fourier Discrète.

Nous allons nous limiter, pour l'établissement des propriétés, à la Transformée de Fourier continue des signaux à temps continu.

1.2.1.2 Définitions

1. **Transformée de Fourier** : soit un signal $x(t)$ à temps continu, tel que $\int_{-\infty}^{+\infty} |x(t)| dt$ converge^{1.4}. On définit alors la *transformée de Fourier* de x , notée $X(\nu)$ ou $\text{TF}[x(t)]$, par :

$$X(\nu) \triangleq \int_{-\infty}^{+\infty} x(t) e^{-j2\pi\nu t} dt \quad (1.3)$$

où j est tel que $j^2 = -1$ ^{1.5}. La transformée de Fourier permet de mesurer le « contenu fréquentiel » d'un signal, à savoir la manière dont on peut le décomposer en une somme de sinusoides de fréquences ν .

2. **Transformée de Fourier inverse** : si **de plus** x est à *énergie finie*^{1.6}, cette relation est inversible en

$$x(t) = \int_{-\infty}^{+\infty} X(\nu) e^{+j2\pi\nu t} d\nu \quad (1.4)$$

1.4. On dit alors que « $x \in \mathcal{L}_1$ », ou que x est d'intégrale « absolument convergente ».

1.5. On utilise la lettre j et non i comme en mathématiques pour désigner la racine carrée « classique » de -1 pour éviter la confusion avec le courant i en électricité.

1.6. Rappel : $\int_{-\infty}^{+\infty} |x(t)|^2 dt$ converge.

L'opération correspondante est appelée *transformation de Fourier inverse* : elle permet de revenir au signal temporel x à partir de son contenu fréquentiel.

Ces deux définitions permettent de disposer de deux manières de définir complètement un signal qui satisfait aux conditions d'inversibilité de la transformée de Fourier. On peut le définir :

- soit par sa représentation temporelle ;
- soit par sa représentation fréquentielle.

Ces deux domaines sont souvent appelés « duaux » car leurs variables t et f sont liées par $f = 1/t$.

Spectre : on appelle *spectre de x* le module de la transformée de Fourier de x :

$$S(\nu) = |X(\nu)|$$

1.2.2 Propriétés

Pour toutes les démonstrations suivantes, les signaux x et y sont d'intégrales absolument convergentes. On notera indifféremment $X(\nu)$ ou $TF_x(\nu)$ la transformée de Fourier du signal x .

1.2.2.1 Linéarité

Soient λ et μ deux nombres complexes quelconques. La linéarité de l'équation 1.3 entraîne facilement que :

$$\boxed{TF(\lambda x + \mu y) = \lambda TF(x) + \mu TF(y)} \quad (1.5)$$

1.2.2.2 Décalage en temps/fréquence

Soit t_0 un réel strictement positif. Calculons $TF[x(t - t_0)]$:

$$TF[x(t - t_0)] = \int_{-\infty}^{+\infty} x(t - t_0) e^{-j2\pi\nu t} dt$$

On effectue le changement de variable^{1.7} $u = t - t_0$, et il vient :

$$TF[x(t - t_0)] = \int_{-\infty}^{+\infty} x(u) e^{-j2\pi\nu(u+t_0)} du$$

D'où :

$$TF[x(t - t_0)] = e^{-2j\pi\nu t_0} \int_{-\infty}^{+\infty} x(u) e^{-j2\pi\nu u} du$$

Et donc :

$$\boxed{TF[x(t - t_0)] = e^{-j2\pi\nu t_0} X(\nu)} \quad (1.6)$$

Par symétrie dans les relations 1.3 et 1.4, on obtient de même :

$$\boxed{TF[e^{+j2\pi\nu_0 t} x(t)] = X(\nu - \nu_0)} \quad (1.7)$$

1.7. Il faut vérifier son caractère C^1 , c'est-à-dire continu et dérivable, et également s'assurer qu'il soit strictement monotone.

1.2.2.3 Dérivation

On note $x'(t) = dx/dt$. Alors :

$$\text{TF}[x'(t)] = \int_{-\infty}^{+\infty} x'(t) e^{-j2\pi\nu t} dt$$

On effectue une intégration par parties^{1.8} en intégrant $x'(t)$ et en dérivant l'exponentielle complexe. On obtient alors :

$$\text{TF}[x'(t)] = [x(t)e^{-j2\pi\nu t}]_{-\infty}^{+\infty} + j2\pi\nu \int_{-\infty}^{+\infty} x(t)e^{-j2\pi\nu t} dt$$

Comme x est, physiquement, nécessairement nul à $\pm\infty$ ^{1.9} et que l'exponentielle complexe y reste bornée, le premier terme de la somme devient nul et donc :

$$\boxed{\text{TF}[x'(t)] = j2\pi\nu X(\nu)} \quad (1.8)$$

1.2.2.4 Dilatation en temps/fréquence

Soit λ un réel non nul. Calculons $\text{TF}[x(\lambda t)]$:

$$\text{TF}[x(\lambda t)] = \int_{-\infty}^{+\infty} x(\lambda t) e^{-j2\pi\nu t} dt$$

Effectuons le changement de variable^{1.10} $u = \lambda t$. Deux cas se présentent alors :

– Soit $\lambda > 0$; alors

$$\text{TF}[x(\lambda t)] = \frac{1}{\lambda} \int_{-\infty}^{+\infty} x(u) e^{-j2\pi\frac{\nu}{\lambda}u} du$$

Et donc

$$\boxed{\text{TF}[x(\lambda t)] = \frac{1}{\lambda} X\left(\frac{\nu}{\lambda}\right)} \text{ avec } \lambda > 0 \quad (1.9)$$

– Soit $\lambda < 0$; alors

$$\text{TF}[x(\lambda t)] = \frac{1}{\lambda} \int_{+\infty}^{-\infty} x(u) e^{-j2\pi\frac{\nu}{\lambda}u} du = -\frac{1}{\lambda} \int_{-\infty}^{+\infty} x(u) e^{-j2\pi\frac{\nu}{\lambda}u} du$$

Et donc

$$\boxed{\text{TF}[x(\lambda t)] = -\frac{1}{\lambda} X\left(\frac{\nu}{\lambda}\right)} \text{ avec } \lambda < 0 \quad (1.10)$$

Remarque : si on applique la formule 1.10 en posant $\lambda = -1$, on obtient $\text{TF}[x(-t)] = X(-\nu)$. On en déduit donc que la parité de la Transformée de Fourier est la même que celle du signal original.

1.8. Rappel sur l'intégration par parties : Soient f et g deux fonctions dérivables et définies sur l'intervalle $[a, b]$, dont les dérivées sont continues sur $]a, b[$. Alors :

$$\int_a^b f'(t)g(t)dt = [f(t)g(t)]_a^b - \int_a^b f(t)g'(t)dt$$

avec $[f(t)g(t)]_a^b = f(b)g(b) - f(a)g(a)$

1.9. Un signal observé est toujours nul en $-\infty$ car il n'était alors pas encore observé, et nul en $+\infty$ car il ne l'est plus.

1.10. cf. note 1.7.

1.2.2.5 Conjugaison complexe

On note x^* le signal conjugué de x ^{1.11}. Calculons $\text{TF}[x^*(t)]$:

$$\text{TF}[x^*(t)] = \int_{-\infty}^{+\infty} x^*(t) e^{-j2\pi\nu t} dt = \left(\int_{-\infty}^{+\infty} x(t) e^{+j2\pi\nu t} dt \right)^*$$

Et donc :

$$\boxed{\text{TF}[x^*(t)] = X^*(-\nu)} \quad (1.11)$$

Remarque : si x est un signal réel, alors $x(t) = x^*(t)$, donc $X(\nu) = X^*(-\nu)$. Si de plus x est pair (ou impair), alors $x(t) = x(-t)$ (respectivement $x(t) = -x(-t)$) et en utilisant la remarque du paragraphe 1.2.2.4, il vient $X^*(-\nu) = X^*(\nu)$ (respectivement $X^*(-\nu) = -X^*(\nu)$) d'où $X(\nu) = X^*(\nu)$ et X est réelle (respectivement imaginaire pure). En définitive, on obtient le tableau récapitulatif suivant :

Signal x	Pair	Impair
Réel	X réelle paire	X imaginaire pure impaire
Imaginaire pur	X imaginaire pure paire	X réelle impaire

1.2.2.6 Convolution

Définition : Soient deux signaux x et y à valeurs continues et à temps continu. On définit le *produit de convolution* des deux signaux, ou plus simplement leur *convolution*, par :

$$\boxed{(x * y)(t) \triangleq \int_{-\infty}^{+\infty} x(\theta) y(t - \theta) d\theta} \quad (1.12)$$

On vérifie aisément que $(x * y)(t) = (y * x)(t)$, c'est-à-dire que la convolution est *commutative*, et donc que :

$$\int_{-\infty}^{+\infty} x(\theta) y(t - \theta) d\theta = \int_{-\infty}^{+\infty} x(t - \theta) y(\theta) d\theta \quad (1.13)$$

Transformée de Fourier : Calculons $\text{TF}[(x * y)(t)]$...

$$\text{TF}[(x * y)(t)] = \int_{-\infty}^{+\infty} \left(\int_{-\infty}^{+\infty} x(\theta) y(t - \theta) d\theta \right) e^{-j2\pi\nu t} dt$$

Ou :

$$\text{TF}[(x * y)(t)] = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} x(\theta) y(t - \theta) e^{-j2\pi\nu t} d\theta dt$$

On écrit $e^{-j2\pi\nu t} = e^{-j2\pi\nu(t-\theta)} e^{-j2\pi\nu\theta}$ et on obtient, en regroupant :

$$\text{TF}[(x * y)(t)] = \int_{-\infty}^{+\infty} \left(\int_{-\infty}^{+\infty} y(t - \theta) e^{-j2\pi\nu(t-\theta)} dt \right) x(\theta) e^{-j2\pi\nu\theta} d\theta$$

Dans l'intégrale centrale, on effectue le changement de variable $u = t - \theta$; il vient alors :

$$\text{TF}[(x * y)(t)] = \int_{-\infty}^{+\infty} \left(\int_{-\infty}^{+\infty} y(u) e^{-j2\pi\nu u} du \right) x(\theta) e^{-j2\pi\nu\theta} d\theta$$

On peut ensuite séparer les variables, et on obtient :

$$\text{TF}[(x * y)(t)] = \left(\int_{-\infty}^{+\infty} y(u) e^{-j2\pi\nu u} du \right) \left(\int_{-\infty}^{+\infty} x(\theta) e^{-j2\pi\nu\theta} d\theta \right)$$

1.11. Autrement dit, si on écrit $x(t)$ sous la forme $x(t) = x_1(t) + jx_2(t)$, alors $x^*(t) = x_1(t) - jx_2(t)$.

Et donc :

$$\boxed{\text{TF}[(x * y)(t)] = X(\nu)Y(\nu)} \quad (1.14)$$

Par symétrie dans les relations 1.3 et 1.4, on obtient de même :

$$\boxed{\text{TF}[(x.y)(t)] = (X * Y)(\nu)} \quad (1.15)$$

La transformée de Fourier de la convolution de deux signaux est le produit de leurs transformées de Fourier, et la transformée de Fourier inverse d'une convolution de deux TF est le produit des deux transformées de Fourier inverses.

1.2.3 Représentation de Fourier des signaux d'énergie infinie

Les signaux d'énergie infinie sont ceux pour lesquels l'intégrale 1.1 ne converge pas.

1.2.3.1 Impulsion de Dirac

Définition : on introduit $\delta(t)$, noté *impulsion de Dirac*^{1.12}, défini par sa transformée de Fourier, tel que :

$$\boxed{\text{TF}[\delta(t)] \triangleq \mathbb{1}} \quad (1.16)$$

où $\mathbb{1}$ désigne la fonction uniformément égale à 1 sur $\mathbb{R}$.

Plus « physiquement », δ est la limite quand $T \rightarrow 0$ du signal suivant :

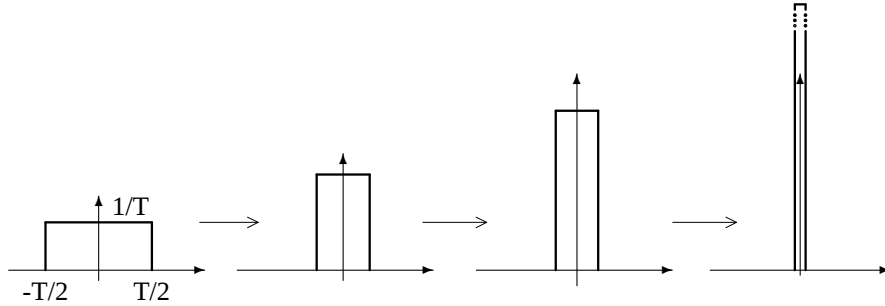


FIG. 1.4 – Construction d'une impulsion de Dirac

On représente graphiquement cette impulsion ainsi :

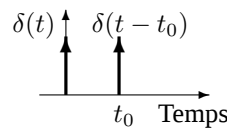


FIG. 1.5 – Représentation schématique d'une impulsion de Dirac

1.12. On dit parfois aussi « pic » de Dirac.

Propriétés : soit x un signal à temps continu, d'énergie finie.

1. Calculons $\text{TF}[x(t)\delta(t)]$: il s'agit de la transformée de Fourier d'un produit, donc en appliquant la formule 1.15, le résultat est la convolution des deux transformées de Fourier :

$$\text{TF}[x(t)\delta(t)] = \int_{-\infty}^{+\infty} X(\nu') \mathbb{1}(\nu - \nu') d\nu' = \int_{-\infty}^{+\infty} X(\nu') d\nu'$$

On écrit $1 = e^{+j2\pi\nu'0}$, et on obtient :

$$\text{TF}[x(t)\delta(t)] = \int_{-\infty}^{+\infty} X(\nu') e^{+j2\pi\nu'0} d\nu'$$

Or le membre de droite n'est autre que la valeur prise par $x(t)$ en $t = 0$ (cf. définition 1.4 de la transformée de Fourier inverse). Il vient donc :

$$\boxed{\text{TF}[x(t)\delta(t)](\nu) = x(0)} \quad (1.17)$$

En particulier, pour $\nu = 0$, on obtient facilement :

$$\int_{-\infty}^{+\infty} x(t)\delta(t)dt = x(0) \quad (1.18)$$

En généralisant, on obtient également facilement par un changement de variable :

$$\boxed{\int_{-\infty}^{+\infty} x(t)\delta(t - t_0)dt = x(t_0)} \quad (1.19)$$

2. Calculons également $(x * \delta)(t)$:

$$(x * \delta)(t) = \text{TF}^{-1}[\text{TF}(x * \delta)] = \text{TF}^{-1}[X(\nu) \cdot \mathbb{1}] = \text{TF}^{-1}[X(\nu)] = x$$

L'impulsion de Dirac est donc l'*élément neutre de la convolution*.

3. La définition 1.16 se traduit par :

$$\int_{-\infty}^{+\infty} e^{+j2\pi\nu t} d\nu = \delta(t)$$

mais également par symétrie entre les relations 1.3 et 1.4, par :

$$\int_{-\infty}^{+\infty} e^{-j2\pi\nu t} dt = \delta(\nu) \quad (1.20)$$

4. Impulsion de Dirac et échelon de Heaviside. L'échelon de Heaviside est défini comme suit :

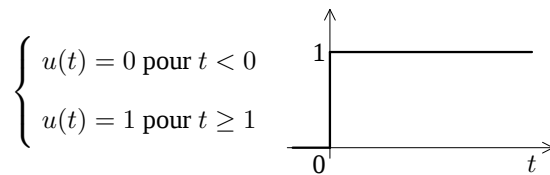


FIG. 1.6 – Échelon de Heaviside

Soient a et b deux réels non nuls, $b > a$. Calculons $I = \int_a^b u'(t)x(t)dt$:

$$\int_a^b u'(t)x(t)dt = [u(t)x(t)]_a^b - \int_a^b u(t)x'(t)dt$$

en utilisant une intégration par parties (cf. note 1.8). Trois cas se présentent alors :

- (a) $a > 0$ et $b > 0$: alors $u(b) = u(a) = 1$, et

$$I = u(b)x(b) - u(a)x(a) - (x(b) - x(a)) = 0$$

(b) $a < 0$ et $b < 0$: alors $u(b) = u(a) = 0$, et

$$I = 0 - 0 + \int_a^b 0 \cdot x(t) dt = 0$$

(c) $a < 0$ et $b > 0$: alors $u(b) = 1$ et $u(a) = 0$, et

$$I = x(b) - \int_0^b x'(t) dt = x(b) - (x(b) - x(0)) = x(0)$$

Cette relation devant être vérifiée quels que soient a et b , on obtient :

$$\int_{-\infty}^{+\infty} u'(t)x(t) dt = x(0)$$

En comparant avec la relation 1.18, et ces égalités devant être vérifiées quel que soit le signal x , il vient donc que

$$\boxed{u'(t) = \delta(t)} \quad (1.21)$$

La dérivée de l'échelon de Heaviside est l'impulsion de Dirac.

1.2.3.2 Spectre des signaux périodiques

Soit $x(t)$ un signal à temps continu, de période T . On *admet* que x est « développable en série de Fourier » sous la forme :

$$\boxed{x(t) = \sum_{n \in \mathbb{Z}} x_n e^{j2\pi n \frac{t}{T}}} \quad (1.22)$$

avec

$$\boxed{x_n = \frac{1}{T} \int_{(T)} x(t) e^{-j2\pi n \frac{t}{T}} dt} \quad (1.23)$$

Pour un signal x impair, son développement en série de Fourier se simplifie en

$$x(t) = \sum_{n \in \mathbb{N}} \alpha_n \sin\left(2\pi n \frac{t}{T}\right)$$

Si x est pair, on peut de même écrire

$$x(t) = \sum_{n \in \mathbb{N}} \alpha_n \cos\left(2\pi n \frac{t}{T}\right)$$

Dans les deux cas, le coefficient α_1 est l'« amplitude du fondamental » et pour $n > 1$ les coefficients α_n sont les amplitudes des « harmoniques ». On peut alors définir le *taux d'harmoniques* τ par

$$\tau = \frac{\alpha_1}{\sum_{n>1} \alpha_n}$$

Calculons la Transformée de Fourier de x :

$$X(\nu) = \int_{-\infty}^{+\infty} x(t) e^{-j2\pi \nu t} dt = \int_{-\infty}^{+\infty} \left(\sum_{n \in \mathbb{Z}} x_n e^{j2\pi n \frac{t}{T}} \right) e^{-j2\pi \nu t} dt$$

En admettant la validité de la permutation des symboles de somme et d'intégration, on obtient :

$$X(\nu) = \sum_{n \in \mathbb{Z}} x_n \left(\int_{-\infty}^{+\infty} e^{j2\pi n \frac{t}{T}} e^{-j2\pi \nu t} dt \right) = \sum_{n \in \mathbb{Z}} x_n \left(\int_{-\infty}^{+\infty} e^{j2\pi \left(\frac{n}{T} - \nu\right)t} dt \right)$$

Or la relation 1.20 donne $\int_{-\infty}^{+\infty} e^{j2\pi \left(\frac{n}{T} - \nu\right)t} dt = \delta\left(\nu - \frac{n}{T}\right)$ donc :

$$\boxed{X(\nu) = \sum_{n \in \mathbb{Z}} x_n \delta\left(\nu - \frac{n}{T}\right)} \quad (1.24)$$

Exemple : cas d'un signal carré.

On considère le signal T -périodique $x(t)$ tel que :

$$\begin{cases} x(t) = 1 & \text{pour } -T/4 < t < +T/4 \\ x(t) = 0 & \text{pour } +T/4 < |t| < T \end{cases}$$

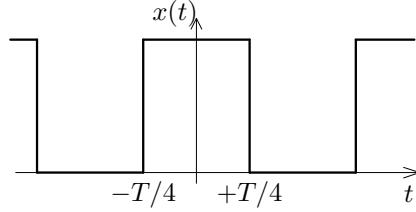


FIG. 1.7 – Exemple de signal carré

On a alors

$$x_n = \frac{1}{T} \int_{(T)} x(t) e^{-j2\pi n \frac{t}{T}} dt = \frac{1}{T} \int_{-T/4}^{+T/4} e^{-j2\pi n \frac{t}{T}} dt = \frac{-1}{j2\pi n} (e^{-jn\pi/2} - e^{+jn\pi/2}) = \frac{1}{\pi n} \sin n \frac{\pi}{2}$$

En remarquant que seuls les termes d'ordre n impair sont non nuls, et en écrivant dans ce cas $n = 2k + 1$, on obtient

$$X(\nu) = \frac{1}{\pi} \sum_{k \in \mathbb{Z}} \frac{(-1)^k}{2k+1} \delta\left(\nu - \frac{n}{T}\right)$$

1.2.3.3 Cas particulier : peigne de Dirac

Définition : on définit le *peigne de Dirac* de période T par la relation suivante :

$$\boxed{\delta_T(t) \triangleq \sum_{n \in \mathbb{Z}} \delta(t - nT)} \quad (1.25)$$

Il se représente graphiquement comme suit :

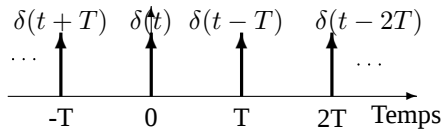


FIG. 1.8 – Peigne de Dirac

Propriété : le peigne de Dirac est un signal périodique, de période T ; il est donc « développable en série de Fourier » :

$$\delta_T(t) = \sum_{n \in \mathbb{Z}} \delta_n e^{j2\pi n \frac{t}{T}}$$

Chacun des coefficients δ_n vaut en vertu de la formule 1.23 :

$$\delta_n = \frac{1}{T} \int_{-T/2}^{+T/2} \delta_T(t) e^{-j2\pi n \frac{t}{T}} dt$$

Soit :

$$\delta_n = \frac{1}{T} \sum_{n \in \mathbb{Z}} \int_{-T/2}^{+T/2} \delta(t - nT) e^{-j2\pi n \frac{t}{T}} dt$$

Dans cette somme infinie, seul le terme pour $n = 0$ est non nul (les autres « $\delta(t - nT)$ » sont nuls sur l'intervalle $[-\frac{T}{2}, +\frac{T}{2}]$). Il vient donc :

$$\delta_n = \frac{1}{T} \int_{-T/2}^{+T/2} \delta(t) e^{-j2\pi n \frac{t}{T}} dt$$

On peut alors augmenter l'intervalle de calcul de l'intégrale sur l'ensemble $\mathbb{R}$ entier, car $\delta(t)$ y est nul ; on obtient alors :

$$\delta_n = \frac{1}{T} \int_{-\infty}^{+\infty} \delta(t) e^{-j2\pi n \frac{t}{T}} dt$$

Et en utilisant la formule 1.18 il vient :

$$\delta_n = \frac{1}{T}$$

En notant $\Delta_T(\nu)$ la Transformée de Fourier du peigne δ_T , il vient donc :

$$\Delta_T(\nu) = \sum_{n \in \mathbb{Z}} \delta_n \delta\left(\nu - \frac{n}{T}\right) = \frac{1}{T} \sum_{n \in \mathbb{Z}} \delta\left(\nu - \frac{n}{T}\right) = \frac{1}{T} \delta_{\frac{1}{T}}(\nu)$$

On peut alors retenir le résultat suivant :

La transformée de Fourier d'un peigne de Dirac (en temps) est un peigne de Dirac (en fréquence).

Corollaire : Autre formule du peigne de Dirac. Utilisons la relation 1.4 de la transformée de Fourier inverse :

$$\delta_T(t) = \int_{-\infty}^{+\infty} \Delta_T(\nu) e^{+j2\pi\nu t} d\nu = \frac{1}{T} \sum_{n \in \mathbb{Z}} \left(\int_{-\infty}^{+\infty} \delta\left(\nu - \frac{n}{T}\right) e^{+j2\pi\nu t} d\nu \right)$$

On applique alors la propriété 1.19, et il vient :

$$\delta_T(t) = \frac{1}{T} \sum_{n \in \mathbb{Z}} e^{j2\pi \frac{n}{T} t}$$

Ce qu'il faut retenir

- La définition de la Transformée de Fourier ;
- Le spectre d'un signal est le module de sa transformée de Fourier ;
- Les propriétés de la TF, et plus spécialement la propriété liée à la convolution :

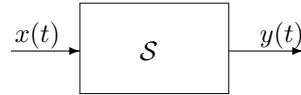
$$\text{TF}[x * y(t)] = \text{TF}[x(t)] \cdot \text{TF}[y(t)]$$

- L'élément neutre de la convolution est l'impulsion de Dirac ; sa transformée de Fourier est la fonction continûment égale à 1.

1.3 Notion de filtre linéaire

1.3.1 Linéarité

On considère un système $\mathcal{S}$ quelconque, représenté sous une forme de « boîte noire », d'entrée x et de sortie y :



Par définition, $\mathcal{S}$ est un *système linéaire* s'il existe une fonction de deux variables $h(t, \theta)$ telle que :

– Si on est à temps continu :

$$y(t) = \int_{-\infty}^{+\infty} h(t, \theta) x(\theta) d\theta \quad (1.26)$$

– Si on est à temps discret :

$$y(n) = \sum_{m \in \mathbb{Z}} h(n, m) x(m)$$

h est appelée *réponse impulsionnelle* du système. En effet, en étudiant la réponse du système à une impulsion, dans le cas par exemple de signaux à temps continu, avec par exemple une impulsion retardée d'un temps τ $x(t) = \delta(t - \tau)$, on obtient facilement en utilisant la formule 1.19 : $y(t, \tau) = h(t, \tau)$. *A priori*, la réponse du système dépend donc du moment de l'excitation.

Dans la suite du cours, on se limitera une fois encore aux signaux à temps continu pour l'établissement des équations.

1.3.2 Invariance

Comme il a été souligné dans le paragraphe précédent, la réponse du système dépend a priori de l'instant où il est excité. L'invariance est la traduction du fait que l'on désire que cette réponse ne dépende plus de cet instant. Autrement dit, si $y(t)$ est la réponse au signal $x(t)$, alors le signal $x(t - \tau)$ doit entraîner la réponse $y(t - \tau)$.

Soit donc le signal $x_1(t)$; son image par le système $\mathcal{S}$ est le signal $y_1(t)$. On considère le signal $x_2(t) = x_1(t - \tau)$ (il s'agit du signal x_1 retardé du temps τ) ; son image est le signal $y_2(t)$. On cherche à avoir $y_2(t) = y_1(t - \tau)$. Traduisons cette égalité en utilisant la relation 1.26 :

$$\int_{-\infty}^{+\infty} h(t, \theta) x_2(\theta) d\theta = \int_{-\infty}^{+\infty} h(t - \tau, \theta) x_1(\theta) d\theta$$

Soit :

$$\int_{-\infty}^{+\infty} h(t, \theta) x_1(\theta - \tau) d\theta = \int_{-\infty}^{+\infty} h(t - \tau, u) x_1(u) du$$

On effectue dans la première intégrale le changement de variable $u = \theta - \tau$; il vient alors :

$$\int_{-\infty}^{+\infty} h(t, u + \tau) x_1(u) du = \int_{-\infty}^{+\infty} h(t - \tau, u) x_1(u) du$$

Cette égalité devant être vérifiée quel que soit le signal x_1 , on a donc nécessairement :

$$\text{Quels que soient } t, \tau, u : h(t, \tau + u) = h(t - \tau, u)$$

En particulier, pour $u = 0$ on obtient :

$$h(t, \tau) = h(t - \tau, 0)$$

La fonction de deux variables $h(t, \theta)$ peut donc se mettre sous la forme d'une fonction de la *différence* de ces deux variables. Par la suite, pour un système linéaire invariant, nous écrirons donc plus simplement $h(t, \theta) = h(t - \theta)$. En remplaçant dans l'équation 1.26, on obtient :

$$\boxed{\mathcal{S} \text{ est un système linéaire invariant} \iff y(t) = \int_{-\infty}^{+\infty} h(t - \theta)x(\theta)d\theta} \quad (1.27)$$

Soit plus simplement, en comparant avec la formule 1.12 :

$$\boxed{\mathcal{S} \text{ est un système linéaire invariant} \iff y(t) = (h * x)(t)}$$

La réponse d'un système linéaire invariant à une entrée quelconque est la *convolution* de cette entrée par la réponse impulsionnelle du système.

1.3.3 Fonction de transfert

Soit $\mathcal{S}$ un système linéaire invariant, et h sa réponse impulsionnelle. Appliquons à l'entrée de $\mathcal{S}$ le signal $x(t) = x_0 e^{st}$, avec $s \in \mathbb{C}$. En utilisant la relation 1.27, il vient :

$$y(t) = \int_{-\infty}^{+\infty} h(t - \theta)x_0 e^{s\theta} d\theta$$

Soit encore, en utilisant la commutativité de la convolution (formule 1.13) :

$$y(t) = x_0 \int_{-\infty}^{+\infty} h(\theta)e^{s(t-\theta)} d\theta$$

On peut alors « sortir » e^{st} de l'intégrale :

$$y(t) = (x_0 e^{st}) \cdot \left(\int_{-\infty}^{+\infty} h(\theta)e^{-s\theta} d\theta \right)$$

Le premier terme du produit est en fait $x(t)$. Le deuxième terme du produit ne dépend pas du temps, mais seulement de la variable s . Pour les mathématiciens, ces deux remarques se traduisent par la constatation que les signaux de la forme e^{st} sont des *signaux propres* du système $\mathcal{S}$. On note le deuxième terme $H(s)$:

$$H(s) = \int_{-\infty}^{+\infty} h(\theta)e^{-s\theta} d\theta$$

H est appelée *fonction de transfert* de $\mathcal{S}$. Dans le cas particulier où $s = j2\pi\nu$, on reconnaît dans l'expression précédente la transformée de Fourier de la réponse impulsionnelle h , et on parle alors de la *fonction de transfert en régime harmonique*.

La fonction de transfert en régime harmonique est la transformée de Fourier de la réponse impulsionnelle, soit :

$$\boxed{H(\nu) = \int_{-\infty}^{+\infty} h(\theta)e^{-j2\pi\nu\theta} d\theta} \quad (1.28)$$

Fonction de transfert et représentation complexe : On a démontré que pour un système linéaire invariant $\mathcal{S}$, de fonction de transfert $H(s)$, dans le cas où l'entrée était de la forme $x(t) = x_0 e^{st}$, on avait la relation suivante entre l'entrée x et la sortie y : $y(t) = H(s)x(t)$. Lorsque l'on utilise la représentation complexe, en écrivant $x(t)$ sous la

forme $x_0 e^{j\omega t}$, la relation qui apparaît lie directement les représentations complexes de l'entrée et de la sortie, et la fonction de transfert en régime harmonique :

$$y(t) = y_0 e^{j\omega t} = x_0 e^{j\omega t} H(j\omega) = x(t) H(j\omega)$$

Ce qu'il faut retenir

- Dans un filtre linéaire invariant, la sortie est la convolution de l'entrée par la réponse impulsionnelle du système. La Transformée de Fourier de la sortie est donc égale à la Transformée de Fourier de l'entrée *multipliée* par celle de la réponse impulsionnelle, appelée *fonction de transfert* ;
-

Chapitre 2

Généralités

2.1 Le circuit électrique

Le but de cette partie est d'introduire quelques notions de base de l'électricité dans son ensemble.

2.1.1 Circuits électriques

Un circuit électrique est un ensemble de *composants électriques* interconnectés d'une manière quelconque par des *conducteurs*.

- Un composant électrique est :
 - dans le cas le plus simple un élément à deux bornes (on dit aussi un *dipôle*), que l'on représente sous la forme suivante :



Les bornes a et b servent à la connexion avec d'autres composants. Dans cette catégorie on trouve par exemple les résistors^{2.1}, condensateurs^{2.1}, bobines^{2.1}, piles, etc.);

- dans certains cas un élément à plus de deux bornes. Par exemple, un transistor possède 3 bornes, un transformateur peut en avoir 4 voire 6. Un composant à quatre bornes est appelé *quadripôle*.
- Un conducteur est constitué d'un matériau transportant bien le courant électrique. Pour des raisons physiques, un bon conducteur électrique est également un bon conducteur thermique. On en trouve ainsi réalisé en métal, et surtout en cuivre. Mais il est également possible d'utiliser un liquide conducteur, appelé *électrolyte* : l'exemple le plus classique est l'eau salée.

2.1.2 Courant, tension, puissance

2.1.2.1 Courant électrique

Un courant électrique est un déplacement d'ensemble ordonné de charges électriques dans un conducteur. On le caractérise par une grandeur, l'*intensité*, définie comme étant le débit de charges électriques dans le conducteur^{2.2}.

2.1. voir section 2.2.

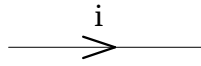
2.2. L'unité légale de charge électrique est le Coulomb (symbole C). Par exemple, un électron porte une charge élémentaire négative, notée e , et valant $e \approx -1,610^{-19}\text{C}$.

Cette grandeur est souvent notée I . Quand, pendant un temps dt , il passe dq Coulombs, l'intensité vaut

$$I = \frac{dq}{dt}$$

L'unité légale dans laquelle s'exprime l'intensité du courant électrique est l'*ampère* (symbole A). Le courant dans le schéma d'un circuit électrique est représenté par une flèche. Il est à noter que du fait de la définition de l'intensité ($I = +\frac{dq}{dt}$) et de la charge de l'électron (charge négative), le sens de déplacement effectif des électrons est l'opposé du sens positif du courant^{2.3}.

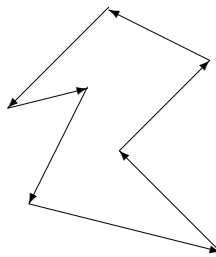
On représente un courant électrique par une flèche sur un conducteur, indiquant le sens positif de l'intensité :



Cette flèche indique que si les électrons passent de droite à gauche, on comptera une intensité positive ; négative s'ils vont de gauche à droite.

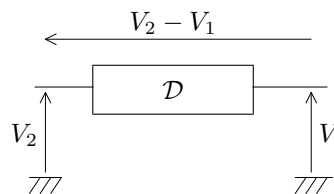
2.1.2.2 Différence de potentiel

Au repos, les charges électriques d'un conducteur sont en mouvement continu sous l'effet de l'agitation thermique :



Cependant, ce mouvement, à une vitesse non nulle, ne se traduit pas par un déplacement global susceptible de se traduire en courant électrique. Pour mettre en mouvement ces charges dans une direction donnée, il est nécessaire d'appliquer un *champ électrique* aux bornes du conducteur. En appliquant le potentiel électrique V_1 et le potentiel V_2 à ces deux bornes, on crée une *différence de potentiel* qui met les électrons^{2.4} en mouvement^{2.5}.

La valeur de la différence de potentiel est appelée la *tension*, et son unité est le *Volt* (symbole V). Le Volt est défini de telle manière qu'une charge d'un Coulomb accélérée sous une tension de 1V acquiert une énergie de 1J : $1V=1J/C$. On représente une différence de potentiel par une flèche à côté du composant, comme sur le schéma suivant :



Dans le bas de ce schéma, les symboles rayés indiquent la *référence de potentiel nulle*, appelée la *masse*, par rapport à laquelle sont définis les potentiels V_1 et V_2 .

2.3. Cette petite incohérence a des origines historiques, l'électron ayant été découvert après la formalisation du phénomène électrique.

2.4. Là où des électrons « manquent » dans la structure cristalline du métal, on trouve des « trous », ou absence d'électrons, que l'on considère comme étant de petites charges positives, également susceptibles d'être mises en mouvement.

2.5. Rappel sur la *force de Laplace* : quand une charge électrique q est placée dans un champ électrique $\vec{E}$, elle est soumise à une force $\vec{F} = q \vec{E}$.

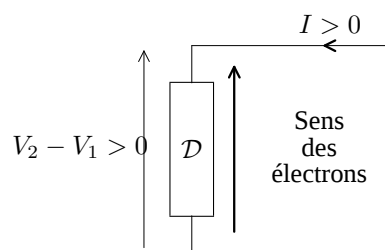
2.1.2.3 Energie, puissance

Ainsi qu'on l'a souligné au paragraphe précédent, l'application d'une différence de potentiel aux bornes d'un conducteur permet de mettre en mouvement les charges électriques libres qu'il renferme. Ce faisant, on leur a communiqué de l'énergie cinétique en apportant de l'énergie électrostatique sous la forme de la différence de potentiel imposée. En se ramenant à une unité de temps, on peut introduire une *puissance électrique* définie comme étant le produit de la tension par le flux de charges par unité de temps dans le conducteur, autrement dit par l'intensité. Il est facile de vérifier que ce produit est effectivement homogène à une puissance : $1V.1A=1(J/C).1(C/s)=1(J/s)=1W$.

2.1.2.4 Conventions générateur/récepteur

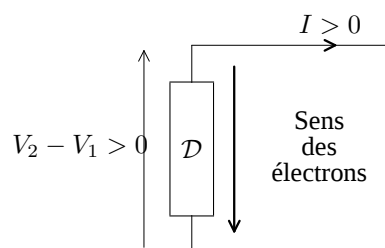
Il est possible de « raffiner » cette notion de puissance électrique en distinguant les composants « générateurs » de puissance de ceux qui se « contentent » de la recevoir.

- **Convention récepteur :** considérons un dipôle que l'on qualifiera de « passif », uniquement capable de recevoir de l'énergie électrique. On impose aux bornes de ce dipôle une ddp $V_2 - V_1$, avec $V_2 > V_1$. Les électrons, de charges négatives, vont se diriger vers le pôle de potentiel le plus élevé. Par conséquent, le courant sera positif dans le sens contraire. Il s'ensuit que l'on peut définir une *convention récepteur* pour les sens positifs des courant et tensions, comme suit :



On notera que la flèche de la tension et celle du courant sont de sens opposés.

- **Convention générateur :** cette convention est la « duale » de la précédente. Il s'agit cette fois-ci pour le dipôle d'imposer la tension à ses bornes *et* l'intensité du courant qui le traverse. En fait, on définit la *convention générateur* d'après la convention récepteur. Si l'on veut pouvoir brancher l'un en face l'autre un récepteur et un générateur, il faut nécessairement que les conventions de signe pour ce dernier soient les suivantes, pour qu'il n'y ait pas d'incompatibilité entre les définitions :

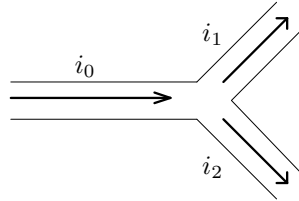


On notera que cette fois-ci, les deux flèches sont dans le même sens.

2.1.3 Lois de Kirchhoff

2.1.3.1 Loi des nœuds

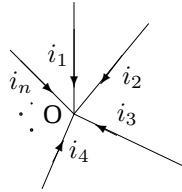
Cette loi se déduit facilement de la notion de courant électrique. Supposons que l'on ait un flux $i_0 = \frac{dq_1}{dt}$ d'électrons dans un conducteur arrivant à un « embranchement » d'un circuit électrique :



Les électrons venant de la « gauche » partiront soit dans la première, soit dans la deuxième branche. Mais le nombre total d'électrons par seconde restera le même que celui qui arrive en permanence par la gauche, et donc $i_0 = i_1 + i_2$ (avec les sens des courants définis suivant la figure précédente).

Dans la théorie des réseaux de Kirchhoff, un *nœud* est un point de convergence de plusieurs conducteurs.

Plus généralement, si on considère n conducteurs arrivant au même point O, avec les sens positifs des courants i_n définis comme suit, vers O...

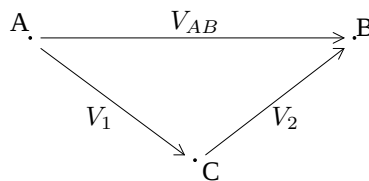


La loi des nœuds stipule alors que *la somme algébrique des courants arrivant à un nœud est constamment nulle* :

$$\sum_{k=1}^n i_k = 0$$

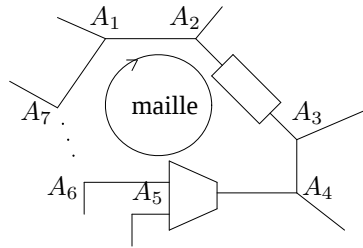
2.1.3.2 Loi des mailles

Cette loi découle de la remarque selon laquelle entre deux points quelconques, la différence de potentiel est bien définie. Considérons par exemple trois points A, B et C. On mesure entre A et B la tension $V_{AB} = V_B - V_A$, entre A et C la tension V_1 et entre C et B la tension V_2 :



Par définition de V_1 , on a $V_1 = V_C - V_A$ et de même pour V_2 , $V_2 = V_B - V_C$. Il s'ensuit que $V_1 + V_2 = (V_C - V_A) + (V_B - V_C) = V_B - V_A = V_{AB}$. Cela s'apparente à une relation vectorielle.

Dans la théorie des réseaux de Kirchhoff, une *maille* est une « chaîne » de conducteurs et de composants électriques, partant d'un point, et arrivant à ce même point, par exemple :



La loi des mailles stipule que *la somme algébrique des tensions le long de la maille est constamment nulle* :

$$\sum_{k=2}^n V_{A_k A_{k-1}} = 0$$

Ce qu'il faut retenir

- ce que sont le courant électrique (un flux d'électrons), sa mesure (l'intensité), et la tension ;
 - la notion d'énergie et de puissance électriques ;
 - les lois des nœuds et des mailles.
-

2.2 Dipôles électriques

2.2.1 Le résistor

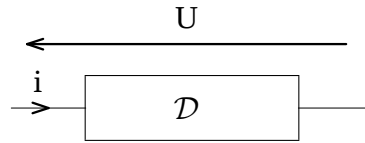
2.2.1.1 L'effet résistif

On considère un conducteur, aux bornes duquel on impose une différence de potentiel. On a déjà indiqué que ce conducteur serait alors traversé par un courant électrique, un flux d'électrons. Cependant, tous les matériaux ne « conduisent » pas l'électricité aussi facilement : certains offrent plus ou moins de *résistance* au passage des électrons. C'est ce phénomène que l'on appelle l'*effet résistif*.

2.2.1.2 Loi d'Ohm

Cette loi exprime que certains matériaux ont une réponse linéaire en courant à une différence de potentiel imposée. Si l'on considère un tel dipôle, noté $\mathcal{D}$ aux bornes duquel on impose la différence de potentiel U , et traversé par le

courant i . Ce dipôle est un *résistor* :



Quel que soit l'instant t , U et i vérifient la relation de proportionnalité

$$U(t) = R.i(t)$$

où R est appelée *résistance* du résistor, et s'exprime en *Ohms*, en abrégé Ω . L'inverse de la résistance est la *conductance*, souvent notée G , et s'exprime en *Siemens* (abréviation S) : $G = 1/R$.

2.2.1.3 Aspect énergétique

On a déjà dit que la résistance traduisait la « difficulté » avec laquelle les électrons peuvent circuler dans le matériau. Cette difficulté s'accompagne d'un échauffement : c'est ce qu'on appelle l'*effet Joule*. Cet échauffement, du point de vue du circuit électrique, est une perte d'énergie par dissipation thermique. Pour une résistance R , parcourue par un courant i et aux bornes de laquelle on mesure la tension U , cette puissance perdue P_J est égale à :

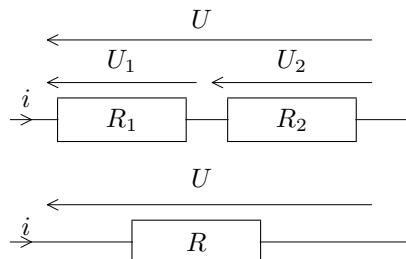
$$P_J = Ri^2 = \frac{U^2}{R}$$

Par exemple, une résistance $R = 10 \Omega$ parcourue par un courant de $i = 0,5 \text{ A}$ dissipe 2,5 W.

2.2.1.4 Associations de résistors

Considérons deux résistances R_1 et R_2 . On peut les associer de deux manières : soit elles sont parcourues par le même courant (association en *série*), soit elles sont soumises à la même différence de potentiel (association en *parallèle*). On cherche dans chaque cas la résistance R équivalente à l'ensemble de R_1 et R_2 .

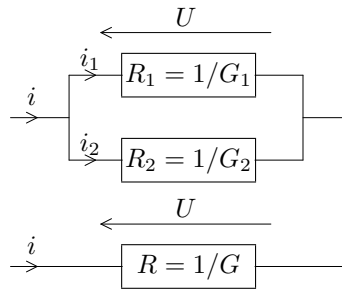
1. **Association en série** ; les deux résistances sont associées ainsi :



La loi des mailles (paragraphe 2.1.3.2) nous permet d'écrire $U = U_1 + U_2$. Or on a aussi $U_1 = R_1 i$ et $U_2 = R_2 i$. Il vient donc $U = (R_1 + R_2)i$, soit $R = R_1 + R_2$:

La résistance équivalente à deux résistances mises en série est égale à la somme des résistances.

2. **Association en parallèle** ; les deux résistances sont associées ainsi :



On note leurs conductances respectives G_1, G_2 et la conductance équivalente G . La loi des nœuds (paragraphe 2.1.3.1) nous permet d'écrire $i = i_1 + i_2$. Or on a aussi $i_1 = G_1 U$ et $i_2 = G_2 U$. Il vient donc $i = (G_1 + G_2)U$, soit $G = G_1 + G_2$:

La conductance équivalente à deux conductances mises en parallèle est égale à la somme des conductances.

Autrement dit, l'inverse de la résistance équivalente est égale à la somme des inverses des résistances.

2.2.2 La bobine

2.2.2.1 Les effets inductif et auto-inductif

Considérons deux conducteurs. On fait circuler dans l'un de ces conducteurs un courant électrique :



Ce courant crée un champ d'induction magnétique. Si de plus le courant est variable, le champ ainsi créé est lui-même variable et est responsable de l'apparition d'un courant dit *induit* dans le deuxième conducteur : c'est l'*effet inductif*. Dans le même temps, le champ d'induction magnétique réagit sur le courant qui l'a créé, en ralentissant sa vitesse de variation. C'est l'*effet auto-inductif*.

2.2.2.2 Caractéristique tension/courant d'une bobine

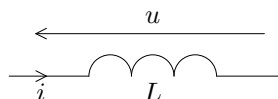
On définit le coefficient d'induction magnétique de la bobine par le rapport entre le flux d'induction magnétique à travers le circuit^{2.6}, et le courant qui lui donne naissance ; on le note L :

$$L(t) = \frac{\Phi(t)}{i(t)}$$

Or la différence de potentiel u apparaissant grâce à l'effet auto-inductif aux bornes de la bobine est égale à $u = \frac{d\Phi}{dt}$. Il vient donc

$$u(t) = L \frac{di}{dt}$$

où L est appelée l'*inductance* de la bobine et s'exprime en Henri (H). Dans un circuit électrique, on représente une bobine sous la forme suivante :



^{2.6}. En résumé, le produit du champ magnétique par la surface enserrée par le circuit.

2.2.2.3 Aspect énergétique

Le phénomène physique correspond au stockage d'énergie sous forme magnétique. Le stockage est momentané et l'énergie est restituée au circuit en courant. L'énergie accumulée par la bobine vaut :

$$E_{\text{mag}}(t) = \frac{1}{2} Li(t)^2$$

2.2.3 Le condensateur

2.2.3.1 L'effet capacitif

Lorsqu'on applique une différence de potentiel à deux conducteurs *isolés*, on assiste à une accumulation de charges par effet électrostatique. C'est l'*effet capacitif*. Il peut être recherché et dans ce cas on fabrique des composants spécialisés qui lui font appel, les *condensateurs*, ou bien n'être qu'un parasite. Il tend à retarder les signaux.

2.2.3.2 Caractéristique tension/courant d'un condensateur

Pour un circuit donné, on définit sa *capacité* C comme le rapport de la charge accumulée sur la tension appliquée à ses bornes :

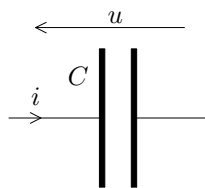
$$C = \frac{q}{u}$$

L'unité de C est le Farad (F).

Or le courant est la dérivée de la charge par unité de temps (cf 2.1.2.1) : $i(t) = \frac{dq}{dt}$ donc il vient :

$$i(t) = C \frac{du}{dt}$$

On représente un condensateur sous la forme suivante :



2.2.3.3 Aspect énergétique

Le phénomène physique correspond au stockage d'énergie sous forme électrostatique. Le stockage est momentané et cette énergie est restituée au circuit sous forme de tension. L'énergie accumulée par le condensateur vaut :

$$E_{\text{stat}} = \frac{1}{2} Cu(t)^2$$

Ce qu'il faut retenir

- résistor et résistance ; condensateur et capacité ; bobine et inductance ;

- les lois d'association en série et en parallèle des résistances.

2.3 Régime sinusoïdal, ou *harmonique*

2.3.1 Définitions

Un signal *harmonique*, ou en utilisant une analogie avec la lumière, *monochromatique*, est un signal sinusoïdal, de fréquence ν donnée. La représentation « classique » de ce signal se fait sous la forme réelle :

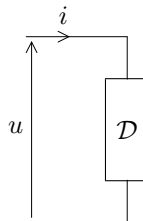
$$x(t) = x_0 \sin 2\pi\nu t \text{ ou } x(t) = X\sqrt{2} \sin 2\pi\nu t$$

x_0 est appelé *amplitude* et X *valeur efficace* de x . On peut poser $\omega = 2\pi f$: ω est appelée pulsation (ou vitesse angulaire pour certaines applications).

2.3.2 Puissance en régime sinusoïdal

2.3.2.1 Puissance en régime périodique

On considère un dipôle $\mathcal{D}$ en convention récepteur :



On définit la *puissance instantanée* dissipée dans le dipôle par

$$p(t) = u(t)i(t)$$

En régime périodique^{2.7}, avec tension et courant de période T , on peut définir également la *puissance moyenne* par

$$P = \frac{1}{T} \int_{(T)} p(t) dt = \frac{1}{T} \int_{(T)} u(t)i(t) dt$$

2.3.2.2 Puissance instantanée en régime sinusoïdal

Supposons que $u(t) = U\sqrt{2} \cos \omega t$ et $i(t) = I\sqrt{2} \cos (\omega t - \phi)$. Il vient alors, après quelques calculs :

$$p(t) = UI[\cos \phi + \cos (2\omega t - \phi)]$$

La puissance instantanée est donc la somme d'un terme constant ($UI \cos \phi$) et d'un terme variable à fréquence double de la fréquence initiale ($UI \cos (2\omega t - \phi)$). Il s'ensuit que dans le cas général ($\phi \neq 0$ et $\phi \neq \pi$), le signe de $p(t)$ varie au cours du temps : le dipôle est tour à tour générateur puis récepteur de puissance électrique.

^{2.7} Sinusoïdal ou non...

2.3.2.3 Puissance moyenne en régime sinusoïdal

1. **Puissance active :** On la définit par $P = UI \cos \phi$. On l'appelle *puissance active* car c'est elle qui est réellement utile (par exemple, dans un moteur, c'est la puissance active qui est transformée en puissance mécanique, aux pertes près). Deux cas se présentent :

- $-\pi/2 < \phi < +\pi/2$: $P > 0$, ce qui signifie que le dipôle est *récepteur de puissance* ;
- $+\pi/2 < \phi < +\pi$: $P < 0$, ce qui signifie que le dipôle est *émetteur de puissance*.

Cas d'un condensateur ou d'une bobine :

- Condensateur : on a $i(t) = C \frac{du(t)}{dt}$ donc si $u(t) = U\sqrt{2} \cos \omega t$, alors

$$i(t) = -C\omega U\sqrt{2} \sin \omega t = \underbrace{(C\omega U)}_I \sqrt{2} \cos [\omega t - (-\pi/2)]$$

On en déduit que $\phi = -\pi/2$, et donc que dans le cas d'un condensateur parfait, la puissance active est nulle.

- Bobine : on a de même $u(t) = L \frac{di(t)}{dt}$, qui nous amène facilement à $\phi = +\pi/2$, et donc également à une puissance active nulle.

2. **Puissance réactive :** On ne peut pas faire de différence, simplement en examinant le bilan de puissance active, entre un condensateur et une bobine. Par symétrie avec la définition de la puissance active, on définit la *puissance réactive*, souvent notée Q , par $Q = UI \sin \phi$. L'unité de puissance réactive est le *Volt Ampère Réactif* (VAR). Quand $0 < \phi \leq \pi/2$, $Q > 0$ et on dit que le dipôle est de type inductif. Quand $-\pi/2 \leq \phi < 0$, $Q < 0$ et le dipôle est dit capacitif^{2.8}.

3. **Puissance apparente :** $P = UI \cos \phi$ et $Q = UI \sin \phi$ amènent naturellement à définir la quantité $S = \sqrt{P^2 + Q^2} = UI$, appelée *puissance apparente*.

Il vient alors $P = S \cos \phi$: $\cos \phi$ est donc un facteur mesurant l'efficacité de production de puissance active du système, et est appelé *facteur de puissance*.

2.3.3 Représentation complexe d'un signal harmonique

On considère un signal harmonique $x(t) = x_0 \cos \omega t$. On définit alors sa représentation complexe $\tilde{x}$ sous la forme

$$\tilde{x}(t) = x_0 e^{j\omega t}$$

On identifiera par la suite x et $\tilde{x}$, et on écrira donc souvent par abus de notation : $x(t) = x_0 e^{j\omega t}$. On verra plus tard que l'utilisation de la représentation complexe permet de simplifier les calculs. Pour repasser ensuite dans le domaine réel, il suffit de prendre la partie imaginaire^{2.9} du résultat des calculs^{2.10} : $x(t) = \Im[\tilde{x}(t)]$.

Dérivation : A partir de la forme complexe, il est aisé d'établir une relation entre un signal $\tilde{x}(t)$ et sa dérivée par rapport au temps. En effet, si $x(t) = x_0 e^{j\omega t}$, alors $\frac{dx}{dt} = (j\omega)x_0 e^{j\omega t} = j\omega (x_0 e^{j\omega t})$, soit :

$$\frac{dx}{dt} = j\omega x$$

De même, pour *intégrer* un signal, il suffit de *diviser* sa représentation complexe par $j\omega$.

Expression de la puissance en notation complexe : L'expression utilisable en notations réelles et donnée dans le paragraphe 2.3.2.2 ne l'est plus quand on manipule les représentations complexes. La puissance instantanée devient

$$p(t) = \frac{1}{2} \Re[u(t)\tilde{i}^*(t)]$$

2.8. Attention : si l'on change la définition du déphasage et que l'on pose par exemple $i(t) = I\sqrt{2} \cos(\omega t + \phi)$, alors la puissance réactive est définie par $Q = -UI \sin \phi$.

2.9. On peut également définir la représentation complexe à partir de $x(t) = x_0 \cos 2\pi\nu t$, auquel cas pour revenir à la représentation réelle du signal il faut prendre la partie réelle.

2.10. Les signes $\Re(z)$ et $\Im(z)$ désignent respectivement les parties réelle et imaginaire du nombre complexe z .

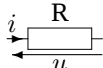
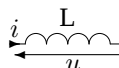
où $\tilde{i}^*(t)$ désigne la quantité complexe conjuguée du courant. La puissance moyenne s'écrit alors

$$P = \frac{1}{2T} \Re \left[\int_{(T)} p(t) dt \right] = \frac{1}{2T} \Re \left[\int_{(T)} u(t) \tilde{i}^*(t) dt \right]$$

2.3.4 Impédances

2.3.4.1 Rappel : caractéristiques tension/courant

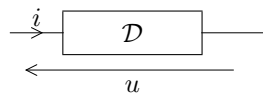
On considère un dipôle, parcouru par un courant i , et aux bornes duquel on mesure la tension u :

Nom	Résistance	Condensateur	Bobine
Schéma			
Expression de la loi d'Ohm	$u = Ri$	$i = C \frac{du}{dt}$	$u = L \frac{di}{dt}$

TAB. 2.1 – Relations entre u et i en réel

2.3.4.2 Impédance complexe

Pour un dipôle $\mathcal{D}$, parcouru par le courant i et aux bornes duquel on mesure la tension u , l'*impédance complexe* est définie comme étant le rapport de la représentation complexe de u par celle de i :



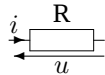
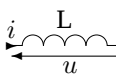
$$Z = \frac{u}{i}$$

L'inverse de l'impédance est appelée *admittance*, et est souvent notée Y .

Dans le cas général, un dipôle quelconque n'a pas une impédance « purement » réelle ou imaginaire. De plus, a priori, cette impédance dépend de la fréquence, comme on peut le remarquer par exemple pour une bobine ou un condensateur. Une impédance peut également avoir une partie imaginaire négative (comme un condensateur, par exemple) et on dit alors qu'elle est de type *capacitif*, ou une partie imaginaire positive (par exemple une bobine) : elle est alors de type *inductif*. En revanche, pour des composants passifs^{2.11}, la partie réelle, qui correspond à une résistance, est dite *résistive* et est toujours positive.

Le tableau 2.1 se traduit alors en :

2.11. Pour simplifier, les composants dits « actifs » sont alimentés : par exemple, un transistor ou un amplificateur opérationnel (qui n'est autre qu'un ensemble de transistors et de composants passifs !), et les composants « passifs » sont... les autres : résistances, condensateurs, diodes, etc.

Nom	Résistance	Condensateur	Bobine
Schéma			
Expression de la loi d'Ohm	$u = Ri$	$u = \frac{1}{jC\omega}i$	$u = jL\omega i$

TAB. 2.2 – Relations entre u et i en complexe

2.3.4.3 Associations d'impédances

Il est facile de vérifier que :

- L'impédance équivalente à deux impédances mises en série est égale à la somme des deux impédances :

$$\begin{array}{c} Z_1 \\ \boxed{} \end{array} \text{---} \begin{array}{c} Z_2 \\ \boxed{} \end{array} \iff \begin{array}{c} Z = Z_1 + Z_2 \\ \boxed{} \end{array}$$

- L'impédance équivalente à deux impédances mises en parallèle est égale à l'inverse de la somme des inverses des impédances (autrement dit, les admittances s'ajoutent) :

$$\begin{array}{c} Z_1 \\ \boxed{} \\ Z_2 \\ \boxed{} \end{array} \iff \begin{array}{c} Z = \frac{Z_1 Z_2}{Z_1 + Z_2} \\ \boxed{} \end{array}$$

Ce qu'il faut retenir

- la définition de la représentation complexe d'un signal harmonique ;
- puissances active, réactive et apparente ;
- la définition de l'impédance complexe ;
- impédances d'un résistor, d'un condensateur, d'une bobine et leurs règles d'association.

2.4 Spectre et fonction de transfert

Ce paragraphe est une reformulation simplifiée de ce qui a déjà été fait en 1.3.

2.4.1 Spectre d'un signal

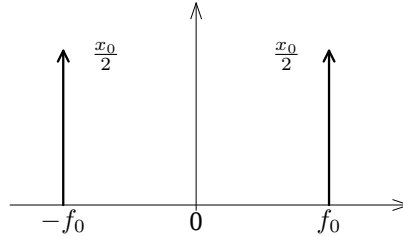
2.4.1.1 Introduction

Nous avons déjà défini ce qu'était un signal « harmonique », ou monochromatique, dans le paragraphe 2.3. Un tel signal ne présente qu'une unique fréquence. Mais on peut imaginer un signal présentant 2, 3 voire une centaine de fréquences différentes. On pourrait représenter ce signal par son évolution temporelle ; il existe néanmoins une *autre* manière de le représenter, en mettant en évidence son contenu fréquentiel. Pour introduire cette nouvelle représentation, nous allons pour un temps revenir à un signal monochromatique, de la forme $x(t) = x_0 \cos \omega_0 t$. On peut

également écrire, en utilisant une formule d'Euler

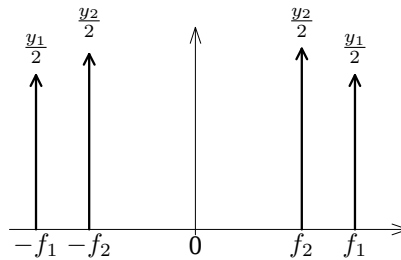
$$x(t) = \frac{x_0}{2} (e^{+j\omega_0 t} + e^{-j\omega_0 t})$$

Cette dernière formulation met en évidence le fait que $x(t)$ peut s'écrire comme la somme de deux exponentielles complexes, associées aux pulsations ω_0 et $-\omega_0$. On représente ces deux composantes sur un axe gradué en pulsations, ou, mieux, en fréquences, par deux « flèches^{2.12} » affectées de leurs poids respectifs (en l'occurrence, les deux composantes ont un poids égal à $x_0/2$) :



Cette représentation est la représentation *fréquentielle* du signal, et la fonction $X(f)$ correspondante, ici limitée à deux pics à $\pm f_0$, est sa « transformée de Fourier ». Le *module* de $X(f)$, noté $S_x(f) = |X(f)|$, est le *spectre* du signal.

Lorsque l'on a un signal présentant *deux* fréquences, comme par exemple $y(t) = y_1 \cos \omega_1 t + y_2 \cos \omega_2 t$, on obtient facilement de même :



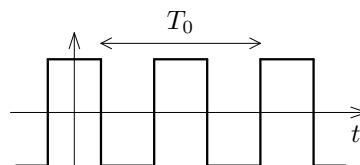
Un problème (qui n'en est bien sûr pas un...) semble se poser pour un signal de la forme $z(t) = z_1 \cos \omega_1 t + z_2 \sin \omega_2 t$. Revenons à la décomposition que nous avons déjà utilisée :

$$z(t) = \frac{z_1}{2} (e^{+j\omega_1 t} + e^{-j\omega_1 t}) + \frac{z_2}{2j} (e^{+j\omega_2 t} - e^{-j\omega_2 t})$$

Cette fois-ci, il apparaît que les « poids » des impulsions de Dirac sont des nombres complexes : pour $\pm f_1$ il s'agit de $z_1/2$, pour $-f_2$ de $(z_2/2)e^{+j\pi/2}$ et pour $+f_2$ de $(z_2/2)e^{-j\pi/2}$. Par conséquent, le spectre de $z(t)$ est rigoureusement égal à celui de $y(t)$, bien que ces deux signaux ne soient pas égaux. En effet, seules les *phases* de leurs transformées de Fourier diffèrent, et elles n'apparaissent pas dans le spectre.

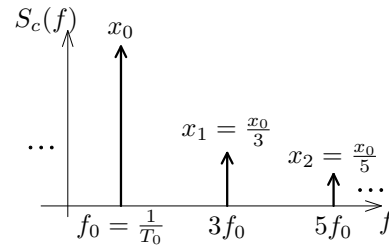
2.4.1.2 Signaux multipériodiques et apériodiques

1. **Spectre d'un signal multipériodique** : il existe des signaux périodiques dont le contenu fréquentiel est infini. Par exemple, le signal carré $c(t)$ suivant...

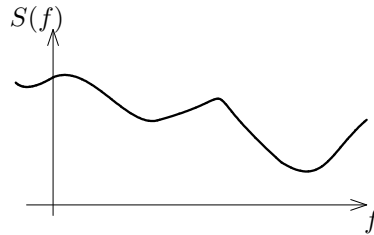


2.12. Cette représentation est en fait celle d'une « impulsion de Dirac » : cf. 1.2.3.1.

... présente ce que l'on appelle un « spectre de raies » : on montre qu'il peut s'écrire sous la forme^{2.13} $c(t) = x_0 \sum_{k=0}^{+\infty} x_k \cos[2\pi(2k+1)t/T_0]$, $x_0 \neq 0$, avec $x_k = x_0/(2k+1)$. Son spectre présente donc de l'énergie aux fréquences du type $f_k = (2k+1)/T_0$, avec k entier naturel :



2. **Spectre d'un signal apériodique :** comme son nom l'indique, un tel signal n'a pas un nombre « dénombrable »^{2.14} de fréquences, mais une infinité. Le spectre d'un tel signal n'est pas un spectre de raies, mais présente des parties continues, par exemple :



On peut considérer qu'il s'agit de la juxtaposition d'un nombre infini d'impulsions de Dirac. On appelle *support fréquentiel* d'un signal l'intervalle de fréquences entre lesquelles son spectre présente de l'énergie.

2.4.2 Fonction de transfert

On considère une « boîte noire », à l'entrée et à la sortie de laquelle on mesure respectivement les tensions v_e et v_s , dont on prend les représentations complexes $\tilde{v}_e$ et $\tilde{v}_s$:



On définit la *fonction de transfert* en régime harmonique du système, notée $H(j\omega)$, par

$$H(j\omega) = \frac{\tilde{v}_s}{\tilde{v}_e}$$

La fonction de transfert est une caractéristique du système, dont la valeur ne dépend que de la fréquence du signal d'entrée.

Remarque : Il est également possible de définir une fonction de transfert comme le rapport de la tension de sortie et du courant d'entrée par exemple, auquel cas cette grandeur a la dimension d'une résistance, mais le plus souvent il s'agit du rapport de deux tensions, quantité sans dimension.

2.13. Il est *décomposable en série de Fourier*. Mathématiquement en fait, il faut de plus qu'un tel signal soit continu, ce qui n'a pas été supposé ; néanmoins, tous les signaux « physiques » en électricité sont continus.

2.14. Que l'on peut compter...

Ce qu'il faut retenir

- la définition du spectre d'un signal ;
 - la notion de fonction de transfert comme le rapport d'une grandeur complexe de sortie sur une grandeur complexe d'entrée.
-

Chapitre 3

Du semi-conducteur aux transistors

Remarque : ce chapitre est très largement inspiré de la partie correspondante du remarquable *Cours d'électronique pour ingénieurs physiciens* de l'Ecole polytechnique fédérale de Lausanne, accessible par Internet à <http://c3iwww.epfl.ch/teach>

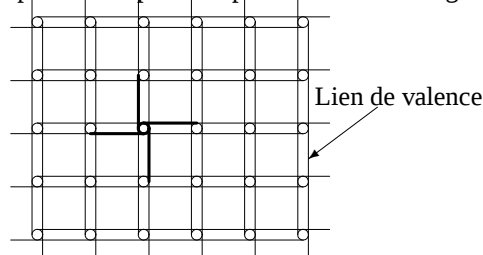
3.1 Les semi-conducteurs

Cette partie va présenter quelques modèles simples de semi-conducteurs, en vue d'expliquer rapidement le fonctionnement des dispositifs les utilisant, tels que diode, transistor à effet de champ, transistor bipolaire, etc.

3.1.1 Semi-conducteurs intrinsèques

3.1.1.1 Réseau cristallin

Un cristal de semi-conducteur intrinsèque est un solide dont les noyaux atomiques sont disposés aux noeuds d'un réseau géométrique régulier. La cohésion de cet édifice est assurée par les liens de valence qui résultent de la mise en commun de deux électrons appartenant chacun à deux atomes voisins de la maille cristalline. Les atomes de semi-conducteur sont tétravalents^{3.1} et le cristal peut être représenté par le réseau de la figure suivante :



3.1.1.2 Définitions

L'électron qui possède une énergie suffisante peut quitter la liaison de valence pour devenir un *électron libre*. Il laisse derrière lui un *trou* qui peut être assimilé à une charge libre positive ; en effet, l'électron quittant la liaison de

3.1. Chaque atome peut former quatre liaisons de valence. Un atome trivalent peut former trois liaisons, et un atome pentavalent peut former cinq liaisons.

valence à laquelle il appartenait démasque une charge positive du noyau correspondant. Le trou peut être occupé par un autre électron de valence qui laisse, à son tour, un trou derrière lui : tout se passe comme si le trou s'était déplacé, ce qui lui vaut la qualification de charge libre. La création d'une paire électron libre-trou est appelée *génération* alors qu'on donne le nom de *recombinaison* au mécanisme inverse.

La température étant une mesure de l'énergie cinétique moyenne des électrons dans le solide, la concentration en électrons libres et en trous en dépend très fortement.

3.1.1.3 Exemples

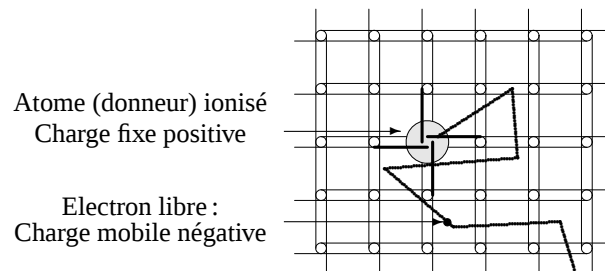
Le silicium a un nombre volumique d'atomes de $5 \cdot 10^{22}$ par cm^3 . A 300K (27°C), le nombre volumique des électrons libres et des trous est de $1,5 \cdot 10^{10} \text{ cm}^{-3}$, soit une paire électron libre-trou pour $3,3 \cdot 10^{12}$ atomes.

Le nombre volumique des atomes dans le germanium est de $4,4 \cdot 10^{22}$ par cm^3 . A 300K, le nombre volumique des électrons libres et des trous est $2,5 \cdot 10^{13} \text{ cm}^{-3}$, soit une paire électron libre-trou pour $1,8 \cdot 10^9$ atomes.

3.1.2 Semi-conducteurs extrinsèques de type n

3.1.2.1 Réseau cristallin

Un semiconducteur dans lequel on aurait substitué à quelques atomes tétravalents des atomes pentavalents est dit *extrinsèque de type n* :



Quatre électrons de la couche périphérique de l'atome pentavalent prennent part aux liens de valence alors que le cinquième, sans attache, est libre de se mouvoir dans le cristal. L'électron libre ainsi créé neutralise la charge positive, solidaire du réseau cristallin, qu'est l'atome pentavalent ionisé.

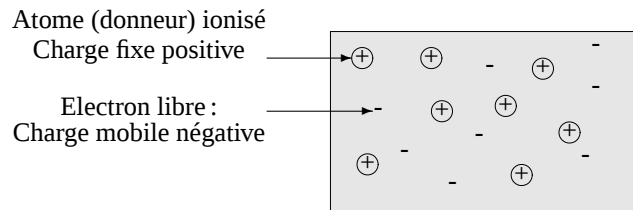
3.1.2.2 Définitions

Le *dopage* est l'action qui consiste à rendre un semiconducteur extrinsèque. Par extension, ce terme qualifie également l'existence d'une concentration d'atomes étrangers : on parle de dopage de type n . On donne le nom d'*impuretés* aux atomes étrangers introduits dans la maille cristalline. Dans le cas d'un semiconducteur extrinsèque de type n , les impuretés sont appelées *donneurs* car chacune d'entre elles donne un électron libre.

3.1.2.3 Modèle

Les dopages courants sont d'environ 10^{16} à 10^{18} atomes par cm^3 . On peut admettre que le nombre volumique des électrons libres est égal au nombre volumique des impuretés et que le nombre volumique des trous (charges libres

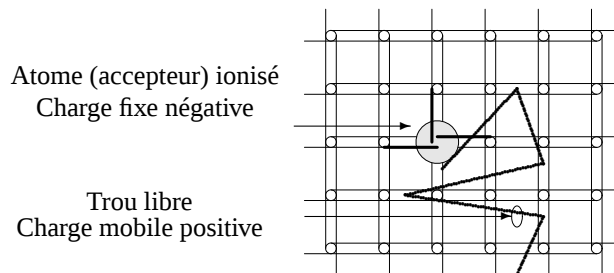
positives) est négligeable. Etant données ces considérations, on établit le modèle de semiconducteur représenté ci-dessous dans lequel n'apparaissent que les charges essentielles, à savoir les électrons libres et les donneurs ionisés. Les charges fixes sont entourées d'un cercle.



3.1.3 Semi-conducteurs extrinsèques de type *p*

3.1.3.1 Réseau cristallin

Si l'on introduit des atomes trivalents dans le réseau cristallin du semiconducteur, les trois électrons de la couche périphérique de l'impureté prennent part aux liens de valence, laissant une place libre. Ce trou peut être occupé par un électron d'un autre lien de valence qui laisse, à son tour, un trou derrière lui. L'atome trivalent est alors ionisé et sa charge négative est neutralisée par le trou (voir figure ci-dessous). Le semi-conducteur est alors dit *extrinsèque de type *p**. Les impuretés, pouvant accepter des électrons, sont appelées *accepteurs*.

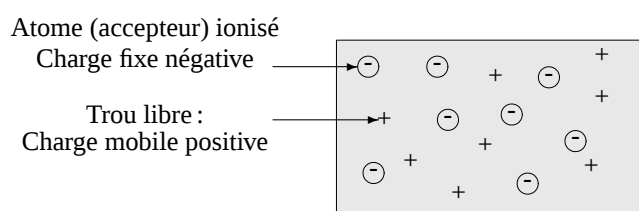


3.1.3.2 Définition

Les impuretés, dans un semi-conducteur extrinsèque de type *p*, sont appelées *accepteurs* au vu de leur propriété d'accepter un électron situé dans un lien de valence.

3.1.3.3 Modèle

On peut faire les mêmes considérations qu'au paragraphe 3.1.2.3 concernant le nombre volumique des trous : il est approximativement égal au nombre volumique des impuretés. Le nombre volumique des électrons libres est alors considéré comme négligeable. Il s'ensuit un modèle, représenté à la figure ci-dessous, dans lequel n'apparaissent que les charges prépondérantes : les trous et les accepteurs ionisés.



Remarque : il faut remarquer que le semiconducteur extrinsèque, type p ou type n , est globalement neutre. On peut le comparer à un réseau géométrique dont certains nœuds sont chargés et dans lequel stagne un « gaz » de charges mobiles qui neutralise les charges fixes du réseau. On élargit, par la suite, la notion de semiconducteur de type n à un semiconducteur dont le nombre volumique des donneurs l'emporte sur celui des accepteurs et celle de semiconducteur de type p à un semiconducteur dans lequel le nombre volumique des accepteurs est prépondérant.

Ce qu'il faut retenir

- la nature d'un semi-conducteur intrinsèque ;
- le dopage (type n et p) et ses conséquences.

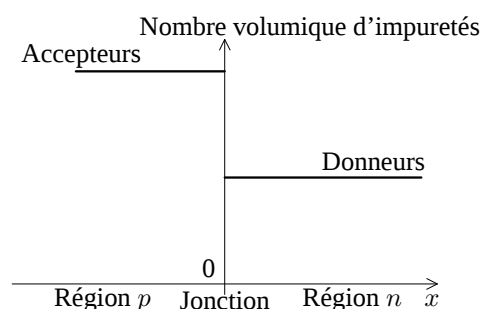
3.2 La jonction PN

3.2.1 Introduction

Le dopage non uniforme d'un semiconducteur, qui met en présence une région de type n et une région de type p , donne naissance à une jonction pn . Une telle jonction est aussi appelée *diode*. Dans la présente section, on étudie, qualitativement, les phénomènes qui ont pour siège la jonction pn . On donne également la relation exponentielle qui lie courant et tension dans une telle jonction.

3.2.2 Description

Soit le semiconducteur à dopage non uniforme ci-dessous qui présente une région p à nombre volumique d'atomes accepteurs constant, suivie immédiatement d'une région n à nombre volumique de donneurs constant également.



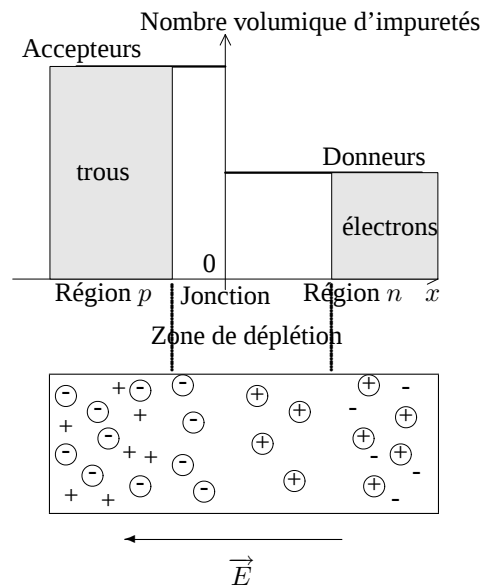
La surface de transition entre les deux régions est appelée jonction pn abrupte. Du fait de la continuité du réseau cristallin, les « gaz » de trous de la région p et d'électrons de la région n ont tendance à uniformiser leur concentration dans tout le volume à disposition. Cependant, la diffusion des trous vers la région n et des électrons libres vers la région p provoque un déséquilibre électrique si bien que, dans la zone proche de la jonction, la neutralité électrique n'est plus satisfaite. On trouve, dans la région p , des atomes accepteurs et des électrons, soit une charge locale négative, et dans la région n , des atomes donneurs et des trous, soit une charge locale positive. Il s'est donc créé un dipôle aux abords de la jonction et, conjointement, un champ électrique. Une fois l'équilibre atteint, ce champ électrique est tel qu'il s'oppose à tout déplacement global de charges libres.

3.2.3 Définitions

La région dans laquelle la neutralité n'est pas satisfaite est appelée *zone de déplétion* ou *zone de charge spatiale* alors que les autres régions sont dites *régions neutres*.

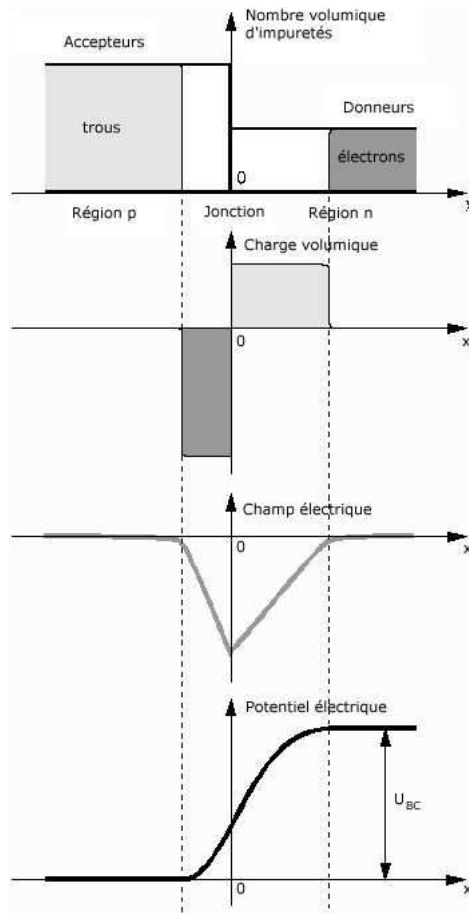
Le champ électrique interne créé par le dipôle est nommé *champ de rétention de la diffusion* car il s'oppose à toute diffusion des charges mobiles.

Remarque : généralement, la concentration des charges mobiles dans la zone de charge spatiale est négligeable vis-à-vis du nombre volumique des charges fixes. On idéalise cet état de fait et l'on admet qu'il n'y a pas de charges mobiles dans la zone de déplétion :



3.2.4 Barrière de potentiel

Il existe, entre la région p et la région n, une barrière de potentiel U_{B0} énergétique pour les charges mobiles. L'existence de cette barrière se traduit par une différence de potentiel électrique liée au champ de rétention de la diffusion :



Exemple : pour une jonction pn au silicium avec un dopage $N_A = 10^{18} \text{ cm}^{-3}$ dans la région p et un dopage $N_D = 10^{17} \text{ cm}^{-3}$ dans la région n, la hauteur de la barrière de potentiel à 300 K (27°C) à l'équilibre vaut 872mV.

Remarque : la hauteur de la barrière de potentiel à l'équilibre est telle que les trous qui sont dans la région p ont une énergie moyenne qui est juste assez insuffisante pour leur interdire de passer la barrière de potentiel. Il en va de même pour les électrons qui se trouvent dans la région n.

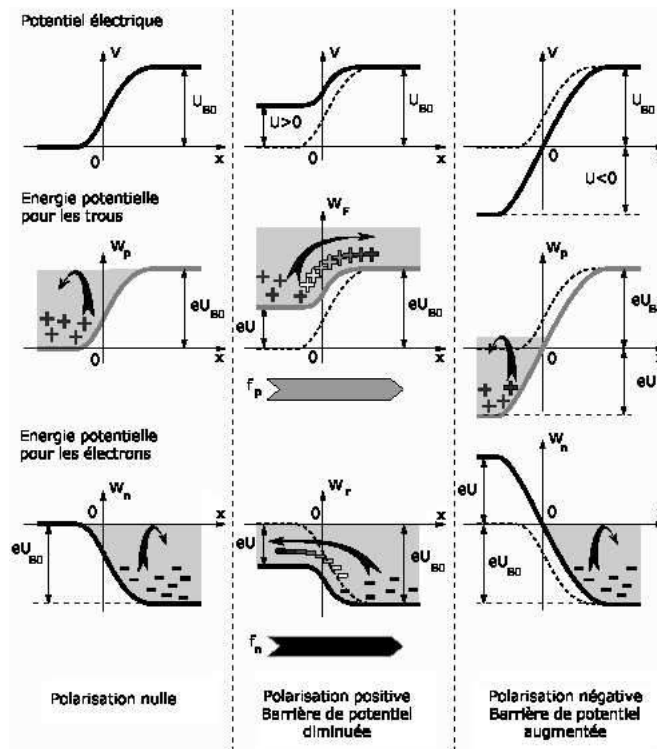
3.2.5 Caractéristique électrique

3.2.5.1 Description

Si l'on applique une tension U à la jonction, cette tension se reporte presque entièrement à la zone de déplétion qui présente une résistivité très grande due à la quasi-absence de charges mobiles. Une tension U négative renforce le champ de rétention de la diffusion et augmente, par conséquent, la hauteur de la barrière de potentiel, de telle sorte qu'aucune charge libre ne traverse la zone de charge spatiale.

Au contraire, si l'on applique une tension U positive, le champ électrique de rétention de la diffusion est diminué et les charges mobiles qui ont une énergie supérieure à celle que représente la hauteur de la barrière de potentiel peuvent traverser la zone de charge spatiale.

Ces situations sont résumées dans le schéma ci-dessous :



3.2.5.2 Définitions

L'application d'une tension qui diminue la hauteur de la barrière de potentiel par rapport à l'équilibre est appelée *polarisation directe* par opposition à la *polarisation inverse* qui augmente la hauteur de la barrière de potentiel par rapport à l'équilibre.

3.2.5.3 Caractéristique et définitions

Une polarisation directe permet le passage d'un courant électrique dans la jonction alors qu'une polarisation inverse l'empêche. *Simultanément*, un « courant de trous » et un « courant d'électrons » se superposent. Le résultat en est un courant unique, et l'on peut montrer qu'il peut s'exprimer sous la forme :

$$I = I_s \left[\exp\left(\frac{U}{nU_T}\right) - 1 \right]$$

... où

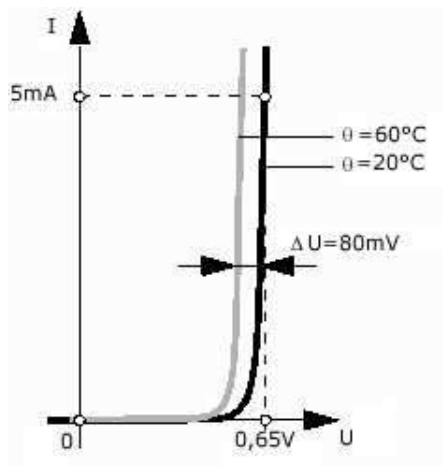
- le courant I_s est appelé *courant inverse de saturation*. C'est la valeur asymptotique du courant traversant la jonction en polarisation inverse ;
- U_T est la *tension thermodynamique* qui vaut

$$U_T = \frac{kT}{e}$$

où k est la constante de Boltzmann, T la température absolue en K et e la charge électrique élémentaire. A 25°C, $U_T = 25\text{mV}$;

- n est le *coefficient d'émission*. Il dépend du matériau, voisin de 1 dans les jonctions de transistors au silicium et dans les diodes au germanium, et compris entre 1 et 2 dans les diodes au silicium.

On obtient donc la caractéristique suivante :



Remarque : le courant inverse de saturation des jonctions au silicium est de l'ordre de grandeur de 10^{-12} à 10^{-15} A de telle sorte qu'on peut généralement le considérer comme nul en polarisation inverse.

Ce qu'il faut retenir

- le principe de la jonction PN ;
- la notion de polarisation (directe, inverse) ;
- la caractéristique courant-tension d'une diode.

3.3 Le transistor bipolaire

3.3.1 Généralités

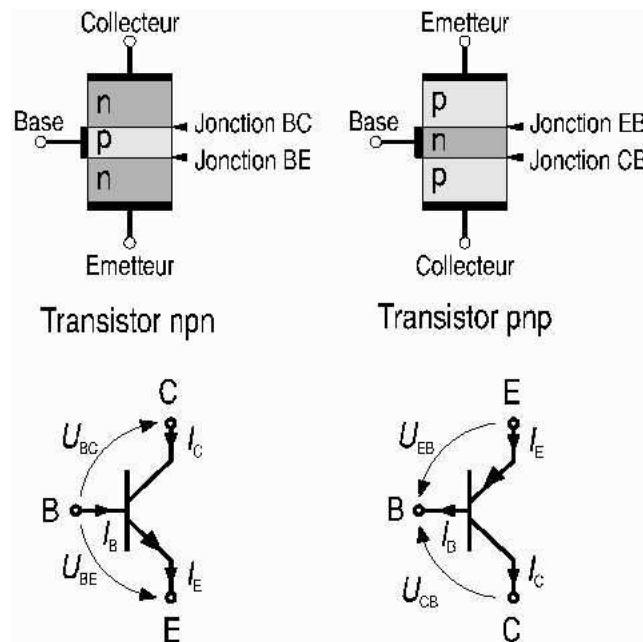
3.3.1.1 Introduction

Le transistor bipolaire est l'un des dispositifs à semiconducteur les plus utilisés à l'heure actuelle dans les rôles d'amplificateur et d'interrupteur. C'est un élément composé de deux jonctions pn.

3.3.1.2 Définitions

Le *transistor bipolaire*^{3.2} est un dispositif présentant trois couches à dopages alternés npn ou pnp :

^{3.2.} Ou *Bipolar Junction Transistor*.



La couche médiane est appelée *base*. Leur géométrie et leur nombre volumique en impuretés distinguent les deux couches externes : *émetteur* et *collecteur*. Par extension, on appelle également base, émetteur et collecteur les trois électrodes qui donnent accès aux trois couches correspondantes.

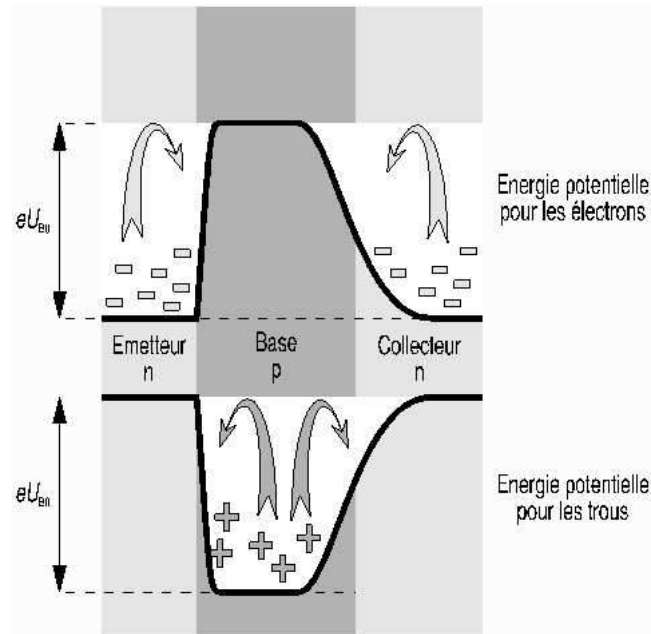
Les deux jonctions qui apparaissent dans le transistor sont désignées par le nom des deux régions entre lesquelles elles assurent la transition : on trouve, par conséquent, la jonction base-émetteur (BE) également dénommée *jonction de commande* et la jonction base-collecteur.

3.3.1.3 Hypothèse

Le principe de superposition s'applique aux charges injectées par la jonction BE et aux charges injectées par la jonction BC. On peut donc étudier séparément l'effet de chaque jonction.

3.3.1.4 Transistor au repos

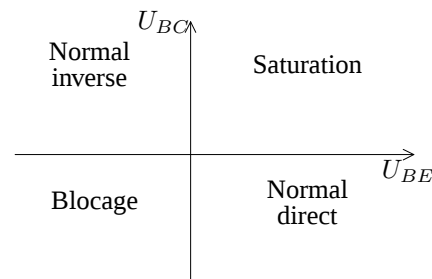
La figure suivante montre les barrières de potentiel énergétique pour les électrons et les trous. Au repos, elles sont telles que ni les électrons de l'émetteur et du collecteur, ni les trous de la base ne peuvent les franchir.



3.3.2 Modes de fonctionnement du transistor

3.3.2.1 Définitions

Les divers cas de fonctionnement du transistor dépendent des valeurs des tensions aux jonctions BE et BC. Si l'on considère l'état bloqué et l'état passant de chaque jonction, on dénombre quatre modes de fonctionnement possibles :



- le transistor est *bloqué* lorsque ses deux jonctions sont en polarisation inverse ;
- le transistor est en *fonctionnement normal direct* lorsque la jonction de commande BE est en polarisation directe et que la jonction BC est en polarisation inverse ;
- le transistor est en *fonctionnement normal inverse* lorsque la jonction de commande BE est en polarisation inverse et que la jonction BC est en polarisation directe ;
- le transistor est *saturé* lorsque ses deux jonctions sont en polarisation directe.

3.3.2.2 Blocage

Aucun courant ne circule dans un transistor bloqué puisque ses deux jonctions sont polarisées en sens inverse. Le transistor se comporte comme un circuit ouvert de telle sorte que le collecteur est isolé de l'émetteur.

3.3.2.3 Fonctionnement normal inverse

La jonction BE détermine le débit des électrons. La jonction BC, polarisée en inverse, n'influence d'aucune manière le débit des électrons. On peut montrer qu'un courant circule alors de l'émetteur vers le collecteur, de la forme

$$I_E = I_{sE} \left[\exp \left(\frac{U_{BE}}{U_T} \right) - 1 \right]$$

où U_T désigne la tension thermodynamique (cf. 3.2.5.2). Un courant s'installe aussi entre base et collecteur :

$$I_B = I_{sB} \left[\exp \left(\frac{U_{BE}}{U_T} \right) - 1 \right]$$

On montre également que :

- le courant base-collecteur est négligeable devant le courant émetteur-collecteur, et que par conséquent le courant « sortant » par le collecteur est approximativement égal au courant « entrant » par l'émetteur ;
- le rapport β entre le courant de collecteur et le courant de base est une constante, caractéristique du transistor, et est appelé *gain de courant en mode direct*, ou *en mode F* (F pour *forward*).

Lors de la fabrication des transistors on met tout en œuvre pour que le courant de base en mode direct soit le plus faible possible. En particulier, l'émetteur est dopé beaucoup plus fortement que la base pour que les électrons injectés dans cette dernière soient plus nombreux que les trous injectés dans l'émetteur. De plus, on réalise des bases aussi étroites que possible de telle sorte que, pendant leur transit, les électrons n'aient que peu de chance de s'y recombiner. Le gain de courant en mode direct atteint des valeurs se situant entre 100 et 1000 pour des transistors de petite puissance (inférieure à 1W).

3.3.2.4 Fonctionnement normal inverse

La jonction BC détermine l'injection dans la base puis dans l'émetteur, indépendamment de la jonction BE. Les électrons de l'émetteur ne peuvent franchir la barrière de potentiel de la jonction BE ; il n'y aura par conséquent aucune influence de la tension U_{BE} sur le débit des électrons. Les relations liant tension et courant sont similaires à celles du mode normal direct, à ceci près que la tension à considérer est U_{BC} . On définit de même le gain β_R (R pour *reverse*) entre le courant de base et celui de collecteur, gain que l'on appelle *gain de courant inverse*, ou *gain de courant en mode R*.

Le gain de courant inverse, du fait de la technologie, est plus petit que le gain de courant direct : dans un transistor discret de petite puissance il peut être compris entre 1 et 10 alors qu'il devient beaucoup plus petit que 1 dans les transistors *intégrés*, c'est-à-dire fabriqués à partir de la même matrice en silicium.

3.3.2.5 Saturation

En saturation, les deux jonctions du transistor conduisent. Il est à remarquer que le courant qui circule de l'émetteur au collecteur est *inférieur* à celui qui circulerait si seulement l'une ou l'autre jonction était polarisée en sens direct sous même tension.

Ce qu'il faut retenir

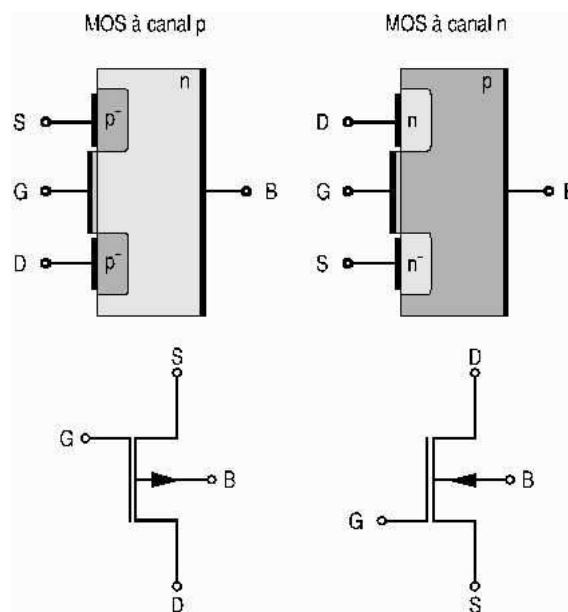
- la nature d'un transistor npn (juxtaposition de deux jonctions) ;
- les modes de fonctionnement (surtout blocage et saturation).

3.4 Le transistor MOS

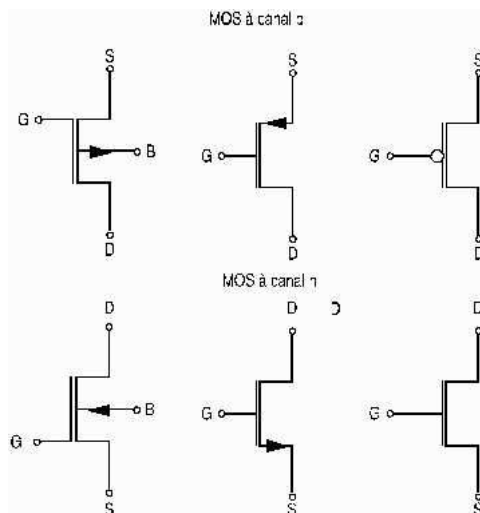
3.4.1 Introduction

En 1930, L. LILIENTHAL de l'Université de Leipzig dépose un brevet dans lequel il décrit un élément qui ressemble au transistor MOS (Metal Oxide Semiconductor) actuel. Cependant, ce n'est que vers 1960 que, la technologie ayant suffisamment évolué, de tels transistors peuvent être réalisés avec succès. En particulier, les problèmes d'interface oxyde-semiconducteur ont pu être résolus grâce à l'affinement de la technologie dans le domaine bipolaire, affinement requis pour obtenir des transistors de meilleure qualité. Aujourd'hui le transistor MOS constitue, par sa simplicité de fabrication et ses petites dimensions, l'élément fondamental des circuits intégrés numériques à large échelle.

3.4.2 Définitions et principe de fonctionnement



Le transistor MOS est un transistor dit « à effet de champ » constitué d'un substrat semiconducteur (B) recouvert d'une couche d'oxyde sur laquelle est déposée l'électrode de *grille* (G). Par le biais d'une différence de potentiel appliquée entre grille et substrat, on crée, dans le semiconducteur, un champ électrique qui a pour effet de repousser les porteurs majoritaires loin de l'interface oxyde-semiconducteur et d'y laisser diffuser des minoritaires venus de deux îlots de type complémentaire au substrat, la source (S) et le drain (D). Ceux-ci forment une couche pelliculaire de charges mobiles appelée *canal*. Ces charges sont susceptibles de transiter entre le drain et la source situés aux extrémités du canal (*cf* figure ci-dessus). Dans cette même figure, on a également représenté les symboles des transistors MOS à canal n et à canal p. La flèche indique le sens de conduction des jonctions substrat-source (BS) et substrat-drain (BD). Sauf près de l'interface oxyde-semiconducteur, ces jonctions sont polarisées en sens inverse. Dans la figure ci-dessous, on a représenté différents symboles couramment utilisés pour les transistors MOS.



Ce qu'il faut retenir

- le principe de fonctionnement d'un transistor à effet de champ ;
 - les symboles d'un transistor MOS.
-

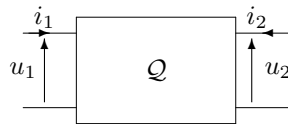
Chapitre 4

Systèmes analogiques

4.1 Représentation quadripolaire

4.1.1 Introduction

Si on veut « cascader »^{4.1} des systèmes, il peut être utile de les modéliser sous la forme de boîtes noires, dont on ne connaîtrait que les paramètres d'entrée/sortie. En pratique, un « signal », en électronique ou en électrotechnique, est soit un courant électrique, soit une tension. Pour pouvoir facilement introduire les paramètres d'entrée dans le cas où le signal est représenté par une tension, on est amené à introduire une représentation dite *quadripolaire*, selon le schéma suivant :



Notez les sens *entrants* des courants !

4.1.2 Matrice de transfert

Si on considère que les grandeurs utiles à connaître sont les tensions d'entrée/sortie, alors les relations d'entrée-sortie peuvent s'écrire :

$$\begin{cases} u_1 &= z_{11}i_1 + z_{12}i_2 \\ u_2 &= z_{21}i_1 + z_{22}i_2 \end{cases}$$

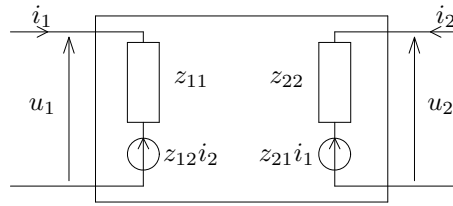
ou, sous une forme matricielle :

$$\begin{pmatrix} u_1 \\ u_2 \end{pmatrix} = \begin{pmatrix} z_{11} & z_{12} \\ z_{21} & z_{22} \end{pmatrix} \cdot \begin{pmatrix} i_1 \\ i_2 \end{pmatrix}$$

- z_{11} est appelée *impédance d'entrée* ;
- z_{12} est appelée *transimpédance directe* ;
- z_{21} est appelée *transimpédance inverse* ;
- z_{22} est appelée *impédance de sortie*

4.1. Autrement dit, brancher des systèmes les uns à la suite des autres...

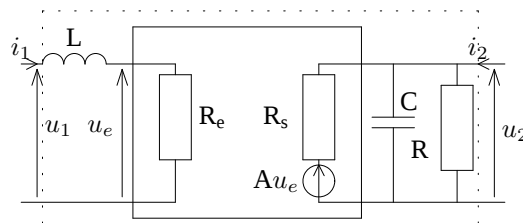
On peut alors modéliser le quadripôle avec le schéma :



Attention ! Il existe des systèmes que l'on ne *peut pas* mettre sous forme quadripolaire.

4.1.3 Exemple

Etudions les relations entrées/sorties du quadripôle suivant :



Pour calculer z_{11} et z_{21} , on fait $i_2 = 0$. On obtient alors :

$$\begin{cases} u_1 = (R_e + jL\omega)i_1 \\ u_e = R_e i_1 \\ u_2 = \frac{A}{(1 + \frac{R_s}{R}) + jR_s C\omega} u_e \end{cases}$$

Et donc :

$$\begin{cases} z_{11} = R_e + jL\omega \\ z_{21} = \frac{A R_e}{(1 + \frac{R_s}{R}) + jR_s C\omega} \end{cases}$$

De même, pour calculer z_{12} et z_{22} on pose $i_1 = 0$ et il vient :

$$z_{12} = 0$$

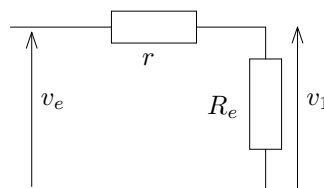
et :

$$z_{22} = (R \text{ en parallèle à } C \text{ en parallèle à } R_s) = \frac{R R_s}{(R + R_s) + jR R_s C\omega}$$

4.1.4 Impédances d'entrée/sortie

On peut essayer de dégager quelques caractéristiques nécessaires pour avoir un « bon » quadripôle.

1. **En entrée :** le quadripôle ne doit pas dégrader les performances de son alimentation. Considérons par exemple un quadripôle, de résistance d'entrée R_e , alimenté par un générateur non idéal, de résistance interne r :



Déterminons R_e de manière à réduire les pertes par effet Joule dans le générateur. Ces pertes valent :

$$P_J = \frac{(V_1 - V_e)^2}{r}$$

(V désignant la valeur efficace de la tension v). Or d'après le théorème du diviseur de tension (cf. paragraphe B.1.1 en annexes) :

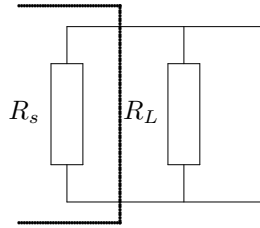
$$V_1 = \frac{R_e}{r + R_e} V_e$$

On remarque d'ailleurs sur cette expression que si l'on veut éviter les chutes de tension parasites (autrement dit, être assuré que la tension en entrée du quadripôle est toujours imposée par la force électromotrice du générateur, et ne dépend pas de la résistance interne de celui-ci), il faut que R_e soit grande devant r .

$$P_J = \frac{r}{(r + R_e)^2} V_e^2$$

Pour minimiser les pertes par effet Joule dans le générateur, on retrouve la même conclusion : *la résistance d'entrée du quadripôle doit être grande devant la résistance du générateur*^{4.2}.

2. **En sortie :** le quadripôle doit présenter la même caractéristique de sortie, quelle que soit la manière dont il est chargé. Considérons donc un quadripôle de résistance de sortie R_s , chargé par la résistance R_L :

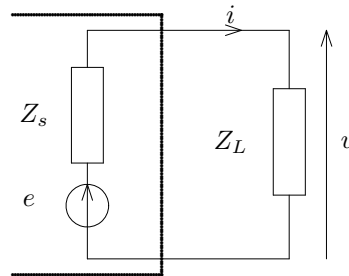


La résistance équivalente à ces deux résistances mises en parallèle vaut :

$$R_{equiv} = \frac{R_s R_L}{R_s + R_L}$$

Si on veut que cette résistance diffère peu de R_L , il est nécessaire que *la résistance de sortie du quadripôle soit petite devant R_L* .

3. **Adaptation en puissance :** les calculs précédents étaient destinés à faciliter l'insertion « transparente » du quadripôle dans un circuit en cascade. Mais on peut également désirer optimiser le transfert de puissance entre la sortie du quadripôle et sa « charge », c'est-à-dire le composant branché en aval. Supposons donc que le quadripôle présente une impédance de sortie Z_s , un comportement en tension en sortie modélisé par la source e , et qu'il débite un courant i dans une charge Z_L aux bornes de laquelle est mesurée la tension v .



On *admet* que la puissance dissipée dans Z_L par effet Joule est égale à :

$$P = \frac{1}{2} \Re(vi^*)$$

où $\Re(z)$ désigne la partie réelle du nombre complexe z , et z^* son conjugué. On a d'après le théorème du diviseur de tension (cf. paragraphe B.1.1) :

$$v = \frac{Z_L}{Z_s + Z_L} e$$

4.2. En fait, on vient d'exprimer fondamentalement la même chose de deux manières différentes. En effet, s'il y a chute de tension, c'est qu'il y a perte d'énergie quelque part *en amont* du quadripôle, c'est-à-dire dans la résistance interne du générateur, sous forme d'un dégagement de chaleur.

et

$$i = \frac{e}{Z_s + Z_L}$$

Il vient donc :

$$P = \frac{1}{2} \Re \left(\frac{Z_L}{|Z_L + Z_s|^2} \right) e e^*$$

Si on écrit $Z_s = R_s + jX_s$ et de même $Z_L = R_L + jX_L$ on obtient :

$$P = \frac{1}{2} \frac{R_L}{(R_s + R_L)^2 + (X_s + X_L)^2} |e|^2$$

Nous avons vu qu'une impédance complexe pouvait avoir une partie imaginaire positive ou négative. On peut donc faire en sorte que $X_s = -X_L$, autrement dit si le quadripôle doit débiter dans une charge inductive, son impédance de sortie doit être de type capacitif, et réciproquement. Ceci étant établi, il ne reste plus alors qu'à maximiser le facteur $\frac{R_L}{(R_L + R_s)^2}$. On montre facilement que cela est réalisé quand $R_s = R_L$ ^{4.3}.

En résumé, pour une adaptation en puissance, il faut que $Z_s = Z_L^*$. Cette relation est en contradiction avec les résultats précédents, et il est à noter également qu'elle ne peut être rigoureusement vérifiée qu'à une fréquence de fonctionnement fixée. En effet, supposons que la charge vaille $1/jC\omega$. Il faut donc que l'impédance de sortie du quadripôle soit égale à $-1/jC\omega$. Cela est possible si on trouve une bobine d'inductance L telle que $L\omega = 1/C\omega$; or il n'est possible de réaliser cette égalité qu'en se plaçant à une pulsation déterminée ω_0 , avec $L = 1/C(\omega_0)^2$. Pour une autre pulsation ω_1 , il n'y aura plus adaptation.

4. **Ligne de transmission:** se reporter à l'annexe D.

Ce qu'il faut retenir

- la représentation quadripolaire : deux fils d'entrée, deux fils de sortie ;
- les règles d'adaptation : faible résistance de sortie et grande résistance d'entrée ou adaptation en puissance.

4.2 Contreréaction

4.2.1 Généralités

4.2.1.1 Introduction

Ce cours a pour but de présenter des notions propres aux systèmes dits « bouclés », que l'on peut rencontrer dans maints domaines. Il s'agit d'une introduction permettant d'établir un certain nombre de concepts qui sont couramment utilisés.

4.3. Considérons en effet la fonction

$$f(x) = \frac{x}{(x+a)^2} \text{ avec } x > 0, a > 0$$

La dérivée de cette fonction vaut

$$f'(x) = \frac{a-x}{(x+a)^3}$$

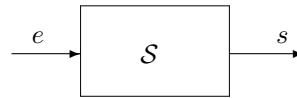
et sa dérivée seconde :

$$f''(x) = \frac{2(x-2a)}{(x+a)^4}$$

La dérivée s'annule en $x = a$ et comme en ce point la dérivée seconde est négative, il s'agit d'un maximum local.

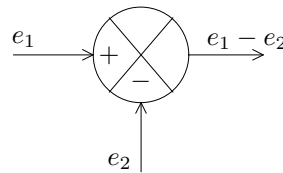
4.2.1.2 Conventions

1. **Les systèmes :** On va dans toute la suite manipuler des « boîtes noires » avec une entrée e et une sortie s , représentées ainsi :



Entrée et sortie peuvent être des tensions ou des courants électriques (cas le plus souvent rencontré) ou bien toute autre sorte de signal : son, onde électromagnétique, déformation mécanique, etc. On symbolise à l'intérieur de la « boîte noire » sa fonction.

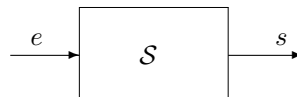
2. **Les opérateurs :** Un opérateur a plusieurs entrées et une sortie. Il réalise une *opération* arithmétique sur les entrées. On représentera ainsi par exemple l'opération « soustraction » de deux signaux e_1 et e_2 par l'élément suivant :



Parmi les opérateurs figurent également l'additionneur et le multiplieur (avec souvent un facteur multiplicatif k supplémentaire).

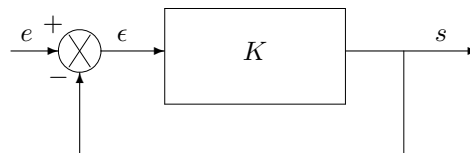
4.2.1.3 Un exemple d'intérêt du bouclage

Considérons un cas simple d'amplification, tel que la sortie s suive précisément les variations de l'entrée e à un facteur de proportionnalité près. On dit que s est *asservie* à e . Soit S un système d'amplification de *gain* K , à l'entrée duquel est injecté le signal e , et à la sortie duquel est mesuré le signal s :



On a donc $s = K.e$. Supposons que l'on ne connaisse le gain K qu'à 10% près : $\frac{\delta K}{K} = 0.1$. Il est immédiat alors que s ne peut être déterminé qu'à 10% près également : $\frac{\delta s}{s} = 0.1$.

L'idée à l'origine du concept de contreréaction est de *comparer* s et e . On construit donc le signal différence entre e et s : $\epsilon = e - s...$



On rappelle que l'opérateur à l'entrée du système est un « soustracteur » et réalise l'opération $e - s$. On a alors les relations suivantes :

$$\begin{cases} \epsilon = e - s \\ s = K.\epsilon \end{cases}$$

On en déduit :

$$s = \frac{K}{1 + K} e$$

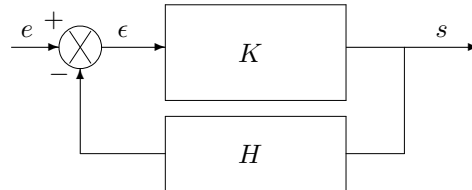
Si on veut que s suive e avec une erreur la plus faible possible, il faut que $K \gg 1$.

Calculons dans ces nouvelles conditions l'influence d'une imprécision portant sur la valeur de K , sur la détermination de s :

$$\frac{\delta s}{s} = \frac{\delta K}{(1+K)^2} \frac{1+K}{K} = \frac{\delta K}{K} \frac{1}{1+K}$$

Pour K suffisamment grand, on voit donc que l'incertitude sur s diminue d'un facteur $1/K$. Par exemple, avec toujours $\frac{\delta K}{K} = 0,1$ et $K = 100$, on obtient $\frac{\delta s}{s} \approx 0,1\%$ seulement.

En règle générale cependant, s et e n'ont aucune raison d'être de même nature physique (par exemple, il peut s'agir d'une tension et d'un courant). Il est nécessaire alors d'introduire un capteur de gain H :



Cet avantage des systèmes bouclés n'est pas le seul. L'idée à l'origine de leur introduction est qu'il est plus facile d'« asservir » un signal à un autre quand on peut les comparer.

4.2.2 Un peu de vocabulaire...

On se servira du schéma précédent pour établir quelques définitions.

4.2.2.1 Les signaux

e est parfois appelé la *commande* ou la *consigne*. En général, on distingue entre ces deux termes en considérant qu'une *commande* est « dynamique » : on est intéressé par la capacité de la sortie à suivre les variations de l'entrée (commande de l'accélération d'un véhicule par la pédale, par exemple), alors qu'une *consigne* est « statique » : ce qui importe est alors que la sortie soit placée dans un état donné (arrêt d'un ascenseur à un étage donné, ou pour rester dans le domaine de l'automobile, garage sur une place de parking).

ϵ , signal différence entre l'entrée et la sortie (à un gain près) est naturellement appelé l'*erreur*.

4.2.2.2 Les « branches » de la boucle

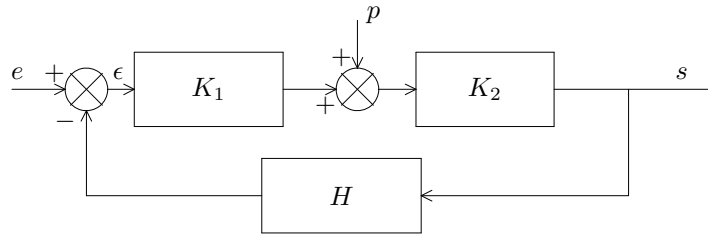
La branche où se trouve K est la *branche* ou *chaîne directe*, celle où se trouve H est la *branche* ou *chaîne inverse*, ou *de retour*.

4.2.2.3 Les gains

Le rapport s/ϵ est le *gain en boucle ouverte* ou *gain direct* et le rapport s/e le *gain en boucle fermée*. Le gain de la chaîne de retour est le *gain de contre-réaction*. Dans le schéma précédent, K est ainsi le gain en boucle ouverte, H le gain de contre-réaction et $K/(1+KH)$ le gain en boucle fermée.

4.2.3 Influence d'une perturbation

Supposons que « quelque part » dans le montage soit introduite une perturbation p *additive*, par exemple :



On a :

$$\begin{cases} \epsilon = e - Hs \\ s = K_2 p + K_1 K_2 \epsilon \end{cases}$$

On en déduit :

$$s = \underbrace{\frac{K_1 K_2}{1 + K_1 K_2 H}}_{\text{terme d'asservissement}} e + \underbrace{\frac{K_2}{1 + K_1 K_2 H}}_{\text{terme de régulation}} p$$

Pour que p soit négligeable, il est nécessaire que le terme d'asservissement soit grand devant celui de régulation, et donc que $K_1 \gg 1$. On aura donc tout intérêt à faire en sorte que le premier « étage » de la chaîne directe soit un étage de pure amplification, avec un gain important. Ceci est *une* des raisons pour lesquelles on place souvent un « préamplificateur » en tête de chaîne.

4.2.4 Exemples de systèmes à contreréaction

4.2.4.1 Exemple détaillé : une file de voitures sur l'autoroute

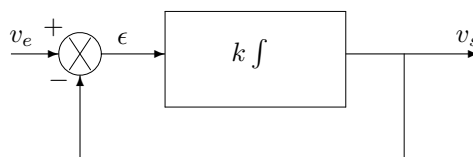
Prenons par exemple une file de voitures roulant sur une autoroute, et considérons le système composé d'une voiture et de son conducteur. Le conducteur observe en permanence la différence de vitesse entre sa propre voiture et celle qui le précède. Si on note v_e la vitesse de cette dernière, et v_s la vitesse de son propre véhicule, on peut dire qu'au bout d'un temps dt , le conducteur corrige sa propre vitesse d'un terme proportionnel à la différence de vitesses qu'il observe :

$$v_s(t + dt) = v_s(t) + k(v_e - v_s)dt \text{ avec } k > 0$$

Cette relation devient :

$$\frac{dv_s}{dt} = k(v_e - v_s)$$

Le système ainsi constitué peut se mettre sous la forme du schéma suivant, comportant une contreréaction :



... où $(k \int)$ désigne un intégrateur-multiplieur par k .

4.2.4.2 Autres exemples

On peut signaler également :

- la régulation du nombre de prédateurs par celui des proies disponibles ;
- le contrôle de la fréquence d'un quartz dans une horloge ;
- la synthèse d'un oscillateur électronique ;
- l'écriture (le retour d'informations se fait par la vue) ;
- et bien d'autres encore...

Ce qu'il faut retenir

- les conventions pour les schémas-blocs: systèmes, opérateurs... ;
- le vocabulaire de la contre-réaction

4.3 Diagramme de Bode ; Gabarit

4.3.1 Diagramme de Bode

4.3.1.1 Définition

Partons d'une fonction de transfert en régime harmonique, $H(j\omega)$, liant deux tensions^{4.4}. Comment représenter cette fonction de la fréquence ? Pour des raisons de commodité, on est amené à écarter toute représentation tridimensionnelle (par exemple, en x , la fréquence, en y , la partie réelle et en z , la partie imaginaire de la fonction). Par ailleurs, on est plus intéressé par la manière dont le filtre *amplifie* le signal à une fréquence donnée, et de combien il le déphase, plutôt que par la partie réelle et la partie imaginaire de la fonction de transfert, dont la signification physique pour un signal quelconque est moins évidente. On est donc amené à décrire la fonction de transfert sous la forme de deux diagrammes :

- un diagramme présentant son module en fonction de la fréquence ;
- un deuxième diagramme présentant sa phase.

Pour que le diagramme recouvre plus facilement la totalité du spectre, et qu'il soit également plus lisible en ordonnée, on utilise en abscisse et en ordonnée des coordonnées logarithmiques, et on définit le *gain en décibels*, G_{dB} , par :

$G_{dB} \triangleq 20 \log |H(j\omega)|$. On introduit aussi parfois l'*atténuation* dans le cas de systèmes à gain inférieur à 1, définie comme l'inverse du gain, $A = 1/G$, soit en dB $A_{dB} = -G_{dB}$.

Remarque: On peut également introduire la notion d'une fonction de transfert liant deux puissances, définie par $H_P(j\omega) = P_s/P_e$. Comme la puissance est proportionnelle au carré de la tension, cette fonction de transfert est proportionnelle à $H(j\omega)^2$, et on définit alors le gain en puissance en décibels par $G_{dB} = 10 \log |H_P(j\omega)|$.

4.4. cf. paragraphe 1.3.3.

4.3.1.2 Exemple

Considérons la fonction de transfert suivante ($A > 0$) :

$$H(j\omega) = A \frac{1 + j \frac{\omega}{\omega_0}}{1 + j \frac{\omega}{\omega_1}}$$

On a donc :

$$G_{dB} = 20 \log A + 10 \log \left[1 + \left(\frac{\omega}{\omega_0} \right)^2 \right] - 10 \log \left[1 + \left(\frac{\omega}{\omega_1} \right)^2 \right]$$

et :

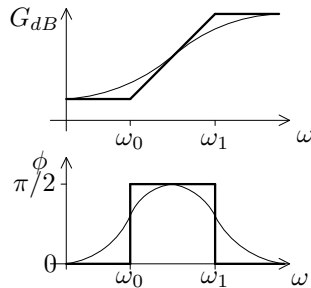
$$\phi = \arctan \frac{\omega}{\omega_0} - \arctan \frac{\omega}{\omega_1}$$

On suppose également pour simplifier que les deux pulsations ω_0 et ω_1 sont très différentes. Deux cas se présentent alors :

– $\omega_0 \ll \omega_1$

- Pour $\omega \ll \omega_0$, on a $G_{dB} \approx 20 \log A$;
- Pour $\omega_0 \ll \omega \ll \omega_1$, on a $G_{dB} \approx 20 \log A + 20 \log \frac{\omega}{\omega_0}$;
- Pour $\omega_1 \ll \omega$, on a $G_{dB} \approx 20 \log A + 20 \log \frac{\omega}{\omega_0} - 20 \log \frac{\omega}{\omega_1}$.

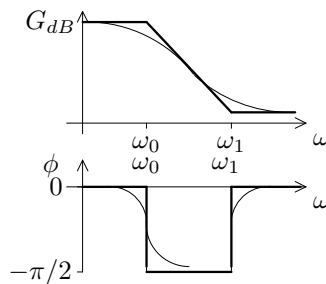
On obtient donc les diagrammes suivants :



– $\omega_0 \gg \omega_1$

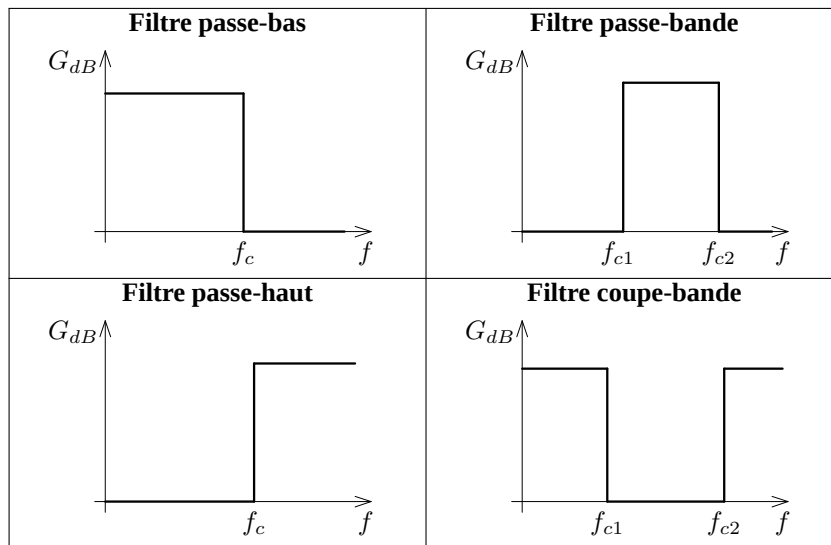
- Pour $\omega \ll \omega_1$, on a $G_{dB} \approx 20 \log A$;
- Pour $\omega_1 \ll \omega \ll \omega_0$, on a $G_{dB} \approx 20 \log A - 20 \log \frac{\omega}{\omega_0}$;
- Pour $\omega_0 \ll \omega$, on a $G_{dB} \approx 20 \log A + 20 \log \frac{\omega}{\omega_0} - 20 \log \frac{\omega}{\omega_1}$.

On obtient donc les diagrammes suivants :



4.3.1.3 Les types de filtres

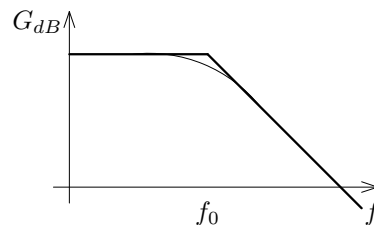
Les filtres fréquentiels sont principalement de 4 types :



1. **Filtre passe-bas** : la fonction de transfert la plus simple pour un tel filtre est du type

$$H(j\omega) = \frac{A}{1 + j\omega/\omega_0}$$

Etude asymptotique : Quand $\omega \ll \omega_0$, on a $G_{dB} \approx 20 \log A = \text{cte}$ et quand $\omega \gg \omega_0$ il vient $G_{dB} \approx 20 \log A - 20 \log (\omega/\omega_0)$. On obtient le diagramme suivant :



La fréquence f_0 , pour laquelle le module de la fonction de transfert vaut $(\text{Valeur max})/\sqrt{2}$, est appelée *fréquence de coupure*. Elle correspond à une perte d'un facteur 2 sur le gain maximal en puissance. D'autre part, quand la fréquence est augmentée d'un facteur 10, le gain en décibels est diminué de 20 unités : on dit que la pente du filtre est de *-20dB par décade*, et le filtre est dit du *premier ordre*. Un filtre passe-bas du deuxième ordre possède une fonction de transfert de la forme suivante :

$$H(j\omega) = \frac{A}{(1 + j\omega/\omega_0)^2}$$

La pente est alors de -40dB par décade.

2. **Filtre passe-bande** : la fonction de transfert la plus simple pour un tel filtre est du type

$$H(j\omega) = A \frac{1 + j\omega/\omega_0}{(1 + j\omega/\omega_1)(1 + j\omega/\omega_2)}$$

avec $\omega_0 < \omega_1 < \omega_2$. Il y a cette fois-ci *deux* fréquences de coupure ; l'intervalle de fréquence entre ces deux bornes est la *bande passante*.

3. **Filtre passe-haut** : la fonction de transfert la plus simple pour un tel filtre est du type

$$H(j\omega) = A(1 + j\omega/\omega_0)$$

C'est un *filtre passe-haut du premier ordre* (pente de +20dB par décade).

4. **Filtre coupe-bande** : La fonction de transfert la plus simple pour un tel filtre est du type

$$H(j\omega) = A \frac{(1 + j\omega/\omega_1)(1 + j\omega/\omega_2)}{1 + j\omega/\omega_0} \text{ avec } \omega_0 < \omega_1 < \omega_2$$

Il est à noter que dans la réalité, tout filtre coupe les hautes fréquences, même un filtre dit passe-haut. Dans les filtres actifs^{4.5}, les transistors présentent toujours des capacités parasites qui ont pour effet d'introduire de hautes fréquences de coupure. Dans les filtres passifs (à base uniquement de circuits R,L,C), il y a aussi toujours des capacités parasites, aux points de soudure des composants par exemple.

4.3.2 Gabarit

Afin de pouvoir spécifier clairement et sans ambiguïté les besoins des utilisateurs, par exemple, ou les caractéristiques techniques d'un système, des définitions ont été énoncées pour les filtres fréquentiels. Nous avons déjà parlé de quelques-unes d'entre elles dans le paragraphe précédent (fréquence de coupure, bande passante). Nous allons maintenant en faire une description plus détaillée, en prenant pour exemple un filtre passe-bas.

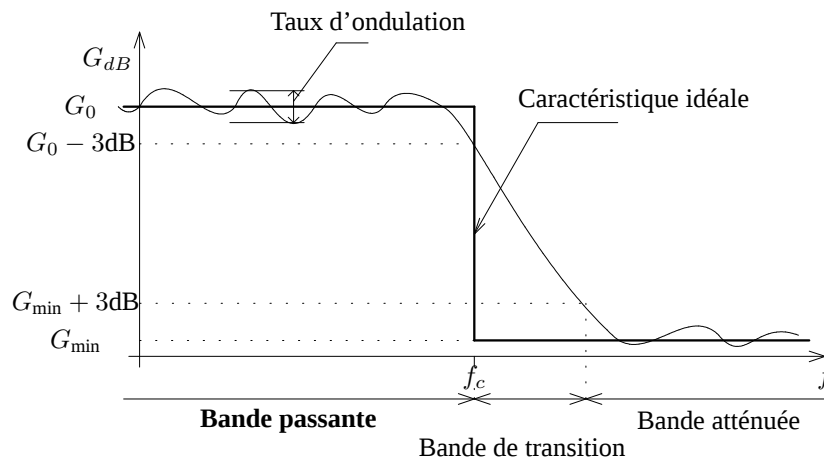


FIG. 4.1 – Gabarit de filtre passe-bas.

Notations :

- H_0 désigne le gain maximal et $G_0 = 20 \log H_0$;
- $H_{\min}$ désigne le gain minimal éventuel^{4.6} et $G_{\min} = 20 \log H_{\min}$;
- f_c désigne la fréquence de coupure.

Définitions : Voici les paramètres que l'on doit spécifier en général pour définir un filtre passe-bas :

- La *bande passante* est définie comme étant le domaine de fréquences où le gain en décibel est supérieur ou égal à $G_0 - 3\text{dB}$;
- La *bande de transition* est définie comme étant le domaine de fréquences où le gain en décibels est compris entre $G_{\min} + 3\text{dB}$ et $G_0 - 3\text{dB}$;
- La *bande atténuée* est définie comme étant le domaine de fréquences où le gain en décibels est inférieur ou égal à $G_{\min} + 3\text{dB}$;
- L'*atténuation de la bande atténuée* vaut $G_0 - G_{\min}$ (en dB) ;

4.5. cf. Note 2.11.

4.6. Certains filtres possèdent un gain minimal en dB égal à $-\infty$.

- La *fréquence de coupure* est définie comme étant la fréquence pour laquelle le gain vaut $G_0 - 3\text{dB}$. On peut parler de fréquence de coupure principale et de fréquence de coupure secondaire si le gain présente deux plateaux de hauteurs différentes ;
- On doit également préciser le *taux d'ondulation* (différence, en dB, entre l'amplitude des oscillations et le gain dans la bande passante) dans la bande passante et éventuellement dans la bande atténuée.

Ce qu'il faut retenir

- La définition du diagramme de Bode ;
- Les différents types de filtres : passe-bas, passe-bande, coupe-bande, passe-haut ;
- les définitions des termes utilisés : fréquence de coupure, bande passante... ;
- La notion de gabarit utilisé pour dresser la liste des spécifications d'un filtre.

4.4 Bruit dans les composants

Les informations qui nous parviennent sont souvent détériorées par des parasites, qui peuvent être dûs à plusieurs causes. Des outils ont été développés afin de pouvoir mieux estimer les contributions parasites, et essayer de s'en affranchir. Ces outils sont basés sur des notions de statistiques, les bruits étant généralement en effet des processus aléatoires.

4.4.1 Densité spectrale de puissance

On rappelle que le spectre^{4.7} d'un signal est le module de sa transformée de Fourier. On définit la *densité spectrale de puissance* comme étant le *carré du module de la transformée de Fourier*. Ainsi, si x est un signal et X sa transformée de Fourier, sa densité spectrale de puissance vaut $D_x = |X(\nu)|^2$.

Il existe une autre expression de la densité spectrale de puissance. Introduisons la notion de fonction d'*autocorrélation* d'un signal x à temps continu :

$$\gamma(\tau) = \int_{-\infty}^{+\infty} x^*(t)x(t+\tau)dt \text{ où } * \text{ est la conjugaison complexe.}$$

Prise au point τ , cette fonction mesure en quelque sorte la manière dont les structures que l'on peut voir dans un signal se répètent sur des échelles de temps de τ . Calculons sa transformée de Fourier $\Gamma(\nu)$:

$$\Gamma(j\omega) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} x^*(t)x(t+\tau)e^{-j\omega\tau}dtd\tau$$

Cette expression peut se mettre sous la forme :

$$\Gamma(j\omega) = \int_{-\infty}^{+\infty} \left(\int_{-\infty}^{+\infty} x(t+\tau)e^{-j\omega(t+\tau)}d\tau \right) x^*(t)e^{+j\omega t}dt$$

On effectue dans l'intégrale centrale le changement de variable $u = t + \tau$ et il vient :

$$\Gamma(j\omega) = \int_{-\infty}^{+\infty} \left(\int_{-\infty}^{+\infty} x(u)e^{-j\omega u}du \right) x^*(t)e^{+j\omega t}dt$$

4.7. cf. le paragraphe 1.2.1.2.

Soit encore :

$$\Gamma(j\omega) = X(j\omega) \int_{-\infty}^{+\infty} x^*(t) e^{+j\omega t} dt$$

On effectue le changement de variable $u = -t$ et on obtient :

$$\Gamma(j\omega) = X(j\omega) \int_{-\infty}^{+\infty} x^*(-u) e^{-j\omega u} du$$

On reconnaît dans le deuxième terme la transformée de Fourier de $x^*(-t)$. Or d'après la propriété 1.11, la transformée de Fourier de x^* vaut $X^*(-\nu)$, et d'après 1.10, la transformée de Fourier de $x(-t)$ vaut $X(-\nu)$. Le deuxième terme vaut donc $X^*(j\omega)$, donc $\Gamma(j\omega) = X(j\omega)X^*(j\omega) = |X(j\omega)|^2$: la densité spectrale de puissance est aussi la transformée de Fourier de l'autocorrélation.

4.4.2 Les types de bruit

Nous nous limiterons dans tout ce qui suit aux seuls bruits *additifs* : la puissance totale transportée est égale à la somme de la puissance de signal utile et à la puissance transportée par le bruit.

4.4.2.1 Bruit thermique

- Egalement nommé *bruit de résistance*, ou *bruit Johnson*, du nom du physicien Johnson qui l'a mis en évidence en 1927 .
- L'étude théorique en a été faite en 1928 par Nyquist. Quand un corps est porté à une certaine température, les noyaux atomiques mais surtout les électrons qui le composent (en raison de leur plus faible masse) sont agités, et dotés d'une vitesse en moyenne nulle (ils ne vont en moyenne dans aucune direction particulière), mais dont la moyenne quadratique (c'est-à-dire la racine carrée de la moyenne des carrés des vitesses) est proportionnelle au produit de la température, exprimée en degrés Kelvin, et d'une constante k , appelée constante de Boltzmann, qui vaut $k = 1.38.10^{-23} J/K$:

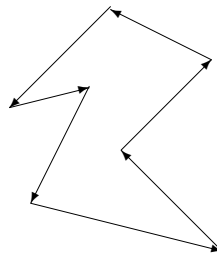
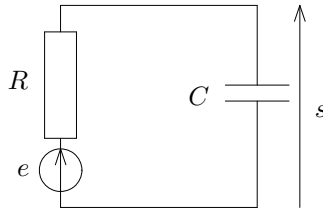


FIG. 4.2 – Déplacement d'un électron : l'électron revient *en moyenne* à son point de départ, donc sa vitesse moyenne est nulle ; mais il se déplace, donc en moyenne, sa vitesse n'est pas nulle : on estime cette dernière valeur par la vitesse quadratique moyenne.

- Pour une résistance R portée à la température T , la densité spectrale de puissance du bruit vaut $D_R = 2kRT$ ^{4.8}. Elle s'exprime en Volts au carré par Hertz (V^2/Hz). Ce bruit est dit *blanc*, par analogie avec la lumière, car toutes les fréquences sont également représentées dans le spectre. Cela n'est pas rigoureusement exact (l'énergie transportée par un tel signal serait infinie), mais cette approximation est tout à fait valable dans les domaines de fréquences où l'on travaille habituellement.

4.8. On peut aussi trouver $D_R = 4kRT$. Tout est une question de définition de la transformée de Fourier. Nous utilisons dans ce cours la définition dans laquelle les bornes d'intégration de l'intégrale généralisée sont $-\infty$ et $+\infty$ (cf paragraphe 1.2.1.2). Cette définition est dite *bilatérale*. On peut également définir une autre TF, où le domaine d'intégration s'étend de 0 à $+\infty$ seulement. Cette TF est dite *monolatérale*, mais ne vérifie pas tout à fait les mêmes propriétés.

- Exemple : circuit RC . Considérons le circuit suivant :



Il est facile de montrer que la fonction de transfert de ce filtre vaut :

$$H(j\omega) = \frac{1}{1 + j\omega/\omega_0}$$

avec $\omega_0 = 1/RC$. La résistance « bruyante » peut être modélisée comme étant la mise en série d'une résistance parfaite, non bruyante, et d'une source e_b délivrant une tension dont la densité spectrale de puissance est celle du bruit. Cette source de tension est filtrée de la même manière par le circuit. La composante bruitée s_b de la sortie vaut donc à la fréquence ν :

$$s_b(t) = \frac{1}{1 + j\nu/\nu_0} e_b(t)$$

Calculons la transformée de Fourier ; il vient :

$$\text{TF}[s_b(t)] = \text{TF}\left(\frac{1}{1 + j\nu/\nu_0} e_b(t)\right) = \frac{1}{1 + j\nu/\nu_0} \text{TF}[e_b(t)]$$

Prenons-en le module au carré ; le terme de gauche devient la densité spectrale de puissance du bruit en sortie, et la Transformée de Fourier de droite la densité spectrale de puissance de l'entrée bruitée, donc du bruit thermique dû à la résistance :

$$D_s(\nu) = \frac{1}{1 + (\nu/\nu_0)^2} D_e(\nu)$$

Application numérique : dans notre cas : $R = 10 \text{ k}\Omega$, $T = 300 \text{ K}$, $C = 1,6 \text{ nF}$. La fréquence de coupure vaut alors $\nu_0 \approx 10 \text{ kHz}$. La puissance totale transportée par le bruit vaut $\int_{-\infty}^{+\infty} D_s(\nu) d\nu$, soit $2kRT\nu_0\pi$.

En règle générale, on dira en fait que la puissance de bruit totale vaut en première approximation *la densité spectrale de puissance de bruit en entrée, multipliée par la bande passante du système* (ici ν_0). Avec les valeurs numériques choisies, on obtient donc environ $2kRT\nu_0 \approx 10^{-12} \text{ V}^2$. Le bruit uniquement dû à cette résistance est donc équivalent à une source de tension moyenne d'environ $1 \text{ }\mu\text{V}$.

4.4.2.2 Bruit de grenaille

- Egalement nommé « shot noise ».
- Il est causé par des discontinuités du débit des porteurs de charge (le plus souvent des électrons^{4.9}), dues à des effets quantiques.
- Il est modélisé par une source de courant, placée en parallèle du composant idéal non bruyant, et de densité spectrale de puissance $D_I = eI$ où I désigne le courant moyen qui parcourt le composant^{4.10}.

4.4.2.3 Bruit en $1/f$

- Egalement nommé « flicker noise », bruit de scintillement ou de papillotement, bruit en excès, bruit de basse fréquence, ou bruit rose.
- Il est toujours présent dans les composants actifs et dans certains composants passifs. Ses origines sont variées : il peut être dû à des impuretés dans le matériau pour un transistor, par exemple, qui libèrent aléatoirement des porteurs de charge, ou bien à des recombinaisons électron-trou parasites (cf. note 4.9), etc.

4.9. Dans les composants actifs, il existe d'autres porteurs de charge : les « trous », qui ne sont en fait que des absences d'électrons, se déplacent également et sont porteurs d'une charge élémentaire positive. On appelle « recombinaison électron-trou » la rencontre d'un électron et d'un trou, résultant en la disparition des deux porteurs de charge.

4.10. ... ou $D_I = 2eI$ dans le cas d'une TF monolatérale.

- La densité spectrale est de la forme

$$D_{1/f} = K_1 \frac{I^\alpha}{f^\beta}$$

avec $0.5 < \alpha < 2$ et $0.8 < \beta < 1.3$, cet exposant étant le plus souvent voisin de 1. K_1 est une caractéristique du composant.

- Remarques :

- si $\beta = 1$, la puissance de bruit par décade est constante. En effet, pour toute fréquence f_0 ,

$$P_{\text{décade}} = \int_{f_0}^{f_0 \cdot 10} D_{1/f} df = K_1 I^\alpha \int_{f_0}^{f_0 \cdot 10} \frac{df}{f} = K_1 I^\alpha (\ln \left(\frac{10f_0}{f_0} \right))$$

Soit $P_{\text{décade}} = K_1 I^\alpha \ln 10 = \text{cte}$

- On note un apparent paradoxe en $f = 0$ Hz, où la densité spectrale de puissance devient infinie. En pratique, aucun système n'a une bande passante s'étendant jusqu'à la fréquence nulle : cela signifierait que la durée de fonctionnement de ce système est infinie.

4.4.2.4 Bruit en créneaux

- Egalement nommé « burst noise », ou bruit popcorn, ou crépitement .
- L'origine de ce bruit est mal comprise. Il semblerait lié à la contamination par des ions métalliques des semi-conducteurs qui composent les éléments actifs.
- Ce bruit est appelé « bruit en créneaux » car les formes d'onde qu'il produit ressemblent à des signaux carrés bruités, de fréquence variable.
- La plus grande partie du spectre de ce bruit se situe dans le domaine des fréquences audibles (de quelques centaines de Hz à quelques dizaines de kHz). La densité spectrale de puissance est de la forme suivante :

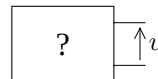
$$D_b = K_2 \frac{I^\gamma}{1 + (f/f_c)^2}$$

où $0,5 < \gamma < 2$, la fréquence de coupure f_c et la constante K_2 étant caractéristiques du composant.

4.4.3 Bruit dans un dipôle

4.4.3.1 Température équivalente de bruit

On considère une « boîte noire » à la sortie de laquelle on observe la différence de potentiel à vide v (sans excitation extérieure) :



La tension mesurée est uniquement due aux différentes sources de bruit internes au dipôle. On introduit la *température équivalente de bruit* T_{eq} dans le dipôle en posant que la puissance totale du bruit dans la bande de travail B ^{4.11} est égale à la puissance de bruit thermique dissipée par la résistance équivalente R_{eq} du dipôle (ie la partie réelle de l'impédance du schéma équivalent de Thévenin^{4.12}), dans cette même bande, portée à une certaine température T_{eq} :

$$(\bar{v}^2)_B = 2kT_{\text{eq}}R_{\text{eq}}B$$

4.11. C'est-à-dire la gamme de fréquences à laquelle le montage est limité.

4.12. cf. le paragraphe B.3.1 en annexes.

4.4.3.2 Rapport de bruit

On introduit aussi parfois la notion de *rapport de bruit* $\mathcal{R}$ pour un dipôle, en le définissant comme le rapport entre la puissance de bruit mesurée dans la bande B et une source de bruit purement thermique étalon, portée à une température de référence T_0 (souvent 300 K) :

$$\mathcal{R} = \frac{(\bar{v}^2)_B}{(\bar{e}^2)_{B,T_0}} = \frac{T_{\text{eq}}}{T_0}$$

4.4.4 Facteur de bruit

4.4.4.1 Définition

On considère un quadripôle $\mathcal{Q}$. On note $\mathcal{S}_e$ (respectivement $\mathcal{S}_s$) la puissance utilisable de signal à l'entrée (resp. en sortie), et $\mathcal{N}_e$ (resp. $\mathcal{N}_s$) la puissance de bruit à l'entrée (resp. en sortie).



On définit le *rapport Signal sur Bruit*, ou rapport Signal à bruit^{4.13} σ par :

$$\sigma = \frac{\mathcal{S}}{\mathcal{N}}$$

Ce rapport permet d'estimer le degré de « contamination » du signal intéressant par le bruit. On définit ensuite le *facteur de bruit* $\mathcal{F}$ par :

$$\mathcal{F} = \frac{\mathcal{S}_e \mathcal{N}_s}{\mathcal{N}_e \mathcal{S}_s} = \frac{\sigma_e}{\sigma_s}$$

4.4.4.2 Température de bruit

On peut exprimer le facteur de bruit d'une autre façon :

$$\mathcal{F} = \frac{\mathcal{S}_e \mathcal{N}_s}{\mathcal{S}_s \mathcal{N}_e}$$

Or $\mathcal{S}_e/\mathcal{S}_s = 1/G$ où G est le gain *en puissance* du quadripôle^{4.14}. De plus, $\mathcal{N}_s$ peut s'écrire sous la forme $\mathcal{N}_s = G\mathcal{N}_e + \mathcal{N}_Q$, où $\mathcal{N}_Q$ désigne la puissance de bruit uniquement produite par le quadripôle. Il vient donc :

$$\mathcal{F} = \frac{1}{G} \left(G + \frac{\mathcal{N}_Q}{\mathcal{N}_e} \right) = 1 + \frac{1}{\mathcal{N}_e} \left(\frac{\mathcal{N}_Q}{G} \right)$$

$\mathcal{N}_Q/G$ désigne en fait une puissance de bruit *interne* au quadripôle, mais « ramenée en entrée » par le gain G . On note cette puissance de bruit $\mathcal{N}_r$: $\mathcal{N}_r = \mathcal{N}_Q/G$.

On considère alors en entrée une source de bruit de référence, délivrant la puissance $\mathcal{N}_e = kT_0B$, définie sur une bande passante unité (ie $B = 1$ Hz), T_0 étant une température de référence (la plupart du temps 300 K). Cette

4.13. De l'anglais *Signal to Noise Ratio* ; on le note d'ailleurs souvent SNR.

4.14. Ce gain *en puissance* est le carré du gain *en tension*.

puissance de bruit est supposée être transmise intégralement au quadripôle^{4.15}. Il vient alors :

$$\mathcal{F} = 1 + \frac{\mathcal{N}_r}{kT_0B}$$

On définit enfin la *température de bruit* T_r ramenée en entrée du quadripôle par $\mathcal{N}_r = kT_rB$, et il vient

$$\mathcal{F} = 1 + \frac{T_r}{T_0} \quad (4.1)$$

Remarques :

- Plus un quadripôle est « bruyant », plus la puissance de bruit qu’il produit est importante, et plus sa température de bruit est grande ;
- *A priori*, la température de bruit dépend de la fréquence.

4.4.4.3 Facteur de bruit d’un quadripôle passif

Pour un quadripôle passif, les composants utilisés sont les seules sources de bruit, puisqu’il n’y a pas d’alimentation électrique. Le bruit résultant est donc dans la majorité des cas purement d’origine thermique. Qui plus est, si l’ensemble du système étudié est placé à la même température, cette température est nécessairement la température de bruit.

Quand le quadripôle de gain G est porté à la température de référence T_0 , le facteur de bruit vaut :

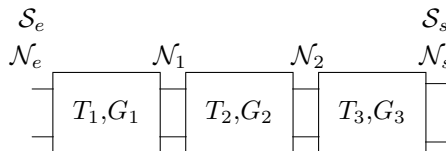
$$\mathcal{F} = \frac{S_e}{\mathcal{N}_e} \frac{\mathcal{N}_s}{S_s} = \frac{\sigma_e}{\sigma_s} = \frac{1}{G} \frac{kT_0B}{kT_0B} = \frac{1}{G}$$

Si on note L l’atténuation du quadripôle ($L = 1/G$), alors le facteur de bruit vaut simplement

$$\mathcal{F} = L$$

4.4.4.4 Théorème de Friiss

Considérons par exemple une chaîne de trois quadripôles en cascade, chacun d’eux étant caractérisé par sa température de bruit et son gain en puissance. Ces quadripôles sont supposés adaptés en puissance en entrée et en sortie (cf. paragraphe 4.1.4) :



Hypothèses et définitions : chaque quadripôle possède une température de bruit T_i , une température de bruit ramenée en entrée T_{ri} , un gain en puissance G_i . La puissance totale du bruit en sortie du premier quadripôle est $\mathcal{N}_1$, en sortie du deuxième $\mathcal{N}_2$ et en sortie de la chaîne $\mathcal{N}_s$. $\mathcal{N}_{Qi}$ désigne le bruit produit par le quadripôle i . La puissance de signal vaut S_e en entrée, et S_s en sortie. La bande passante de travail est B . Les quadripôles sont adaptés en puissance^{4.16}. T_0 est la température de référence d’une source de bruit étalon, T_s la température associée au bruit total mesuré en sortie.

4.15. Ce qui signifie qu’il y a adaptation en puissance de l’entrée du quadripôle.

4.16. cf. paragraphe 4.1.4.

On cherche à déterminer le gain en puissance total G et la température de bruit ramenée en entrée T_r du quadripôle équivalent aux trois montés en série.

1. Il est facile de démontrer que $G = \prod_{k=1}^3 G_k$ 4.17.
2. On a :

$$\begin{cases} \mathcal{N}_s &= \mathcal{N}_2 G_3 + \mathcal{N}_{Q3} \\ \mathcal{N}_2 &= \mathcal{N}_1 G_2 + \mathcal{N}_{Q2} \\ \mathcal{N}_1 &= \mathcal{N}_e G_1 + \mathcal{N}_{Q1} \end{cases}$$

On en déduit

$$\mathcal{N}_s = \mathcal{N}_e G_1 G_2 G_3 + G_2 G_3 \mathcal{N}_{Q1} + G_3 \mathcal{N}_{Q2} + \mathcal{N}_{Q3}$$

Avec $\mathcal{N}_e = kBT_0$, $\mathcal{N}_{Q1} = kBT_1$, $\mathcal{N}_{Q2} = kBT_2$ et $\mathcal{N}_{Q3} = kBT_3$ il vient

$$\mathcal{N}_s = kBT_0 G_1 G_2 G_3 + kBT_1 G_2 G_3 + kBT_2 G_3 + kBT_3$$

Soit

$$\mathcal{N}_s = G \mathcal{N}_e \left(1 + \frac{T_1}{G_1 T_0} + \frac{T_2}{G_1 G_2 T_0} + \frac{T_3}{G_1 G_2 G_3 T_0} \right)$$

Calculons le facteur de bruit du quadripôle équivalent :

$$\mathcal{F} = \left(\frac{\mathcal{S}_e}{\mathcal{S}_s} \right) \left(\frac{\mathcal{N}_s}{\mathcal{N}_e} \right) = \left(\frac{1}{G} \right) \left[G \left(1 + \frac{T_1}{G_1 T_0} + \frac{T_2}{G_1 G_2 T_0} + \frac{T_3}{G_1 G_2 G_3 T_0} \right) \right]$$

En identifiant avec la relation 4.1, on obtient

$$T_r = T_{r1} + \frac{T_{r2}}{G_1} + \frac{T_{r3}}{G_1 G_2}$$

En résumé, sous réserve d'adaptation en puissance :

- Le gain en puissance d'une chaîne de quadripôles est égal au produit des gains en puissance ;
- La température de bruit ramenée en entrée est égale à

$$T_r = \sum_i \frac{T_{ri}}{\prod_{j=0}^{i-1} G_j}$$

avec la convention $G_0 = 1$.

On remarque qu'une température joue un plus grand rôle que les autres : celle du premier quadripôle. Une conséquence remarquable est que **si on veut diminuer le bruit à la sortie d'une chaîne de quadripôles, il est nécessaire de placer en tête de la chaîne un quadripôle de faible température de bruit, et de surcroît amplificateur, pour diminuer encore plus l'influence des bruits dûs aux autres quadripôles.**

Ce qu'il faut retenir

- La densité spectrale de puissance (DSP) est le carré du module du spectre d'un signal ;
- La DSP d'un bruit blanc est constante en fonction de la fréquence ;
- Une résistance R portée à la température T produit un bruit blanc dit « thermique », de DSP $2kTR$, avec $k = 1,38 \cdot 10^{-23}$ J/K. Quand ce bruit est filtré par un filtre de bande passante B , la puissance de bruit résultante est de l'ordre de $kTRB$;
- La définition de la température équivalente de bruit d'un système quelconque ; plus un système est bruyant, plus sa température de bruit est grande ;
- Le théorème de Friiss et sa conséquence principale : dans une chaîne de réception, le premier étage doit être un amplificateur peu bruyant.

4.17. Le symbole Π indique le produit : ici, $G = G_1 G_2 G_3$.

4.5 Parasites radioélectriques

Nous avons étudié dans le paragraphe précédent les sources de « pollution » d'un signal dues aux composants utilisés pour son analyse ou sa production. On peut un peu arbitrairement distinguer maintenant la catégorie des « parasites » radioélectriques, que l'on définira grossièrement comme le bruit causé par des sources plus ou moins distantes du système utilisé.

4.5.1 Les sources de parasites

Ce paragraphe n'a pas pour ambition de présenter un catalogue exhaustif des différentes sources de parasites, mais simplement d'en faire un inventaire indicatif. Principalement, on distingue :

1. Des parasites d'origine purement électriques :

- Ouverture et fermeture d'interrupteurs, qui s'accompagnent d'arcs électriques, sources de craquements ;
- Présence d'harmoniques de la fréquence du secteur. Par exemple, l'harmonique $350\text{Hz}=7\times 50\text{Hz}$ est souvent relativement puissant. Si un fil électrique mal isolé est placé près d'un câble téléphonique, et que cet harmonique est présent, un ronflement se fera entendre dans le combiné, car cette fréquence se trouve dans des domaines audibles ;
- Variations de fréquence du secteur ;
- Variations d'amplitude du secteur, ces variations pouvant aller jusqu'à de véritables « microcoupures ».

2. Des parasites d'origine météorologique :

- Les éclairs, grandes décharges électriques entre le sol et la base des nuages, ou bien entre la base et le sommet des nuages, provoquent des crépitements (par exemple, sur une radio réglée en grandes ondes). Les perturbations du champ électrique peuvent se faire sentir bien au-delà de la zone où l'éclair est visible ;
- Une gouttelette d'eau est globalement neutre; mais si elle est fragmentée, elle se sépare en deux parties, l'une chargée positivement, l'autre négativement. A grande échelle, le champ électrique ainsi créé peut engendrer des parasites. Le même phénomène est observable avec les cristaux de glace. Une antenne dans une tempête de neige, frappée en permanence par des cristaux de glace de polarités opposées, sera soumise à un champ électrique source de parasites. Dans une moindre mesure, cela est vrai également pour le sable.

3. Des parasites d'origine électrostatique : produits par le frottement sur une moquette, ou bien par certains vêtements. Les micro-décharges provoquées sont sources de craquements dans les postes de radio ou les vieilles télévisions ;

4. Des parasites d'origine chimique :

- Un chalumeau, par exemple, ionise l'air qu'il chauffe. Les ions ainsi créés produisent un champ électrique perturbateur ;
- La terre renferme des métaux, du fait ou non de l'activité humaine; la corrosion de ces métaux par l'infiltration des eaux de pluie, par exemple, produit des courants souterrains capables d'introduire de légères fluctuations des prises de terre électriques.

5. Sans compter les autres...

- La Terre est traversée de courants dits « telluriques », de faible intensité et dont l'origine est en grande partie inconnue. Ces courants produisent le même phénomène que la corrosion souterraine ;
- Il existe également des courants de fuite industriels ou particuliers. Ces courants, qui partent dans la terre, introduisent des déséquilibres dans les masses électriques des appareils ;
- La foudre qui tombe au loin fait circuler de manière brève mais intense un courant dans le sol. Selon la nature du terrain, ce courant peut être transporté plus ou moins loin.

4.5.2 Classification des parasites...

La liste précédente est difficilement utilisable telle quelle. Pour permettre un traitement global des sources de parasites, il est plus pratique de les classer suivant la nature de leur propagation, ou leurs effets.

4.5.2.1 ... par leur propagation

On distingue les parasites qui sont transportés par *conduction* de ceux qui le sont par *rayonnement*.

1. **par conduction** : la source de parasites et l'appareil perturbé sont reliés par des conducteurs électriques (fil, masse métallique quelconque, terre...);
2. **par rayonnement** : la présence d'un champ électrique produit par le perturbateur *induit* des courants dans le perturbé. Ce phénomène peut même être localement amplifié : toute masse métallique se comporte plus ou moins comme une antenne. Pour peu que par malchance cette antenne se trouve « optimisée » pour la réception d'un parasite, ce dernier sera amplifié.

4.5.2.2 ... par leurs effets

On peut distinguer suivant ce critère trois classes d'effets :

1. **destructifs** : le perturbé est détruit par la survenue du parasite, qui peut être une tension ou un courant de crête trop important, ou bien une puissance électrique trop grande ;
2. **non destructifs mais nuisibles** : l'effet désagréable du parasite disparaît avec celui-ci, comme par exemple la « neige » sur un écran de télévision, ou bien les craquements et sifflements dans une radio au moment d'un orage, etc ;
3. **perturbateurs des systèmes logiques** : « macroscopiquement », l'utilisateur peut ne pas remarquer d'effet. Mais la survenue d'un « rayon cosmique »^{4.18} peut perturber ponctuellement le fonctionnement d'un unique transistor, et le placer dans un état erroné. Ce genre d'accident peut arriver plus souvent aux satellites, car ils ne sont pas protégés par l'épaisseur de l'atmosphère.

4.5.3 Les parades

Pour mettre au point les parades, il faut tout d'abord évaluer la *nécessité* ou non d'un dispositif anti-parasites : un parasite transitoire peut en effet être supportable, ou bien le coût d'installation du dispositif être supérieur au préjudice subi. De plus, la « sensibilité » aux parasites, en particulier en matière de réception radio ou télévisuelle, est tout à fait subjective. Un individu donné ressentira le besoin d'un déparasitage, mais son voisin n'en verra pas l'utilité.

Les parasites de conduction étant transportés le plus souvent par les fils entrant dans l'appareil perturbé, une solution couramment retenue est l'utilisation d'un filtre fréquentiel. Souvent en effet, les parasites de conduction ont un « support fréquentiel » assez marqué. Un filtre coupe-bande ou passe-bas^{4.19} permet de s'en débarrasser. Une autre méthode consiste à réduire la bande passante du récepteur : le parasite, s'il est d'une fréquence voisine de celle-ci, pourra éventuellement se trouver alors en-dehors de la nouvelle bande. Les deux méthodes sont en fait similaires, mais diffèrent dans leur endroit d'application : le filtre est inséré dans le montage, alors que le réglage de la bande passante se fait sur le récepteur lui-même. De plus, l'utilisation de condensateurs ou de bobines pour réaliser le filtre peut ajouter des problèmes : lorsqu'un courant circule dans une bobine, par exemple, un champ électromagnétique est créé. Si on insère un filtre contenant une bobine *dans* un système à déparasiter, on insère en fait à l'intérieur même de ce système une source de parasites.

Les parasites de rayonnement, transportés par le champ électrique, sont interceptés par une cage métallique, appelée *cage de Faraday*. Les conditions de continuité du champ électrique entre l'intérieur et l'extérieur imposent alors à celui-ci d'être nul à l'intérieur de la cage, si l'on excepte bien sûr toute source interne. Ce carter métallique, pour des raisons de sécurité pour l'utilisateur, doit impérativement être relié à la terre. Mais il est parfois nécessaire de ménager des trous dans cette cage : alimentation électrique, arrivée d'un signal, aération, etc. Ces trous sont autant

4.18. Il s'agit en fait d'une particule chargée venant de l'espace (Soleil, autres étoiles, etc.), qui arrive dans l'atmosphère avec une grande vitesse. Elle crée alors une « gerbe » de particules chargées énergétiques en interagissant avec les molécules de l'atmosphère, qui peuvent amener le « changement d'état » d'un transistor.

4.19. cf. paragraphe 4.3.1.3.

de failles par lesquelles les parasites peuvent pénétrer. Cependant, il n'est souvent pas nécessaire d'exiger une cage parfaitement étanche aux parasites : une cage conductrice grillagée suffit parfois. La taille des trous du grillage permet de sélectionner les longueurs d'onde capables de passer ou non, et dans le cas où les parasites sont bien localisés en fréquence, cette solution, moins onéreuse, est préférée. De plus, une cage permet de limiter l'émission de rayonnement dans l'environnement par le système lui-même.

Ce qu'il faut retenir

- Les parasites peuvent être classés en deux catégories : les parasites par conduction, et les parasites par rayonnement ;
 - Les parasites par conduction sont éliminés par filtrage fréquentiel, les parasites par rayonnement avec un blindage par une cage conductrice.
-

Chapitre 5

Systèmes numériques

5.1 Introduction

5.1.1 Généralités

Au début du cours^{5.1}, nous avons distingué les *signaux à valeurs discrètes* des signaux à valeurs continues. Ces signaux ne peuvent prendre qu'un nombre fini de valeurs dans un intervalle. Ils ne sont pas à confondre avec les *signaux à temps discret*. Dans cette partie du cours d'ailleurs, nous considérerons aussi bien des signaux à temps discret qu'à temps continu.

En pratique, l'intervalle dans lequel les signaux peuvent prendre leurs valeurs est souvent de la forme $(2^n - 1)u_0$, u_0 étant une valeur de référence, nommée le *pas de l'échantillonnage* (par exemple $255u_0$, ou $1023u_0$). Ce choix d'une base 2 est lié à des contraintes logiques (cela permet une représentation aisée de deux valeurs vrai/faux), et matérielles, comme on le verra plus loin.

On distingue deux « branches » de l'électronique logique, ou numérique : la logique *combinatoire* et la logique *séquentielle*. La première fait référence aux composants dont l'état de la sortie dépend uniquement des états en entrée ; la deuxième aux composants dont l'état de sortie dépend des états des entrées *et* d'une horloge externe ou interne qui « cadence » le fonctionnement.

5.1.2 Représentation logique

Si on considère un intervalle $2^n - 1$, chaque valeur p entière de cet intervalle peut s'écrire sous la forme de son développement en base 2 :

$$p = \sum_{k=0}^{n-1} p_k 2^k \text{ avec } p_{k=0 \dots n-1} = 0 \text{ ou } 1$$

Les coefficients p_k sont appelés *digits* en français, ou *bits* en anglais. p_0 est appelé *bit de plus faible poids*, et p_n *bit de plus fort poids*. 8 bits forment un *octet*, ou *byte*, 2 octets forment un *mot*, ou *word* et 4 octets forment un *mot long*, ou *longword*.

On utilise souvent également la représentation hexadécimale. Dans cette représentation, les nombres de 0 à 15 sont notés 0, 1, 2, 3, 4, 5, 6, 7, 8, 9, A, B, C, D, E, F. On utilise cette représentation quand il s'agit de noter simplement de longues séries de bits, en regroupant ces bits par paquets de 4. Par exemple, le nombre binaire 11000011, qui

5.1. cf. paragraphe 1.1.2.

correspond en décimal à 195, sera noté C3 en hexadécimal : 1100 correspond en effet à 12, donc C en hexa, et 0011 à 3, donc 3 aussi en hexa. FF codera donc de même 255.

Remarque : souvent on associe 1 à la valeur « vrai » et 0 à « faux ». Mais en *logique négative*, c'est le contraire : l'état par défaut d'un niveau logique est le 1 (potentiel « haut »), et l'information est considérée comme présente quand le niveau passe à 0.

5.1.3 Familles de portes logiques

Le composant de base est le transistor, fonctionnant en mode « interrupteur commandé »^{5.2}. On distingue deux grandes familles :

- **TTL** : à base de transistors bipolaires. Les niveaux logiques sont souvent 0/15V.
- **CMOS** : à base de transistors à effet de champ. Les niveaux logiques sont souvent 0/5V, mais peuvent aller jusqu'à 20V. Les portes CMOS ont une consommation électrique plus faible en général que les portes TTL.

Ce qu'il faut retenir

- En électronique numérique, on utilise la représentation binaire ;
 - Le vocabulaire : bit, octet.
-

5.2 Logique combinatoire

Ce chapitre va aborder l'étude des composants logiques dont la sortie dépend uniquement de l'état des entrées.

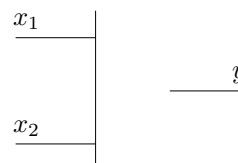
5.2.1 Les opérateurs de base

Ces opérateurs sont souvent appelés « portes » .

5.2.1.1 Les opérateurs simples

1. Porte ET (AND)

– *Schéma :*



^{5.2} Le niveau logique est représenté par la différence de potentiel aux bornes du transistor, qui fonctionne en interrupteur (modes passant ou bloqué: cf. paragraphe 7.4.2).

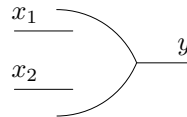
– Table de vérité :

$x_2 \backslash x_1$	0	1
0	0	0
1	0	1

– Notation : $y = x_1.x_2$

2. Porte OU (OR)

– Schéma :



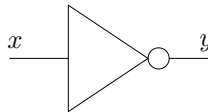
– Table de vérité :

$x_2 \backslash x_1$	0	1
0	0	1
1	1	1

– Notation : $y = x_1 \vee x_2$ ou $y = x_1 + x_2$. Attention, cette dernière notation n'est pas une réelle addition ; ainsi, par exemple, 1 ou 1, qui est égal à 1, s'écrit $1+1=1$.

3. Porte NON (NOT^{5.3})

– Schéma :



– Table de vérité :

x	0	1
y	1	0

– Notation : $y = \bar{x}$

5.2.1.2 Propriétés

On les donne sans démonstration : il suffit d'examiner les tables de vérité.

1. Commutativité :

$$\begin{cases} a + b &= b + a \\ a.b &= b.a \end{cases}$$

2. Associativité :

$$\begin{cases} (a + b) + c &= a + (b + c) &= a + b + c \\ (a.b).c &= a.(b.c) &= a.b.c \end{cases}$$

3. Distributivité :

$$\begin{cases} a.(b + c) &= a.b + a.c \\ a + (b.c) &= (a + b).(a + c) \end{cases}$$

4. Idempotance :

$$\begin{cases} a + a &= a \\ a.a &= a \end{cases}$$

5.3. Cette opération est aussi appelée *complémentation*.

5. Éléments neutres :

$$\begin{cases} a.1 &= a \\ a+0 &= a \\ a+1 &= 1 \\ a.0 &= 0 \end{cases}$$

6. Complémentarité :

$$\begin{cases} a + \bar{a} &= 1 \\ a.\bar{a} &= 0 \end{cases}$$

7. Double négation :

$$\bar{\bar{a}} = a$$

8. Formules de De Morgan :

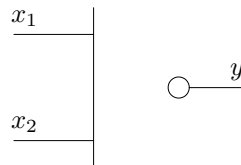
$$\begin{cases} \bar{a}.\bar{b} &= \overline{a+b} \\ \bar{a} + \bar{b} &= \overline{a.b} \end{cases}$$

5.2.1.3 Les opérateurs « intermédiaires »

Ces opérateurs ne sont pas ceux qui viennent en premier à l'esprit, mais figurent pourtant parmi les plus utilisés.

1. Porte ET complémenté (NAND)

– Schéma :



– Table de vérité :

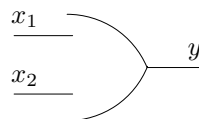
$x_2 \backslash x_1$	0	1
0	1	1
1	1	0

– Notation : $y = \overline{x_1.x_2} = x_1 \uparrow x_2$

– Remarque : Il est possible de remplacer tout type de porte par des portes NAND^{5.4}.

2. Porte OU exclusif (XOR)

– Schéma :



– Table de vérité :

$x_2 \backslash x_1$	0	1
0	0	1
1	1	0

– Notation : $y = \overline{x_1}.x_2 + x_1.\overline{x_2} = x_1 \oplus x_2$

– Associativité : On montre facilement que $a \oplus (b \oplus c) = (a \oplus b) \oplus c = a \oplus b \oplus c$.

^{5.4} Il est en effet possible de montrer que la complémentation est réalisable à l'aide d'une unique porte NAND ($\bar{x} = x \uparrow x$), que le ET logique nécessite 3 NAND ($x_1.x_2 = (x_1 \uparrow x_2) \uparrow (x_1 \uparrow x_2)$), ainsi que le OU logique ($x_1 + x_2 = (x_1 \uparrow x_1) \uparrow (x_2 \uparrow x_2)$).

5.2.2 Table de Karnaugh

Nous avons déjà utilisé des « tables de vérité ». Il existe une autre représentation possible.

5.2.2.1 Principe

On utilise la distributivité et la propriété $a + \bar{a} = 1$.

Par exemple, supposons que $y = (a.\bar{b}.c\bar{d}) + (\bar{a}.\bar{b}.c.\bar{d})$. Cette expression se distribue en $y = \bar{b}.c.\bar{d}.(a + \bar{a}) = \bar{b}.c.\bar{d}$.

L'utilisation de cette propriété est plus aisée avec un autre outil...

5.2.2.2 Code binaire réfléchi

Vous connaissez le code binaire « naturel » : 0, 1, 10, 11, 100, 101... qui est simplement la suite des entiers écrits en base 2. Le *code binaire réfléchi* se construit en commençant par 0, puis en changeant la valeur d'un *unique* bit à chaque fois, de la droite vers la gauche. Sur 4 bits, on obtient ainsi la suite : 0000, 0001, 0011, 0010, 0110, 0111, 0101, 0100, 1100, 1101, 1111, 1110, 1010, 1011, 1001, 1000. Pour plus de clarté, on inscrit cette suite dans une table :

a_3	$a_2 \backslash \begin{matrix} a_1 \\ a_0 \end{matrix}$	0	0	1	1
		0	1	1	0
0	0		x	x	
0	1				
1	1	o			o
1	0				

Cette table possède de nombreux axes de symétrie ou d'anti-symétrie. Il est possible de les utiliser afin de créer des regroupements, par exemple entre les cases « x », où on constate que la valeur de a_1 est indifférente, ce qui permet de regrouper ces cases en $\bar{a}_3\bar{a}_2a_0$, ou bien entre les cases « o », regroupables en $a_3a_2\bar{a}_0$.

Remarque : Dans certains cas, la sortie n'est pas définie pour toutes les combinaisons possibles des entrées. On place alors un « X » dans les cases correspondantes de la table. Ces valeurs pourront éventuellement être fixées à 1 ou 0 pour faciliter des regroupements.

5.2.2.3 Exemple

Pour un système à 4 bits $a_3a_2a_1a_0$, on crée le bit s défini par :

$$s = \bar{a}_3.\bar{a}_2.\bar{a}_1.a_0 + \bar{a}_3.\bar{a}_2.\bar{a}_1.a_0 + \bar{a}_3.a_2\bar{a}_1.a_0 + \bar{a}_3.a_2.a_1.a_0 + a_3.\bar{a}_2.a_1.a_0 + a_3a_2\bar{a}_1.a_0 + a_3.a_2.a_1.a_0$$

Ce bit est défini comme un OU logique sur 7 termes (appelés *mintermes*). La table de Karnaugh définissant s est la suivante :

a_3	$a_2 \backslash \begin{matrix} a_1 \\ a_0 \end{matrix}$	0	0	1	1
		0	1	1	0
0	0	1	1	0	0
0	1	0	1	1	0
1	1	0	1	1	0
1	0	0	0	1	0

On peut alors rassembler :

- les 1 qui se trouvent en haut à gauche du tableau peuvent être rassemblés dans le minterme $\overline{a_3}.\overline{a_2}.\overline{a_1}$;
- les 1 du milieu du tableau peuvent être rassemblés en a_2a_0 ;
- le 1 du bas du tableau s'exprime directement en $a_3\overline{a_2}a_1a_0$, ou bien peut être rassemblé avec le 1 au-dessus de lui en $a_3a_1a_0$.

Et donc $s = \overline{a_3}.\overline{a_2}.\overline{a_1} + a_2a_0 + a_3\overline{a_2}a_1a_0$. Avec un peu d'habitude, les simplifications de ce genre se font beaucoup plus rapidement qu'en partant de l'expression initiale de s .

5.2.3 Quelques fonctions plus évoluées de la logique combinatoire

5.2.3.1 Codage, décodage, transcodage

1. Codage

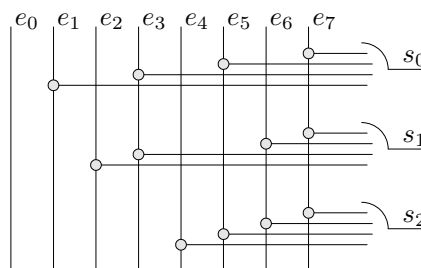
- *Définition* : un codeur est une « boîte noire », avec 2^n entrées, dont une seule est active à la fois, et n sorties. L'état des sorties indique quelle entrée est active.
- *Exemple* : codeur en binaire naturel à 3 bits. On a 8 entrées $e_{k=0\dots7}$ et 3 sorties $s_{k=0\dots2}$, vérifiant la table suivante :

Entrée active	s_2	s_1	s_0
e_0	0	0	0
e_1	0	0	1
e_2	0	1	0
e_3	0	1	1
e_4	1	0	0
e_5	1	0	1
e_6	1	1	0
e_7	1	1	1

Cette table se traduit par les relations suivantes :

$$\begin{cases} s_0 = e_7 + e_5 + e_3 + e_1 \\ s_1 = e_7 + e_6 + e_3 + e_2 \\ s_2 = e_7 + e_6 + e_5 + e_4 \end{cases}$$

– *Schéma* :

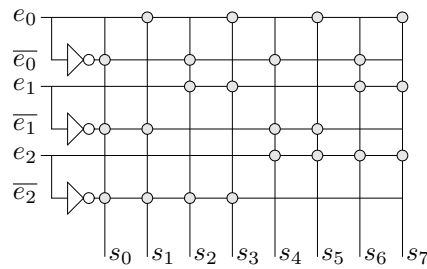


Chaque point indique une connexion physique.

2. Décodage

- *Définition* : c'est l'opération inverse du codage... Un décodeur est une « boîte noire », avec n entrées et 2^n sorties. Chaque sortie est activée par une combinaison particulière des entrées.
- *Exemple* : On prend l'exemple inverse du codeur précédent, avec 3 entrées $e_{k=0\dots2}$ et 8 sorties $s_{k=0\dots7}$. La table de fonctionnement est la même que pour l'exemple précédent, sous réserve de permuter entrées et sorties.

– Schéma :



Ici, chaque point indique une connexion « ET ».

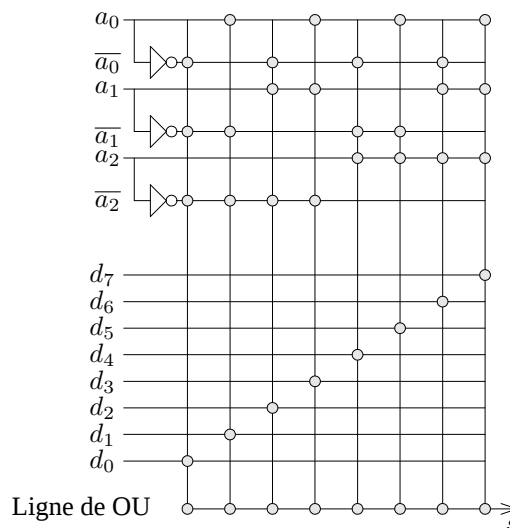
3. Transcodage

- *Définition* : La « boîte noire » permet de passer d'un code 1 à un code 2. Il y a n entrées et n sorties.
- *Exemple* : Transcodage code binaire naturel $\rightarrow$ code binaire réfléchi, ou le contraire, etc.
- *Remarque* : On peut également réaliser un transcodeur à n entrées et m sorties, comme par exemple un transcodeur binaire naturel $\rightarrow$ binaire codé décimal^{5.5}.

5.2.3.2 Multiplexage, démultiplexage

1. Multiplexage

- *Définition* : Un multiplexeur est une « boîte noire », avec 2^n entrées de données d_i , n entrées d'adresse a_k et une sortie. Celle-ci reproduit l'entrée de données dont le numéro est codé par les entrées d'adresse.
- *Exemple* : avec $n = 3$. On considère un multiplexeur à 8 entrées de données $d_{i=0\dots7}$ et à 3 entrées d'adresse $a_{k=0\dots2}$. Si on a $d_0d_1d_2d_3d_4d_5d_6d_7 = 01001110$ en entrée, et si on code l'adresse à $a_2a_1a_0 = 001$, on obtient en sortie l'entrée d_1 , soit $s = d_1 = 1$. Si on code $a_2a_1a_0 = 111$, on a en sortie $s = d_7 = 0$.
- Schéma :



Ici, chaque point indique une connexion « ET », sauf les points de la dernière ligne.

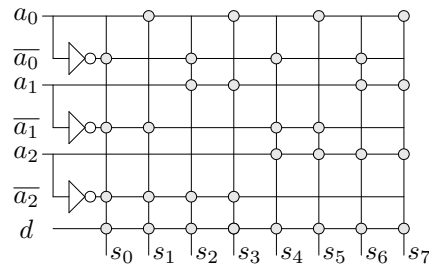
- *Remarque* : On peut aussi réaliser un multiplexeur avec m sorties, si on lui injecte en entrée n mots de m bits.

2. Démultiplexage

- *Définition* : C'est l'inverse de la précédente... Un démultiplexeur est une « boîte noire » avec 1 entrée de donnée, n entrées d'adresse et 2^n sorties. L'entrée est « aiguillée » vers la sortie dont le numéro est codé dans les adresses.

5.5. cf. l'exemple dans le paragraphe 5.2.6.2.

– Schéma :



Chaque point indique une connexion « ET ».

- *Remarque* : D'après le schéma précédent, les sorties non actives sont dans l'état 0. Mais il est possible de faire en sorte que leur état par défaut soit 1. Ce problème de l'indétermination des sorties inactives d'un démultiplexeur doit être pris en compte en aval lors d'un câblage.

3. Exemple d'application : conversion parallèle↔série.

- Considérons un mot de n bits. Il peut être transmis soit sur un fil unique, bit après bit (transmission *série*), soit sur plusieurs fils à la fois, un fil par bit (transmission *parallèle*).
- *Conversion parallèle→série* : elle est effectuée à l'aide d'un multiplexeur : on envoie en entrée les n bits du mot à transmettre, et sur les entrées d'adresse successivement **00**, **01**, **10**, **11**. En sortie on obtient la série des n bits du mot.
- *Conversion série→parallèle* : elle est effectuée à l'aide d'un démultiplexeur. On envoie en entrée successivement les n bits du mot, et en même temps, on fait varier les bits d'adresse en les incrémentant^{5.6}. En sortie, les fils doivent être reliés à une mémoire, qui stocke l'un après l'autre les bits du mot.

5.2.4 Fonctions arithmétiques

On définit une *opération logique* comme une opération réalisée sur un mot binaire, et une *opération arithmétique* comme une opération réalisée sur un mot binaire codé. Par exemple, si on considère le mot binaire A , on peut effectuer des opérations logiques comme des ET ou des OU. Mais si ce mot A désigne un certain nombre suivant un codage défini (comme le binaire naturel), alors les opérations réalisées seront dites arithmétiques (par exemple, addition de deux nombres, négation, etc.).

5.2.4.1 Fonctions logiques

Soient A et B deux mots logiques de n bits, dont les bits sont notés respectivement a_i et b_i .

- Le ET entre les deux mots est noté $A.B = C$, et est tel que $c_i = a_i.b_i$.
- Le OU entre les deux mots est noté $A \vee B = A + B = C$, et est tel que $c_i = a_i \vee b_i = a_i + b_i$.
- Le NON du mot A est noté $\overline{A} = C$, et est tel que $c_i = \overline{a_i}$.

5.2.4.2 Fonctions arithmétiques

On se limite à des entiers relatifs, en codage binaire naturel. On note A et B deux nombres binaires de n bits : $A = a_{n-1}...a_0$ et $B = b_{n-1}...b_0$. En décimal, le nombre A s'écrit $A = \sum_{k=0}^{n-1} a_k 2^k$.

1. Addition de deux entiers naturels

On note C la somme de A et de B : $C = c_n...c_0$, C étant *a priori* un nombre de $n + 1$ bits, quand on tient compte de la retenue.

5.6. C'est-à-dire en les augmentant d'une unité. La *décrément* est l'opération inverse : diminution d'une unité.

Réalisons par exemple $1011+1101$, et posons l'opération :

$$\begin{array}{r} 1011 \\ + 1101 \\ \hline 10100 \end{array}$$

retenue

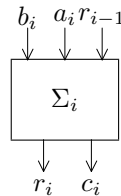
Dans le cas général, exprimons le bit c_i en fonction de a_i , b_i , et de la retenue r_{i-1} :

$r_{i-1} \backslash b_i$	0	0	1	1
a_i	0	1	1	0
0	0	1	0	1
1	1	0	1	0

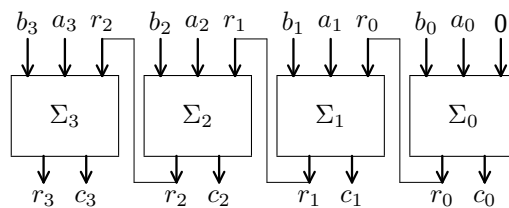
On en déduit que $c_i = a_i \oplus b_i \oplus r_{i-1}$. Etablissons de même l'expression de la retenue r_i :

$r_{i-1} \backslash b_i$	0	0	1	1
a_i	0	1	1	0
0	0	0	1	0
1	0	1	1	1

On montre alors que $r_i = a_i \cdot b_i + r_{i-1} \cdot (a_i \oplus b_i)$. On peut alors rassembler ces entrées/sorties sous la forme d'un bloc unique :



Il est ensuite possible de cascader n de ces blocs pour réaliser un additionneur à n bits ; par exemple avec $n = 3$ on peut faire :



Cet additionneur est simple à construire, mais peu performant car la retenue doit être « propagée » d'un bloc à l'autre.

2. Soustraction de deux entiers naturels

On considère deux nombres de n bits, auxquels on adjoint un bit de signe.

On introduit le *complément à 2* du nombre A , noté $\mathcal{C}(A)$: $\mathcal{C}(A) = \overline{A} + 1$. Par exemple, si $A = 1001$, son complément à 2 vaut $\mathcal{C}(A) = 0110 + 1 = 0111$.

Ce complément à 2 présente une propriété remarquable : calculons $A + \mathcal{C}(A)$:

$$A + \mathcal{C}(A) = A + \overline{A} + 1 = \underbrace{1}_{\text{retenue}} \underbrace{00\dots 0}_{n \text{ bits}} = 2^n = 0 \text{ modulo } 2^n$$

Il s'ensuit que :

$$B + \mathcal{C}(A) = B + \overline{A} + 1 = (B - A) + (A + \overline{A} + 1) = (B - A) \text{ modulo } 2^n$$

Pour réaliser un soustracteur, on peut donc utiliser un additionneur, et lui adjoindre un « complémenteur à 2 ».

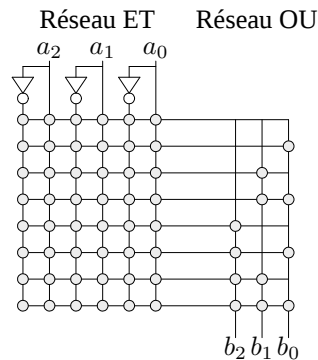
5.2.5 Mémoire morte

Cette mémoire est souvent désignée sous le nom de ROM (Read Only Memory) . C'est une « boîte noire » qui possède n entrées d'adresse et p sorties de données. A chaque combinaison des n bits d'adresse est associée une combinaison des p bits de sortie.

5.2.6 Le PAL et le PLA

5.2.6.1 Le PAL

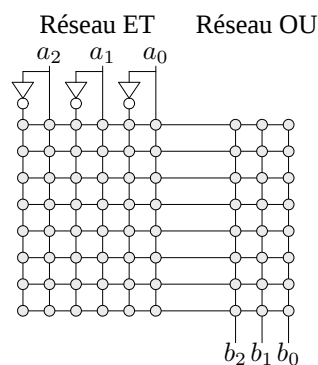
Le *réseau logique programmable*^{5.7} est une matrice à réseau ET et réseau OU. Le réseau OU est fixe et le réseau ET est programmable. On choisit alors les connexions en fonction des besoins :



Dans le schéma d'un PAL à 3 bits ci-dessus, les cercles indiquent des fusibles, *ie* des points où les connexions peuvent être établies ou non. La programmation est réalisée en appliquant des différences de potentiel capables de « claquer » ces fusibles, et la déprogrammation en exposant le réseau à des ultraviolets qui les « régénèrent ».

5.2.6.2 Le PLA

Dans un PLA^{5.8}, toutes les connexions des réseaux ET et OU sont programmables : c'est le plus polyvalent des circuits intégrés. Voici par exemple le schéma de principe d'un PLA à 3 bits :



Les cercles indiquent également des fusibles programmables.

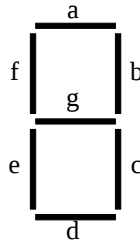
Exemple d'application : Transcodeur BCD-7 segments.

- *Le code BCD* : BCD pour Binaire Codé Décimal. Ce codage est fondé sur la constatation suivante : un chiffre quelconque entre 0 et 9 peut être codé en binaire naturel sur 4 bits, de 0000 à 1001. Un nombre à deux chiffres peut être codé sur 8 bits : par exemple, 23 sera codé par 00100011. Bien sûr, ce codage est beaucoup moins efficace que le binaire naturel (dans ce dernier, 23 s'écrit avec 5 bits seulement : 10110) ;

5.7. Programmable Array Logic en anglais.

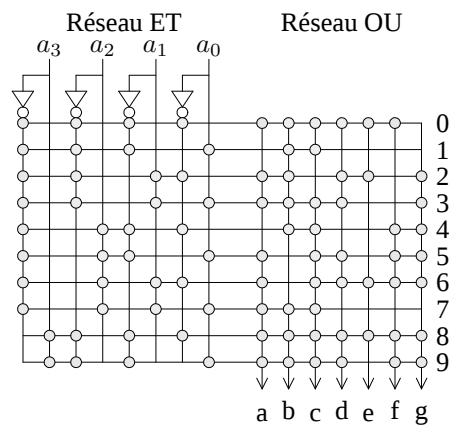
5.8. Programmable Logic Array.

- *L’afficheur 7 segments*: Il est composé de 7... segments, et se trouve dans beaucoup de systèmes numériques d’affichage. Ces segments sont notés a, b, c, d, e, f et g suivant le schéma suivant :



Par exemple, le 7 est codé par $abc\bar{d}\bar{e}\bar{f}\bar{g}$.

- On cherche à réaliser un composant qui prend en entrée 4 fils de données correspondant aux 4 bits de codage d’un chiffre en BCD, et qui délivre 7 sorties $abcdefg$, correspondant aux 7 fils d’un afficheur 7 segments. Ce composant est facile à réaliser avec un PLA comme suit, en ne gardant que certaines connexions :



Ce qu’il faut retenir

- Les tables de vérité des opérateurs simples : ET, OU, NON, NAND ;
- Tous les opérateurs simples peuvent être remplacés par des portes NAND ;
- Les Tables de Karnaugh et le code binaire réfléchi permettent de simplifier les expressions logiques ;
- Le multiplexage est une sorte d’« aiguillage » des informations ;
- Le PAL et le PLA sont des circuits logiques programmables, utilisables pour remplacer tout système de logique combinatoire.

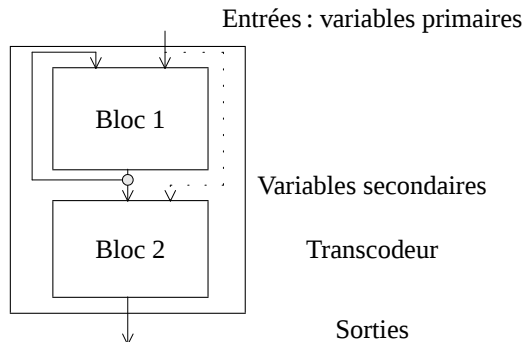
5.3 Logique séquentielle

5.3.1 Généralités

Pour un système séquentiel, l’histoire antérieure intervient dans la définition de l’état actuel.

5.3.1.1 Le caractère séquentiel

Il y a nécessité de créer des variables internes secondaires, produites par le système lui-même, permettant de caractériser l'effet des états antérieurs : ces variables sont introduites par une rétroaction interne. Un système séquentiel peut être représenté sous la forme du schéma suivant :



Le fonctionnement de principe est le suivant :

- Un circuit combinatoire (Bloc 2) élabore les sorties à partir :
 - des variables secondaires ;
 - des variables primaires (*ie* les entrées) éventuellement.
- Un circuit combinatoire (Bloc 1) élabore les variables secondaires à partir :
 - des variables primaires ;
 - des variables secondaires.

Les nouvelles valeurs des variables secondaires attaquant le Bloc 1 sont mises à jour après un certain retard, qui peut être dû :

- à des délais causés par les transports des informations dans les circuits utilisés ; on parle alors de logique asynchrone ;
- à des délais imposés par des circuits spécialisés ; on parle alors de logique synchrone.

5.3.1.2 Systèmes synchrones et asynchrones

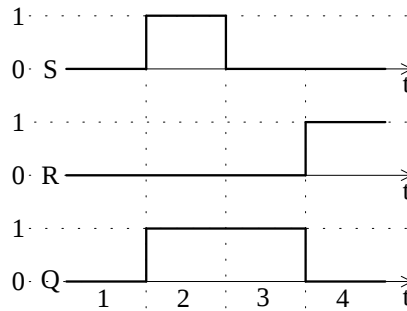
1. **Définition :** Dans un système synchrone, les nouvelles valeurs des variables secondaires *et des sorties* ne peuvent apparaître qu'aux instants définis par un signal périodique (horloge).
2. **Propriétés :** Un système asynchrone est plus rapide (il n'a pas besoin d'attendre le signal de l'horloge pour se mettre à jour), mais sa souplesse même d'utilisation le rend particulièrement vulnérable : on n'est jamais certain si le signal de sortie est stabilisé ou s'il est encore susceptible d'évoluer car le calcul n'est pas terminé. Les « gros » systèmes logiques ne sont donc pas réalisés en logique asynchrone.

5.3.1.3 Exemple : bascule RS asynchrone

1. **Cahier des charges :** on cherche à réaliser un système numérique à deux entrées R et S et une sortie Q. L'activation de S (*ie* passage de S à 1) porte la sortie à 1 ou la laisse dans cet état si elle y est déjà. L'activation de R (*ie* passage de R à 1) porte la sortie à 0 ou la laisse dans cet état si elle y est déjà^{5.9}. On formule de plus la contrainte supplémentaire que les entrées R et S ne doivent pas être activées simultanément.

5.9. L'entrée S est l'entrée Set, et l'entrée R est l'entrée Reset.

2. **Analyse des états possibles** : on distingue 4 états possibles du système :



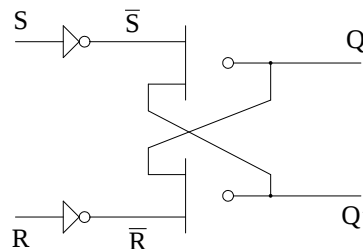
Une variable secondaire est nécessaire pour distinguer les états 1 et 3 malgré les valeurs identiques des entrées. On choisit Q comme variable secondaire ; le Bloc 2 est alors réduit à un simple « fil ».

3. **Tableau des états** : on note q la valeur « future » de la sortie Q. On cherche à pouvoir construire la Table de Karnaugh de q en fonction de R, S et Q :

état	S	R	Q	q
1	0	0	0	0
2	1	0	0	1
3	0	0	1	1
4	0	1	1	0

4. **Table de Karnaugh de q et équation** : on déduit facilement de la table précédente la table de Karnaugh de q, puis $q = S + Q\bar{R} = \bar{S} \uparrow (Q \uparrow \bar{R})$.

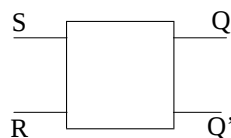
5. **Table de Karnaugh de Q et équation** : en l'occurrence, la sortie est égale à la variable secondaire : $Q = q$. On en tire donc le schéma de câblage suivant, où Q' désigne une sortie « bonus » facile à obtenir :



5.3.2 Fonctions importantes de la logique séquentielle

5.3.2.1 Bascules simples

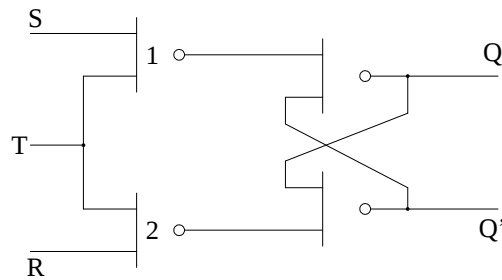
1. **Bascule RS asynchrone** : cf. paragraphe précédent :



On a « pour le même prix » également la sortie $Q' = \bar{Q}$.

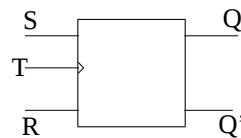
2. Bascules simples synchrones :

- *Bascule RST* : cette bascule présente une entrée T de plus que la bascule RS, qui est une entrée d'horloge, suivant le schéma :



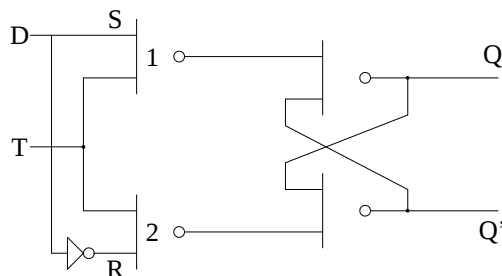
- si $T=0$, les portes 1 et 2 sont bloquées ; la bascule est alors inhibée, et $Q'=\bar{Q}$;
- si $T=1$, le fonctionnement est le même que celui de la bascule RS. Bien sûr, il faut que les entrées restent stables tant que T reste égal à 1 ; mais le fonctionnement est synchrone car il ne peut y avoir « basculement » que si T passe à 1. Notez la présence d'un *alea* : si R et S passent simultanément à 1, l'état de la sortie n'est pas défini^{5.10}.

On la représente par le schéma suivant :



Le triangle correspondant à l'entrée T est le symbole habituel de l'entrée d'horloge.

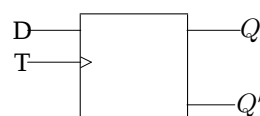
- *Bascule D* : il s'agit en fait d'une bascule RST où l'entrée R est le complément de S : $R=\bar{S}$. Le schéma de câblage est le suivant :



La contrainte est la même que celle de la bascule précédente : en effet, si l'entrée variait alors que $T=1$, le temps de propagation dans la porte NON interviendrait, et l'état des sorties ne serait pas défini. Le fonctionnement de cette bascule est en fait très simple, puisqu'il ne s'agit que d'une « propagation » de l'information. Quand $T=1$,

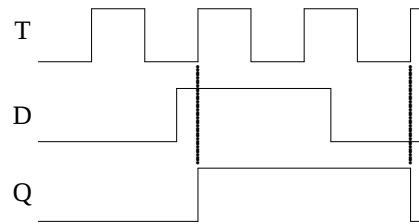
- si $D=0$, la sortie Q passe à 0 ;
- si $D=1$, la sortie Q passe à 1.

On la représente par le schéma suivant :



5.10. Ou plutôt, il est défini par le câblage électrique « microscopique » des composants, et dépend donc du choix de ceux-ci, ce qui est inacceptable en logique.

Voici un chronogramme résumant ses propriétés :

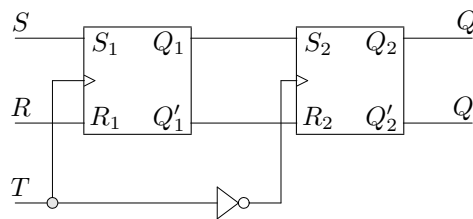


5.3.2.2 Bascules à fonctionnement en deux temps

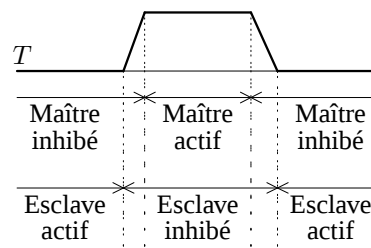
Ces bascules sont aussi appelées « maître-esclave ». Elles sont conçues afin d'éviter de rencontrer les contraintes sur l'absence de variation des entrées tant que l'horloge est dans un état actif. Dès que les entrées sont prises en compte, la bascule devient insensible jusqu'au prochain signal d'horloge, avant que les sorties aient pris leur nouvelles valeurs. La réalisation se fait en cascade de deux bascules fonctionnant en deux temps disjoints.

1. Bascules déclenchées par une impulsion (*pulse triggered*) :

- On trouve par exemple des bascules RST maître-esclave selon le schéma suivant :

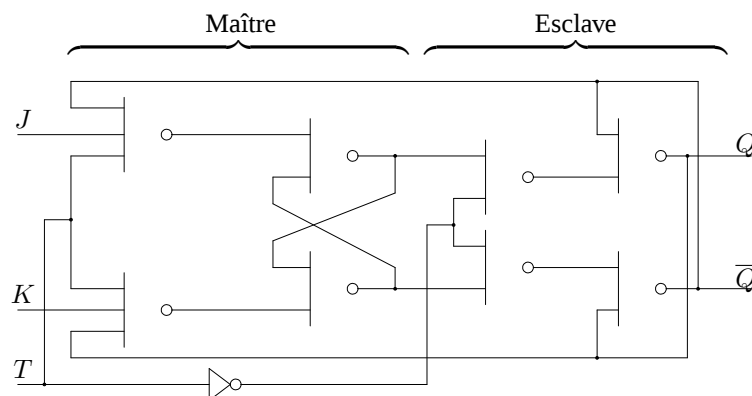


Le chronogramme de fonctionnement est :



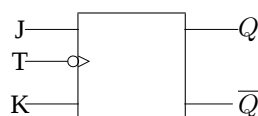
Dans ce cas, l'aléa lié à $RS = 11$ existe toujours.

- On trouve aussi la bascule JK^{5.11}, dont voici le schéma :



5.11. Pour Joker- King.

On la représente ainsi :



Le symbole $\circ$ devant l'entrée d'horloge indique que celle-ci est active sur le front descendant. On a toujours $Q' = \overline{Q}$. Cette fois-ci, l'indétermination liée à $JK = 11$ est levée. En notant Q_{n-1} et Q_n deux états successifs de la sortie, cette bascule présente la table de vérité :

J	K	Q_n
0	0	Q_{n-1}
0	1	0
1	0	1
1	1	$\overline{Q_{n-1}}$

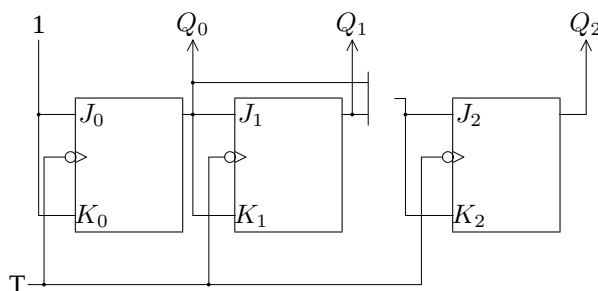
2. **Autres bascules fonctionnant en deux temps** : on trouve également :

- des bascules déclenchées par un front montant ou descendant (dites *edge triggered*) ;
- des bascules à verrouillage de données (dites *data lock out*).

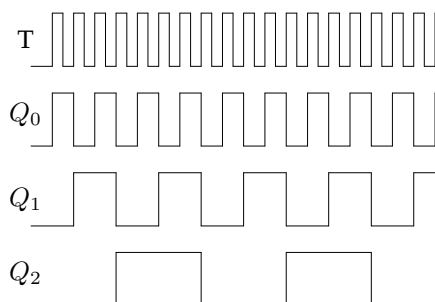
5.3.2.3 Registres (ensembles de bascules)

1. Compteurs

- (a) *Comptage, décomptage*. On cascade des bascules, en récupérant à chaque fois les sorties (cf. schémas plus loin). Si on note A_n le mot constitué par les sorties des bascules après la n -ième impulsion d'horloge, alors : $A_n = A_{n-1} \pm 1$. Si $A_n = A_{n-1} + 1$, on parle de comptage ; dans le cas opposé, de décomptage. On parle également de *cycle complet* pour m bascules si A_n peut varier entre 0 et $2^m - 1$ ^{5.12}.
- (b) *Compteurs à cycle complet (bascules JK)*. On distingue les compteurs synchrones des compteurs asynchrones.
- i. Synchrones : les entrées J_i et K_i des bascules sont connectées entre elles de telle sorte que $J_i = K_i = \prod_{j=0}^{i-1} Q_j$, c'est-à-dire en reliant par un *et* logique les sorties des bascules précédentes. Pour un décompteur, on réalise $J_i = K_i = \prod_{j=0}^{i-1} \overline{Q_j}$. Par exemple, un compteur sur 3 bits peut être réalisé comme ceci :

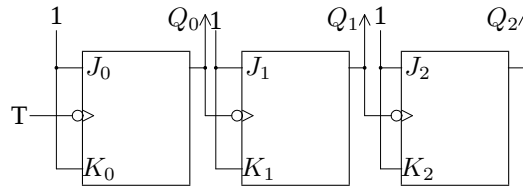


L'horloge est appliquée de façon synchrone sur les entrées d'horloge de toutes les bascules :



5.12. Notez qu'alors le passage au zéro de la sortie du compteur permet de définir une nouvelle période par rapport à celle de l'horloge : la fréquence de cette dernière est divisée par 2^m .

- ii. Asynchrones : cette fois-ci, $J_i = K_i = 1$ pour tout i . L'horloge est appliquée sur l'entrée d'horloge de la bascule qui délivre le bit de plus faible poids Q_0 . Q_i (pour un compteur) ou $\overline{Q_{i-1}}$ (pour un décompteur) est appliqué sur l'entrée d'horloge de Q_i . Ce compteur présente donc des états transitoires erronés :



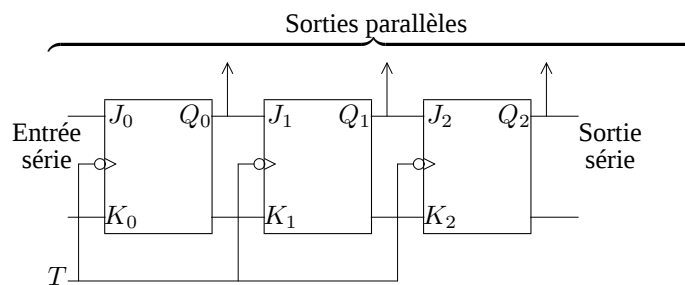
- (c) *Compteurs à cycles incomplets. Définition :* un tel compteur réalisé avec m bascules revient à 0 après p impulsions d'horloge, avec $p < 2^m$. On utilise pour ce faire des bascules JK modifiées, auxquelles on a adjoint des entrées « Preset » et « Clear », asynchrones, qui ont pour effet immédiat une fois activées de mettre respectivement la sortie à 1 ou 0.

Les compteurs ont de nombreuses applications :

- (a) comptage direct : par exemple comptage d'objets sur un tapis roulant, du nombre de pilules à mettre dans un flacon, etc.
- (b) division de fréquence par une puissance de 2 : lorsqu'on regarde les chronogrammes d'un compteur, comme par exemple le compteur trois bits plus haut, il est évident que chaque bit du compteur produit un signal en créneau dont la fréquence est égale à la fréquence de l'horloge divisée par une puissance de 2 dépendant du poids du bit ;
- (c) mesure de fréquence : il est possible d'utiliser le compteur pour compter le nombre de passages à zéro d'un signal donné pendant 1s, par exemple. La valeur indiquée par le compteur au bout d'une seconde est proportionnelle à la fréquence du signal analysé ;
- (d) il existe encore beaucoup d'autres applications : mesures de temps et donc de distance, de vitesse, suivi des opérations dans un calculateur numérique, multiplexage temporel, etc.

2. Registres à décalage

- (a) *Définition, structure.* Constitué de m bascules, il y a décalage de l'information d'une bascule à l'autre à chaque impulsion d'horloge. Par exemple :



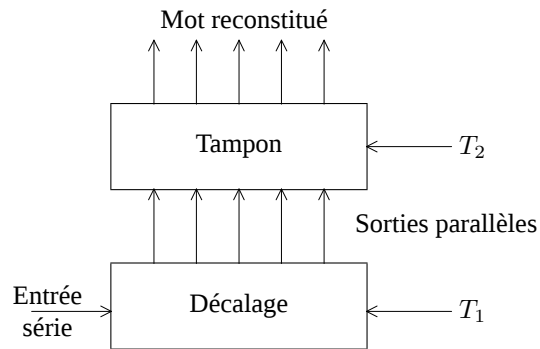
- (b) *Exemples d'application.*

- Conversions série↔parallèle et parallèle↔série ;
- Division/multiplication par des puissances de deux : pour diviser un nombre binaire par 2, en effet, il suffit de décaler ses bits d'un rang vers la droite.

3. Registres tampon (ou latches)

- (a) *Définition, structure. Association de bascules.* C'est une association de bascules (D, le plus souvent), sans interactions directes les unes avec les autres : un registre tampon met en mémoire sur ses sorties le mot présent sur ses entrées à la dernière impulsion d'horloge, et le garde jusqu'à la prochaine impulsion d'horloge.
- (b) *Exemple d'application :* En association avec un registre à décalage, un registre tampon permet la conversion série↔parallèle : le mot parallèle est transmis sur les sorties du registre tampon quand il est en place

sur les sorties du registre à décalage. Par exemple :



Dans cet exemple, la fréquence de l'horloge 2 doit être le cinquième de celle de l'horloge 1.

4. **Mémoire vive (RAM)**^{5.13} Ce sont des « boîtes noires », avec des entrées d'adresse, des entrées de données et des sorties de données. On y trouve aussi une entrée de commande lecture/écriture, permettant de choisir le mode de fonctionnement :

- soit écrire une donnée à l'adresse définie par le mot d'adresse ;
- soit lire la donnée présente à l'adresse définie par le mot d'adresse, et qui y a été écrite antérieurement.

Il s'agit d'une amélioration de la fonction de mémoire temporaire d'un registre tampon : on a la possibilité de stocker plusieurs mots simultanément.

5.3.3 Synthèse des systèmes séquentiels synchrones

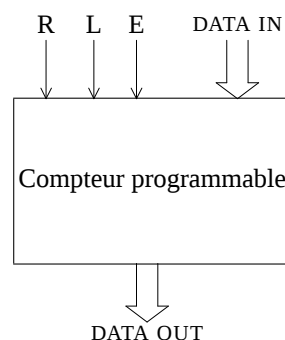
Il existe principalement trois voies systématiques permettant de synthétiser les deux schémas-blocs du paragraphe 5.3.1.1. C'est en fonction de la complexité de l'action à réaliser par le système que s'effectue le choix de l'une ou l'autre méthode.

5.3.3.1 Registres de bascules

C'est la manière la plus simple : il « suffit » de traduire le cahier des charges en une suite d'états différents. Cette suite d'états est alors traduite à l'aide de registres de bascules (par exemple des JK).

5.3.3.2 Compteur programmable

Dans cette solution, les bascules du registre sont organisées pour constituer un compteur possédant certaines propriétés. Il possède des entrées de commande synchrones : entrée d'inhibition (ENABLE), entrée de chargement (LOAD), entrées de données (DATA IN) et entrée de remise à zéro (RESET).



5.13. RAM pour Random Access Memory.

- *Entrée d'inhibition* : le compteur est inhibé quand cette entrée est activée; le mot binaire qu'il délivre reste inchangé à chaque impulsion d'horloge ;
- *Entrée de chargement* : quand cette entrée est activée et que survient l'impulsion d'horloge, le mot délivré par les sorties du compteur n'est pas incrémenté^{5.14} mais remplacé par le mot présent sur les entrées de données ;
- *Entrées de données* : le mot appliqué sur les entrées est transféré sur les sorties de façon synchrone quand l'entrée de chargement est activée ;
- *Entrée de remise à zéro*: quand elle est activée, le compteur est remis à zéro de façon synchrone.

Un programme est une suite d'instructions binaires envoyées sur les entrées de contrôle (R, L et E) et les entrées de données du compteur programmable, appelé pour l'occasion « compteur de programme ».

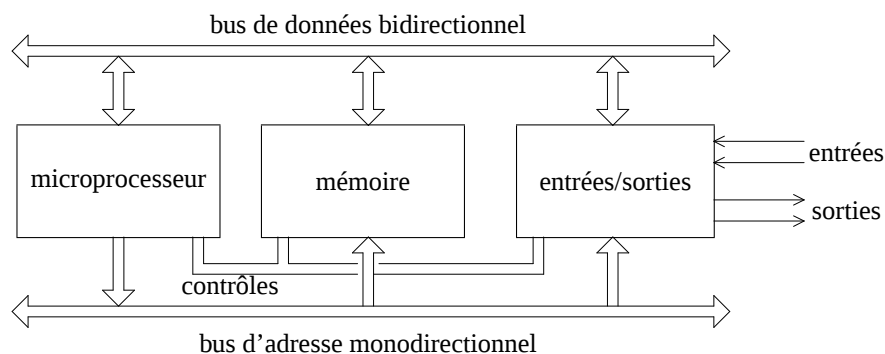
5.3.3.3 Unité centrale de contrôle et de traitement (CPU) : microprocesseur

CPU signifie Central Processing Unit. Le microprocesseur est un circuit intégré qui comporte :

1. un circuit séquentiel qui réalise les actions demandées par les instructions. Chaque instruction est un mot binaire qui est appliqué sur les entrées de ce circuit. Les actions qu'il accomplit se limitent à des transferts de données entre des registres et un compteur de programme ;
2. une unité arithmétique et logique. Les données et les instructions transitent par l'intermédiaire d'un bus de données interne. Le circuit séquentiel contrôle ces transferts par l'intermédiaire des lignes de contrôle^{5.15}.

Le microprocesseur communique avec l'extérieur par l'intermédiaire d'un bus de données bidirectionnel, venant de la mémoire programme, d'un bus d'adresse monodirectionnel et de lignes de contrôle. Les données sont émises par le compteur de programme.

Le microprocesseur doit fonctionner avec un certain nombre de circuits associés. Le programme est contenu dans une mémoire extérieure. Une mémoire peut de plus être nécessaire pour contenir des résultats qui devront être réutilisés. Les entrées et les sorties du système se font le plus souvent par l'intermédiaire de circuits d'entrée/sortie spécialisés. Les circuits associés sont connectés au microprocesseur par les bus (données et adresses) et les lignes de contrôle :



Ce qu'il faut retenir

- la différence entre logiques synchrone et asynchrone ;
- le vocabulaire : bascule, registre, compteur, microprocesseur...

5.14. C'est-à-dire augmenté d'une unité.

5.15. Read/Write, Enable, Set, Reset, ClOCK, etc.

5.4 Numérisation de l'information

La nature qui nous environne est perceptible par nos sens. Mais l'archivage de ces perceptions (la mémoire) et son traitement (la pensée) ne sont pas des moyens adaptés respectivement au *partage* de ces informations, non plus qu'à leur traitement *massif*. Pour ce faire, nous utilisons des outils numériques. Mais ces outils, par leur nature même, requièrent une interface avec le monde « réel », qui est lui fondamentalement analogique, pour autant que nous le sachions. La réalisation de cette interface est un problème complexe, dont nous n'allons aborder que quelques idées.

On peut distinguer plusieurs composants nécessaires à cette réalisation :

- un échantillonneur ;
- un convertisseur analogique/numérique ;
- le système numérique de traitement ;
- un convertisseur numérique/analogique éventuellement.

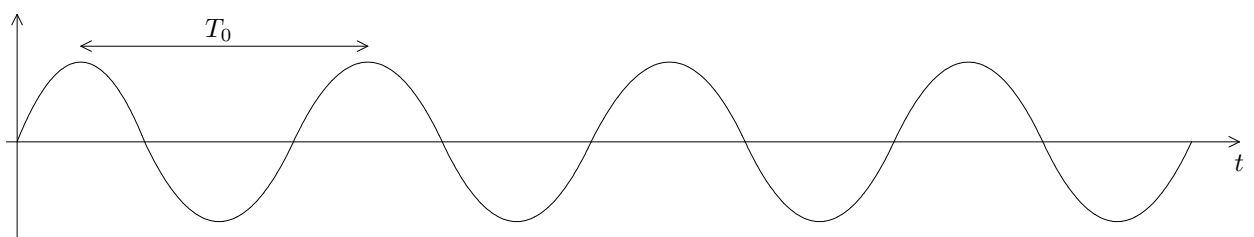
5.4.1 Le théorème de Shannon

5.4.1.1 Nécessité de l'échantillonnage

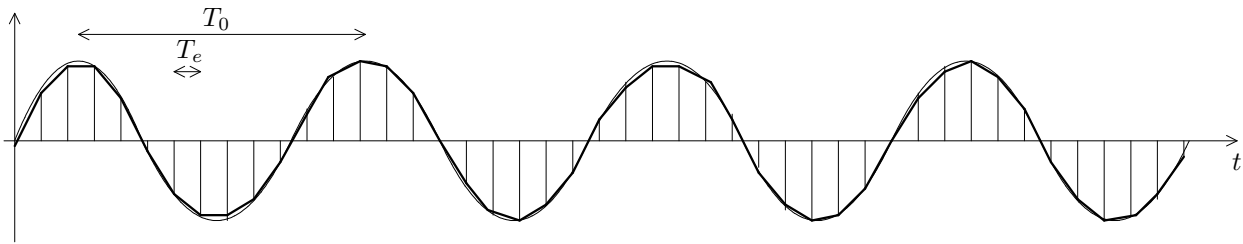
Les « gros » systèmes de traitement de l'information sont aujourd'hui fondamentalement séquentiels synchrones, pour des raisons liées aux indéterminations sur les états des sorties dans un circuit purement combinatoire (cf. paragraphe 5.3.1.2). Il faut donc pouvoir définir les entrées à des instants particuliers : on appelle cette opération l'*échantillonnage*.

5.4.1.2 Exemple : échantillonnage d'une sinusoïde

On considère une sinusoïde $x(t) = x_0 \sin \omega_0 t$, avec $\omega_0 = 2\pi f_0 = 2\pi/T_0$. On échantillonne cette sinusoïde aux instants t_n définis par $t_n = n/f_e$, où $f_e = 1/T_e$ désigne la *fréquence d'échantillonnage*.

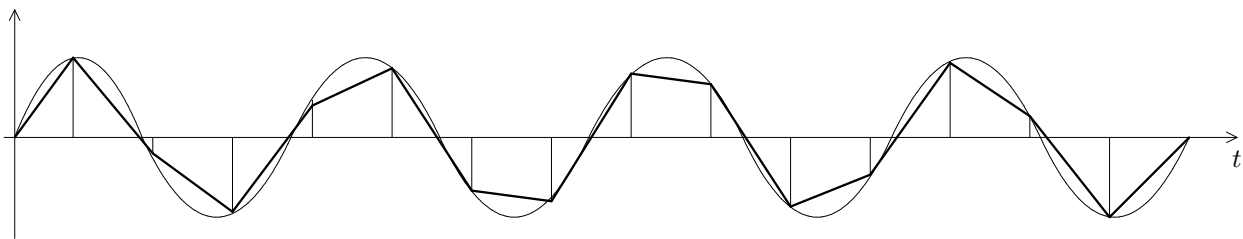


Si f_e est grande devant la fréquence de la sinusoïde, comme sur le schéma suivant,

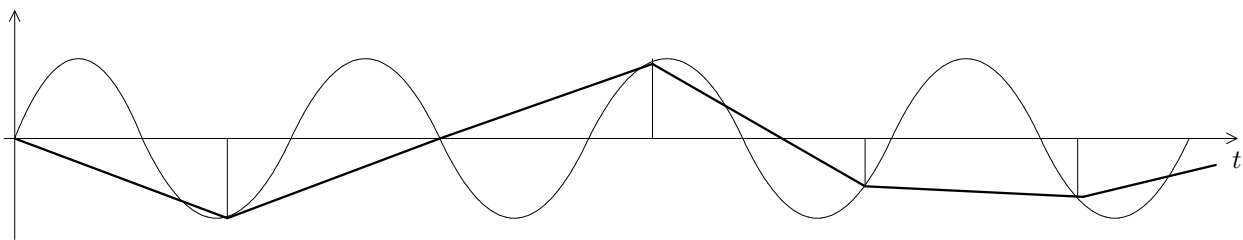


la reconstruction de celle-ci peut se faire facilement.

Si on diminue la fréquence d'échantillonnage, le signal reconstruit ressemble « moins » à une sinusoïde, mais reste quand même reconnaissable. Plus précisément, si on *sait* a priori que le signal est une sinusoïde, on peut encore en déterminer la période et l'amplitude, et il est donc entièrement reconstituible :



En revanche, si on diminue trop cette fréquence, le signal reconstruit ne ressemble en rien au signal original, et on perd l'information sur la fréquence du signal échantillonné :



En fait, il est possible de démontrer que l'on arrive à reconstruire mathématiquement le signal à partir de ses échantillons^{5.16} lorsque $f_e \geq 2f_0$.

5.4.1.3 Cas général

Dans la réalité, un signal physique *observé* est *toujours* à « support fréquentiel infini », autrement dit il transporte de l'énergie à toutes les fréquences réelles comprises entre 0 et $+\infty$ ^{5.17}.

5.16. On utilise dans la démonstration une forme de la transformée de Fourier adaptée aux signaux à temps continus échantillonnés, et la formule de la transformée de Fourier du peigne de Dirac (cf. paragraphe 1.2.3.3).

5.17. En effet, si on observe ce signal x physique, on ne le fait que pendant un intervalle de temps déterminé T . Cela revient à observer en fait le signal x multiplié par un signal « porte », valant 1 pour $0 < t < T$, et 0 ailleurs. D'après le paragraphe 1.2.2.6 sur la convolution, la Transformée de Fourier du signal observé est la convolution de la TF de x et de celle du signal porte. Or cette dernière est un « sinus cardinal » (cf. annexe A), dont le « support fréquentiel » est infini. Le support fréquentiel du signal *observé* est donc en toute rigueur infini.

En pratique néanmoins, les contributions des fréquences au-delà d'une certaine limite f_{max} sont négligeables devant les autres. On peut donc écrire en partant de la Transformée de Fourier inverse :

$$x(t) = \int_{-\infty}^{+\infty} X(\nu) e^{j2\pi\nu t} d\nu \approx \int_{-f_{max}}^{+f_{max}} X(\nu) e^{j2\pi\nu t} d\nu$$

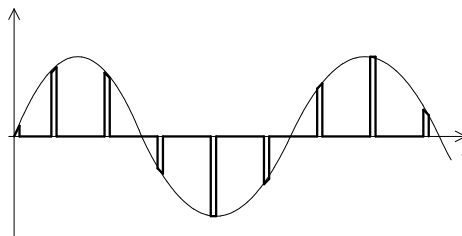
Tout se passe comme si x était une somme (l'intégrale) d'une infinité de sinusoides de fréquences ν et d'amplitudes $X(\nu)$, ces fréquences variant jusqu'à f_{max} . Pour assurer la reconstruction de x après un échantillonnage, il suffit d'assurer la reconstruction de toutes ses composantes fréquentielles, et pour cela d'échantillonner à une fréquence telle que même la sinusoïde de plus haute fréquence sera reconstruite correctement. On en déduit le *théorème de Shannon* :

Si x est un signal à support fréquentiel limité à f_{max} , alors x peut être entièrement reconstruit à partir de ses échantillons pris à la fréquence f_e si f_e vérifie $f_e \geq 2f_{max}$.

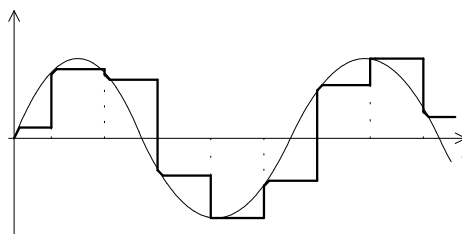
5.4.2 Les échantillonneurs

L'échantillonneur est un composant essentiel destiné à relier le « monde analogique » au « monde numérique ». Il prend en entrée un signal analogique, et est branché en sortie sur un système numérique de traitement de données via un convertisseur analogique/numérique (cf. paragraphe 5.4.3). On distingue deux types de circuits échantillonneurs :

1. **Echantillonneur simple** : entre chaque prise d'échantillon, le signal en aval de l'échantillonneur revient à 0 :



2. **Echantillonneur bloqueur** : entre chaque prise d'échantillon, le signal en aval de l'échantillonneur est maintenu à son niveau au moment du blocage :



Le deuxième type d'échantillonneur présente l'avantage de garder au signal sa « forme » initiale ; de plus, le système est plus robuste car le convertisseur analogique/numérique en aval a plus de temps pour réaliser sa conversion. En revanche, si l'on doit transporter sur une relativement grande distance le signal issu de l'échantillonneur avec le minimum d'énergie, le premier type doit être préféré. En effet, le signal échantillonné revient à 0, et l'énergie qu'il transporte est d'autant plus faible qu'il reste longtemps à cette valeur.

5.4.3 Convertisseur analogique/numérique (CAN)

5.4.3.1 Généralités

Un convertisseur analogique/numérique^{5.18}, comme son nom l'indique, convertit un signal analogique en un signal numérique. Il est à noter qu'on perd de la précision dans l'opération : alors qu'en analogique on peut espérer avoir une précision infinie, en numérique on est limité par le *pas de l'échantillonnage*. En règle générale, on aura à établir un compromis entre la précision désirée et la rapidité de conversion.

Un CAN aura une sortie série ou parallèle. Il sera nécessairement composé :

- d'une horloge, car les bits à la sortie du composant auront une durée déterminée ;
- de signaux de contrôle (début de conversion, fin de conversion, enable, reset...) pour lui permettre de dialoguer avec le système numérique situé en aval ;
- d'un système « producteur de bits » : registre, compteur ou simple circuit de logique combinatoire ;

5.4.3.2 Les caractéristiques d'un CAN

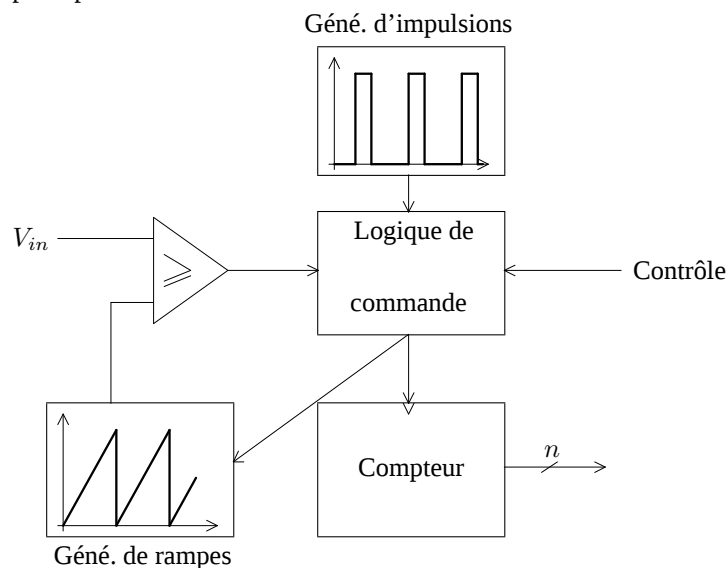
Un CAN sera caractérisé par :

- sa *résolution*, définie comme étant la plus petite variation de l'entrée entraînant une variation d'un bit en sortie ;
- son *temps de conversion* ;

5.4.3.3 Quelques CAN

1. Convertisseurs à comptage d'impulsions :

- *Convertisseurs à rampe* : L'idée est de comparer le signal à convertir à une tension augmentant régulièrement. Le schéma de principe est le suivant :



Les signaux de contrôle déclenchent le début de conversion. La logique de commande met en marche le générateur de rampes. La sortie de celui-ci est comparée à la tension à convertir. Tant qu'elle lui reste inférieure, la sortie du comparateur reste à 1 par exemple, et la logique de commande transfère le signal du

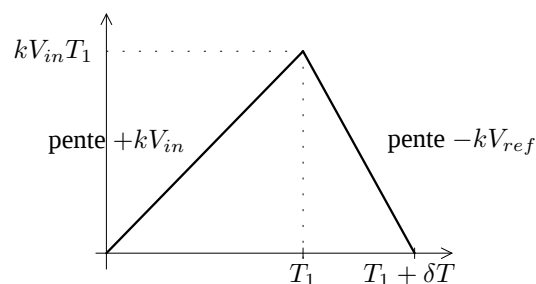
5.18. ou Analog to Digital Converter, soit ADC.

générateur d'impulsions au compteur, sur l'entrée d'horloge de celui-ci. Le compteur s'incrémente donc. Quand la tension de sortie du générateur de rampe devient supérieure à la tension à convertir, la sortie du comparateur change d'état, ce qui entraîne le blocage de la sortie de la logique de commande vers le compteur : celui-ci ne reçoit plus son horloge, et son état est *fixé*, jusqu'aux mises à zéro préalables à toute nouvelle conversion.

Ce convertisseur très simple est cependant lent et sensible au bruit (des fluctuations sur la sortie du générateur de rampes peuvent entraîner des erreurs de conversion).

On peut introduire un convertisseur à double rampe. Le principe est alors d'intégrer la tension à convertir V_{in} pendant une durée déterminée, ce qui produit un signal linéaire, pendant cette durée, dont la pente est proportionnelle à V_{in} . On commute alors l'entrée de l'intégrateur sur une tension de référence V_{ref} , de signe opposé à V_{in} , jusqu'à ce que le signal revienne à zéro. On a donc produit un signal composé de deux rampes :

- une rampe de durée fixe et de pente proportionnelle à V_{in} ;
- une rampe de pente fixe, de signe opposé et de durée proportionnelle à la valeur atteinte à la fin de la rampe précédente.

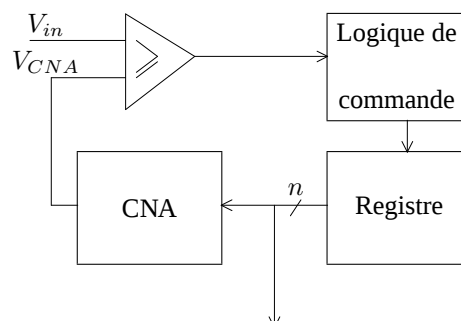


Il suffit de mesurer le temps de retour au zéro δT pour avoir une mesure proportionnelle à V_{in} : $\delta T = \frac{T_1}{V_{ref}} V_{in}$. Ce convertisseur est peu coûteux, remarquablement insensible au bruit, mais est très lent ce qui le rend inadapté à l'acquisition rapide de données. En revanche, ce principe de fonctionnement est utilisable pour des appareils de mesure.

- **Convertisseurs tension/fréquence** : Un *convertisseur tension/fréquence* ou VCO^{5.19} délivre des impulsions à une fréquence proportionnelle à sa tension d'entrée. Il suffit alors d'utiliser cette sortie comme signal d'horloge d'un compteur. Le nombre de changements d'état de ce dernier en un temps déterminé est proportionnel à la tension d'entrée. Ce convertisseur est simple car il nécessite peu de composants, et peu sensible au bruit car les variations rapides de la tension d'entrée n'ont que peu d'influence sur le nombre d'impulsions comptées ; mais il est peu précis car il est difficile de réaliser un VCO de précision meilleure que 1%.

2. Autres types de convertisseurs :

- *Convertisseur à approximations successives* :



Le CNA est un convertisseur numérique/analogique. On calcule les bits par comparaisons successives, en partant du bit de plus fort poids. Par exemple, prenons un tel convertisseur sur 4 bits (faible précision), conçu pour convertir des tensions de 0 à 10V (on dit alors que la *pleine échelle* est de 10V), et supposons qu'il ait à convertir 2,9V. Le pas d'échantillonnage vaut $10/2^4 = 0,625V$. La sortie du comparateur vaut 1

5.19. Pour Voltage Controlled Oscillator.

si $V_{CNA} < V_{in}$, 0 sinon.

Registre	CNA	Comparateur	Actions
0000	0V	1	bit suivant
1000	5V	0	remise à zéro, bit suivant
0100	2,5V	1	bit suivant
0110	3,75V	0	remise à zéro, bit suivant
0101	3,125V	0	remise à zéro, fin conversion
0100	2,5V	1	état final

Ce convertisseur est rapide, et surtout son temps de conversion est constant, au contraire des convertisseurs à rampe. Il est cependant complexe et très sensible au bruit (un pic transitoire à 3,2V dans l'exemple précédent entre les étapes 4 et 5 aurait entraîné une sortie à 0101 au lieu de 0100).

- *Convertisseur flash* : On compare la tension V_{in} à n tensions de référence *simultanément*. Par exemple, si on a à réaliser un convertisseur 3 bits avec une pleine échelle de 4V (d'où un pas d'échantillonnage de $4/2^3 = 0.5V$), on comparera les tensions en entrées à toutes les tensions comprises entre 0 et 4V par pas de 0,5V : 0, 0.5, 1, 1.5, 2, 2.5, 3, 3.5, 4V. La sortie du convertisseur est simplement la réunion des sorties des comparateurs. Ce convertisseur est très rapide, mais il nécessite un bon nombre de composants.

5.4.4 Convertisseur numérique/analogique (CNA)

5.4.4.1 Généralités

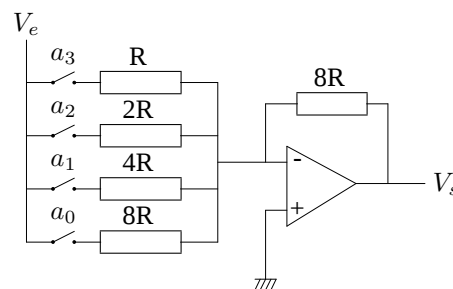
Si on veut pouvoir utiliser les capacités de calcul des ordinateurs pour commander des machines, par exemple, il faut avoir à disposition des systèmes capables d'assurer la traduction des sorties numériques de circuits logiques en signaux utilisables par ces machines. C'est le rôle des convertisseurs numériques/analogiques.

On distingue deux problèmes :

- le choix du code binaire : un CNA conçu pour fonctionner avec en entrée un code binaire naturel, par exemple, ne pourra pas être utilisé s'il reçoit en entrée un code binaire réfléchi ;
- le temps de conversion devra être adapté aux besoins.

5.4.4.2 Un exemple de CNA

L'idée générale est d'utiliser les signaux logiques à 1 ou 0 en commande d'interrupteurs. Par exemple on peut réaliser le montage suivant, avec 4 bits : a_3 et a_0 désignent respectivement les bits de plus fort et de plus faible poids, et le codage utilisé est le binaire naturel. L'interrupteur i est fermé quand $a_i = 1$.



Il est facile de vérifier que V_s vérifie $V_s = -V_e \sum_{i=0}^3 a_i 2^i$. La précision de ce convertisseur simple est limitée par celles sur V_e et les résistances.

5.4.4.3 Applications des CNA

On les trouve à chaque fois qu'il s'agit de reconstituer un signal à partir de valeurs numériques : lecteur CD, oscilloscope numérique, etc. Les CNA sont également utilisés dans certains convertisseurs analogique/numérique (cf. paragraphe 5.4.3.3), et aussi évidemment dans les asservissements numériques.

Ce qu'il faut retenir

- Le théorème de Shannon ($f_e \geq 2f_{max}$);
 - La différence entre un échantillonneur et un échantillonneur-bloqueur ;
 - Les fonctions d'un CAN et d'un CNA ;
-

Chapitre 6

Transmission de l'information

6.1 Généralités

6.1.1 Quelques dates

On peut dire que l'étude « raisonnée » de la manière dont l'information est transmise date d'un bon siècle et demi :

- **1831** : découverte de l'induction par Faraday (production d'effets électriques à distance sans liaison galvanique) ;
- **1887** : mise en évidence par Hertz de la propagation des ondes électromagnétiques ;
- **1901** : première liaison transatlantique par Marconi et début de la télégraphie ;
- **1906** : invention de la « triode » : on peut produire et amplifier des ondes électromagnétiques ;
- **1915** : première liaison transpacifique avec un relais à Honolulu ;
- **1920** : découverte de l'ionosphère entre 80 et 300km d'altitude. Les ondes électromagnétiques de fréquences supérieures à 30MHz s'y réfléchissent, ce qui rend possibles des liaisons par réflexions ionosphériques ;
- **années 30** : invention de la télévision ;
- **années 40** : invention du radar ;
- **années 50** : liaisons « troposphériques » : les ondes sont réfléchies par la turbulence des basses couches de l'atmosphère terrestre ;
- **années 60** : premiers satellites défilants et géostationnaires.

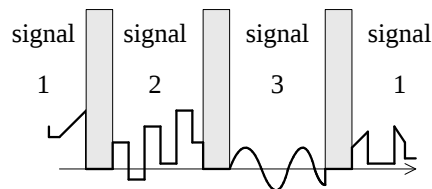
6.1.2 Nécessité d'un conditionnement de l'information

On a vu qu'un signal électrique pouvait être traité de manière analogique et/ou numérique. Cette relative facilité de manipulation peut être exploitée quand il s'agit de transmettre et d'analyser des informations. Mais le nombre de celles-ci et surtout la grande diversité de leurs supports (ondes radio, ondes lumineuses, courants électriques, etc.) impliquent qu'il soit procédé à leur pré-traitement, un codage, avant leur émission, et donc à un décodage à leur réception. Nous allons nous intéresser dans la suite de ce chapitre aux moyens de transmettre les informations par modulation.

6.1.3 Transports simultanés des informations

Il faut parfois transporter plusieurs signaux à la fois. Il est alors nécessaire de pouvoir bien les séparer à la réception. Ce but est atteint en délimitant précisément le « support » de chaque signal :

- par multiplexage temporel : si on doit transmettre trois signaux par exemple, on transmet un échantillon du signal 1, puis un échantillon du signal 2, enfin du signal 3, pour revenir ensuite au signal 1, etc. :



On se réserve un intervalle de temps de sécurité entre chaque transmission de signal.

- par « découpage fréquentiel » : supposons par exemple qu'un signal x_1 ait un spectre compris entre f_{11} et f_{12} , et qu'un signal x_2 ait un spectre compris entre f_{21} et f_{22} , avec $f_{11} < f_{12} < f_{21} < f_{22}$. On peut transmettre *simultanément* les deux signaux en les additionnant. A la réception, il suffira de connaître les bandes de fréquence où sont respectivement contenus x_1 et x_2 pour extraire ces deux signaux par filtrages passe-bande.

6.1.4 Introduction sur les modulations

L'idée générale des modulations est de *mélanger* le signal électrique contenant l'information avec un autre signal, dit *modulant*, tel que le signal modulé que l'on obtient en sortie contienne encore l'information sous une forme ou une autre, et se propage « bien ».

Prenons pour simplifier un signal modulant sinusoïdal de la forme

$$a(t) = a_0 \cos(\omega_0 t + \phi_0) = a_0 \cos[\phi(t)]$$

On peut agir

- sur a_0 : c'est la modulation d'amplitude ;
- sur ω_0 : c'est la modulation de fréquence ;
- sur ϕ_0 : c'est la modulation de phase.

Ce qu'il faut retenir

- Nécessité d'un « fractionnement » de l'information et d'un codage ;
 - Multiplexage temporel et découpage fréquentiel ;
 - Les types de modulations : d'amplitude, de fréquence, de phase.
-

6.2 Emission d'informations

6.2.1 Modulation d'amplitude

6.2.1.1 Introduction

Beaucoup d'informations transitent par des ondes électromagnétiques. Mais certaines fréquences ne sont pas correctement transmises dans l'atmosphère. Par ailleurs il existe parfois des contraintes légales d'« encombrement spectral » : certaines bandes de fréquence sont interdites. Une des parades trouvées pour décaler en fréquence un signal que l'on veut transmettre par voie hertzienne est la *modulation d'amplitude* : il s'agit de modifier l'amplitude d'un signal de fréquence f donnée, que l'on veut transmettre, par un signal de *plus basse fréquence*^{6.1}.

6.2.1.2 Modulation à porteuse conservée

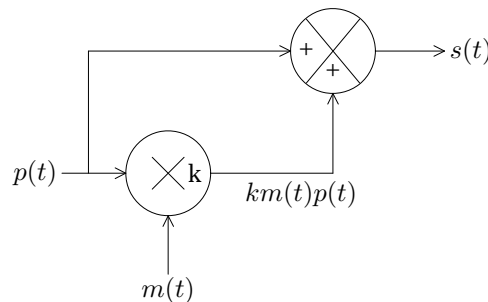
On considère un signal à transmettre $m(t)$, appelé *modulant*, et un signal $p(t) = A_p \sin \omega_p t$, appelé *porteuse*.

1. Aspect temporel

La modulation à porteuse conservée consiste à construire le signal

$$s(t) = A_p \left(1 + k \frac{m(t)}{A_p} \right) \sin \omega_p t$$

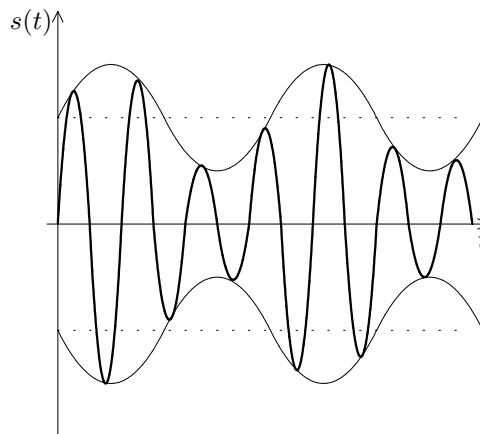
L'opération précédente se schématise sous forme de schéma-bloc comme suit :



Si on écrit que A_m est l'amplitude maximale atteinte par m , on définit l'*indice de modulation* β par $\beta = k \frac{A_m}{A_p}$.

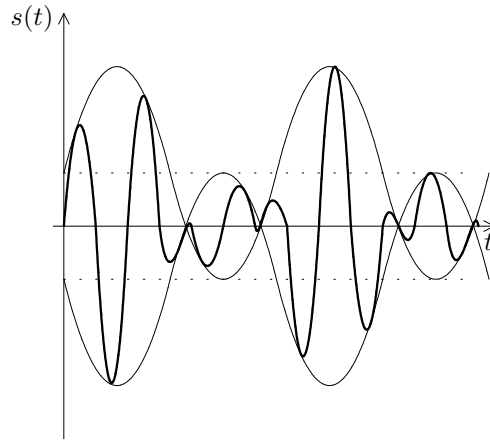
On parle alors :

- de modulation « classique » si $\beta < 1$: exemple avec $\beta = 0.5$ et un modulant $m(t) = A_m \sin \omega_m t$:



6.1. Abrégée en AM, soit amplitude modulation, en anglais.

- de « surmodulation » si $\beta > 1$: exemple avec $\beta = 2$ et un modulant $m(t) = A_m \sin \omega_m t$:

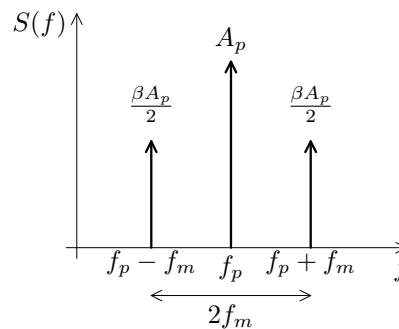


2. Aspect fréquentiel

Considérons un exemple simple où porteuse et modulant sont des signaux sinusoïdaux : $p(t) = A_p \sin \omega_p t$ et $m(t) = A_m \sin \omega_m t$. Il est facile de montrer en utilisant la relation $\cos a \cos b = [\cos(a+b) + \cos(a-b)]/2$, que :

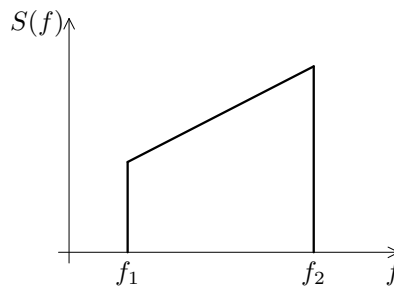
$$s(t) = A_p \cos \omega_p t + \frac{\beta A_p}{2} [\cos(\omega_m + \omega_p)t + \cos(\omega_m - \omega_p)t]$$

Le spectre de s est donc le suivant :

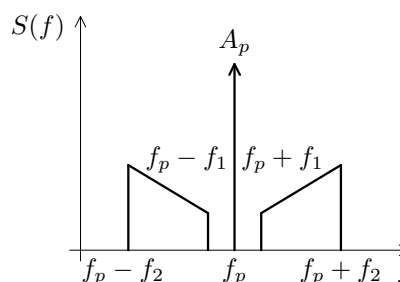


En pratique, on choisit $f_m \ll f_p$; la composante à $f_p - f_m$ est appelée *onde latérale inférieure* et la composante à $f_p + f_m$ *onde latérale supérieure*.

Remarque : en général, le modulant peut avoir un certain « encombrement spectral », ou support fréquentiel, entre deux fréquences f_1 et f_2 , encombrement que l'on représente par le schéma ci-dessous :



Cela se traduit par un spectre de s de la forme suivante :

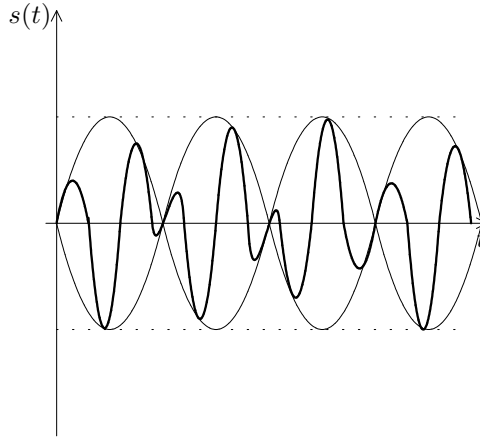


6.2.1.3 Modulation à porteuse supprimée

On construit cette fois-ci le signal $s(t) = A_p m(t) \sin \omega_p t$. On ne peut plus alors définir d'indice de modulation.

1. Aspect temporel

Par exemple, avec un modulant $m(t) = A_m \sin \omega_m t$ on obtient :

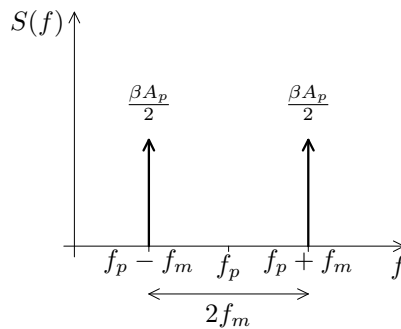


2. Aspect spectral

On a cette fois-ci :

$$s(t) = \frac{A_p A_m}{2} [\cos(\omega_m + \omega_p)t + \cos(\omega_m - \omega_p)t]$$

Cela se traduit par un spectre de s de la forme suivante :



Bien qu'encore virtuellement présente, la porteuse est cette fois-ci plus difficile à détecter. Cependant, l'énergie nécessaire pour transporter le signal est moindre que dans le cas de la modulation à porteuse conservée.

6.2.2 Modulations angulaires

6.2.2.1 Introduction

On considère un signal $s(t)$ que l'on peut écrire sous la forme $s(t) = s_0 \sin[\phi_i(t)]$. On appelle :

- $\phi_i(t)$ la *phase instantanée* à l'instant t du signal $s(t)$;
- $f_i(t)$ définie par $f_i(t) = \frac{1}{2\pi} \frac{d\phi_i}{dt}$ la *fréquence instantanée* du signal $s(t)$.

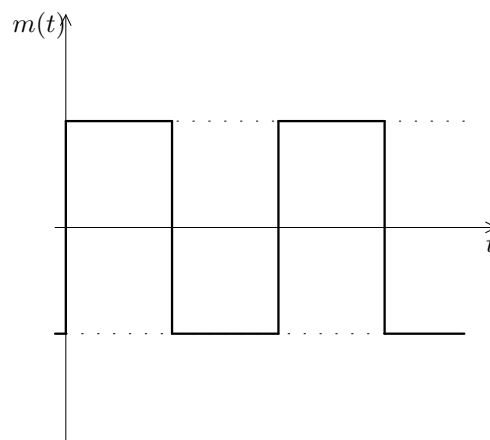
Pour un signal m à transmettre, on peut alors distinguer deux types de modulations angulaires :

- la modulation de phase, où on réalise $\phi_i(t) = \phi_0 + K_\phi m(t)$;

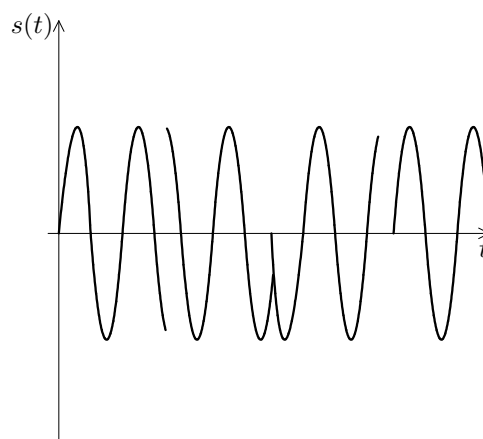
- la modulation de fréquence, où on réalise $f_i(t) = f_0 + K_f m(t)$ ^{6.2}.

6.2.2.2 Aspect temporel

Considérons par exemple un signal modulant $m(t)$ carré :

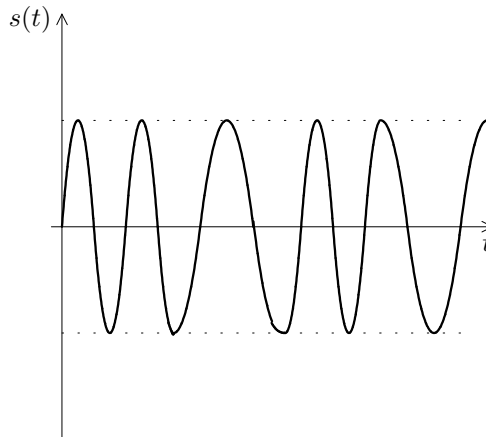


La modulation de phase donnera...



6.2. FM, soit frequency modulation, en anglais.

... et la modulation de fréquence :



6.2.2.3 Aspect fréquentiel de la modulation de fréquence

Dans le cas d'une modulation de fréquence, quand $m(t)$ est un signal sinusoïdal $A_m \cos \omega_m t$, on montre que le spectre de s est un « spectre de raies » : on ne trouve du signal qu'aux fréquences $f_0 \pm n f_m$, où n est un entier naturel. D'autre part, si on pose $\delta f = K_f f_m$ (δf est l'« excursion en fréquences »), on montre que la *bande utile* B de signal, c'est-à-dire la portion de s qu'il suffit de garder par filtrage pour reconstituer m dans de bonnes conditions, vaut : $B \approx 2(\delta f + f_m)$ (règle de Carson).

On distingue alors deux cas selon la valeur du rapport $\beta = \delta f / f_m$:

- $\beta \ll 1$: alors $B \approx 2f_m$, et on se rapproche du cas d'une « simple » modulation d'amplitude ;
- $\beta \gg 1$: alors $B \approx 2\delta f$, et il faut utiliser des méthodes spécifiques de démodulation de fréquence.

Ce qu'il faut retenir

- L'allure d'un signal modulé en amplitude, en phase ou en fréquence ;
- Modulation à porteuse conservée, à porteuse supprimée.

6.3 Réception d'informations

Après une émission marquée par une information codée en utilisant un type de modulation, il est nécessaire de pouvoir retrouver cette information à la réception ; on appelle naturellement cette opération la *démodulation*.

6.3.1 Démodulation d'amplitude

On supposera dans ce paragraphe que la porteuse est de la forme $p(t) = A_p \sin \omega_p t$, et que le signal modulé est :

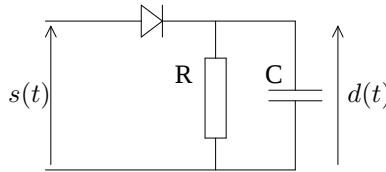
- de la forme $s(t) = A_p [1 + \frac{k}{A_p} m(t)] \sin \omega_p t$ dans le cas d'une transmission à porteuse conservée ;
- de la forme $s(t) = A_p k m(t) \sin \omega_p t$ dans le cas d'une transmission à porteuse supprimée.

Si de plus le signal modulant $m(t)$ est sinusoïdal, de la forme $A_m \sin \omega_m t$, on rappelle que l'indice de modulation β vaut $\beta = \frac{k A_m}{A_p}$.

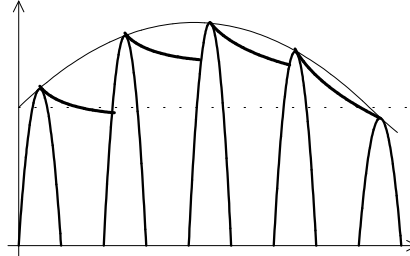
6.3.1.1 Démodulation incohérente

Ce type de démodulation est réservé à la transmission par porteuse conservée, avec un indice de modulation inférieur à 1.

1. **Détection d'enveloppe** : c'est le plus simple montage démodulateur. Il suffit de câbler le schéma suivant^{6.3} :



On doit calculer les valeurs de la capacité et de la résistance pour que la décroissance de la tension de sortie entre deux maxima locaux permette d'« accrocher » la partie ascendante suivante de la sinusoïde du modulant :

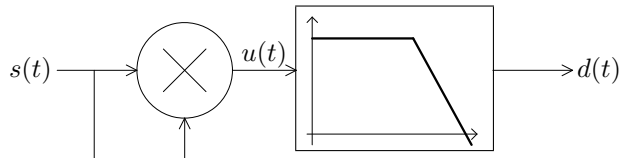


En fait, on peut montrer que la détection d'enveloppe est efficace si l'on respecte les conditions suivantes :

$$\frac{1}{\omega_p} \ll RC \ll \frac{\sqrt{1 - \beta^2}}{\beta \omega_m}$$

2. **Détection par filtrage et non-linéarité**. Sous cette dénomination on rassemble notamment deux types de détection, qui toutes les deux font appel à des processus non linéaires :

- la *détection quadratique*. On calcule d'abord le carré du signal à démoduler, puis on le filtre par un filtre passe-bas de fréquence de coupure judicieusement choisie :



Avec une porteuse sinusoïdale, on obtient :

$$u(t) = \frac{\alpha_1 A_p^2}{2} [1 + k^2 m(t)^2 + 2k m(t)] (1 + \cos 2\omega_p t)$$

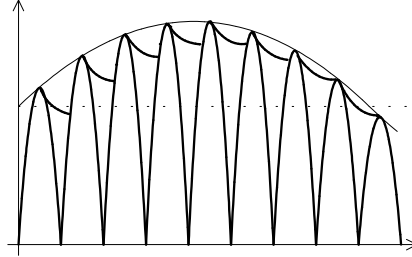
6.3. On trouvera la caractéristique tension-courant de la diode dans le paragraphe 7.4.2 sur les interrupteurs.

Si $m(t)$ est faible, autrement dit si l'indice de modulation est petit, alors $k^2 m(t)^2 \ll km(t)$ et après filtrage on obtient

$$d(t) \approx \frac{\alpha_2 A_p^2}{2} [1 + 2km(t)]$$

On retrouve donc le signal modulant $m(t)$.

- la *détection par redressement*. On « redresse » le signal à démoduler en en prenant la valeur absolue, puis on filtre par un filtre passe-bas. Le principe est le même que dans le cas de la détection d'enveloppe, mais cette fois-ci le fait de prendre la valeur absolue du signal modulé permet de garder toutes les arches de sinusoïdes :

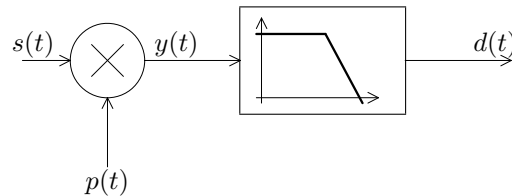


6.3.1.2 Détection synchrone

On peut utiliser une méthode de détection synchrone quand la porteuse est conservée ou non.

1. **La porteuse est disponible :** soit par filtrage passe-bande sélectif quand la porteuse est conservée, soit en utilisant un oscillateur local produisant un signal de fréquence *exactement* égale à celle de la porteuse à un déphasage près ϕ (les deux options sont difficiles), ou bien encore simplement en *transmettant* la porteuse en parallèle du signal.

Si la porteuse $p(t)$ est disponible, sans déphasage résiduel par rapport à l'émission, alors il suffit de la multiplier par le signal à démoduler, puis d'appliquer un filtre passe-bas pour obtenir le signal modulant. Si elle est disponible à un déphasage ϕ près, on montre que le signal démodulé $d(t)$ obtenu après filtrage est multiplié par $\cos \phi$, ce qui dégrade la restitution du signal modulant.



En effet, avec $s(t) = A_p [1 + \frac{k}{A_p} m(t)] \sin \omega_p t$, si on multiplie par $p_1(t) = A_p \sin(\omega_p t + \phi)$, on obtient $y(t) = s(t) \cdot p_1(t) = \frac{A_p}{2} [1 + \frac{k}{A_p} m(t)] [\cos \phi - \cos(2\omega_p t + \phi)]$. Un filtrage passe-bas coupant la composante à la fréquence $2f_p$ permet de retrouver le signal modulant, basse fréquence, pondéré par $\cos \phi$.

2. **Reconstitution de la porteuse :** lorsque la porteuse ne peut être transmise avec le signal, ou que sa « récupération » à partir du signal modulé par filtrage sélectif est trop difficile, on peut la reconstruire à partir du signal modulé en utilisant un dispositif appelé *boucle à verrouillage de phase*. On dispose alors d'un signal de même fréquence que la porteuse, et surtout *en phase* avec elle : il n'y a pas d'atténuation par un cosinus du déphasage.

6.3.2 Démodulation angulaire

La méthode la plus simple est d'utiliser un convertisseur fréquence-tension, composant délivrant en sortie un signal proportionnel à la fréquence du signal qu'on lui présente en entrée.

- pour une modulation de fréquence, il suffit d'extraire la partie non constante du signal de sortie. On récupère

en effet un signal démodulé de la forme $d(t) = d_0 + \tilde{d}(t) = K[f_0 + K_fm(t)]$. Un filtrage passe-haut, ou passe-bande suffira.

- pour une modulation de phase, on utilise le fait que la fréquence instantanée est la dérivée de la phase instantanée (cf. paragraphe 6.2.2.1). Il suffit alors d'*intégrer* le signal issu du convertisseur fréquence-tension.

Ce qu'il faut retenir

- La différence entre démodulation incohérente et détection synchrone ;
 - La détection d'enveloppe qui est le plus simple démodulateur d'amplitude ;
 - L'utilisation d'un convertisseur fréquence-tension pour les démodulations angulaires.
-

Chapitre 7

Notions d'électrotechnique

7.1 Le transformateur monophasé

7.1.1 Description, principe

7.1.1.1 Nécessité du transformateur

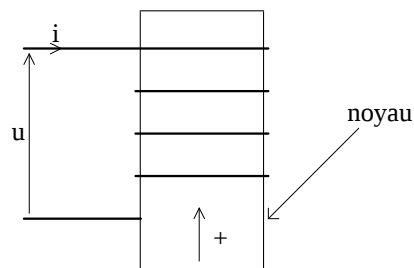
La *distribution* de l'énergie électrique (domestique ou industrielle) se fait généralement sous tension faible (220 ou 380V) pour des raisons de sécurité et de commodité d'emploi.

Le *transport* de cette énergie, en revanche, se fait sous tension élevée, afin que le courant de ligne soit faible pour minimiser les pertes par effet Joule le long de la ligne.

Ces deux exigences contradictoires rendent nécessaire une machine capable de modifier les caractéristiques (U,I) de l'énergie électrique tout en consommant le moins possible.

7.1.1.2 Principe du transformateur statique

1. **Induction** : lorsque l'on fait parcourir un fil électrique par du courant, un champ magnétique se forme^{7.1}.
2. **Bobine à noyau de fer** : une bobine à noyau de fer est un *circuit magnétique* dans lequel le champ magnétique est créé par un bobinage parcouru par un courant :



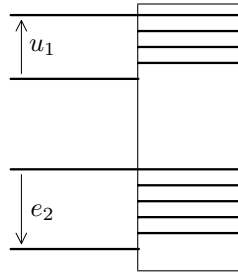
Il faut déterminer des conventions de signes pour les grandeurs algébriques intervenant dans la bobine. On choisit *arbitrairement* le sens positif pour le champ magnétique créé dans le noyau. L'orientation de i est alors définie de telle sorte que $i > 0$ crée un champ magnétique positif. L'orientation de u est telle que (u, i) correspond à une

^{7.1} Cas idéal : à la distance r d'un fil électrique infini parcouru par un courant I , le champ magnétique créé est « orthoradial », c'est-à-dire perpendiculaire au fil, et vaut $B(r) = \mu_0 I / 2\pi r$.

convention récepteur (cf. paragraphe 2.1.2.4). Dans un cas simple, u s'exprime en fonction de i par $u = ri + L \frac{di}{dt}$, où r désigne la résistance de la bobine et L son inductance propre. Dans un transformateur, les bobines vérifient $u = ri \pm n \frac{d\phi}{dt}$ (le signe exact dépend du sens d'enroulement du bobinage), où ϕ désigne le flux du champ magnétique^{7.2} dans le noyau et n le nombre de spires. Une bobine peut donc fonctionner en deux modes :

- *mode récepteur* quand la tension u est imposée de l'extérieur, par un générateur par exemple. Le flux est alors entièrement déterminé, contraint par u ;
- *mode émetteur* quand c'est le *flux* qui est directement imposé par l'expérimentateur : ses variations déterminent alors la tension de sortie.

3. **Transformateur statique** : si on dispose un second bobinage composé de n_2 spires sur le même circuit magnétique, il sera le siège d'une force électromotrice $e_2 = -n_2 \frac{d\phi}{dt}$ ^{7.3} et sera en mesure d'alimenter une impédance quelconque sous une différence de potentiel fixée arbitrairement par le choix de n_2 .

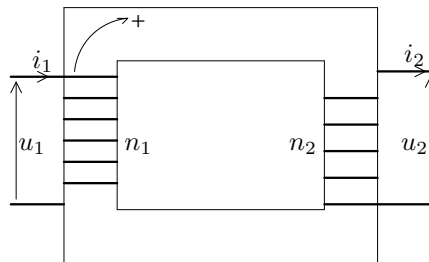


Un transformateur statique se composera donc d'un circuit magnétique sur lequel sont bobinés deux enroulements distincts.

7.1.2 Les équations du transformateur

7.1.2.1 Conventions algébriques

Le seul choix arbitraire est celui de $\vec{B}$ dans le circuit magnétique :



Sur chaque bobinage, i est tel que i positif entraîne un champ magnétique positif (d'où les sens de i_1 et i_2 sur l'exemple ci-dessus).

7.2. Le flux du champ $\vec{B}$ à travers une surface fermée S est égal à l'intégrale sur cette surface du produit scalaire de $\vec{B}$ avec les vecteurs normaux à la surface, par élément de surface :

$$\phi = \int \int_{(S)} \vec{B} \cdot \vec{n} \, d^2s$$

Dans le cas particulier d'un champ magnétique uniforme orthogonal à la surface, le flux ϕ est simplement le produit de l'intensité du champ par l'aire.

7.3. Rappel : on peut dire qu'une *force électromotrice* (en abrégé f.é.m.) est une tension imposée par un générateur (au sens large : pile électrique, générateur de tension sinusoïdale, bobine dans un champ magnétique variable, machine électrique, etc.). Une *force contre-électromotrice* (en abrégé f.c.é.m.) est aussi une tension produite par un générateur, mais s'opposant à une force électromotrice : la plus « forte » des deux gagne et est appelée électromotrice. Notez que si l'on branche en série un générateur de tension sinusoïdale (ou GBF, pour Générateur Basse Fréquence) délivrant une tension d'amplitude 5V, et une pile électrique de 12V, la moitié du temps le GBF fonctionnera en force électromotrice et l'autre moitié en force contre-électromotrice.

L'un des bobinages, appelé *primaire*, voit sa tension (d'entrée) définie d'après la convention récepteur ; l'autre, appelé *secondaire*, voit sa tension définie d'après la convention générateur.

On note n_1 et n_2 le nombre de spires respectivement du primaire et du secondaire, R_1 et R_2 leurs résistances.

En appelant e_1 et e_2 les forces électromotrices induites dans le primaire et le secondaire, les lois des mailles dans les deux bobinages s'écrivent :

$$\begin{cases} u_1 + e_1 &= R_1 i_1 \\ -u_2 + e_2 &= R_2 i_2 \end{cases}$$

7.1.2.2 Détermination des forces électromotrices induites

Il est possible de montrer que si ϕ désigne le flux passant à travers *une* spire (du primaire *ou* du secondaire), et respectivement l_1 et l_2 les inductances propres des bobinages primaire et secondaire, alors :

$$\begin{cases} e_1 &= -n_1 \frac{d\phi}{dt} - l_1 \frac{di_1}{dt} \\ e_2 &= -n_2 \frac{d\phi}{dt} - l_2 \frac{di_2}{dt} \end{cases}$$

7.1.2.3 Le transformateur parfait

On a à déterminer les valeurs de 4 variables : u_1 , u_2 , i_1 et i_2 . Pour ce faire, il nous faut 4 équations. On en a déjà deux avec 7.1.2.1 et 7.1.2.2. On peut y ajouter :

$$n_1 i_1 + n_2 i_2 = \mathcal{R} \phi$$

$\mathcal{R}$ est appelée *réluctance* et est une caractéristique du transformateur (plus précisément, $\mathcal{R}$ dépend du matériau et de la géométrie du transformateur). Il ne faut pas oublier non plus que le transformateur sera branché sur une « charge », et que donc il existera une relation supplémentaire entre u_2 et i_2 ^{7.4} :

$$f(u_2, i_2) = 0$$

Dans le cas du transformateur parfait, on fait les approximations suivantes :

- les résistances de bobinages sont nulles : $R_2 = R_1 = 0$;
- le circuit magnétique est parfait : il n'y a pas de « fuite de champ magnétique » et $l_2 = l_1 = 0$;
- le circuit magnétique a une réluctance nulle : $\mathcal{R} = 0$.

Les équations du transformateur parfait sont donc simplement :

$$\begin{cases} u_1 &= n_1 \frac{d\phi}{dt} \\ u_2 &= -n_2 \frac{d\phi}{dt} \\ n_1 i_1 + n_2 i_2 &= 0 \\ f(u_2, i_2) &= 0 \end{cases}$$

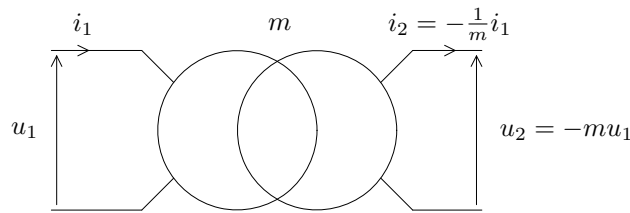
On en déduit la relation fondamentale :

$$\boxed{\frac{u_2}{u_1} = \frac{-n_2}{n_1} = -m = \frac{i_1}{i_2}}$$

où $m = \frac{n_2}{n_1}$ est appelé le *rapport de transformation* .

7.4. Par exemple, dans le cas où le secondaire est chargé par une impédance z_2 , la relation supplémentaire sera $u_2 - z_2 i_2 = 0$.

On constate également que $u_1 i_1 = u_2 i_2$: le dispositif restitue intégralement au secondaire la puissance électrique qu'il reçoit au primaire. Le transformateur parfait est représenté par le schéma suivant :



Une propriété importante du transformateur parfait est le *transfert d'impédance* : par exemple, si $u_2 = Z_2 i_2$, alors

$$u_1 = -\frac{1}{m} u_2 = -\frac{1}{m} Z_2 i_2 = \frac{1}{m} Z_2 \frac{i_1}{m}$$

D'où

$$u_1 = \left(\frac{Z_2}{m^2} \right) i_1$$

Une impédance Z_2 au secondaire est équivalente à Z_2/m^2 au primaire, et Z_1 placée au primaire est équivalente à $Z_1 m^2$ au secondaire.

Ce qu'il faut retenir

- Le transformateur est un moyen de transférer de la puissance en modifiant les caractéristiques tension/courant d'un réseau ;
- Les relations $\frac{u_2}{u_1} = \frac{i_1}{i_2} = -m$ où $m = -\frac{n_2}{n_1}$ est le *rapport de transformation*.

7.2 Systèmes triphasés

7.2.1 Définition et classification

7.2.1.1 Définition d'un système polyphasé

Un système n-phasé équilibré est constitué de n grandeurs sinusoïdales (tensions ou courants), de même amplitude, déphasées régulièrement de $m \cdot 2\pi/n$ entre elles. Il est donc composé de :

$$\begin{cases} y_1(t) &= Y\sqrt{2} \cos(\omega t - \phi) \\ y_2(t) &= Y\sqrt{2} \cos(\omega t - 1 \cdot m \frac{2\pi}{n} - \phi) \\ &(\dots) \\ y_k(t) &= Y\sqrt{2} \cos(\omega t - (k-1)m \frac{2\pi}{n} - \phi) \end{cases}$$

Y est la *valeur efficace* commune aux trois grandeurs, et m est l'*ordre* de ce système n-phasé équilibré^{7.5}.

Dans la suite de ce cours, on se limitera sauf mention contraire aux *systèmes triphasés*, avec $n = 3$.

^{7.5}. Dans un système non équilibré, les valeurs efficaces des différentes grandeurs sont différentes.

7.2.1.2 Systèmes direct, inverse et homopolaire

1. $m = 1$: Système d'ordre 1

$$\begin{cases} y_1(t) = Y\sqrt{2}\cos(\omega t - \phi) \\ y_2(t) = Y\sqrt{2}\cos(\omega t - \frac{2\pi}{3} - \phi) \\ y_3(t) = Y\sqrt{2}\cos(\omega t - \frac{4\pi}{3} - \phi) \end{cases}$$

On associe à ces trois grandeurs les représentations complexes :

$$\begin{cases} Y_1(t) = Y\sqrt{2}e^{j(\omega t - \phi)} \\ Y_2(t) = Y\sqrt{2}e^{j(\omega t - \phi)}e^{-j\frac{2\pi}{3}} \\ Y_3(t) = Y\sqrt{2}e^{j(\omega t - \phi)}e^{-j\frac{4\pi}{3}} \end{cases}$$

Dans la suite, on posera^{7.6} « classiquement » $a = e^{+j\frac{2\pi}{3}}$. On peut constater que

$$a^2 = e^{+j\frac{4\pi}{3}} \text{ et que } a^3 = 1$$

a , a^2 et 1 sont les trois « racines cubiques de 1 ». On peut donc écrire :

$$\begin{cases} Y_1 = Y_1 \\ Y_2 = a^2 Y_1 \\ Y_3 = a Y_1 \end{cases}$$

Ce système est appelé *système direct* : la troisième composante est en avance sur la deuxième, qui est en avance sur la première.

2. $m = 2$: Système d'ordre 2

On a cette fois-ci :

$$\begin{cases} y_1(t) = Y\sqrt{2}\cos(\omega t - \phi) \\ y_2(t) = Y\sqrt{2}\cos(\omega t - 2\frac{2\pi}{3} - \phi) \\ y_3(t) = Y\sqrt{2}\cos(\omega t - 4\frac{2\pi}{3} - \phi) \end{cases}$$

On associe alors les représentations complexes :

$$\begin{cases} Y_1(t) = Y\sqrt{2}e^{j(\omega t - \phi)} \\ Y_2(t) = Y\sqrt{2}e^{j(\omega t - \phi)}e^{+j\frac{2\pi}{3}} \\ Y_3(t) = Y\sqrt{2}e^{j(\omega t - \phi)}e^{-j\frac{2\pi}{3}} \end{cases}$$

Ce que l'on peut résumer en :

$$\begin{cases} Y_1 = Y_1 \\ Y_2 = a Y_1 \\ Y_3 = a^2 Y_1 \end{cases}$$

Ce système est appelé *système inverse* : la première composante est en avance sur la deuxième, puis sur la troisième.

3. $m = 3$: Système d'ordre 3

On a cette fois-ci :

$$\begin{cases} y_1(t) = Y\sqrt{2}\cos(\omega t - \phi) \\ y_2(t) = Y\sqrt{2}\cos(\omega t - 3\frac{2\pi}{3} - \phi) \\ y_3(t) = Y\sqrt{2}\cos(\omega t - 6\frac{2\pi}{3} - \phi) \end{cases}$$

On a donc $Y_1 = Y_2 = Y_3$. Ce système est appelé *système homopolaire*.

7.2.1.3 Propriétés des systèmes triphasés équilibrés

1. La somme des grandeurs du système triphasé équilibré est nulle.

En effet, que le système soit direct ou inverse :

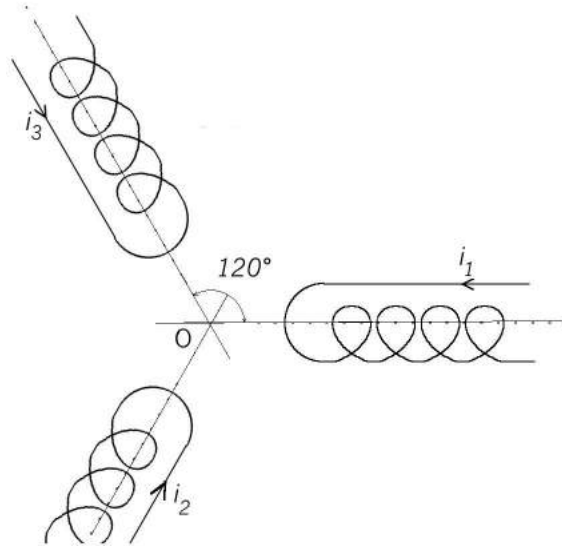
$$Y_1 + Y_2 + Y_3 = Y_1(1 + a + a^2)$$

Or 1, a et a^2 sont les trois racines cubiques de l'unité donc leur somme est nulle et : $Y_1 + Y_2 + Y_3 = 0$. Comme $y_1 + y_2 + y_3 = \Im(Y_1 + Y_2 + Y_3)$, on a donc $y_1 + y_2 + y_3 = 0$.

7.6. Ce nombre, racine cubique de 1, est noté j en mathématiques, mais on comprendra facilement la nécessité de changer ici de notation.

2. L'utilisation de grandeurs sinusoïdales équilibrées permet la création de champs magnétiques tournants.

On explicitera cette propriété plus loin, mais elle peut être considérée comme l'inverse de la façon d'obtenir des grandeurs triphasées (cf. théorème de Ferraris, paragraphe 7.3.1.2). Pour ce faire, supposons que l'on dispose de trois bobinages identiques, placés de telle manière que leurs axes de symétrie se coupent en un unique point, et décalés en position de $2\pi/3$. On place un aimant au point central.

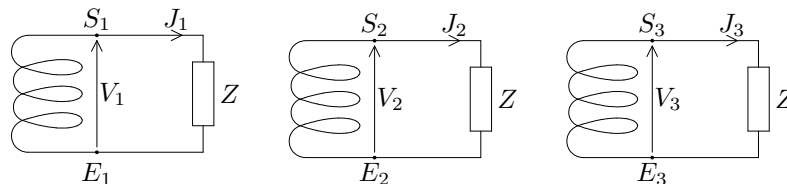


La rotation de l'aimant central à la vitesse angulaire ω produit dans ces trois bobinages des f.é.m. induites qui forment un système triphasé équilibré (à condition que le circuit magnétique et la forme des bobinages soient tels que les f.é.m. induites soient sinusoïdales).

7.2.2 Associations étoile et triangle

7.2.2.1 Position du problème

On dispose de trois bobinages du type de ceux introduits au paragraphe précédent. S'ils sont tous trois chargés par des impédances Z identiques, ils sont parcourus par des courants J_1, J_2, J_3 qui forment un système triphasé équilibré :

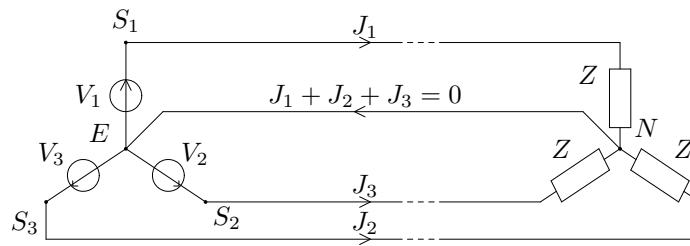


Un tel système n'offre pas d'intérêt par rapport à ce qu'on peut faire avec trois sources monophasées quelconques. Mais on peut coupler ces trois bobinages de deux manières différentes montrant les avantages du triphasé.

On assimilera par la suite, pour la simplicité des schémas, les bobinages à leurs sources de tension V_k équivalentes.

7.2.2.2 Association étoile

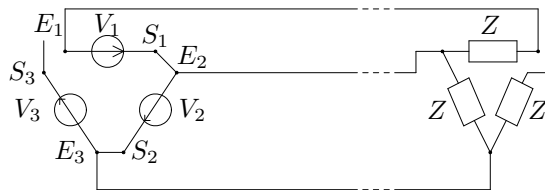
Elle est souvent notée Y ou *. On peut imposer le même potentiel aux bornes E_i . Le montage devient alors :



Il apparaît que le quatrième conducteur, où circule le courant $J_1 + J_2 + J_3 = 0$, peut être supprimé. Lorsqu'il existe, ce conducteur est appelé *fil neutre*, mais il est important de retenir que dans le cas où le système est *équilibré*, les points E et N sont au même potentiel, qui est le potentiel du neutre, identique pour tous les « N » du même réseau, qu'ils soient explicitement reliés entre eux ou non.

7.2.2.3 Association triangle

Elle est souvent notée D ou Δ . On peut attribuer à E_3 le potentiel de S_2 , puis à E_2 le potentiel de S_1 :



Or $V_{S_3} - V_{E_1} = (V_{S_3} - V_{E_3}) + (V_{S_2} - V_{E_2}) + (V_{S_1} - V_{E_1}) = V_3 + V_2 + V_1 = 0$. On peut donc là encore économiser un quatrième fil conducteur, en reliant S_3 et E_1 , ainsi que les points correspondants dans le triangle formé par les impédances Z .

7.2.2.4 Bilan

Il apparaît que dans le cas d'un système triphasé équilibré, trois conducteurs suffisent pour distribuer l'énergie électrique. De plus, le choix d'un type d'association pour la source (par exemple étoile) n'a pas d'influence sur le choix du type d'association utilisé sur le récepteur. On pourra donc faire des associations mixtes, et réaliser des couplages étoile-triangle ou triangle-étoile.

Au total, un système de distribution triphasé pourra comporter au maximum 5 fils :

- 3 fils dits *de phase*, porteurs du courant ;
- le fil neutre ;
- un fil de terre (ou prise de terre, ou ligne de terre).

Mais une ligne électrique triphasée ne montre souvent que trois conducteurs et non cinq : le fil neutre, quand il est défini (à savoir pour un système branché en triangle) conduit en fait normalement un courant nul dans le cas d'un système équilibré, et peut donc être éliminé. La prise de terre, quant à elle, ne peut et ne *doit pas* être transportée sur de grandes distances, ainsi qu'on l'a indiqué dans le paragraphe 4.5.3 portant sur les parasites. Une ligne électrique transporte donc les trois phases.

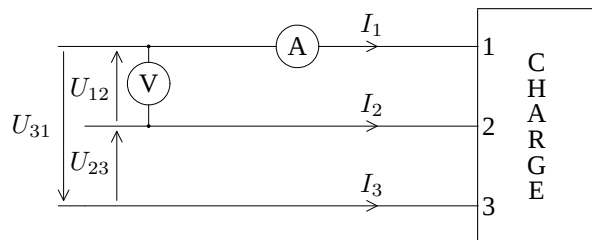
Par ailleurs, une prise électrique montre en général trois emplacements :

- deux phases prises parmi les trois ;
- le neutre qui en France est relié à une terre prise *localement*, et qui permet de fournir la prise de terre aux appareils branchés.

7.2.3 Grandeurs de phase et grandeurs de ligne

7.2.3.1 Définitions

Comme on l'a remarqué au paragraphe précédent, une ligne électrique triphasée ne laisse apparaître, souvent, que trois conducteurs, et son fonctionnement est caractérisé par les grandeurs U et I que l'on peut mesurer :



Les symboles $\textcircled{A}$ et $\textcircled{V}$ désignent respectivement un ampèremètre et un voltmètre, mesurant courant et différence de potentiel. On appellera dans la suite du cours :

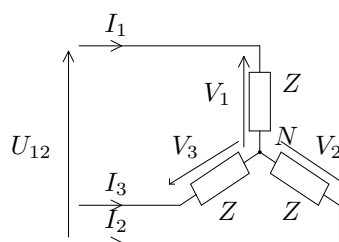
- I la valeur efficace du courant de ligne ;
- U la valeur efficace de la tension de ligne.

Ces grandeurs sont donc définies indépendamment du type d'association de la charge. La charge triphasée (récepteur ou générateur) est constituée de trois dipôles identiques, associés en étoile ou en triangle. Chaque dipôle est parcouru par un courant J et supporte une différence de potentiel à ses bornes V . On appellera donc :

- J la valeur efficace du courant de phase ;
- V la valeur efficace de la tension de phase.

Comme on le voit sur les schémas du paragraphe 7.2.2, les grandeurs de phase et de ligne ne sont pas identiques en général.

7.2.3.2 Relations dans le montage étoile



Les courants de ligne et de phase sont identiques : $I = J$. Par contre U et V sont différents :

$$U_{12}(t) = V_{1N}(t) - V_{2N}(t) = V_1(t) - V_2(t)$$

Si le système des tensions de phase est par exemple direct, $V_2 = a^2 V_1$ d'où

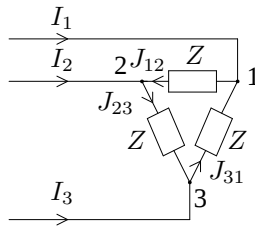
$$U_{12} = V_1(1 - a^2) = V_1 \left(1 + \frac{1}{2} + j\frac{\sqrt{3}}{2} \right) = V_1 \left(\frac{3}{2} + j\frac{\sqrt{3}}{2} \right) = V_1 \sqrt{3} e^{j\frac{\pi}{6}}$$

On en déduit :

- $U = V\sqrt{3}$;
- la tension de ligne est en avance de 30° sur la tension de phase.

Dans le cas de la distribution d'électricité domestique par EDF, la tension de phase vaut 220V entre chaque phase et le neutre, et on retrouve que la tension de ligne vaut $220\sqrt{3} = 380\text{V}$ ^{7.7}.

7.2.3.3 Relations dans le montage triangle



Les tensions de ligne et de phase sont identiques : $U = V$. Par contre I et J sont différents. La loi des nœuds appliquée au point 1 conduit à $I_1 = J_{12} - J_{31}$. Si le système des courants est direct, alors $J_{31} = a J_{12}$ d'où $I_1 = J_{12}(1 - a)$. Le même type de calculs que celui du paragraphe précédent permet d'établir la relation $I_1 = J_{12}\sqrt{3}e^{-j\frac{\pi}{6}}$.

On en déduit :

- $I = J\sqrt{3}$;
- le courant de ligne est en retard de 30° sur le courant de phase.

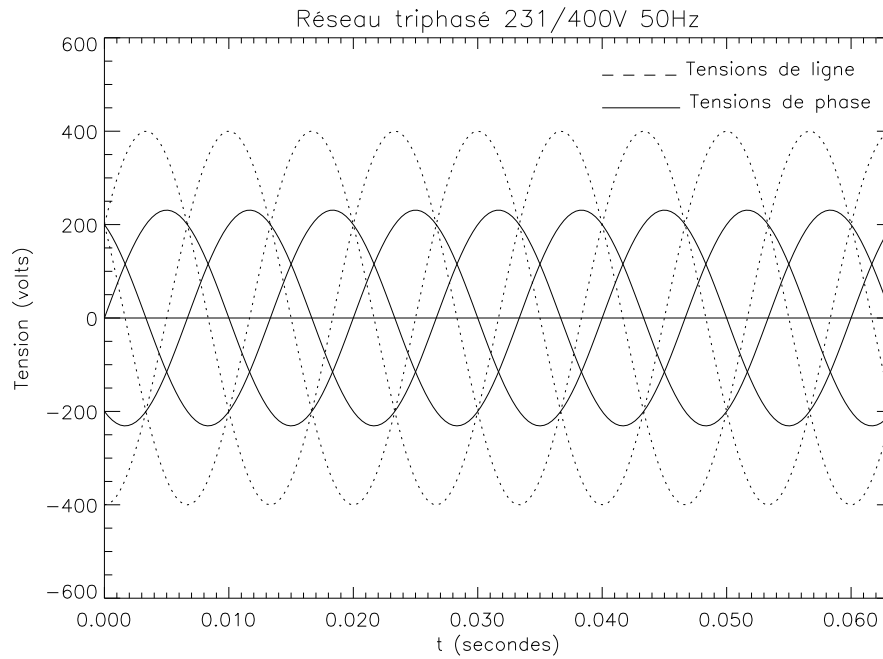
7.2.3.4 Bilan

Pour un système triphasé équilibré :

Montage	Courants	Tensions
étoile	$I = J$	$U = V\sqrt{3}$
triangle	$I = J\sqrt{3}$	$U = V$

Voici un schéma des tensions de ligne et de phase d'un réseau triphasé 231/400V, 50Hz.

7.7. De plus en plus, EDF tend vers un réseau de distribution de 231/400V et non 220/381V.



Ce qu'il faut retenir

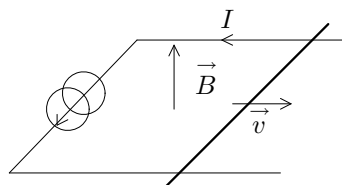
- Les définitions d'un système triphasé : direct, inverse, homopolaire ;
- Les associations étoile/triangle et leurs conséquences sur les valeurs des grandeurs en ligne et en phase ;
- Ce qu'il y a dans une prise électrique (!) : deux phases et un neutre relié à la terre.

7.3 Les machines électriques

7.3.1 Généralités

7.3.1.1 Mouvement d'un conducteur dans un champ d'induction magnétique uniforme

On considère un conducteur rectiligne, placé dans un champ magnétique uniforme $\vec{B}$, et parcouru par un courant constant I , comme dans le schéma ci-dessous.



On montre alors qu'il est soumis à une force (la force de Laplace) qui le met en mouvement selon la direction $\vec{v}$ indiquée.

Symétriquement, quand un conducteur est en mouvement dans un champ magnétique $\vec{B}$, un courant dit *induit* circule.

7.3.1.2 Le théorème de Ferraris

Trois bobinages identiques, décalés dans l'espace de 120° , sont alimentés par trois courants J_1, J_2, J_3 qui forment un système triphasé équilibré (direct, par exemple). Chaque bobinage crée en O un champ de direction fixe et de mesure variable, que l'on peut représenter par :

$$\begin{aligned} B_1(0) &= B_{max} e^{j0} \left(\frac{e^{j\omega t} + e^{-j\omega t}}{2} \right) \\ B_2(0) &= B_{max} e^{j\frac{2\pi}{3}} \left(\frac{e^{j(\omega t - \frac{2\pi}{3})} + e^{-j(\omega t - \frac{2\pi}{3})}}{2} \right) \\ B_3(0) &= B_{max} e^{j\frac{4\pi}{3}} \left(\frac{e^{j(\omega t - \frac{4\pi}{3})} + e^{-j(\omega t - \frac{4\pi}{3})}}{2} \right) \end{aligned}$$

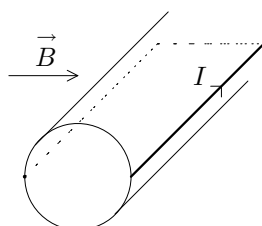
La somme de ces trois composantes est $\frac{3B_{max}}{2} e^{j\omega t}$, qui représente un vecteur de norme constante, tournant à la vitesse angulaire ω . On en déduit le théorème :

Trois bobinages identiques, décalés régulièrement dans l'espace, et parcourus par des courants triphasés équilibrés (par exemple directs), créent un champ tournant à la vitesse angulaire ω (dans le sens direct), de norme constante égale à $\frac{3B_{max}}{2}$. Ce champ est sur l'axe d'un bobinage lorsque le courant dans ce bobinage est maximum.

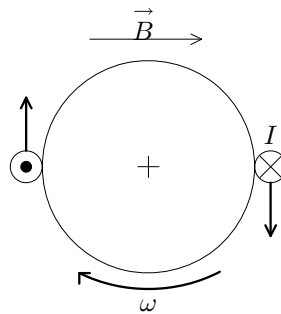
7.3.2 La machine à courant continu (MCC)

7.3.2.1 Principe de la machine

On considère maintenant un conducteur faisant le tour d'une pièce cylindrique placée dans un champ $\vec{B}$ uniforme :



Ce système peut se représenter de face comme suit, le sens positif de I étant donné par la convention rappelée ci-dessus^{7.8} :



Ce conducteur est soumis de part et d'autre de la pièce cylindrique à une force tangentielle. Ces deux forces opposées créent un couple qui fait tourner l'ensemble.

Pendant, alimenter un tel conducteur en rotation est impossible sans dispositif spécial : les câbles d'alimentation seraient en quelques secondes vrillés et casseraient. Il a donc fallu imaginer un système de « récupération » du courant électrique, les *balais* ; l'ensemble des balais et des conducteurs formant le *collecteur*. Il s'agit de patins, frottant de part et d'autre du rotor^{7.9}, et en contact successivement avec les différentes parties du bobinage :

Par symétrie, si on suppose maintenant que l'on place le même système dans un champ magnétique $\vec{B}$, et qu'on le fait tourner autour de son axe, le conducteur sera le siège d'un courant induit.

Dans le premier cas, le système fonctionne en *moteur* (on fournit de la puissance électrique par le courant I) ; dans le second, il fonctionne en *générateur* (on fournit de la puissance mécanique en faisant tourner l'arbre).

7.3.2.2 Réalisation

Dans la pratique, le champ $\vec{B}$ est créé à partir de bobinages fixés sur une partie immobile, le *stator*, appelé l'*inducteur*, et parcourus par un courant I_e continu^{7.10}. Le conducteur placé sur la partie mobile, appelée le *rotor*, bobiné, forme l'*induit*^{7.11}.

En fonctionnement générateur, le *courant d'induit*, créé par le champ $\vec{B}$, est loin d'être constant dans le conducteur. Pour améliorer la situation, on agit sur la forme de la partie entre le stator et le rotor, que l'on appelle l'*entrefer*, qui conditionne le profil du champ magnétique inducteur.

7.3.2.3 Modèle

On peut établir principalement deux équations liant les paramètres de la machine à courant continu. Si Φ désigne le flux créé dans l'inducteur, flux proportionnel au courant d'inducteur, Ω la vitesse de rotation de l'arbre et I le courant d'induit et si on pose de plus que :

- E est la force électromotrice en fonctionnement générateur, ou la force contre-électromotrice en fonctionnement moteur ;

7.8. Rappel : conventions de représentation des vecteurs dont la direction se confond avec l'axe de visée. On représente par $\otimes$ un vecteur « fuyant » l'observateur, et par $\odot$ un vecteur se dirigeant vers lui.

7.9. Le rotor est la partie tournante.

7.10. En électrotechnique, un courant est dit *alternatif* quand il est périodique et à moyenne nulle, et *continu* quand il ne change jamais de signe. Une conséquence est qu'un courant continu n'est pas forcément constant !

7.11. On appelle couramment le courant circulant dans l'inducteur, le *courant d'inducteur*, et le courant circulant dans l'induit le... *courant d'induit*.

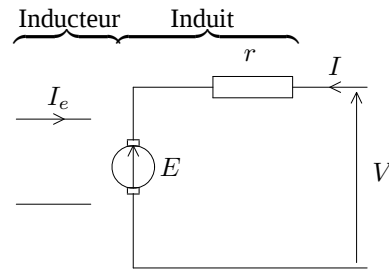
- C le couple sur l'arbre (ie la partie tournante, cylindrique, constituée notamment des accouplements des rotors des machines présentes)...

... on obtient :

$$E = k\Phi\Omega \text{ et } C = k\Phi I$$

k étant un coefficient de proportionnalité (le même dans les deux expressions).

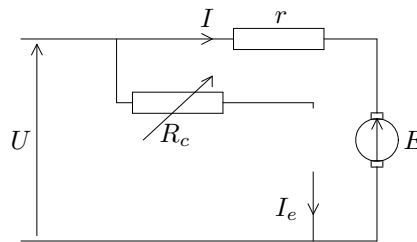
On modélise alors la MCC par le schéma suivant, r désignant la résistance du bobinage de l'induit :



Les deux circuits (inducteur et induit) étant tous les deux alimentés en continu, on peut envisager de les alimenter avec un seul générateur, soit en parallèle, soit en série.

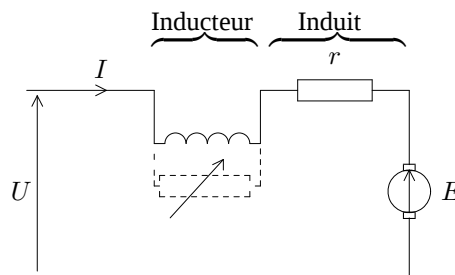
7.3.2.4 Excitation parallèle, excitation série

1. **Excitation parallèle** : aussi appelée machine *shunt* ou à *excitation indépendante*. Les deux circuits sont câblés selon le schéma suivant :



Les courants I et I_e sont indépendants. Le rhéostat ^{7.12} R_c permet de régler I_e et donc le flux dans l'inducteur.

2. **Excitation série** : aussi appelée *moteur universel*. Les deux circuits sont câblés selon le schéma suivant :



Les courants I et I_e ne sont plus indépendants, mais il est encore possible de régler le flux en ajoutant un rhéostat en parallèle sur le bobinage inducteur.

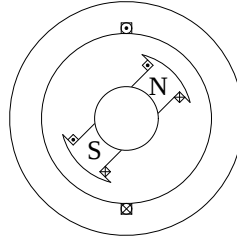
Ce branchement présente en outre une propriété remarquable : on peut montrer que le sens de rotation de l'arbre ne change pas quand la tension de commande U change de signe. Cette machine peut donc être également alimentée en alternatif, d'où son nom de *moteur universel*. Néanmoins, dans ce dernier mode de fonctionnement, son rendement chute. Son faible coût de fabrication comparé aux « vraies » machines alternatives la rend cependant attractive, et largement répandue quand il s'agit d'utiliser un petit moteur électrique : moulin à café, perceuse électrique, etc.

7.12. Résistance variable.

7.3.3 La machine synchrone

On parle aussi d'alternateur.

Cette fois-ci, le rotor porte l'inducteur et le stator porte l'induit. Considérons le schéma suivant :



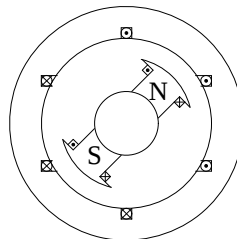
On a représenté sur ce schéma une situation où l'induit ne porte qu'une spire. L'inducteur (aimant permanent ou bobine alimentée par un courant continu I_e) crée un champ d'induction magnétique dans l'entrefer, champ dont on a représenté sur le schéma précédent les pôles nord et sud. Un flux Φ est donc lui aussi induit et son mouvement de rotation à la vitesse angulaire ω crée dans la spire de l'induit une f.é.m (plus ou moins approximativement) sinusoïdale. On peut :

- soit multiplier le nombre de spires dans chaque encoche de l'induit, et les placer en série : on augmente alors l'intensité du courant induit recueilli ;
- soit créer d'autres paires d'encoches régulièrement réparties sur le stator. Si on utilise trois paires d'encoches réparties tous les 120° , on créera un système triphasé équilibré en courant, pour peu que ces trois bobinages alimentent des impédances de charge identiques.

Symétriquement, si on alimente l'inducteur par un courant continu, et les bobinages d'induit par un système triphasé équilibré, la machine fonctionnera en moteur.

7.3.4 La machine asynchrone

Reprenons le schéma précédent, et supposons que les spires du stator sont alimentées par un système triphasé équilibré.



Le champ tournant ainsi créé entre le stator et le rotor (cf. théorème de Ferraris, paragraphe 7.3.1.2) induit des courants dans les bobinages du rotor. Si d'autre part le rotor tourne (car par exemple une autre machine est montée sur le même arbre, fonctionnant en moteur), les spires portées par le rotor vont voir un champ tournant résultant de la combinaison :

- du champ tournant créé par les spires du stator ;
- de la rotation du rotor.

Cette combinaison de deux rotations entraîne que le courant approximativement sinusoïdal créé dans chacune des spires du rotor n'est *pas* synchrone avec le courant dans les spires du stator : ils n'ont pas la même fréquence, d'où le nom de machine *asynchrone*.

En résumé :

- on alimente la machine par un système triphasé équilibré sur le stator ;
- fonctionnement générateur :
 - on fait tourner le rotor ;
 - des courants induits apparaissent dans les spires du rotor, de fréquence différente de celle des courants du stator.
- fonctionnement moteur :
 - on alimente le rotor par un système triphasé équilibré de fréquence différente de celle présente au stator ;
 - le rotor tourne à une vitesse angulaire correspondant à la différence des pulsations de ces deux systèmes de courants ;
 - si le système triphasé « injecté » au rotor a la même fréquence que celui du stator, le rotor est immobile.

Ce qu'il faut retenir

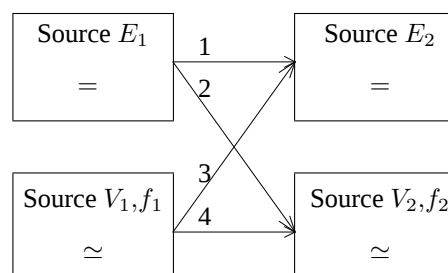
- Le théorème de Ferraris : trois bobinages, séparés de 120° , dans lesquels circulent des courants formant un système triphasé équilibré direct ou inverse créent un champ magnétique tournant ;
- Les principes de fonctionnement des machines à courant continu, machines synchrones et machines asynchrones.

7.4 Conversion d'énergie

7.4.1 Introduction

Le réseau de distribution électrique fournit un courant à 50Hz. Or on a besoin ou de courant continu, ou de courant alternatif mais à une autre fréquence, pour alimenter par exemple les machines dont nous avons dressé un rapide panorama dans les paragraphes précédents. Il est nécessaire de pouvoir disposer d'outils permettant de faire les transformations entre réseaux continu et alternatif, mais également d'alternatif à alternatif (changement de fréquence) ou de continu à continu.

On peut dénombrer quatre possibilités de conversion^{7.13} :



7.13. Par convention, les symboles $=$ et $\simeq$ désignent respectivement un réseau continu et un réseau alternatif.

- 1 : conversion continu/continu : changement de valeur ($E_1 \neq E_2$) : hacheur ;
- 2 : conversion continu/alternatif : onduleur ;
- 3 : conversion alternatif/continu : redresseur ;
- 4 : on distingue :
 - changement de valeur efficace : gradateur ;
 - changement de fréquence : cycloconvertisseur.

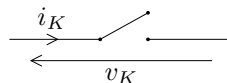
Ces cinq dispositifs sont des *convertisseurs statiques*. Nous nous limiterons à un survol de deux d'entre eux (redresseur et onduleur).

7.4.2 Les interrupteurs

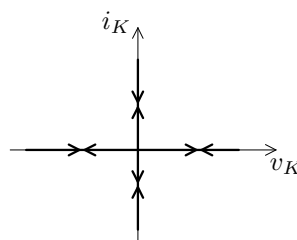
Un convertisseur statique est principalement constitué d'interrupteurs pour « aiguiller » le courant. Plutôt que des interrupteurs mécaniques, des relais, on préfère utiliser des interrupteurs électroniques, moins encombrants, rapidement maniables et résistants.

7.4.2.1 Principe de fonctionnement

On classe les interrupteurs en utilisant un diagramme sur lequel est tracée leur caractéristique. Lorsqu'un interrupteur est *fermé* (on dit aussi *passant*), il laisse passer le courant et la tension à ses bornes est nulle ; lorsqu'il est *ouvert* (ou *bloqué*), il ne laisse passer aucun courant et la tension à ses bornes est non nulle. Un interrupteur idéal laisse passer tous les courants, positifs ou négatifs, lorsqu'il est fermé, et accepte toute tension à ses bornes quand il est ouvert. En notant i_K et v_K respectivement le courant qui le traverse et la tension à ses bornes,



sa caractéristique est la suivante :



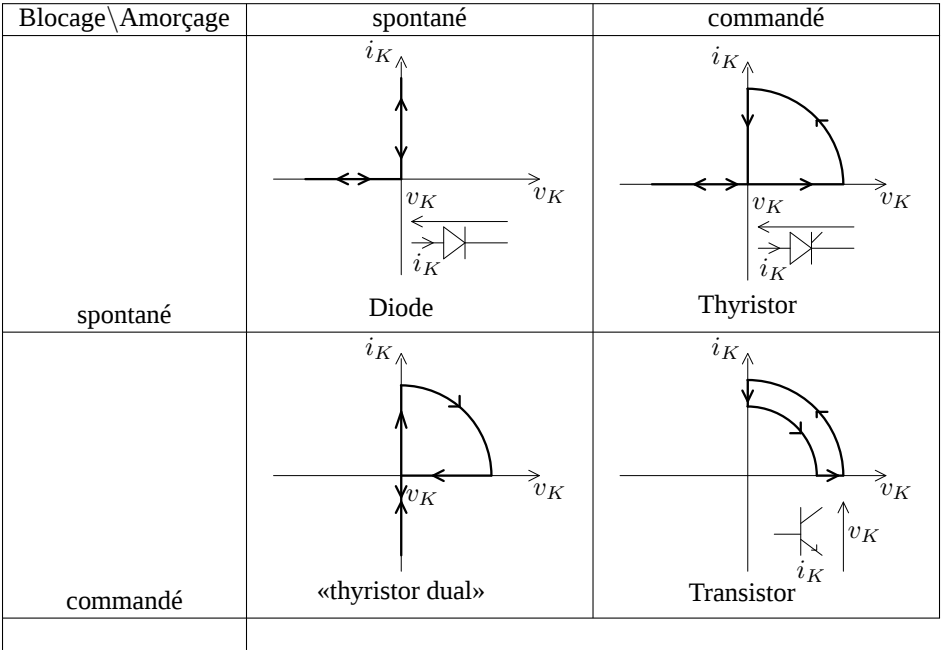
Elle comporte 4 « branches ». Chacun des quatre secteurs délimités par les axes des abscisses et des ordonnées est appelé *quadrant*. Il existe des interrupteurs :

- à 2, 3 ou 4 branches ;
- dont la fermeture (ou amorçage) ou l'ouverture (ou blocage) peut être commandée par l'opérateur ou non ; si elle n'est pas commandée, on parle de fonctionnement spontané. Le changement d'état d'un interrupteur est appelée *commutation*.

Dans le cas d'une commutation spontanée, l'ensemble du circuit environnant déclenche l'ouverture ou la fermeture, parce qu'un courant ou une tension change de signe par exemple. Lorsqu'il s'agit d'une commutation commandée, la commande se fait par l'envoi d'une impulsion sur une borne de l'interrupteur.

7.4.2.2 Les types d'interrupteurs

Dans le diagramme ci-dessous, l'envoi d'une commande se représente par un arc de cercle dans le quadrant correspondant.

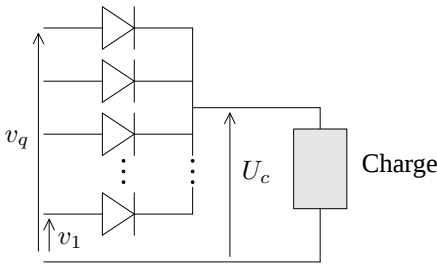


Le « thyristor dual » est un composant de synthèse, fabriqué à l'aide d'une diode et d'un transistor.

7.4.3 Le redressement

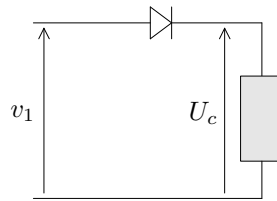
7.4.3.1 Montages à diodes

1. **Montage demi-onde ou simple alternance :** On s'intéresse au cas général à q phases. Chacune de ces phases est portée à un potentiel sinusoïdal de valeur efficace V : pour k variant de 1 à q , $v_k(t) = V\sqrt{2} \sin\left(\omega t - k\frac{2\pi}{q}\right)$. On monte alors une diode par phase, et les cathodes de ces q diodes sont reliées à la charge :

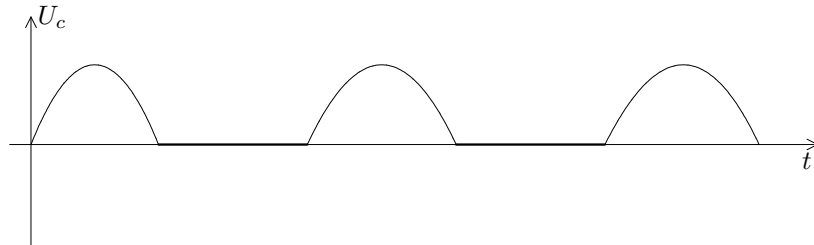


Dans ce dispositif, une seule diode conduit à un instant donné (celle dont le potentiel d'anode v_k est le plus élevé). Les autres sont bloquées. On dit que ce montage est à *cathode commune*.

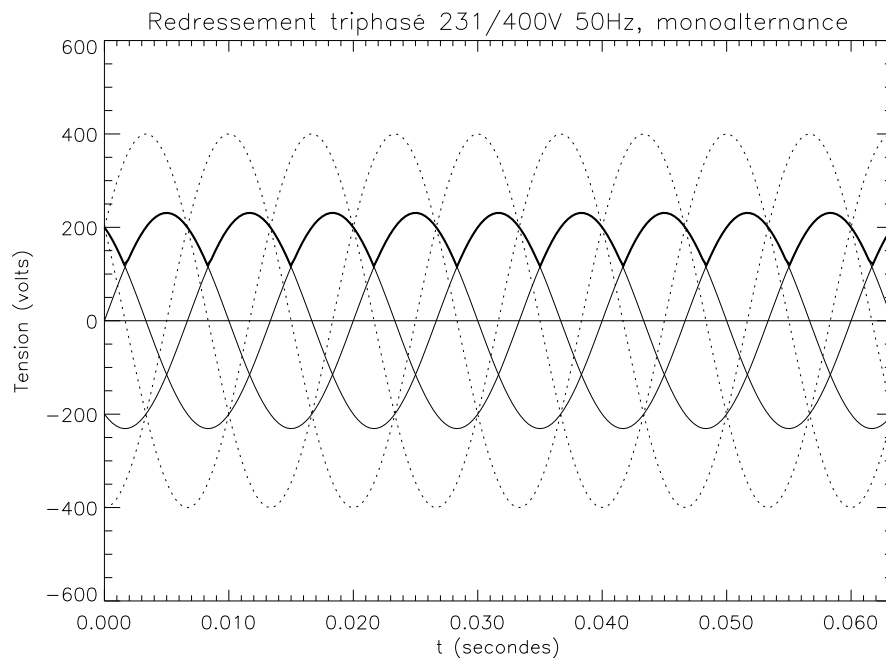
Pour une phase unique ($q=1$), le schéma est le suivant :



On obtient sur la charge la tension U_c :



On trouvera ci-après un schéma pour un système triphasé.

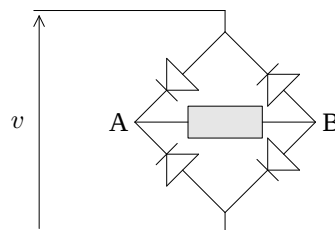


On montre que la valeur moyenne de la tension sur la charge vaut dans le cas général d'un système à q phases :

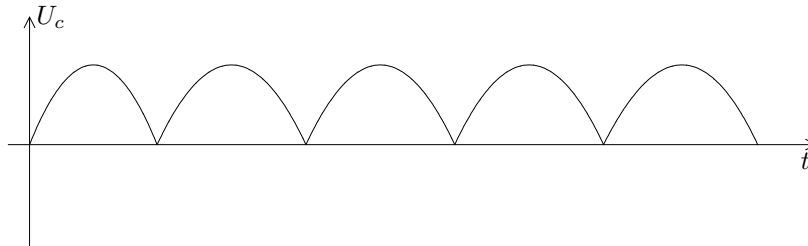
$$\langle u_c \rangle = \frac{q}{\pi} V \sqrt{2} \sin \frac{\pi}{q}$$

On peut obtenir une tension continue négative en inversant le sens des diodes, qui présentent dans ce cas leur anode à la charge.

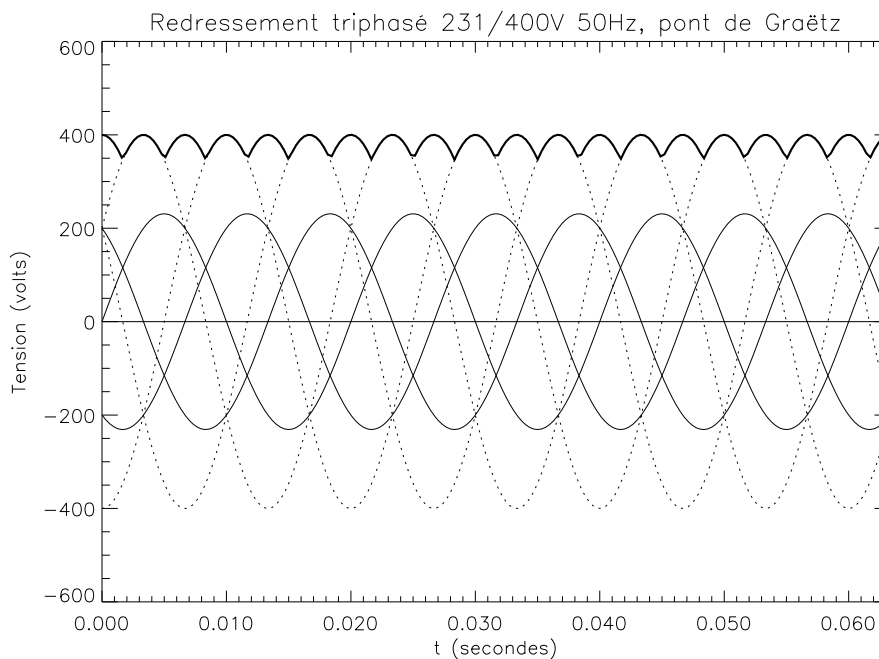
2. **Montage en pont de Graëtz :** prenons l'exemple du redresseur monophasé en pont de Graëtz, conforme au schéma ci-dessous :



Il s'agit en fait de la combinaison de deux redresseurs demi-onde : un à cathode commune, dont le point commun est A (cf. le schéma donné en exemple du redresseur demi-onde), et un à anode commune, dont le point commun est B. La différence de potentiel aux bornes de la charge, $V_A - V_B$, se calcule en remarquant que V_A est donné par un montage à cathode commune, et V_B par un montage à anode commune. La tension redressée aux bornes de la charge est donc la suivante :



On trouvera ci-après un schéma pour un système triphasé.



On montre que la valeur moyenne de la tension sur la charge vaut dans le cas général d'un système à q phases :

$$\langle u_c \rangle = 2 \frac{q}{\pi} V \sqrt{2} \sin \frac{\pi}{q}$$

Il existe bien sûr d'autres types de redresseurs pleine onde (redresseur avec transformateur « à point milieu » notamment)...

7.4.3.2 Montage à thyristors

Il s'agit de remplacer les diodes des circuits précédents par des thyristors, qui présentent la caractéristique de devoir être commandés à l'amorçage *et* au blocage (cf. paragraphe 7.4.2). En fait, on choisit de les commander avec un retard δt par rapport à l'instant où ils commuteraient « naturellement » s'ils étaient des diodes. On définit l'angle de retard ψ par $\psi = \omega \delta t$, où $\omega = 2\pi \cdot (50\text{Hz}) = 100\pi \text{ rad/s}$ dans le cas d'un redresseur fonctionnant à partir du réseau EDF.

Dans le cas d'un montage à pont de Graëtz à thyristors, on montre que la valeur de la tension redressée vaut :

$$\langle u_c \rangle = 2 \frac{q}{\pi} V \sqrt{2} \sin \frac{\pi}{q} \cos \psi$$

Ce montage permet le réglage de la valeur de la tension redressée par l'intermédiaire du retard ψ .

Remarque : si l'angle ψ devient supérieur à $\pi/2$, la valeur moyenne de la tension redressée devient négative, ce qui se traduit par le fait que le transfert de puissance se fait de la « charge » continue au réseau alternatif. On parle alors d'*onduleur assisté par le réseau*. On effectue en fait l'opération inverse du redressement : voir paragraphe suivant.

7.4.4 L'ondulation

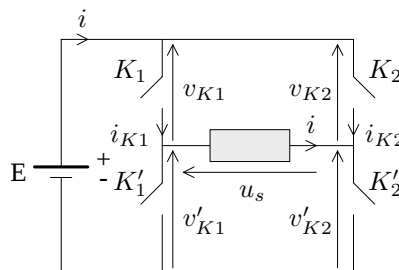
7.4.4.1 Généralités

Le principe est l'inverse de celui du redresseur : il s'agit cette fois-ci, à partir d'une source continue, de fabriquer une source alternative, toujours en se servant d'interrupteurs. En règle générale, les « formes d'onde » (ie les formes des tensions de sortie) ne sont pas sinusoïdales, et il est nécessaire de les filtrer par un filtre passe-bas en aval.

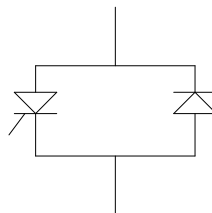
7.4.4.2 Exemple d'onduleur

On va étudier un peu plus précisément l'*onduleur monophasé en pont*.

1. **Principe :** On actionne deux paires d'interrupteurs $K_1 K'_1$ et $K_2 K'_2$ de telle manière que leurs commandes soient « complémentaires » : il est indispensable d'éviter de fermer K_1 et K'_1 en même temps, par exemple, car dans ce cas on court-circuiterait la source continue, notée par la double barre sur la gauche du schéma :

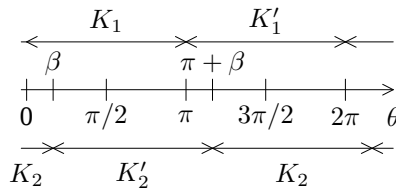


Chaque interrupteur est en fait composé d'une diode et d'un thyristor placés en parallèle :



2. **Détermination des formes d'onde :** Il est interdit de fermer en même temps K_i et K'_i , mais on peut introduire un décalage dans les fermetures des interrupteurs. Si on représente une période de fonctionnement par sa « durée angulaire », c'est-à-dire une période de 2π , on peut choisir de décaler la commande d'un angle β , autrement dit (les flèches indiquent les intervalles de temps au cours desquels l'interrupteur concerné est fermé ; il est ouvert

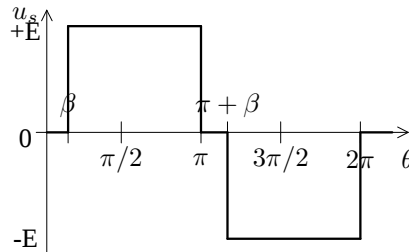
le reste du temps) :



Quatre cas sont à examiner, cas qui vont se répéter sur chaque période :

- $0 < \theta < \beta$: $v'_{K1} = v'_{K2} = E$ donc $u_s = 0$;
- $\beta < \theta < \pi$: $v'_{K1} = v_{K2} = E$ donc $u_s = E$;
- $\pi < \theta < \pi + \beta$: $v_{K1} = v_{K2} = E$ donc $u_s = 0$;
- $\pi + \beta < \theta < 2\pi$: $v_{K1} = v'_{K2} = E$ donc $u_s = -E$;

On en déduit la forme d'onde suivante pour la tension u_s :



On peut régler l'angle de commande β de manière à ce que cette courbe se rapproche le plus d'une sinusoïde. On dit que β permet de « régler le taux d'harmoniques »^{7.14}.

Comme dans le cas des redresseurs, il existe beaucoup de types d'onduleurs. En règle générale, on peut dire que multiplier le nombre d'interrupteurs commandables peut améliorer la forme d'onde, et donc la facilité avec laquelle on peut filtrer le signal en aval pour le rendre plus « sinusoïdal ». Evidemment, les améliorations se font au détriment de la simplicité et de la robustesse de la commande. Un type particulier d'ondeur est l'« ondeur à modulation de largeur d'impulsions » (ou ondeur à MLI). Pour améliorer le taux d'harmoniques, on choisit dans ce dernier cas d'agir plutôt sur la commande des interrupteurs que sur leur nombre...

Ce qu'il faut retenir

- Les types d'interrupteurs : diode, thyristor, thyristor dual, transistor et leurs caractéristiques ;
- Ce qu'est le redressement ; le pont de Graëtz à diodes et à thyristors ;
- Ce qu'est l'ondulation de courant : principe, détermination des formes d'onde.

7.14. D'après le paragraphe 1.2.3.2 sur le spectre d'un signal périodique, ce qui est le cas de la tension ici, on peut définir le taux d'harmoniques τ si le signal est de plus pair ou impair. C'est également le cas ici si on procède artificiellement à un décalage de l'origine des temps de $\beta/2$: le signal ainsi obtenu est impair. Le taux d'harmoniques que l'on peut alors calculer vaut

$$\tau = \frac{\sqrt{1 - \frac{\beta}{\pi} - \frac{8}{\pi^2} (\cos \beta/2)^2}}{\frac{2\sqrt{2}}{\pi} \cos \beta/2}$$

Annexe A

Table de transformées de Fourier usuelles

A.1 Définitions

Pour un signal à temps continu $x(t)$, tel que $\int_{-\infty}^{+\infty} |x(t)|^2 dt$ converge, on définit sa *transformée de Fourier* $X(\nu)$ ^{A.1} par :

$$X(\nu) = \int_{-\infty}^{\infty} x(t) e^{-j2\pi\nu t} dt$$

ou avec $\omega = 2\pi\nu$:

$$X(j\omega) = \int_{-\infty}^{\infty} x(t) e^{-j\omega t} dt$$

On définit également son *spectre* comme étant le module de sa transformée de Fourier..

$$S_x(\nu) = |X(\nu)|$$

... et sa *densité spectrale de puissance* comme étant le carré du spectre :

$$D_x(\nu) = S_x(\nu)^2 = |X(\nu)|^2$$

A.1. L'opérateur qui à $x(t)$ associe $X(\nu)$ est la Transformée de Fourier, avec une majuscule. Sans majuscule, on parle de $X(\nu)$ lui-même.

A.2 Table

Original	Image	Remarques
$\lambda x(t) + \mu y(t)$	$\lambda X(\nu) + \mu Y(\nu)$	
$x(t - t_0)$	$e^{-j2\pi\nu t_0} X(\nu)$	Décalage en temps avec t_0 réel > 0
$e^{+j2\pi\nu_0 t} x(t)$	$X(\nu - \nu_0)$	Décalage en fréquence avec ν_0 réel > 0
$x(\lambda t)$	$X(\nu/\lambda)/\lambda$	Dilatation en temps avec λ réel > 0
$x(\lambda t)$	$X^*(\nu/\lambda)/\lambda$	λ réel < 0
$x'(t)$	$j2\pi\nu X(\nu)$	Dérivation
$\int_0^t x(\theta) d\theta$	$X(\nu)/j2\pi\nu$	Intégration
$x^*(t)$	$X^*(-\nu)$	Conjugaison complexe
$(x * y)(t)$	$X(\nu)Y(\nu)$	Convolution en temps
$x(t)y(t)$	$(X * Y)(\nu)$	Convolution en fréquence
$\delta(t)$	1	Impulsion de Dirac
$\cos 2\pi\nu_0 t$	$[\delta(\nu_0) + \delta(-\nu_0)]/2$	ν_0 réel
$\sin 2\pi\nu_0 t$	$[\delta(\nu_0) - \delta(-\nu_0)]/2j$	ν_0 réel
$e^{j2\pi\nu_0 t}$	$\delta(\nu_0)$	ν_0 réel
$x(t) = 1$ si $ t < t_0/2$ $= 0$ sinon	$t_0 \text{sinc}(\nu t_0)$	- Fonction porte - $\text{sinc}(x) = \sin(\pi x)/\pi x$

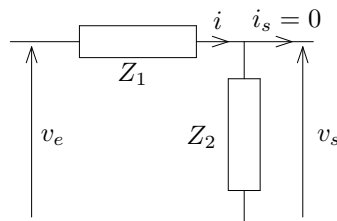
Annexe B

Quelques théorèmes généraux de l'électricité

B.1 Diviseur de tension, diviseur de courant

B.1.1 Diviseur de tension

On considère la situation suivante :



On a :

$$\begin{cases} v_e = (Z_1 + Z_2)i \\ v_s = Z_2 i \end{cases}$$

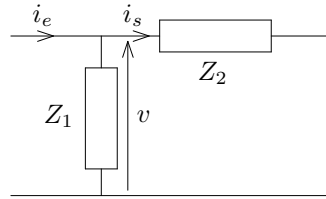
On en déduit la relation du *diviseur de tension* :

$$v_s = \frac{Z_2}{Z_1 + Z_2} v_e$$

On notera avec soin que le courant i_s doit être *nul*.

B.1.2 Diviseur de courant

On considère la situation suivante :



On a :

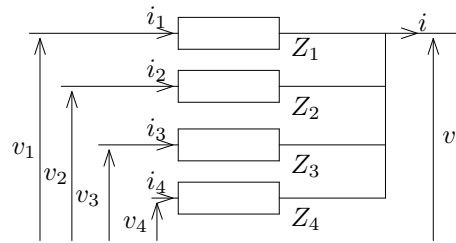
$$\begin{cases} i_e = \frac{v}{Z_1} + i_s \\ i_s = \frac{v}{Z_2} \end{cases}$$

On en déduit la relation du *diviseur de courant* :

$$i_s = \frac{Z_1}{Z_1 + Z_2} i_e$$

B.2 Théorème de Millman

On considère la situation suivante :



... avec la condition supplémentaire que $i = 0$. On a donc :

$$\begin{cases} v_s &= v_1 - Z_1 i_1 \\ &= v_2 - Z_2 i_2 \\ &= v_3 - Z_3 i_3 \\ &= v_4 - Z_4 i_4 \end{cases}$$

D'où :

$$\begin{cases} i_1 &= (v_1 - v_s)/Z_1 \\ i_2 &= (v_2 - v_s)/Z_2 \\ i_3 &= (v_3 - v_s)/Z_3 \\ i_4 &= (v_4 - v_s)/Z_4 \end{cases}$$

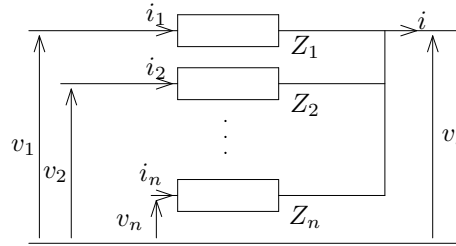
Et d'autre part :

$$i_1 + i_2 + i_3 + i_4 = 0$$

On obtient donc la relation suivante :

$$\left(\frac{1}{Z_1} + \frac{1}{Z_2} + \frac{1}{Z_3} + \frac{1}{Z_4} \right) v_s = \left(\frac{v_1}{Z_1} + \frac{v_2}{Z_2} + \frac{v_3}{Z_3} + \frac{v_4}{Z_4} \right)$$

En règle générale, si on considère le schéma suivant :



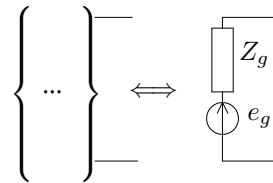
avec la condition supplémentaire $i = 0$, il vient :

$$v_s \sum_{k=1}^n \frac{1}{Z_k} = \sum_{k=1}^n \frac{v_k}{Z_k}$$

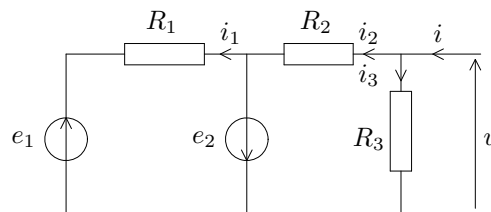
B.3 Théorèmes de Thévenin et Norton

B.3.1 Théorème de Thévenin

Toute association de dipôles linéaires, actifs ou non, peut se modéliser sous la forme de la mise en série d'une source de tension et d'une impédance :



Par exemple, considérons le circuit suivant :



où l'on veut exprimer u en fonction de i et d'éventuelles forces électromotrices. On a :

$$\begin{cases} u = R_3 i_3 \\ u = R_2 i_2 - e_2 \\ i = i_2 + i_3 \end{cases}$$

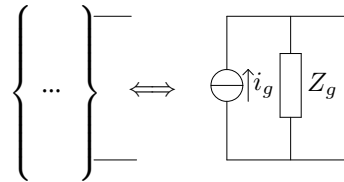
Ce système équivaut à :

$$\begin{cases} \frac{u}{R_3} = i_3 \\ \frac{u}{R_2} + \frac{e_2}{R_2} = i_2 \\ i = i_2 + i_3 \end{cases}$$

On en déduit facilement $u \left(\frac{1}{R_2} + \frac{1}{R_3} \right) + \frac{e_2}{R_2} = i$. On peut alors exprimer u sous la forme $u = e_g + Z_g i$ en posant $Z_g = R_2 R_3 / (R_2 + R_3)$ et $e_g = -R_3 / (R_2 + R_3) e_2$.

B.3.2 Théorème de Norton

Toute association de dipôles linéaires, actifs ou non, peut se modéliser sous la forme de la mise en parallèle d'une source de courant et d'une impédance :



B.3.3 Relation entre les deux théorèmes

On peut facilement passer d'une représentation à l'autre avec la relation aisément démontrable :

$$e_g = Z_g i_g$$

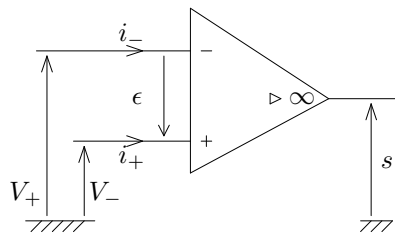
Annexe C

L'Amplificateur Opérationnel (AO)

L'amplificateur opérationnel est un amplificateur de tension. Il peut fonctionner en deux modes : linéaire et non-linéaire.

C.1 L'AO idéal en fonctionnement linéaire

C.1.1 Représentation



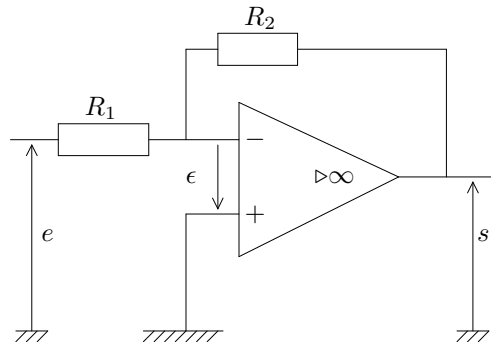
C.1.2 Caractéristiques

L'entrée « - » est appelée *entrée inverseuse*, l'entrée « + » est l'*entrée non inverseuse*. L'amplificateur opérationnel idéal vérifie les conditions suivantes :

- $\epsilon = 0$: quelle que soit la sortie s , la différence de potentiel en entrée est nulle : *le gain de l'AO idéal est infini* ;
- $i_+ = i_- = 0$: les courants d'entrée sur chacune des bornes de l'AO sont nuls : cela revient à dire que *l'impédance d'entrée de l'AO idéal est infinie*.

C.1.3 Exemple : montage amplificateur

Schéma du montage :



On a $\epsilon = 0$, et $i_+ = i_- = 0$ A. En appliquant le théorème de Millman (cf. Annexe B, paragraphe B.2), comme $V_+ = V_- = 0$ V, il vient :

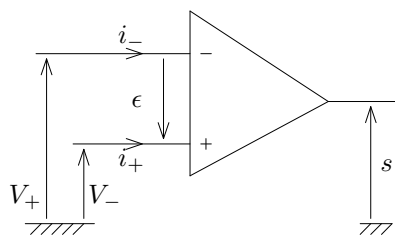
$$0 = \frac{e}{R_1} + \frac{s}{R_2}$$

D'où :

$$\frac{s}{e} = -\frac{R_2}{R_1}$$

C.2 L'AO non idéal en fonctionnement linéaire

C.2.1 Représentation



C.2.2 Caractéristiques

Plusieurs défauts peuvent s'ajouter à l'AO idéal ; dans l'ordre d'importance en général on compte :

1. Gain non infini : le gain de l'AO vaut A_0 ; cela se traduit par $s = A_0 \epsilon$ avec bien sûr $\epsilon \neq 0$;
2. Impédance d'entrée non infinie : les courants i_- et i_+ ne sont pas nuls, et $\epsilon = Z_e(i_- - i_+)$;
3. La réponse en fréquence n'est pas parfaite : la fonction de transfert de l'AO est celle d'un filtre passe-bas : $H(j\omega) = A_0 / (1 + j\frac{\omega}{\omega_0})$.

C.2.3 Exemples : montage amplificateur

Dans les trois cas, le montage est le même que celui du paragraphe C.1.3.

C.2.3.1 Gain non infini

On a $i_+ = i_- = 0$. En appliquant le théorème de Millman, il vient :

$$V_- = \frac{e}{R_1} + \frac{s}{R_2}$$

D'autre part $\epsilon = V_+ - V_-$ et $s = A_0 \epsilon$. On obtient donc :

$$s = -A_0 \left(\frac{e}{R_1} + \frac{s}{R_2} \right)$$

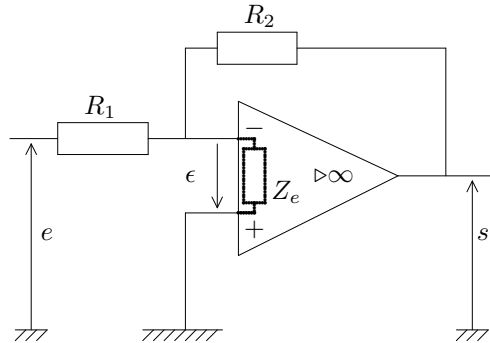
D'où :

$$\boxed{\frac{s}{e} = -\frac{R_2}{R_1} \frac{A_0}{A_0 + (1 + \frac{R_2}{R_1})}} \quad (C.1)$$

On vérifie que dans le cas où $A_0 \rightarrow +\infty$, on retrouve le cas du paragraphe C.1.3.

C.2.3.2 Impédance d'entrée non infinie

Le schéma équivalent du montage est le suivant :



En appliquant le théorème de Millman à l'entrée inverseuse, il vient :

$$-\epsilon \left(\frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{Z_e} \right) = \frac{e}{R_1} + \frac{s}{R_2} + \frac{0}{Z_e}$$

Avec d'autre part $s = A_0 \epsilon$ on obtient :

$$\boxed{\frac{s}{e} = -\frac{R_2}{R_1} \frac{A_0}{A_0 + (1 + \frac{R_2}{R_1} + \frac{R_2}{Z_e})}}$$

Supposons par exemple que $Z_e = R_e // C_e$, autrement dit :

$$Z_e = \frac{R_e}{1 + jR_e C_e \omega}$$

Il est facile alors de montrer que la fonction de transfert du système est celle d'un filtre passe-bas, de pulsation de coupure

$$\omega_c = \frac{R_2 + R_e(A_0 + \frac{R_2}{R_1} + 1)}{R_2 R_e C_e}$$

En pratique, R_e et A_0 sont grands respectivement devant R_2 et 1, donc

$$\omega_c \approx \frac{A_0}{R_2 C_e}$$

Exemple : $R_2 = 10 \text{ k}\Omega$, $A_0 = 1000$ et $C_e = 10 \text{ nF}$ donne $\omega_c \approx 10^7 \text{ rad/s}$ d'où une fréquence de coupure de l'ordre de 1 MHz.

C.2.3.3 Réponse en fréquence imparfaite

Supposons cette fois-ci (avec une impédance d'entrée infinie) que le gain $\frac{s}{e}$ de l'AO vaut

$$\frac{A_0}{1 + j\omega/\omega_0}$$

On peut alors reprendre directement la formule C.1 en

$$\frac{s}{e} = -\frac{R_2}{R_1} \frac{A_0}{A_0 + (1 + \frac{R_2}{R_1})(1 + j\omega/\omega_0)}$$

On peut alors montrer que la fonction de transfert est un filtre du premier ordre, de pulsation de coupure

$$\omega_1 = \frac{A_0 + 1 + \frac{R_2}{R_1}}{1 + \frac{R_2}{R_1}} \omega_0$$

Souvent, $A_0 \gg (1 + \frac{R_2}{R_1})$ et donc la pulsation de coupure vaut

$$\omega_1 \approx \frac{A_0}{1 + \frac{R_2}{R_1}} \omega_0$$

Le produit $A_0 \omega_0$, ou plutôt le produit $A_0 f_0$ est une *caractéristique* de l'AO, et est appelé *produit gain-bande* de l'amplificateur. On note que pour un montage amplificateur, la fréquence de coupure est égale au produit gain-bande divisé par l'amplification : un « bon » AO a donc un produit gain-bande aussi grand que possible, afin de ne pas introduire de limitation dans ce genre de montage.

C.3 L'AO en fonctionnement non linéaire

Ce n'est pas le mode de fonctionnement nominal de l'amplificateur opérationnel. Il se distingue du fonctionnement linéaire par le fait que la différence de potentiel d'entrée est non nulle ($\epsilon \neq 0$). L'amplificateur fonctionne alors en saturation, et la tension qu'il délivre en sortie peut être égale soit en valeur inférieure, à la tension négative d'alimentation, soit en valeur supérieure, à la valeur positive d'alimentation (par exemple, un AO alimenté en $\pm 15\text{V}$, en fonctionnement non linéaire, délivrera $+15$ ou -15V). Physiquement, un AO « sort » du fonctionnement linéaire lorsque le circuit lui « demande » de débiter une puissance supérieure à celle que lui fournit son alimentation.

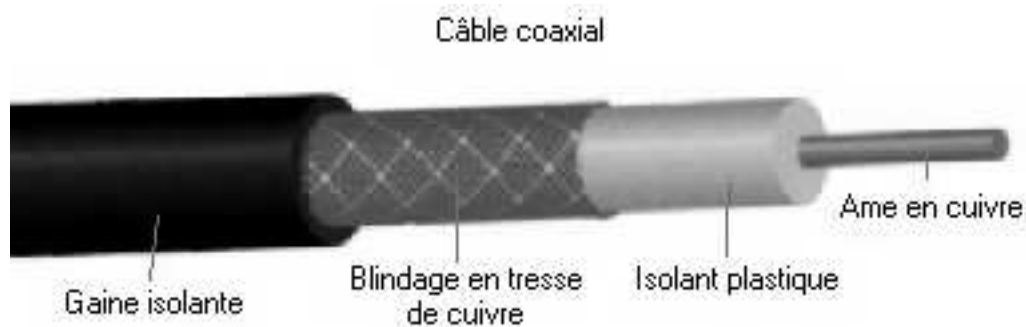
Annexe D

Lignes de transmission

D.1 Lignes sans perte

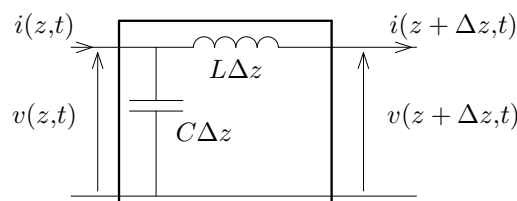
D.1.1 Quelques types de lignes

Il en existe plusieurs types. Les plus répandues sont les lignes bifilaires (composées de deux fils parallèles) et les câbles coaxiaux, où le signal est transporté par un fil central, entouré d'une gaine métallique servant de référence de tension (de masse).



D.1.2 Equation de propagation

Supposons pour simplifier que la ligne à étudier soit bifilaire, ou un câble coaxial. Par construction, il existe une inductance de couplage et une capacité entre les deux conducteurs, réparties sur la longueur de la ligne. On note C et L respectivement les capacité et inductance linéiques (ie par unité de longueur). Considérons un élément de longueur Δz de la ligne, à la position z ; sur cette portion, il existe une capacité $C\Delta z$ et une inductance $L\Delta z$, suivant le schéma :



On a les relations, approchées au deuxième ordre près :

$$\begin{aligned} i(z + \Delta z, t) - i(z, t) &= C \Delta z \frac{\partial v(z, t)}{\partial t} \\ v(z + \Delta z, t) - v(z, t) &= L \Delta z \frac{\partial i(z + \Delta z, t)}{\partial t} \approx L \Delta z \frac{\partial i(z, t)}{\partial t} \end{aligned}$$

On en déduit les équations aux dérivées partielles :

$$\begin{cases} \frac{\partial v(z, t)}{\partial z} = -L \frac{\partial i(z, t)}{\partial t} \\ \frac{\partial i(z, t)}{\partial z} = -C \frac{\partial v(z, t)}{\partial t} \end{cases}$$

On pose alors $LC = \frac{1}{c_0^2} \cdot v$ et i vérifient alors l'« équation de propagation » suivante :

$$\frac{\partial^2 f}{\partial z^2} - \frac{1}{c_0^2} \frac{\partial^2 f}{\partial t^2} = 0$$

Cette équation est également appelée l'« équation des télégraphistes ».

D.1.3 Résolution de l'équation

Il est aisé de vérifier que les solutions sont de la forme

$$\begin{cases} v(z, t) = f(z - v_0 t) + g(z + v_0 t) \\ i(z, t) = C c_0 [f(z - v_0 t) - g(z + v_0 t)] = \frac{1}{L c_0} [f(z - v_0 t) - g(z + v_0 t)] \end{cases}$$

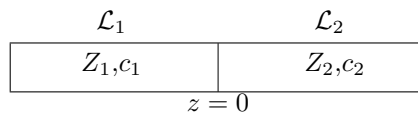
Ces solutions correspondent à des *ondes de tension et de courant* se déplaçant à la vitesse c_0 dans la ligne, pour f dans le sens des z positifs (onde *progressive*), pour g dans le sens des z négatifs (onde *régressive*).

On définit alors l'*impédance caractéristique* Z_0 de la ligne par $Z_0 = \sqrt{\frac{L}{C}}$, ou plus généralement $Z_0 = (\frac{L}{C})^{1/2}$, car Z_0 peut prendre une valeur complexe.

D.2 Interface entre deux lignes

D.2.1 Coefficients de réflexion/transmission

On considère deux lignes $\mathcal{L}_1$ et $\mathcal{L}_2$ placées en série, d'impédances respectives Z_1 et Z_2 , et de vitesses caractéristiques c_1 et c_2 . L'interface entre les deux lignes est à $z = 0$.



La tension et le courant dans la ligne k sont notés respectivement v_k et i_k . Les conditions aux limites^{D.1} sont :

$$\begin{cases} v_1(0, t) = v_2(0, t) \\ i_1(0, t) = i_2(0, t) \end{cases}$$

D.1. Elles sont liées aux conditions de continuité entre deux diélectriques : cf. un cours d'électromagnétisme.

On suppose qu'initialement ne circule qu'une onde progressive dans $\mathcal{L}_1$. Les ondes circulant dans les deux lignes vont être réfléchies à cette interface. On note Γ et T les coefficients de réflexion et de transmission entre $\mathcal{L}_1$ et $\mathcal{L}_2$:

$$\begin{cases} \Gamma &= \frac{\text{Amplitude de l'onde réfléchie}}{\text{Amplitude de l'onde incidente}} \\ T &= \frac{\text{Amplitude de l'onde transmise}}{\text{Amplitude de l'onde incidente}} \end{cases}$$

On a alors les relations :

$$\begin{cases} v_1(z,t) &= f\left(t - \frac{z}{c_1}\right) + \Gamma f\left(t + \frac{z}{c_1}\right) \\ i_1(z,t) &= \frac{1}{Z_1} \left[f\left(t - \frac{z}{c_1}\right) - \Gamma f\left(t + \frac{z}{c_1}\right) \right] \\ v_2(z,t) &= T f\left(t - \frac{z}{c_2}\right) \\ i_2(z,t) &= \frac{T}{Z_2} f\left(t - \frac{z}{c_2}\right) \end{cases}$$

Les conditions aux limites se traduisent par les relations :

$$\begin{cases} 1 + \Gamma &= T \\ \frac{1}{Z_1}(1 - \Gamma) &= \frac{T}{Z_2} \end{cases}$$

On obtient donc :

$$\begin{cases} \Gamma &= \frac{Z_2 - Z_1}{Z_2 + Z_1} \\ T &= \frac{2Z_2}{Z_2 + Z_1} \end{cases}$$

D.2.2 Cas particuliers

Dans le cas où la ligne $\mathcal{L}_1$ est branchée en sortie sur un circuit ouvert (par exemple, un câble coaxial débranché à une extrémité), Z_2 vaut alors $+\infty$ et donc $\Gamma = 1$. On peut alors montrer qu'en *régime harmonique* (ie quand v est supposée être une fonction sinusoïdale du temps), la forme de l'onde dans la ligne est *stationnaire*^{D.2} : il ne peut plus y avoir de propagation dans la ligne.

Dans le cas où la ligne $\mathcal{L}_1$ est court-circuitée, $Z_2 = 0$ et donc $\Gamma = -1$. De même que dans le cas précédent, en régime harmonique il n'y a plus de propagation^{D.3}.

Il n'y a qu'un seul cas où une tension appliquée à l'entrée de la ligne n'est *jamais* perturbée par une réflexion parasite : c'est celui où le coefficient de réflexion est nul. Dans ce cas $Z_1 = Z_2$. Un cas pratique est celui, par exemple, où un câble BNC (coaxial) est branché en sortie d'un générateur, et le relie à une platine de montage. L'impédance caractéristique d'une ligne coaxiale étant souvent de 50Ω , si on veut limiter les réflexions parasites au niveau du branchement sur la platine de montage, réflexions qui pourraient nuire à la « propreté » du signal délivré, il est nécessaire que l'impédance d'entrée de celui-ci soit égale à 50Ω . Or comme nous l'avons remarqué dans le corps du cours (paragraphe 4.1.4), un montage présente souvent une grande impédance d'entrée. Il suffit alors de placer en parallèle de l'entrée une petite résistance de 50 ohm .

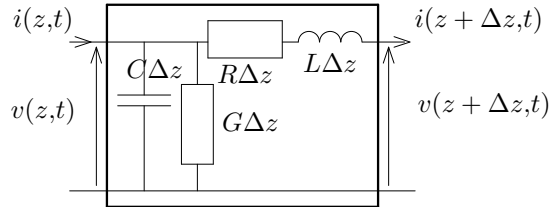
D.2. A la fréquence f , on obtient une tension de la forme $v(z,t) = V_0 \cos \omega t \cos kz$, avec $\omega = 2\pi f$ et $k = \omega/c$.

D.3. Cette fois-ci, la tension est de la forme $v(z,t) = V_0 \sin \omega t \sin kz$ lorsque le court-circuit est en $z = 0$.

D.3 Ligne avec pertes

D.3.1 Equation de propagation

On tient compte de résistances parasites entre les deux fils (le milieu les séparant est légèrement conducteur), et de leur résistance linéique, en modélisant la ligne sous la forme suivante :



Si on pose $Z_s = R + jL\omega$ et $Y_p = G + jC\omega$, les équations aux dérivées partielles géant le système se retrouvent facilement par simple transposition des relations du paragraphe D.1.2. On obtient donc :

$$\begin{cases} \frac{\partial v(z,t)}{\partial z} = -Z_s \frac{\partial i(z,t)}{\partial t} \\ \frac{\partial i(z,t)}{\partial z} = -Y_p \frac{\partial v(z,t)}{\partial t} \end{cases}$$

D'où l'équation des ondes avec pertes :

$$\frac{\partial^2 f}{\partial z^2} - Z_s Y_p \frac{\partial^2 f}{\partial t^2} = 0$$

D.3.2 Résolution de l'équation

Les solutions sont de la forme :

$$\begin{cases} v(z,t) = Ae^{-\gamma z} + Be^{+\gamma z} \\ i(z,t) = Y_0(Ae^{-\gamma z} + Be^{+\gamma z}) \end{cases}$$

avec :

$$\begin{cases} Z_0 = \sqrt{\frac{L}{C}} + j \frac{G}{2\sqrt{LC}\omega} \left(\frac{L}{C} - \frac{R}{G} \right) = \alpha + j\beta \\ \gamma = \left(\frac{R}{2} \sqrt{\frac{C}{L}} + \frac{G}{2} \sqrt{\frac{L}{C}} \right) + j\omega \sqrt{LC} \end{cases}$$

α est la constante d'atténuation et β la constante de phase. Ces solutions correspondent à des « ondes évanescentes » : il y a dissipation d'énergie sur le trajet dans le guide d'onde. L'affaiblissement est en $-\alpha z$. Il est couramment exprimé en dB/m sous la forme

$$A_m = 20 \log e^\alpha = 20\alpha \log e \approx 8,7\alpha \text{ dB/m}$$

Le comportement d'un câble coaxial dépend de la fréquence d'utilisation, mais typiquement, l'atténuation est de l'ordre de quelques dizaines de dB pour 100m.

Dans le cas particulier où $\frac{L}{C} = \frac{R}{G}$, l'impédance caractéristique est la même que dans le cas d'une ligne sans perte, et est purement résistive. En général néanmoins, l'impédance résultante est résistive *et* légèrement capacitive.

Annexe E

Rappels sur les nombres complexes

E.1 Introduction

L'ensemble des nombres réels est noté $\mathbb{R}$. Il regroupe tous les nombres entiers relatifs, les nombres fractionnels (sous-ensemble $\mathbb{Q}$) et les nombres transcendants (ceux qui ne peuvent s'exprimer sous la forme d'une fraction, comme π , e , $\sqrt{2}$, etc.).

Les nombres complexes ont été introduits pour des raisons pratiques liées à la résolution d'équations du second degré, au XVI^e siècle en Italie. Ils s'écrivent sous la forme $z = x+iy$, où x et y sont deux nombres réels et i un nombre tel que $i^2 = -1$.

E.2 Représentations algébrique et polaire

E.2.1 Représentation algébrique

E.2.1.1 Vocabulaire

La représentation algébrique d'un nombre complexe z est la forme précédemment évoquée ($z = x+iy$). On désigne alors sous le nom de :

- partie *réelle* le nombre x ; elle est notée $x = \Re(z)$;
- partie *imaginaire* le nombre y ; elle est notée $y = \Im(z)$.

Un nombre complexe dont la partie imaginaire est nulle est un nombre réel ; un nombre complexe dont la partie réelle est nulle est un *imaginaire pur*.

Si z est un nombre complexe de la forme $z = x+iy$, alors le nombre $z' = x-iy$ est appelé son *conjugué*. On le note souvent $\bar{z}$, ou z^* .

E.2.1.2 Règles de calcul

Considérons deux nombres complexes $z = x+iy$ et $z' = x'+iy'$. Alors :

- $z + z' = (x+iy) + (x'+iy') = (x+x') + i(y+y')$;
- $zz' = (x+iy).(x'+iy') = xx' + i(x'y + xy') + i^2 yy' = (xx' - yy') + i(xy' + x'y)$;
- $z = z'$ si et seulement si $x = x'$ et $y = y'$.

E.2.1.3 Conjugaison

Considérons un nombre complexe $z = x+iy$ et son conjugué $\bar{z} = x-iy$. Alors :

- $z.\bar{z} = (x+iy).(x-iy) = x^2 + y^2$;
- $z + \bar{z} = (x+iy) + (x-iy) = 2x = 2\Re(z)$;
- $z - \bar{z} = (x+iy) - (x-iy) = 2iy = 2i\Im(z)$;
- $\overline{\bar{z}} = \overline{x-iy} = \overline{x} + i\overline{y} = x + iy$ soit $\overline{\bar{z}} = z$.

Conjugué de $1/z$:

$$\overline{\left(\frac{1}{z}\right)} = \overline{\left(\frac{1}{x+iy}\right)} = \overline{\left(\frac{x-iy}{x^2+y^2}\right)} = \frac{x+iy}{x^2+y^2} = \frac{1}{x-iy}$$

Soit :

$$\overline{\left(\frac{1}{z}\right)} = \frac{1}{\bar{z}}$$

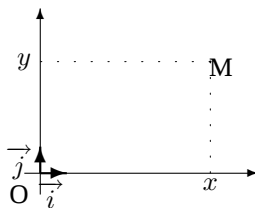
Considérons maintenant deux nombres complexes $z = x+iy$ et $z' = x'+iy'$. Il est alors facile de montrer que :

- $\overline{(z+z')} = \bar{z} + \bar{z}'$;
- $\overline{zz'} = \bar{z}\bar{z}'$;
- $\overline{\left(\frac{z}{z'}\right)} = \frac{\bar{z}}{\bar{z}'}$.

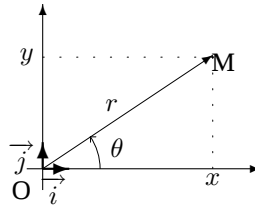
E.2.2 Représentation polaire

E.2.2.1 Interprétation géométrique

Soit un nombre complexe $z = x+iy$. On considère le plan doté d'un repère orthonormé $(O, \vec{i}, \vec{j})$, et le point M d'abscisse x et d'ordonnée y :



On dit que M est d'*abscisse* x . Le vecteur $\overrightarrow{OM}$ peut s'écrire $\overrightarrow{OM} = x\vec{i} + y\vec{j}$. Mais on peut également le caractériser entièrement par la donnée de sa *norme* r , c'est-à-dire sa « longueur », et de l'angle θ qu'il fait avec l'axe des abscisses :



D'une part le théorème de Pythagore nous dit que $r^2 = x^2 + y^2$, et d'autre part on peut déterminer de manière univoque θ par les relations suivantes :

$$\begin{cases} \cos \theta &= x / \sqrt{x^2 + y^2} \\ \sin \theta &= y / \sqrt{x^2 + y^2} \end{cases}$$

On vérifie que l'on peut facilement exprimer z en fonction de r et θ . En effet,

$$z = x + iy = \sqrt{x^2 + y^2} \left(\frac{x}{\sqrt{x^2 + y^2}} + i \frac{y}{\sqrt{x^2 + y^2}} \right)$$

On en déduit :

$$z = r(\cos \theta + i \sin \theta)$$

On peut montrer que $(\cos \theta + i \sin \theta)$ peut s'écrire $e^{i\theta}$.

E.2.2.2 Représentation polaire

On a montré dans le paragraphe précédent qu'un nombre complexe $z = x + iy$ pouvait s'écrire sous la forme $z = re^{i\theta}$, avec :

$$\begin{cases} r &= \sqrt{x^2 + y^2} \\ \cos \theta &= x / \sqrt{x^2 + y^2} \\ \sin \theta &= y / \sqrt{x^2 + y^2} \end{cases}$$

On appelle r le *module*, et θ l'*argument* ou la *phase* de z . On les note respectivement $|z|$ et $\text{Arg}(z)$.

E.2.2.3 Règles de calcul et conjugaison

Considérons deux nombres complexes $z = re^{i\theta}$ et $z' = r'e^{i\theta'}$. Alors :

- $zz' = rr'(\cos \theta + i \sin \theta)(\cos \theta' + i \sin \theta') = rr'[(\cos \theta \cos \theta' - \sin \theta \sin \theta') + i(\sin \theta \cos \theta' + \cos \theta \sin \theta')]$, soit en utilisant les formules de trigonométrie $zz' = rr'[\cos(\theta + \theta') + i \sin(\theta + \theta')] = rr'e^{i(\theta + \theta')}$: le module du produit est égal au produit des modules, et l'argument du produit est égal à la somme des arguments. On aurait obtenu directement le résultat sur les arguments en utilisant le fait que le produit des exponentielles de deux termes est égal à l'exponentielle de la somme de ces deux termes : $e^{i\theta}e^{i\theta'} = e^{i(\theta + \theta')}$
- $\bar{z} = \overline{r(\cos \theta + i \sin \theta)} = r(\cos \theta - i \sin \theta)$ donc $\bar{z} = re^{-i\theta}$: le module de $\bar{z}$ est égal au module de z , l'argument de $\bar{z}$ est l'opposé de l'argument de z ;
- $1/z = 1/(re^{i\theta}) = e^{-i\theta}/r$: le module de $1/z$ est égal à l'inverse du module de z , l'argument de $1/z$ est égal à l'opposé de l'argument de z .

E.3 Tables récapitulatives

E.3.1 Quelques nombres complexes remarquables

Nombre	Représentation algébrique	Représentation polaire	Remarques
0	$0+i0$	module=0, argument indéterminé	L'argument de 0 ne peut être déterminé, car la notion d'argument n'a alors pas de sens.
j	$\frac{1}{2}(1+i\sqrt{3})$	$1.e^{i\frac{2\pi}{3}}$	Racine « cubique » de l'unité, car $j^3=1$, ce nombre vérifie en outre $\bar{j} = j^2 = -j$.
1	$1+0i$	$1.e^{i0}$	
-1	$-1+0i$	$1.e^{i\pi}$	
i	$0+1i$	$1.e^{i\frac{\pi}{2}}$	$i^2 = -1$
-i	$0-1i$	$1.e^{-i\frac{\pi}{2}}$	

Remarque : Afin de ne pas confondre le nombre i avec ce qui pourrait être une intensité, en électricité ce nombre est souvent noté j . En contrepartie, quand on a besoin d'écrire la racine cubique de l'unité, on la note souvent simplement ω . On écrira par exemple en électricité $\frac{1}{2}(\sqrt{2}+i\sqrt{2}) = e^{j\frac{\pi}{4}}$.

E.3.2 Règles de calcul et propriétés

Dans tout le tableau suivant, on considère deux nombres complexes non nuls $z = x+iy = re^{i\theta}$ et $z' = x'+iy' = r'e^{i\theta'}$.

Expression	Représentation algébrique	Représentation polaire	Autre expression
$\bar{z}$	$x-iy$	$re^{-i\theta}$	
$\frac{\bar{z}}{z}$			$\frac{z}{z}$
$z + \bar{z}$	$2x$	$2r \cos \theta$	$2\Re(z)$
$z - \bar{z}$	$2iy$	$2ir \sin \theta$	$2i\Im(z)$
$ z $	$\sqrt{x^2 + y^2}$	r	$\sqrt{z\bar{z}}$
$\text{Arg}(z)$	$\arctan \frac{y}{x}$ si $x > 0$ $\pi + \arctan \frac{y}{x}$ si $x < 0$ $+\pi/2$ si $x = 0$ et $y > 0$ $-\pi/2$ si $x = 0$ et $y < 0$	θ	
$\frac{1}{z}$		$e^{-i\theta}/r$	$\bar{z}/ z ^2$
zz'	$(xx' - yy') + i(xy' + x'y)$	$rr'e^{i(\theta+\theta')}$	
z/z'		$\frac{r}{r'}e^{i(\theta-\theta')}$	
$z + z'$	$(x + x') + i(y + y')$		
$\bar{z} + \bar{z}'$			$\overline{z + z'}$
zz'			$\overline{z\bar{z}'}$
$ zz' $		rr'	$ z z' $
$\left \frac{z}{z'}\right $		r/r'	$\frac{ z }{ z' }$
$\text{Arg}(zz')$		$\theta + \theta'$	$\text{Arg}(z) + \text{Arg}(z')$
$\text{Arg}(z/z')$		$\theta - \theta'$	$\text{Arg}(z) - \text{Arg}(z')$

Annexe F

Liste d'abréviations usuelles en électricité

Ω : Ohm

- A -

A : Ampère

ACL : affichage à cristaux liquides

ADC : analog to digital converter

AM : amplitude modulation

AO(P) : amplificateur opérationnel

ASIC : application-specific integrated circuit

- B -

BVP : boucle à verrouillage de phase

- C -

C : Coulomb

CAN : convertisseur analogique/numérique

CNA : convertisseur numérique/analogique

CPU : central processing unit

- D -

DAC : digital to analog converter

DEL : diode électro-luminescente

- F -

F : Farad

FET : field effect transistor

FM : frequency modulation

- G -

GBF : générateur basse fréquence

G_{dB} : gain en décibels

- H -

H : Henri

- I -

I/O : input/output

IC : integrated circuit

- J -**J** : Joule**- K -****LCD** : liquid-crystal display**LED** : light-emitting diode**- M -****MAS** : machine asynchrone**MCC** : machine à courant continu**MOS** : metal-oxyde semiconductor**MS** : machine synchrone**MUX** : multiplexeur**- O -****OCT** : osceillateur contrôlé en tension**- P -****PAL** : programmable array logic**PLA** : programmable logic array**PLL** : phase locked loop**- R -****RAM** : random access memory**ROM** : read only memory**- S -****S** : Siemens**SNR** : signal to noise ratio**- T -****TEC** : transistor à effet de champ**TF** : transformée de Fourier**TTL** : transistor-transistor logic**- U -****UHF** : ultra high frequency**- V -****V** : Volt**VAR** : Volt Ampère réactif**VCO** : voltage controled oscillator**VHF** : very high frequency**- W -****W** : Watt

Index

- accepteur, *voir* atome, accepteur
- adaptation, 57
- ADC, 98
- additionneur, 83
- admittance, 37
- afficheur 7 segments, 86
- alternateur, 125
- amorçage, 127
- Ampère, 28
- amplificateur opérationnel, 139–142
- amplitude, 35
- angle de retard, 130
- asservissement, 59, 61
- atome
 - accepteur, 44
 - donneur, 43
- atténuation, 62
- autocorrélation, 66

- balai, 123
- bande
 - atténuée, 65
 - de transition, 65
 - passante, 64, 65, 68, 74
- bascule, 87–91
 - maître esclave, 90
 - D, 89
 - JK, 90
 - RS asynchrone, 87
 - RST, 89
- base, 50
- bit*, 76
- blocage, 51, 127
- bobinage
 - primaire, 114
 - secondaire, 114
- bobine, 33, 36, 112
- bruit, 66–72
 - blanc, 67
 - de basse fréquence, 68
 - de grenaille, 68
 - de papillotement, 68
 - de résistance, 67
 - de scintillement, 68
 - en $1/f$, 68
 - en créneaux, 69
 - en excès, 68
 - et filtrage, 67
 - Johnson, 67
 - popcorn, 69
 - rose, 68
 - thermique, 67
- burst noise*, 69
- byte*, *voir* octet

- câble coaxial, 143
- cage de Faraday, 74
- CAN, 98
- capacité, 34
- Carson (règle de), 108
- cascadage, 55, 71
- champ de rétention de la diffusion, 46, 47
- charge spatiale, *voir* déplétion
- circuit électrique, 27–31
- CNA, 99, 100
- codage, 81
 - BCD, 85
 - binaire codé décimal, 85
 - binaire naturel, 76
 - binaire réfléchi, 80
 - hexadécimal, 76
- coefficient
 - d'émission, 48
 - de réflexion, 144
 - de transmission, 144
- collecteur, 50, 123
- commande, 60
- commutation, 127
- compteur, 91
 - programmable, 93
- condensateur, 34, 36
- conductance, 32
- conducteur électrique, 27
- conjugaison, 148
- consigne, 60
- contréréaction, 58–62
- convertisseur
 - à approximations successives, 99
 - à comptage d'impulsions, 98
 - à rampe, 98
 - analogique/numérique, 98
 - flash, 100
 - numérique/analogique, 99, 100
 - statique, 127
 - tension/fréquence, 99
- convolution, 18, 25
- Coulomb, 27
- courant continu, 123
- courant inverse de saturation, 48

- courant électrique, 27
- CPU, 94
- crépitement, 69
- cycloconvertisseur, 127

- DAC, 99, 100
- De Morgan (formules de), 79
- de retour
 - branche, *voir* inverse, branche
- décodage, 81
- démodulation
 - angulaire, 110
 - d'amplitude, 109–110
 - incohérente, 109
- démultiplexage, 82
- densité spectrale de puissance, 66, 133
- déplétion, 46
- détection
 - d'enveloppe, 109
 - par redressement, 110
 - quadratique, 109
 - synchrone, 110
- diagramme de Bode, 62–65
- différence de potentiel, 28
- digit, *voir bit*
- diode, 45, 128
- dipôle, 27
 - actif, 37
 - capacitif, *voir* condensateur
 - générateur, 29
 - impédance d'un, 37
 - inductif, *voir* bobine
 - passif, 29, 37
 - récepteur, 29, 35
- Dirac
 - impulsion de, 19, 24, 39
 - peigne de, 22, 40
- directe
 - branche, 60
 - chaîne, *voir* directe, branche
 - polarisation, 48
- diviseur
 - de courant, 136
 - de tension, 135
- donneur, *voir* atome, donneur
- dopage, 43
- drain, 53

- échantillonnage, 95–101
- échantillonneur, 97
- échelon de Heaviside, 20
- effet Joule, 32, 57
- électron libre, 42
- émetteur, 50
- énergie, 14, 29
 - électrostatique, 34
 - magnétique, 34
- entrée (non)inverseuse, 139
- équation des télégraphistes, 144

- état
 - bloqué, 127
 - passant, 127
- étoile, 118, 119

- facteur de bruit, 70
- facteur de puissance, 36
- Farad, 34
- fermeture, 127
- Ferraris (théorème de), 122
- filtre, 64
 - coupe-bande, 65, 74
 - ordre d'un, 64
 - passe-bande, 64
 - passe-bas, 64, 74
 - passe-haut, 64
- flicker noise*, 68
- fonction de transfert, 25, 40, 62–66
- force (contre)-électromotrice, 113
- Fourier, *voir* Transformée de Fourier
- fréquence, 14, 35, 38
 - de coupure, 64, 66
 - instantanée, 106

- gabarit, 65–66
- gain, 62–66
 - de contre réaction, 60
 - de courant
 - en mode direct, 52
 - en mode F, 52
 - en mode inverse, 52
 - en mode R, 52
 - direct, *voir* gain, en boucle ouverte
 - en boucle fermée, 60
 - en boucle ouverte, 60
 - en décibels, 62
- générateur, 29
- génération, *voir* paire électron-trou, génération
- gradateur, 127
- grille, 53

- hacheur, 127
- Henri, 33
- hexa, hexadécimal, 76
- horloge, 87

- impédance, 37, 55
- impureté, *voir* dopage
- inductance, 33
- inducteur, 123, 125
- induction, 33, 112, 117, 121
- induit, 123, 125
- intensité, 27
- interrupteur, 127
- inverse
 - branche, 60
 - chaîne, *voir* inverse, branche
 - polarisation, 48
- jonction

- de commande, 50
- PN, 45–49
- latch*, 92
- ligne
 - courant de, 119
 - de transmission, 143–146
 - tension de, 119
- logique
 - combinatoire, 77–86
 - négative, 77
 - séquentielle, 86–95
- loi
 - d’Ohm, 31, 37
 - de Kirchhoff, 30–31
 - des mailles, 30
 - des nœuds, 30
- machine
 - à courant continu, 122–124
 - à excitation indépendante, 124
 - à excitation parallèle, 124
 - à excitation série, 124
 - asynchrone, 125
 - shunt*, 124
 - synchrone, 125
- MAS, 125
- masse, 28
- matrice de transfert, 55
- MCC, 122–124
- mémoire
 - morte, 84
 - vive, 93
- microprocesseur, 94
- Millman (théorème de), 136
- mode
 - bloqué, 51
 - normal direct, 51, 52
 - normal inverse, 51, 52
 - saturé, 51, 52
- modulant, 104
- modulation
 - à porteuse conservée, 104, 109
 - à porteuse supprimée, 106
 - angulaire, 106–108
 - d’amplitude, 104–106
 - de fréquence, 107, 108
 - de phase, 106
 - indice de, 104
- moteur universel, 124
- MS, 125
- multiplexage, 82, 103
- neutre, 118
- nombre complexe, 147–150
- Norton (théorème de), 138
- nnp, 49
- Nyquist, 67
- octet, 76
- Ohm, 32
- onde latérale, 105
- onduleur, 127, 131–132
 - monophasé en pont, 131
- opérateur, 59, 77
- ordre
 - d’un filtre, 64
 - d’un système polyphasé, 115
- ouverture, 127
- paire électron-trou
 - génération, 43
 - recombinaison, 43
- PAL, 85
- parallèle, 33, 38
- parasites, 73–75
 - par conduction, 74
 - par rayonnement, 74
- partie
 - imaginaire, 147
 - réelle, 147
- pas d’échantillonnage, 98
- pas d’échantillonnage, 76
- période, 14, 21
- phase
 - courant de, 119
 - instantanée, 106
 - tension de, 119
- PLA, 85
- pnp, 49
- polarisation, 48
- pont
 - de Graëtz, 129, 131
 - diviseur, voir diviseur (de tension, de courant)
- portes, 77
 - ET, AND, 77
 - NON, NOT, 78
 - NON-ET, NAND, 79
 - OU exclusif, XOR, 79
 - OU, OR, 78
- porteuse, 104
- prise électrique, 119
- puissance, 14, 29, 35–36, 57, 62
 - active, 36
 - apparente, 36
 - de bruit, 68, 70
 - facteur de, 36
 - instantanée, 35
 - instantanée(complexe), 36
 - moyenne, 36
 - moyenne(complexe), 37
 - réactive, 36
- pulsation, 35
- quadrant, 127
- quadripôle, 27, 55–58, 70
- RAM, 93

- rapport de bruit, 70
- rapport de transformation, 114
- rapport signal à/sur bruit, 70
- récepteur, *voir* dipôle récepteur
- recombinaison, *voir* paire électron-trou, recombinaison
- redresseur, 127–131
 - à cathode commune, 128
 - à diodes, 128
 - à pont de Graëtz, 129, 131
 - demi-onde, 128
 - simple alternance, 128
- registre, 91–93
 - à décalage, 92
 - tampon, 92
- régulation, 61
- réluctance, 114
- réponse impulsionnelle, 24
- représentation
 - algébrique, 147
 - complexe, 36
 - polaire, 148
- réseau
 - EDF, 120
 - logique programmable, 85
 - polyphasé, 115
 - programmable logique, 85
 - triphasé, 115–121
- résistance, 31–33
- résistor, *voir* résistance
- rhéostat, 124
- ROM, 84
- rotor, 123
- saturation, 51, 52
- semi-conducteur, 42–45
 - extrinsèque, 43–45
 - intrinsèque, 42
- série, 32, 38
- Shannon, 95
- shot noise*, 68
- Siemens, 32
- signal, 12–15
 - à valeurs continues, 13
 - à valeurs discrètes, 13, 76–77
 - harmonique, 35
 - monochromatique, 35
- signal to noise ratio*, *voir* rapport signal à bruit
- source, 53
- soustracteur, 84
- spectre, 16, 38–40, 133
 - de raies, 40, 108
- stator, 123
- support fréquentiel, 40
- surmodulation, 105
- système, 59
 - asynchrone, 87
 - bouclé, 58–62
 - direct, 116
 - homopolaire, 116
 - invariant, 24
 - inverse, 116
 - linéaire, 24
 - polyphasé, 115
 - équilibré, 115, 122
 - synchrone, 87
 - triphasé, 115–121
- table de Karnaugh, 79
- taux
 - d'ondulation, 66
 - d'harmoniques, 21, 132
- température
 - de bruit, 70
 - équivalente de bruit, 69
 - ramenée en entrée, 71
- temps
 - continu, 12
 - discret, 12
- tension thermodynamique, 48
- tension électrique, 28
- théorème de Friss, 71
- thyristor, 128, 130
- Thévenin (théorème de), 137
- transcodage, 82
- transfert, *voir* fonction de transfert
- transformateur, 112–115
- Transformée de Fourier, 15–23, 39, 66, 133–134
 - inverse, 15
- transimpédance, 55
- transistor, 128
 - à effet de champ, 53
 - bipolaire, 49–52
 - MOS, 53–54
- triangle, 118, 120
- trou, 42
- valeur efficace, 35, 115
- variable
 - primaire, 87
 - secondaire, 87
- VCO, 99
- Volt, 28
- Volt Ampère réactif, 36
- zone
 - de charge spatiale, *voir* déplétion
 - de déplétion, *voir* déplétion