



Département d'Electricité, Electronique et Informatique
(Institut Montefiore)

Notes théoriques du cours ELEC0029

ANALYSE ET FONCTIONNEMENT DES SYSTEMES D'ENERGIE ELECTRIQUE

Thierry VAN CUTSEM

directeur de recherches FNRS
professeur adjoint ULg

janvier 2015

Table des matières

1	Puissances en régime sinusoïdal	5
1.1	Puissance fournie ou absorbée par un dipôle	5
1.2	Puissance transitant dans plusieurs conducteurs	6
1.3	Régime sinusoïdal et phaseurs	7
1.4	Puissances instantanée, active, réactive, fluctuante et apparente	8
1.5	Puissance complexe	10
1.6	Expressions relatives aux dipôles	11
1.7	Facteur de puissance et compensation des charges	11
2	Systèmes et régimes triphasés équilibrés	14
2.1	Principe	14
2.2	Tensions de ligne (ou composées)	17
2.3	Connexions en étoile et en triangle	18
2.4	Analyse par phase	19
2.5	Puissances en régime triphasé	22
3	Quelques propriétés du transport de l'énergie électrique	24
3.1	Transit de puissance et chute de tension dans une liaison	24

3.2	Caractéristique QV à un jeu de barres d'un réseau	28
3.3	Puissance de court-circuit	29
3.4	Puissance maximale transmissible à une charge	30
4	La ligne de transport	36
4.1	Paramètres linéiques d'une ligne	36
4.2	Caractéristiques des câbles	46
4.3	La ligne traitée comme un composant distribué	47
4.4	Quelques propriétés liées à l'impédance caractéristique	50
4.5	Schéma équivalent d'une ligne	51
4.6	Limite thermique	52
5	Le système "per unit"	55
5.1	Passage en per unit d'un circuit électrique	56
5.2	Passage en per unit de deux circuits magnétiquement couplés	57
5.3	Passage en per unit d'un système triphasé	58
5.4	Changement de base	60
6	Le transformateur de puissance	61
6.1	Transformateur monophasé	61
6.2	Transformateur triphasé	68
6.3	Valeurs nominales, système per unit et ordres de grandeur	76
6.4	Autotransformateur	77
6.5	Ajustement du nombre de spires d'un transformateur	79
6.6	Transformateur à trois enroulements	81

6.7	Transformateur déphaseur	82
7	Le calcul de répartition de charge (ou load flow)	85
7.1	Les équations de load flow	85
7.2	Spécification des données du load flow	87
7.3	Résolution numérique des équations de load flow	90
7.4	Extensions de la formulation	94
7.5	Découplage électrique	98
7.6	L'approximation du courant continu (ou DC load flow)	100
7.7	Analyse de sensibilité	102
8	La machine synchrone	106
8.1	Principe de fonctionnement	106
8.2	Les deux types de machines synchrones	111
8.3	Modélisation au moyen de circuits magnétiquement couplés	114
8.4	Transformation et équations de Park	118
8.5	Energie, puissance et couple	122
8.6	La machine synchrone en régime établi	124
8.7	Valeurs nominales, système per unit et ordres de grandeur	129
8.8	Courbes de capacité	130
9	Comportement des charges	133
9.1	Comportement du moteur asynchrone en tant que charge	133
9.2	Modèles simples des variations des charges avec la tension et la fréquence	141

10 Régulation de la fréquence	146
10.1 Régulateur de vitesse	147
10.2 Régulation primaire	150
10.3 Régulation secondaire	153
11 Régulation de la tension	160
11.1 Contrôle de la tension par condensateur ou inductance shunt	161
11.2 Régulation de tension des machines synchrones	162
11.3 Compensateurs synchrones	170
11.4 Compensateurs statiques de puissance réactive	171
11.5 Régulation de tension par les régulateurs en charge	179
12 Analyse des défauts équilibrés	184
12.1 Phénomènes liés aux défauts	184
12.2 Comportement de la machine synchrone pendant un court-circuit	187
12.3 Calcul des courants de court-circuit triphasé	196
12.4 Appendice. Phénomènes physiques liés à la foudre	201
13 Analyse des systèmes et régimes triphasés déséquilibrés	205

Chapitre 1

Puissances en régime sinusoïdal

Dans ce chapitre nous rappelons quelques définitions et relations fondamentales, essentielles pour l'analyse des systèmes électriques de puissance. L'accent est mis sur les notions de puissance en régime sinusoïdal.

1.1 Puissance fournie ou absorbée par un dipôle

Considérons le dipôle représenté à la figure 1.1, avec ses deux bornes d'extrémité.

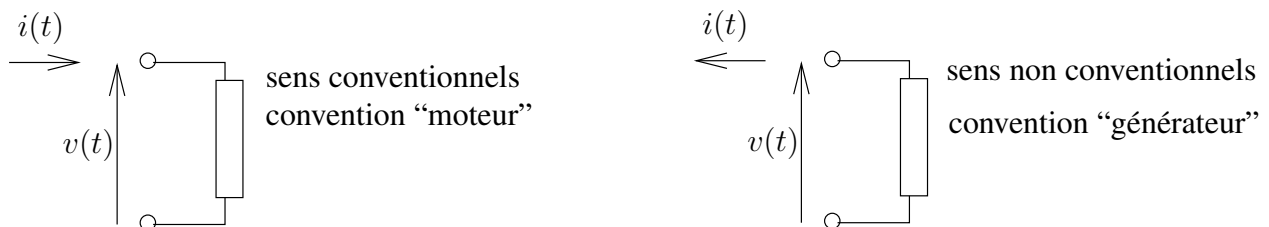


FIGURE 1.1 – dipôle : conventions d'orientation du courant par rapport à la tension

La tension $v(t)$ aux bornes du dipôle est la différence entre le potentiel de la borne repérée par l'extrémité de la flèche et le potentiel de la borne repérée par son origine.

Deux conventions sont possibles en ce qui concerne l'orientation du courant i par rapport à la tension v (cf figure 1.1) :

- la convention *moteur* correspond aux *sens conventionnels* de la Théorie des circuits. Le courant est compté comme positif s'il *entre* dans le dipôle par la borne correspondant à l'extrémité de la flèche repérant la tension. Dans ce cas, le produit $p(t) = v(t) i(t)$ représente la puissance instantanée *absorbée* par le dipôle. Une valeur positive (resp. négative) indique donc que le dipôle consomme (resp. fournit) de la puissance à l'instant t ;

- la convention *générateur* correspond aux *sens non conventionnels* de la Théorie des circuits. Le courant est considéré comme positif s'il *sort* du dipôle par la borne correspondant à l'extrémité de la flèche repérant la tension. Le produit $p(t) = v(t) i(t)$ représente la puissance instantanée *générée* par le dipôle. Une valeur positive (resp. négative) indique donc que le dipôle produit (resp. consomme) de la puissance à l'instant t .

1.2 Puissance transitant dans plusieurs conducteurs

L'expression de la puissance consommée (ou produite) par un dipôle se généralise aisément à ensemble de n conducteurs ($n \geq 2$), à condition de supposer que la somme des courants i_j dans ces conducteurs est nulle :

$$\sum_{j=1}^n i_j = 0 \quad (1.1)$$

ce qui était évidemment le cas du dipôle considéré à la section précédente.

Considérons que ces conducteurs relient des circuits A et B, comme représenté à la figure 1.2, entre lesquels de l'énergie électrique est échangée. S'il n'y a pas d'autre liaison entre A et B (auquel cas, à la figure 1.2, Σ est une "coupe"), la première loi de Kirchhoff impose la relation 1.1. Les circuits A et B pourraient être reliés par d'autres conducteurs (non représentés). Dans ce cas, nous supposons que le fonctionnement est tel que la relation (1.1) est satisfaite.

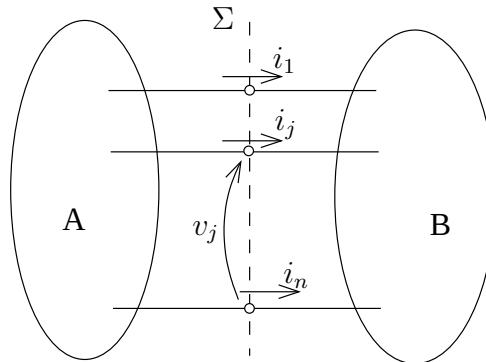


FIGURE 1.2 – puissance transitant dans plusieurs conducteurs

Nous prenons le n -ème conducteur comme référence des tensions. Par exemple, ce conducteur pourrait être relié à la terre, auquel cas celle-ci sera considérée comme référence des tensions. Notons toutefois que ce choix n'est pas obligatoire pour ce qui suit.

Soit v_j ($j = 1, \dots, n-1$) la tension entre le j -ème et le n -ème conducteur, comme représenté à la figure 1.2. Avec le choix des courants indiqué à la figure 1.2, le circuit B est dans la convention moteur. Dans ces conditions, l'expression :

$$p(t) = \sum_{j=1}^{n-1} v_j i_j \quad (1.2)$$

représente la puissance instantanée traversant les n conducteurs (ou passant dans la section Σ) de A vers B, c'est-à-dire la puissance absorbée par B, ou encore la puissance produite par A.

1.3 Régime sinusoïdal et phaseurs

En régime sinusoïdal, toute tension se présente sous la forme :

$$v(t) = \sqrt{2} V \cos(\omega t + \phi) \quad (1.3)$$

où $\sqrt{2}V$ est l'amplitude (ou *valeur de crête*) de la tension, V sa *valeur efficace*¹ et ω la pulsation, reliée à la fréquence f et la période T par :

$$\omega = 2\pi f = \frac{2\pi}{T} \quad (1.4)$$

De même, tout courant se présente sous la forme :

$$i(t) = \sqrt{2} I \cos(\omega t + \psi) \quad (1.5)$$

Définissons les grandeurs complexes suivantes :

$$\bar{V} = V e^{j\phi} \quad (1.6)$$

$$\bar{I} = I e^{j\psi} \quad (1.7)$$

où les lettres surlignées désignent des nombres complexes, afin de les différencier des nombres réels. $\bar{V}$ est le *phaseur* relatif à la tension $v(t)$ tandis que $\bar{I}$ est le *phaseur* relatif au courant $i(t)$.

On a évidemment :

$$v(t) = \sqrt{2} \operatorname{re} (V e^{j(\omega t + \phi)}) = \sqrt{2} \operatorname{re} (\bar{V} e^{j\omega t}) \quad (1.8)$$

$$i(t) = \sqrt{2} \operatorname{re} (I e^{j(\omega t + \psi)}) = \sqrt{2} \operatorname{re} (\bar{I} e^{j\omega t}) \quad (1.9)$$

Dans le plan complexe, aux nombres $\bar{V} e^{j\omega t}$ et $\bar{I} e^{j\omega t}$, on peut associer des vecteurs tournants. Chaque vecteur part de l'origine $0 + j0$ et aboutit au nombre complexe en question. Chaque grandeur sinusoïdale est, au facteur $\sqrt{2}$ près, la projection sur l'axe réel du vecteur tournant correspondant.

Usuellement, pour représenter ces vecteurs tournants, on considère leur position en $t = 0$. A cet instant, le vecteur tournant relatif à la tension n'est rien d'autre que le phaseur $\bar{V}$ et celui relatif au courant est le phaseur $\bar{I}$.

Une représentation graphique des phaseurs est donnée à la figure 1.3 (dont une partie sera utilisée dans un développement ultérieur). On désigne ce type de schéma sous le terme de *diagramme de phaseur*.

1. la pratique a consacré l'usage des valeurs efficaces pour caractériser les grandeurs sinusoïdales : lorsque l'on donne la valeur d'une tension alternative, il s'agit, sauf mention contraire, de la valeur efficace. Rappelons que V est la valeur de la tension continue qui, appliquée à une résistance, y dissipe la même puissance que la tension sinusoïdale (1.3) en moyenne

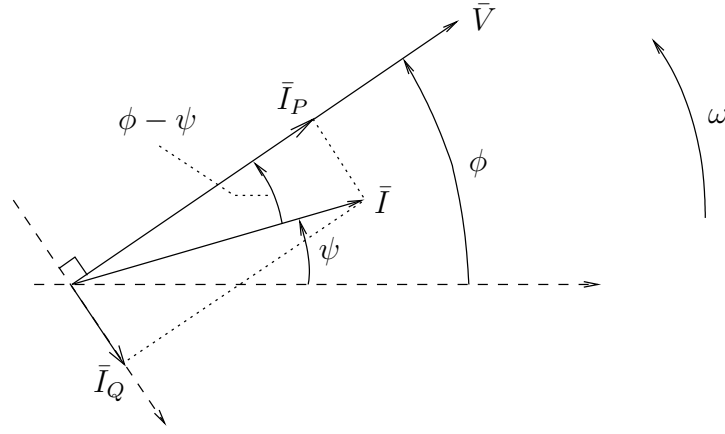


FIGURE 1.3 – diagramme de phaseur

1.4 Puissances instantanée, active, réactive, fluctuante et appa- rente

Considérons un dipôle soumis à une tension $\bar{V}$ et parcouru par un courant $\bar{I}$.

Projetons le vecteur $\bar{I}$ sur l'axe défini par le vecteur $\bar{V}$ et de même orientation que ce dernier (cf figure 1.3). Soit $\bar{I}_P$ le vecteur projeté ainsi obtenu. On peut écrire :

$$\bar{I}_P = I_P e^{j\phi} \quad (1.10)$$

où I_P est un nombre réel, positif si le vecteur $\bar{I}_P$ est de même sens que $\bar{V}$ et négatif dans le cas contraire. I_P est appelé *courant actif*. On a :

$$I_P = I \cos(\phi - \psi) \quad (1.11)$$

Projetons à présent le vecteur $\bar{I}$ sur un axe perpendiculaire au vecteur $\bar{V}$ et *en retard* sur ce dernier (cf figure 1.3). Soit $\bar{I}_Q$ le vecteur projeté ainsi obtenu. On peut écrire :

$$\bar{I}_Q = I_Q e^{j(\phi - \frac{\pi}{2})} \quad (1.12)$$

où I_Q est un nombre réel, positif si $\bar{I}_Q$ est en retard sur $\bar{V}$ et négatif dans le cas contraire. I_Q est appelé *courant réactif*. On a :

$$I_Q = I \sin(\phi - \psi) \quad (1.13)$$

Exprimons $i(t)$ en fonction des courants actif et réactif. On a successivement :

$$\begin{aligned} i(t) &= \sqrt{2} \operatorname{re}(\bar{I}e^{j\omega t}) = \sqrt{2} \operatorname{re}(\bar{I}_P e^{j\omega t} + \bar{I}_Q e^{j\omega t}) = \sqrt{2} \operatorname{re}(I_P e^{j(\omega t + \phi)} + I_Q e^{j(\omega t + \phi - \frac{\pi}{2})}) \\ &= \sqrt{2} I_P \cos(\omega t + \phi) + \sqrt{2} I_Q \sin(\omega t + \phi) \end{aligned} \quad (1.14)$$

On voit que la variation sinusoïdale du courant peut se décomposer en la somme de deux variations sinusoïdales, l'une relative au courant actif, l'autre au courant réactif.

En utilisant l'expression (1.3) de la tension et l'expression (1.14) du courant, la puissance instantanée vaut :

$$\begin{aligned} p(t) = v(t) i(t) &= 2V I_P \cos^2(\omega t + \phi) + 2V I_Q \cos(\omega t + \phi) \sin(\omega t + \phi) \\ &= V I_P [1 + \cos 2(\omega t + \phi)] + V I_Q \sin 2(\omega t + \phi) \end{aligned} \quad (1.15)$$

On en déduit les propriétés importantes suivantes :

- la puissance instantanée est la somme de deux composantes, l'une relative au courant actif, l'autre au courant réactif
- la composante relative au courant actif se présente elle-même sous forme d'une somme d'un terme constant et d'un terme oscillatoire de pulsation 2ω , changeant donc de signe quatre fois par période. Toutefois, la somme de ces deux termes ne change jamais de signe et correspond donc à une puissance allant toujours dans le même sens
- la composante relative au courant réactif ne comporte qu'un terme oscillatoire de pulsation 2ω
- sur une période, les composantes oscillatoires ont une moyenne nulle. La valeur moyenne de la puissance $p(t)$ est donc la constante présente dans la composante relative au courant actif. Cette valeur, notée P , est appelée *puissance active*. On a donc :

$$P = V I_P \quad (1.16)$$

et en utilisant (1.11) :

$$P = V I \cos(\phi - \psi) \quad (1.17)$$

- l'amplitude de la composante relative au courant réactif, notée Q , est appelée *puissance réactive*. On a donc :

$$Q = V I_Q \quad (1.18)$$

et en utilisant (1.13) :

$$Q = V I \sin(\phi - \psi) \quad (1.19)$$

- on sait que dans un circuit RLC, le déphasage du courant par rapport à la tension, c'est-à-dire l'existence du courant réactif I_Q , est dû aux éléments L et C. La puissance $V I_Q \sin 2(\omega t + \phi)$ est le taux de variation de l'énergie emmagasinée sous forme magnétique et électrique. Cette énergie est toujours positive. A certains moments, les éléments L et C accumulent de l'énergie et à d'autres ils en restituent. Comme il s'agit d'éléments passifs, il ne peuvent en restituer plus qu'ils n'en ont reçu.
- la somme des termes oscillatoires, notée $p_f(t)$, est appelée *puissance fluctuante*. On a :

$$p_f(t) = V I \cos(\phi - \psi) \cos 2(\omega t + \phi) + V I \sin(\phi - \psi) \sin 2(\omega t + \phi) = V I \cos(2\omega t + \phi + \psi) \quad (1.20)$$

résultat que l'on obtient plus directement en multipliant (1.3) par (1.5). Etant de moyenne nulle, la puissance fluctuante ne correspond à aucun travail utile. La puissance active est la seule composante utile.

- le produit $V I$ est appelé *puissance apparente*. On voit que puissances apparente et active coïncident quand il n'y a pas de déphasage entre la tension et le courant, c'est-à-dire pas de courant réactif.

Les grandeurs $p(t)$, $p_f(t)$, P , Q et S ont toutes la dimension d'une puissance et devraient donc s'exprimer en watt. Cependant, étant donné la nature très différente de ces grandeurs, on utilise des unités séparées :

- $p(t)$, $p_f(t)$ et P s'expriment en watt, dont le symbole est W . Dans le cadre des réseaux d'énergie électrique, il est plus confortable d'exprimer les grandeurs en kilowatt (kW) et en mégawatt (MW)
- Q s'exprime en vars (abréviation pour volt ampère réactif), dont le symbole est VAr, Var ou var (nous retiendrons ce dernier). En pratique, on utilise plutôt le kvar et le Mvar
- S s'exprime en volt.ampère (VA). En pratique, on utilise plutôt le kVA et le MVA.

1.5 Puissance complexe

La *puissance complexe* est définie par :

$$S = \bar{V} \bar{I}^* \quad (1.21)$$

où $*$ désigne le conjugué d'un nombre complexe. En remplaçant $\bar{V}$ par (1.6) et $\bar{I}$ par (1.7) on trouve :

$$S = V e^{j\phi} I e^{-j\psi} = V I e^{j(\phi-\psi)} = V I \cos(\phi - \psi) + j V I \sin(\phi - \psi) = P + jQ$$

La partie réelle de la puissance complexe est donc la puissance active tandis que sa partie imaginaire est la puissance réactive. Le module de la puissance complexe vaut quant à lui :

$$|S| = \sqrt{P^2 + Q^2} = V I \quad (1.22)$$

c'est-à-dire la puissance apparente.

L'intérêt de la puissance complexe réside dans le fait que P et Q se calculent souvent plus aisément en passant par S .

Lorsque l'on travaille avec la puissance complexe, on utilise très souvent le

Théorème de conservation de la puissance complexe. Dans un circuit alimenté par des sources sinusoïdales fonctionnant toutes à la même fréquence, la somme des puissances complexes entrant dans n'importe quelle partie du circuit est égale à la somme des puissances complexes reçues par les branches de cette partie du circuit.

Appliqué à la figure 1.4, par exemple, ce théorème fournit la relation :

$$S_1 + S_2 + S_3 = \sum_i S_{bi}$$

où le membre de droite représente la somme des puissances complexes reçues par toutes les branches du circuit $\mathcal{C}$ (voir fig. 1.4). En décomposant en parties réelles et imaginaires, on obtient :

$$\begin{aligned} P_1 + P_2 + P_3 &= \sum_i P_{bi} \\ Q_1 + Q_2 + Q_3 &= \sum_i Q_{bi} \end{aligned}$$

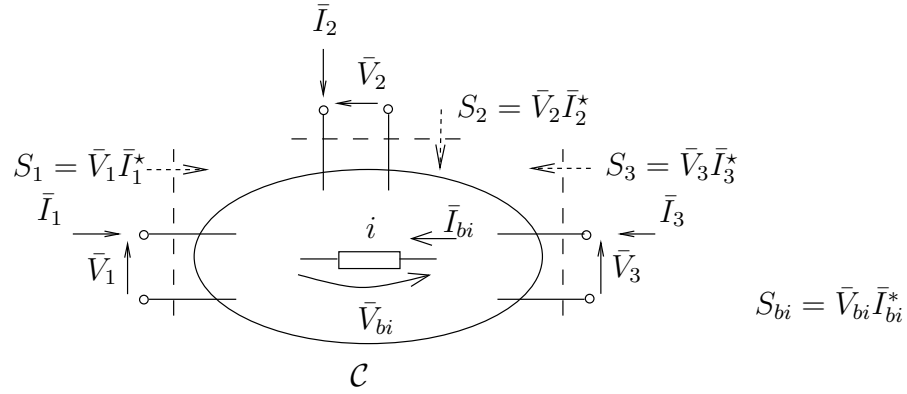


FIGURE 1.4 – illustration du théorème de la conservation de la puissance complexe

Dans le cas d'un réseau d'énergie électrique, on peut classer les branches en deux catégories : les éléments de transport (lignes, câbles, transformateurs, etc. . .) et les charges. Les relations ci-dessus expriment que la puissance fournie au réseau est égale à la consommation de toutes les charges, augmentée des pertes dans les éléments de transport.

Un tel équilibre des puissances est une notion naturelle en ce qui concerne la puissance instantanée : il traduit le principe de conservation de l'énergie, dont la puissance est la dérivée temporelle. Il est naturel qu'il s'applique à la puissance active, qui est la valeur moyenne de la puissance instantanée. Mais le fait plus remarquable est qu'il s'applique également à la puissance réactive, pour laquelle on donc peut écrire le bilan de puissance "production = consommation + pertes" comme pour la puissance active.

1.6 Expressions relatives aux dipôles

La table 1.1 donne les relations entre tension, courant et puissances pour un dipôle tandis que la table 1.2 donne les expressions des puissances actives et réactives consommées par les dipôles élémentaires. Dans les deux cas, on a considéré la convention moteur.

On notera qu'une inductance *consomme* de la puissance réactive, tandis qu'une capacité *produit*.

1.7 Facteur de puissance et compensation des charges

Considérons une charge alimentée par une source de tension (cf figure 1.5.a). Rappelons que la puissance active P correspond à la puissance utile consommée par la charge.

TABLE 1.1 – tension, courant et puissances dans un dipôle (convention moteur)

$\bar{V} = Z \bar{I} = (R + jX) \bar{I}$ Z : impédance R : résistance X : réactance	$\bar{I} = Y \bar{V} = (G + jB) \bar{V}$ Y : admittance G : conductance B : susceptance
$S = Z I^2$ $P = R I^2$ $Q = X I^2$	$S = Y^* V^2$ $P = G V^2$ $Q = -B V^2$

TABLE 1.2 – puissances absorbées par les dipôles élémentaires (convention moteur)

	résistance R	inductance L	capacité C
$\phi - \psi$	0	$\pi/2$	$-\pi/2$
P	$R I^2 = \frac{V^2}{R}$	0	0
Q	0	$\omega L I^2 = \frac{V^2}{\omega L}$	$-\frac{I^2}{\omega C} = -\omega C V^2$

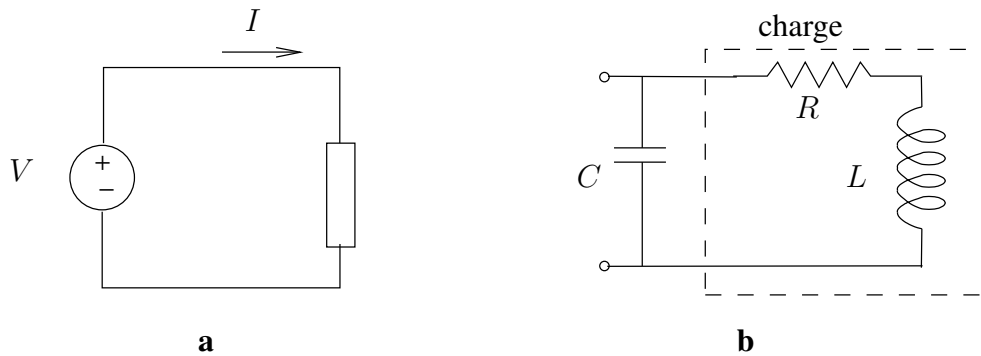


FIGURE 1.5 – compensation d’une charge pour amélioration de son facteur de puissance

De (1.17) on tire l’expression du courant parcourant le circuit :

$$I = \frac{P}{V \cos(\phi - \psi)}$$

Cette relation montre que, pour une même puissance utile P et sous une tension V constante, le courant augmente d’autant plus que $\cos(\phi - \psi)$ est faible.

On désigne $\cos(\phi - \psi)$ sous le vocable de *facteur de puissance*². Le facteur de puissance est d’autant plus faible que le courant est fortement déphasé par rapport à la tension. Dans le cas d’une charge résistive, le facteur de puissance est égal à l’unité.

2. Le facteur de puissance est très souvent désigné par le vocable “cosinus phi”, ce qui correspond au cas $\psi = 0$

On a d'ailleurs à partir de (1.22) :

$$I = \frac{\sqrt{P^2 + Q^2}}{V}$$

qui montre que pour une même puissance utile P et sous une tension V constante, le courant augmente avec la puissance réactive, consommée ou produite par la charge.

L'augmentation du courant I requiert d'utiliser des sections de conducteurs plus importantes, d'où un investissement plus important. Elle entraîne également des pertes RI^2 par effet Joule plus élevées dans les résistances des conducteurs traversés par le courant, d'où un coût de fonctionnement plus élevé.

Nous verrons ultérieurement que la consommation de puissance réactive entraîne également une chute des tensions, susceptible de gêner le bon fonctionnement de la charge.

La plupart des charges étant inductives (à cause de la présence de circuits magnétiques), donc consommatrices de puissance réactive, il y a intérêt à compenser ces dernières, c'est-à-dire à produire de la puissance réactive de sorte que l'ensemble présente un facteur de puissance aussi proche que possible de l'unité. Le moyen le plus simple consiste à brancher des condensateurs en parallèle sur la charge.

Considérons à titre d'exemple le cas d'une charge RL, comme représenté à la figure 1.5.b. Le facteur de puissance vaut :

$$\cos(\phi - \psi) = \frac{P}{\sqrt{P^2 + Q^2}} = \frac{RI^2}{\sqrt{R^2I^4 + \omega^2L^2I^4}} = \frac{R}{\sqrt{R^2 + \omega^2L^2}}$$

Pour avoir une compensation idéale, il faut que la puissance réactive Q_c produite par le condensateur égale la puissance réactive Q_ℓ consommée par la charge, soit :

$$\begin{aligned} Q_c &= -Q_\ell \\ \Leftrightarrow \omega CV^2 &= \frac{\omega L V^2}{R^2 + \omega^2 L^2} \\ \Leftrightarrow C &= \frac{L}{R^2 + \omega^2 L^2} \end{aligned}$$

Notons que si la charge varie au cours du temps, il est nécessaire d'adapter le volume de compensation de manière à conserver un facteur de puissance aussi proche que possible de l'unité. Ceci peut être réalisé en disposant plusieurs condensateurs en parallèle et en enclenchant le nombre adéquat.

Pour des charges variant très rapidement, il peut devenir difficile de déclencher/enclencher les condensateurs au moyen de disjoncteurs, condamnés à une usure prématurée. On peut alors faire appel à l'électronique de puissance.

Notons enfin qu'une surcompensation conduit à une augmentation du courant au même titre qu'une absence de compensation.

Chapitre 2

Systemes et régimes triphasés équilibrés

Si l'on excepte la présence de liaisons haute tension à courant continu, la quasi-totalité du transport et de la distribution d'énergie électrique est réalisée au moyen de systèmes triphasés. Les avantages principaux de ce système sont l'économie de conducteurs et la possibilité de générer des champs magnétiques tournants dans les générateurs et dans les moteurs.

Dans ce chapitre, nous rappelons le principe de fonctionnement d'un tel système, en régime équilibré, ainsi que les grandeurs et les relations qui le caractérisent.

2.1 Principe

Un circuit triphasé équilibré est constitué de trois circuits identiques, appelés *phases*. Le régime triphasé équilibré est tel que les tensions et les courants aux points des trois phases qui se correspondent sont de même amplitude mais décalés dans le temps d'un tiers de période d'une phase à l'autre.

La figure 2.1 donne un exemple de système triphasé qui pourrait représenter un générateur alimentant une charge par l'intermédiaire d'une ligne de transport que nous supposons idéale, pour simplifier. On a pour les tensions indiquées sur cette figure :

$$\begin{aligned}v_a(t) &= \sqrt{2}V \cos(\omega t + \phi) \\v_b(t) &= \sqrt{2}V \cos(\omega(t - \frac{T}{3}) + \phi) = \sqrt{2}V \cos(\omega t + \phi - \frac{2\pi}{3}) \\v_c(t) &= \sqrt{2}V \cos(\omega(t - \frac{2T}{3}) + \phi) = \sqrt{2}V \cos(\omega t + \phi - \frac{4\pi}{3})\end{aligned}$$

et pour les courants :

$$i_a(t) = \sqrt{2}I \cos(\omega t + \psi) \quad (2.1)$$

$$i_b(t) = \sqrt{2}I \cos(\omega(t - \frac{T}{3}) + \psi) = \sqrt{2}I \cos(\omega t + \psi - \frac{2\pi}{3}) \quad (2.2)$$

$$i_c(t) = \sqrt{2}I \cos(\omega(t - \frac{2T}{3}) + \psi) = \sqrt{2}I \cos(\omega t + \psi - \frac{4\pi}{3}) \quad (2.3)$$

relations dans lesquelles on a tenu compte de (1.4).

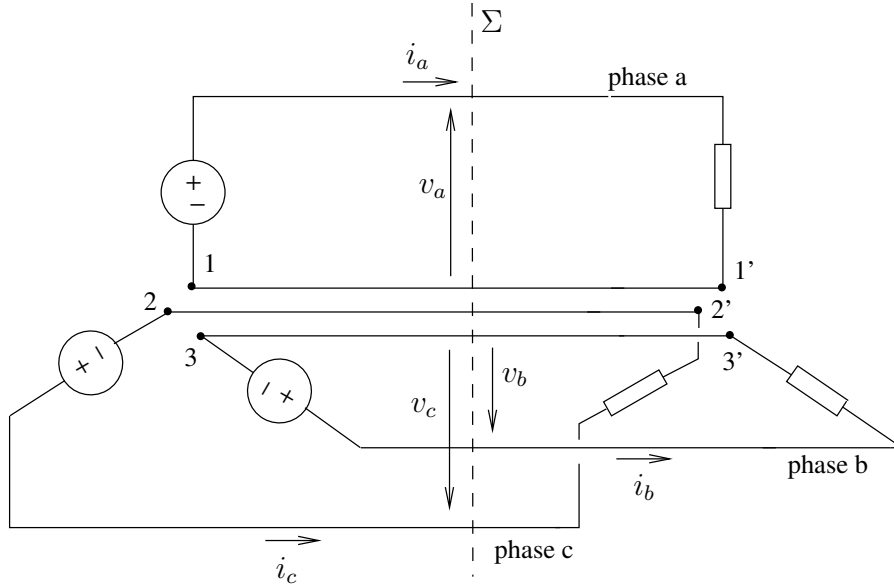


FIGURE 2.1 – circuit triphasé constitué de trois circuits monophasés

Les diagrammes de phaseur relatifs aux tensions et aux courants se présentent sous forme d'étoiles aux branches de même amplitude et déphasées l'une par rapport à l'autre de $2\pi/3$ radians (120 degrés), comme représenté à la figure 2.2. On a donc pour les tensions :

$$\begin{aligned} \bar{V}_a &= V e^{j\phi} \\ \bar{V}_b &= V e^{j(\phi - \frac{2\pi}{3})} = \bar{V}_a e^{-j\frac{2\pi}{3}} \\ \bar{V}_c &= V e^{j(\phi - \frac{4\pi}{3})} = \bar{V}_a e^{-j\frac{4\pi}{3}} = \bar{V}_b e^{-j\frac{2\pi}{3}} \end{aligned}$$

et pour les courants :

$$\begin{aligned} \bar{I}_a &= I e^{j\psi} \\ \bar{I}_b &= I e^{j(\psi - \frac{2\pi}{3})} = \bar{I}_a e^{-j\frac{2\pi}{3}} \\ \bar{I}_c &= I e^{j(\psi - \frac{4\pi}{3})} = \bar{I}_a e^{-j\frac{4\pi}{3}} = \bar{I}_b e^{-j\frac{2\pi}{3}} \end{aligned}$$

Il est clair que :

$$\bar{V}_a + \bar{V}_b + \bar{V}_c = 0 \quad (2.4)$$

$$\bar{I}_a + \bar{I}_b + \bar{I}_c = 0 \quad (2.5)$$

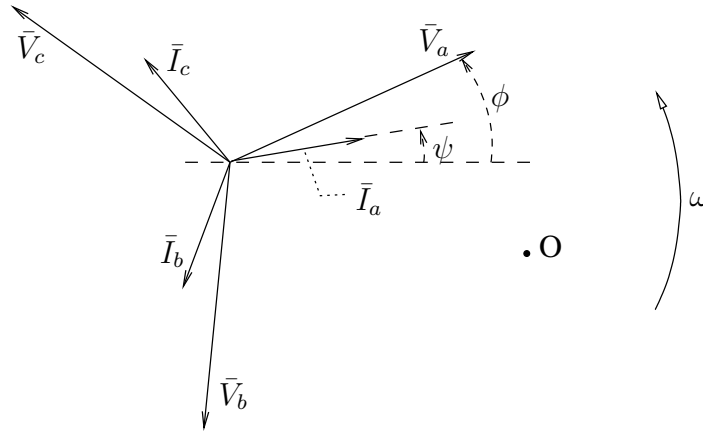


FIGURE 2.2 – diagramme de phaseur des tensions et courants en régime triphasé équilibré

Nous avons supposé que l'onde de tension de la phase b est en retard sur celle de la phase a et celle de la phase c en retard sur celle de la phase b . Dans le diagramme de la figure 2.2, un observateur placé en O voit passer les vecteurs tournants dans l'ordre a, b, c . On dit que les tensions $\bar{V}_a, \bar{V}_b, \bar{V}_c$ forment une *séquence directe*.

En fait, la configuration de la figure 2.1 présente peu d'intérêt. On peut obtenir un montage plus intéressant en regroupant les conducteurs de retour $11', 22'$ et $33'$ en un conducteur unique. Ce dernier est parcouru par le courant total $\bar{I}_a + \bar{I}_b + \bar{I}_c = 0$. On peut donc supprimer cette connexion sans modifier le fonctionnement du système, ce qui donne le circuit de la figure 2.3, typique des réseaux de transport à haute tension.

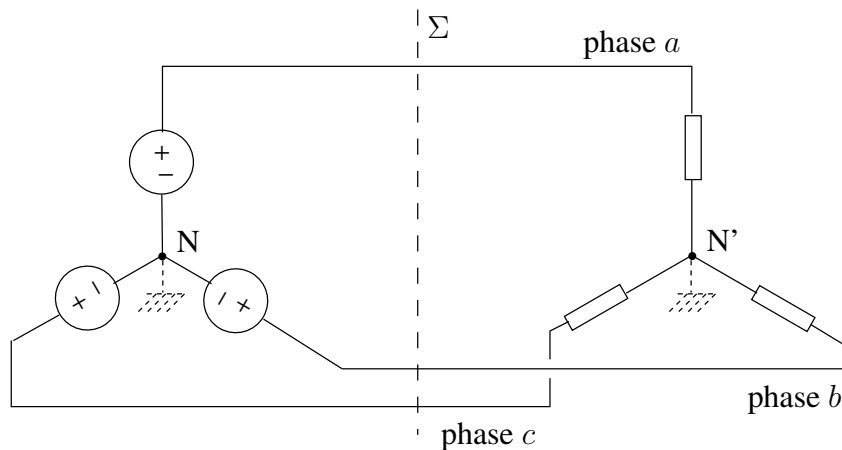


FIGURE 2.3 – un authentique circuit triphasé !

L'avantage du système triphasé de la figure 2.3 par rapport à un système monophasé est évident : la puissance transmise par le système triphasé à travers la coupe Σ vaut 3 fois celle transmise par une de ses phases, pour seulement 1,5 fois le nombre de conducteurs. De façon équivalente, le système triphasé de la figure 2.3 transporte autant de puissance que celui de la figure 2.1 mais avec moitié moins de conducteurs.

Les points tels que N et N' sont appelés *neutres*. En régime parfaitement équilibré, tous les neutres sont au même potentiel.

Les tensions $\bar{V}_a$, $\bar{V}_b$ ou $\bar{V}_c$ sont appelées *tensions de phase* ou *tensions phase-neutre*.

2.2 Tensions de ligne (ou composées)

Définissons à présent les différences :

$$\bar{U}_{ab} = \bar{V}_a - \bar{V}_b \quad (2.6)$$

$$\bar{U}_{bc} = \bar{V}_b - \bar{V}_c \quad (2.7)$$

$$\bar{U}_{ca} = \bar{V}_c - \bar{V}_a \quad (2.8)$$

Ces tensions sont appelées *tensions composées* ou *tensions entre phases* ou *tensions de ligne*.

Le diagramme de phaseur correspondant, représenté à la figure 2.4, fournit :

$$\bar{U}_{ab} = \sqrt{3} \bar{V}_a e^{j\frac{\pi}{6}} = \sqrt{3} V e^{j(\phi + \frac{\pi}{6})} \quad (2.9)$$

$$\bar{U}_{bc} = \sqrt{3} \bar{V}_b e^{j\frac{\pi}{6}} = \sqrt{3} V e^{j(\phi + \frac{\pi}{6} - \frac{2\pi}{3})} \quad (2.10)$$

$$\bar{U}_{ca} = \sqrt{3} \bar{V}_c e^{j\frac{\pi}{6}} = \sqrt{3} V e^{j(\phi + \frac{\pi}{6} - \frac{4\pi}{3})} \quad (2.11)$$

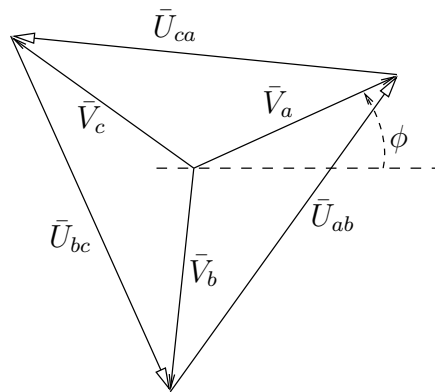


FIGURE 2.4 – tensions de phase et tensions de ligne

On voit que l'amplitude de la tension de ligne vaut $\sqrt{3}$ fois celle de la tension de phase et que $\bar{U}_{ab}$, $\bar{U}_{bc}$ et $\bar{U}_{ca}$ forment aussi une séquence directe.

Il est à noter qu'en pratique, quand on spécifie la tension d'un équipement triphasé, il s'agit, sauf mention contraire, de la *valeur efficace de la tension de ligne*. C'est le cas lorsque l'on parle, par exemple, d'un réseau à 380, 150, 70, etc... kV.

2.3 Connexions en étoile et en triangle

Il existe deux modes de connexion d'un équipement triphasé : en étoile ou en triangle, comme représenté à la figure 2.5.

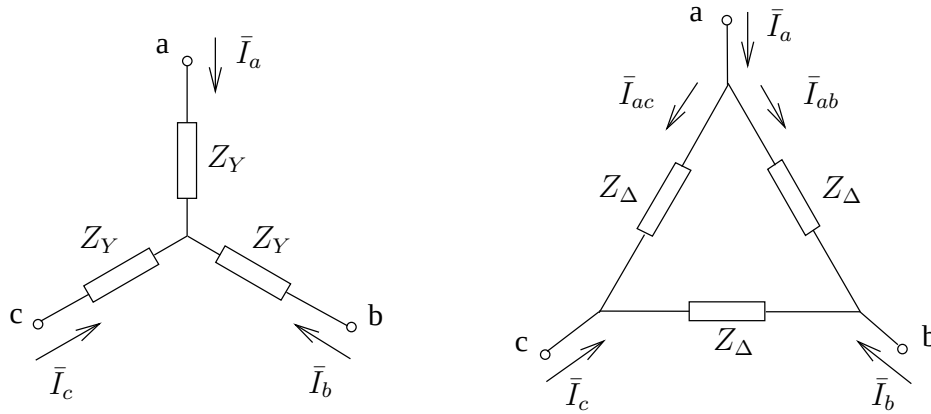


FIGURE 2.5 – connexion d'une charge triphasée en étoile et en triangle

Recherchons la relation entre les courants $\bar{I}_{ab}$ et $\bar{I}_a$ dans le montage en triangle. On a successivement :

$$\bar{I}_a = \bar{I}_{ab} + \bar{I}_{ac} = \frac{\bar{U}_{ab} + \bar{U}_{ac}}{Z_{\Delta}} = \frac{\bar{U}_{ab} - \bar{U}_{ca}}{Z_{\Delta}} = \frac{\bar{U}_{ab} - \bar{U}_{ab}e^{-j\frac{4\pi}{3}}}{Z_{\Delta}} = \frac{\bar{U}_{ab}}{Z_{\Delta}}(1 - e^{-j\frac{4\pi}{3}}) = \sqrt{3} e^{-j\frac{\pi}{6}} \bar{I}_{ab}$$

dont on tire évidemment :

$$\bar{I}_{ab} = \frac{1}{\sqrt{3}} e^{j\frac{\pi}{6}} \bar{I}_a \quad (2.12)$$

Le cours de Circuits électriques (et plus précisément la méthode par transfiguration) a montré que, si l'on applique les mêmes tensions de phase $\bar{V}_a$, $\bar{V}_b$ et $\bar{V}_c$ aux deux montages, les courants de phase $\bar{I}_a$, $\bar{I}_b$ et $\bar{I}_c$ sont identiques à condition que :

$$Z_{\Delta} = 3 Z_Y \quad (2.13)$$

Une charge alimentée sous tension monophasée doit donc être placée dans une branche d'étoile ou de triangle, selon la valeur de la tension en question.

Les distributeurs d'électricité veillent à connecter les différentes charges monophasées de manière à équilibrer les trois phases. C'est pourquoi il est raisonnable de considérer que les charges vues du réseau de transport sont équilibrées.

Au niveau d'une habitation alimentée en triphasé (380 V entre phases), les équipements monophasés fonctionnant sous 220 V sont placés entre phase et neutre. On veille à répartir les équipements (p.ex. les pièces d'habitation) sur les phases de la manière la plus équilibrée possible. Evidemment,

au niveau d'une habitation, il existe un déséquilibre. Les câbles d'alimentation sont dotés d'un conducteur de neutre et ce dernier est parcouru par un certain courant. Les neutres des différents consommateurs sont regroupés. Au fur et à mesure de ce groupement, le courant total de neutre devient négligeable devant les courants de phases. Notons que le câble d'alimentation peut être doté d'un cinquième conducteur, destiné à mettre les équipements à la terre.

Certaines charges, alimentées sous une tension sinusoïdale, produisent des harmoniques de courant. Ces derniers ont des effets indésirables telles que pertes supplémentaires, vibrations dans les machines, perturbations des équipements électroniques, etc... Il convient donc de prendre des mesures pour limiter leur propagation dans le réseau. Etant donné que dans un spectre de Fourier, l'énergie contenue dans une harmonique diminue quand le rang de cette harmonique (c'est-à-dire la fréquence) augmente, ce sont principalement les harmoniques de rang le plus bas qu'il faut supprimer (ou du moins atténuer).

La connexion des charges en triangle permet la suppression de certaines harmoniques, comme le montre l'exercice 2.

2.4 Analyse par phase

La symétrie qui existe entre les différentes phases permet de simplifier l'analyse d'un système triphasé équilibré. Il suffit en effet de déterminer tensions et courants dans une phase, pour obtenir automatiquement les tensions et courants dans les autres phases, par simple déphasage de $\pm 2\pi/3$ radians.

Pour pouvoir déterminer l'état électrique d'une phase en se passant des deux autres, deux opérations sont toutefois nécessaires :

- remplacer les charges connectées en triangle par leur schéma équivalent en étoile, en utilisant simplement la relation (2.13) ;
- s'affranchir des couplages inductifs et capacitifs entre phases. Cette opération simple est détaillée dans les deux sous-sections qui suivent.

2.4.1 Traitement des couplages inductifs entre phases

Considérons le circuit triphasé de la figure 2.6, dans lequel chaque phase possède une résistance, une self-inductance et un couplage inductif avec les autres phases.

Les tensions d'extrémité sont liées aux courants par :

$$\begin{bmatrix} \bar{V}_a \\ \bar{V}_b \\ \bar{V}_c \end{bmatrix} = \begin{bmatrix} \bar{V}_{a'} \\ \bar{V}_{b'} \\ \bar{V}_{c'} \end{bmatrix} + \begin{bmatrix} Z_s & Z_m & Z_m \\ Z_m & Z_s & Z_m \\ Z_m & Z_m & Z_s \end{bmatrix} \begin{bmatrix} \bar{I}_a \\ \bar{I}_b \\ \bar{I}_c \end{bmatrix}$$

relation dans laquelle on a supposé (idéalement) un parfait équilibre entre les phases (même

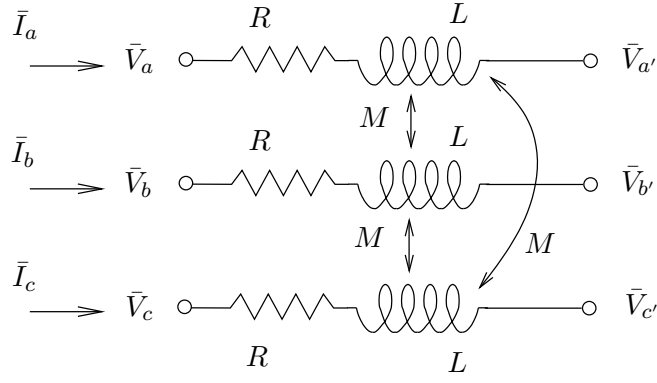


FIGURE 2.6 – couplage inductif entre phases

terme diagonal dans chaque phase et même terme non-diagonal quelle que soit la paire de phases considérée).

La première composante de cette relation matricielle donne :

$$\bar{V}_a = \bar{V}_{a'} + Z_s \bar{I}_a + Z_m \bar{I}_b + Z_m \bar{I}_c$$

et en tenant compte du fait que le régime est équilibré :

$$\begin{aligned} \bar{V}_a &= \bar{V}_{a'} + \left[Z_s + Z_m(e^{-j\frac{2\pi}{3}} + e^{-j\frac{4\pi}{3}}) \right] \bar{I}_a \\ &= \bar{V}_{a'} + [Z_s - Z_m] \bar{I}_a \end{aligned}$$

Tout se passe donc comme si la phase a était seule mais présentait une impédance :

$$Z_{eq} = Z_s - Z_m \quad (2.14)$$

qui est appelée *impédance cyclique*. Insistons sur le fait que ce résultat n'est valable qu'en régime équilibré !

Dans le cas représenté à la figure 2.6, on trouve aisément que :

$$Z_{eq} = R + j\omega(L - M) \quad (2.15)$$

Le schéma équivalent “par phase” est donc celui de la figure 2.7.a. Dans ce schéma, la première loi de Kirchhoff impose un courant de retour $\bar{I}_a$. Comme on l'a dit plus haut, celui-ci n'existe pas dans le circuit triphasé.

2.4.2 Traitement des couplages capacitifs entre phases

Considérons à présent le circuit de la figure 2.8, dans lequel chaque phase possède un couplage capacitif avec la terre (capacité C) et avec les autres phases (capacité C_m).

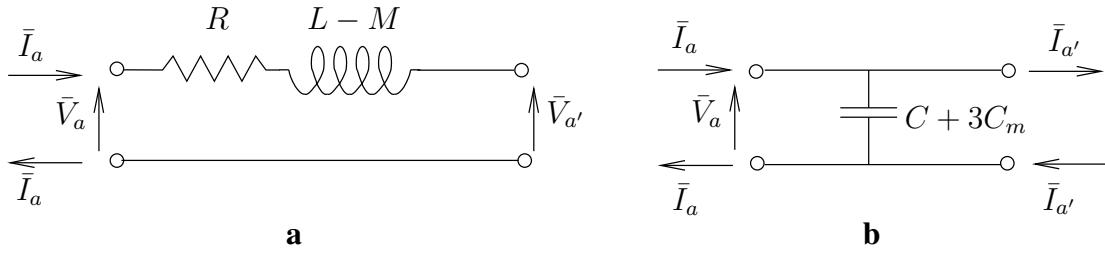


FIGURE 2.7 – schémas équivalents par phase des circuits des figures 2.6 et 2.8

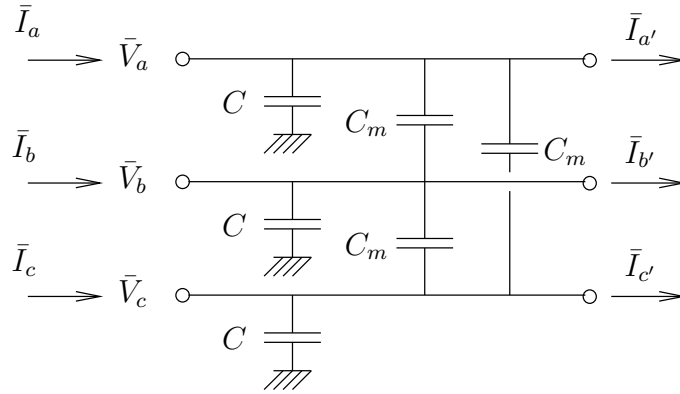


FIGURE 2.8 – couplage capacitif entre phases et avec la terre

La relation entre courants et tensions est :

$$\begin{bmatrix} \bar{I}_a - \bar{I}_{a'} \\ \bar{I}_b - \bar{I}_{b'} \\ \bar{I}_c - \bar{I}_{c'} \end{bmatrix} = \begin{bmatrix} Y_s & Y_m & Y_m \\ Y_m & Y_s & Y_m \\ Y_m & Y_m & Y_s \end{bmatrix} \begin{bmatrix} \bar{V}_a \\ \bar{V}_b \\ \bar{V}_c \end{bmatrix}$$

Notons ici encore l'hypothèse de parfaite symétrie triphasée. La première composante de cette relation matricielle donne :

$$\bar{I}_a - \bar{I}_{a'} = Y_s \bar{V}_a + Y_m \bar{V}_b + Y_m \bar{V}_c$$

et en tenant compte du fait que le régime est équilibré :

$$\begin{aligned} \bar{I}_a - \bar{I}_{a'} &= \left[Y_s + Y_m (e^{-j\frac{2\pi}{3}} + e^{-j\frac{4\pi}{3}}) \right] \bar{V}_a \\ &= [Y_s - Y_m] \bar{V}_a \end{aligned}$$

On voit que tout se passe comme si la phase a était seule mais avec une admittance entre phase et neutre :

$$Y_{eq} = Y_s - Y_m \quad (2.16)$$

Dans le cas représenté à la figure 2.8, on a $Y_s = j\omega C + 2j\omega C_m$ et $Y_m = -j\omega C_m$ et donc :

$$Y_{eq} = j\omega(C + 3C_m) \quad (2.17)$$

Le schéma équivalent par phase est donc celui de la figure 2.7.b.

2.4.3 Schéma unifilaire

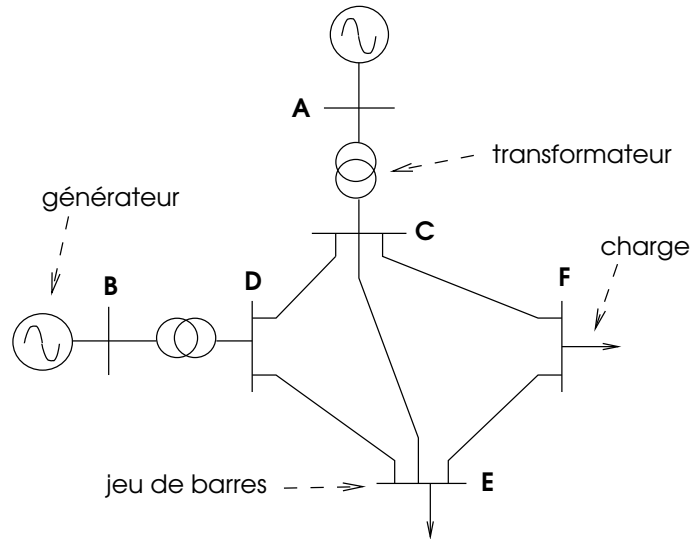


FIGURE 2.9 – schéma unifilaire d'un système de puissance

L'analyse par phase se concrétise en particulier dans l'utilisation du *schéma unifilaire*. Il s'agit d'un diagramme monophasé, sans conducteur de retour, représentant les équipements qui composent un système de puissance. Un exemple est donné à la figure 2.9.

Les équipements tels que lignes, câbles, transformateurs, générateurs, charges, etc... sont reliés entre eux, dans les postes à haute tension, par l'intermédiaire de barres conductrices. Une barre est considérée comme un équipement équipotentiel. L'ensemble des trois barres relatives aux trois phases est appelé un *jeu de barres*. Les jeux de barres du système de la figure 2.9 sont A, B, ..., F.

2.5 Puissances en régime triphasé

La puissance instantanée traversant la coupe Σ des figures 2.1 et 2.3 vaut :

$$\begin{aligned}
 p(t) &= v_a i_a + v_b i_b + v_c i_c \\
 &= 2VI \left[\cos(\omega t + \phi) \cos(\omega t + \psi) + \cos\left(\omega t + \phi - \frac{2\pi}{3}\right) \cos\left(\omega t + \psi - \frac{2\pi}{3}\right) + \right. \\
 &\quad \left. + \cos\left(\omega t + \phi - \frac{4\pi}{3}\right) \cos\left(\omega t + \psi - \frac{4\pi}{3}\right) \right] \\
 &= 3VI \cos(\phi - \psi) \\
 &\quad + VI \left[\cos(2\omega t + \phi + \psi) + \cos\left(2\omega t + \phi + \psi - \frac{4\pi}{3}\right) + \cos\left(2\omega t + \phi + \psi - \frac{2\pi}{3}\right) \right] \\
 &= 3VI \cos(\phi - \psi) = 3P
 \end{aligned}$$

On voit que la puissance instantanée est une constante, égale à trois fois la puissance active P transférée à une des phases. Il n'y a donc pas de puissance fluctuante en régime triphasé équilibré.

Puisque la puissance réactive a été définie comme l'amplitude d'un des termes de la puissance fluctuante (cf (1.15,1.18)), on pourrait penser que la notion de puissance réactive n'est pas appelée à jouer un rôle en régime triphasé équilibré. Il n'en est rien. En fait, dans chaque phase, il y a une puissance fluctuante ; une de ses composantes correspond à l'énergie emmagasinée dans les bobines et les condensateurs de cette phase et son amplitude est la puissance réactive Q relative à la phase considérée. Simplement, les puissances fluctuantes des différentes phases sont décalées temporellement d'un tiers de période, de sorte que leur somme est nulle à tout instant.

La *puissance complexe triphasée* vaut, par extension de la formule monophasée :

$$S_{3\phi} = \bar{V}_a \bar{I}_a^* + \bar{V}_b \bar{I}_b^* + \bar{V}_c \bar{I}_c^* = \bar{V}_a \bar{I}_a^* + \bar{V}_a e^{-j\frac{2\pi}{3}} \bar{I}_a^* e^{j\frac{2\pi}{3}} + \bar{V}_a e^{-j\frac{4\pi}{3}} \bar{I}_a^* e^{j\frac{4\pi}{3}} = 3\bar{V}_a \bar{I}_a^*$$

La partie réelle de $S_{3\phi}$ est la *puissance active triphasée* :

$$P_{3\phi} = 3VI \cos(\phi - \psi) = 3P \quad (2.18)$$

tandis que la partie imaginaire est la *puissance réactive triphasée* :

$$Q_{3\phi} = 3VI \sin(\phi - \psi) = 3Q \quad (2.19)$$

La notion de puissance réactive triphasée est artificielle dans la mesure où il n'y a pas de puissance fluctuante triphasée. En fait, seule la puissance réactive par phase Q a une interprétation. $Q_{3\phi} = 3Q$ est une grandeur aussi artificielle que le serait un "courant triphasé $3I$ ". Cependant, cette notion est universellement utilisée, pour des raisons de symétrie avec la puissance active.

En vertu de (2.9), on a :

$$P_{3\phi} = \sqrt{3}UI \cos(\phi - \psi) \quad (2.20)$$

$$Q_{3\phi} = \sqrt{3}UI \sin(\phi - \psi) \quad (2.21)$$

où U est la valeur efficace de la tension de ligne. Ces formules sont souvent utilisées parce qu'elle font intervenir U , elle-même utilisée pour désigner la tension. Notons toutefois que ces formules sont hybrides dans la mesure où $\phi - \psi$ est le déphasage entre le courant et la tension de phase (et non la tension de ligne).

Chapitre 3

Quelques propriétés du transport de l'énergie électrique

Dans ce chapitre nous montrons quelques propriétés fondamentales du fonctionnement des réseaux d'énergie électrique en régime établi et nous introduisons quelques notions importantes.

3.1 Transit de puissance et chute de tension dans une liaison

3.1.1 Modèle et relations principales

Considérons le système simple de la figure 3.1. Il comporte deux jeux de barres (ou noeuds électriques, ou simplement noeuds) reliés par une ligne ou un câble, dont nous supposons que le schéma par phase consiste simplement en une résistance R en série avec une réactance X . Les conclusions qui suivent restent d'application quand on inclut des capacités shunt dans le modèle de la ligne (voir chapitre 4) ou si l'on considère un transformateur à la place de la ligne (cf chapitre 6).

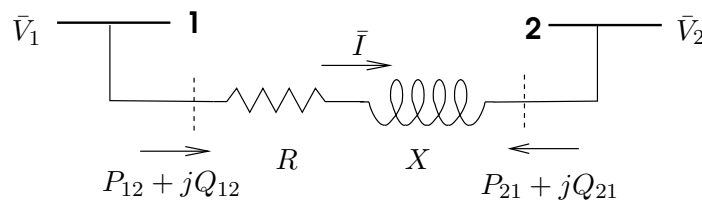


FIGURE 3.1 – système simple à deux jeux de barres

Par un choix approprié de l'origine des temps, on peut supposer que le phaseur de la tension au

noeud 1 a une phase nulle. Posons :

$$\bar{V}_1 = V_1 e^{j0} = V_1 \quad \text{et} \quad \bar{V}_2 = V_2 e^{j\theta_2} = V_2 \angle \theta_2$$

Soit $\bar{I}$ le courant parcourant la ligne. Soient P_{12} et Q_{12} les puissances active et réactive *par phase* entrant dans la ligne par le noeud 1 (cf figure 3.1). On a évidemment :

$$\bar{V}_2 = \bar{V}_1 - (R + jX)\bar{I} \quad (3.1)$$

Il y correspond le diagramme de phaseur de la figure 3.2.

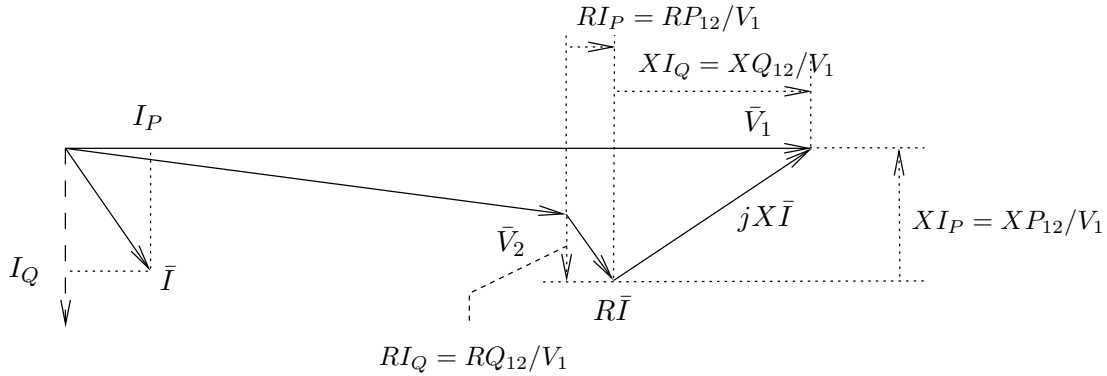


FIGURE 3.2 – diagramme de phaseur relatif au système de la figure 3.1

Etablissons l'expression de la tension $\bar{V}_2$ en fonction de la tension $\bar{V}_1$ et des transits de puissance P_{12} et Q_{12} . On a :

$$P_{12} + jQ_{12} = \bar{V}_1 \bar{I}^* \quad (3.2)$$

d'où l'on tire :

$$\bar{I} = \frac{P_{12} - jQ_{12}}{\bar{V}_1^*} = \frac{P_{12} - jQ_{12}}{V_1}$$

En introduisant cette dernière relation dans (3.1), on obtient :

$$\bar{V}_2 = \bar{V}_1 - (R + jX) \left[\frac{P_{12} - jQ_{12}}{V_1} \right] = \bar{V}_1 - \frac{RP_{12} + XQ_{12}}{V_1} - j \frac{XP_{12} - RQ_{12}}{V_1} \quad (3.3)$$

On peut retrouver ce résultat au départ de la figure 3.2, en considérant que :

- la projection de $\bar{I}$ sur $\bar{V}_1$ est le courant actif $I_P = P_{12}/V_1$
- la projection de $\bar{I}$ sur la perpendiculaire à $\bar{V}_1$ est le courant réactif $I_Q = Q_{12}/V_1$.

3.1.2 Effet du transport de puissance active et réactive

Comme nous le verrons, dans les réseaux de transport à Très Haute Tension (THT), la valeur de la résistance R est faible devant celle de la réactance X . Il est à noter que cette simplification ne s'applique pas aux réseaux de distribution à Moyenne Tension (MT) où R est du même ordre de grandeur que X . Si l'on suppose donc $R = 0$, la relation (3.3) se simplifie en :

$$\bar{V}_2 = \bar{V}_1 - \frac{XQ_{12}}{V_1} - j \frac{XP_{12}}{V_1} \quad (3.4)$$

Le diagramme de phaseur correspondant est donné à la figure 3.3.

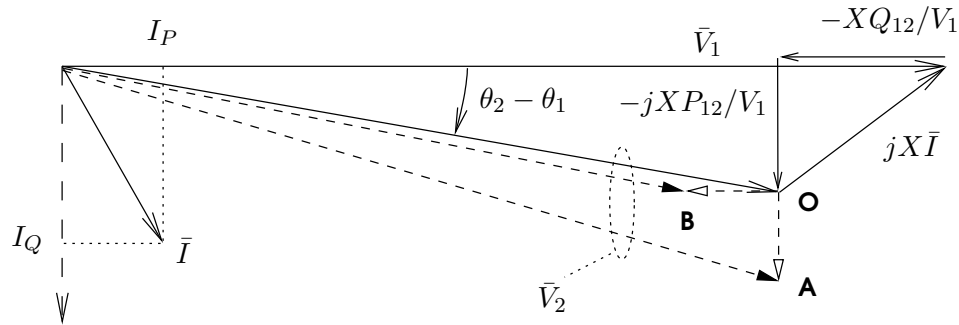


FIGURE 3.3 – diagramme de phaseur de la fig. 3.2 quand $R = 0$

Cette figure montre de plus la variation de la tension $\bar{V}_2$ sous l'effet de variations supplémentaires de la puissance active (passage du point O au point A) et de la puissance réactive (passage de O en B), la tension V_1 étant supposée constante. On peut en conclure que :

- le transfert de puissance active crée une chute de tension en quadrature avec $\bar{V}_1$. Si l'on suppose, comme c'est le cas en pratique, que $\|\bar{V}_2 - \bar{V}_1\|$ est faible devant V_1 , on peut conclure que *le transport de puissance active induit principalement un déphasage des tensions* ;
- le transfert de puissance réactive crée une chute de tension en phase avec $\bar{V}_1$. On peut en conclure que *le transport de puissance réactive induit principalement une chute des (modules des) tensions*.

3.1.3 Transport de puissance réactive à longue distance

Dans les réseaux de transport à THT, il est d'usage de dire que *la puissance réactive ne se transporte pas aisément sur de longues distances*. Ce fait peut être illustré comme suit sur notre exemple à deux noeuds.

Le bilan de puissance complexe de la liaison fournit :

$$P_{12} = -P_{21} + RI^2 \quad (3.5)$$

$$Q_{12} = -Q_{21} + XI^2 \quad (3.6)$$

Comme X est beaucoup plus grand que R , on voit que les pertes réactives sont nettement plus élevées que les pertes actives. Ainsi, si les puissances active et réactive entrent en quantités égales dans la liaison, il sort à l'autre extrémité nettement moins de puissance réactive que de puissance active.

Par ailleurs, nous venons de voir que le transfert de puissance réactive va de pair avec une variation des (modules des) tensions. Transférer beaucoup de puissance réactive requiert des chutes de tension importantes. En pratique, ceci n'est pas acceptable car les tensions aux différents noeuds d'un réseau doivent rester dans une plage de quelques pour-cents autour des valeurs nominales, sous peine de fonctionnement incorrect des matériels.

Une telle limitation n'existe par pour la puissance active car le déphasage des tensions n'a pas de conséquence directe pour les équipements.

3.1.4 Expressions des transits en fonction des tensions

Etablissons à présent l'expression des puissances P_{12} et Q_{12} en fonction des modules et des phases des tensions aux extrémités. Pour plus de généralité, nous considérerons le cas où $\theta_1 \neq 0$.

On obtient à partir des relations (3.1,3.2) :

$$\begin{aligned} P_{12} + jQ_{12} &= \bar{V}_1 \bar{I}^* = \bar{V}_1 \frac{\bar{V}_1^* - \bar{V}_2^*}{R - jX} = V_1 e^{j\theta_1} \frac{V_1 e^{-j\theta_1} - V_2 e^{-j\theta_2}}{R - jX} = \frac{V_1^2 - V_1 V_2 e^{j(\theta_1 - \theta_2)}}{R - jX} \\ &= \frac{[V_1^2 - V_1 V_2 \cos(\theta_1 - \theta_2) - jV_1 V_2 \sin(\theta_1 - \theta_2)] (R + jX)}{R^2 + X^2} \end{aligned}$$

En développant le numérateur et en égalant parties réelle et imaginaire des deux membres, on obtient les relations recherchées :

$$P_{12} = V_1^2 \frac{R}{R^2 + X^2} - V_1 V_2 \left[\frac{R}{R^2 + X^2} \cos(\theta_1 - \theta_2) - \frac{X}{R^2 + X^2} \sin(\theta_1 - \theta_2) \right] \quad (3.7)$$

$$Q_{12} = V_1^2 \frac{X}{R^2 + X^2} - V_1 V_2 \left[\frac{X}{R^2 + X^2} \cos(\theta_1 - \theta_2) + \frac{R}{R^2 + X^2} \sin(\theta_1 - \theta_2) \right] \quad (3.8)$$

Par simple permutation des indices 1 et 2, on obtient l'expression des puissances *entrant* dans la ligne du côté du noeud 2 :

$$\begin{aligned} P_{21} &= V_2^2 \frac{R}{R^2 + X^2} - V_2 V_1 \left[\frac{R}{R^2 + X^2} \cos(\theta_2 - \theta_1) - \frac{X}{R^2 + X^2} \sin(\theta_2 - \theta_1) \right] \\ Q_{21} &= V_2^2 \frac{X}{R^2 + X^2} - V_2 V_1 \left[\frac{X}{R^2 + X^2} \cos(\theta_2 - \theta_1) + \frac{R}{R^2 + X^2} \sin(\theta_2 - \theta_1) \right] \end{aligned}$$

Sous l'hypothèse $R = 0$, les relations ci-dessus deviennent simplement :

$$P_{12} = \frac{V_1 V_2 \sin(\theta_1 - \theta_2)}{X} \quad (3.9)$$

$$Q_{12} = \frac{V_1^2 - V_1 V_2 \cos(\theta_1 - \theta_2)}{X} \quad (3.10)$$

$$P_{21} = \frac{V_2 V_1 \sin(\theta_2 - \theta_1)}{X} \quad (3.11)$$

$$Q_{21} = \frac{V_2^2 - V_2 V_1 \cos(\theta_2 - \theta_1)}{X} \quad (3.12)$$

Ces relations sont utilisées dans de très nombreux raisonnements.

Rappelons que P_{12} et Q_{12} sont des puissances *par phase*. La puissance triphasée s'obtient en multipliant ces relations par un facteur trois.

3.2 Caractéristique QV à un jeu de barres d'un réseau

Dans cette section, nous nous intéressons à la relation entre la puissance réactive Q injectée en un jeu de barres et la tension V à celui-ci, toute autre chose restant constante. Choisissons de compter Q positif quand la puissance entre dans le réseau. Le développement qui suit est limité à une seule source de puissance réactive et ne rend pas compte des interactions entre deux sources voisines.

Dans une certaine plage de variation, on peut représenter un réseau vu d'un de ses jeux de barres par un schéma équivalent de Thévenin (cf figure 3.4.a). Rappelons le

Théorème de Thévenin. Vu d'un accès, un circuit linéaire peut être remplacé par un schéma équivalent composé d'une source de tension en série avec une impédance. La f.e.m. de la source équivalente est la tension apparaissant à vide à l'accès considéré. L'impédance équivalente est l'impédance vue de l'accès après avoir passifié le circuit, c'est-à-dire avoir annulé les f.e.m. (resp. les courants) des sources de tension (resp. de courant) indépendantes.

Nous supposons que l'impédance de Thévenin est essentiellement inductive, hypothèse déjà discutée. Quant à la f.e.m. de Thévenin, dans le cas qui nous occupe, c'est la tension relevée au jeu de barres lorsqu'aucune puissance n'y est produite ni consommée.

Considérons à présent l'injection d'une puissance réactive Q en ce jeu de barres. Comme aucune puissance active n'est injectée, la relation (3.9) montre qu'il n'y a pas de déphasage entre la tension du jeu de barres et la f.e.m. de Thévenin, tandis que (3.10) fournit l'expression de la puissance réactive Q entrant *par phase* dans l'équivalent :

$$Q = \frac{V^2 - V E_{th}}{X_{th}} \quad (3.13)$$

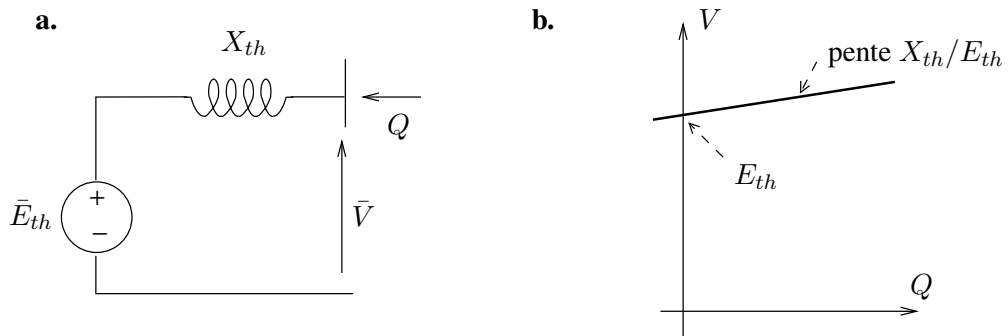


FIGURE 3.4 – schéma équivalent de Thévenin et caractéristique QV d'un réseau

Sous les hypothèses adoptées plus haut, l'équation (3.13) nous indique que la relation entre Q et V est quadratique. Cependant, pour des variations de tension suffisamment faibles autour de E_{th} , cette relation peut être linéarisée. Le coefficient angulaire de la droite correspondante (cf figure 3.4.b) est donné par :

$$\frac{1}{\left(\frac{\partial Q}{\partial V}\right)_{V=E_{th}}} = \frac{X_{th}}{2V - E_{th}} \bigg|_{V=E_{th}} = \frac{X_{th}}{E_{th}}$$

On voit donc que, suite à des variations de la puissance réactive en un jeu de barres, les variations de tension y sont d'autant plus faibles que la réactance de Thévenin vue de ce jeu de barres est faible.

La représentation d'un ensemble aussi complexe qu'un système d'énergie électrique par un simple schéma équivalent de Thévenin est évidemment une abstraction assez forte. Des remarques s'imposent à ce sujet :

- les résultats ci-dessus ne sont pas valables pour de grandes variations de V et/ou de Q . En effet, dans ce cas, la caractéristique n'est plus linéaire, non seulement à cause de la relation (3.13) mais surtout à cause du passage en limite de production réactive des générateurs (voir chapitre 11), ce qui modifie les paramètres de Thévenin ;
- après une variation brutale de la puissance réactive, la réactance de Thévenin varie dans le temps. Ceci provient du fait que le système le siège de dynamiques provenant de ses composants et de ses régulations. La réactance de Thévenin vue dans les tout premiers instants doit être calculée en tenant compte du comportement des composants (surtout les générateurs : voir chapitre 12) ; elle diffère de la réactance de Thévenin qui caractérise le passage d'un point de fonctionnement en régime établi à un autre ;
- lorsqu'un réseau perd un de ses composants (ligne, transformateur, générateur), les paramètres de Thévenin se modifient. Dans de nombreux cas, E_{th} diminue et X_{th} augmente suite à un tel incident.

3.3 Puissance de court-circuit

La notion de *puissance de court-circuit* est très utilisée dans l'analyse des réseaux d'énergie électrique. Elle est définie par :

$$S_{cc} = 3V_N I_{cc} = \sqrt{3}U_N I_{cc} \quad (3.14)$$

où V_N est la valeur nominale de la tension de phase, U_N celle de la tension de ligne et I_{cc} le courant circulant dans (chaque phase d')un court-circuit triphasé sans impédance au jeu de barres considéré.

Notons que S_{cc} ne représente pas une puissance au sens physique du terme. En effet, les grandeurs intervenant dans cette formule ne se rapportent pas à la même configuration, puisque V_N est la tension *avant* court-circuit et I_{cc} le courant *pendant* le court-circuit.

La puissance de court-circuit est utilisée dans le dimensionnement des disjoncteurs. En effet :

- un disjoncteur doit être capable d'éteindre l'arc électrique qui apparaît entre ses contacts au fur et à mesure que ceux-ci s'éloignent l'un de l'autre. Plus le courant de court-circuit I_{cc} est élevé, plus le disjoncteur doit être puissant pour éteindre cet arc ;
- une fois le courant interrompu, le disjoncteur doit être capable de tenir la tension qui se rétablit à ces bornes sans qu'il y ait rupture diélectrique du gaz situé entre ses contacts. Cette tension est d'autant plus élevée que V_N est élevé.

Il est donc raisonnable de dimensionner un disjoncteur sur la base du produit de I_{cc} et de V_N .

Il existe une relation simple entre la puissance de court-circuit en un jeu de barres et le schéma équivalent de Thévenin du réseau vu de ce jeu de barres. En effet, si l'on suppose que la tension au jeu de barres avant court-circuit est égale à la tension nominale V_N , c'est également la valeur de la f.e.m. de Thévenin et l'amplitude du courant de court-circuit est donné par :

$$I_{cc} = \frac{V_N}{|Z_{th}|} \quad (3.15)$$

Il en résulte que la puissance de court-circuit est donnée par :

$$S_{cc} = 3 \frac{V_N^2}{|Z_{th}|} = \frac{U_N^2}{|Z_{th}|} \quad (3.16)$$

La puissance de court-circuit donne également une indication sur la tenue de la tension en un jeu de barres. En effet, plus S_{cc} est élevée, plus $|Z_{th}|$ est faible et, comme on l'a vu à la section précédente, plus les variations de la tension avec la puissance réactive sont faibles. C'est pourquoi il importe que des charges fluctuant rapidement soient connectées à des jeux de barres où la puissance de court-circuit est suffisamment élevée.

Lorsque l'impédance de Thévenin tend vers zéro, la puissance de court-circuit tend vers l'infini. A la limite, on parle de *jeu de barres infini*.

3.4 Puissance maximale transmissible à une charge

Certains des développements qui précèdent supposaient une variation faible des puissances autour d'un point de fonctionnement "normal". Dans cette section, par contre, nous considérons de plus grandes variations et recherchons une limite de fonctionnement. Celle-ci correspond à la puissance maximale transmissible à une charge, notion à laquelle il convient de prêter une grande attention dans certains réseaux.

3.4.1 Modélisation

Considérons le système simple à deux noeuds de la figure 3.5, dans lequel un générateur alimente une charge via une ligne représentée simplement par sa réactance série X . Le jeu de barres du générateur est pris comme référence des phases.

Notons P (resp. Q) la puissance active (resp. réactive) *par phase* consommée par la charge. Les équations (3.12, 3.12) s'écrivent :

$$-P = VV_g \frac{1}{X} \sin \theta \quad \Leftrightarrow \quad P = -\frac{VV_g}{X} \sin \theta \quad (3.17)$$

$$-Q = \frac{1}{X} V^2 - \frac{1}{X} VV_g \cos \theta \quad \Leftrightarrow \quad Q = -\frac{V^2}{X} + \frac{VV_g}{X} \cos \theta \quad (3.18)$$

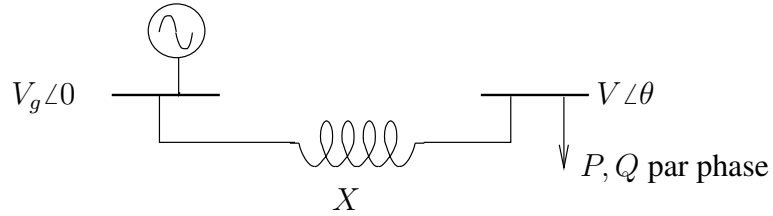


FIGURE 3.5 – système simple à 2 noeuds

où V_g est le module de la tension au noeud générateur et $V \angle \theta$ est la tension au noeud charge.

Dans ce qui suit, nous supposons que la limite thermique de la ligne n'est pas atteinte et que le générateur est capable de fournir les puissances active et réactive qui lui sont demandées.

3.4.2 Conditions d'existence d'une solution

Les équations (3.17, 3.17) sont très simples et peuvent être résolues analytiquement.

Considérons d'abord les conditions sous lesquelles elles ont une solution. En éliminant θ de ces relations et en regroupant les termes, on trouve :

$$(V^2)^2 + (2QX - V_g^2)V^2 + X^2(P^2 + Q^2) = 0 \quad (3.19)$$

En posant $V^2 = y$, cette dernière relation prend la forme d'une équation du second degré en y . Pour avoir (au moins) une solution, le discriminant doit être positif. Après calcul, cette condition s'écrit :

$$-\left(\frac{PX}{V_g^2}\right)^2 - \frac{QX}{V_g^2} + 0.25 \geq 0 \quad (3.20)$$

Dans le plan (P, Q) , la courbe correspondant à l'égalité dans (3.20) est une parabole d'axe vertical, comme le montre la figure 3.6, où l'on a considéré les grandeurs adimensionnelles PX/V_g^2 et QX/V_g^2 plutôt que P et Q .

Si le point (P, Q) se situe "en dessous" de la parabole, le discriminant est positif et l'équation en y a deux solutions. Si le point (P, Q) se situe "au-dessus" de la parabole, l'équation n'a pas de solution.

Un point de la parabole situé dans le premier quadrant ($P, Q > 0$) correspond à une puissance maximale transmissible à une charge inductive, tandis qu'un point situé dans le quatrième quadrant ($P > 0, Q < 0$) correspond à une puissance maximale transmissible à une charge capacitive.

Cette puissance maximale est intimement liée au fonctionnement en courant alternatif. En fait, elle correspond à un résultat bien connu de la Théorie des circuits : l'existence d'une charge maximale extractible d'un dipôle (dans le cas présent, l'impédance jX en série avec la source de tension V_g). Toutefois, en Théorie des circuits, le problème de l'adaptation de la charge est formulé différemment. On suppose que le dipôle présente une impédance $Z = R + jX$. La charge

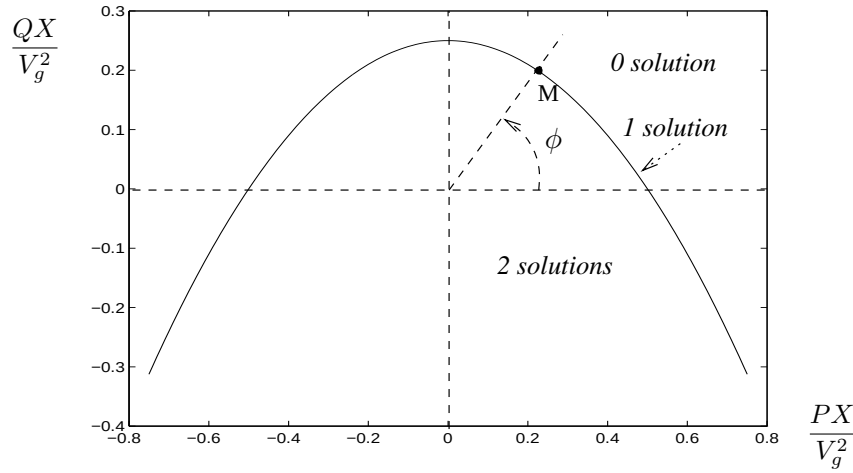


FIGURE 3.6 – région du plan (P, Q) où il existe une solution au problème

“adaptée”, qui tire le maximum de puissance active du dipôle, est l’impédance $Z^* = R - jX$. Ce résultat ne s’applique pas au cas particulier où R est nul. Dans ce cas, il est possible de tirer une puissance infinie du dipôle (si l’on fait abstraction des limites physiques, bien entendu). Dans le développement ci-dessus, on suppose précisément $R = 0$ et la figure 3.6 confirme qu’on peut tirer une puissance active infinie.

Dans le contexte d’un réseau électrique, il est plus réaliste de rechercher la puissance active maximale, *sous contrainte d’un facteur de puissance $\cos \phi$ donné*. Dans ces conditions, les puissances active et réactive sont liées par :

$$Q = P \operatorname{tg} \phi \quad (3.21)$$

Au fur et à mesure que la charge augmente, dans le plan (P, Q) , le point de fonctionnement se déplace sur le segment de droite en pointillé à la figure 3.6, jusqu’à atteindre le point M au delà duquel le fonctionnement n’est plus possible.

La partie de la parabole située dans les deuxième et troisième quadrants correspond évidemment à une production de puissance active au noeud PQ. On voit qu’il existe également une limite de puissance transmissible.

La figure révèle une dissymétrie fondamentale entre puissances active et réactive : en théorie, il est possible d’atteindre n’importe quelle valeur de P , éventuellement au prix d’une production de puissance réactive ($Q < 0$), alors que la puissance réactive ne peut jamais dépasser la valeur $V_g^2/4X$. L’équation (3.17) montre en effet que pour une valeur V_g imposée on peut atteindre n’importe quelle valeur de P à condition d’augmenter V . En pratique, les tensions ne peuvent pas dépasser un plafond ; ceci limite également la puissance transmissible. Cette dissymétrie entre puissances active et réactive provient du modèle de la ligne, qui comporte une réactance et non une résistance.

La parabole, lieu des puissances maximales transmissibles, est symétrique par rapport à l’axe $P = 0$. Si l’on ajoute une résistance série au modèle de la ligne, le lieu reste une parabole mais dont l’axe se déplace.

3.4.3 Evolution de la tension avec P et Q

Si l'on suppose que la condition (3.20) est satisfaite, les solutions de l'équation (3.19) sont données par :

$$y = \frac{V_g^2}{2} - QX \pm \sqrt{\frac{V_g^4}{4} - X^2 P^2 - XQV_g^2}$$

La tension V est donnée par $\pm\sqrt{y}$ mais, comme cette grandeur est positive par définition, les solutions sont finalement :

$$V = +\sqrt{y} = \sqrt{\frac{V_g^2}{2} - QX \pm \sqrt{\frac{V_g^4}{4} - X^2 P^2 - XQV_g^2}} \quad (3.22)$$

La variation de V avec P et Q est montrée à la figure 3.7, pour des valeurs positives de P . La partie supérieure (resp. inférieure) de la surface $V(P, Q)$ correspond à la solution $+$ (resp. $-$) dans (3.22).

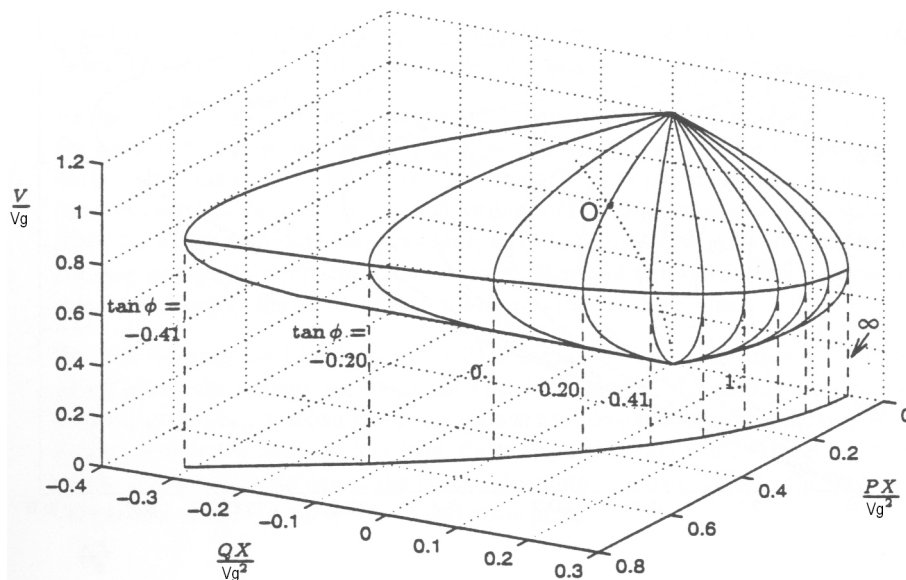


FIGURE 3.7 – variation de la tension V avec les puissances active P et réactive Q de la charge

L'“équateur” de la surface $V(P, Q)$ est le lieu des points où la charge est maximale ; sa projection sur le plan (P, Q) est la parabole de la figure 3.6.

Pour un facteur de puissance donné, le point de fonctionnement évolue le long des courbes obtenues en coupant la surface $V(P, Q)$ par des plans verticaux dont l'équation est (3.21). Ces diverses courbes, appelées usuellement “courbes PV”, sont reprises dans le diagramme à deux dimensions de la figure 3.8.

Les figures 3.7 et 3.8 montrent clairement l'existence de deux solutions pour une valeur donnée de P et de ϕ . L'existence de ces deux solutions s'explique intuitivement par le fait que la puissance

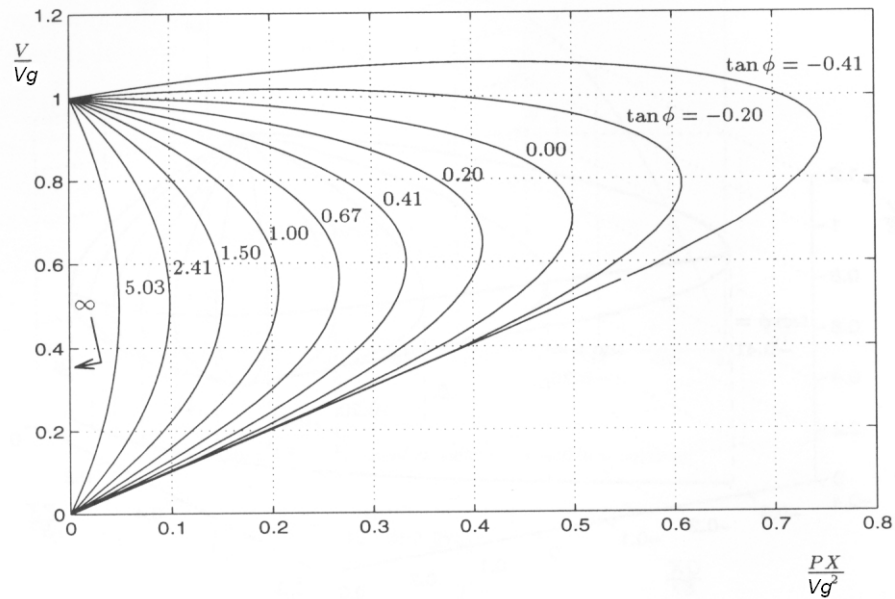


FIGURE 3.8 – courbes PV

est le produit de la tension par le courant et qu'il y a donc deux façons d'atteindre une puissance donnée : avec une tension élevée et un courant faible ou avec une tension faible et un courant élevé.

Le fonctionnement "normal" d'un réseau est sur la partie supérieure des courbes PV. En effet, lorsque $P = Q = 0$ la solution avec le signe $-$ dans (3.22) est $V = 0$ tandis que celle avec le signe $+$ est $V = V_g$. La première correspond à un court-circuit au noeud de la charge, tandis que la seconde correspond à une chute de tension nulle dans la réactance X . C'est donc la solution avec le signe $+$ qui doit être retenue, qui correspond à la partie supérieure des courbes PV.

Si la charge était purement "statique" il serait possible de fonctionner sur la partie inférieure des courbes PV (à supposer bien entendu que la tension et le courant dans la ligne aient des valeurs admissibles). Cependant, en pratique, les charges sont alimentées via des transformateurs qui ajustent automatiquement leurs rapports de transformation afin de maintenir les tensions des réseaux de distribution près de valeurs de consigne (nous reviendrons sur ce point aux sections 6.5 et 11.5). On montre que le fonctionnement de ces dispositifs est instable sur la partie inférieure des courbes PV. Le point "critique" correspondant au maximum de puissance est donc aussi une limite de stabilité. Toute tentative de fonctionner à une puissance supérieure entraîne une *instabilité de tension*. Celle-ci est étudiée plus en détail dans le cours ELEC0047.

La figure 3.8 montre clairement qu'au fur et à mesure que l'on compense la charge (valeurs de $\tan \phi$ de plus en plus négatives), il est possible de lui fournir davantage de puissance active en maintenant la tensions dans une plage de valeurs correcte. Cependant, la figure montre aussi que cette façon de procéder a pour effet d'amener la tension au point critique de plus en plus près d'une valeur normale en exploitation. En conséquence, une tension normale n'est pas la garantie qu'il reste une marge de sécurité suffisante par rapport au point critique.

On arrive à la même conclusion si l'on place des condensateurs shunt de compensation en parallèle

avec la charge dont le facteur de puissance est constant.

La notion de courbe PV s'applique aux réseaux réels (qui comportent évidemment beaucoup plus de noeuds que notre exemple simple). Dans un réseau exposé au risque d'instabilité de tension, le gestionnaire est amené à déterminer des courbes PV pour évaluer les marges de puissance consommable.

Dans l'exemple simple vu plus haut, le générateur était supposé idéal et capable de tenir sa tension constante. En pratique, les calculs de marges de puissance consommable doivent prendre en compte les limites imposées à la production de puissance réactive par les générateurs (cf section 7.4.1).

Chapitre 4

La ligne de transport

Dans ce chapitre, nous nous intéressons au comportement d'une ligne de transport de l'énergie électrique en régime sinusoïdal établi. Après avoir rappelé comment peuvent être calculés les paramètres linéiques, nous étudions le comportement de la ligne en tant que composant distribué¹. Nous en déduisons le schéma équivalent à éléments localisés utilisé dans les calculs de réseaux usuels. Nous terminons par des considérations relatives à la limite thermique. Les considérations de ce chapitre s'appliquent également aux câbles à haute tension.

4.1 Paramètres linéiques d'une ligne

Les paramètres linéiques sont les paramètres (inductance, capacité, résistance, conductance) relatifs à un tronçon de longueur infinitésimale dx , divisés par cette longueur dx . Il s'agit donc de paramètres par unité de longueur. Dans ce qui suit, nous considérons que la longueur dx est d'un mètre, valeur très petite pour une ligne électrique.

4.1.1 Inductance série

La ligne est entourée d'air, dont la perméabilité magnétique est :

$$\mu = \mu_0 \mu_r \simeq \mu_0 = 4\pi 10^{-7} \text{ H/m} \quad (4.1)$$

Le métal dont est constitué chaque conducteur est caractérisé par une perméabilité relative μ_r très proche de 1 en pratique.

1. par opposition à "localisé" : voir cours de Circuits électriques

Ligne triphasée simple

Nous considérons une ligne composée de trois conducteurs, chacun relatif à une phase. Les dimensions sont définies à la figure 4.1.

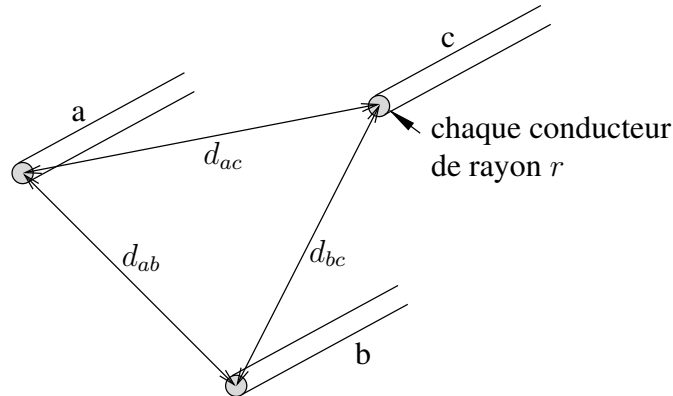


FIGURE 4.1 – ligne triphasée simple : géométrie et distances

On s'appuyant sur des notions fondamentales d'électromagnétisme, on peut établir la relation suivante entre flux et courants :

$$\begin{bmatrix} \psi_a \\ \psi_b \\ \psi_c \end{bmatrix} = \frac{\mu_0}{2\pi} \underbrace{\begin{bmatrix} \frac{\mu_r}{4} + \ln \frac{1}{r} & \ln \frac{1}{d_{ab}} & \ln \frac{1}{d_{ac}} \\ & \frac{\mu_r}{4} + \ln \frac{1}{r} & \ln \frac{1}{d_{bc}} \\ & & \frac{\mu_r}{4} + \ln \frac{1}{r} \end{bmatrix}}_{\mathbf{L}} \begin{bmatrix} i_a \\ i_b \\ i_c \end{bmatrix} \quad (4.2)$$

où ψ_a désigne le flux magnétique embrassé par un mètre du conducteur de la phase a , i_a le courant circulant dans cette phase, et de même pour les deux autres phases. La matrice $\mathbf{L}$ est la matrice d'inductance. Cette matrice est symétrique ; pour alléger l'écriture, les termes laissés en blanc sont identiques à ceux situés symétriquement par rapport à la diagonale. Le terme $\frac{\mu_0 \mu_r}{8\pi}$ correspond au champ magnétique existant à l'intérieur du conducteur.

On notera que l'expression ci-dessus est établie sous l'hypothèse :

$$i_a + i_b + i_c = 0 \quad (4.3)$$

ce qui suppose qu'il n'y a pas de retour de courant par un conducteur autre que les trois phases considérées.

Ligne triphasée transposée

Dans bon nombre de cas, les positions des conducteurs sur les pylônes sont telles que les distances d_{ab} , d_{ac} et d_{bc} ne sont pas toutes trois égales. Il en résulte un certain déséquilibre entre phases. Celui-ci peut être compensé en transposant les phases comme représenté à la figure 4.2. La matrice

d'inductance s'obtient alors comme la moyenne arithmétique des matrices relatives à chacune des trois configurations. On trouve :

$$\mathbf{L} = \frac{\mu_0}{2\pi} \begin{bmatrix} \frac{\mu_r}{4} + \ln \frac{1}{r} & \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} & \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \\ & \frac{\mu_r}{4} + \ln \frac{1}{r} & \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \\ & & \frac{\mu_r}{4} + \ln \frac{1}{r} \end{bmatrix} \quad (4.4)$$

L'expression $\sqrt[3]{d_{ab}d_{ac}d_{bc}}$ est appelée *distance moyenne géométrique*².

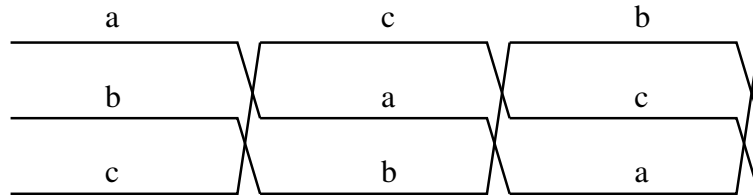


FIGURE 4.2 – transposition des conducteurs d'une ligne triphasée

A présent que les trois inductances mutuelles sont égales, on peut calculer l'inductance linéique par phase (en H/m), c'est-à-dire la partie imaginaire de l'impédance cyclique (2.14) relative à un tronçon d'un mètre de longueur, divisée par la pulsation ω . On obtient :

$$\ell = \frac{\mu_0}{2\pi} \left(\frac{\mu_r}{4} + \ln \frac{1}{r} - \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \right) = \frac{\mu_0}{2\pi} \left(\frac{\mu_r}{4} + \ln \frac{\sqrt[3]{d_{ab}d_{ac}d_{bc}}}{r} \right) \quad (4.5)$$

Ligne triphasée à faisceaux de conducteurs

A proximité d'un conducteur de faible section porté à un potentiel élevé (par rapport à la terre), les lignes équipotentielles sont très rapprochées et le champ électrique est très intense. Ceci produit une ionisation de l'air ambiant, connue sous le nom d'*effet couronne*. Ce dernier est responsable de pertes, d'interférences radio et d'une gêne acoustique (bruit audible à proximité des lignes, surtout par temps humide).

C'est la raison pour laquelle, pour des tensions nominales supérieures ou égales à 220 kV, chaque conducteur de phase est remplacé par un *faisceau* de plusieurs conducteurs maintenus à distance constante les uns des autres par des *entretoises* disposées à intervalle régulier. Le faisceau se comporte comme un conducteur dont le rayon serait nettement plus grand que celui des conducteurs qui le composent, comme le confirme un calcul ci-après. Le champ électrique est donc moins intense. En Belgique, les lignes à 380 kV (et certaines à 220 kV) comportent deux conducteurs par phase ; dans certains pays, surtout pour des tensions nominales supérieures à 380 kV, on en utilise jusqu'à quatre.

Considérons la ligne à faisceau de deux conducteurs dont la géométrie et les dimensions sont définies à la figure 4.3. En pratique, la distance d entre conducteurs d'une même phase est très

2. en anglais : Geometrical Mean Distance (GMD)

faible par rapport aux distances entre phases, de sorte que l'on peut considérer que chacun des conducteurs de la phase a est à la distance d_{ab} de chacun des conducteurs de la phase b, et de même pour les autres phases.

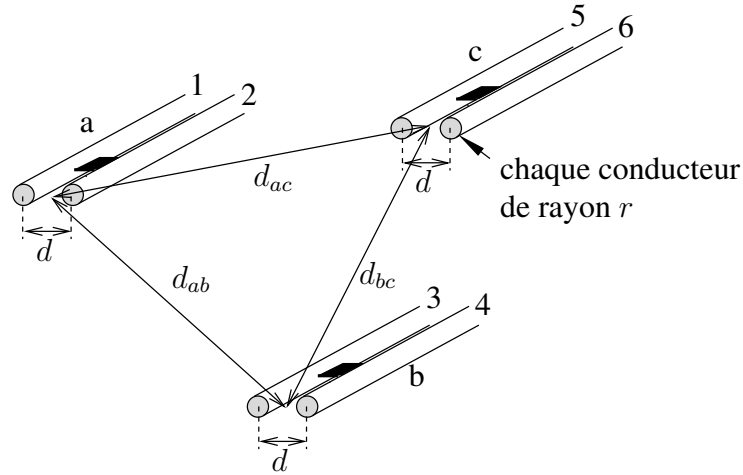


FIGURE 4.3 – ligne triphasée à faisceaux de deux conducteurs : géométrie et distances

Sous cette hypothèse, la relation entre flux et courants des six conducteurs de la figure 4.3 se présente sous la forme :

$$\begin{bmatrix} \psi_1 \\ \psi_2 \\ \psi_3 \\ \psi_4 \\ \psi_5 \\ \psi_6 \end{bmatrix} = \frac{\mu_0}{2\pi} \begin{bmatrix} \frac{\mu_r}{4} + \ln \frac{1}{r} & \ln \frac{1}{d} & \ln \frac{1}{d_{ab}} & \ln \frac{1}{d_{ab}} & \ln \frac{1}{d_{ac}} & \ln \frac{1}{d_{ac}} \\ & \frac{\mu_r}{4} + \ln \frac{1}{r} & \ln \frac{1}{d_{ab}} & \ln \frac{1}{d_{ab}} & \ln \frac{1}{d_{ac}} & \ln \frac{1}{d_{ac}} \\ & & \frac{\mu_r}{4} + \ln \frac{1}{r} & \ln \frac{1}{d} & \ln \frac{1}{d_{bc}} & \ln \frac{1}{d_{bc}} \\ & & & \frac{\mu_r}{4} + \ln \frac{1}{r} & \ln \frac{1}{d_{bc}} & \ln \frac{1}{d_{bc}} \\ & & & & \frac{\mu_r}{4} \ln \frac{1}{r} & \ln \frac{1}{d} \\ & & & & & \frac{\mu_r}{4} + \ln \frac{1}{r} \end{bmatrix} \begin{bmatrix} i_1 \\ i_2 \\ i_3 \\ i_4 \\ i_5 \\ i_6 \end{bmatrix} \quad (4.6)$$

Les conducteurs qui composent un faisceau étant reliés entre eux par les entretoises conductrices, ils sont maintenus au même potentiel tandis que chacun reprend une partie du courant de la phase. Un tronçon d'un mètre de long de la phase *a* peut donc être représenté par le circuit de la figure 4.4, où r' est la résistance d'un mètre de conducteur.

Etant donné que les deux conducteurs sont identiques et supposés à la même distance des autres phases, il est raisonnable de supposer que les flux sont égaux :

$$\psi_1 = \psi_2 \quad \psi_3 = \psi_4 \quad \psi_5 = \psi_6 \quad (4.7)$$

On déduit du circuit de la figure 4.4 que :

$$i_1 = i_2 = \frac{i_a}{2} \quad i_3 = i_4 = \frac{i_b}{2} \quad i_5 = i_6 = \frac{i_c}{2} \quad (4.8)$$

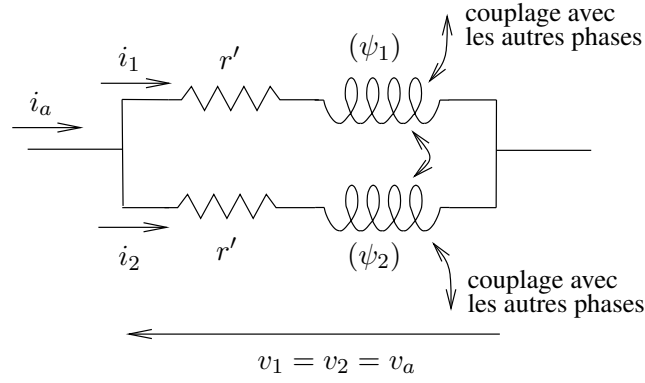


FIGURE 4.4 – schéma équivalent d'un tronçon de ligne pour le calcul de l'inductance série

On a donc :

$$v_a = v_1 = v_2 = r' \frac{i_a}{2} + \frac{d\psi_1}{dt} = r' \frac{i_a}{2} + \frac{d\psi_2}{dt} \quad (4.9)$$

Pour la phase a globalement, on peut écrire :

$$v_a = r_a i_a + \frac{d\psi_a}{dt} \quad (4.10)$$

En identifiant cette dernière relation avec (4.9), on déduit :

$$r_a = \frac{r'}{2} \quad \text{et} \quad \psi_a = \psi_1 = \psi_2$$

qui montre que le flux à considérer pour la phase a est ψ_1 (ou ψ_2). Il en va bien entendu de même pour les autres phases.

En considérant une ligne sur deux dans (4.6) et en regroupant les colonnes, on obtient aisément :

$$\begin{aligned} \begin{bmatrix} \psi_a \\ \psi_b \\ \psi_c \end{bmatrix} &= \frac{\mu_0}{2\pi} \begin{bmatrix} \frac{1}{2} \left(\frac{\mu_r}{4} + \ln \frac{1}{dr} \right) & \ln \frac{1}{d_{ab}} & \ln \frac{1}{d_{ac}} \\ & \frac{1}{2} \left(\frac{\mu_r}{4} + \ln \frac{1}{dr} \right) & \ln \frac{1}{d_{bc}} \\ & & \frac{1}{2} \left(\frac{\mu_r}{4} + \ln \frac{1}{dr} \right) \end{bmatrix} \begin{bmatrix} i_a \\ i_b \\ i_c \end{bmatrix} \\ &= \frac{\mu_0}{2\pi} \begin{bmatrix} \left(\frac{\mu_r}{8} + \ln \frac{1}{\sqrt{dr}} \right) & \ln \frac{1}{d_{ab}} & \ln \frac{1}{d_{ac}} \\ & \left(\frac{\mu_r}{8} + \ln \frac{1}{\sqrt{dr}} \right) & \ln \frac{1}{d_{bc}} \\ & & \left(\frac{\mu_r}{8} + \ln \frac{1}{\sqrt{dr}} \right) \end{bmatrix} \begin{bmatrix} i_a \\ i_b \\ i_c \end{bmatrix} \end{aligned} \quad (4.11)$$

L'expression $\sqrt{dr}$ est appelée *rayon moyen géométrique* ³.

En comparant (4.2) et (4.11), on voit que l'utilisation des deux conducteurs au lieu d'un seul, toute autre chose restant égale, n'affecte pas les inductances mutuelles mais diminue la self inductance d'une phase. En effet, le terme de self-induction à l'intérieur de chaque conducteur est divisé par deux et, surtout, le rayon r est remplacé par le rayon moyen géométrique, qui est nécessairement plus grand (vu que $d > r$).

3. en anglais : Geometric Mean Radius (GMR)

Ligne triphasée transposée à faisceau de conducteurs

Lorsque l'on combine les techniques de transposition et de faisceau, la matrice d'inductance de la ligne devient :

$$\mathbf{L} = \frac{\mu_0}{2\pi} \begin{bmatrix} \left(\frac{\mu_r}{8} + \ln \frac{1}{\sqrt{dr}} \right) & \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} & \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \\ \left(\frac{\mu_r}{8} + \ln \frac{1}{\sqrt{dr}} \right) & \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} & \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \\ \left(\frac{\mu_r}{8} + \ln \frac{1}{\sqrt{dr}} \right) & \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} & \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \end{bmatrix} \quad (4.12)$$

qui fait intervenir la distance et le rayon moyens géométriques.

Les inductances mutuelles étant à nouveau toutes égales, on peut calculer l'inductance linéique par phase (en H/m) :

$$\ell = \frac{\mu_0}{2\pi} \left(\frac{\mu_r}{8} + \ln \frac{1}{\sqrt{dr}} - \ln \frac{1}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \right) = \frac{\mu_0}{2\pi} \left(\frac{\mu_r}{8} + \ln \frac{\sqrt[3]{d_{ab}d_{ac}d_{bc}}}{\sqrt{dr}} \right) \quad (4.13)$$

qui est plus petite que celle de la ligne triphasée simple (donnée par (4.5)).

L'impédance que présente un réseau de transport contribue à limiter la puissance transmissible par celui-ci, à cause de la chute de tension qu'elle entraîne. Les résultats ci-dessus montrent que, pour diminuer l'impédance cyclique, on a intérêt à rapprocher les phases le plus possible, toutes autres choses restant égales. Cependant, il importe de maintenir une distance d'isolation minimale entre celles-ci. Cette distance est d'autant plus grande que la tension nominale du réseau est élevée.

4.1.2 Capacité shunt

La ligne est entourée d'air, dont la permittivité diélectrique est :

$$\epsilon = \epsilon_0 \epsilon_r \simeq \epsilon_0 = \frac{1}{36\pi} 10^{-9} \text{ F/m} \quad (4.14)$$

Ligne triphasée simple

Considérons à nouveau la géométrie décrite à la figure 4.1.

On considère que le sol est une surface plane, parallèle aux conducteurs et de potentiel nul. Les conducteurs portent des charges q_a , q_b et q_c par unité de longueur. Pour tenir compte du sol, on fait appel à la méthode "des images", bien connue en Electrostatique. Celle-ci consiste à remplacer le sol par des conducteurs supplémentaires a' , b' et c' . a' est placé symétriquement à a par rapport au plan du sol et porte une charge $-q_a$, et de même pour les phases b et c .

En s'appuyant sur des notions fondamentales d'électromagnétisme, on peut établir la relation suivante entre potentiels et charges électriques :

$$\begin{aligned}
 \begin{bmatrix} v_a \\ v_b \\ v_c \end{bmatrix} &= \frac{1}{2\pi\epsilon_o\epsilon_r} \begin{bmatrix} \ln \frac{1}{r} & \ln \frac{1}{d_{ab}} & \ln \frac{1}{d_{ac}} & \ln \frac{1}{d_{aa'}} & \ln \frac{1}{d_{ab'}} & \ln \frac{1}{d_{ac'}} \\ \ln \frac{1}{d_{ab}} & \ln \frac{1}{r} & \ln \frac{1}{d_{bc}} & \ln \frac{1}{d_{ba'}} & \ln \frac{1}{d_{bb'}} & \ln \frac{1}{d_{bc'}} \\ \ln \frac{1}{d_{ac}} & \ln \frac{1}{d_{bc}} & \ln \frac{1}{r} & \ln \frac{1}{d_{ca'}} & \ln \frac{1}{d_{cb'}} & \ln \frac{1}{d_{cc'}} \end{bmatrix} \begin{bmatrix} q_a \\ q_b \\ q_c \\ -q_a \\ -q_b \\ -q_c \end{bmatrix} \\
 &= \underbrace{\frac{1}{2\pi\epsilon_o\epsilon_r} \begin{bmatrix} \ln \frac{d_{aa'}}{r} & \ln \frac{d_{ab'}}{d_{ab}} & \ln \frac{d_{ac'}}{d_{ac}} \\ \ln \frac{d_{ba'}}{d_{ab}} & \ln \frac{d_{bb'}}{r} & \ln \frac{d_{bc'}}{d_{bc}} \\ \ln \frac{d_{ca'}}{d_{ac}} & \ln \frac{d_{cb'}}{d_{bc}} & \ln \frac{d_{cc'}}{r} \end{bmatrix}}_S \begin{bmatrix} q_a \\ q_b \\ q_c \end{bmatrix} \quad (4.15)
 \end{aligned}$$

où v_a, v_b et v_c sont les différences de potentiel entre les conducteurs et le sol, $d_{aa'}$ est la distance entre le conducteur a et son image a' (soit le double de la hauteur de a par rapport au sol), $d_{ab'}$ est la distance entre les conducteurs a et b' , et ainsi de suite pour les autres combinaisons. S est appelée *matrice d'inélastance*. Ses termes ont la dimension de l'inverse d'une capacité par unité de longueur.

Avant de poursuivre, on peut simplifier comme suit les expressions des termes de S . Etant donné que les conducteurs sont à une distance relativement importante du sol, on peut approximer la distance entre un conducteur et un conducteur image par $2H$, où H est la hauteur moyenne des conducteurs par rapport au sol. On suppose donc :

$$d_{aa'} \simeq d_{ab'} \simeq d_{ac'} \simeq d_{ba'} \simeq d_{bb'} \simeq d_{bc'} \simeq d_{ca'} \simeq d_{cb'} \simeq d_{cc'} \simeq 2H \quad (4.16)$$

et la relation (4.15) devient :

$$\begin{aligned}
 \begin{bmatrix} v_a \\ v_b \\ v_c \end{bmatrix} &= \frac{1}{2\pi\epsilon_o\epsilon_r} \underbrace{\begin{bmatrix} \ln \frac{2H}{r} & \ln \frac{2H}{d_{ab}} & \ln \frac{2H}{d_{ac}} \\ & \ln \frac{2H}{r} & \ln \frac{2H}{d_{bc}} \\ & & \ln \frac{2H}{r} \end{bmatrix}}_S \begin{bmatrix} q_a \\ q_b \\ q_c \end{bmatrix} \quad (4.17)
 \end{aligned}$$

Avec cette approximation, la matrice d'inélastance est devenue symétrique ; pour alléger l'écriture, les termes laissés en blanc sont identiques à ceux situés symétriquement par rapport à la diagonale.

Ligne triphasée transposée

En procédant comme pour les inductances, on établit l'expression suivante de la matrice d'inélastance d'une ligne triphasée transposée :

$$S = \frac{1}{2\pi\epsilon_o\epsilon_r} \begin{bmatrix} \ln \frac{2H}{r} & \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} & \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \\ & \ln \frac{2H}{r} & \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \\ & & \ln \frac{2H}{r} \end{bmatrix} \quad (4.18)$$

dans laquelle on retrouve la distance moyenne géométrique.

Les termes non diagonaux de S étant tous égaux, on peut calculer la capacité shunt linéique par phase, c'est-à-dire la capacité $C + 3C_m$ de la figure 2.7.b, relative à un tronçon d'un mètre de longueur. Les capacités C et C_m proviennent de la figure 2.8.

Pour ce faire, nous faisons l'hypothèse que la charge totale portée par les trois phases est nulle :

$$q_a + q_b + q_c = 0 \quad (4.19)$$

En fait, il est possible d'obtenir le résultat sans calculer au préalable les capacités C et C_m . En effet, de (4.18) on tire pour la phase a , par exemple :

$$\begin{aligned} v_a &= \frac{1}{2\pi\epsilon_o\epsilon_r} \left(\ln \frac{2H}{r} q_a + \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} (q_b + q_c) \right) \\ &= \frac{1}{2\pi\epsilon_o\epsilon_r} \left(\ln \frac{2H}{r} - \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \right) q_a \end{aligned}$$

On en déduit la capacité recherchée (en F/m) :

$$C = 2\pi\epsilon_o\epsilon_r \frac{1}{\ln \frac{\sqrt[3]{d_{ab}d_{ac}d_{bc}}}{r}} \quad (4.20)$$

résultat qui ne dépend pas de H . Ceci laisse penser que l'approximation (4.16) revient à négliger l'effet du sol. En effet, si l'on recommence le développement en négligeant la présence du sol dès le départ, on trouve également l'expression (4.20).

Ligne triphasée à faisceaux de conducteurs

Revenons à la géométrie détaillée à la figure 4.3. Nous considérons à nouveau que : (i) chacun des conducteurs de la phase a est à la distance d_{ab} de chacun des conducteurs de la phase b , et ainsi de suite pour les autres combinaisons ; (ii) la distance entre un conducteur et un conducteur image vaut $2H$.

Sous ces hypothèses, la relation entre potentiels et charges des six conducteurs de la figure 4.3 se présente sous la forme :

$$\begin{bmatrix} v_1 \\ v_2 \\ v_3 \\ v_4 \\ v_5 \\ v_6 \end{bmatrix} = \frac{1}{2\pi\epsilon_o\epsilon_r} \begin{bmatrix} \ln \frac{2H}{r} & \ln \frac{2H}{d} & \ln \frac{2H}{d_{ab}} & \ln \frac{2H}{d_{ab}} & \ln \frac{2H}{d_{ac}} & \ln \frac{2H}{d_{ac}} \\ & \ln \frac{2H}{r} & \ln \frac{2H}{d_{ab}} & \ln \frac{2H}{d_{ab}} & \ln \frac{2H}{d_{ac}} & \ln \frac{2H}{d_{ac}} \\ & & \ln \frac{2H}{r} & \ln \frac{2H}{d} & \ln \frac{2H}{d_{bc}} & \ln \frac{2H}{d_{bc}} \\ & & & \ln \frac{2H}{r} & \ln \frac{2H}{d_{bc}} & \ln \frac{2H}{d_{bc}} \\ & & & & \ln \frac{2H}{r} & \ln \frac{2H}{d} \\ & & & & & \ln \frac{2H}{r} \end{bmatrix} \begin{bmatrix} q_1 \\ q_2 \\ q_3 \\ q_4 \\ q_5 \\ q_6 \end{bmatrix} \quad (4.21)$$

On suppose de plus que la charge d'une phase se répartit de manière égale sur les deux conducteurs (identiques) qui la composent :

$$q_1 = q_2 = \frac{q_a}{2} \quad q_3 = q_4 = \frac{q_b}{2} \quad q_5 = q_6 = \frac{q_c}{2}$$

On suppose enfin que les potentiels des conducteurs d'une même phase (reliés par des entretoises) sont égaux :

$$v_1 = v_2 = v_a \quad v_3 = v_4 = v_b \quad v_5 = v_6 = v_c$$

En considérant une ligne sur deux dans (4.21) et en regroupant les colonnes, on obtient aisément :

$$\begin{aligned} \begin{bmatrix} v_a \\ v_b \\ v_c \end{bmatrix} &= \frac{1}{2\pi\epsilon_o\epsilon_r} \begin{bmatrix} \frac{1}{2} \left(\ln \frac{4H^2}{dr} \right) & \ln \frac{2H}{d_{ab}} & \ln \frac{2H}{d_{ac}} \\ & \frac{1}{2} \left(\ln \frac{4H^2}{dr} \right) & \ln \frac{2H}{d_{bc}} \\ & & \frac{1}{2} \left(\ln \frac{4H^2}{dr} \right) \end{bmatrix} \begin{bmatrix} q_a \\ q_b \\ q_c \end{bmatrix} \\ &= \frac{1}{2\pi\epsilon_o\epsilon_r} \begin{bmatrix} \ln \frac{2H}{\sqrt{dr}} & \ln \frac{2H}{d_{ab}} & \ln \frac{2H}{d_{ac}} \\ & \ln \frac{2H}{\sqrt{dr}} & \ln \frac{2H}{d_{bc}} \\ & & \ln \frac{2H}{\sqrt{dr}} \end{bmatrix} \begin{bmatrix} q_a \\ q_b \\ q_c \end{bmatrix} \end{aligned} \quad (4.22)$$

dans laquelle on retrouve le rayon moyen géométrique.

Ligne triphasée transposée à faisceau de conducteurs

Lorsque l'on combine les techniques de transposition et de faisceau, la matrice d'inélastance de la ligne devient :

$$\mathbf{S} = \frac{1}{2\pi\epsilon_o\epsilon_r} \begin{bmatrix} \ln \frac{2H}{\sqrt{dr}} & \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} & \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \\ & \ln \frac{2H}{\sqrt{dr}} & \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \\ & & \ln \frac{2H}{\sqrt{dr}} \end{bmatrix} \quad (4.23)$$

qui fait intervenir la distance et le rayon moyens géométriques.

Les capacités mutuelles étant à nouveau toutes égales, on peut calculer la capacité shunt par phase, toujours sous l'hypothèse (4.19). De (4.22) on tire pour la phase a , par exemple :

$$\begin{aligned} v_a &= \frac{1}{2\pi\epsilon_o\epsilon_r} \left(\ln \frac{2H}{\sqrt{dr}} q_a + \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} (q_b + q_c) \right) \\ &= \frac{1}{2\pi\epsilon_o\epsilon_r} \left(\ln \frac{2H}{\sqrt{dr}} - \ln \frac{2H}{\sqrt[3]{d_{ab}d_{ac}d_{bc}}} \right) q_a \end{aligned} \quad (4.24)$$

On en déduit la capacité recherchée (en F/m) :

$$c = 2\pi\epsilon_o\epsilon_r \frac{1}{\ln \frac{\sqrt[3]{d_{ab}d_{ac}d_{bc}}}{\sqrt{dr}}} \quad (4.25)$$

résultat, ici encore, indépendant de H .

4.1.3 Résistance série

Les bons conducteurs utilisés pour le transport de l'énergie électrique sont le cuivre et l'aluminium. Le cuivre a une résistivité de $18 \cdot 10^{-9} \Omega \cdot m$ tandis que celle de l'aluminium est de $29 \cdot 10^{-9} \Omega \cdot m$, soit 60 % de plus. Cependant la densité du cuivre est nettement plus élevée (8900 kg/m^3) que celle de l'aluminium (2700 kg/m^3). Il en résulte qu'à résistance égale, le poids d'un conducteur en aluminium est de $29/18 \times 2700/8900 = 0.488$ fois celui du cuivre. Ceci conduit à utiliser l'aluminium afin de moins solliciter les pylônes et les chaînes d'isolateurs. De plus, l'aluminium est moins cher que le cuivre.

En revanche, l'aluminium n'a pas la résistance mécanique requise pour les longues portées entre pylônes d'une ligne THT. On utilise donc un alliage d'aluminium plus résistant ou l'on équipe le conducteur d'une âme en acier.

Les conducteurs sont constitués de brins qui leur confèrent la flexibilité requise pour les enrouler en bobines en vue de les transporter sur les lieux de leur utilisation. En présence d'une âme en acier, les brins d'acier sont au centre, ceux d'aluminium en périphérie. Les brins sont torsadés, chaque couche étant torsadée en sens inverse de la précédente pour éviter que les brins se déroulent. A cause de leur disposition en spirale les brins sont environ 1 à 2 % plus longs que le conducteur assemblé.

Aux fréquences de 50 ou 60 Hz, il y a un certain effet pelliculaire (ou effet de peau). Rappelons que ce dernier est présent quand un courant alternatif circule dans un conducteur. La densité de courant n'est pas uniforme dans la section ; le courant se répartit davantage en périphérie du conducteur qu'au centre. Il en résulte que la résistance linéique (en Ω/m) est plus élevée qu'en courant continu, laquelle est donnée par :

$$r = \frac{\rho}{s} \quad (4.26)$$

où ρ est la résistivité du matériau (en $\Omega \cdot m$) et s la section du conducteur (en m^2). Dans le cas d'un conducteur avec âme en acier, l'effet pelliculaire conduit le courant à se concentrer dans la partie en aluminium et la partie en acier contribue relativement peu à la résistance de l'ensemble.

4.1.4 Conductance shunt

La conductance shunt (ou "latérale") d'une ligne est très faible. En fait, il existe des courants de fuite, principalement à la surface des isolateurs et surtout quand l'atmosphère est poussiéreuse (en milieu industriel) ou saline (à proximité de la mer). Toutefois les pertes associées à ces courants sont très faibles devant les puissances véhiculées par les lignes et l'on néglige très souvent cette conductance en pratique.

4.1.5 Ordres de grandeur

Le tableau ci-après donne l'ordre de grandeur des résistances série, réactances série et admittances shunt, par phase, linéiques et à 50 Hz, pour un échantillon représentatif de lignes HT et THT présentes dans le réseau belge.

Le tableau reprend également les limites thermiques considérées en fin de chapitre.

	tensions nominales (kV)			
	380	220	150	70
r (Ω/km)	0.03	0.04 – 0.09	0.05 – 0.12	0.09 – 0.35
ωl (Ω/km)	0.3 ⁽²⁾	0.3 ⁽²⁾ ou 0.4 ⁽¹⁾	0.4 ⁽¹⁾	0.2 – 0.4 ⁽¹⁾
ωc ($\mu S/\text{km}$)	3.0	3.0	3.0	3.0
S_{max} (MVA)	1350 ou 1420	250–500	150 – 350	30 – 100
⁽¹⁾ 1 conducteur par phase		⁽²⁾ 2 conducteurs par phase		

On voit qu'il y a une assez grande dispersion dans les valeurs des résistances, correspondant à une assez grande variété de sections de conducteurs.

On notera qu'en THT la résistance est faible devant la réactance.

La susceptance shunt est relativement constante pour les différents niveaux de tension considérés dans le tableau ci-dessus.

4.2 Caractéristiques des câbles

Pour des raisons évidentes, les câbles sont utilisés en milieu urbain et en milieu aquatique. Sous la pression de l'opinion et des pouvoirs publics, par souci du respect du paysage, on tend à les substituer aux lignes aériennes HT ou THT, du moins lorsqu'il s'agit de remplacer une ligne arrivée en fin de vie ou de renforcer le réseau existant.

Les câbles sont en général plus fiables que les lignes aériennes car ils ne sont pas exposés aux perturbations provenant de la foudre, du vent ou de la glace. Toutefois, il faut noter qu'un câble coûte entre 6 et 10 fois plus qu'une ligne aérienne, à puissance égale. De plus, la maintenance est plus malaisée en ce sens qu'une inspection visuelle n'est pas possible comme pour les lignes aériennes et que la réparation nécessite d'ouvrir le sol. Enfin, il existe des limitations de nature électrique, comme mentionné plus loin.

En ce qui concerne le choix du métal, par rapport à une ligne aérienne, le critère le plus contraignant n'est pas le poids mais plutôt le dégagement de chaleur dû aux pertes Joule. En effet le câble est dans un espace confiné et n'est pas ventilé comme les conducteurs d'une ligne aérienne. Dans ce cas, le cuivre est plus avantageux que l'aluminium car, à section égale, la résistance est plus faible.

Les développements de la section 4.1 s'appliquent aussi aux câbles, moyennant quelques adaptations dans le détail desquelles nous n'entrerons pas. Cependant, les valeurs des paramètres linéiques sont très différentes. Le tableau ci-après donne l'ordre de grandeur des résistances série, réactances série et admittances shunt, par phase, par km et à 50 Hz, pour un échantillon représentatif de câbles utilisés dans le réseau belge.

	tensions nominales (kV)	
	150	36
r (Ω/km)	0.03 – 0.12	0.06 – 0.16
ωl (Ω/km)	0.12 – 0.22	0.10 – 0.17
ωc ($\mu\text{S}/\text{km}$)	30 – 70	40 – 120
S_{max} (MVA)	100 – 300	10 – 30

Par rapport à une ligne aérienne de même tension nominale et de section comparable :

- dans le cas d'un câble, la permittivité ϵ du matériau isolant est beaucoup plus élevée que celle de l'air qui entoure une ligne aérienne. Les phases sont donc beaucoup plus proches l'une de l'autre. Il en résulte que l'inductance cyclique d'un câble est nettement plus faible que celle d'une ligne aérienne de même tension nominale et de section comparable. Ceci explique les valeurs nettement plus faibles de la réactance série par phase observées dans le tableau ci-dessus ;
- la permittivité plus élevée du milieu isolant conduit à une capacité shunt par phase plus élevée pour le câble. Les distances plus faibles entre phases contribuent également à une valeur plus élevée de cette capacité. Il s'en suit qu'un câble présente une susceptance shunt par phase nettement plus élevée, comme le montre le tableau ci-dessus.

La valeur élevée de cette susceptance shunt est un obstacle à l'utilisation de câbles HT ou THT sur de longues distances. En effet :

- plus la longueur augmente, plus le courant capacitif total augmente. Il existe même une longueur à laquelle ce courant pourrait atteindre la limite thermique admissible pour le câble, auquel cas ce dernier fonctionnerait à sa limite rien que par le fait d'être mis sous tension, avant même d'y faire transiter une puissance ! (voir exercice No 8)
- plus la longueur augmente, plus le câble produit de la puissance réactive, ce qui peut provoquer des surtensions dans le réseau.

Ceci conduit à utiliser le transport à courant continu (sous haute tension) au delà d'une certaine longueur de câble, par exemple pour des liaisons sous-marines.

4.3 La ligne traitée comme un composant distribué

La figure 4.5 représente le schéma par phase d'une ligne de longueur d . Nous désignons par $z = r + j\omega l$ l'impédance série linéique (en Ω/m) et par $y = g + j\omega c$ l'admittance shunt linéique (entre phase et neutre, en S/m)⁴. Nous considérons la présence d'une conductance shunt, dans un souci de généralité.

4. Nous utilisons des lettres minuscules pour désigner des grandeurs linéiques

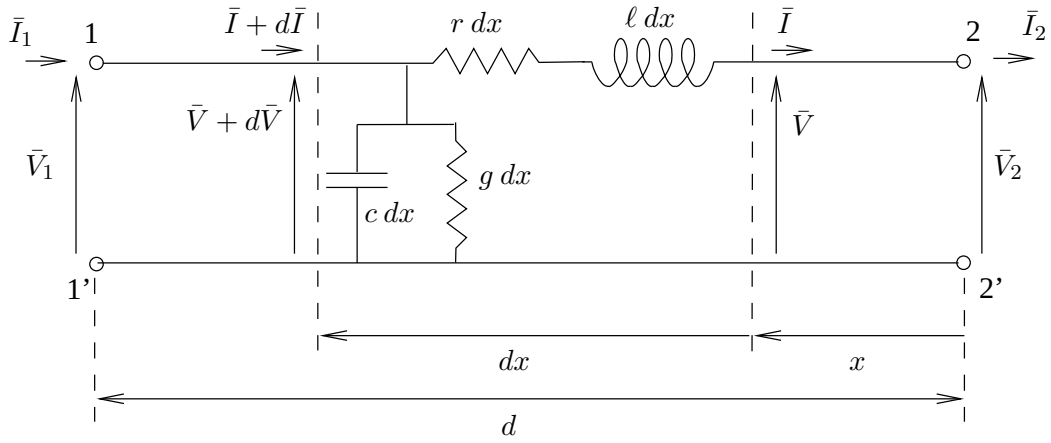


FIGURE 4.5 – schéma par phase d'une ligne en régime sinusoïdal

Désignons par x la position d'un point de la ligne, repérée par rapport à l'extrémité 22' ⁵. Les impédance, admittance, tensions et courants relatifs à une section de longueur infinitésimale dx sont indiqués à la figure 4.5.

L'application des lois d'Ohm et de Kirchhoff à cette section infinitésimale donne :

$$\begin{aligned} d\bar{V} &= \bar{I} z dx \\ d\bar{I} &= (\bar{V} + d\bar{V}) y dx \simeq \bar{V} y dx \end{aligned}$$

où le produit $d\bar{V} dx$ a été négligé. Ceci conduit aux deux équations différentielles du premier ordre :

$$\frac{d\bar{V}}{dx} = z\bar{I} \quad (4.27)$$

$$\frac{d\bar{I}}{dx} = y\bar{V} \quad (4.28)$$

qui peuvent être combinées en une équation différentielle du second ordre :

$$\text{soit } \frac{d^2\bar{V}}{dx^2} = yz\bar{V} = \bar{\gamma}^2\bar{V} \quad (4.29)$$

$$\text{soit } \frac{d^2\bar{I}}{dx^2} = yz\bar{I} = \bar{\gamma}^2\bar{I} \quad (4.30)$$

où l'on a posé :

$$\gamma = \sqrt{yz} \quad (4.31)$$

Cette grandeur est appelée la *constante de propagation* de la ligne et s'exprime en m^{-1} .

L'équation caractéristique relative à (4.29) est $s^2 - \gamma^2 = 0$, dont les racines sont $\pm\gamma$. La solution de l'équation (4.29) est donc de la forme :

$$\bar{V} = k_1 e^{\gamma x} + k_2 e^{-\gamma x} \quad (4.32)$$

⁵. ce choix simplifie certains des calculs qui suivent

La solution de l'équation (4.30) est de la même forme.

Avant de poursuivre le développement, un commentaire s'impose sur la signification des deux termes de (4.32). Décomposons k_1, k_2 et γ comme suit :

$$\begin{aligned} k_1 &= |k_1| e^{j\nu_1} \\ k_2 &= |k_2| e^{j\nu_2} \\ \gamma &= \alpha + j\beta \end{aligned} \quad (4.33)$$

La relation (4.32) devient :

$$\bar{V} = |k_1| e^{\alpha x} e^{j(\beta x + \nu_1)} + |k_2| e^{-\alpha x} e^{j(-\beta x + \nu_2)}$$

La tension à l'instant t et à la coordonnée x vaut donc :

$$v(x, t) = \underbrace{|k_1| \sqrt{2} e^{\alpha x} \cos(\omega t + \beta x + \nu_1)}_{v_1(x, t)} + \underbrace{|k_2| \sqrt{2} e^{-\alpha x} \cos(\omega t - \beta x + \nu_2)}_{v_2(x, t)}$$

Le terme $v_1(x, t)$ correspond à une onde qui se propage de la gauche vers la droite, en s'atténuant. En effet, pour un x fixé, $v_1(x, t)$ est une fonction sinusoïdale du temps et, pour un t fixé, c'est une fonction sinusoïdale de la position x . Cette onde est appelée *onde incidente*, tandis que α est appelé *constante d'atténuation* et β *constante de phase*. De même $v_2(x, t)$ correspond à une onde qui se propage, en s'atténuant, de la droite vers la gauche. Il s'agit de l'*onde réfléchie*.

La vitesse de propagation de ces ondes, soit ω/β , est celle de la lumière dans l'air qui entoure la ligne, soit un peu moins de 300.000 km/s. La longueur d'onde λ est la distance entre deux maxima voisins de la cosinusoïde, à un instant donné. On trouve aisément que $\lambda = 2\pi/\beta$. En combinant ces deux informations, on conclut que la longueur d'onde d'un signal à 50 Hz est d'environ 6.000 km. Même les lignes les plus longues utilisées dans le monde sont donc courtes par rapport à cette longueur d'onde.

L'interprétation ci-dessus prend tout son sens lorsque l'on étudie les transitoires électromagnétiques se produisant suite à un coup de foudre sur la ligne ou suite à une manoeuvre (mise sous tension par exemple). Ainsi, si une onde de tension due à la foudre se propage sur une ligne et atteint une extrémité ouverte, elle se réfléchit entièrement, ce qui peut conduire à une tension double à cette extrémité ouverte. De tels phénomènes doivent évidemment être pris en compte lors du design de l'isolation des équipements. Leur étude requiert de résoudre des équations aux dérivées partielles (équation "*des télégraphistes*"), qui sortent du cadre de ce cours.

Revenons à l'expression (4.32) et transformons-la en l'expression suivante, plus pratique :

$$\bar{V} = (k_1 + k_2) \frac{e^{\gamma x} + e^{-\gamma x}}{2} + (k_1 - k_2) \frac{e^{\gamma x} - e^{-\gamma x}}{2} = K_1 \operatorname{ch} \gamma x + K_2 \operatorname{sh} \gamma x \quad (4.34)$$

Notons $\bar{V}_1, \bar{V}_2$ (resp. $\bar{I}_1, \bar{I}_2$) les tensions (resp. courants) aux extrémités 11' et 22' de la ligne. On identifie les constantes K_1 et K_2 en considérant les conditions en $x = 0$:

- $\bar{V} = \bar{V}_2$ ce qui fournit : $K_1 = \bar{V}_2$

- $\bar{I} = \bar{I}_2$, c'est-à-dire, en vertu de (4.27) : $\left. \frac{d\bar{V}}{dx} \right|_{x=0} = z\bar{I}_2$

$$\Leftrightarrow K_1\gamma \operatorname{sh} \gamma x + K_2\gamma \operatorname{ch} \gamma x \Big|_{x=0} = z\bar{I}_2$$

$$\Leftrightarrow K_2 = \frac{z\bar{I}_2}{\gamma} = \sqrt{\frac{z}{y}}\bar{I}_2$$

En remplaçant dans (4.34), on obtient donc l'expression de la tension en un point d'abscisse x :

$$\bar{V} = \bar{V}_2 \operatorname{ch} \gamma x + Z_c \bar{I}_2 \operatorname{sh} \gamma x \quad (4.35)$$

où l'on a posé :

$$Z_c = \sqrt{\frac{z}{y}} \quad (4.36)$$

Cette grandeur est appelée *impédance caractéristique* de la ligne et s'exprime en Ω .

Nous laissons au lecteur le soin d'établir l'expression correspondante du courant en un point d'abscisse x :

$$\bar{I} = \frac{\bar{V}_2}{Z_c} \operatorname{sh} \gamma x + \bar{I}_2 \operatorname{ch} \gamma x \quad (4.37)$$

Enfin, évaluées en $x = d$, ces équations fournissent les relations entre tensions et courants aux extrémités de la ligne :

$$\bar{V}_1 = \bar{V}_2 \operatorname{ch} \gamma d + Z_c \bar{I}_2 \operatorname{sh} \gamma d \quad (4.38)$$

$$\bar{I}_1 = \frac{\bar{V}_2}{Z_c} \operatorname{sh} \gamma d + \bar{I}_2 \operatorname{ch} \gamma d \quad (4.39)$$

4.4 Quelques propriétés liées à l'impédance caractéristique

Considérons le cas d'une ligne sans pertes : $r = 0$, $g = 0$, hypothèse justifiée par le fait que g est tout à fait négligeable et r faible pour une ligne THT. On a successivement :

$$z = j\omega\ell$$

$$y = j\omega c$$

$$\gamma = j\beta = j\omega\sqrt{\ell c}$$

$$Z_c = |Z_c| = \sqrt{\frac{\ell}{c}}$$

$$\bar{V} = \bar{V}_2 \cos \beta x + jZ_c \bar{I}_2 \sin \beta x \quad (4.40)$$

$$\bar{I} = \bar{I}_2 \cos \beta x + j\frac{\bar{V}_2}{Z_c} \sin \beta x \quad (4.41)$$

L'impédance caractéristique est donc une résistance pure. Si l'on ferme la ligne sur cette résistance, c'est-à-dire si $\bar{V}_2 = Z_c \bar{I}_2$, le régime qui s'installe possède plusieurs propriétés remarquables. En effet, les relations (4.40, 4.41) fournissent :

$$\begin{aligned}\bar{V} &= \bar{V}_2 e^{j\beta x} \\ \bar{I} &= \bar{I}_2 e^{j\beta x}\end{aligned}$$

En comparant avec (4.32), on voit qu'il n'y a pas d'onde réfléchie.

On en déduit que :

- la tension (resp. le courant) a une amplitude V_2 (resp. I_2) constante tout au long de la ligne
- la tension et le courant sont en phase en tout point de la ligne. L'impédance $\bar{V}/\bar{I}$ vue en n'importe lequel de ses points est la résistance Z_c
- en conséquence, la ligne ne consomme ni ne produit de puissance réactive. Les productions $\omega c V^2$ équilibrent les pertes $\omega \ell I^2$
- la puissance triphasée qui transite au droit de n'importe quel point de la ligne est donc la puissance active fournie à la résistance Z_c , soit :

$$P_c = 3 \frac{V_2^2}{Z_c} \quad (4.42)$$

où V_2 est la tension de phase. Lorsque l'on prend pour V_2 la tension *nominale* V_N de la ligne, c'est-à-dire la tension pour laquelle elle a été conçue, la valeur de P_c est appelée *puissance naturelle* de la ligne

- ces propriétés s'appliquent quelle que soit la longueur d de la ligne !

Nous laissons au lecteur le soin de montrer que si la ligne est fermée sur une résistance inférieure (resp. supérieure) à Z_c , c'est-à-dire si elle reçoit une puissance active supérieure (resp. inférieure) à P_c , elle consomme (resp. produit) de la puissance réactive. En particulier une ligne ouverte à une de ses extrémités se comporte comme un condensateur à l'autre.

Remarque. Dans la transmission d'information, on s'arrange généralement pour fermer les câbles (coaxiaux p.ex.) sur leurs impédances caractéristiques (p.ex. 75Ω) afin de minimiser la réflexion d'onde. Par contre, ce n'est jamais le cas pour les lignes de transport d'énergie électrique. En effet, en pratique, celles-ci ne fonctionnent jamais à leurs puissances naturelles, les flux de puissance étant variables car dictés par les consommations et les productions.

4.5 Schéma équivalent d'une ligne

La structure la plus employée pour représenter une ligne est le schéma équivalent en pi, représenté à la figure 4.6. Déterminons la valeur à donner à l'impédance Z_{ser} et à l'admittance Y_{sh} de ce *circuit à éléments condensés* pour que, vu des accès 11' et 22', il ait le même comportement que le *composant distribué* considéré à la figure 4.5.

Définissons au préalable les grandeurs suivantes, relatives à la totalité de la ligne :

$$R = r d \quad L = \ell d \quad G = g d \quad C = c d \quad Z = z d \quad Y = y d$$

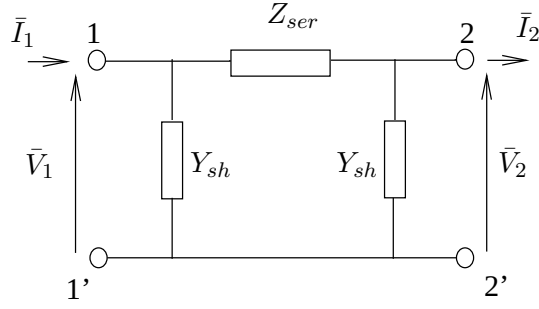


FIGURE 4.6 – schéma équivalent en pi de la ligne (éléments condensés)

On établit aisément à partir de la figure 4.6 que :

$$\bar{V}_1 = \bar{V}_2 + Z_{ser}(\bar{I}_2 + Y_{sh}\bar{V}_2) = (1 + Z_{ser}Y_{sh})\bar{V}_2 + Z_{ser}\bar{I}_2$$

Une identification terme à terme avec (4.38) fournit :

$$Z_{ser} = Z_c \operatorname{sh} \gamma d \quad (4.43)$$

$$1 + Z_{ser}Y_{sh} = \operatorname{ch} \gamma d \quad \Leftrightarrow \quad Y_{sh} = \frac{\operatorname{ch} \gamma d - 1}{Z_c \operatorname{sh} \gamma d} = \frac{1}{Z_c} \operatorname{th} \frac{\gamma d}{2} \quad (4.44)$$

Pour des lignes d'une longueur inférieure à 150 km, on considère que $|\gamma d|$ est suffisamment faible pour pouvoir remplacer les fonctions hyperboliques par leurs développements en série limités au premier ordre :

$$\begin{aligned} \operatorname{sh} \gamma d &\simeq \gamma d + \dots \\ \operatorname{th} \frac{\gamma d}{2} &\simeq \frac{\gamma d}{2} + \dots \end{aligned}$$

Une substitution dans (4.43, 4.44) donne alors :

$$\begin{aligned} Z_{ser} &= Z_c \gamma d = \sqrt{\frac{z}{y}} \sqrt{zy} d = zd = Z \\ Y_{sh} &= \frac{1}{Z_c} \frac{\gamma d}{2} = \frac{1}{2} \sqrt{\frac{y}{z}} \sqrt{zy} d = \frac{1}{2} yd = \frac{Y}{2} \end{aligned}$$

En conclusion, une ligne de transport peut toujours être modélisée par un schéma équivalent en pi. Pour une longueur inférieure à 150 km, les paramètres de ce schéma équivalent sont obtenus en multipliant simplement les valeurs linéiques par la longueur de la ligne. On parle de *ligne courte*. Au-delà de 150 km, il convient d'utiliser les expressions (4.43, 4.44).

4.6 Limite thermique

Le passage du courant dans un conducteur de ligne y entraîne des pertes par effet Joule, qui échauffent le matériau.

Un tel échauffement doit être limité pour deux raisons :

- il entraîne une dilatation du conducteur, qui le fait se rapprocher du sol, d'où un risque de court-circuit ou d'électrocution ;
- au delà d'une certaine température, l'aluminium subit une dégradation irréversible par un effet de "recuit", diminuant sa résistance mécanique.

La température maximale typique des conducteurs d'une ligne aérienne est de 75° C.

Chaque ligne est caractérisée par un courant maximal admissible en permanence, dans une quelconque de ses phases. Nous noterons ce dernier I_{max} . Cette valeur est souvent désignée sous le vocable d'*ampacité*.

C'est principalement la densité de courant maximale (en A/mm²) qui détermine la valeur de I_{max} . Evidemment, plus la section du conducteur augmente, plus le courant I_{max} est élevé.

I_{max} dépend des conditions de refroidissement de la ligne. Ainsi, en hiver une ligne peut supporter un courant plus élevé qu'en été, car l'air ambiant la refroidit davantage. Par ailleurs, à section de métal égale, un faisceau offre une plus grande surface de contact avec l'air, d'où une meilleure évacuation de la chaleur et donc un courant maximal admissible plus élevé.

De manière à tirer le meilleur parti possible des lignes, certains exploitants de réseau s'efforcent d'estimer la valeur de I_{max} en fonction des conditions climatiques du moment. Une difficulté réside dans le fait que tous les tronçons de la ligne ne sont pas nécessairement exposés aux mêmes conditions de refroidissement (p.ex. vitesse du vent) et que c'est le tronçon le moins favorisé qui limite le courant que la ligne peut véhiculer. Pour plus de détails, le lecteur est invité à se reporter au cours ELEC0014 "Transport et distribution de l'énergie électrique".

En ce qui concerne la limite thermique d'un câble :

- à section de métal égale, elle est plus faible que pour une ligne, étant donné que l'évacuation de la chaleur se fait beaucoup moins bien et qu'une température excessive dégraderait l'isolant entourant les conducteurs ;
- un obstacle à l'augmentation de la section des conducteurs est la perte de souplesse mécanique du câble, rendant son enroulement impossible.

Très souvent on caractérise la capacité thermique par la puissance apparente triphasée S_{max} qui traverse le composant lorsque la tension est à sa valeur nominale et le courant égal à I_{max} . On a donc :

$$S_{max} = 3V_N I_{max} = \sqrt{3}U_N I_{max}$$

où V_N est la valeur nominale de la tension de phase et U_N celle entre phases. Des exemples de valeurs sont donnés à la section 4.1.5.

Notons que, par inertie thermique, la montée en température de la ligne ou du câble n'est pas instantanée. Une surcharge thermique au delà de I_{max} est donc tolérable durant un certain temps. Ce dernier est d'autant plus court que la surcharge est forte. Certains exploitants définissent des limites thermiques admissibles à 1, 10 ou 20 minutes, par exemple. Au-delà d'une certaine valeur du courant, la ligne est mise hors service par les protections.

Mentionnons enfin le développement au cours de la dernière décennie de lignes aériennes équipées de conducteurs HTLS, abréviation pour “High Temperature Low Sag ⁶”. Dans ces derniers, l’âme aussi bien que les conducteurs peuvent supporter des températures plus élevées - jusqu’à 210° C - sans se détériorer, ni se dilater autant que les matériaux traditionnels (coefficient de dilatation de 3 à 6 fois plus faible). L’âme en acier, en alliage spécial ou en matériau composite reprend tout l’effort mécanique au delà d’une certaine température. Ceci permet de transporter des courants environ 2 fois plus élevés, soit un gain considérable en termes de limite S_{max} . Des variations de résistance importantes accompagnent de telles excursions de température : la résistance est de 1.5 à 2 fois plus élevée à 210° C qu’à 25° C. Il convient donc d’ajuster R dans le modèle de la ligne, en fonction du courant véhiculé par celle-ci. Pour une ligne traditionnelle où la température se situe entre 25° C et 75° C, cette variation de R est nettement moins importante.

6. flèche

Chapitre 5

Le système “per unit”

La plupart des calculs dans les systèmes électriques de puissance se font en traitant des grandeurs adimensionnelles. Ces dernières s’obtiennent en divisant chaque grandeur (tension, courant, puissance, etc...) par une grandeur de même dimension, appelée *base*. On dit que les grandeurs sans dimension ainsi obtenues sont exprimées en *per unit*, ce que l’on note par *pu*.

Cette pratique universellement répandue offre principalement les avantages suivants :

1. En per unit, les paramètres des équipements construits d’une manière semblable ont des valeurs assez proches, quelle que soit leur puissance nominale. Les valeurs des paramètres étant prévisibles, on peut :
 - vérifier plus aisément la plausibilité de données ou de résultats
 - affecter des valeurs par défaut à des paramètres manquants, lorsque l’on désire chiffrer en première approximation tel ou tel phénomène.
2. En per unit, les tensions sont, en régime de fonctionnement normal, proches de l’unité (càd proches de 1 pu). Ceci conduit généralement à un meilleur conditionnement numérique des calculs, par suite d’une moins grande dispersion des valeurs numériques.
3. Le passage en per unit fait disparaître les transformateurs idéaux qui sont présents dans les schémas équivalents des transformateurs réels. En d’autres termes, le système per unit permet de faire abstraction des différents niveaux de tension.

Exemple. La réactance interne d’une machine synchrone vaut typiquement entre 1.5 et 2.5 pu (dans la base de la machine). Pour une machine de caractéristiques 20 kV et 300 MVA, une réactance de 2.667Ω est-elle normale ? Même question pour une machine de caractéristiques 15 kV et 30 MVA.

Pour la première machine, l’impédance de base Z_B vaut, comme on le verra ci-après, $20^2/300 = 1.333 \Omega$. La réactance en per unit vaut donc $2.667/1.333 = 2 \text{ pu}$, soit une valeur tout à fait normale.

Pour la seconde machine, Z_B vaut $15^2/30 = 7.5 \Omega$. La réactance en per unit vaut donc $2.667/7.5 = 0.356 \text{ pu}$, soit une valeur anormalement faible.

5.1 Passage en per unit d'un circuit électrique

La mise en per unit des équations qui régissent un circuit électrique requiert le choix de *trois* grandeurs de base. Par exemple, si nous choisissons (arbitrairement) une puissance, une tension et un temps de base, que nous notons respectivement S_B , V_B et t_B , les autres grandeurs de base s'en déduisent en utilisant les lois fondamentales de l'électricité :

- courant de base : $I_B = \frac{S_B}{V_B}$
- impédance de base : $Z_B = \frac{V_B}{I_B} = \frac{V_B^2}{S_B}$
- flux de base : $\psi_B = V_B t_B$
- inductance de base : $L_B = \frac{\psi_B}{I_B} = \frac{V_B^2 t_B}{S_B}$
- pulsation de base : $\omega_B = \frac{Z_B}{L_B} = \frac{1}{t_B}$

Notons que, conformément à l'usage, V_B et I_B sont des valeurs *efficaces*.

On peut évidemment choisir une pulsation plutôt qu'un temps de base, tous deux étant liés par la dernière relation ci-dessus. Dans ce cours, nous choisissons pour ω_B la pulsation ω_N correspondant à la fréquence nominale f_N :

$$\omega_B = \omega_N = 2\pi 50 \quad \text{ou} \quad 2\pi 60$$

et donc :

$$t_B = \frac{1}{\omega_B} = \frac{1}{2\pi f_N} = \frac{1}{100\pi} \quad \text{ou} \quad \frac{1}{120\pi}$$

Notons au passage que moyennant ce choix, une réactance à la fréquence f_N a la même valeur que l'inductance correspondante, puisque :

$$X_{pu} = \frac{X}{Z_B} = \frac{\omega_B L}{\omega_B L_B} = \frac{L}{L_B} = L_{pu}$$

Considérons à présent le passage en per unit d'une relation typique du régime sinusoïdal :

$$S = V I \cos(\phi - \psi) + j V I \sin(\phi - \psi)$$

On a successivement :

$$S_{pu} = \frac{S}{S_B} = \frac{V I}{V_B I_B} \cos(\phi - \psi) + j \frac{V I}{V_B I_B} \sin(\phi - \psi) = V_{pu} I_{pu} \cos(\phi - \psi) + j V_{pu} I_{pu} \sin(\phi - \psi)$$

Comme cette relation ne fait pas intervenir le temps, t_B n'est pas utilisé. Seule la puissance et la tension de base sont utilisées en régime sinusoïdal.

Considérons ensuite la mise en per unit d'une équation différentielle typique du régime dynamique :

$$v = R i + L \frac{di}{dt}$$

On a successivement :

$$v_{pu} = \frac{v}{V_B} = \frac{R i}{Z_B I_B} + \frac{L}{\omega_B L_B I_B} \frac{d i}{d t} = R_{pu} i_{pu} + L_{pu} \frac{1}{\omega_B} \frac{d i_{pu}}{d t} = R_{pu} i_{pu} + L_{pu} \frac{d i_{pu}}{d t_{pu}}$$

Dans ce second exemple, le temps apparaît explicitement. On voit qu'il y a deux possibilités :

- soit toutes les grandeurs sont mises en per unit, y compris le temps : l'équation est alors strictement identique en unités physiques et en per unit ;
- soit on préfère conserver le temps en secondes : il apparaît alors un facteur $1/\omega_B$ devant l'opérateur de dérivation.

5.2 Passage en per unit de deux circuits magnétiquement couplés

Considérons deux bobines magnétiquement couplées, possédant respectivement n_1 et n_2 spires. Les flux totaux ψ_1 et ψ_2 embrassés par ces bobines sont reliés aux courants i_1 et i_2 qui les traversent par :

$$\psi_1 = L_{11} i_1 + L_{12} i_2 \quad (5.1)$$

$$\psi_2 = L_{21} i_1 + L_{22} i_2 \quad (5.2)$$

En principe, la mise en per unit de ces deux circuits requiert de choisir 6 grandeurs de base (4 en régime sinusoïdal). Il existe toutefois deux contraintes pratiques, qui ne laissent en fait que 4 degrés de liberté (3 en régime sinusoïdal) :

1. **Temps identiques.** Pour des raisons de simplicité, on désire avoir le même temps en pu dans les deux circuits. On choisit donc :

$$t_{1B} = t_{2B} \quad (5.3)$$

2. **Symétrie des matrices d'inductances.** En Henrys, on a toujours $L_{12} = L_{21}$. Il est indiqué de conserver cette propriété après passage en per unit.

La relation (5.1) se met en per unit comme suit :

$$\psi_{1pu} = \frac{\psi_1}{\psi_{1B}} = \frac{L_{11}}{L_{1B}} \frac{i_1}{I_{1B}} + \frac{L_{12}}{L_{1B}} \frac{i_2}{I_{1B}} = L_{11pu} i_{1pu} + \frac{L_{12} I_{2B}}{L_{1B} I_{1B}} i_{2pu}$$

On en déduit la valeur de L_{12} en per unit :

$$L_{12pu} = \frac{L_{12} I_{2B}}{L_{1B} I_{1B}} \quad (5.4)$$

On obtient de même :

$$L_{21pu} = \frac{L_{21} I_{1B}}{L_{2B} I_{2B}}$$

Pour avoir $L_{12pu} = L_{21pu}$, il faut donc que :

$$\frac{I_{2B}}{L_{1B} I_{1B}} = \frac{I_{1B}}{L_{2B} I_{2B}}$$

soit après calcul :

$$S_{1B} t_{1B} = S_{2B} t_{2B}$$

Etant donné que l'on a choisi le même temps de base dans les deux circuits, il faut, pour conserver la symétrie de la matrice d'inductances, choisir également la même puissance de base :

$$S_{1B} = S_{2B} \quad (5.5)$$

Un système per unit qui satisfait à (5.3, 5.5) est dit *réciproque*. En effet, la matrice d'inductance des deux bobines étant symétrique, le quadripôle correspondant est réciproque.

En application de ce raisonnement, dans un circuit comportant plusieurs niveaux de tension reliés par des transformateurs, on choisira partout un même temps de base t_B et une même puissance de base S_B . Ensuite, à chaque niveau de tension, on choisira une tension de base V_B .

5.3 Passage en per unit d'un système triphasé

Un circuit triphasé n'est jamais qu'un cas particulier de circuit et l'on peut lui appliquer ce qui précède, c'est-à-dire choisir :

- partout : un temps de base t_B et une puissance de base S_B
- par niveau de tension : une tension de base V_B entre phase et neutre, à laquelle on va rapporter toutes les tensions entre phase et neutre.

On notera que le choix d'une même tension de base pour les 3 phases a du sens puisque le système est conçu avec la même tension nominale dans chaque phase.

En ce qui concerne le choix de S_B , il convient de distinguer régimes équilibrés et déséquilibrés.

5.3.1 Régime déséquilibré - analyse des 3 phases

Si le régime est déséquilibré, il y a lieu d'analyser chacune des trois phases. Dans ce cas, il est confortable de prendre une puissance *monophasée* comme base S_B . On en déduit le courant de base :

$$I_B = \frac{S_B}{V_B} \quad (5.6)$$

l'impédance de base :

$$Z_B = \frac{V_B}{I_B} = \frac{V_B^2}{S_B} \quad (5.7)$$

et ainsi de suite pour les autres grandeurs.

Que devient, en per unit, l'expression de la puissance complexe transitant dans les trois phases ?
On a :

$$S = \bar{V}_a \bar{I}_a^* + \bar{V}_b \bar{I}_b^* + \bar{V}_c \bar{I}_c^*$$

et en divisant par S_B , on obtient la puissance en per unit :

$$S_{pu} = \frac{S}{S_B} = \frac{\bar{V}_a}{V_B} \frac{\bar{I}_a^*}{I_B} + \frac{\bar{V}_b}{V_B} \frac{\bar{I}_b^*}{I_B} + \frac{\bar{V}_c}{V_B} \frac{\bar{I}_c^*}{I_B} = \bar{V}_{apu} \bar{I}_{apu}^* + \bar{V}_{bpu} \bar{I}_{bpu}^* + \bar{V}_{c pu} \bar{I}_{c pu}^*$$

On voit que l'expression est la même, que l'on travaille en unités physiques ou per unit. Ce confort justifie le choix de S_B . Il est bien évident qu'on peut faire un autre choix, à condition d'ajuster les relations en per unit par rapport aux expressions en unités physiques.

5.3.2 Régime équilibré - analyse par phase

Si le régime est équilibré, l'analyse du système triphasé peut se ramener à l'analyse d'une de ses phases, comme on la montré à la section 2.4. Dans ce cas, comme montré ci-après, il est plus confortable de prendre une puissance *triphasee* comme base S_B . On en déduit le courant de base :

$$I_B = \frac{S_B}{3V_B}, \quad (5.8)$$

l'impédance de base :

$$Z_B = \frac{V_B}{I_B} = \frac{3V_B^2}{S_B} \quad (5.9)$$

et ainsi de suite pour les autres grandeurs.

L'expression de la puissance complexe dans les trois phases devient, compte tenu de l'équilibre entre phases :

$$S = \bar{V}_a \bar{I}_a^* + \bar{V}_b \bar{I}_b^* + \bar{V}_c \bar{I}_c^* = 3 \bar{V}_a \bar{I}_a^*$$

où le choix de la phase a est évidemment arbitraire. En per unit, cette expression devient :

$$S_{pu} = \frac{S}{S_B} = \frac{3 \bar{V}_a \bar{I}_a^*}{3 V_B I_B} = \bar{V}_{apu} \bar{I}_{apu}^*$$

c'est-à-dire la puissance calculée dans le circuit monophasé relatif à une phase. Le système per unit prolonge la technique de l'analyse par phase en ce sens que les calculs en per unit ne doivent plus du tout tenir compte de la présence des deux autres phases. Le facteur 3 intervenant dans la puissance disparaît.

Evidemment, une fois effectuée l'analyse de la phase a , en per unit, on revient à la puissance triphasée, en MW, Mvar ou MVA, par la multiplication par S_B , cette dernière valeur incluant le facteur 3 qui tient compte de la présence de 3 phases.

Remarques

1. Comme on l'a déjà mentionné, l'usage est de caractériser la tension d'un système triphasé par la

valeur efficace de la tension *entre phases*. En désignant par U_B la tension de base de cette nature, les grandeurs de base I_B et Z_B sont données par :

$$I_B = \frac{S_B}{\sqrt{3}U_B} \quad Z_B = \frac{U_B^2}{S_B} \quad (5.10)$$

2. La puissance de base triphasée valant trois fois la puissance monophasée, les impédances de base (5.7) et (5.9) ont la même valeur dans les deux systèmes. Il en est donc de même des impédances en per unit.

5.4 Changement de base

Enfin, en pratique, on est souvent amené à transférer d'une base à une autre des paramètres fournis en per unit. Les formules s'établissent aisément comme suit.

Pour un système triphasé, une impédance Z (en ohms) vaut en per unit :

$$\text{dans la première base : } Z_{pu1} = \frac{Z}{Z_{B1}} = \frac{Z S_{B1}}{3V_{B1}^2}$$

$$\text{dans la seconde base : } Z_{pu2} = \frac{Z}{Z_{B2}} = \frac{Z S_{B2}}{3V_{B2}^2}$$

En divisant une relation par l'autre, on trouve aisément :

$$Z_{pu2} = Z_{pu1} \frac{S_{B2}}{S_{B1}} \left(\frac{V_{B1}}{V_{B2}} \right)^2 \quad (5.11)$$

Chapitre 6

Le transformateur de puissance

Au-delà d'une certaine distance et/ou d'une certaine puissance, le transport d'énergie électrique doit se faire sous une tension suffisamment élevée. En effet, la puissance est le produit de la tension par le courant ; pour une puissance donnée, plus la tension est élevée, plus le courant est faible. Il en résulte donc des pertes par effet Joule et des sections de conducteurs plus faibles. A l'heure actuelle, l'énergie électrique se transporte sous des tensions allant de 70 à 380 kV en Europe, jusque 765 kV en Amérique du Nord. Il existe des lignes à 1 MV ailleurs dans le monde.

Or, la tension aux bornes d'un alternateur ne dépasse pas 25 kV en pratique. Il s'agit en effet d'une machine relativement compacte et son fonctionnement sous des tensions plus élevées poserait des problèmes d'isolation.

Le transformateur est le composant permettant d'élever l'amplitude de la tension alternative disponible à la sortie de l'alternateur pour l'amener aux niveaux requis pour le transport. A l'autre bout de la chaîne, du côté des consommateurs, les transformateurs sont utilisés pour abaisser la tension et la ramener aux valeurs utilisées dans les réseaux de distribution.

Enfin, en plus de transmettre de l'énergie électrique d'un niveau de tension à un autre, les transformateurs peuvent être utilisés pour contrôler la tension et les flux de puissance dans le réseau.

6.1 Transformateur monophasé

6.1.1 Principe

Un transformateur monophasé est constitué :

- d'un *noyau* magnétique feuilleté, obtenu par empilement de tôles réalisées dans un matériau à haute perméabilité magnétique
- de deux bobinages enroulés autour du noyau magnétique de manière à assurer un bon couplage

magnétique entre ces deux circuits.

Un des enroulements est qualifié de *primaire* et sera repéré par l'indice 1 dans ce qui suit ; l'autre est qualifié de *secondaire* et sera repéré par l'indice 2. Si la tension secondaire est supérieure (resp. inférieure) à la tension primaire, on parle de transformateur *élévateur* (resp. *abaisseur*).

Un schéma de principe est donné à la figure 6.1. Insistons sur le fait qu'il s'agit d'un schéma idéalisé. Ainsi, les enroulements primaire et secondaire ont été représentés séparés pour des raisons de lisibilité mais dans un transformateur réel, ils se présentent généralement sous forme de deux cylindres concentriques, ou parfois de galettes alternées, de manière à assurer le meilleur couplage possible.

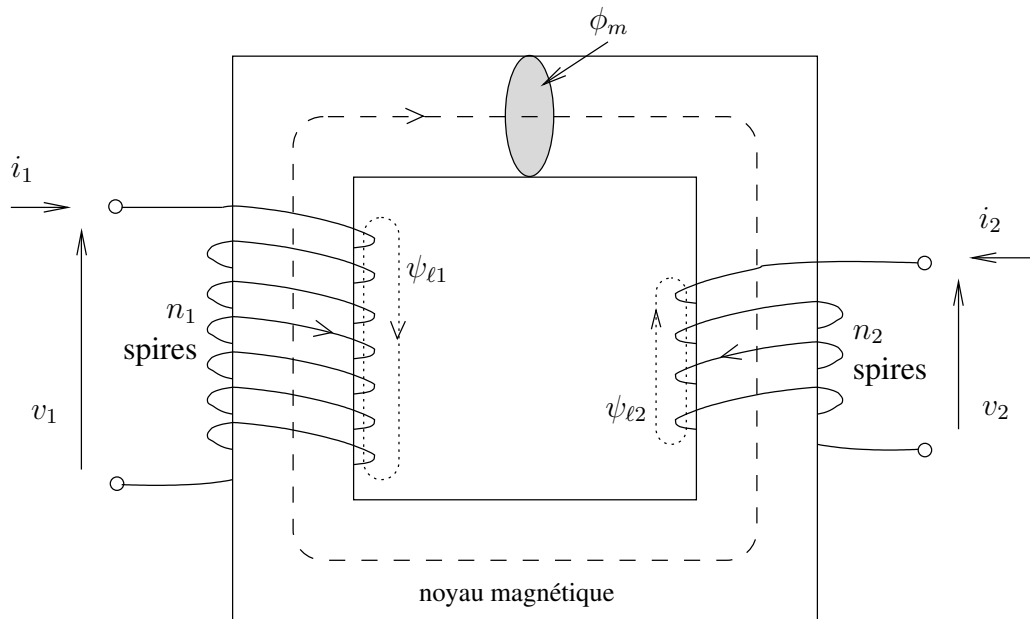


FIGURE 6.1 – transformateur monophasé (schéma de principe)

Le principe du transformateur est simple. Lorsque le primaire est alimenté par une source de tension alternative, il circule un courant i_1 qui crée dans le noyau magnétique un champ également alternatif dont l'amplitude dépend du nombre de spires n_1 du primaire et de la tension appliquée. Ce champ coupe les spires de l'enroulement secondaire et y crée un flux d'induction variable. Ceci induit une tension proportionnelle au nombre de spires n_2 de cet enroulement. La fermeture du circuit secondaire sur une charge (par exemple) provoque la circulation d'un courant i_2 dans cet enroulement. Ce courant génère à son tour un champ magnétique dans le noyau.

6.1.2 Modélisation

Lignes de champ et flux

Les lignes du champ magnétique créé par les courants i_1 et i_2 sont esquissées à la figure 6.1. Comme le suggère la figure, la majeure partie des lignes de champ sont contenues dans le noyau

magnétique et coupent les deux enroulements. Cependant, certaines lignes de champ se ferment à l'extérieur du noyau en ne coupant les spires que d'un seul enroulement.

Notons ψ_1 (resp. ψ_2) le flux total embrassé par l'enroulement primaire (resp. secondaire). Notons également ϕ_m le flux créé par le champ dans une section du noyau magnétique. Ce flux est relié aux courants i_1 et i_2 par :

$$n_1 i_1 + n_2 i_2 = \mathcal{R} \phi_m \quad (6.1)$$

où $\mathcal{R}$ est la reluctance du noyau magnétique.

Compte tenu des deux types de lignes de champ, on peut décomposer ψ_1 en :

$$\psi_1 = \psi_{\ell 1} + n_1 \phi_m \quad (6.2)$$

où $\psi_{\ell 1}$ est le *flux de fuite* créé par les lignes de champ qui ne passent pas par le noyau et ne coupent que l'enroulement 1, tandis que $n_1 \phi_m$ est le flux créé par les lignes du champ qui passent par le circuit magnétique et sont communes aux deux enroulements. On peut de même décomposer ψ_2 en :

$$\psi_2 = \psi_{\ell 2} + n_2 \phi_m \quad (6.3)$$

Transformateur idéal

Supposons d'abord que le transformateur est construit de façon parfaite, c'est-à-dire :

- avec des enroulements sans résistance
- ne présentant aucun flux de fuite
- enroulés autour d'un matériau de perméabilité infinie.

La troisième hypothèse simplificatrice conduit à une reluctance nulle. La relation (6.1) fournit directement :

$$i_2 = -\frac{n_1}{n_2} i_1 \quad (6.4)$$

En l'absence de résistances et de flux de fuite, on a par ailleurs :

$$\begin{aligned} v_1 &= \frac{d\psi_1}{dt} = n_1 \frac{d\phi_m}{dt} \\ v_2 &= \frac{d\psi_2}{dt} = n_2 \frac{d\phi_m}{dt} \end{aligned}$$

d'où on tire évidemment :

$$v_2 = \frac{n_2}{n_1} v_1 \quad (6.5)$$

Les relations (6.4) et (6.5) correspondent bien au transformateur idéal tel qu'utilisé en théorie des circuits et repris à la figure 6.2.

Ces relations montrent clairement que, dans le cas d'un transformateur abaisseur, on a $n_2 < n_1$, $v_2 < v_1$ et $i_2 > i_1$. L'enroulement secondaire présente moins de spires mais une section plus grande. C'est évidemment l'inverse pour un transformateur élévateur¹.

1. il suffit d'ailleurs de permuter les accès 1 et 2 !

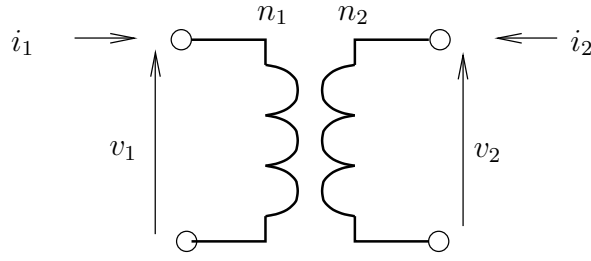


FIGURE 6.2 – transformateur idéal

On déduit encore de (6.4, 6.5) que :

$$v_1 i_1 = -v_2 i_2 \quad (6.6)$$

Le transformateur idéal est sans pertes ; la puissance électrique qui entre par un enroulement sort entièrement par l'autre.

Transformateur réel

Nous allons à présent relâcher les trois hypothèses simplificatrices faites plus haut et bâtir un schéma équivalent du transformateur réel au départ du transformateur idéal.

Définissons au préalable les *inductances de fuite* relatives aux deux enroulements :

$$L_{\ell 1} = \frac{\psi_{\ell 1}}{i_1} \quad L_{\ell 2} = \frac{\psi_{\ell 2}}{i_2} \quad (6.7)$$

ainsi que l'*inductance magnétisante (vue du primaire)* :

$$L_{m1} = \frac{n_1^2}{\mathcal{R}} \quad (6.8)$$

En introduisant (6.1) et (6.7) dans (6.2), on obtient pour l'enroulement primaire :

$$\begin{aligned} \psi_1 &= L_{\ell 1} i_1 + n_1 \frac{n_1 i_1 + n_2 i_2}{\mathcal{R}} \\ &= L_{\ell 1} i_1 + \frac{n_1^2}{\mathcal{R}} i_1 + \frac{n_1 n_2}{\mathcal{R}} i_2 \\ &= L_{\ell 1} i_1 + L_{m1} i_1 + \frac{n_2}{n_1} L_{m1} i_2 \end{aligned} \quad (6.9)$$

et pour l'enroulement secondaire :

$$\begin{aligned} \psi_2 &= L_{\ell 2} i_2 + n_2 \frac{n_1 i_1 + n_2 i_2}{\mathcal{R}} \\ &= L_{\ell 2} i_2 + \frac{n_2^2}{\mathcal{R}} i_2 + \frac{n_1 n_2}{\mathcal{R}} i_1 \\ &= L_{\ell 2} i_2 + \left(\frac{n_2}{n_1} \right)^2 L_{m1} i_2 + \frac{n_2}{n_1} L_{m1} i_1 \end{aligned} \quad (6.10)$$

Les tensions aux bornes des enroulements s'obtiennent alors comme suit :

$$v_1 = R_1 i_1 + \frac{d\psi_1}{dt} = R_1 i_1 + L_{\ell 1} \frac{di_1}{dt} + L_{m1} \frac{di_1}{dt} + \frac{n_2}{n_1} L_{m1} \frac{di_2}{dt}$$

$$v_2 = R_2 i_2 + \frac{d\psi_2}{dt} = R_2 i_2 + L_{\ell 2} \frac{di_2}{dt} + \left(\frac{n_2}{n_1}\right)^2 L_{m1} \frac{di_2}{dt} + \frac{n_2}{n_1} L_{m1} \frac{di_1}{dt}$$

Le lecteur est invité à vérifier que le schéma équivalent de la figure 6.3 correspond parfaitement à ces deux dernières relations.

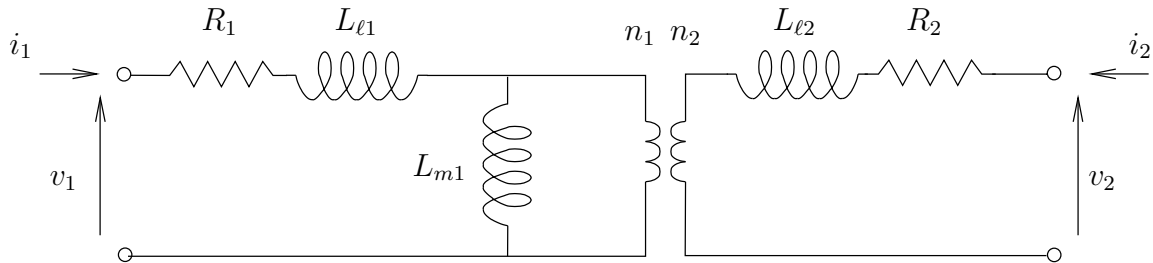


FIGURE 6.3 – schéma équivalent du transformateur monophasé (1ère version)

Ce dernier schéma peut être modifié en faisant passer la résistance R_2 et l'inductance $L_{\ell 2}$ de l'autre côté du transformateur idéal, moyennant multiplication par $(n_1/n_2)^2$. Ceci conduit au schéma équivalent de la figure 6.4.

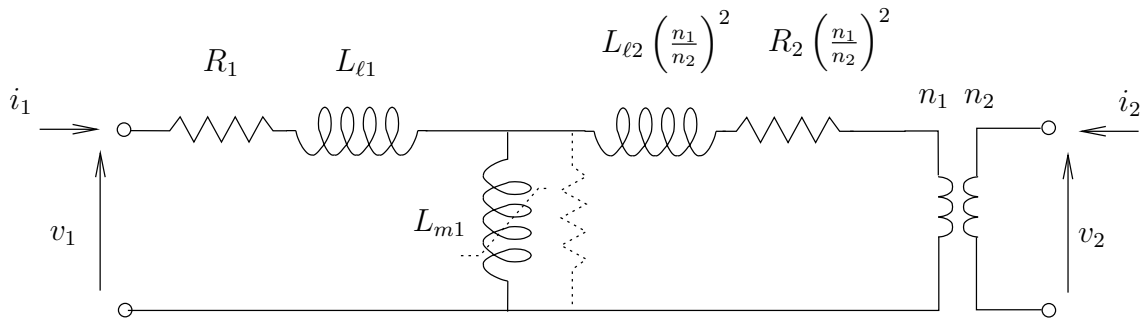


FIGURE 6.4 – schéma équivalent du transformateur monophasé (2ème version)

Le modèle pourrait être raffiné en considérant (voir trait pointillé dans la figure 6.4) :

1. les pertes Joule causées par les courants de Foucault induits dans le matériau magnétique. On désigne couramment ces pertes sous le nom de “pertes fer”, par opposition aux “pertes cuivre” $R_1 i_1^2 + R_2 i_2^2$ subies dans les enroulements. Celles-ci peuvent être prise en compte en plaçant une résistance en parallèle sur l'inductance magnétisante L_{m1}
2. en considérant la saturation du matériau magnétique, via une inductance L_{m1} non linéaire.

Les pertes évoquées au point 1 sont toutefois tout à fait négligeables devant les puissances transitant dans un transformateur de puissance.

En pratique, on utilise couramment un schéma équivalent simplifié dérivé de celui de la figure 6.4. La simplification tient compte du fait que la réactance magnétisante ωL_{m1} est très grande devant

les résistances R_1, R_2 et les réactances de fuite $\omega L_{\ell 1}, \omega L_{\ell 2}$. Rappelons en effet que L_{m1} est associée aux lignes de champ qui passent par le noyau, à forte perméabilité magnétique, tandis que $L_{\ell 1}$ et $L_{\ell 2}$ correspondent aux lignes de champ à l'extérieur du noyau, c'est-à-dire dans un milieu non magnétique (isolant entourant les conducteurs, huile du transformateur). Les ordres de grandeur donnés à la section 6.3 confirmeront cette assertion. Dans ces conditions, on commet une erreur très faible en déplaçant la branche magnétisante à l'entrée 1 du schéma équivalent, ce qui conduit au schéma de la figure 6.5, dans lequel :

$$\begin{aligned} n &= \frac{n_2}{n_1} \\ R &= R_1 + \frac{R_2}{n^2} \\ X &= \omega L_{\ell 1} + \frac{\omega L_{\ell 2}}{n^2} \\ X_m &= \omega L_{m1} \end{aligned}$$

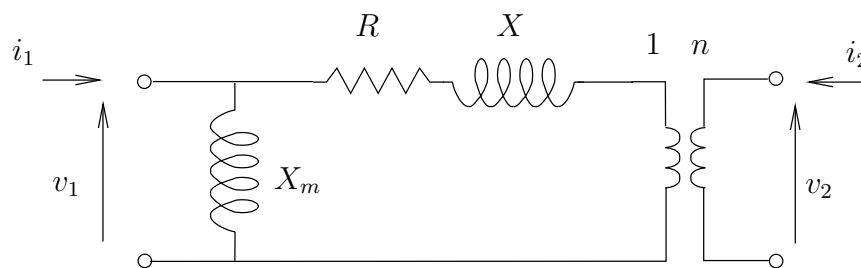


FIGURE 6.5 – schéma équivalent simplifié usuel du transformateur monophasé

Une des raisons pratiques d'utiliser ce schéma est que les mesures classiquement effectuées sur un transformateur (essais à vide et en court-circuit : voir travaux pratiques) ne permettent d'identifier que les paramètres combinés R et X de la figure 6.5 et non les paramètres individuels apparaissant à la figure 6.4.

X (resp. R) est la réactance de fuite (resp. la résistance) “cumulée”. Compte tenu du fait que $X_m \gg X$, $R + jX$ est l'impédance vue de l'accès 1 lorsque l'accès est court-circuité. C'est la raison pour laquelle X est aussi appelée *réactance de court-circuit* du transformateur.

6.1.3 Matrice d'inductance

Le transformateur est évidemment constitué de deux circuits magnétiquement couplés, dont la représentation en Théorie des circuits est donnée à la figure 6.6. Rappelons que dans ce schéma les points noirs • signifient que si les courants entrent par les bornes ainsi marquées, leurs contributions au flux commun ϕ_m s'ajoutent (au lieu de se soustraire l'un de l'autre). En revenant à la figure 6.1, on vérifie (p.ex. par la règle “du tire-bouchon”) que les courants i_1 et i_2 donnent dans le noyau des flux qui s'ajoutent². Le marquage dans la figure 6.6 correspond donc bien à la disposition de la figure 6.1.

2. c'est la raison pour laquelle on a considéré le signe + dans l'équation (6.1)

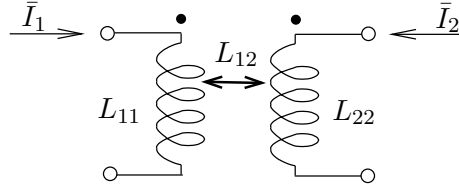


FIGURE 6.6 – circuits magnétiquement couplés

Rappelons aussi que l'inductance mutuelle a un signe. Si les courants sont comptés positivement en entrant par les bornes marquées du point noir •, l'inductance mutuelle est positive. C'est le cas dans la figure 6.6.

Enfin, on montre aisément que le marquage • indique des tensions qui sont en phase lorsque le transformateur est supposé idéal. La différence entre le potentiel de la borne marquée et celui de la borne non marquée au primaire est en phase avec la différence de potentiel correspondante au secondaire.

En régime non saturé, les flux reliés aux courants i_1, i_2 via la matrice d'inductance :

$$\begin{bmatrix} \psi_1 \\ \psi_2 \end{bmatrix} = \begin{bmatrix} L_{11} & L_{12} \\ L_{12} & L_{22} \end{bmatrix} \begin{bmatrix} i_1 \\ i_2 \end{bmatrix}$$

Une comparaison avec les relations (6.9,6.10) fournit directement les termes de la matrice en question :

$$L_{11} = L_{\ell 1} + L_{m1} = L_{\ell 1} + \frac{n_1^2}{\mathcal{R}} \quad (6.11)$$

$$L_{12} = \frac{n_2}{n_1} L_{m1} = \frac{n_1 n_2}{\mathcal{R}} \quad (6.12)$$

$$L_{22} = L_{\ell 2} + \left(\frac{n_2}{n_1} \right)^2 L_{m1} = L_{\ell 2} + \frac{n_2^2}{\mathcal{R}} \quad (6.13)$$

6.1.4 Matrice d'admittance et schéma équivalent en pi

En régime sinusoïdal, de nombreux calculs de réseaux font appel à la matrice d'admittance. Celle du quadripole de la figure 6.5 s'établit aisément comme suit.

Les relations entre tensions et courants tirées de la figure en question sont :

$$\bar{I}_1 = \frac{\bar{V}_1}{jX_m} + \frac{\bar{V}_1 - \frac{\bar{V}_2}{n}}{R + jX} = \left(\frac{1}{jX_m} + \frac{1}{R + jX} \right) \bar{V}_1 - \frac{1}{n} \frac{1}{R + jX} \bar{V}_2 \quad (6.14)$$

$$\bar{I}_2 = \frac{1}{n} \frac{\frac{\bar{V}_2}{n} - \bar{V}_1}{R + jX} \quad (6.15)$$

dont on déduit la matrice d'admittance :

$$\mathbf{Y} = \begin{bmatrix} \frac{1}{jX_m} + \frac{1}{R+jX} & -\frac{1}{n} \frac{1}{R+jX} \\ -\frac{1}{n} \frac{1}{R+jX} & \frac{1}{n^2} \frac{1}{R+jX} \end{bmatrix} \quad (6.16)$$

Il est parfois opportun de remplacer le schéma de la figure 6.5 par un schéma équivalent en pi. Le lecteur est invité à vérifier que l'on obtient le schéma de la figure 6.7 où les paramètres sont cette fois des admittances.

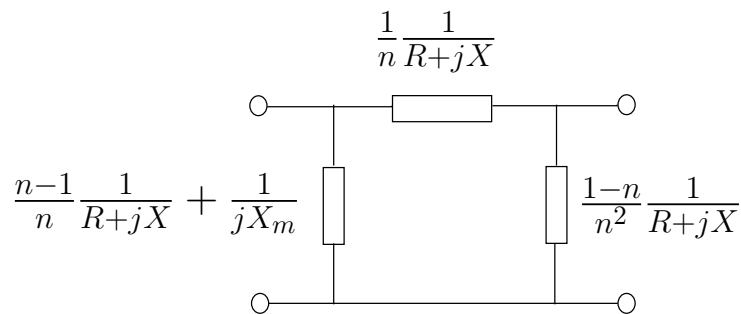


FIGURE 6.7 – schéma équivalent en pi du transformateur

6.2 Transformateur triphasé

6.2.1 Constitution

Un transformateur triphasé est un assemblage de trois transformateurs monophasés. Cet assemblage peut être réalisé de deux manières.

La première configuration fait appel à trois transformateurs monophasés séparés. La figure 6.8 se rapporte à un montage en étoile d'un côté et en triangle de l'autre. En régime équilibré, chaque transformateur voit évidemment transiter un tiers de la puissance totale. Cet agencement offre deux avantages : (i) en cas de défaillance d'une des phases, il suffit de remplacer le transformateur monophasé concerné et non la totalité du transformateur³ ; (ii) dans le cas de postes situés dans des régions difficilement accessibles (p.ex. à cause du gabarit routier), il est plus aisé de transporter trois unités monophasées qu'une unité triphasée.

L'autre agencement utilise un noyau magnétique commun aux trois phases. Le principe est le suivant. Considérons trois transformateurs monophasés du type représenté à la figure 6.9, où les parties en gris clair et gris foncé représentent les enroulements primaire et secondaire, respectivement. La figure 6.10.a les représente vus du dessus. A la figure 6.10.b ils ont été fusionnés en

3. dans un poste comportant plusieurs transformateurs identiques, on peut disposer d'une unité monophasée de réserve

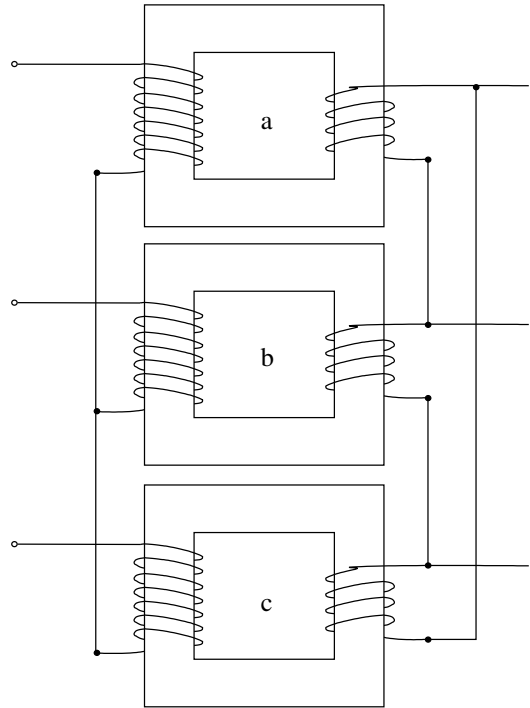


FIGURE 6.8 – assemblage de trois transformateurs monophasés

un seul noyau en respectant une parfaite symétrie. Cependant, dans cette configuration, la colonne centrale n'est le siège d'aucun flux magnétique, car les champs créés par les différents enroulements s'annulent sous l'effet du déphasage des courants⁴. On peut donc retirer cette colonne, ce qui conduit au schéma de la figure 6.10.c.

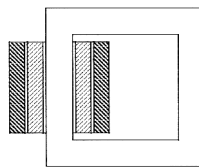


FIGURE 6.9 – disposition des enroulements sur le noyau d'un transformateur monophasé

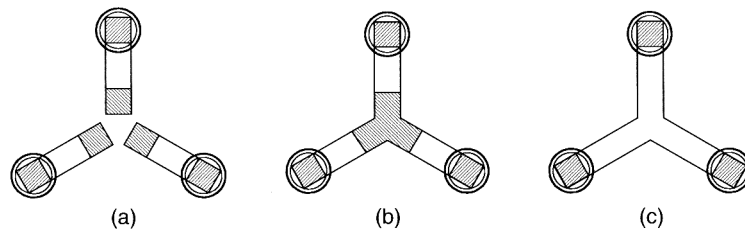


FIGURE 6.10 – création d'un noyau unique à partir de trois noyaux séparés

4. c'est le même raisonnement qui a conduit à éliminer le conducteur de retour dans le schéma de la figure 2.3

Cependant, le transformateur montré à la figure 6.10.c serait difficile à réaliser et encombrant à transporter. On préfère disposer les trois colonnes dans un même plan, ce qui conduit à la configuration de la figure 6.11.a, qui est désignée en anglais par le vocable “core”. Toutefois, cette configuration présente l’inconvénient de ne pas être symétrique vis-à-vis des trois phases. Cette dissymétrie est atténuée en adoptant la structure montrée à la figure 6.11.b, qui est désignée en anglais par le vocable “shell”.

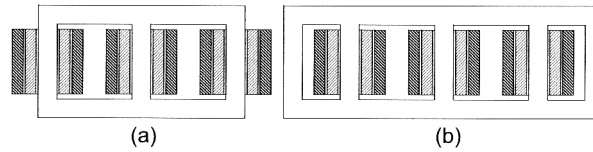


FIGURE 6.11 – disposition des enroulements dans un transformateur triphasé : (a) core, (b) shell

Avec un noyau commun aux trois phases, le volume de matériau magnétique est inférieur au volume total des trois noyaux séparés, d’où un coût de fabrication moindre.

6.2.2 Choix du montage

Il y a quatre configurations possibles, selon que l’on connecte les enroulements primaires et secondaires en triangle ou en étoile. Parmi les critères qui guident le choix d’une configuration, citons :

- quand le transformateur est destiné à un réseau de transport THT, on préfère le montage en étoile car les tensions aux bornes de ses enroulements sont $\sqrt{3}$ fois plus petites ;
- le montage en étoile offre la possibilité de connecter le neutre à la terre ;
- il est aussi préféré pour le placement d’un régulateur en charge : voir section 6.5 ;
- quand le transformateur doit être parcouru par des courants importants (p.ex. le côté générateur d’un transformateur élévateur) on préfère le montage en triangle, dans les branches duquel les courants sont $\sqrt{3}$ fois plus petits ;
- le montage en triangle est utilisé pour éliminer les harmoniques d’ordre multiple de 3 (voir exercice 2 du chapitre 2).

6.2.3 Schémas équivalents monophasés

Dans la disposition de la figure 6.8, les phases ne sont pas couplées magnétiquement ; par contre, elles le sont dans le cas de la figure 6.11. Dans ce dernier cas, il convient de procéder à une analyse par phase en établissant la matrice d’impédance relatives aux trois phases et en se ramenant à une seule d’entre elles, comme expliqué à la section 2.4. Cette technique étant bien acquise, nous ne ferons pas de développement dans ce sens et supposons pour simplifier que le transformateur est constitué de trois transformateurs monophasés séparés.

Par contre, nous allons considérer l’impact sur le schéma équivalent par phase des montages en étoile ou en triangle.

Montage étoile-étoile

Le schéma équivalent du transformateur intervenant dans chaque phase est donné à la figure 6.5. En assemblant ces schémas équivalents en étoile-étoile, on aboutit aisément au schéma de la figure 6.12. Cette figure donne également les phaseurs des tensions au primaire et au secondaire de chaque transformateur idéal.

Remarques. Pour les besoins du dessin, les deux branches symbolisant un transformateur idéal ont été éloignées l'une de l'autre. Cependant, les deux branches qui se correspondent ont été dessinées parallèles. De plus, leur direction rappelle celle du phaseur des tensions à leurs bornes.

Le schéma équivalent monophasé de l'ensemble s'obtient en ne considérant tout simplement qu'une des trois phases, ce qui nous ramène évidemment au schéma de la figure 6.5.

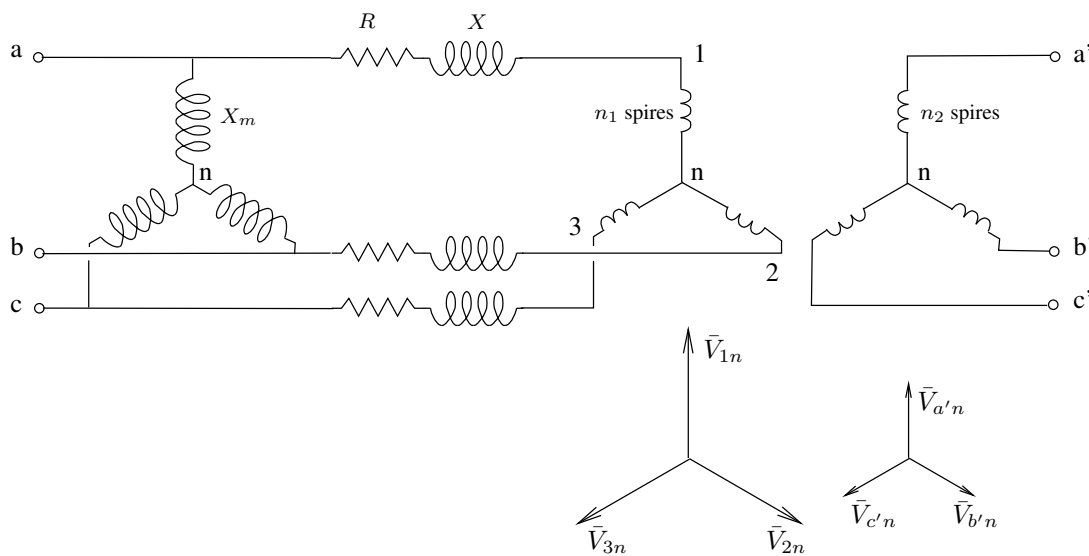


FIGURE 6.12 – schéma équivalent du transformateur triphasé étoile-étoile

Montage triangle-triangle

En partant à nouveau du schéma équivalent du transformateur intervenant dans chaque phase (cf fig. 6.5) et en les assemblant en triangle-triangle, on aboutit au schéma de la figure 6.13.

On constate que ce schéma est en fait la mise en parallèle de deux triangles : le premier comporte dans chacune de ses branches une impédance X_m , le second une impédance $R + jX$ en série avec un transformateur idéal de rapport $n = n_2/n_1$.

Pour se ramener au schéma équivalent monophasé, il faut transformer ces deux triangles en étoiles. Conformément à la formule (2.13), l'impédance qui intervient dans cette étoile est le tiers de celle intervenant dans le triangle. Par ailleurs, les tensions phase-neutre au primaire et au secondaire

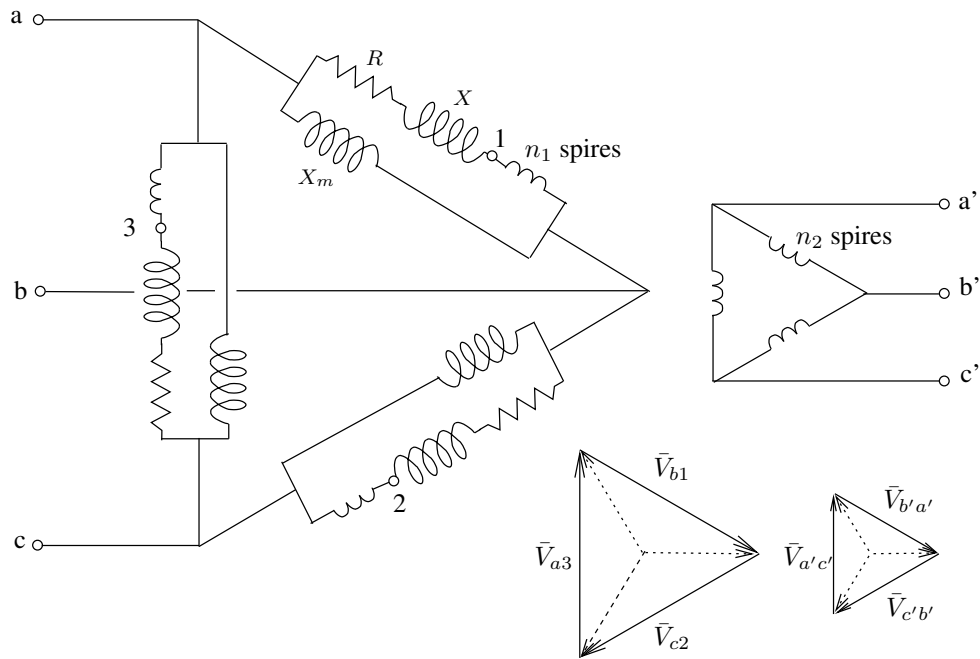


FIGURE 6.13 – schéma équivalent du transformateur triphasé triangle-triangle

sont dans le même rapport que les tensions entre phases correspondantes, soit n_2/n_1 . Ceci conduit au schéma équivalent de la figure 6.14.

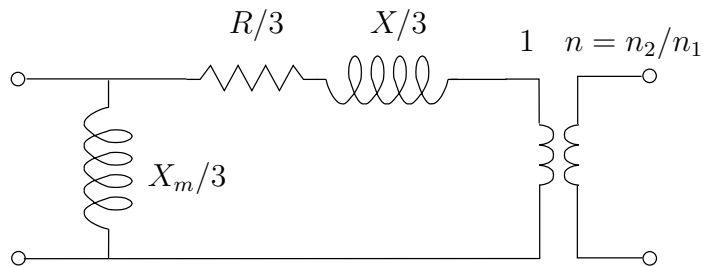


FIGURE 6.14 – schéma équivalent monophasé du transformateur triphasé triangle-triangle

Montage étoile-triangle

Par combinaison des deux configurations précédentes, on obtient directement le schéma de la figure 6.15. Cette dernière montre également les phaseurs des tensions et des courants.

On déduit du phaseur des tensions que :

$$\bar{V}_{a'} = \frac{1}{\sqrt{3}} e^{j\pi/6} \bar{V}_{a'c'} = \frac{n_2}{\sqrt{3}n_1} e^{j\pi/6} \bar{V}_{1n} = \bar{n} \bar{V}_{1n} \quad (6.17)$$

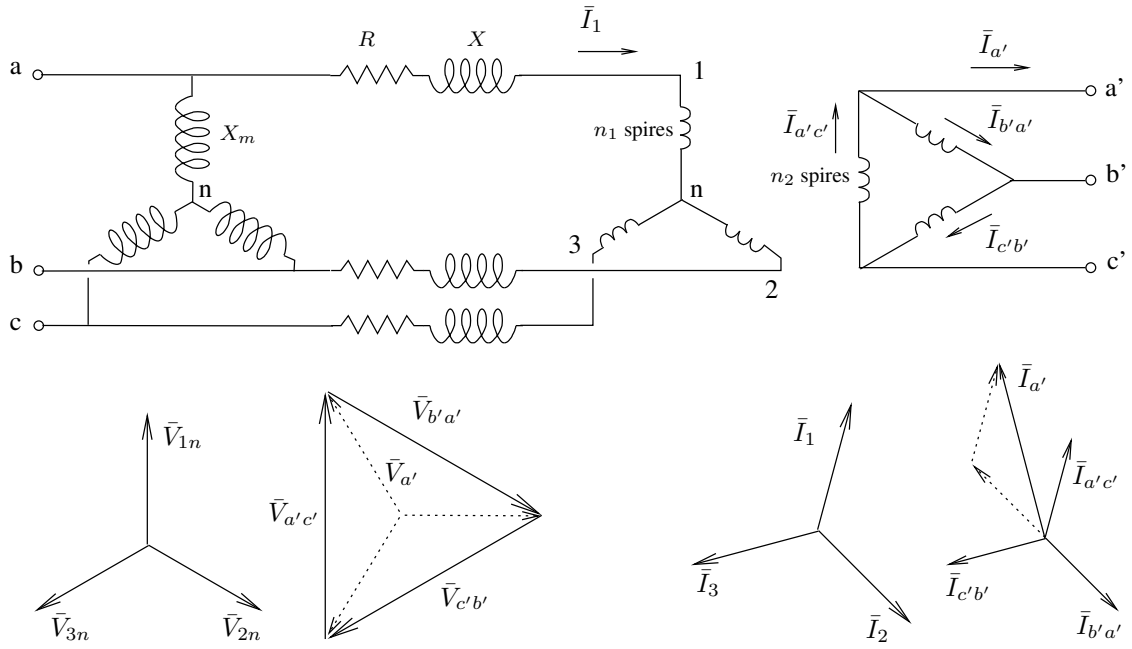


FIGURE 6.15 – schéma équivalent du transformateur triphasé étoile-triangle

où l'on a posé :

$$\bar{n} = \frac{n_2}{\sqrt{3} n_1} e^{j\pi/6} \quad (6.18)$$

On voit qu'il existe une avance de phase de 30° et une division par $\sqrt{3}$ lorsque l'on passe de la tension primaire (phase-neutre) à la tension secondaire (phase-neutre) du transformateur idéal. Nous reviendrons plus loin sur les conséquences de ce déphasage.

Pour les courants, on établit de même que :

$$\bar{I}_{a'} = \bar{I}_{a'c'} - \bar{I}_{b'a'} = \sqrt{3} e^{j\pi/6} \bar{I}_{a'c'} = \frac{\sqrt{3} n_1}{n_2} \frac{1}{e^{-j\pi/6}} \bar{I}_1 = \frac{1}{\bar{n}^*} \bar{I}_1 \quad (6.19)$$

Ces relations étant établies, on obtient aisément le schéma équivalent monophasé de la figure 6.16. Notons que ce dernier fait intervenir un *transformateur idéal à rapport $\bar{n}$ complexe*, dont les relations caractéristiques sont rappelées en marge de la figure.

Le transformateur à rapport complexe est une abstraction qui permet de tenir compte du déphasage de 30° créé par l'utilisation d'un montage mixte étoile-triangle. Notons les propriétés suivantes :

1. dans le cas où n est réel, on retrouve le transformateur idéal "classique"
2. il y a conservation de la puissance complexe :

$$\bar{V}_{a'} \bar{I}_{a'}^* = \bar{n} \bar{V}_1 \frac{1}{\bar{n}} \bar{I}_1^* = \bar{V}_1 \bar{I}_1^*$$

3. le quadripôle de la figure 6.16 n'est *pas réciproque* :

$$\bar{I}_a]_{\bar{V}_a=0, \bar{V}_{a'}=1} \neq -\bar{I}_{a'}]_{\bar{V}_a=1, \bar{V}_{a'}=0}$$

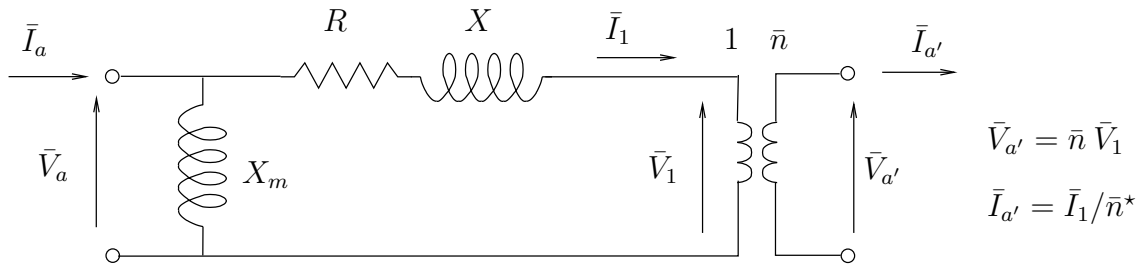


FIGURE 6.16 – schéma équivalent monophasé du transformateur triphasé étoile-triangle

Il en résulte que la matrice d'admittance de ce même quadripôle n'est *pas symétrique* :

$$Y_{a a'} \neq Y_{a' a}$$

et qu'il n'est pas possible d'établir un schéma équivalent en pi du type de celui de la figure 6.7 (relative au cas où n est réel).

Montage triangle-étoile

Ce cas est semblable au précédent, sauf que le schéma équivalent monophasé comporte à présent le rapport de transformation complexe :

$$\bar{n} = \frac{\sqrt{3} n_2}{n_1} e^{-j\pi/6} \quad (6.20)$$

et les résistances et réactances $R/3$, $X/3$, $X_m/3$ (comme à la figure 6.14).

6.2.4 Désignation des transformateurs et indice horaire

Pour caractériser complètement le couplage des enroulements d'un transformateur triphasé, on utilise une abréviation standardisée par l'I.E.C.⁵, caractérisée au minimum par 3 symboles :

- une lettre majuscule pour le côté haute tension : Y pour étoile ou D pour triangle
- une lettre minuscule pour le côté basse tension : y pour étoile ou d pour triangle
- un nombre compris entre 0 et 11, appelé *indice horaire*, caractérisant le déphasage entre tensions primaire et secondaire relatives à une même phase, le transformateur étant supposé idéal. Ce nombre est obtenu en plaçant sur un cadran d'horloge, le phaseur de la haute tension sur le nombre 12 et en lisant le nombre pointé par le phaseur de la basse tension.

Dans le cas d'un montage en étoile, on ajoute la lettre "n" après "Y" ou "y" pour indiquer que le neutre est mis à la terre.

Lorsqu'un même sous-réseau est alimenté par deux transformateurs (ou plus) "en parallèle" (c'est-à-dire qu'il existe au moins une maille incluant ces deux transformateurs), ceux-ci doivent avoir

5. International Electrotechnical Commission

le même indice horaire, sous peine de créer des déphasages et donc des transits de puissance très importants que l'équipement ne pourrait supporter. Par exemple, la figure 6.17 montre deux situations autorisées et une interdite.

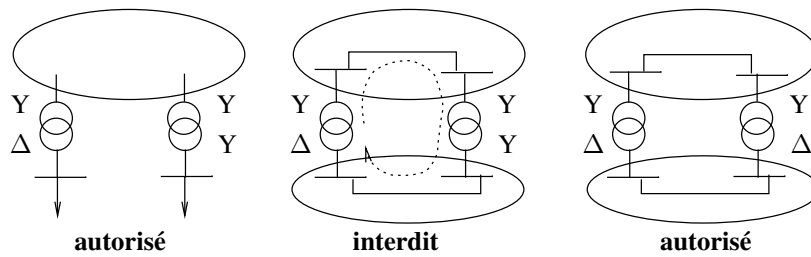


FIGURE 6.17 – connexion de transformateurs en parallèle

6.2.5 Simplification des calculs

Considérons le cas d'un réseau R_1 relié à un réseau R_2 via deux transformateurs. Les modules n_A et n_B de leurs rapports de transformation peuvent être différents mais, comme indiqué précédemment, les deux transformateurs doivent avoir le même groupe horaire ; les phases des rapports de transformation sont donc égales : $\varphi_A = \varphi_B = \varphi$.

La figure 6.18 fait apparaître le schéma équivalent de chacun d'entre eux, dans lequel le transformateur idéal à rapport complexe $\bar{n}_A = n_A e^{j\varphi_A}$ (resp. $\bar{n}_B = n_B e^{j\varphi_B}$) a été remplacé par un transformateur idéal à rapport réel $n_A e^{j0}$ (resp. $n_B e^{j0}$) en cascade avec un transformateur idéal à rapport complexe $e^{j\varphi_A}$ (resp. $e^{j\varphi_B}$).

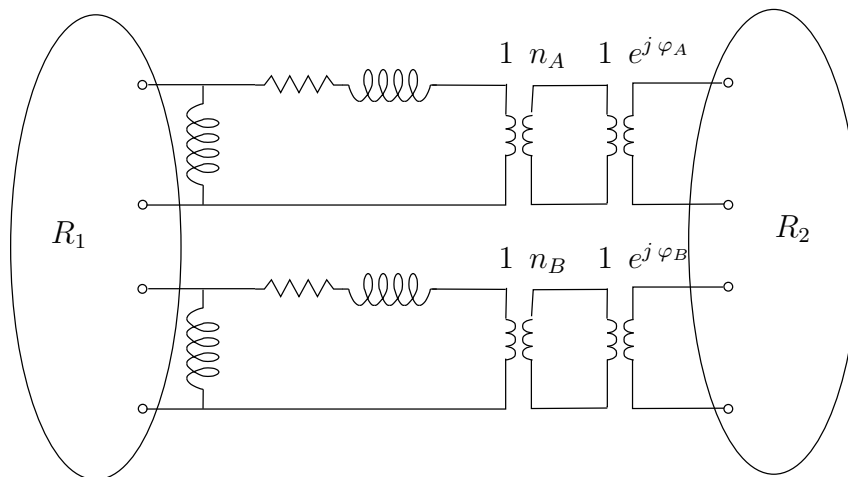


FIGURE 6.18 – schéma équivalent d'un ensemble de transformateurs en parallèle

Lors de l'analyse d'un tel système, en régime triphasé équilibré, on peut *supprimer tous les transformateurs idéaux de rapport $e^{j\varphi}$* sans modifier les amplitudes des courants et des tensions, ni les puissances transitant dans les branches. En effet, cette simplification conduit à déphaser la

tension $\bar{V}_i$ en tout noeud de R_2 , de φ radians par rapport à la réalité. Le courant $\bar{I}_j$ dans une quelconque branche de R_2 , fonction de la différence des tensions à ses extrémités, est déphasé de la même quantité φ . Les puissances complexes $\bar{V}_i \bar{I}_j^*$ transitant dans les branches de R_2 sont donc inchangées. Enfin, les puissances complexes transitant dans les transformateurs reliant R_1 et R_2 sont elles aussi inchangées, les transformateurs idéaux supprimés étant sans pertes.

6.3 Valeurs nominales, système per unit et ordres de grandeur

Un transformateur est caractérisé par diverses valeurs nominales :

- les tensions nominales primaire U_{1N} et secondaire U_{2N} . Ce sont les tensions pour lesquelles l'isolation des enroulements a été prévue. Un certain dépassement de ces valeurs est toutefois admissible en pratique. Sauf mention contraire, il s'agit de tensions entre phases en valeurs efficaces ;
- les courants nominaux primaire I_{1N} et secondaire I_{2N} . Ce sont les courants pour lesquels les sections des conducteurs ont été prévues. Il s'agit donc des courants maxima admissibles pendant un temps infini. Sauf mention contraire, il s'agit de courants de phase en valeurs efficaces ;
- la puissance apparente nominale S_N . Il s'agit de la puissance triphasée liée aux grandeurs précédentes par la relation (6.6) du transformateur idéal :

$$S_N = \sqrt{3} U_{1N} I_{1N} = \sqrt{3} U_{2N} I_{2N}$$

La conversion des paramètres du transformateur en per unit se fait conformément aux considérations des sections 5.2 et 5.3 :

- on choisit la puissance de base $S_B = S_N$
- dans la partie connectée au primaire, on choisit la tension de base $V_{1B} = U_{1N}/\sqrt{3}$ (valeur efficace entre phase et neutre)
- dans la partie connectée au secondaire, on choisit la tension de base $V_{2B} = U_{2N}/\sqrt{3}$ (valeur efficace entre phase et neutre)
- les impédances du schéma équivalent de la figure 6.5 se situant au primaire, on les divise par l'impédance de base $Z_{1B} = 3V_{1B}^2/S_B$
- le rapport $n = n_2/n_1$ se convertit comme suit. Le transformateur idéal étant caractérisé par :

$$v_2 = \frac{n_2}{n_1} v_1$$

on divise cette relation par V_{2B} , ce qui donne :

$$v_{2pu} = \frac{v_2}{V_{2B}} = \frac{n_2}{n_1} \frac{V_{1B}}{V_{2B}} \frac{v_1}{V_{1B}} = \frac{n_2}{n_1} \frac{V_{1B}}{V_{2B}} v_{1pu}$$

En valeurs unitaires, le rapport de transformation vaut donc :

$$n_{pu} = \frac{n_2}{n_1} \frac{V_{1B}}{V_{2B}}$$

Si les tensions nominales sont dans le rapport des nombres de spires, c'est-à-dire si :

$$\frac{V_{2B}}{V_{1B}} = \frac{n_2}{n_1}$$

on a simplement :

$$n_{pu} = 1$$

et le transformateur idéal disparaît du schéma de la figure 6.5. C'est un des avantages, déjà mentionné, du système per unit.

En pratique, le rapport V_{2B}/V_{1B} est proche, mais pas égal à n_2/n_1 , surtout lorsque le transformateur est doté d'un réglage du nombre de spires, comme expliqué à la section 6.5. Dans ce cas, il reste dans le schéma équivalent un transformateur idéal avec un rapport n_{pu} proche de l'unité.

On peut citer les ordres de grandeur suivants pour les transformateurs utilisés dans les réseaux de transport :

résistance R	< 0.005 pu
réactance de fuite ωL	de 0.06 à 0.20 pu
réactance magnétisante ωL_m	de 20 à 50 pu
rapport de transformation $n = n_2/n_1$	de 0.85 à 1.15 pu

Ces valeurs s'entendent dans la base du transformateur, telle que définie plus haut. Dans un calcul de réseau utilisant une autre base, il y a lieu de procéder à une conversion, en utilisant par exemple la formule (5.11).

6.4 Autotransformateur

Un autotransformateur est un transformateur dont on connecte le primaire et le secondaire de sorte qu'ils aient un enroulement en commun. Un schéma de principe est donné à la figure 6.19, avec le nombre de spires de chaque enroulement.

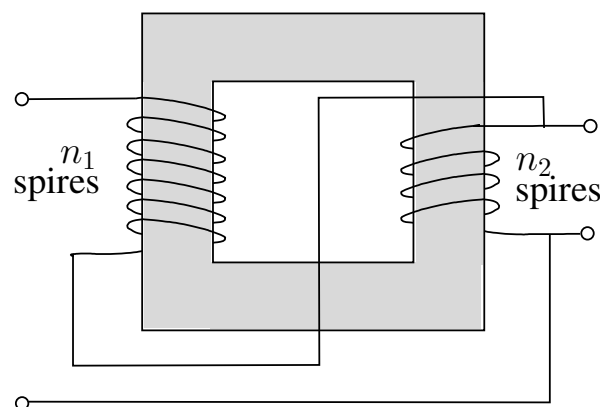


FIGURE 6.19 – schéma de principe d'un autotransformateur monophasé

Considérons le fonctionnement du transformateur sous ses tensions et courant nominaux et déterminons les valeurs nominales de l'autotransformateur correspondant. Comme nous nous intéressons seulement au régime nominal, nous supposons le transformateur idéal.

La situation est détaillée à la figure 6.20. Au régime nominal, les tensions aux bornes des enroulements sont respectivement V_{1N} et V_{2N} tandis que les courants qui les parcourent sont I_{1N} et I_{2N} . On en déduit les autres tensions et courants indiqués par des flèches en pointillé à la figure 6.20.

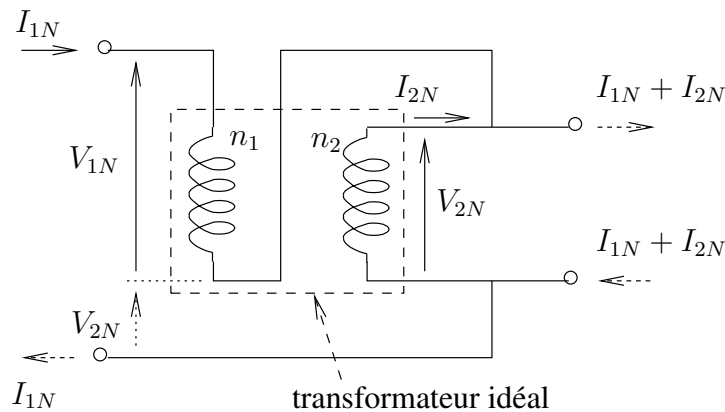


FIGURE 6.20 – régime nominal de l'autotransformateur

On en déduit les courants et tensions nominaux de l'autotransformateur :

$$\begin{aligned} V_{1N}^{auto} &= V_{1N} + V_{2N} = \left(1 + \frac{n_2}{n_1}\right)V_{1N} \\ I_{1N}^{auto} &= I_{1N} \\ V_{2N}^{auto} &= V_{2N} \\ I_{2N}^{auto} &= I_{1N} + I_{2N} = \left(\frac{n_2}{n_1} + 1\right)I_{2N} \end{aligned}$$

Le rapport de transformation de l'autotransformateur idéal vaut donc :

$$n^{auto} = \frac{V_{2N}^{auto}}{V_{1N}^{auto}} = \frac{V_{2N}}{V_{1N} + V_{2N}} = \frac{\frac{n_2}{n_1}}{1 + \frac{n_2}{n_1}} \quad (6.21)$$

On voit que, quelles que soient les valeurs de n_1 et de n_2 , le transformateur est du type abaisseur, pour le primaire et le secondaire choisis à la figure 6.20.

La puissance nominale apparente S_N^{auto} de l'autotransformateur est reliée à la puissance correspondante S_N du transformateur par :

$$S_N^{auto} = V_{1N}^{auto} I_{1N}^{auto} = \left(1 + \frac{n_2}{n_1}\right)V_{1N} I_{1N} = \left(1 + \frac{n_2}{n_1}\right)S_N \quad (6.22)$$

Cette dernière relation montre qu'il est possible de faire transiter dans l'autotransformateur une puissance plus grande que dans le transformateur sur lequel s'appuie le montage. Il en résulte évidemment des économies au niveau du coût de fabrication mais aussi des pertes. Ceci est vrai quelles que soient les valeurs de n_1 et de n_2 , mais pour obtenir une "amplification" plus grande, il

faut que n_2 soit nettement supérieur à n_1 . Cependant, dans ce cas, la relation (6.21) montre que le rapport de transformation tend vers l'unité. Il est alors impossible de relier des niveaux de tensions nominales fort différentes.

On utilise un autotransformateur lorsqu'il est nécessaire de faire transiter de grandes puissances entre deux niveaux de tension relativement proches. Par exemple, c'est le cas, en Belgique, des transformateurs de 550 MVA qui connectent les réseaux à 400 et à 150 kV et, en France, de ceux qui relient les niveaux à 400 et 225 kV.

L'autotransformateur présente cependant l'inconvénient d'avoir une connexion métallique entre ses deux enroulements : des perturbations de tension peuvent donc passer plus facilement d'un enroulement à l'autre que dans le cas du seul couplage par induction.

On obtient un autotransformateur triphasé en assemblant trois autotransformateurs monophasés.

6.5 Ajustement du nombre de spires d'un transformateur

6.5.1 Principe

Il est souvent utile de pouvoir modifier le nombre de spires d'un transformateur, en vue d'ajuster les tensions au voisinage de ce dernier.

Ce réglage est discret par nature : un transformateur présente typiquement entre 15 et 25 *prises de réglage*, comme symbolisé à la figure 6.21.

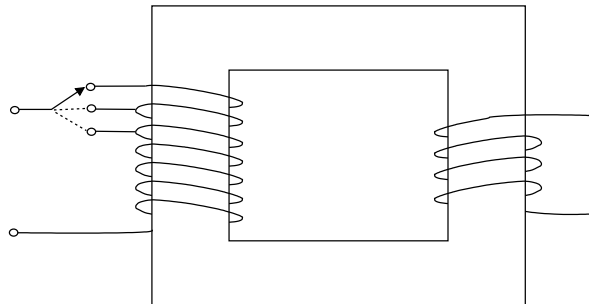


FIGURE 6.21 – principe de la modification de la prise de réglage

Dans certains transformateurs, la modification requiert de mettre l'appareil hors service et de changer manuellement ses connexions.

Dans de nombreux transformateurs, cependant, cette modification peut être effectuée *en charge* c'est-à-dire sans interrompre le courant qui parcourt l'enroulement dont on modifie le nombre de spires. Le dispositif correspondant, appelé *régleur en charge*, comporte un contacteur conçu

pour éviter la formation d'arcs électriques (susceptibles d'endommager les contacts) et un moteur électrique pour entraîner ce contacteur.

Notons enfin qu'un régleur en charge peut être commandé manuellement (en fait télécommandé par l'opérateur depuis un centre de conduite) ou automatiquement (système asservi commandant localement le régleur en charge).

En général, le régleur en charge modifie les spires de l'enroulement dont la tension nominale est la plus élevée (primaire d'un transformateur abaisseur) parce que :

- les courants y sont plus faibles ; la commutation est donc plus aisée
- les spires plus nombreuses ; le réglage est donc plus fin.

Il peut toutefois exister des exceptions à cette règle.

Dans un transformateur triphasé, on place généralement le régleur en charge dans un montage en étoile et du côté du neutre de manière à ce que l'appareil soit soumis à des tensions plus basses.

6.5.2 Prise en compte dans le schéma équivalent

A chaque changement de prise, il est clair que le rapport de transformation n_2/n_1 se modifie. En principe, tous les paramètres du schéma équivalent du transformateur se modifient. Cette variation est surtout significative pour l'inductance de fuite L (cf figure 6.5) car la résistance R étant faible et l'inductance L_m très élevée, les variations de ces deux paramètres sont sans grand effet.

Dans un logiciel de calcul, il est possible de spécifier des valeurs de L et de n_2/n_1 pour chaque prise. Une simplification est toutefois possible, comme expliqué ci-après.

Supposons que l'on ajuste les spires de l'enroulement secondaire. En admettant, en première approximation, que l'inductance de fuite $L_{\ell 2}$ varie comme le carré du nombre de spires n_2 , on a :

$$L_{\ell 2} = L_{\ell 2}^o \left(\frac{n_2}{n_2^o} \right)^2$$

où $L_{\ell 2}^o$ est la valeur de l'inductance de fuite lorsque $n_2 = n_2^o$. Supposons de plus que la résistance R_2 varie de la même manière, soit :

$$R_2 = R_2^o \left(\frac{n_2}{n_2^o} \right)^2$$

Cette dernière hypothèse est beaucoup plus difficile à admettre mais, comme a l'a indiqué plus haut, les conséquences de cette simplification sont mineures.

Remplaçons, dans le schéma équivalent de la figure 6.3, R_2 et $L_{\ell 2}$ par les expressions ci-dessus et faisons-les passer de l'autre côté du transformateur idéal, moyennant multiplication par $(n_1/n_2)^2$. On obtient le schéma équivalent de la figure 6.22, analogue à celui de la figure 6.4. On voit que les impédances à gauche du transformateur idéal sont toutes indépendantes du nombre de spires n_2 .

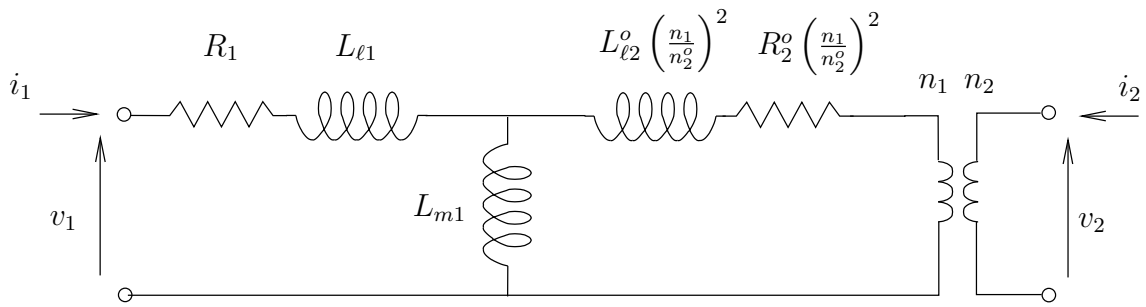


FIGURE 6.22 – autre forme du schéma équivalent de la figure 6.4

En conclusion, sous les hypothèses énoncées plus haut, on peut garder constantes les impédances du schéma équivalent, à condition de les placer du côté du transformateur idéal où le nombre de spires n'est pas modifié⁶. Seul le rapport de transformation change alors avec la prise.

6.6 Transformateur à trois enroulements

On utilise fréquemment des transformateurs “à trois enroulements”. C'est un raccourci de langage pour “trois enroulements par phase”.

Un transformateur monophasé à trois enroulements est représenté schématiquement à la figure 6.23. Cet appareil permet de transférer de la puissance entre trois niveaux de tension, le sens de transfert de la puissance dépendant de ce qui est connecté au transformateur. Dans un transformateur triphasé, on retrouve ces trois enroulements dans chaque phase.

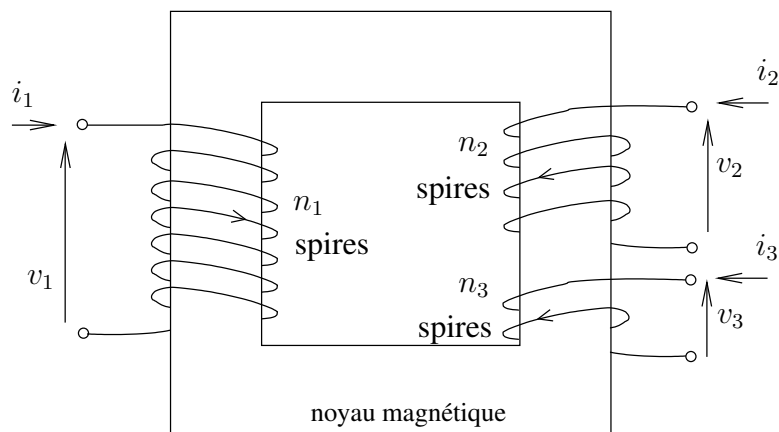


FIGURE 6.23 – transformateur monophasé à trois enroulements

Dans un transformateur monophasé et dans chaque phase d'un transformateur triphasé “à deux enroulements”, les enroulements primaire et secondaire ont la même puissance nominale étant

6. ce qui est toujours possible, pourvu que l'on choisisse le primaire du bon côté

donné que la puissance qui entre par l'un ressort par l'autre, aux pertes près. Dans un transformateur à trois enroulements, les puissances nominales des divers enroulements sont généralement différentes.

Le troisième enroulement peut également servir :

- à alimenter des auxiliaires dans un poste. Il s'agit alors d'un enroulement de petite puissance nominale par rapport aux deux autres ;
- à connecter une inductance ou un condensateur shunt de compensation ;
- à améliorer le fonctionnement en régime déséquilibré (transformateur triphasé à deux enroulements doté d'un troisième, dont les phases sont connectées en triangle et ne sont parcourues par aucun courant en régime équilibré).

Un schéma équivalent usuel de transformateur à trois enroulements est donné à la figure 6.24. Le noeud O est fictif. Les rapports de transformation permettent de représenter des régleurs en charge. Ils peuvent être complexes pour tenir compte des couplages.

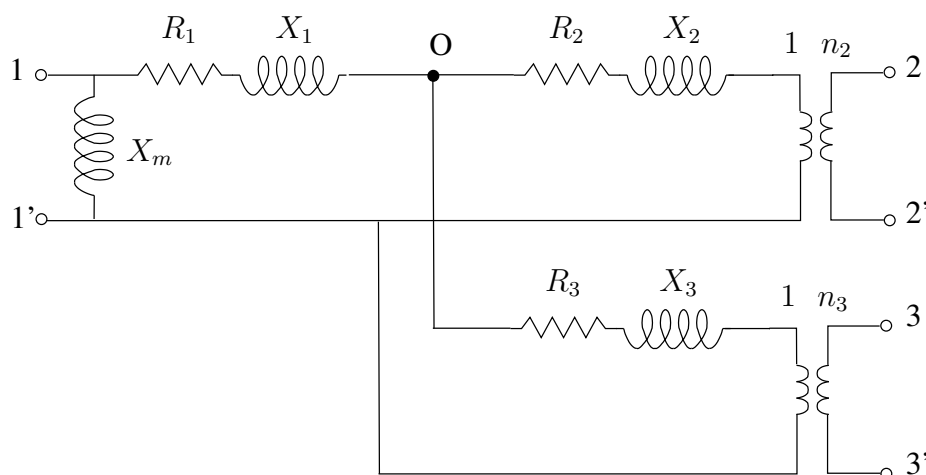


FIGURE 6.24 – schéma équivalent d'un transformateur à trois enroulements

En supposant X_m infinie, on établit aisément que $R_1 + R_2 + j(X_1 + X_2)$ est l'impédance vue de l'accès 1 lorsque l'accès 2 est court-circuité et l'accès 3 laissé ouvert. De même, $R_1 + R_3 + j(X_1 + X_3)$ est l'impédance vue de l'accès 1 lorsque l'accès 3 est court-circuité et l'accès 2 laissé ouvert. Ces propriétés sont utilisées lors de l'établissement du schéma équivalent à partir de mesures (voir Exercice No 3). On notera que dans ce schéma équivalent, certaines réactances peuvent être négatives (dans le cas où les puissances nominales des enroulements sont très différentes).

6.7 Transformateur déphaseur

Le transformateur déphaseur⁷ est un dispositif destiné à déphaser plus ou moins fortement la tension secondaire par rapport à la tension primaire, afin d'ajuster les transits de puissance active dans les branches du réseau. Il peut s'agir :

7. en anglais : “phase shifting transformer” ou “phase shifter”

- soit d'un transformateur triphasé connectant deux niveaux de tensions nominales différentes auquel on ajoute un dispositif de réglage
- soit d'un transformateur triphasé de mêmes tensions nominales au primaire et au secondaire, ne servant qu'à effectuer le réglage de phase.

La figure 6.25 présente un premier schéma de principe. Les enroulements dessinés parallèlement sont montés sur le même noyau magnétique. Comme le montre le diagramme de phaseur, on ajoute à la tension phase-neutre d'entrée une fraction de la tension composée prise entre les deux autres phases. Cette dernière est déphasée de 90 degrés par rapport à la tension d'entrée, d'où le nom de *régler en quadrature* fréquemment utilisé. Dans ce dispositif, le module de la tension de sortie varie légèrement avec le déphasage, d'autant plus que ce dernier est important. Il existe des montages plus élaborés où le module de la tension reste constant au fur et à mesure du réglage du déphasage.

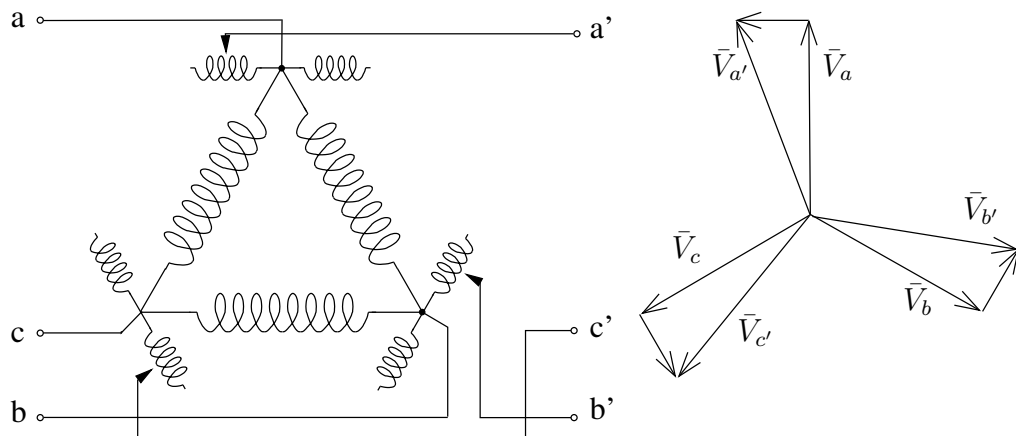


FIGURE 6.25 – transformateur avec réglage en quadrature (1er type)

L'inconvénient du montage de la figure 6.25 réside dans le fait que le régleur en charge (utilisé pour ajuster l'amplitude de la tension en série dans chaque phase) supporte le plein courant de ligne, d'où d'éventuels problèmes de commutation. Pour éviter cet inconvénient, on peut utiliser le montage représenté à la figure 6.26. Comme dans la figure précédente, les enroulements dessinés parallèlement sont montés sur le même noyau magnétique. Dans ce montage, moyennant l'utilisation d'un transformateur supplémentaire en série avec la ligne, le régleur en charge véhicule des courants plus faibles.

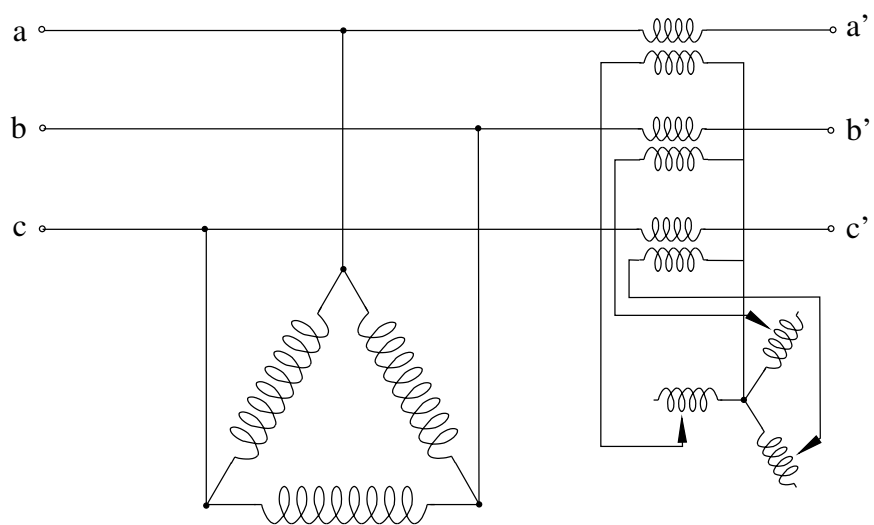


FIGURE 6.26 – transformateur avec réglage en quadrature (2nd type)

Chapitre 7

Le calcul de répartition de charge (ou load flow)

Le *calcul de répartition de charge*, ou encore *calcul d'écoulement de charge* (ou de puissance) est sans aucun doute le calcul le plus fréquemment effectué dans les réseaux d'énergie électrique.

En termes simples, son objectif est de déterminer l'état électrique complet du réseau, à savoir les tensions à tous les noeuds, les transits de puissance dans toutes les branches, les pertes, etc... à partir des consommations et des productions spécifiées en ses noeuds.

On utilise couramment la traduction anglaise "*load flow*". En anglais, le terme "*power flow*" est préféré.

7.1 Les équations de load flow

7.1.1 Quadripôle universel

Afin de simplifier les développements analytiques, nous supposerons toutes les branches du réseau modélisées par le quadripôle représenté à la figure 7.1.

On vérifie aisément que :

- le schéma équivalent en pi de la ligne (ou du câble) de la figure 4.6 s'obtient en posant $n_{ij} = 1$ et $\varphi_{ij} = 0$
- le schéma équivalent du transformateur de la figure 6.16 s'obtient en posant $B_{sji} = 0$; seuls les transformateurs déphaseurs ont un paramètre ϕ_{ij} différent de zéro.

Supposons que le quadripôle relie les noeuds i et j . Le courant $\bar{I}_{ij}$ qui y entre du côté du noeud i

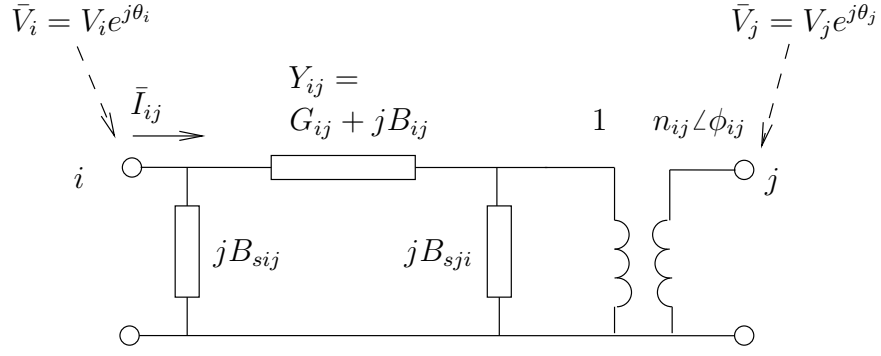


FIGURE 7.1 – quadripôle universel

vaut :

$$\bar{I}_{ij} = jB_{sij}\bar{V}_i + \bar{Y}_{ij}\left(\bar{V}_i - \frac{\bar{V}_j}{n_{ij}}e^{-j\phi_{ij}}\right) = jB_{sij}\bar{V}_i + (G_{ij} + jB_{ij})\left(\bar{V}_i - \frac{\bar{V}_j}{n_{ij}}e^{-j\phi_{ij}}\right) \quad (7.1)$$

7.1.2 Bilans de puissance nodaux

Soit N le nombre de jeux de barres du réseau. On note $\mathcal{N}(i)$ l'ensemble des noeuds reliés au i -ème noeud ($i = 1, \dots, N$) par au moins une branche (cf figure 7.2).

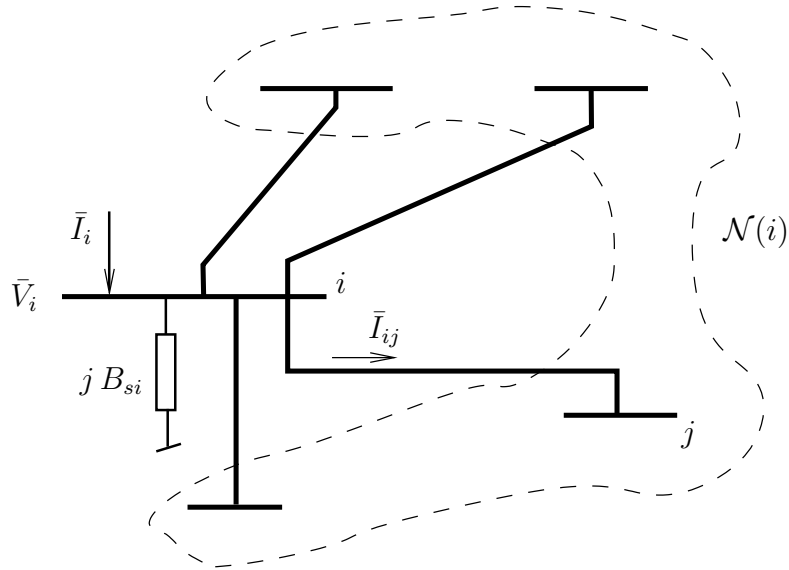


FIGURE 7.2 – topologie au voisinage du i -ème noeud

La première loi de Kirchhoff donne :

$$\bar{I}_i = jB_{si}\bar{V}_i + \sum_{j \in \mathcal{N}(i)} \bar{I}_{ij} \quad i = 1, \dots, N$$

où le courant $\bar{I}_i$ est compté positivement quand il entre dans le réseau.

La puissance complexe entrant dans le réseau au i -ème jeu de barres vaut :

$$P_i + jQ_i = \bar{V}_i \bar{I}_i^* = -jB_{si}V_i^2 + \sum_{j \in \mathcal{N}(i)} \bar{V}_i \bar{I}_{ij}^* \quad (7.2)$$

En remplaçant $\bar{I}_{ij}$ par son expression (7.1), on obtient successivement :

$$\begin{aligned} P_i + jQ_i &= -jB_{si}V_i^2 + \sum_{j \in \mathcal{N}(i)} \bar{V}_i \bar{I}_{ij}^* \\ &= -jB_{si}V_i^2 - j \sum_{j \in \mathcal{N}(i)} B_{sij}V_i^2 + \sum_{j \in \mathcal{N}(i)} \bar{V}_i (G_{ij} - jB_{ij}) \left(V_i^* - \frac{\bar{V}_j^*}{n_{ij}} e^{j\phi_{ij}} \right) \\ &= -jB_{si}V_i^2 - j \sum_{j \in \mathcal{N}(i)} B_{sij}V_i^2 + \sum_{j \in \mathcal{N}(i)} (G_{ij} - jB_{ij}) \left(V_i^2 - \frac{V_i V_j}{n_{ij}} e^{j(\theta_i - \theta_j + \phi_{ij})} \right) \end{aligned}$$

et une décomposition en parties réelle et imaginaire fournit :

$$P_i = \underbrace{\sum_{j \in \mathcal{N}(i)} G_{ij}V_i^2 - \sum_{j \in \mathcal{N}(i)} \frac{V_i V_j}{n_{ij}} [G_{ij} \cos(\theta_i - \theta_j + \phi_{ij}) + B_{ij} \sin(\theta_i - \theta_j + \phi_{ij})]}_{f_i(\dots, V_i, \theta_i, \dots)} \quad (7.3)$$

$$\begin{aligned} Q_i &= -(B_{si} + \sum_{j \in \mathcal{N}(i)} B_{sij} + \sum_{j \in \mathcal{N}(i)} B_{ij})V_i^2 + \\ &+ \underbrace{\sum_{j \in \mathcal{N}(i)} \frac{V_i V_j}{n_{ij}} [B_{ij} \cos(\theta_i - \theta_j + \phi_{ij}) - G_{ij} \sin(\theta_i - \theta_j + \phi_{ij})]}_{g_i(\dots, V_i, \theta_i, \dots)} \end{aligned} \quad (7.4)$$

7.2 Spécification des données du load flow

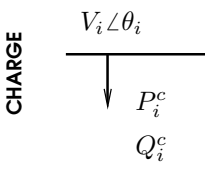
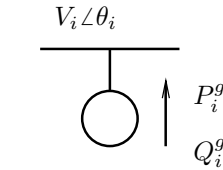
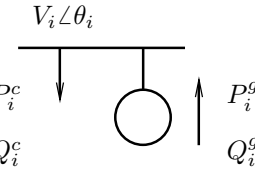
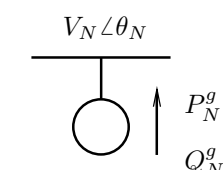
7.2.1 Données nodales

Le réseau est décrit par les $2N$ équations (7.3, 7.4). En chaque noeud du réseau, ces équations font intervenir quatre grandeurs : le module V_i et la phase θ_i de la tension, les puissances active P_i et réactive Q_i . Pour qu'inconnues et équations soient en nombre égal, il faut donc spécifier deux de ces quatre grandeurs en chaque noeud.

La figure 7.3 détaille les différentes données nodales que l'on spécifie en pratique ainsi que les équations et les inconnues correspondantes.

En un jeu de barres où est connectée une charge, on spécifie les puissances active et réactive consommées par celle-ci, car ces informations sont généralement disponibles au départ de mesures. Les équations relatives à un tel noeud sont données par (7.3, 7.4) où P_i, Q_i sont les consommations changées de signe. En un tel noeud, les inconnues sont donc V_i et θ_i .

FIGURE 7.3 – load flow : données, équations et inconnues nodales

EQUIPEMENT CONNECTE	EQUATIONS	INCONNUES	en remplaçant modules et phases dans l'équ. de bilan de puissance non utilisée, on en déduit :
CHARGE 	$-P_i^c = f_i(\dots)$ $-Q_i^c = g_i(\dots)$ NOEUD "PQ"	V_i θ_i	
GENERATEUR 	$P_i^g = f_i(\dots)$ $Q_i^g = g_i(\dots)$ NOEUD "PQ"	V_i θ_i	
	$P_i^g = f_i(\dots)$ $V_i = V_i^o$ NOEUD "PV"	θ_i	Q_i^g
GENERATEUR + CHARGE 	$P_i^g - P_i^c = f_i(\dots)$ $Q_i^g - Q_i^c = g_i(\dots)$ NOEUD "PQ"	V_i θ_i	
	$P_i^g - P_i^c = f_i(\dots)$ $V_i = V_i^o$ NOEUD "PV"	θ_i	Q_i^g
"BALANCIER" 	$\theta_N = 0$ $V_N = V_N^o$		P_N^g Q_N^g

Les mêmes informations sont généralement spécifiées pour les générateurs de faible puissance. La production active P est, aux pertes près, la puissance générée par la turbine, tandis qu'un asservissement maintient le facteur de puissance $P/\sqrt{P^2 + Q^2}$ à une valeur spécifiée, ce qui fournit la puissance réactive générée.

Ces noeuds où l'on spécifie P et Q sont souvent désignés sous le vocable de "noeuds PQ".

Comme nous le verrons au chapitre 11, les générateurs des grandes centrales sont dotés de régulateurs de tension qui maintiennent (quasiment) constantes les tensions à leurs bornes. En un tel jeu de barres, il est plus naturel de spécifier la tension que la puissance réactive. Les données sont donc P_i et V_i . Le module de la tension étant directement spécifié, il ne reste que θ_i comme inconnue. L'équation (7.4) n'est donc pas utilisée pour calculer les tensions aux noeuds du système. Cependant, une fois ces tensions connues, cette équation est utilisée pour calculer la puissance réactive produite par le générateur.

Ces noeuds où l'on spécifie P et V sont désignés sous le vocable de "noeuds PV".

Certains jeux de barres peuvent recevoir une charge et un générateur¹. Dans ce cas, ce sont les données relatives au générateur qui dictent le type du noeud : PQ ou PV selon le cas. L'injection de puissance active P_i (resp. réactive Q_i) est évidemment la différence entre la puissance générée et la puissance consommée.

7.2.2 Le générateur balancier

A ce stade, deux remarques s'imposent :

1. on ne peut spécifier les puissances P_i à tous les noeuds. En effet, le bilan de puissance active du réseau s'écrit :

$$\sum_{i=1}^N P_i = p \quad (7.5)$$

où p représente les pertes actives totales dans le réseau. Spécifier toutes les valeurs P_i reviendrait donc à spécifier les pertes. Or, ces dernières sont fonction des courants dans les branches et donc des tensions aux noeuds, lesquelles ne sont pas connues à ce stade ;

2. seules des différences angulaires interviennent dans les équations (7.3, 7.4) ; on peut ajouter une même constante à toutes les phases sans changer l'état électrique du réseau. Il convient en fait de calculer les déphasages de $N - 1$ noeuds par rapport à l'un d'entre eux pris comme référence.

Pour traiter ce double problème, un des jeux de barres du réseau est désigné comme référence angulaire et se voit imposer la phase de sa tension, plutôt que la puissance active. Il est d'usage de spécifier une phase nulle, mais c'est arbitraire. En ce noeud, l'équation (7.3) n'est pas utilisée.

Ce jeu de barres est qualifié de *balancier*². Nous supposons dans ce qui suit qu'il s'agit du N -ème noeud.

Au balancier, la relation (7.5) donne :

$$P_N = - \sum_{i=1}^{N-1} P_i + p \quad (7.6)$$

1. c'est le cas quand les auxiliaires d'une centrale sont alimentés via le jeu de barres où est connecté le générateur

2. en anglais : *slack bus*

où les différents termes de la somme sont spécifiés dans les données, tandis que, comme indiqué plus haut, p n'est connu qu'à l'issue du calcul. Pour disposer d'une flexibilité au niveau de l'injection de puissance active au noeud balancier, il est d'usage de choisir un jeu de barres où est connecté un générateur.

Qu'en est-il des pertes réactives et du rôle du balancier ? En fait, en chaque noeud PV, la puissance réactive n'est pas spécifiée ; elle est le résultat du calcul. Le problème de la spécifier indûment à tous les noeuds ne se pose donc pas. En quelque sorte, chaque noeud PV joue (localement) le rôle de balancier pour la puissance réactive. Evidemment, le problème réapparaît si l'on spécifie la puissance réactive de tous les générateurs³. Si l'on met ce cas de côté, il est permis de spécifier soit la tension, soit la puissance réactive au noeud balancier. Dans la mesure où une centrale est connectée à ce noeud, il sera plus naturel (mais pas obligatoire) d'y spécifier la tension V , pour les raisons déjà mentionnées. C'est le choix qui a été considéré à la figure 7.3. Dans ce cas, il n'y a aucune inconnue à calculer au noeud balancier. Les équations (7.3) et (7.4) y sont utilisées pour calculer les puissances active et réactive injectées.

7.3 Résolution numérique des équations de load flow

7.3.1 Equations de load flow sous forme vectorielle

Soient N_{PV} le nombre de noeuds PV et N_{PQ} le nombre de noeuds PQ, avec :

$$N_{PV} + N_{PQ} + 1 = N$$

Soient :

$\mathbf{v}$ le vecteur des modules des tensions aux noeuds (dimension N)

$\boldsymbol{\theta}$ le vecteur des phases des tensions aux noeuds (dimension N)

$\mathbf{p}^o$ le vecteur des injections actives spécifiées aux noeuds PV et PQ (dimension $N - 1$)

$\mathbf{q}^o$ le vecteur des injections réactives spécifiées aux noeuds PQ (dimension N_{PQ})

$\mathbf{v}^o$ le vecteur des modules des tensions spécifiés aux noeuds PV (dimension N_{PV}).

Les équations de load flow peuvent s'écrire sous forme vectorielle de la façon suivante :

- $N - 1$ équations de puissance active (7.3) aux noeuds PV et PQ :

$$\mathbf{f}(\mathbf{v}, \boldsymbol{\theta}) - \mathbf{p}^o = \mathbf{0} \quad (7.7)$$

où les composantes de $\mathbf{f}$ sont les fonctions f_i définies par (7.3) ;

- une équation de phase au noeud balancier :

$$\theta_N = 0 \quad (7.8)$$

3. une telle situation ne se présente jamais en pratique dans les calculs de réseaux de transport mais elle peut se présenter dans un réseau de distribution hébergeant des sources d'énergie renouvelables dont aucune ne participe au réglage de la tension, pratique très répandue à l'heure actuelle

- N_{PQ} équations de puissance réactive (7.4) aux noeuds PQ :

$$\mathbf{g}(\mathbf{v}, \boldsymbol{\theta}) - \mathbf{q}^o = \mathbf{0} \quad (7.9)$$

où les composantes de $\mathbf{g}$ sont les fonctions g_i définies par (7.4) ;

- $N_{PV} + 1$ équations de tension aux noeuds PV et au balancier :

$$\mathbf{v} - \mathbf{v}^o = \mathbf{0} \quad (7.10)$$

Ces $2N$ équations font intervenir $2N$ variables (ouf !).

Sauf cas particulier très simple (voir p.ex. section 3.4), les équations (7.7, 7.9) ne peuvent pas être résolues analytiquement. Une résolution numérique s'impose. La méthode la plus répandue est celle de Newton (ou Newton-Raphson), dont nous rappelons le principe avant de l'appliquer au cas qui nous occupe.

7.3.2 Méthode de Newton : rappel

Cas d'une fonction scalaire à une variable

Soit à résoudre :

$$\varphi(x) = 0 \quad \text{avec } \varphi : R \rightarrow R$$

Nous notons $\varphi_x = \frac{d\varphi}{dx}$ la dérivée de φ par rapport à x .

Ayant choisi une valeur initiale $x^{(0)}$, la méthode de Newton consiste à calculer la suite de points :

$$x^{(k+1)} = x^{(k)} - \frac{\varphi(x^{(k)})}{\varphi_x(x^{(k)})} \quad k = 0, 1, 2, \dots$$

jusqu'à ce que :

$$|\varphi(x^{(k+1)})| < \epsilon$$

où ϵ est une tolérance.

Ce procédé itératif est représenté graphiquement à la figure 7.4.

Pour autant que $x^{(0)}$ soit "suffisamment proche" de la solution, cette méthode offre une convergence rapide (quadratique).

Cas d'une fonction vectorielle à plusieurs variables

Soit à résoudre :

$$\boldsymbol{\varphi}(\mathbf{x}) = \mathbf{0} \quad \text{avec } \boldsymbol{\varphi} : R^n \rightarrow R^n$$

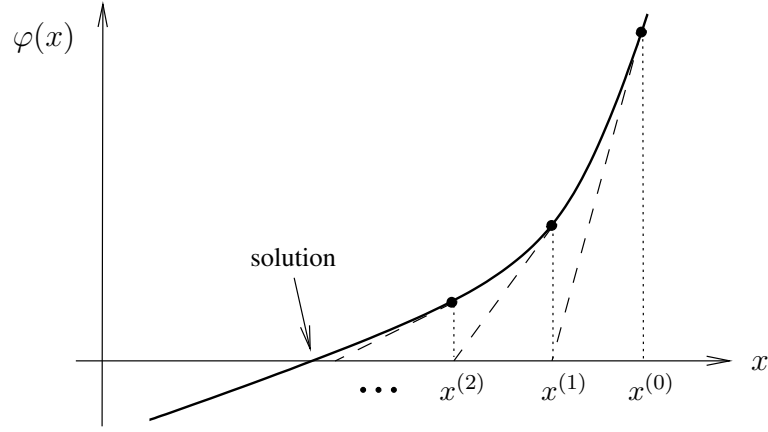


FIGURE 7.4 – illustration graphique de la méthode de Newton

Nous notons φ_x la matrice *jacobienne* de φ par rapport à x , c'est-à-dire la matrice des dérivées partielles telle que :

$$[\varphi_x]_{ij} = \frac{\partial \varphi_i}{\partial x_j} \quad i, j = 1, \dots, n$$

Ayant choisi une valeur initiale $x^{(0)}$, on calcule la suite de points :

$$x^{(k+1)} = x^{(k)} - [\varphi_x(x^{(k)})]^{-1} \varphi(x^{(k)}) \quad k = 0, 1, 2, \dots \quad (7.11)$$

jusqu'à ce que :

$$\max_i |\varphi_i(x^{(k+1)})| < \epsilon$$

où ϵ est une tolérance.

En pratique, plutôt que d'inverser la matrice jacobienne, on résout le système linéaire :

$$\varphi_x(x^{(k)}) \Delta x = -\varphi(x^{(k)}) \quad (7.12)$$

et on incrémente x selon :

$$x^{(k+1)} = x^{(k)} + \Delta x \quad (7.13)$$

La résolution du système linéaire se fait en deux étapes :

1. *factorisation* (aussi appelée *LDU-décomposition*) : elle consiste à décomposer la matrice φ_x en :

$$\varphi_x = L D U \quad (7.14)$$

où L est une matrice triangulaire inférieure à diagonale unitaire, U est une matrice triangulaire supérieure à diagonale unitaire et D une matrice diagonale, comme représenté à la figure 7.5.

2. *Substitution* du membre de droite. En introduisant (7.14) dans (7.12), le système à résoudre devient :

$$L D U \Delta x = -\varphi$$

Ce dernier est résolu à son tour en deux étapes :

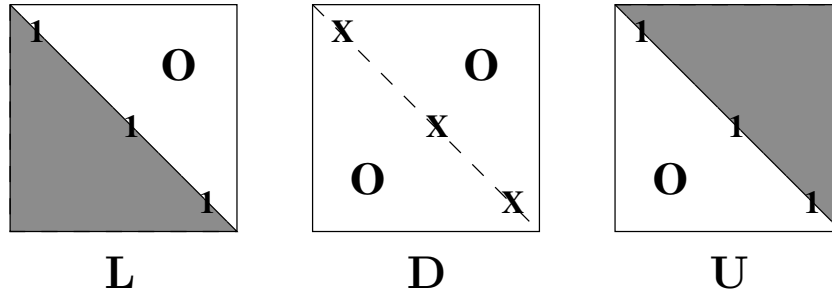


FIGURE 7.5 – structure des matrices L, D et U

(a) On résoud d'abord :

$$\mathbf{L} \mathbf{y} = -\varphi$$

par rapport à $\mathbf{y}$,

(b) puis on résoud :

$$\mathbf{U} \Delta \mathbf{x} = \mathbf{D}^{-1} \mathbf{y}$$

par rapport à $\Delta \mathbf{x}$.

Ces deux opérations sont simples étant donné le caractère triangulaire des matrices impliquées.

7.3.3 Application aux équations de load flow

La matrice jacobienne des équations (7.7 - 7.10) par rapport à $\mathbf{v}$ et $\boldsymbol{\theta}$ se décompose comme suit :

$$\varphi_{\mathbf{x}} = \left[\begin{array}{c|c} \mathbf{f}_{\mathbf{v}} & \mathbf{f}_{\boldsymbol{\theta}} \\ \hline \mathbf{g}_{\mathbf{v}} & \mathbf{g}_{\boldsymbol{\theta}} \\ \mathbf{U} & \mathbf{0} \end{array} \right] \quad (7.15)$$

où $\mathbf{e}_N$ est un vecteur ligne unitaire de dimension N et $\mathbf{U}$ une matrice comportant des 0 et des 1. Chacune des quatre sous-matrices dans (7.15) est de dimension $N \times N$.

La caractéristique principale de cette matrice est d'être très *creuse*, c'est-à-dire de comporter une très grande proportion d'éléments nuls. En ce qui concerne $\mathbf{f}_{\mathbf{v}}$, $\mathbf{f}_{\boldsymbol{\theta}}$, $\mathbf{g}_{\mathbf{v}}$ et $\mathbf{g}_{\boldsymbol{\theta}}$, cette propriété vient du fait que chaque équation de load flow (7.3, 7.4) ne fait intervenir que la tension du noeud auquel elle se rapporte et celles des noeuds voisins. Plus le réseau traité est grand, plus la proportion d'éléments nuls augmente, le nombre moyen de voisins d'un noeud restant constant.

Les matrices creuses sont manipulées en faisant appel à des algorithmes spéciaux⁴ dont le principe peut se résumer comme suit :

- seuls les éléments différents de zéro sont stockés. Un système de pointeurs permet de parcourir les éléments non nuls présents dans une ligne ou une colonne de la matrice d'origine ;

4. en anglais : *sparsity programming*

- en ne manipulant que ces éléments, on évite toutes les opérations mathématiques inutiles impliquant des zéros ;
- lors de la factorisation de la matrice, on permute ses lignes et/ou ses colonnes de manière à ce que le nombre de nouveaux éléments non nuls créés par cette opération reste le plus faible possible⁵. L'ordre dans lequel les lignes et/ou colonnes vont être traitées⁶ est décidé par analyse des emplacements des éléments non nuls. Cette analyse est souvent effectuée avant de procéder aux calculs proprement dits. Si nécessaire, les permutations sont combinées avec les opérations de pivotage, destinées à préserver la précision.

L'algorithme de Newton (7.12) comporte en principe la mise à jour de la matrice jacobienne à chaque itération. On peut cependant gagner du temps de calcul en conservant cette matrice constante après quelques itérations. En effet, pour autant que l'algorithme converge, la solution ne dépend pas de la matrice choisie dans le membre de gauche de (7.12). En pratique, pour éviter la divergence, la matrice ne sera bloquée que lorsque l'amplitude des fonctions φ_i sera suffisamment faible, ce qui est l'indice que l'on est proche de la solution. Quand la matrice n'est plus mise à jour, ses facteurs **L**, **D** et **U** ne le sont plus non plus et l'on ne procède plus qu'à l'étape de substitution.

Si l'on ne dispose pas d'une estimation plus précise, la séquence d'itérations est démarrée en initialisant :

- aux noeuds PV : le module de la tension à la valeur spécifiée dans les données
- aux autres noeuds : le module de la tension à 1 pu
- à tous les noeuds : la phase de la tension à 0° (la phase imposée au balancier).

7.4 Extensions de la formulation

7.4.1 Prise en compte de contraintes de fonctionnement - application aux limites réactives des générateurs

Un système électrique de puissance doit satisfaire un certain nombre de contraintes de fonctionnement, qui expriment que des grandeurs physiques ne doivent pas dépasser certaines limites. Mathématiquement, une contrainte se présente sous la forme d'une inégalité :

$$h(\mathbf{v}, \boldsymbol{\theta}) \leq 0 \quad (7.16)$$

où h est une fonction des modules et des phases des tensions. En pratique, h ne fait intervenir qu'un petit nombre de ces variables.

Si la solution d'un premier calcul de load flow ne satisfait pas une contrainte de fonctionnement, il est possible dans certains cas d'imposer la contrainte en question, sous forme de l'égalité :

$$h(\mathbf{v}, \boldsymbol{\theta}) = 0$$

5. rappelons qu'on ne modifie pas la solution d'un système linéaire si l'on permute l'ordre des lignes ou des colonnes, à condition évidemment de répercuter ces permutations sur le vecteur des inconnues et sur le terme indépendant

6. en anglais : *optimal ordering*

à condition de pouvoir retirer une autre équation, de manière à maintenir le nombre d'équations constant. On résout alors le nouvel ensemble d'équations ; on recommence si nécessaire.

Les contraintes les plus importantes pouvant être traitées de cette façon sont les productions réactives des générateurs. Supposons qu'un générateur soit connecté au i -ème noeud et que celui-ci soit du type PV. La production réactive Q_i du générateur doit rester à l'intérieur de limites dictées par l'échauffement, voire la stabilité de son fonctionnement, comme détaillé au chapitre 8. Ces limites s'expriment par :

$$Q_i^{min} \leq Q_i(\mathbf{v}, \boldsymbol{\theta}) \leq Q_i^{max}$$

ou encore :

$$\begin{aligned} Q_i^{min} - Q_i(\mathbf{v}, \boldsymbol{\theta}) &\leq 0 \\ Q_i(\mathbf{v}, \boldsymbol{\theta}) - Q_i^{max} &\leq 0 \end{aligned}$$

où Q_i est une fonction des tensions au i -ème noeud et à tous ses voisins, comme le montre la relation (7.4).

Si à l'issue d'un premier calcul, on trouve que la production Q_i dépasse la borne supérieure Q_i^{max} , on impose :

$$Q_i(\mathbf{v}, \boldsymbol{\theta}) - Q_i^{max} = 0 \quad (7.17)$$

ce qui fait passer le noeud du type PV au type PQ. Il y a donc une équation (7.10) en moins et une équation (7.9) en plus. En fonctionnement normal, V_i prend une valeur inférieure à V_i^o .

On procède de la même manière si la production vient à passer sous la borne inférieure Q_i^{min} . Dans ce cas, V_i prend une valeur supérieure à V_i^o .

L'algorithme ci-après résume les différentes étapes de la méthode de Newton avec passage en limite des générateurs.

Méthode de Newton avec traitement des limites réactives des générateurs

1. $k := 0$
2. donner une valeur initiale à $\mathbf{v}^{(0)}$ et $\boldsymbol{\theta}^{(0)}$
3. calculer $\mathbf{f}(\mathbf{v}^{(k)}, \boldsymbol{\theta}^{(k)}) - \mathbf{p}^o$ et $\mathbf{g}(\mathbf{v}^{(k)}, \boldsymbol{\theta}^{(k)}) - \mathbf{q}^o$
4. SI $\max_i |g_i(\mathbf{v}^{(k)}, \boldsymbol{\theta}^{(k)}) - Q_i^o| < \delta_Q$:
 tester les limites réactives des générateurs
 SI dépassements : basculer les noeuds correspondants de PV en PQ
 aller en 3
5. SI $\max_i |f_i(\mathbf{v}^{(k)}, \boldsymbol{\theta}^{(k)}) - P_i^o| < \epsilon_P$ et $\max_i |g_i(\mathbf{v}^{(k)}, \boldsymbol{\theta}^{(k)}) - Q_i^o| < \epsilon_Q$: STOP
6. SI $\max_i |f_i(\mathbf{v}^{(k)}, \boldsymbol{\theta}^{(k)}) - P_i^o| > \beta_P$ ou $\max_i |g_i(\mathbf{v}^{(k)}, \boldsymbol{\theta}^{(k)}) - Q_i^o| > \beta_Q$:
 calculer et factoriser la matrice jacobienne : $\boldsymbol{\varphi}_x = \mathbf{LDU}$
7. résoudre $\mathbf{LDU} \begin{bmatrix} \Delta \mathbf{v} \\ \Delta \boldsymbol{\theta} \end{bmatrix} = \begin{bmatrix} \mathbf{p}^o - \mathbf{f}(\mathbf{v}^{(k)}, \boldsymbol{\theta}^{(k)}) \\ \mathbf{q}^o - \mathbf{g}(\mathbf{v}^{(k)}, \boldsymbol{\theta}^{(k)}) \end{bmatrix}$

8. incrémenter les inconnues : $\mathbf{v}^{(k+1)} := \mathbf{v}^{(k)} + \Delta \mathbf{v}$ $\boldsymbol{\theta}^{(k+1)} := \boldsymbol{\theta}^{(k)} + \Delta \boldsymbol{\theta}$
9. $k := k + 1$
10. aller en 3.

Dans cet algorithme :

δ_Q est le seuil de puissance réactive en dessous duquel on considère que les productions réactives des générateurs sont connues avec une précision suffisante pour tester les limites de celles-ci
 ϵ_P (resp. ϵ_Q) est la tolérance en dessous de laquelle on considère que les équations de puissance active (resp. réactive) sont résolues (p.ex. 0.1 MW, 0.1 Mvar)
 β_P (resp. β_Q) est un seuil de puissance active (resp. réactive) en dessous duquel on ne rafraîchit plus la matrice jacobienne.

7.4.2 Traitement itératif des pertes

Revenons sur le rôle du balancier vis-à-vis des pertes actives.

Avant d'effectuer un calcul de load flow, on spécifie les différentes charges, on estime les pertes actives et l'on répartit la somme des deux sur les différents générateurs, en ce compris le balancier. Soit $\hat{p}$ la valeur estimée des pertes. La production estimée du balancier est donnée par l'équation (7.6) :

$$\hat{P}_N = - \sum_{i=1}^{N-1} P_i + \hat{p} \quad (7.18)$$

Une fois le calcul effectué, on connaît les pertes p et la production P_N du balancier, liées par :

$$P_N = - \sum_{i=1}^{N-1} P_i + p \quad (7.19)$$

En soustrayant ces deux dernières équations l'une de l'autre, on obtient :

$$P_N - \hat{P}_N = p - \hat{p}$$

qui montre, comme on pouvait s'y attendre, que si l'estimée $\hat{p}$ était trop éloignée de la valeur trouvée à l'issue du calcul, le balancier produit une puissance significativement différente de ce que l'on escomptait.

Le balancier étant un générateur parmi d'autres, cette production inattendue peut être jugée inacceptable. Si c'est le cas, on corrige le schéma de production en répartissant la somme des charges et de la nouvelle valeur p des pertes sur les différents générateurs, en ce compris le balancier. On procède à un nouveau calcul de load flow, qui fournit une nouvelle valeur des pertes. On peut itérer de la sorte jusqu'à ce que les pertes ne varient plus significativement d'une itération à la suivante.

Dans de nombreux cas pratiques, une seule itération, c'est-à-dire un seul calcul de load flow supplémentaire, suffit pour corriger la mauvaise estimation initiale des pertes.

En pratique, un gestionnaire de réseau connaît assez précisément les pertes dans son système ; la première estimation $\hat{p}$ sera assez précise pour ne pas devoir effectuer ce type de correction.

7.4.3 Balanciers distribués

Tel que formulé jusqu'ici, le calcul de load flow confie à un seul générateur, le balancier, le rôle de “boucler le bilan de puissance” active du réseau. Il est aussi possible de répartir cette tâche sur plusieurs générateurs, en ayant recours à un ensemble de “balanciers distribués”.

On définit une variable ΔP_g qui sert à ajuster les productions des divers balanciers. Au i -ème noeud ($i = 1, \dots, N$) on affecte une fraction α_i de ΔP_g selon la formule :

$$P_i = P_i^o + \alpha_i \Delta P_g \quad \text{avec} \quad \sum_{i=1}^N \alpha_i = 1$$

Aux noeuds sans générateurs ou avec générateur ne participant pas au réglage, on pose simplement $\alpha_i = 0$.

Dans cette formulation, on distingue seulement deux types de noeuds : PV et PQ. Un des jeux de barres (nous continuerons à supposer qu'il s'agit du N -ème noeud) sert encore de référence angulaire en posant $\theta_N = 0$.

Les équations de load flow s'écrivent à présent :

- N équations de puissance active (7.3) aux noeuds du réseau (tous les noeuds, y compris les balanciers) :

$$\mathbf{f}(\mathbf{v}, \boldsymbol{\theta}) - \mathbf{p}^o - \Delta P_g \boldsymbol{\alpha} = \mathbf{0} \quad (7.20)$$

où le vecteur $\boldsymbol{\alpha}$ regroupe les différents coefficients α_i ;

- une équation de phase au noeud de référence :

$$\theta_N = 0 \quad (7.21)$$

- N_{PQ} équations de puissance réactive (7.4) aux noeuds PQ :

$$\mathbf{g}(\mathbf{v}, \boldsymbol{\theta}) - \mathbf{q}^o = \mathbf{0} \quad (7.22)$$

- N_{PV} équations de tension aux noeuds PV :

$$\mathbf{v} - \mathbf{v}^o = \mathbf{0} \quad (7.23)$$

Notons que dans (7.20) il y a une équation en plus que dans (7.7) (N au lieu de $N - 1$), ce qui compense l'inconnue ΔP_g supplémentaire. Les équations (7.21, 7.22, 7.23) sont identiques aux équations (7.8, 7.9, 7.10). ΔP_g est calculée par la méthode de Newton en même temps que $\mathbf{v}$ et $\boldsymbol{\theta}$.

7.4.4 Ajustement des rapports de transformation

Les logiciels de calcul de load flow perfectionnés peuvent comporter plusieurs ajustements destinés à reproduire des réglages automatiques ou manuels effectués sur le réseau.

Un de ces ajustements est le passage en limite des générateurs, exposé à la section 7.4.1, dans lequel certains noeuds passent du type PV au type PQ avant de procéder à une nouvelle résolution des équations.

Un autre ajustement très répandu est celui des rapports de transformateurs, destiné à maintenir les tensions de certains noeuds dans des plages de valeurs spécifiées. Ceci peut traduire une action effectuée par l'opérateur d'un centre de conduite ou par un régleur en charge automatique, comme évoqué à la section 6.5.

Dans le calcul, les rapports de transformation peuvent être traités comme des variables continues ou comme des variables discrètes dont les valeurs correspondent aux diverses prises du régleur en charge. Les variations des rapports peuvent être constantes ou proportionnelles à l'écart entre la tension observée et la tension de consigne. Lors des ajustements, le rapport de transformation ne peut pas dépasser les limites inférieure et supérieure correspondant à la première et à la dernière prise du régleur en charge.

Après ajustement, les équations de load flow sont à nouveau résolues, ce qui fournit de nouvelles tensions à contrôler. On itère ainsi jusqu'à ce que plus aucun transformateur ne subisse de variation.

7.5 Découplage électrique

A la section 3.1.2, nous avons mis en évidence le découplage électrique qui existe dans les réseaux de transport entre les puissances actives et les phases des tensions, d'une part, les puissances réactives et les modules des tensions, d'autre part. Cette propriété se marque au niveau de la matrice jacobienne par le fait que, parmi les sous-matrices $\mathbf{f}_\theta$, $\mathbf{g}_v$, $\mathbf{f}_v$ et $\mathbf{g}_\theta$, les deux premières sont dominantes, comme le montrent les calculs ci-après.

A partir de (7.3), on calcule aisément les dérivées partielles suivantes :

$$\begin{aligned}\frac{\partial P_i}{\partial V_i} &= 2V_i \sum_{j \in \mathcal{N}(i)} G_{ij} - \sum_{j \in \mathcal{N}(i)} \frac{V_j}{n_{ij}} [G_{ij} \cos(\theta_i - \theta_j + \phi_{ij}) + B_{ij} \sin(\theta_i - \theta_j + \phi_{ij})] \\ \frac{\partial P_i}{\partial V_j} &= -\frac{V_i}{n_{ij}} [G_{ij} \cos(\theta_i - \theta_j + \phi_{ij}) + B_{ij} \sin(\theta_i - \theta_j + \phi_{ij})] \\ \frac{\partial P_i}{\partial \theta_i} &= \sum_{j \in \mathcal{N}(i)} \frac{V_i V_j}{n_{ij}} [G_{ij} \sin(\theta_i - \theta_j + \phi_{ij}) - B_{ij} \cos(\theta_i - \theta_j + \phi_{ij})] \\ \frac{\partial P_i}{\partial \theta_j} &= -\frac{V_i V_j}{n_{ij}} [G_{ij} \sin(\theta_i - \theta_j + \phi_{ij}) - B_{ij} \cos(\theta_i - \theta_j + \phi_{ij})]\end{aligned}$$

Si l'on suppose que les modules des tensions et les rapports des transformateurs sont proches de l'unité :

$$V_i = V_j = n_{ij} \simeq 1 \text{ pu}$$

et que le déphasage angulaire le long de chaque branche est faible :

$$\theta_i - \theta_j + \phi_{ij} \simeq 0$$

les dérivées partielles ci-dessus deviennent :

$$\frac{\partial P_i}{\partial V_i} \simeq \sum_{j \in \mathcal{N}(i)} G_{ij} \quad \frac{\partial P_i}{\partial V_j} \simeq -G_{ij} \quad \frac{\partial P_i}{\partial \theta_i} \simeq - \sum_{j \in \mathcal{N}(i)} B_{ij} \quad \frac{\partial P_i}{\partial \theta_j} \simeq B_{ij}$$

Etant donné que $G_{ij} \ll |B_{ij}|$ dans les réseaux de transport, on en déduit que :

$$|\frac{\partial P_i}{\partial V_i}|, |\frac{\partial P_i}{\partial V_j}| \ll |\frac{\partial P_i}{\partial \theta_i}|, |\frac{\partial P_i}{\partial \theta_j}|$$

On a de même pour la puissance réactive :

$$\begin{aligned} \frac{\partial Q_i}{\partial V_i} &= -2[B_{si} + \sum_j (B_{ij} + B_{sij})]V_i + \sum_j \frac{V_j}{n_{ij}} [B_{ij} \cos(\theta_i - \theta_j + \phi_{ij}) - G_{ij} \sin(\theta_i - \theta_j + \phi_{ij})] \\ \frac{\partial Q_i}{\partial V_j} &= \frac{V_i}{n_{ij}} [B_{ij} \cos(\theta_i - \theta_j + \phi_{ij}) - G_{ij} \sin(\theta_i - \theta_j + \phi_{ij})] \\ \frac{\partial Q_i}{\partial \theta_i} &= - \sum_{j \in \mathcal{N}(i)} \frac{V_i V_j}{n_{ij}} [B_{ij} \sin(\theta_i - \theta_j + \phi_{ij}) + G_{ij} \cos(\theta_i - \theta_j + \phi_{ij})] \\ \frac{\partial Q_i}{\partial \theta_j} &= \frac{V_i V_j}{n_{ij}} [B_{ij} \sin(\theta_i - \theta_j + \phi_{ij}) + G_{ij} \cos(\theta_i - \theta_j + \phi_{ij})] \end{aligned}$$

En appliquant les mêmes simplifications que ci-dessus :

$$\frac{\partial Q_i}{\partial V_i} \simeq -2[B_{si} + \sum_{j \in \mathcal{N}(i)} (B_{ij} + B_{sij})] \quad \frac{\partial Q_i}{\partial V_j} \simeq B_{ij} \quad \frac{\partial Q_i}{\partial \theta_i} = - \sum_{j \in \mathcal{N}(i)} G_{ij} \quad \frac{\partial Q_i}{\partial \theta_j} \simeq G_{ij}$$

et en considérant à nouveau que $G_{ij} \ll |B_{ij}|$, on trouve :

$$|\frac{\partial Q_i}{\partial V_i}|, |\frac{\partial Q_i}{\partial V_j}| \gg |\frac{\partial Q_i}{\partial \theta_i}|, |\frac{\partial Q_i}{\partial \theta_j}|$$

Ces inégalités confirment bien les propriétés de découplage électrique rappelées plus haut.

La dominance des sous-matrices $\mathbf{f}_\theta$ et $\mathbf{g}_v$ est à la base de la variante *découplée rapide* de la méthode de Newton, qui consiste à incrémenter en alternance les phases et les modules en utilisant les équations actives (7.3) et la sous-matrice jacobienne $\begin{bmatrix} \mathbf{f}_\theta \\ \mathbf{e}_N \end{bmatrix}$ pour les phases, les équations réactives (7.4) et la sous-matrice jacobienne $\begin{bmatrix} \mathbf{g}_v \\ \mathbf{U} \end{bmatrix}$ pour les modules.

7.6 L'approximation du courant continu (ou DC load flow)

L'approximation *du courant continu* est un modèle simplifié de l'écoulement de puissance dans un réseau, qui consiste à :

- approximer la relation entre puissances actives et phases des tensions par une fonction linéaire
- négliger les pertes actives dans toutes les branches
- supposer que les modules des tensions sont tous égaux à 1 pu
- négliger les transits de puissance réactive.

Ce modèle simplifié est utilisé pour alléger des calculs lourds, tels que :

- de très nombreux calculs de répartition de charge (la méthode ne requiert que la résolution d'un seul système linéaire, dont la taille est la moitié de celle d'un calcul complet)
- des problèmes d'optimisation incluant les équations de load flow comme contraintes ("optimal power flow").

Etant donné son caractère linéaire, le modèle convient aussi pour calculer par superposition les effets de plusieurs modifications appliquées au réseau.

En général, l'erreur commise sur les transits de puissance active est de l'ordre de quelques pourcents de la valeur exacte.

Considérons le quadripole universel de la figure 7.1. On suppose donc : $V_i = V_j \simeq 1$ pu. Cette approche s'appliquant exclusivement aux réseaux de transport, on suppose comme précédemment que les conductances sont négligeables :

$$G_{ij} \simeq 0$$

Il en résulte que :

$$B_{ij} = -\frac{1}{X_{ij}}$$

où X_{ij} est la réactance série de la branche ij . Enfin, on néglige l'effet des transformateurs en prenant :

$$n_{ij} \simeq 1$$

En introduisant ces simplifications dans l'expression (7.3) de la puissance active injectée au i -ème noeud, on trouve aisément :

$$P_i \simeq \sum_{j \in \mathcal{N}(i)} \frac{1}{X_{ij}} \sin(\theta_i - \theta_j + \phi_{ij})$$

et, en supposant enfin que les déphasages angulaires sont faibles, on obtient l'expression linéarisée :

$$P_i \simeq \sum_{j \in \mathcal{N}(i)} \frac{\theta_i - \theta_j + \phi_{ij}}{X_{ij}} \quad (7.24)$$

Les conductances G_{ij} étant négligées, les pertes actives le sont aussi et le bilan de puissance active du système s'écrit :

$$\sum_{i=1}^N P_i = 0 \quad \Leftrightarrow \quad P_N = - \sum_{i=1}^{N-1} P_i$$

Les modules des tensions étant supposés égaux à 1 et les transits de puissance réactive étant négligés, le module du courant dans la branche ij vaut *en per unit* :

$$I_{ij} = |P_{ij}| = \left| \frac{\theta_i - \theta_j + \phi_{ij}}{X_{ij}} \right|$$

Supposons pour simplifier qu'il n'y a pas de transformateurs déphaseurs dans le système ($\phi_{ij} = 0$). En regroupant les relations (7.24) sous forme matricielle, les équations de load flow sous l'approximation du courant continu s'écrivent :

$$\mathbf{p}^o = \mathbf{A} \boldsymbol{\theta} \quad (7.25)$$

où $\mathbf{p}^o$ et $\boldsymbol{\theta}$ sont les vecteurs de dimension $N - 1$ définis précédemment et la matrice $\mathbf{A}$ est définie par :

$$\begin{aligned} [\mathbf{A}]_{ij} &= -\frac{1}{X_{ij}} \quad i, j = 1, \dots, N-1; \quad i \neq j \\ [\mathbf{A}]_{ii} &= \sum_{j \in \mathcal{N}(i)} \frac{1}{X_{ij}} \quad i = 1, \dots, N-1 \end{aligned}$$

Exemple

Considérons le réseau à 4 noeuds de la figure 7.6.a. Le noeud 4 est pris comme balancier, avec $P_4 = -P_1 + P_2 + P_3$ et $\theta_4 = 0$.

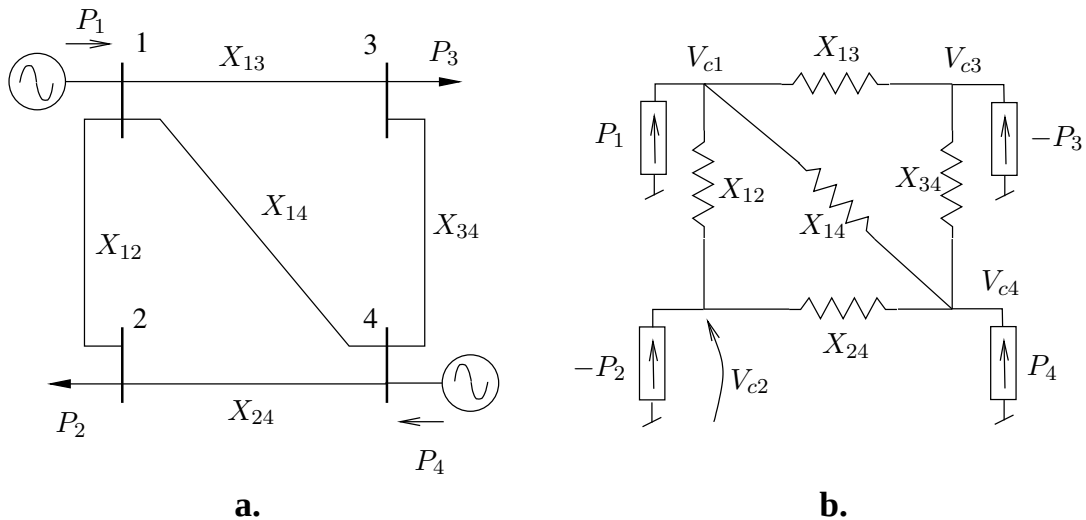


FIGURE 7.6 – exemple à 4 noeuds

Avec les réactances et les puissances actives définies à la figure 7.6.a, les équations de load flow sous l'approximation du courant continu s'écrivent :

$$\begin{bmatrix} P_1 \\ -P_2 \\ -P_3 \end{bmatrix} = \begin{bmatrix} \frac{1}{X_{12}} + \frac{1}{X_{13}} + \frac{1}{X_{14}} & -\frac{1}{X_{12}} & -\frac{1}{X_{13}} \\ -\frac{1}{X_{12}} & \frac{1}{X_{12}} + \frac{1}{X_{24}} & 0 \\ -\frac{1}{X_{13}} & 0 & \frac{1}{X_{13}} + \frac{1}{X_{34}} \end{bmatrix} \begin{bmatrix} \theta_1 \\ \theta_2 \\ \theta_3 \end{bmatrix} \quad (7.26)$$

Il est temps de justifier la terminologie “courant continu”. Considérons pour cela le circuit de la figure 7.6.b, qui a la même topologie mais est purement résistif, la résistance de chacune de ses branches étant égale à la réactance de la branche correspondante du réseau de la figure 7.6.a. Supposons que l'on injecte aux noeuds de ce circuit des courants continus valant respectivement $P_1, -P_2, -P_3, P_4$. En prenant le noeud 4 comme référence des tensions, la méthode des noeuds fournit les relations :

$$\begin{bmatrix} P_1 \\ -P_2 \\ -P_3 \end{bmatrix} = \begin{bmatrix} \frac{1}{X_{12}} + \frac{1}{X_{13}} + \frac{1}{X_{14}} & -\frac{1}{X_{12}} & -\frac{1}{X_{13}} \\ -\frac{1}{X_{12}} & \frac{1}{X_{12}} + \frac{1}{X_{24}} & 0 \\ -\frac{1}{X_{13}} & 0 & \frac{1}{X_{13}} + \frac{1}{X_{34}} \end{bmatrix} \begin{bmatrix} V_{c1} - V_{c4} \\ V_{c2} - V_{c4} \\ V_{c3} - V_{c4} \end{bmatrix}$$

En supposant $V_{c4} = 0$, on voit que les tensions continues aux noeuds du circuit de la figure 7.6.b ne sont rien d'autre que les phases des tensions aux noeuds du réseau de la figure 7.6.a.

7.7 Analyse de sensibilité

La méthode de Newton est non seulement efficace pour résoudre les équations de load flow mais elle offre également la possibilité d'effectuer une analyse de sensibilité, comme expliqué ci-après.

7.7.1 Formule générale de sensibilité

Pour les développements qui suivent, il est confortable de présenter les équations de load flow (7.7 - 7.10) sous la forme compacte :

$$\varphi(\mathbf{x}, \mathbf{p}) = \mathbf{0} \quad (7.27)$$

où $\mathbf{x}$ est le vecteur des modules et phases des tensions et $\mathbf{p}$ un vecteur de paramètres. $\mathbf{x}$ est aussi appelé *vecteur d'état* car, une fois ce vecteur connu, on peut calculer l'état électrique complet du système. Soit n la dimension de $\mathbf{x}$ et m celle de $\mathbf{p}$.

Les développements qui suivent sont souvent utilisés dans le cas où les composantes de $\mathbf{p}$ sont des productions ou des consommations nodales mais on peut également les appliquer à des impédances de branches, des rapports de transformation, etc. . .

Soit $\eta(\mathbf{x}, \mathbf{p})$ une grandeur, fonction du vecteur d'état $\mathbf{x}$ et éventuellement du vecteur de paramètres $\mathbf{p}$. Suivant l'application, il peut s'agir de la tension en un noeud, du transit de puissance dans une branche, de la production réactive d'un générateur, des pertes actives dans le système, etc. . .

Soit $\mathbf{x}^*$ la solution de l'équation (7.27) pour la valeur $\mathbf{p}^*$ des paramètres. Supposons que l'on désire calculer de combien varierait η si l'on imposait une variation $\Delta \mathbf{p}$ des paramètres et que l'on recalculait l'état du système.

Une solution “brutale” consiste à résoudre les équations de load flow pour la nouvelle valeur des paramètres, c'est-à-dire à rechercher la valeur $\Delta \mathbf{x}$ telle que :

$$\varphi(\mathbf{x}^* + \Delta \mathbf{x}, \mathbf{p}^* + \Delta \mathbf{p}) = \mathbf{0} \quad (7.28)$$

et à calculer la valeur correspondante de η : $\eta(\mathbf{x}^* + \Delta \mathbf{x}, \mathbf{p}^* + \Delta \mathbf{p})$.

Cependant, si on se limite à de petites variations, il est possible de calculer *directement* le vecteur des sensibilités de η à $\mathbf{p}$:

$$S_{\eta \mathbf{p}} = \begin{bmatrix} \lim_{\Delta p_1 \rightarrow 0} \frac{\Delta \eta}{\Delta p_1} \\ \vdots \\ \lim_{\Delta p_m \rightarrow 0} \frac{\Delta \eta}{\Delta p_m} \end{bmatrix}$$

A cette fin, remplaçons les variations finies $\Delta \mathbf{p}$ et $\Delta \mathbf{x}$ par des variations infinitésimales $d\mathbf{p}$ et $d\mathbf{x}$ et linéarisons (7.28) :

$$\varphi(\mathbf{x}^* + d\mathbf{x}, \mathbf{p}^* + d\mathbf{p}) \simeq \varphi(\mathbf{x}^*, \mathbf{p}^*) + \varphi_{\mathbf{x}} d\mathbf{x} + \varphi_{\mathbf{p}} d\mathbf{p} = \varphi_{\mathbf{x}} d\mathbf{x} + \varphi_{\mathbf{p}} d\mathbf{p} = \mathbf{0} \quad (7.29)$$

où $\varphi_{\mathbf{x}}$ (resp. $\varphi_{\mathbf{p}}$) est la matrice jacobienne de φ par rapport à $\mathbf{x}$ (resp. $\mathbf{p}$), de dimensions $n \times n$ (resp. $n \times m$).

Supposant que la matrice $\varphi_{\mathbf{x}}$ n'est pas singulière, on tire de (7.29) :

$$d\mathbf{x} = -\varphi_{\mathbf{x}}^{-1} \varphi_{\mathbf{p}} d\mathbf{p} \quad (7.30)$$

Une linéarisation de la fonction $\eta(\mathbf{x}, \mathbf{p})$ donne par ailleurs :

$$d\eta = \sum_i \frac{\partial \eta}{\partial p_i} dp_i + \sum_i \frac{\partial \eta}{\partial x_i} dx_i$$

soit, sous forme matricielle :

$$d\eta = d\mathbf{p}^T \nabla_{\mathbf{p}} \eta + d\mathbf{x}^T \nabla_{\mathbf{x}} \eta \quad (7.31)$$

où :

$$\nabla_{\mathbf{p}} \eta = \begin{bmatrix} \frac{\partial \eta}{\partial p_1} \\ \vdots \\ \frac{\partial \eta}{\partial p_m} \end{bmatrix} \quad \nabla_{\mathbf{x}} \eta = \begin{bmatrix} \frac{\partial \eta}{\partial x_1} \\ \vdots \\ \frac{\partial \eta}{\partial x_n} \end{bmatrix}$$

En introduisant (7.30) dans (7.31), on obtient :

$$d\eta = d\mathbf{p}^T \nabla_{\mathbf{p}} \eta - d\mathbf{p}^T \varphi_{\mathbf{p}}^T (\varphi_{\mathbf{x}}^T)^{-1} \nabla_{\mathbf{x}} \eta = d\mathbf{p}^T \left[\nabla_{\mathbf{p}} \eta - \varphi_{\mathbf{p}}^T (\varphi_{\mathbf{x}}^T)^{-1} \nabla_{\mathbf{x}} \eta \right]$$

qui fournit directement les sensibilités recherchées :

$$S_{\eta \mathbf{p}} = \nabla_{\mathbf{p}} \eta - \boldsymbol{\varphi}_{\mathbf{p}}^T (\boldsymbol{\varphi}_{\mathbf{x}}^T)^{-1} \nabla_{\mathbf{x}} \eta \quad (7.32)$$

En pratique, la procédure pour calculer ces sensibilités est la suivante :

1. calculer $\nabla_{\mathbf{x}} \eta$
2. la matrice $\boldsymbol{\varphi}_{\mathbf{x}}$ étant disponible sous forme factorisée à l'issue de l'algorithme de Newton, résoudre le système linéaire :

$$\boldsymbol{\varphi}_{\mathbf{x}}^T \mathbf{y} = \nabla_{\mathbf{x}} \eta$$

par substitution

3. calculer $S_{\eta \mathbf{p}} = \nabla_{\mathbf{p}} \eta - \boldsymbol{\varphi}_{\mathbf{p}}^T \mathbf{y}$.

Comme on le voit, il est possible d'obtenir le vecteur des sensibilités en effectuant une seule opération de substitution, ce qui est particulièrement efficace.

7.7.2 Exemples

Dans les exemples qui suivent, on calcule la sensibilité de différentes grandeurs aux puissances actives et réactives spécifiées aux noeuds du réseau. Considérant les équations de load flow sous la forme (7.7 - 7.10), on a :

$$\mathbf{p} = \begin{bmatrix} \mathbf{p}^o \\ \mathbf{q}^o \end{bmatrix} \quad \text{et} \quad \boldsymbol{\varphi}_{\mathbf{p}} = - \begin{bmatrix} \mathbf{U}_p \\ \mathbf{0} \\ \mathbf{U}_q \\ \mathbf{0} \end{bmatrix}$$

où $\mathbf{U}_p$ et $\mathbf{U}_q$ sont des matrices comportant des 0 et des 1.

Sensibilité de la tension du i -ème noeud

On a simplement :

$$\eta = V_i \quad \nabla_{\mathbf{p}} \eta = \mathbf{0} \quad \nabla_{\mathbf{x}} \eta = \mathbf{e}_{V_i}$$

où $\mathbf{e}_{V_i}$ est un vecteur unité dont les composantes sont nulles, sauf celle correspondant à V_i , qui vaut 1.

Sensibilité de la production réactive du i -ème générateur

Le générateur considéré doit être du type PV. En effet, s'il était du type PQ, sa production réactive serait spécifiée et la sensibilité de celle-ci à n'importe quel paramètre serait nulle. On a :

$$\eta = Q_{gi}(\mathbf{x}) \quad \nabla_{\mathbf{p}} \eta = \mathbf{0}$$

où i est le numéro du noeud où est connecté le générateur et la fonction $Q_{gi}(\mathbf{x})$ est donnée par (7.4). Pour déterminer $\nabla_{\mathbf{x}}\eta = \nabla_{\mathbf{x}}Q_{gi}$, les dérivées partielles à calculer correspondent aux modules et aux phases des tensions au noeud i et à ses voisins directs. Les autres composantes de $\nabla_{\mathbf{x}}Q_{gi}$ sont nulles.

Sensibilité des pertes actives

Pour obtenir la fonction $\eta(\mathbf{x})$, on pourrait additionner les pertes dans toutes les branches du réseau, chacune étant une fonction des tensions aux extrémités de la branche. Une méthode plus directe et numériquement plus précise consiste à passer par le bilan de puissance (7.5) dont on déduit que les pertes valent :

$$p = P_N + \sum_{i=1}^{N-1} P_i = P_N(\mathbf{x}) + \sum_{i=1}^{N-1} P_i$$

Le premier terme du membre de droite représente la production du générateur balancier, donnée par (7.3). On a donc :

$$\eta = P_N(\mathbf{x}) + \sum_{i=1}^{N-1} P_i \quad \nabla_{\mathbf{p}}\eta = \begin{bmatrix} \mathbf{1} \\ \mathbf{0} \end{bmatrix}$$

où $\mathbf{1}$ est un vecteur de même dimension que $\mathbf{p}^o$ dont toutes les composantes valent 1 et le vecteur nul est de même dimension que $\mathbf{q}^o$. Pour déterminer $\nabla_{\mathbf{x}}\eta = \nabla_{\mathbf{x}}P_N$, les dérivées partielles à calculer correspondent aux modules et aux phases des tensions au noeud balancier et à ses voisins directs. Les autres composantes de $\nabla_{\mathbf{x}}P_N$ sont nulles.

Chapitre 8

La machine synchrone

La majeure partie de l'énergie électrique est produite à l'heure actuelle par les machines synchrones des centrales thermiques et hydrauliques. Les machines synchrones jouent un rôle important : ce sont elles qui imposent la fréquence des tensions et courants alternatifs et elles permettent de produire et absorber de la puissance réactive, nécessaire à la régulation des tensions.

Dans ce chapitre nous rappelons le principe de fonctionnement et nous établissons un modèle dynamique général de ce composant important. Nous considérons ensuite le cas particulier du fonctionnement en régime établi. Nous terminons par les limites de fonctionnement admissible.

8.1 Principe de fonctionnement

8.1.1 Champ magnétique créé par le stator

Une des motivations de l'emploi du système triphasé est la production d'un champ tournant dans les machines à courant alternatif.

Les machines électriques tournantes, tels les générateurs des centrales électriques, sont constituées d'un *stator*, qui est la partie fixe, et d'un *rotor*, qui est la partie tournante, séparée de la première par un mince *entrefer*. Stator et rotor sont tous deux fabriqués dans un matériau à haute perméabilité magnétique.

Le stator d'une machine tournante triphasée est équipé d'un ensemble de trois enroulements, correspondant chacun à une phase. Un de ces enroulements, que nous supposerons relatif à la phase *a*, est représenté en coupe à la figure 8.1.

Si l'on injecte un courant continu dans l'enroulement en question, les lignes du champ magnétique se présentent comme montré à la figure 8.1. La perméabilité magnétique du matériau utilisé

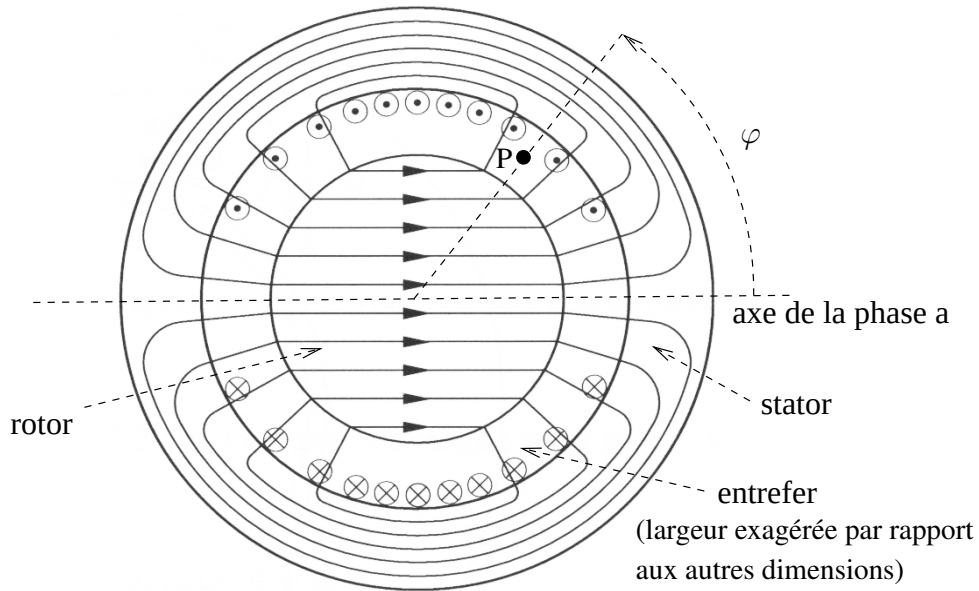


FIGURE 8.1 – ligne du champ magnétique créé par un courant continu parcourant la phase a

dans le rotor et le stator étant beaucoup plus élevée que celle de l'entrefer, les lignes de champ traversent ce dernier radialement (c'est-à-dire perpendiculairement à la surface extérieure du rotor et la surface intérieure du stator).

Repérons un point quelconque P de l'entrefer au moyen de l'angle φ , défini à la figure 8.1 et désignons par $B(\varphi)$ l'amplitude de l'induction magnétique en ce point.

$B(\varphi)$ est une fonction périodique (de période 2π) "en escalier". Afin d'induire des f.e.m. sinusoïdales dans les enroulements statoriques, les constructeurs s'efforcent de rendre cette fonction aussi proche que possible d'une sinusoïde, en répartissant judicieusement les conducteurs d'une même phase le long du stator. La distribution spatiale non uniforme des conducteurs est suggérée à la figure 8.1.

Supposant que la variation $B(\varphi)$ est sinusoïdale et tenant compte du fait que l'induction est proportionnelle au courant, on peut écrire pour la contribution de la phase a :

$$B(\varphi) = ki_a \cos \varphi \quad (8.1)$$

où le facteur k dépend de la géométrie, de l'emplacement et du nombre de conducteurs.

L'enroulement de la phase b (resp. c) est décalé spatialement de $2\pi/3$ (resp. $4\pi/3$) radians par rapport à celui de la phase a ¹. La figure 8.2.a montre la disposition des trois phases, en représentant symboliquement chaque enroulement par une seule spire, pour des raisons de lisibilité.

1. la position relative des phases dicte le sens de rotation du champ et donc du rotor. Dans le cas présent, on suppose que le rotor tourne dans le sens trigonométrique

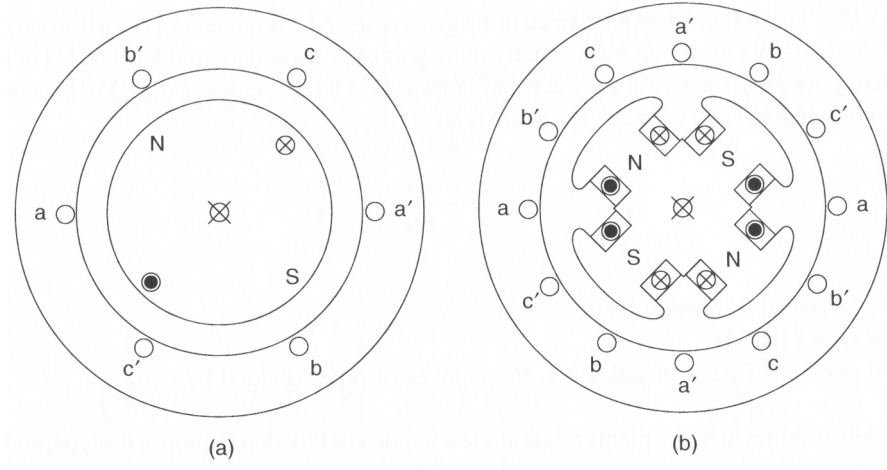


FIGURE 8.2 – disposition des enroulements statoriques triphasés : (a) machine à une paire de pôles ($p = 1$) ; (b) machine à deux paires de pôles ($p = 2$)

L'induction totale créée par les trois phases vaut donc, au point correspondant à l'angle φ :

$$B_{3\phi} = ki_a \cos \varphi + ki_b \cos\left(\varphi - \frac{2\pi}{3}\right) + ki_c \cos\left(\varphi - \frac{4\pi}{3}\right)$$

et si l'on alimente l'ensemble par les courants triphasés équilibrés :

$$\begin{aligned} B_{3\phi} &= \sqrt{2}kI \left[\cos(\omega t + \psi) \cos \varphi + \cos\left(\omega t + \psi - \frac{2\pi}{3}\right) \cos\left(\varphi - \frac{2\pi}{3}\right) + \right. \\ &\quad \left. + \cos\left(\omega t + \psi - \frac{4\pi}{3}\right) \cos\left(\varphi - \frac{4\pi}{3}\right) \right] \\ &= \frac{\sqrt{2}kI}{2} \left[\cos(\omega t + \psi + \varphi) + \cos(\omega t + \psi - \varphi) + \cos\left(\omega t + \psi + \varphi - \frac{4\pi}{3}\right) \right. \\ &\quad \left. + \cos(\omega t + \psi - \varphi) + \cos\left(\omega t + \psi + \varphi - \frac{2\pi}{3}\right) + \cos(\omega t + \psi - \varphi) \right] \\ &= \frac{3\sqrt{2}kI}{2} \cos(\omega t + \psi - \varphi) \end{aligned} \quad (8.2)$$

Cette équation est celle d'une onde qui circule dans l'entrefer à la vitesse angulaire ω , comme représenté à la figure 8.3, dans laquelle l'entrefer a été "déroulé".

Les trois courants triphasés produisent donc un champ magnétique tournant, comme le ferait un aimant tournant à la vitesse angulaire ω . Les pôles Nord et Sud de cet aimant coïncident avec les extrema de la fonction $B(\varphi)$ (cf figure 8.3).

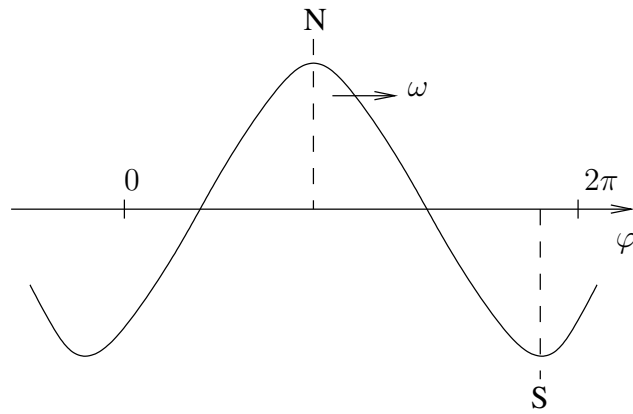


FIGURE 8.3 – onde d'induction circulant dans l'entrefer (déroulé)

8.1.2 Champ magnétique créé par le rotor

Le rotor d'une machine synchrone comporte un enroulement qui, en régime établi, est parcouru par du courant continu. Les lignes du champ magnétique créé par ce courant sont représentées à la figure 8.4, pour une position du rotor à un instant donné. Dans cette figure, l'enroulement d'excitation est représenté symboliquement par une seule spire.

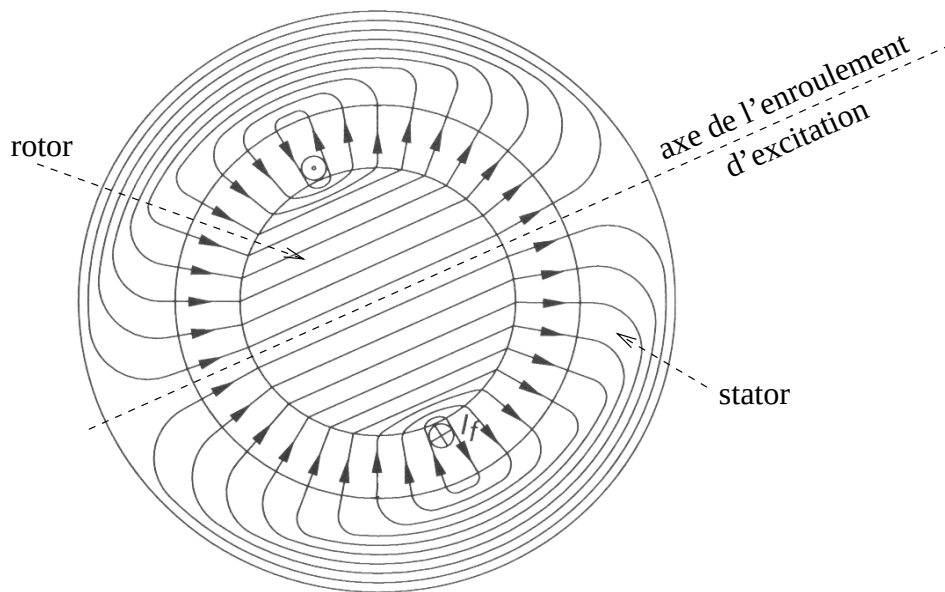


FIGURE 8.4 – lignes du champ magnétique créé par un courant continu I_f parcourant l'enroulement d'excitation

8.1.3 Interaction entre champs et conversion électromécanique

Une machine synchrone est caractérisée par le fait qu'en régime établi le rotor tourne à la même vitesse angulaire ω que le champ produit par le stator. Cette vitesse est appelée *vitesse de synchronisme*. En conséquence, les champs statorique et rotorique sont fixes l'un par rapport à l'autre et tournent tous deux à la vitesse de synchronisme.

Ces deux champs tendent à s'aligner à la façon de deux aimants attirés l'un par l'autre. Si l'on cherche à les écarter, un couple de rappel apparaît et s'y oppose. Ce couple à l'origine de la conversion d'énergie mécanique en énergie électrique et inversement.

Une analogie mécanique est proposée à la figure 8.5, où les deux cercles en mouvement à la vitesse angulaire ω correspondent aux champs rotorique et statorique, respectivement, et sont reliés par un ressort exerçant une force proportionnelle à son élongation. Les deux situations représentées sont les suivantes :

- la figure 8.5.a se rapporte à un générateur synchrone. La turbine entraîne le rotor en lui appliquant un couple mécanique T_m dans le sens de la rotation. Ceci tend le ressort qui applique au rotor un couple de rappel T_e dans le sens opposé à la rotation. En régime établi, la vitesse de rotation est constante ; les deux couples sont donc opposés et de même amplitude. Ainsi, si la puissance mécanique transmise par la turbine au rotor augmente, la vitesse reste constante (en régime établi), le couple mécanique T_m augmente, le ressort est davantage tendu, ce qui correspond à un couple électromagnétique T_e plus grand ;
- la figure 8.5.b se rapporte à un moteur synchrone. Le champ statorique entraîne le ressort qui est tendu et applique au rotor un couple d'entraînement T_e dans le sens de la rotation. La charge mécanique entraînée (ventilateur, pompe, compresseur, etc...) oppose un couple résistant T_m dans le sens opposé à la rotation. En régime établi, la vitesse de rotation est constante ; les deux couples sont donc opposés et de même amplitude. Ainsi, si la puissance mécanique demandée par la charge solidaire du rotor augmente, la vitesse reste constante (en régime établi), le couple mécanique T_m augmente, le ressort est davantage tendu, ce qui correspond à un couple électromagnétique T_e plus grand.

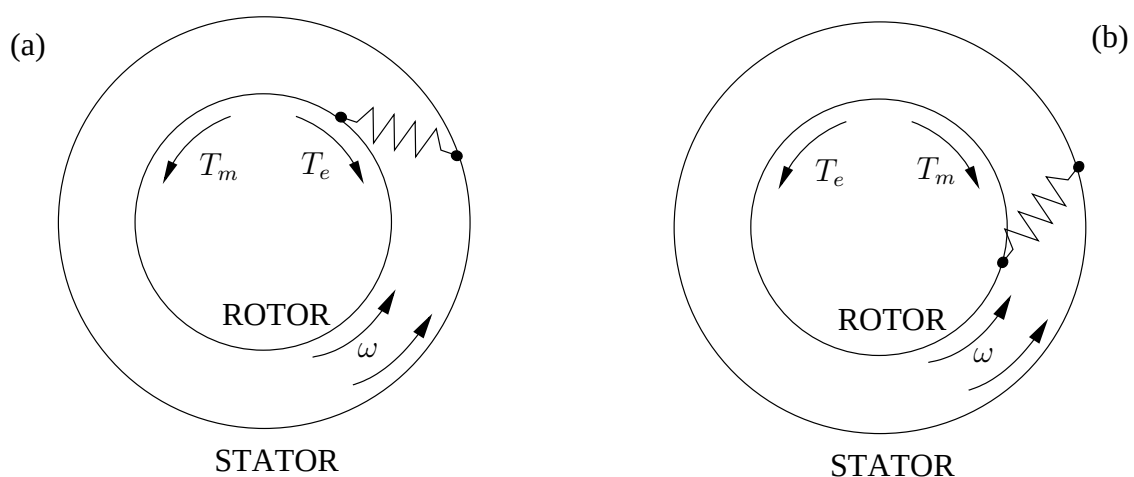


FIGURE 8.5 – couples dans la machine synchrone fonctionnant en : (a) générateur ; (b) moteur

Dans la machine synchrone, comme dans cette analogie mécanique, il existe une valeur maximale du couple électromagnétique T_e . Si le couple T_m tend à dépasser cette valeur, l'équilibre des couples est rompu et le rotor ne peut plus tourner à vitesse constante. Il y a *rupture de synchronisme* ; soit le rotor accélère par rapport au champ statorique, soit il ralentit.

8.1.4 Machines à plusieurs paires de pôles

Dans la machine représentée à la figure 8.2.a, l'enroulement d'excitation crée un électro-aimant avec un pôle Nord et un pôle Sud. Le pôle Nord balaie l'étendue complète d'une phase en une période T du signal, soit 20 millisecondes dans un système à 50 Hz. On dit qu'une telle machine possède une paire de pôles.

Certaines turbines requièrent que le rotor tourne à une vitesse plus faible (voir section 8.2). Or, les tensions et courants alternatifs au stator doivent conserver la même fréquence, ce qui impose que le balayage d'une phase par le rotor se fasse toujours en un temps T .

Pour satisfaire ces deux exigences :

- on utilise un rotor avec p paires de pôles ($p > 1$). Un exemple avec deux paires de pôles ($p = 2$) est donné à la figure 8.2.b. Durant une période T , le rotor effectue $1/p$ tour ;
- au stator, on installe p fois la séquence d'enroulements a, b et c. Un enroulement couvre un angle de π/p radians (au lieu de π radians dans le cas de la figure 8.2.a). De la sorte, chaque enroulement continue à être balayé par un pôle Nord et un pôle Sud durant une période T . Le tableau ci-dessous donne quelques exemples, pour un réseau à 50 ou à 60 Hz.

Les différents enroulements relatifs à une même phase sont reliés (en parallèle ou en série) pour aboutir à trois phases à la sortie de la machine.

nombre p de paires de pôles	vitesse angulaire ω/p en tours/minute	
	$f = 50$ Hz	$f = 60$ Hz
1	3000	3600
2	1500	1800
4	750	900
6	500	600
20	150	180
40	75	90

8.2 Les deux types de machines synchrones

Les machines synchrones ont toutes un stator portant des enroulements triphasés, comme indiqué précédemment. Notons que ce stator est constitué par un empilement de tôles (réalisées dans un

matériau à haute perméabilité magnétique) de manière à réduire le plus possible l'effet des courants de Foucault.

En revanche, on distingue deux types de machines synchrones, en fonction de la structure du rotor.

1. Machines à rotor lisse : cf figure 8.6.

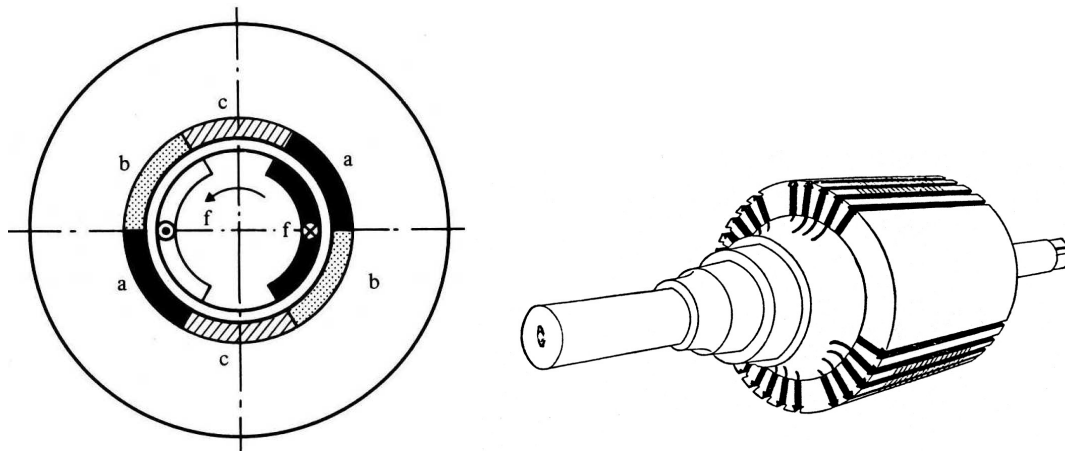


FIGURE 8.6 – machine à rotor lisse ($p = 1$)

Ces machines, appelées également *turbo-alternateurs*, sont généralement entraînées par des turbines à vapeur ou à gaz. Ces dernières fonctionnent de préférence à des vitesses élevées. Conformément à ce qui a été dit à la section 8.1.1, ces machines synchrones comportent une, ou au plus deux, paires de pôles. Dans un système à 50 Hz, elles tournent donc à 3000 tours par minute ($p = 1$, centrales thermiques classiques) ou à 1500 tours par minute ($p = 2$, centrales nucléaires).

Le rotor est constitué d'un cylindre en acier forgé, de forme allongée dont le diamètre est relativement petit par rapport à la longueur, afin de réduire les contraintes mécaniques liées à la force centrifuge. Les conducteurs de l'enroulement d'excitation sont logés dans des encoches, creusées longitudinalement dans le rotor, comme montré à la figure 8.6.

Les machines synchrones modernes ont un excellent rendement, de l'ordre de 99 %. Ceci signifie néanmoins qu'un générateur d'une puissance nominale de 500 MW, par exemple, est le siège de pertes calorifiques d'une puissance de 5 MW, ce qui est loin d'être négligeable ! Il importe de refroidir la machine afin d'éviter des "points chauds" qui peuvent détériorer les isolants entourant les conducteurs. A cette fin, on fait circuler de l'hydrogène ou de l'eau à l'intérieur des barres creuses qui constituent les enroulements statoriques. L'hydrogène a une capacité d'évacuation de la chaleur sept fois supérieure à celle de l'air et l'eau douze fois.

2. Machines à pôles saillants : cf figure 8.7.

Ces machines sont généralement entraînées par des turbines hydrauliques². Ces dernières tournent à des vitesses nettement plus faibles que les turbo-alternateurs. Les machines synchrones qu'elles entraînent doivent donc comporter un nombre de paires de pôles beaucoup plus élevé (au moins quatre en pratique). Or, il serait malaisé de loger de nombreux pôles dans un rotor cylindrique comme celui de la figure 8.6. Il est plus indiqué de les placer en

2. ou, dans les installations de petite puissance, par des moteurs diesels

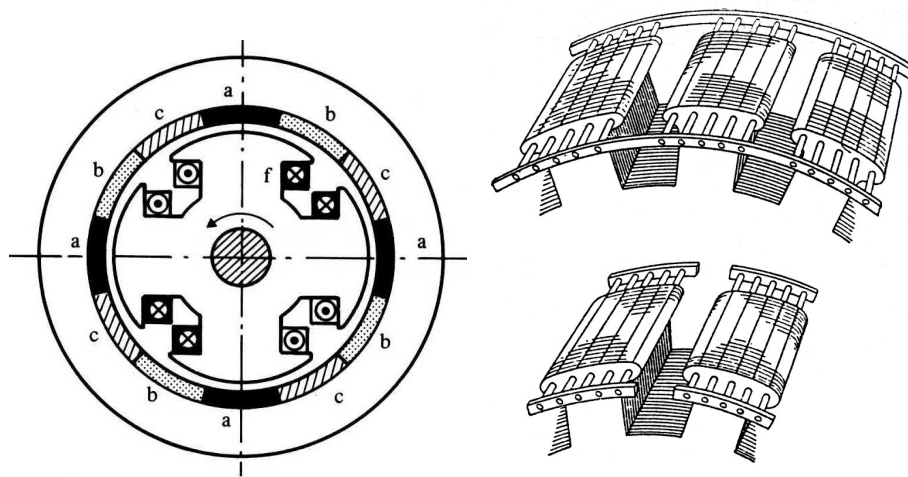


FIGURE 8.7 – machine à pôles saillants ($p = 2$)

“excroissance” comme représenté à la figure 8.7. L’entrefer d’une telle machine n’est pas d’épaisseur constante : il est minimum en face d’un pôle et maximum entre deux pôles.

Comparé à celui d’un turbo-alternateur de même puissance, le rotor à pôles saillants présente un diamètre nettement plus élevé (forces centrifuges plus faibles) et une longueur nettement plus courte.

Le rotor est généralement constitué d’un empilement de tôles magnétiques serrées les unes contre les autres. L’ensemble est calé sur l’axe de la machine, constitué d’un cylindre de diamètre plus faible.

Les générateurs à pôles saillants sont généralement refroidis par circulation d’air.

Les rotors des générateurs synchrones sont également équipés d’amortisseurs. Dans les machines à rotor lisse, il s’agit de conducteurs plats, logés dans les mêmes encoches que ceux de l’enroulement d’excitation et reliés en leurs extrémités pour permettre la circulation des courants induits. Dans les machines à pôles saillants, les amortisseurs sont constitués de barres logées dans les pôles et reliées à leurs extrémités par des anneaux (cf figure 8.7, schéma en haut à droite) ou par des segments (même figure, en bas à droite).

En régime établi parfait, aucun courant ne circule dans les barres d’amortisseur. En effet, les champs statorique et rotorique sont fixes par rapport au rotor ; le flux d’induction magnétique est donc constant dans le circuit constitué par les barres d’amortisseurs et aucune tension n’y est induite. Par contre, suite à une perturbation, il se peut que le rotor oscille par rapport au champ statorique. Des courants sont alors induits dans les barres d’amortisseurs. En vertu de la loi de Lenz, ces courants induits tendent à s’opposer à la cause qui les crée. Il apparaît donc un couple de rappel supplémentaire qui tend à amortir les oscillations du rotor et à réaligner ce dernier avec le champ statorique. Ce *couple d’amortissement* n’existe qu’en régime perturbé.

Dans les turbo-alternateurs, des courants sont induits dans la masse métallique du rotor et ces courants créent également un couple d’amortissement.

8.3 Modélisation au moyen de circuits magnétiquement couplés

Nous allons représenter la machine synchrone par un certain nombre d'enroulements, magnétiquement couplés, dont certains sont en mouvement.

Le modèle électrique d'une machine synchrone ne dépend pas du nombre p de paires de pôles qu'elle comporte (évidemment les valeurs de certains paramètres changent avec p). Pour des raisons de simplicité, nous considérerons donc une machine à une seule paire de pôles ($p = 1$).

8.3.1 Signification des enroulements

La machine idéalisée que nous allons étudier est représentée à la figure 8.8. Le stator est muni de 3 enroulements repérés a, b et c , décalés de 120 degrés. Le rotor comporte un certain nombre d'enroulements équivalents, répartis selon deux axes : l'*axe direct* qui coïncide avec celui de l'enroulement d'excitation et l'*axe en quadrature*, perpendiculaire au précédent. Nous plaçons arbitrairement l'axe en quadrature *en retard* sur l'axe direct par rapport au sens de rotation.

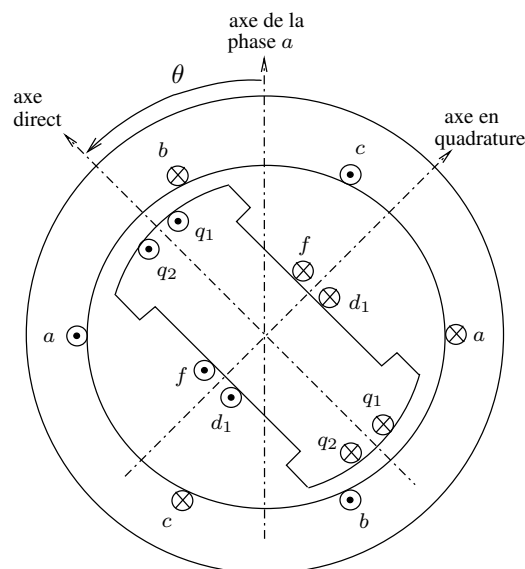


FIGURE 8.8 – machine synchrone idéalisée

Nous avons donné au rotor une forme en pôles saillants mais les développements qui suivent s'appliquent également à une machine à rotor lisse. Pour celle-ci, il suffit de considérer que le rotor présente une parfaite symétrie de révolution.

Le nombre d'enroulements rotoriques caractérise le degré de raffinement du modèle. Toutefois, il ne faut pas perdre de vue qu'un modèle plus sophistiqué requiert davantage de données pour tous les paramètres qui y interviennent et que le gain est marginal si les données ne sont pas fiables. Cette remarque prend tout son sens si l'on considère qu'en pratique seul le circuit d'excitation est accessible aux instruments de mesure. Les paramètres (résistances, inductances, ...) des autres

circuits sont déterminés de manière indirecte (p.ex. réponse de la machine lors d'un essai en court-circuit, réponse en fréquence, identification par calcul numérique).

Compte tenu de ces considérations, le modèle le plus répandu pour la machine synchrone est celui à quatre ou trois enroulements rotoriques. L'axe direct comporte l'enroulement d'excitation, désigné par f ³ et un circuit équivalent désigné par d_1 ⁴. Ce dernier représente l'effet des amortisseurs. L'axe en quadrature comporte deux enroulements, désignés par q_1 et q_2 ⁵. L'un représente l'effet des courants de Foucault induits dans la masse du rotor, l'autre tient compte des amortisseurs. Toutefois, dans les machines à pôles saillants, le rotor est généralement constitué de tôles et les courants de Foucault sont négligeables. Pour ces machines, on ne considère donc qu'un seul enroulement (q_2) dans l'axe en quadrature.

Les développements qui suivent s'appliquent au cas général d'une machine à quatre enroulements rotoriques. Le modèle à trois enroulements s'en déduit par des simplifications assez évidentes.

Notons enfin que l'enroulement d'excitation est soumis à une tension v_f tandis que les circuits d_1 , q_1 et q_2 sont court-circuités en permanence.

8.3.2 Relations tensions-courants-flux

Comme nous nous intéressons principalement à des générateurs, nous adoptons la convention générateur dans chaque enroulement statorique. En revanche, étant donné qu'on fournit de la puissance à l'enroulement d'excitation, nous y adoptons la convention moteur. Rappelons que ces choix sont arbitraires ; leur mérite est de conduire à des puissances positives pour un générateur en régime établi.

Pour les enroulements statoriques, on a :

$$\begin{aligned} v_a(t) &= -R_a i_a(t) - \frac{d\psi_a}{dt} \\ v_b(t) &= -R_a i_b(t) - \frac{d\psi_b}{dt} \\ v_c(t) &= -R_a i_c(t) - \frac{d\psi_c}{dt} \end{aligned}$$

où R_a est la résistance d'une des phases et ψ le flux *total* embrassé par l'enroulement considéré.

Ces relations s'écrivent sous forme matricielle :

$$\mathbf{v}_T = -\mathbf{R}_T \mathbf{i}_T - \frac{d}{dt} \boldsymbol{\psi}_T \quad (8.3)$$

où l'indice T désigne des grandeurs triphasées et où l'on a posé $\mathbf{R}_T = \text{diag}(R_a \ R_a \ R_a)$.

3. "f" pour "field winding" en anglais

4. "d" pour "direct"

5. "q" pour "quadrature"

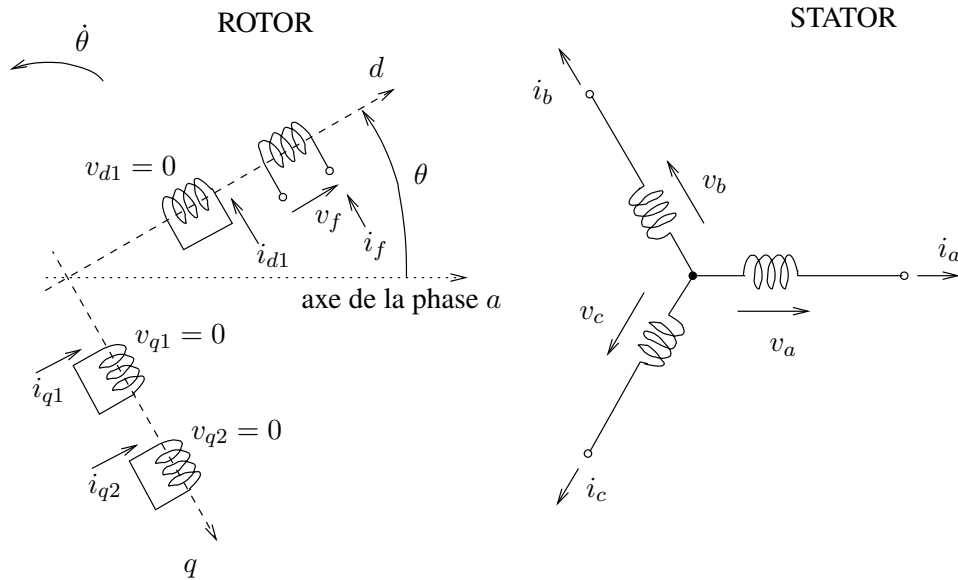


FIGURE 8.9 – enroulements de la machine synchrone : conventions de signe

Pour les enroulements rotoriques on a de même :

$$v_f(t) = R_f i_f(t) + \frac{d\psi_f}{dt} \quad (8.4)$$

$$0 = R_{d1} i_{d1}(t) + \frac{d\psi_{d1}}{dt} \quad (8.5)$$

$$0 = R_{q1} i_{q1}(t) + \frac{d\psi_{q1}}{dt} \quad (8.6)$$

$$0 = R_{q2} i_{q2}(t) + \frac{d\psi_{q2}}{dt} \quad (8.7)$$

et sous forme matricielle :

$$\mathbf{v}_r = \mathbf{R}_r \mathbf{i}_r + \frac{d}{dt} \boldsymbol{\psi}_r \quad (8.8)$$

où l'indice r désigne des grandeurs rotoriques et où l'on a posé $\mathbf{R}_r = \text{diag}(R_f \ R_{d1} \ R_{q1} \ R_{q2})$.

8.3.3 Inductances

Les équations ci-dessus sont tout-à-fait générales ; en particulier, aucune hypothèse n'est faite sur les propriétés du milieu magnétique. Dans ce cours, nous nous limitons toutefois au régime linéaire et négligeons la saturation du matériau magnétique.

Sous cette hypothèse flux et courants sont liés par :

$$\begin{bmatrix} \boldsymbol{\psi}_T \\ \boldsymbol{\psi}_r \end{bmatrix} = \begin{bmatrix} \mathbf{L}_{TT}(\theta) & \mathbf{L}_{Tr}(\theta) \\ \mathbf{L}_{Tr}^T(\theta) & \mathbf{L}_{rr} \end{bmatrix} \begin{bmatrix} \mathbf{i}_T \\ \mathbf{i}_r \end{bmatrix} \quad (8.9)$$

dans laquelle θ est la position angulaire du rotor, définie par convention comme l'angle entre l'axe direct du rotor et l'axe de la phase a (voir figures 8.8 et 8.9).

Les matrices d'inductances $\mathbf{L}_{TT}$ et $\mathbf{L}_{Tr}$ dépendent de la position du rotor. Ce n'est pas le cas de la matrice $\mathbf{L}_{rr}$ étant donné que, vu du rotor, le stator se présente toujours de la même manière, quelle que soit la position du rotor (on néglige ici l'effet des encoches dans lesquelles sont logés les conducteurs).

Les composantes de $\mathbf{L}_{TT}(\theta)$ et $\mathbf{L}_{Tr}(\theta)$ sont évidemment des fonctions périodiques. Développées en série de Fourier, celles-ci comportent, en principe, des harmoniques spatiaux. Comme mentionné à la section 8.1.1, on s'arrange en pratique pour rendre ces harmoniques aussi faibles que possible. Nous les négligerons donc, ce qui conduit au modèle de *machine sinusoïdale* dans lequel les matrices d'inductances prennent la forme suivante :

$$\mathbf{L}_{TT}(\theta) = \begin{bmatrix} L_0 + L_1 \cos 2\theta & -L_m - L_1 \cos 2(\theta + \frac{\pi}{6}) & -L_m - L_1 \cos 2(\theta - \frac{\pi}{6}) \\ -L_m - L_1 \cos 2(\theta + \frac{\pi}{6}) & L_0 + L_1 \cos 2(\theta - \frac{2\pi}{3}) & -L_m - L_1 \cos 2(\theta + \frac{\pi}{2}) \\ -L_m - L_1 \cos 2(\theta - \frac{\pi}{6}) & -L_m - L_1 \cos 2(\theta + \frac{\pi}{2}) & L_0 + L_1 \cos 2(\theta + \frac{2\pi}{3}) \end{bmatrix} \quad (8.10)$$

$$\mathbf{L}_{Tr}(\theta) = \begin{bmatrix} L_{af} \cos \theta & L_{ad1} \cos \theta & L_{aq1} \sin \theta & L_{aq2} \sin \theta \\ L_{af} \cos(\theta - \frac{2\pi}{3}) & L_{ad1} \cos(\theta - \frac{2\pi}{3}) & L_{aq1} \sin(\theta - \frac{2\pi}{3}) & L_{aq2} \sin(\theta - \frac{2\pi}{3}) \\ L_{af} \cos(\theta + \frac{2\pi}{3}) & L_{ad1} \cos(\theta + \frac{2\pi}{3}) & L_{aq1} \sin(\theta + \frac{2\pi}{3}) & L_{aq2} \sin(\theta + \frac{2\pi}{3}) \end{bmatrix} \quad (8.11)$$

$$\mathbf{L}_{rr} = \begin{bmatrix} L_{ff} & L_{fd1} & 0 & 0 \\ L_{fd1} & L_{d1d1} & 0 & 0 \\ 0 & 0 & L_{q1q1} & L_{q1q2} \\ 0 & 0 & L_{q1q2} & L_{q2q2} \end{bmatrix} \quad (8.12)$$

Dans ces expressions, toutes les constantes L sont positives, les signes $-$ adéquats ayant été introduits. Partant de la figure 8.8, ces différentes expressions se justifient comme suit :

- la self-inductance de la phase statorique a est maximale quand l'axe direct coïncide avec l'axe de cette phase ($\theta = 0$). En effet, les lignes de champ (dont le contour est montré à la figure 8.1) trouvent alors le chemin maximal dans le matériau ferromagnétique. Pour la même raison, la self-inductance est minimale quand l'axe en quadrature coïncide avec l'axe de la phase a ($\theta = \pi/2$). Par ailleurs, un retournement de 180 degrés du rotor ne modifie pas cette self-inductance ;
- les self-inductances des phases b et c se déduisent de celle de la phase a en remplaçant simplement θ par $\theta \pm 2\pi/3$;
- l'inductance mutuelle entre deux phases statoriques est maximale quand l'axe direct coïncide avec la bissectrice de l'angle aigu formé par leurs axes. Cette inductance mutuelle est toujours négative car un courant positif i_x crée dans la phase y un flux de sens opposé à celui créé par un courant i_y positif. Ici encore, un retournement du rotor de 180 degrés ne modifie pas cette inductance mutuelle
- l'inductance mutuelle entre un enroulement statorique et un enroulement rotorique est maximale quand ces enroulements sont coaxiaux et nulle quand il sont perpendiculaires. Un retournement de 180 degrés du rotor change le signe du couplage ;
- le caractère constant des termes de la matrice $\mathbf{L}_{rr}$ a déjà été justifié ;
- les termes nuls de la matrice $\mathbf{L}_{rr}$ se justifient par le fait que les enroulements sont perpendiculaires.

De toute évidence, une substitution de ces expressions dans les relations (8.3, 8.8) serait très fastidieuse, à cause de la dépendance angulaire. Ceci justifie le recours à de nouvelles variables, plus

appropriées que les grandeurs de phase $i_a, v_a, \dots$. Ce changement de variables indispensable est la *transformation de Park*⁶.

8.4 Transformation et équations de Park

8.4.1 La transformation de Park

La transformation de Park est définie par une matrice $\mathcal{P}$ qui s'applique aux grandeurs statoriques pour donner les grandeurs de Park correspondantes, de la manière suivante :

$$\mathbf{v}_P = \mathcal{P} \mathbf{v}_T \quad (8.13)$$

$$\boldsymbol{\psi}_P = \mathcal{P} \boldsymbol{\psi}_T \quad (8.14)$$

$$\mathbf{i}_P = \mathcal{P} \mathbf{i}_T \quad (8.15)$$

$$\text{où } \mathcal{P} = \sqrt{\frac{2}{3}} \begin{bmatrix} \cos \theta & \cos(\theta - \frac{2\pi}{3}) & \cos(\theta + \frac{2\pi}{3}) \\ \sin \theta & \sin(\theta - \frac{2\pi}{3}) & \sin(\theta + \frac{2\pi}{3}) \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix} \quad (8.16)$$

Les grandeurs vectorielles de Park sont repérées par l'indice P tandis que les composantes individuelles sont désignées par les indices d, q, o :

$$\begin{aligned} \mathbf{v}_P &= \begin{bmatrix} v_d & v_q & v_o \end{bmatrix}^T \\ \boldsymbol{\psi}_P &= \begin{bmatrix} \psi_d & \psi_q & \psi_o \end{bmatrix}^T \\ \mathbf{i}_P &= \begin{bmatrix} i_d & i_q & i_o \end{bmatrix}^T \end{aligned}$$

On montre aisément que :

$$\mathcal{P} \mathcal{P}^T = \mathbf{I}$$

où $\mathbf{I}$ est la matrice unité. Il en résulte que :

$$\mathcal{P}^{-1} = \mathcal{P}^T \quad (8.17)$$

La matrice de transformation $\mathcal{P}$ est donc orthogonale.

La transformation des variables T en variables P reçoit l'interprétation suivante.

Le courant i_a qui parcourt la phase a crée un champ magnétique d'amplitude $k i_a$ dirigé selon l'axe de la bobine a . De même, les champs créés par les courants i_b et i_c sont dirigés selon les axes des enroulements b et c . La projection sur l'axe d du champ total vaut :

$$k \left(\cos \theta i_a + \cos\left(\theta - \frac{2\pi}{3}\right) i_b + \cos\left(\theta - \frac{4\pi}{3}\right) i_c \right) = k \sqrt{\frac{3}{2}} i_d$$

6. du nom de son auteur. La transformation utilisée ici est une variante de celle proposée originellement par Park. Elle est plus simple, conserve les puissances et la symétrie des matrices d'inductance. Certains auteurs font référence à Blondel plutôt qu'à Park

et la projection sur l'axe q :

$$k \left(\sin \theta i_a + \sin\left(\theta - \frac{2\pi}{3}\right) i_b + \sin\left(\theta - \frac{4\pi}{3}\right) i_c \right) = k\sqrt{\frac{3}{2}} i_q$$

Considérons à présent deux enroulements fictifs d et q , situés respectivement sur l'axe direct et sur l'axe en quadrature (et tournant donc avec le rotor). Le courant i_d produit un champ magnétique dirigé selon l'axe d et d'amplitude $k' i_d$ tandis que le courant i_q produit un champ d'amplitude $k' i_q$ et dirigé selon l'axe q . Les relations ci-dessus montrent qu'à condition d'admettre que $k' = k\sqrt{3/2}$, les enroulements fictifs d et q , solidaires du rotor, produisent le même effet que les enroulements statoriques a , b et c .

On peut également considérer un enroulement fictif o , qui n'est pas couplé aux deux autres. Notons au passage que cet enroulement n'est parcouru par un courant qu'en cas de régime déséquilibré.

Après application de la transformation de Park, les enroulements de la machine synchrone sont ceux représentés à la figure 8.10.

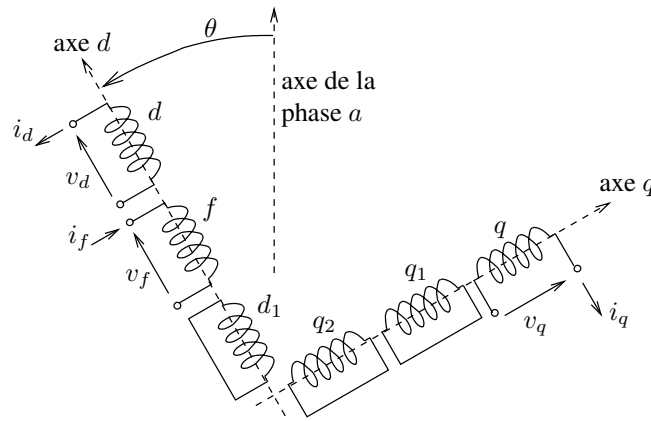


FIGURE 8.10 – enroulements de la machine synchrone après transformation de Park (o non montré)

8.4.2 Equations de Park de la machine synchrone

Partant de (8.3), on a successivement :

$$\begin{aligned} v_T &= -\mathbf{R}_T \mathbf{i}_T - \frac{d}{dt} \psi_T \\ \mathcal{P}^{-1} v_P &= -R_a \mathbf{I} \mathcal{P}^{-1} \mathbf{i}_P - \frac{d}{dt} (\mathcal{P}^{-1} \psi_P) \\ v_P &= -R_a \mathcal{P} \mathcal{P}^{-1} \mathbf{i}_P - \mathcal{P} \left(\frac{d}{dt} \mathcal{P}^{-1} \right) \psi_P - \mathcal{P} \mathcal{P}^{-1} \frac{d}{dt} \psi_P \\ v_P &= -\mathbf{R}_P \mathbf{i}_P - \dot{\theta} \mathbf{P} \psi_P - \frac{d}{dt} \psi_P \end{aligned} \quad (8.18)$$

avec

$$\begin{aligned} \mathbf{R}_P &= \mathbf{R}_T \\ \mathbf{P} &= \begin{bmatrix} 0 & 1 & 0 \\ -1 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix} \end{aligned}$$

$\mathbf{P}$ est un opérateur de rotation de 90 degrés dans le plan (d, q) .

Rappelons que la transformation de Park ne s'applique qu'au stator. Le rotor reste décrit par l'équation (8.8).

En détaillant (8.18), on trouve aisément les *équations de Park* :

$$v_d = -R_a i_d - \dot{\theta} \psi_q - \frac{d\psi_d}{dt} \quad (8.19)$$

$$v_q = -R_a i_q + \dot{\theta} \psi_d - \frac{d\psi_q}{dt} \quad (8.20)$$

$$v_o = -R_a i_o - \frac{d\psi_o}{dt} \quad (8.21)$$

dans lesquelles les termes du type $\dot{\theta} \psi$ sont appelés *forces électromotrices de rotation* et les termes en $d\psi/dt$ *forces électromotrices de transformation*.

8.4.3 Matrice des inductances de Park

Partant de (8.9), on a successivement :

$$\begin{aligned} \begin{bmatrix} \psi_T \\ \psi_r \end{bmatrix} &= \begin{bmatrix} \mathbf{L}_{TT} & \mathbf{L}_{Tr} \\ \mathbf{L}_{Tr}^T & \mathbf{L}_{rr} \end{bmatrix} \begin{bmatrix} \mathbf{i}_T \\ \mathbf{i}_r \end{bmatrix} \\ \begin{bmatrix} \mathcal{P}^{-1} \psi_P \\ \psi_r \end{bmatrix} &= \begin{bmatrix} \mathbf{L}_{TT} & \mathbf{L}_{Tr} \\ \mathbf{L}_{Tr}^T & \mathbf{L}_{rr} \end{bmatrix} \begin{bmatrix} \mathcal{P}^{-1} \mathbf{i}_P \\ \mathbf{i}_r \end{bmatrix} \\ \begin{bmatrix} \psi_P \\ \psi_r \end{bmatrix} &= \begin{bmatrix} \mathcal{P} \mathbf{L}_{TT} \mathcal{P}^{-1} & \mathcal{P} \mathbf{L}_{Tr} \\ \mathbf{L}_{Tr}^T \mathcal{P}^{-1} & \mathbf{L}_{rr} \end{bmatrix} \begin{bmatrix} \mathbf{i}_P \\ \mathbf{i}_r \end{bmatrix} \end{aligned}$$

Posons

$$\begin{bmatrix} \mathcal{P} \mathbf{L}_{TT} \mathcal{P}^{-1} & \mathcal{P} \mathbf{L}_{Tr} \\ \mathbf{L}_{Tr}^T \mathcal{P}^{-1} & \mathbf{L}_{rr} \end{bmatrix} = \begin{bmatrix} \mathbf{L}_{PP} & \mathbf{L}_{Pr} \\ \mathbf{L}_{rP} & \mathbf{L}_{rr} \end{bmatrix}$$

Nous laissons au lecteur le soin de vérifier que cette matrice prend la forme simple suivante :

$$\begin{bmatrix} \mathbf{L}_{PP} & \mathbf{L}_{Pr} \\ \mathbf{L}_{rP} & \mathbf{L}_{rr} \end{bmatrix} = \begin{bmatrix} L_{dd} & & & L_{df} & L_{dd1} & & \\ & L_{qq} & & & & L_{qq1} & L_{qq2} \\ & & L_{oo} & & & & \\ L_{df} & & & L_{ff} & L_{fd1} & & \\ L_{dd1} & & & L_{fd1} & L_{d1d1} & & \\ & L_{qq1} & & & & L_{q1q1} & L_{q1q2} \\ & L_{qq2} & & & & L_{q1q2} & L_{q2q2} \end{bmatrix} \quad (8.22)$$

où l'on a posé :

$$\begin{aligned}
L_{dd} &= L_0 + L_m + \frac{3}{2}L_1 \\
L_{qq} &= L_0 + L_m - \frac{3}{2}L_1 \\
L_{df} &= \sqrt{\frac{3}{2}}L_{af} \\
L_{dd1} &= \sqrt{\frac{3}{2}}L_{ad1} \\
L_{qq1} &= \sqrt{\frac{3}{2}}L_{aq1} \\
L_{qq2} &= \sqrt{\frac{3}{2}}L_{aq2} \\
L_{oo} &= L_0 - 2L_m
\end{aligned}$$

La matrice (8.22) est la *matrice des inductances de Park*. Contrairement à ceux de (8.9), ses termes sont tous indépendants de la position θ du rotor. Ce résultat était prévisible dans la mesure où, contrairement aux enroulements d'origine, les enroulements d, f, d_1, q, q_1 et q_2 sont tous fixes les uns par rapport aux autres (cf figure 8.10).

Par ailleurs, les enroulements se répartissent en deux groupes entre lesquels les inductances mutuelles sont nulles : d, f, d_1 d'une part et q, q_1, q_2 d'autre part. Ce résultat était également prévisible puisque les axes d et q sont perpendiculaires et que deux bobines d'axes perpendiculaires ont une inductance mutuelle nulle.

Abandonnant le circuit o et regroupant les circuits d, f, d_1 d'une part et q, q_1, q_2 d'autre part, les équations de la machine s'écrivent :

$$\begin{bmatrix} v_d \\ -v_f \\ 0 \end{bmatrix} = - \begin{bmatrix} R_a & & \\ & R_f & \\ & & R_{d1} \end{bmatrix} \begin{bmatrix} i_d \\ i_f \\ i_{d1} \end{bmatrix} - \begin{bmatrix} \dot{\theta}\psi_q \\ 0 \\ 0 \end{bmatrix} - \frac{d}{dt} \begin{bmatrix} \psi_d \\ \psi_f \\ \psi_{d1} \end{bmatrix} \quad (8.23)$$

$$\begin{bmatrix} v_q \\ 0 \\ 0 \end{bmatrix} = - \begin{bmatrix} R_a & & \\ & R_{q1} & \\ & & R_{q2} \end{bmatrix} \begin{bmatrix} i_q \\ i_{q1} \\ i_{q2} \end{bmatrix} + \begin{bmatrix} \dot{\theta}\psi_d \\ 0 \\ 0 \end{bmatrix} - \frac{d}{dt} \begin{bmatrix} \psi_q \\ \psi_{q1} \\ \psi_{q2} \end{bmatrix} \quad (8.24)$$

avec les relations flux-courants :

$$\begin{bmatrix} \psi_d \\ \psi_f \\ \psi_{d1} \end{bmatrix} = \begin{bmatrix} L_{dd} & L_{df} & L_{dd1} \\ L_{df} & L_{ff} & L_{fd1} \\ L_{dd1} & L_{fd1} & L_{d1d1} \end{bmatrix} \begin{bmatrix} i_d \\ i_f \\ i_{d1} \end{bmatrix} \quad (8.25)$$

$$\begin{bmatrix} \psi_q \\ \psi_{q1} \\ \psi_{q2} \end{bmatrix} = \begin{bmatrix} L_{qq} & L_{qq1} & L_{qq2} \\ L_{qq1} & L_{q1q1} & L_{q1q2} \\ L_{qq2} & L_{q1q2} & L_{q2q2} \end{bmatrix} \begin{bmatrix} i_q \\ i_{q1} \\ i_{q2} \end{bmatrix} \quad (8.26)$$

8.5 Energie, puissance et couple

8.5.1 Expression du couple électromagnétique

Nous allons à présent établir l'expression du couple électromagnétique T_e , en utilisant les équations (8.19-8.21) et en exprimant le bilan de puissance du stator et du rotor, respectivement.

Le bilan de puissance au stator de la machine s'écrit :

$$p_T + p_{Js} + \frac{dW_{ms}}{dt} = p_{r \rightarrow s} \quad (8.27)$$

où p_T la puissance instantanée sortant du stator, p_{Js} les pertes Joule statoriques, W_{ms} l'énergie magnétique emmagasinée dans les enroulements statoriques et $p_{r \rightarrow s}$ la puissance transférée du rotor au stator. A ce stade, nous ne connaissons pas encore la nature de $p_{r \rightarrow s}$ (puissance mécanique et/ou électrique ?).

La puissance instantanée sortant du stator vaut :

$$p_T(t) = v_a i_a + v_b i_b + v_c i_c = \mathbf{v}_T^T \mathbf{i}_T = \mathbf{v}_P^T \mathcal{P} \mathcal{P}^T \mathbf{i}_P = \mathbf{v}_P^T \mathbf{i}_P = v_d i_d + v_q i_q + v_o i_o$$

Ce résultat montre que la transformation de Park conserve la puissance : les enroulements (d, q, o) produisent la même puissance que les enroulements statoriques (a, b, c) .

En remplaçant les tensions par leurs expressions (8.19-8.21), l'expression ci-dessus devient :

$$p_T(t) = - \underbrace{(R_a i_d^2 + R_a i_q^2 + R_a i_o^2)}_{p_{Js}} - \underbrace{\left(i_d \frac{d\psi_d}{dt} + i_q \frac{d\psi_q}{dt} + i_o \frac{d\psi_o}{dt} \right)}_{dW_{ms}/dt} + \dot{\theta} (\psi_d i_q - \psi_q i_d) \quad (8.28)$$

Une comparaison avec (8.27) fournit directement l'expression de la puissance transférée du rotor au stator :

$$p_{r \rightarrow s} = \dot{\theta} (\psi_d i_q - \psi_q i_d) \quad (8.29)$$

Considérons à présent le bilan de puissance du rotor :

$$P_m + p_f = p_{Jr} + \frac{dW_{mr}}{dt} + p_{r \rightarrow s} + \frac{dW_c}{dt} \quad (8.30)$$

où P_m est la puissance mécanique fournie par la turbine, p_f la puissance électrique fournie à l'enroulement d'excitation, p_{Jr} les pertes Joule rotoriques, W_{mr} l'énergie magnétique dans les enroulements rotoriques et W_c l'énergie cinétique des masses tournantes.

La puissance p_f vaut $p_f = v_f i_f$ mais comme $v_{d1} = v_{q1} = v_{q2} = 0$ on peut encore écrire :

$$p_f = v_f i_f + v_{d1} i_{d1} + v_{q1} i_{q1} + v_{q2} i_{q2}$$

et en remplaçant les tensions par leurs expressions (8.4-8.7) :

$$p_f = \underbrace{(R_f i_f^2 + R_{d1} i_{d1}^2 + R_{q1} i_{q1}^2 + R_{q2} i_{q2}^2)}_{p_{Jr}} + \underbrace{i_f \frac{d\psi_f}{dt} + i_{d1} \frac{d\psi_{d1}}{dt} + i_{q1} \frac{d\psi_{q1}}{dt} + i_{q2} \frac{d\psi_{q2}}{dt}}_{dW_{mr}/dt}$$

En tenant compte de cette dernière relation et de (8.29), le bilan de puissance (8.30) devient simplement :

$$P_m - \frac{dW_c}{dt} = \dot{\theta}(\psi_d i_q - \psi_q i_d) \quad (8.31)$$

Considérons enfin l'équation du mouvement du rotor :

$$\mathcal{I} \frac{d^2\theta}{dt^2} = T_m - T_e \quad (8.32)$$

où $\mathcal{I}$ est le moment d'inertie des masses tournantes, T_m est le couple mécanique appliqué par la turbine et T_e est le couple électromagnétique. En multipliant cette dernière équation par $\dot{\theta}$ on obtient :

$$\mathcal{I} \dot{\theta} \ddot{\theta} = \dot{\theta} T_m - \dot{\theta} T_e \quad (8.33)$$

ou encore :

$$\frac{dW_c}{dt} = P_m - \dot{\theta} T_e \quad (8.34)$$

où P_m est la puissance mécanique communiquée par la turbine. Une comparaison de cette relation avec (8.31) donne directement l'expression (particulièrement simple !) du couple électromagnétique :

$$T_e = \psi_d i_q - \psi_q i_d \quad (8.35)$$

A posteriori, nous voyons que la puissance transmise du rotor au stator est exclusivement de nature mécanique.

8.5.2 Composantes du couple

En remplaçant les flux par leurs expressions tirées de (8.25, 8.26), la relation (8.35) devient :

$$T_e = L_{dd} i_d i_q + L_{df} i_f i_q + L_{dd1} i_{d1} i_q - L_{qq} i_q i_d - L_{qq1} i_{q1} i_d - L_{qq2} i_{q2} i_d$$

On peut distinguer trois composantes dans le couple :

$$T_{e1} = (L_{dd} - L_{qq}) i_d i_q \quad (8.36)$$

Cette composante n'existe que dans une machine à pôles saillants. Elle correspond au fait que, *même sans excitation* ($i_f = 0$), le rotor tend à aligner son axe direct sur le champ magnétique tournant créé par le stator, ce qui crée un certain couple. Dans cette position, les lignes du champ statorique passent au maximum dans le milieu ferromagnétique et au minimum dans l'entrefer.

En d'autres termes, le rotor tend à se positionner de manière à minimiser la reluctance offerte au champ statorique. T_{e1} est appelé *couple synchrone reluctant*⁷. Il est d'autant plus élevé que la *saillance* est marquée, c'est-à-dire que L_{dd} diffère fortement de L_{qq} ;

$$T_{e2} = L_{dd1}i_{d1}i_q - L_{qq1}i_{q1}i_d - L_{qq2}i_{q2}i_d \quad (8.37)$$

Cette composante est nulle en régime établi, car tous les courants d'amortisseurs sont nuls, comme mentionné précédemment. T_{e2} est un couple d'amortissement ;

$$T_{e3} = L_{df}i_fi_q \quad (8.38)$$

Cette composante, la seule dépendant du courant d'excitation, constitue la majeure partie du couple en régime établi. En régime établi, i_f est constant et T_{e3} est le *couple synchrone dû à l'excitation*. En régime perturbé, une composante dynamique de i_f apparaît, du même type que les courants d'amortisseurs, et une partie de T_{e3} contribue au couple d'amortissement total.

Remarque. Dans le cas d'une machine à p paires de pôles, la vitesse de rotation est $\dot{\theta}/p$ et la puissance transmise sous forme de couple est $\dot{\theta}T_e/p$. Le couple électromagnétique vaut donc :

$$T_e = p(\psi_d i_q - \psi_q i_d) \quad (8.39)$$

8.6 La machine synchrone en régime établi

Après avoir établi le modèle dynamique général, nous considérons le cas particulier d'une machine :

- dont le stator est parcouru par des courants triphasés équilibrés, de pulsation $\omega_N = 2\pi f_N$
- dont l'enroulement d'excitation est soumis à une tension continue V_f et est parcouru par un courant continu

$$i_f = \frac{V_f}{R_f} \quad (8.40)$$

- tournant à la vitesse de synchronisme :

$$\theta = \theta_o + \omega_N t \quad (8.41)$$

Pour des raisons déjà mentionnées, on a :

$$i_{d1} = i_{q1} = i_{q2} = 0 \quad (8.42)$$

8.6.1 Fonctionnement à vide

Le stator étant ouvert, on a évidemment :

$$i_a = i_b = i_c = 0$$

7. ce type de couple est à la base du moteur "à reluctance" utilisé dans les applications de positionnement

Il en résulte que :

$$i_d = i_q = i_o = 0$$

et pour les flux :

$$\begin{aligned}\psi_d &= L_{df} i_f \\ \psi_q &= 0\end{aligned}$$

Les équations de Park s'écrivent :

$$\begin{aligned}v_d &= 0 \\ v_q &= \omega_N \psi_d = \omega_N L_{df} i_f\end{aligned}$$

En repassant aux grandeurs statoriques par la transformation de Park inverse, on trouve, pour la phase a par exemple :

$$v_a(t) = \sqrt{\frac{2}{3}} \omega_N L_{df} i_f \sin(\theta_o + \omega_N t) = \sqrt{2} E_q \sin(\theta_o + \omega_N t)$$

où :

$$E_q = \frac{\omega_N L_{df} i_f}{\sqrt{3}} \quad (8.43)$$

est une force électromotrice proportionnelle au courant d'excitation. C'est aussi la tension apparaissant aux bornes de la machine à vide.

8.6.2 Fonctionnement en charge

Considérons à présent le régime défini par :

$$\begin{aligned}v_a(t) &= \sqrt{2} V \cos(\omega_N t + \phi) \\ v_b(t) &= \sqrt{2} V \cos(\omega_N t + \phi - \frac{2\pi}{3}) \\ v_c(t) &= \sqrt{2} V \cos(\omega_N t + \phi + \frac{2\pi}{3}) \\ i_a(t) &= \sqrt{2} I \cos(\omega_N t + \psi) \\ i_b(t) &= \sqrt{2} I \cos(\omega_N t + \psi - \frac{2\pi}{3}) \\ i_c(t) &= \sqrt{2} I \cos(\omega_N t + \psi + \frac{2\pi}{3})\end{aligned}$$

en plus des équations (8.40, 8.41 et 8.42), toujours d'application. On en déduit successivement :

$$\begin{aligned}i_d &= \sqrt{\frac{2}{3}} \sqrt{2} I [\cos(\theta_o + \omega_N t) \cos(\omega_N t + \psi) + \cos(\theta_o + \omega_N t - \frac{2\pi}{3}) \cos(\omega_N t + \psi - \frac{2\pi}{3}) \\ &\quad + \cos(\theta_o + \omega_N t + \frac{2\pi}{3}) \cos(\omega_N t + \psi + \frac{2\pi}{3})] \\ &= \frac{I}{\sqrt{3}} [\cos(\theta_o + 2\omega_N t + \psi) + \cos(\theta_o + 2\omega_N t + \psi - \frac{4\pi}{3}) + \cos(\theta_o + 2\omega_N t + \psi + \frac{4\pi}{3}) \\ &\quad + 3 \cos(\theta_o - \psi)] = \sqrt{3} I \cos(\theta_o - \psi)\end{aligned} \quad (8.44)$$

et par un calcul semblable :

$$i_q = \sqrt{3}I \sin(\theta_o - \psi) \quad (8.45)$$

$$i_o = 0$$

$$v_d = \sqrt{3}V \cos(\theta_o - \phi) \quad (8.46)$$

$$v_q = \sqrt{3}V \sin(\theta_o - \phi) \quad (8.47)$$

$$v_o = 0$$

En régime triphasé équilibré, les courants i_d et i_q sont donc constants. Ce résultat est conforme à l'interprétation de la transformation de Park. En effet, en régime établi, le champ statorique est fixe par rapport au rotor. Pour produire un tel champ avec les enroulements fictifs d et q , il faut injecter dans ces derniers des courants continus.

Les flux dans les enroulements d et q sont également constants et valent :

$$\psi_d = L_{dd}i_d + L_{df}i_f$$

$$\psi_q = L_{qq}i_q$$

L'expression (8.35) montre dès lors que le couple électromagnétique est également constant en régime établi. C'est un avantage supplémentaire important du système triphasé. En effet, au niveau de l'usure mécanique, on préfère que le couple appliqué au rotor d'une machine tournante soit constant, au lieu, par exemple, de présenter une composante alternative.

Les équations de Park (8.19 -8.21) s'écrivent :

$$v_d = -R_a i_d - \omega_N L_{qq} i_q = -R_a i_d - X_q i_q \quad (8.48)$$

$$v_q = -R_a i_q + \omega_N L_{dd} i_d + \omega_N L_{df} i_f = -R_a i_q + X_d i_d + \sqrt{3}E_q \quad (8.49)$$

$$v_o = 0$$

La réactance $X_d = \omega_N L_{dd}$ (resp. $X_q = \omega_N L_{qq}$) est appelée *réactance synchrone dans l'axe direct* (resp. *dans l'axe en quadrature*). Ces deux réactances sont des paramètres importants de la machine synchrone. Dans le cas d'une machine à rotor lisse, elles sont égales : $X_d = X_q$.

Notons que la relation (8.49) fait apparaître la f.e.m. E_q définie à la section précédente.

Montrons à présent que le fonctionnement de la machine peut être décrit par un diagramme de phaseur relativement simple.

Etant donné que le rotor de la machine tourne à la vitesse angulaire ω_N , il est possible de représenter sur une même figure les vecteurs tournants relatifs aux grandeurs sinusoïdales et les axes d et q de la machine, à condition de choisir correctement la référence des angles. Un tel diagramme est représenté à la figure 8.11. Cette figure montre le diagramme de phaseur à $t = 0$. L'axe horizontal représente à la fois l'axe sur lequel on projette les vecteurs tournants pour retrouver l'évolution temporelle des grandeurs sinusoïdales et l'axe par rapport auquel on mesure la position du rotor,

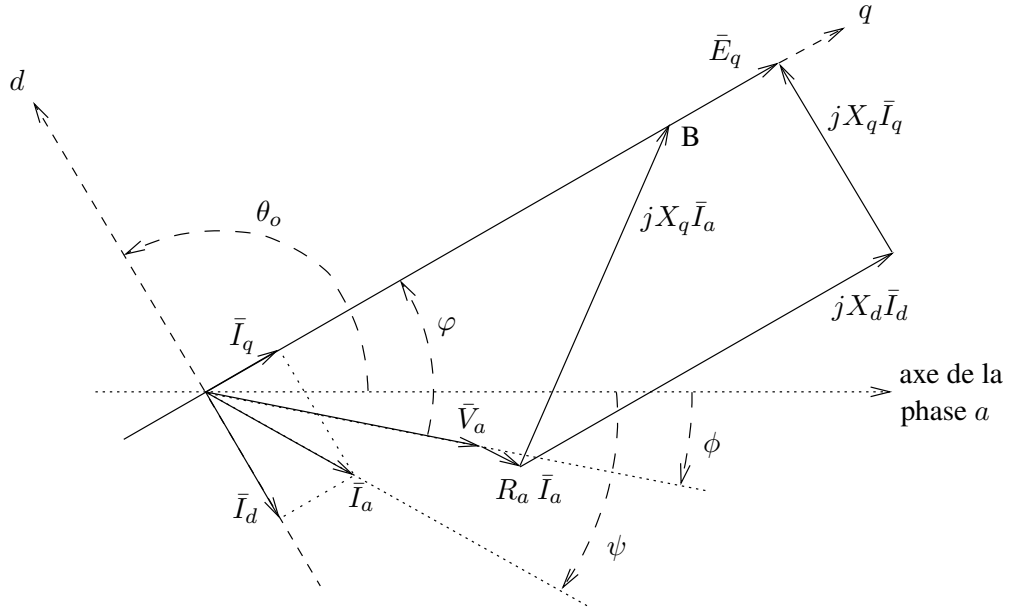


FIGURE 8.11 – diagramme de phaseur de la machine synchrone en régime établi

c'est-à-dire l'axe de la phase statorique a . L'angle entre cet axe horizontal et l'axe direct est donc la valeur en $t = 0$ de l'angle θ , soit θ_o .

En introduisant les relations (8.44-8.47) dans (8.48, 8.49), on obtient :

$$\begin{aligned} V \cos(\theta_o - \phi) &= -R_a I \cos(\theta_o - \psi) - X_q I \sin(\theta_o - \psi) \\ V \sin(\theta_o - \phi) &= -R_a I \sin(\theta_o - \psi) + X_d I \cos(\theta_o - \psi) + E_q \end{aligned}$$

Nous laissons au lecteur le soin de vérifier à partir de la figure 8.11, que ces deux équations sont en fait les projections sur les axes d et q de l'équation complexe :

$$\bar{E}_q = \bar{V}_a + R_a \bar{I}_a + jX_d \bar{I}_d + jX_q \bar{I}_q \quad (8.50)$$

dans laquelle $\bar{E}_q$ est un vecteur dirigé selon l'axe q et $\bar{I}_d$ (resp. $\bar{I}_q$) est la projection de $\bar{I}_a$ sur l'axe d (resp. q). On a donc :

$$\begin{aligned} \bar{E}_q &= E_q e^{j(\theta_o - \frac{\pi}{2})} \\ \bar{I}_d &= I \cos(\theta_o - \psi) e^{j\theta_o} = \frac{i_d}{\sqrt{3}} e^{j\theta_o} \\ \bar{I}_q &= I \sin(\theta_o - \psi) e^{j(\theta_o - \frac{\pi}{2})} = -j \frac{i_q}{\sqrt{3}} e^{j\theta_o} \end{aligned}$$

avec :

$$\bar{I}_a = \bar{I}_d + \bar{I}_q = \left(\frac{i_d}{\sqrt{3}} - j \frac{i_q}{\sqrt{3}} \right) e^{j\theta_o} \quad (8.51)$$

Dans le cas d'une machine à rotor lisse, $X_d = X_q = X$ et (8.50) devient simplement :

$$\bar{E}_q = \bar{V}_a + R_a \bar{I}_a + jX(\bar{I}_d + \bar{I}_q) = \bar{V}_a + R_a \bar{I}_a + jX \bar{I}_a \quad (8.52)$$

Il y correspond le schéma équivalent de la figure 8.12. Notons qu'il n'est pas possible de construire un tel schéma équivalent dans le cas d'une machine à pôles saillants.

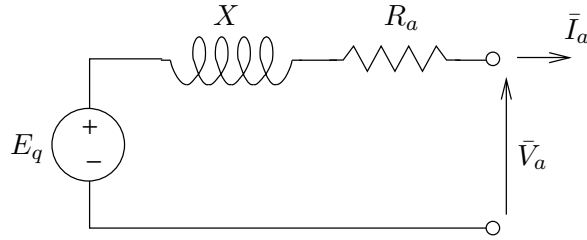


FIGURE 8.12 – schéma équivalent d'une machine à rotor lisse en régime établi

8.6.3 Puissances

Comme pour le courant, définissons les projections du vecteur $\bar{V}$:

$$\bar{V}_d = V \cos(\theta_o - \phi) e^{j\theta_o} = \frac{v_d}{\sqrt{3}} e^{j\theta_o} \quad (8.53)$$

$$\bar{V}_q = V \sin(\theta_o - \phi) e^{j(\theta_o - \frac{\pi}{2})} = -j \frac{v_q}{\sqrt{3}} e^{j\theta_o} \quad (8.54)$$

avec :

$$\bar{V}_a = \bar{V}_d + \bar{V}_q = \left(\frac{v_d}{\sqrt{3}} - j \frac{v_q}{\sqrt{3}} \right) e^{j\theta_o} \quad (8.55)$$

La puissance complexe fournie par la machine vaut :

$$S = 3\bar{V}_a \bar{I}_a^* = 3 \left(\frac{v_d}{\sqrt{3}} - j \frac{v_q}{\sqrt{3}} \right) \left(\frac{i_d}{\sqrt{3}} + j \frac{i_q}{\sqrt{3}} \right) = (v_d - j v_q)(i_d + j i_q)$$

dont on tire directement :

$$P = v_d i_d + v_q i_q \quad (8.56)$$

$$Q = v_d i_q - v_q i_d \quad (8.57)$$

Il est fort utile d'établir les expressions des puissances active et réactive en fonction de la tension V , de la f.e.m. E_q et de l'angle interne φ de la machine (cf figure 8.11). Pour ce faire, nous négligerons la résistance statorique R_a , qui, en pratique, est très faible devant X_d et X_q .

Sous cette hypothèse, les équations (8.48, 8.49) deviennent :

$$v_d = -X_q i_q \Rightarrow i_q = -\frac{v_d}{X_q}$$

$$v_q = X_d i_d + \sqrt{3} E_q \Rightarrow i_d = \frac{v_q - \sqrt{3} E_q}{X_d}$$

tandis que l'on tire de la figure 8.11 et de (8.53, 8.54) :

$$v_d = -\sqrt{3} V \sin \varphi$$

$$v_q = \sqrt{3} V \cos \varphi$$

Une substitution de toutes ces relations dans (8.56, 8.57) fournit les expressions recherchées :

$$P = 3 \frac{E_q V}{X_d} \sin \varphi + \frac{3V^2}{2} \left(\frac{1}{X_q} - \frac{1}{X_d} \right) \sin 2\varphi \quad (8.58)$$

$$Q = 3 \frac{E_q V}{X_d} \cos \varphi - 3V^2 \left(\frac{\sin^2 \varphi}{X_q} + \frac{\cos^2 \varphi}{X_d} \right) \quad (8.59)$$

qui, dans le cas d'une machine à rotor lisse, deviennent évidemment :

$$P = 3 \frac{E_q V}{X} \sin \varphi \quad (8.60)$$

$$Q = 3 \frac{E_q V}{X} \cos \varphi - 3 \frac{V^2}{X} \quad (8.61)$$

8.7 Valeurs nominales, système per unit et ordres de grandeur

8.7.1 Stator

Au stator, une machine synchrone est caractérisée par trois grandeurs nominales :

- la tension nominale U_N . C'est la tension pour laquelle la machine a été conçue. Un écart de quelques pour-cents par rapport à U_N est admissible ;
- le courant nominal I_N . C'est le courant maximal permanent pour lequel la section des enroulements statoriques a été prévue ;
- la puissance apparente nominale S_N . Il s'agit de la puissance triphasée liée aux grandeurs précédentes par :

$$S_N = \sqrt{3} U_N I_N$$

La conversion des paramètres de la machine en valeurs unitaires se fait conformément aux considérations de la section 5.3. On choisit la puissance de base $S_B = S_N$ et la tension de base $V_B = U_N / \sqrt{3}$. On en déduit le courant de base $I_B = S_N / 3V_B$ et l'impédance de base $Z_B = 3V_B^2 / S_B$.

Le tableau ci-après donne les ordres de grandeur des paramètres R_a , X_d et X_q , pour des machines d'une puissance supérieure à 100 MVA. Ces valeurs s'entendent dans la base de la machine, telle que définie plus haut. Dans un calcul de réseau utilisant une autre base, il y a lieu de procéder à une conversion, en utilisant par exemple la formule (5.11).

	machines à	
	rotor lisse	pôles saillants
résistance R_a	0.005 pu	
réactance dans l'axe direct X_d	1.5 - 2.5 pu	0.9 - 1.5 pu
réactance dans l'axe en quadrature X_q	1.5 - 2.5 pu	0.5 - 1.1 pu

Comme expliqué à la section 5.3, après passage en per unit, le coefficient 3 disparaît des formules (8.58 - 8.61) donnant la puissance triphasée.

8.7.2 Enroulements de Park

Nous avons montré que des grandeurs telles que $v_d, v_q, i_d, i_q, \dots$ se rapportent à des enroulements fictifs d et q solidaires du rotor. On peut également utiliser le système per unit dans ces enroulements. A cette fin, nous prenons dans chacun :

- comme puissance de base : S_N , pour les raisons exposées à la section 5.2
- comme tension de base : $\sqrt{3}V_B$, pour la raison expliquée ci-après

On en déduit le courant de base :

$$\frac{S_N}{\sqrt{3}V_B} = \sqrt{3}I_B$$

expression correspondant à un enroulement monophasé (et non triphasé).

Ce choix permet quelques simplifications confortables des relations établies à la section 8.6. Ainsi, par exemple, la relation (8.44) devient en per unit :

$$i_{dpu} = \frac{i_d}{\sqrt{3}I_B} = \frac{\sqrt{3}}{\sqrt{3}} \frac{I}{I_B} \cos(\theta_o - \psi) = I_{pu} \cos(\theta_o - \psi)$$

De même, les relations (8.45, 8.46 et 8.47) deviennent :

$$\begin{aligned} i_{qpu} &= I_{pu} \sin(\theta_o - \psi) \\ v_{dpu} &= V_{pu} \cos(\theta_o - \phi) \\ v_{qpu} &= V_{pu} \sin(\theta_o - \phi) \end{aligned}$$

tandis que (8.51) et (8.55) s'écrivent à présent :

$$\begin{aligned} \bar{I}_a &= \bar{I}_d + \bar{I}_q = (i_d - j i_q) e^{j\theta_o} \\ \bar{V}_a &= \bar{V}_d + \bar{V}_q = (v_d - j v_q) e^{j\theta_o} \end{aligned}$$

On voit qu'en per unit tous les coefficients $\sqrt{3}$ disparaissent et que les courants (resp. tensions) de Park sont directement les projections du vecteur $\bar{I}_a$ (resp. $\bar{V}_a$) sur les axes d et q de la machine.

8.7.3 Enroulements rotoriques

Dans les études dynamiques détaillées, on met également en per unit les grandeurs relatives à chaque enroulement rotorique. Nous ne détaillerons pas ici cette opération, qui n'est pas requise pour l'analyse du régime établi.

8.8 Courbes de capacité

Vu du réseau, le fonctionnement d'un générateur est caractérisé par trois grandeurs : la tension terminale V , la production active P et la production réactive Q .

Comme on l'imagine aisément, il existe des limites sur les valeurs que peuvent prendre P , Q et V , limites dictées par un fonctionnement admissible de la machine.

Les *courbes de capacité* délimitent le lieu des points de fonctionnement admissible dans le plan (P, Q) , à tension V constante. Cette dernière hypothèse est acceptable considérant que les générateurs sont dotés de régulateurs qui maintiennent leurs tensions (quasiment) constantes en fonctionnement normal (cf §11.2.1).

Un exemple de courbes de capacité est donné à la figure 8.13, correspondant à une machine à rotor lisse.

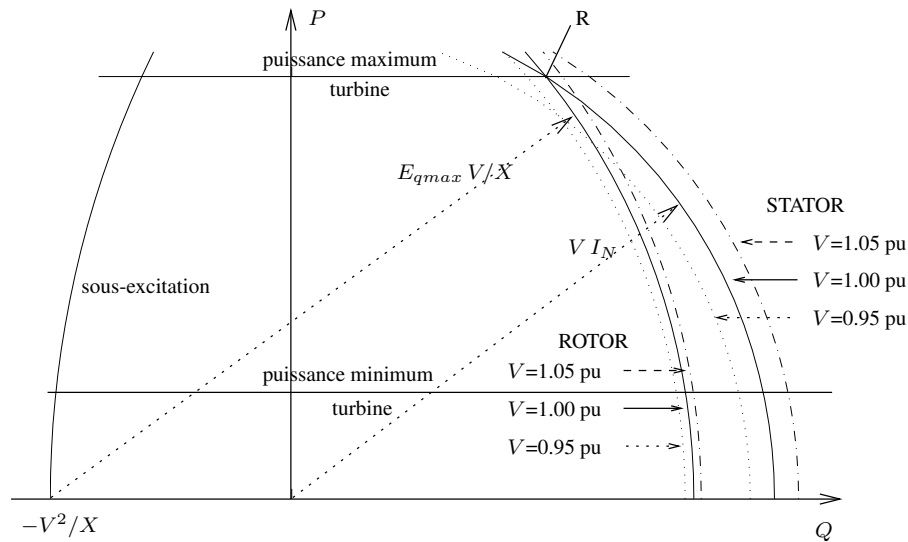


FIGURE 8.13 – courbes de capacité d'une machine à rotor lisse

On distingue les limites suivantes :

1. *puissance maximale turbine* : c'est évidemment aussi la puissance active maximale que le générateur peut fournir ;
2. *puissance minimale turbine*. Dans les centrales thermiques, la nécessité de produire une puissance minimale est liée à des problèmes de stabilité de la combustion ;
3. *limite statorique* : elle correspond à des points de fonctionnement pour lesquels le courant statorique est égal à I_N , défini à la section 8.7.1. On a, en per unit :

$$(S^2 =) P^2 + Q^2 = V^2 I_N^2$$

V étant fixé, cette équation est celle d'un cercle centré à l'origine, de rayon $V I_N$;

4. *limite rotorique* : elle correspond à des points de fonctionnement pour lesquels le courant d'excitation I_f est égal à la valeur maximale permise en régime permanent, pour des raisons d'échauffement.

L'expression de cette courbe s'établit aisément dans le cas d'une machine à rotor lisse, dont on néglige la résistance statorique. Notons I_{fmax} la valeur maximale du courant rotorique. En vertu de (8.43), la f.e.m. E_q est constante et égale à :

$$E_{qmax} = \frac{\omega_N L_{df} I_{fmax}}{\sqrt{3}}$$

Les relations (8.60, 8.61) s'écrivent, après passage en per unit :

$$P = \frac{E_{qmax} V}{X} \sin \varphi$$

$$Q = \frac{E_{qmax} V}{X} \cos \varphi - \frac{V^2}{X}$$

Une élimination de φ donne :

$$\left(\frac{V E_{qmax}}{X} \right)^2 = \left(Q + \frac{V^2}{X} \right)^2 + P^2$$

soit l'équation d'un cercle dont le centre est $(P = 0, Q = -V^2/X)$ et le rayon $V E_{qmax}/X$;

5. *limite en sous-excitation*. Il existe une puissance réactive maximale que la machine peut absorber, sous peine d'une rupture de synchronisme.

La figure 8.13 montre qu'une manière d'augmenter la puissance réactive maximale productible par une machine consiste à diminuer sa production de puissance active.

La même figure montre l'influence d'une variation de la tension terminale V . Pour une valeur donnée de P , une augmentation de V augmente les limites réactives, tant en production qu'en absorption (voir cependant la remarque plus loin).

A la figure 8.13, les courbes 1, 3 et 4 ci-dessus se croisent en un même point, sous $V = 1$ pu. En pratique, ce n'est pas toujours le cas mais cependant les points d'intersection de ces courbes deux à deux sont toujours très proches (ce qui traduit le fait que pour $P = P_{max}$ et $V = 1$ pu, le rotor et le stator sont dimensionnés de manière cohérente).

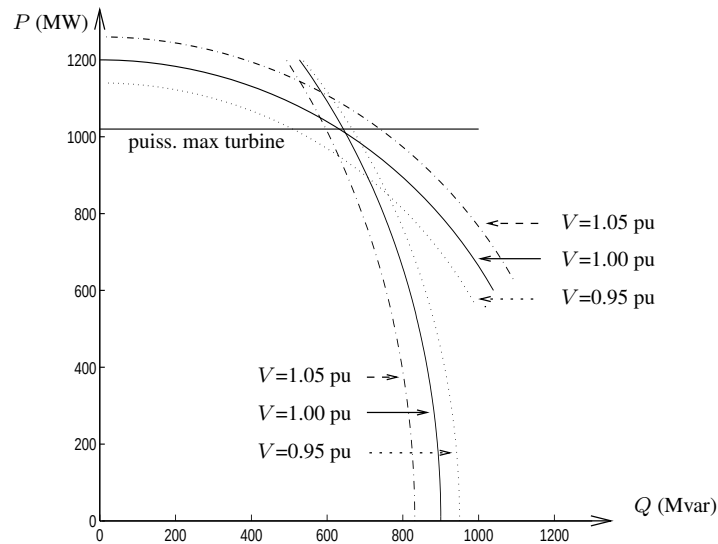


FIGURE 8.14 – courbes de capacité tenant compte de la saturation (courbes 1, 3 & 4, rotor lisse)

La figure 8.14 montre les courbes 1, 3 et 4 d'une machine réelle, compte tenu de la saturation du matériau. L'allure générale des courbes est celle de la figure 8.13. Cependant, les caractéristiques de saturation de cette machine sont telles que la limite rotorique devient moins contraignante quand la tension V diminue.

Chapitre 9

Comportement des charges

Ce chapitre aborde le comportement des charges.

Nous nous intéressons d'abord au moteur asynchrone triphasé. La machine asynchrone est très utilisée, principalement en tant que moteur. Elle est parfois utilisée pour produire de petites quantités d'énergie électrique, par exemple dans les éoliennes de première génération ou dans des centrales hydrauliques au fil de l'eau.

Nous nous intéressons ensuite à des modèles simples convenant à divers types de charges, utilisés couramment dans les études de grands systèmes et permettant de prendre en compte la variation des puissances consommées avec la tension et la fréquence.

9.1 Comportement du moteur asynchrone en tant que charge

9.1.1 Rappel : principe de fonctionnement

Rappelons brièvement le principe de fonctionnement d'une machine asynchrone en régime établi. Nous supposons pour simplifier que la machine possède une seule paire de pôles ($p = 1$) mais la théorie qui suit s'applique au cas où elle en possède plus.

Soit ω_s la pulsation des tensions et courants au stator. Le stator a la même structure que celui d'une machine synchrone. Son rôle est de produire un champ tournant à la vitesse angulaire ω_s .

Le rotor, quant à lui, tourne à une vitesse ω_m différente de la vitesse ω_s du champ tournant statorique, la différence étant caractérisée par le *glissement* :

$$g = \frac{\omega_s - \omega_m}{\omega_s} = 1 - \frac{\omega_m}{\omega_s} \quad (9.1)$$

Les circuits rotoriques, que l'on peut assimiler à des circuits triphasés, sont le siège de courants induits, alternatifs, de pulsation $g\omega_s$. Ces courants créent un champ tournant à la vitesse $g\omega_s$ par rapport au rotor, c'est-à-dire à la vitesse $(g\omega_s) + \omega_m = \omega_s$ par rapport au stator. Les deux champs sont donc fixes l'un par rapport à l'autre et leur interaction est à l'origine du couple électromagnétique, constant en régime établi.

Il existe deux catégories de machines asynchrones :

- moteurs à rotor *en cage (d'écureuil)* : le circuit rotorique est constitué de barres (en Al ou Cu) nues, court-circuitées en leurs extrémités par des anneaux pour permettre une circulation aisée du courant. Les moteurs à cage d'écureuil sont de construction simple, de maintenance aisée et fiables. Ils peuvent être dotés d'une double cage ; l'une sert alors à obtenir un couple électromagnétique plus important au démarrage ;
- moteurs à rotor *bobiné* : le rotor porte des enroulements isolés, généralement triphasés. Ces moteurs se rencontrent dans des applications où il faut avoir accès aux circuits rotoriques (via des bagues et des balais) pour contrôler le courant et/ou le couple de démarrage, ou encore la vitesse de rotation. Une technique très ancienne, par exemple, consiste à insérer des résistances dans les circuits rotoriques pour augmenter le couple de démarrage. Ces moteurs sont plus coûteux que les moteurs à cage d'écureuil eu égard à leur construction et leur maintenance.

Sans nuire à la généralité, la théorie de ce chapitre est établie pour un moteur à une cage.

9.1.2 Rappel : schéma équivalent en régime établi

Le lecteur voudra bien se reporter au cours de Conversion de l'énergie électromagnétique pour l'établissement du schéma équivalent d'une machine asynchrone triphasée, en régime établi, donné à la figure 9.1. Il s'agit bien entendu d'un schéma équivalent par phase. Dans ce dernier :

- R_s représente la résistance d'une phase statorique
- R_r représente la résistance d'une phase rotorique
- $L_{ss} - L_{sr}$ et $L_{rr} - L_{sr}$ sont les inductances de fuite et L_{sr} l'inductance magnétisante
- $\bar{V}$ est la tension phase-neutre aux bornes du stator
- $\bar{I}$ est le courant de phase au stator
- $\bar{I}_r$ est le courant circulant dans une phase rotorique, rapporté au stator : le vecteur tournant correspondant tourne à la vitesse angulaire ω_s .

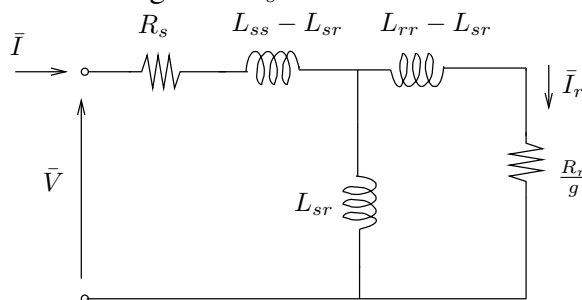


FIGURE 9.1 – schéma équivalent de la machine asynchrone en régime établi

Le tableau ci-après donne les ordres de grandeur des résistances et inductances, en per unit dans la base de la machine ¹.

R_s	0.01 - 0.12 pu	R_r	0.01 - 0.13 pu
$L_{ss} - L_{sr}$	0.07 - 0.15 pu	$L_{rr} - L_{sr}$	0.06 - 0.18 pu
L_{sr}	1.8 - 3.8 pu		

En régime établi, le bilan de puissance au stator s'écrit :

$$P = p_{Js} + p_{s \rightarrow r}$$

où P représente la puissance active consommée par le moteur, p_{Js} les pertes Joule au stator et $p_{s \rightarrow r}$ la puissance passant du stator au rotor (ou puissance “dans l'entrefer”).

On déduit aisément du schéma équivalent que $p_{s \rightarrow r}$ est la puissance dissipée dans la résistance équivalente R_r/g , soit :

$$p_{s \rightarrow r} = \frac{R_r}{g} I_r^2 \quad (9.2)$$

Les pertes Joule dans les résistances rotoriques valent $R_r I_r^2$. Elles représentent donc une fraction g de la puissance passant du stator au rotor. La fraction complémentaire est transformée en puissance mécanique, la machine développant un couple électromagnétique T_e sous une vitesse de rotation $\omega_m = (1 - g)\omega_s$.

On a donc :

$$\omega_m T_e = (1 - g) p_{s \rightarrow r}$$

dont on tire :

$$p_{s \rightarrow r} = \frac{\omega_m}{1 - g} T_e = \omega_s T_e \quad (9.3)$$

9.1.3 Couples et points de fonctionnement

Considérons une machine asynchrone alimentée sous une tension $\bar{V}$, comme représenté à la figure 9.2.a. Le schéma équivalent de Thévenin de la partie à gauche de AA' (voir figure 9.2.b) a pour paramètres :

$$\begin{aligned} \bar{V}_e &= \bar{V} \frac{j\omega_s L_{sr}}{R_s + j\omega_s L_{ss}} \\ R_e + jX_e &= j\omega_s (L_{rr} - L_{sr}) + \frac{j\omega_s L_{sr} (R_s + j\omega_s (L_{ss} - L_{sr}))}{R_s + j\omega_s L_{ss}} = j\omega_s L_{rr} + \frac{\omega_s^2 L_{sr}^2}{R_s + j\omega_s L_{ss}} \end{aligned}$$

1. Rappelons que, en per unit, les réactances à la pulsation nominale ont les mêmes valeurs que les inductances

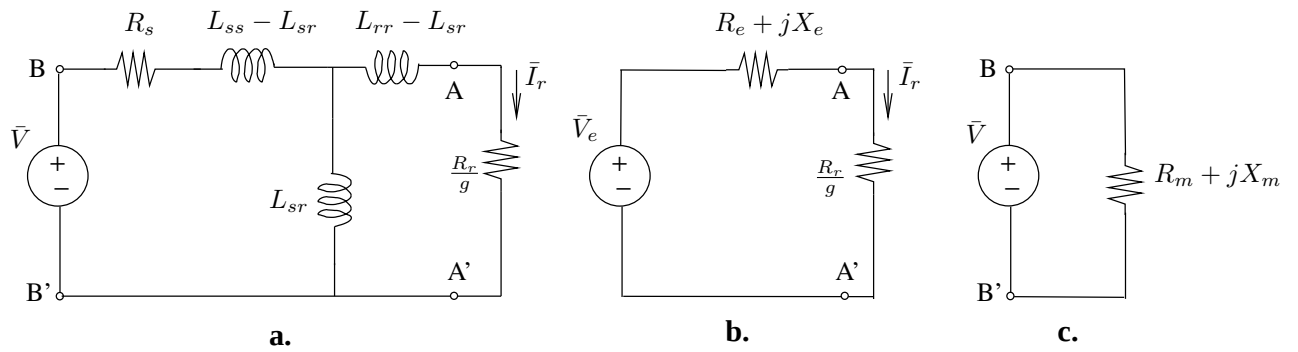


FIGURE 9.2 – manipulations du schéma équivalent

Des relations (9.2) et (9.3) on déduit :

$$T_e = \frac{1}{\omega_s} \frac{R_r}{g} I_r^2 = \frac{1}{\omega_s} \frac{R_r}{g} \frac{V_e^2}{(R_e + \frac{R_r}{g})^2 + X_e^2} \quad (9.4)$$

La variation du couple T_e en fonction du glissement g est montrée à la figure 9.3, pour deux valeurs de V . Ces courbes sont relatives à un moteur industriel de grande puissance, dont les paramètres sont : $L_{ss} = 3.867$, $L_{sr} = 3.800$, $L_{rr} = 3.970$, $R_s = 0.013$, $R_r = 0.009$ pu, $p = 1$.

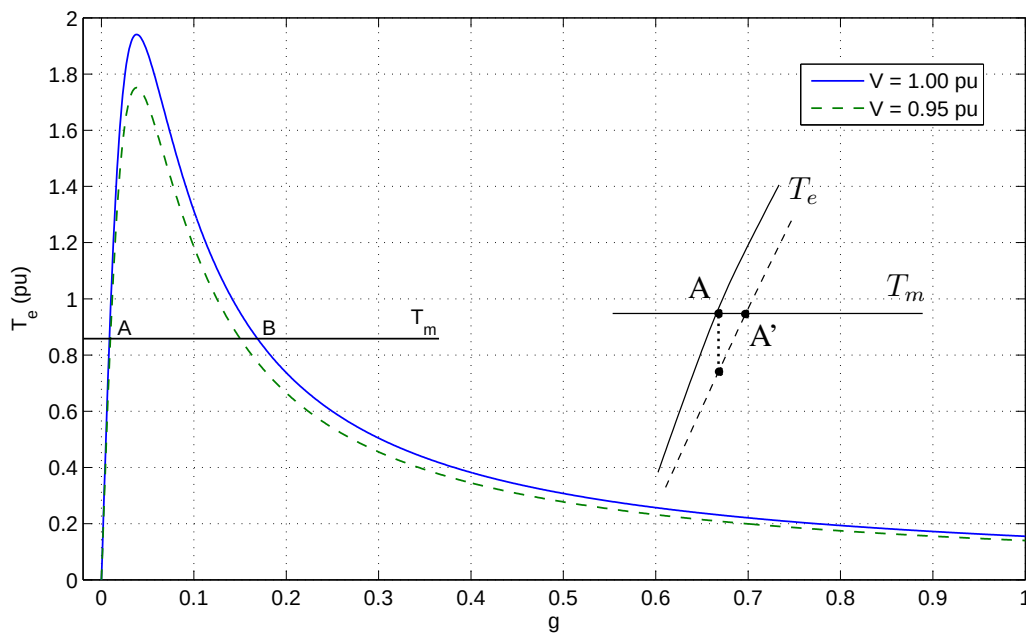


FIGURE 9.3 – variation du couple électromagnétique en fonction du glissement

En régime établi, on a évidemment $T_e = T_m$, où T_m est le couple mécanique (résistant). Ce dernier varie généralement avec la vitesse de rotation (et donc le glissement). La loi de variation dépend du type de charge mécanique entraînée (ventilateur, pompe, compresseur, etc...). Dans ce qui

suit, nous supposons pour simplifier que T_m est constant, ce qui est acceptable pour de faibles variations du glissement.

Pour un couple T_m donné, les points A et B à la figure 9.3 sont des points d'équilibre. Le raisonnement intuitif suivant montre que A est un point d'équilibre stable. En effet, le moteur fonctionnant en A, si l'on suppose qu'une perturbation augmente le glissement g , le couple électromagnétique devient supérieur au couple mécanique ; le moteur accélère, le glissement diminue et le moteur retourne vers son point de fonctionnement initial. De même, le point d'équilibre B est instable. En effet, si l'on applique la même perturbation lorsque le moteur fonctionne en B, le couple électromagnétique devient inférieur au couple mécanique ; le moteur décélère, le glissement augmente et le moteur s'éloigne davantage de son point de fonctionnement. On aboutit à la même conclusion si l'on considère une diminution initiale du glissement plutôt qu'une augmentation.

On voit qu'il existe un couple maximum au delà duquel le fonctionnement n'est pas possible. Ce couple maximum varie comme le carré de la tension V_e et donc le carré de la tension V . Le moteur doit être conçu pour fonctionner avec une marge de sécurité par rapport à ce couple maximum. Le couple T_m choisi à la figure 9.3 correspond aux conditions d'échauffement maximum car la puissance apparente absorbée par le moteur est proche de 1. pu.

9.1.4 Réponse à une chute de tension

Supposons que le moteur fonctionne initialement en A lorsqu'une chute de tension (en échelon) se produit. Un agrandissement du voisinage du point A est montré à la figure 9.3. La courbe du couple électromagnétique après perturbation est représentée en pointillé. Dans les premiers instants qui suivent cette perturbation, à cause de l'inertie des masses tournantes, le glissement du moteur ne peut changer et la résistance R_r/g conserve sa valeur d'avant perturbation. Le moteur se comporte donc comme une admittance constante.

Suite à la chute de tension, le couple T_e est devenu inférieur au couple T_m . Le moteur décélère donc et rejoint son nouveau point d'équilibre stable A'. En vertu de (9.3) la puissance passant du stator au rotor reprend sa valeur avant perturbation.

Comme on le voit, le moteur asynchrone est une charge qui, suite à une perturbation de la tension à ses bornes tend à restaurer à sa valeur avant perturbation une puissance consommée en interne. Ce processus de restauration est rapide : de l'ordre d'une seconde. Un tel comportement est inconfortable car dans les régimes perturbés où la tension chute, il est avantageux que les puissances consommées par les charges diminuent, pour soulager le réseau. Le moteur asynchrone n'a pas un tel comportement, du moins pas en ce qui concerne la puissance active.

Pour une chute de tension très importante, le couple électromagnétique T_e peut devenir inférieur au couple mécanique, auquel cas un "décrochage" du moteur² va se produire : ce dernier va ralentir jusqu'à s'arrêter. Cette augmentation du glissement a pour effet de diminuer la résistance R_r/g ; il s'en suit que le courant consommé augmente considérablement. En pratique, ceci peut amener une

2. en anglais : "motor stalling"

protection à déconnecter le moteur du réseau.

9.1.5 Variation des puissances active et réactive avec la tension et la fréquence

Considérons pour terminer la variation des puissances active et réactive avec la tension et la fréquence. Nous supposons toujours le couple mécanique constant. Les puissances active et réactive consommées par le moteur sont données par :

$$P = \frac{R_m}{R_m^2 + X_m^2} V^2 \quad (9.5)$$

$$Q = \frac{X_m}{R_m^2 + X_m^2} V^2 \quad (9.6)$$

où $R_m + jX_m$ est l'impédance équivalente du moteur vu de l'accès BB', comme représenté à la figure 9.2.c :

$$\begin{aligned} R_m + jX_m &= R_s + j\omega_s(L_{ss} - L_{sr}) + \frac{j\omega_s L_{sr} \left(\frac{R_r}{g} + j\omega_s(L_{rr} - L_{sr}) \right)}{\frac{R_r}{g} + j\omega_s L_{rr}} \\ &= R_s + j\omega_s L_{ss} + \frac{\omega_s^2 L_{sr}^2}{\frac{R_r}{g} + j\omega_s L_{rr}} \end{aligned} \quad (9.7)$$

R_m et X_m dépendent du glissement. Ce dernier s'obtient par la condition d'équilibre des couples :

$$T_m = T_e \quad \Leftrightarrow \quad T_m = \frac{1}{\omega_s} \frac{R_r}{g} \frac{V_e^2}{(R_e + \frac{R_r}{g})^2 + X_e^2} \quad (9.8)$$

où R_e , X_e et V_e ont été définis précédemment.

Pour une paire (V, ω_s) donnée, le glissement s'obtient en résolvant (9.8). A cette fin, il est plus aisé de considérer R_r/g comme inconnue intermédiaire, puis d'en tirer la valeur de g . On peut alors calculer les valeurs de R_m et X_m à partir de (9.7) et les puissances à partir de (9.5, 9.6).

Les figures 9.4 à 9.7 montrent les variations recherchées pour le moteur considéré à la figure 9.3 et pour deux valeurs du couple. On notera la plage de variation de la fréquence par rapport à celle de la tension. Ces figures appellent les commentaires suivants :

- les courbes sont limitées aux tensions basses par le décrochage du moteur ;
- la puissance active varie très peu avec la tension. En fait, elle augmente légèrement quand la tension diminue ;
- la puissance active varie quasi linéairement avec la fréquence. La sensibilité de P à ω_s est d'autant plus prononcée que le moteur est chargé ;
- la puissance réactive évolue significativement avec la tension. Lorsque la tension augmente, la consommation de la réactance magnétisante, proportionnelle au carré de la tension, finit par dominer. Par contre, lorsque la tension diminue, c'est la consommation des réactances de fuite qui finit par l'emporter. Entre les deux, il existe une tension où la sensibilité de Q à V change de signe. Cette tension augmente avec la charge du moteur ;
- la sensibilité de Q à ω_s peut être positive ou négative selon le niveau de charge du moteur.

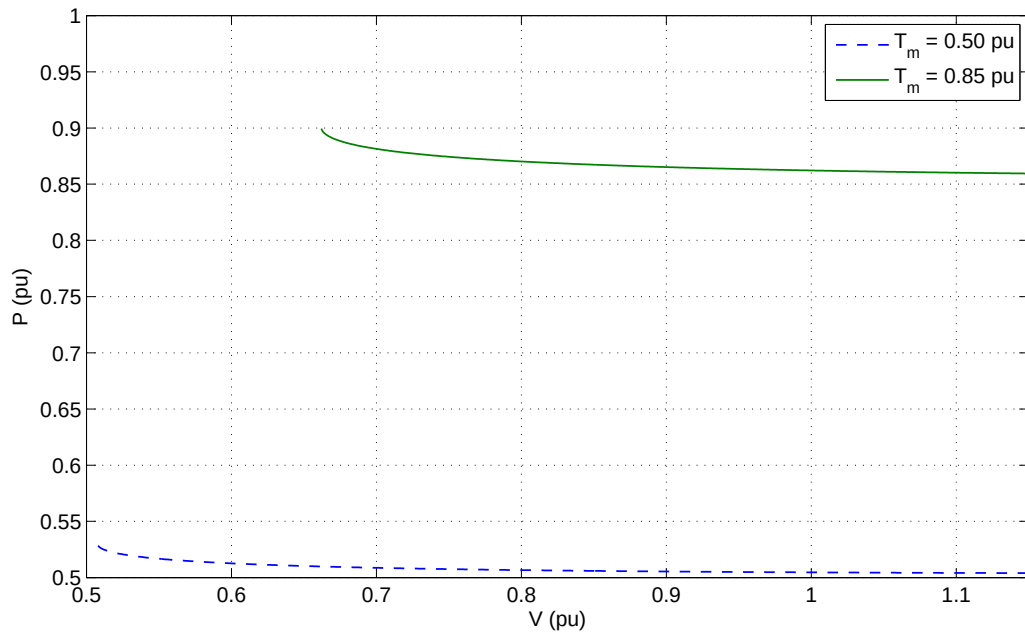


FIGURE 9.4 – variation avec la tension de la puissance active consommée par un moteur

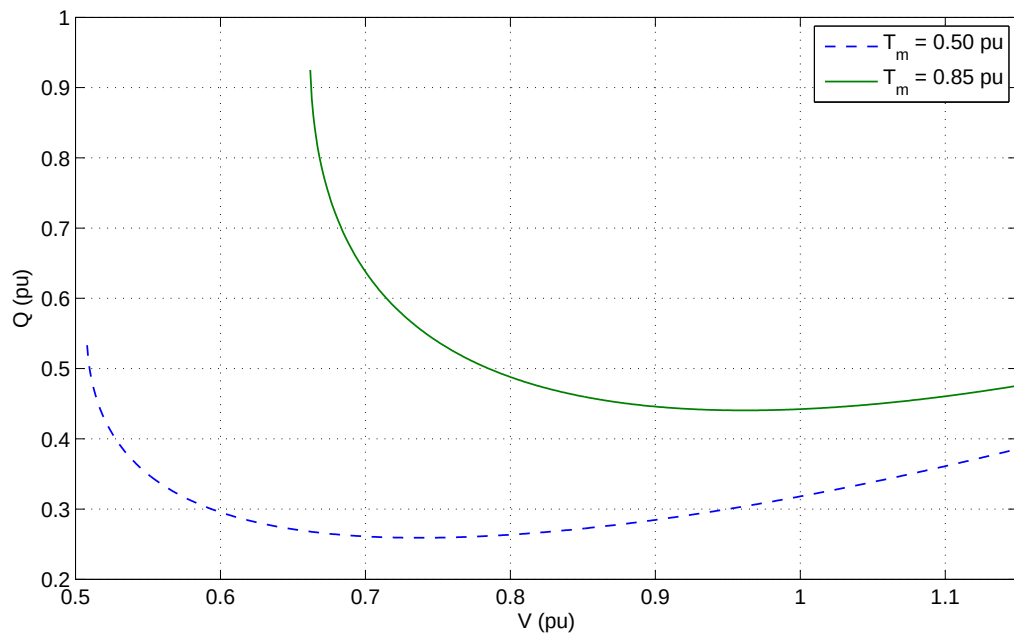


FIGURE 9.5 – variation avec la tension de la puissance réactive consommée par un moteur

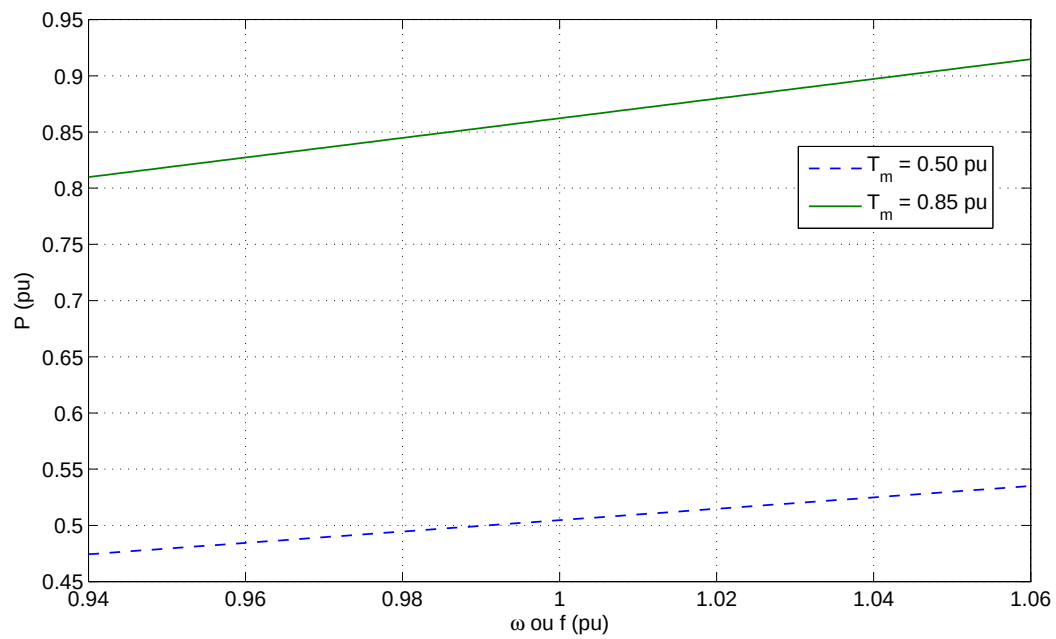


FIGURE 9.6 – variation avec la fréquence de la puissance active consommée par un moteur

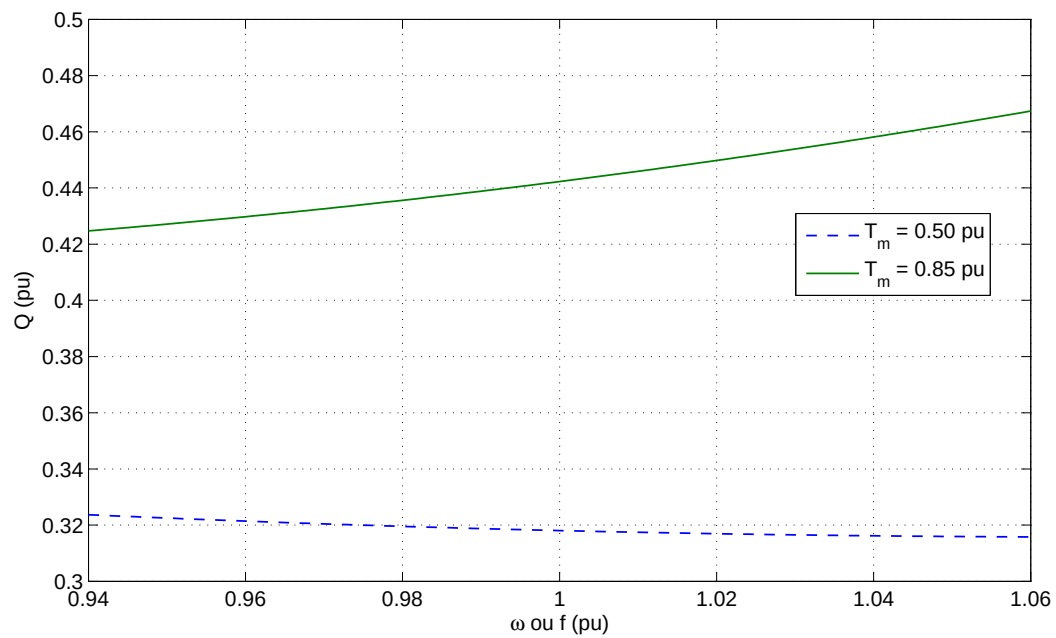


FIGURE 9.7 – variation avec la fréquence de la puissance réactive consommée par un moteur

9.2 Modèles simples des variations des charges avec la tension et la fréquence

9.2.1 Modèles dynamique et statique

Le modèle dynamique général d'une charge peut s'écrire :

$$P = H_P(V, f, \mathbf{x}) \quad (9.9)$$

$$Q = H_Q(V, f, \mathbf{x}) \quad (9.10)$$

$$\dot{\mathbf{x}} = \mathbf{g}(V, f, \mathbf{x}) \quad (9.11)$$

où P (resp. Q) est la puissance active (resp. réactive) consommée, V est le module de la tension aux bornes de la charge, f la fréquence de cette tension et $\mathbf{x}$ un vecteur d'état relatif au processus dynamique pouvant exister à l'intérieur de cette charge.

Dans de nombreux cas, on se contente d'un modèle statique, soit que la dynamique est négligeable, soit qu'on ne s'y intéresse pas, soit que l'on ne dispose pas de données fiables à son sujet.

Le modèle statique est obtenu en considérant que la dynamique interne est à l'équilibre, ce qui se traduit par :

$$\mathbf{g}(V, f, \mathbf{x}) = \mathbf{0} \quad (9.12)$$

En éliminant $\mathbf{x}$ des relations (9.9, 9.10, 9.12), on obtient formellement :

$$P = h_P(V, f) \quad (9.13)$$

$$Q = h_Q(V, f) \quad (9.14)$$

C'est aux modèles de ce type que nous nous intéressons dans le reste de ce chapitre.

9.2.2 Modèle à exposant de charges individuelles

Considérons dans un premier temps la variation de la charge avec la tension. Un modèle statique très utilisé en pratique est le modèle à *exposant* :

$$P = P_o \left(\frac{V}{V_o} \right)^\alpha \quad (9.15)$$

$$Q = Q_o \left(\frac{V}{V_o} \right)^\beta \quad (9.16)$$

dans lequel V_o est une tension de référence et P_o (resp. Q_o) est la puissance active (resp. réactive) consommée sous cette tension. α et β caractérisent le type de la charge, tandis que P_o et Q_o représentent le "volume" d'équipements de ce type.

Notons que le facteur de puissance d'une telle charge dépend de la tension, si $\alpha \neq \beta$, ce qui est souvent le cas en pratique.

Mentionnons quelques cas particuliers :

- $\alpha = \beta = 2$: charge à admittance constante
- $\alpha = \beta = 1$: charge à courant constant
- $\alpha = \beta = 0$: charge à puissance constante.

Choix de la tension de référence. La puissance active consommée sous une tension V_1 vaut :

$$P_1 = P_o \left(\frac{V_1}{V_o} \right)^\alpha$$

En tirant P_o de cette relation et en remplaçant dans (9.15), on trouve :

$$P = P_1 \left(\frac{V}{V_1} \right)^\alpha$$

avec une relation semblable pour la puissance réactive. On voit donc que la tension de référence peut être choisie arbitrairement sans que la caractéristique soit modifiée, à condition de prendre pour P_o et Q_o les puissances consommées sous cette tension de référence.

Interprétation des exposants. α et β peuvent être interprétés comme suit. Considérons une variation de tension ΔV pour laquelle on peut linéariser (9.15) en :

$$\Delta P = \alpha P_o \frac{V^{\alpha-1}}{V_o^\alpha} \Delta V$$

Évaluée à la tension de référence V_o , cette expression donne :

$$\frac{\Delta P}{P_o} = \alpha \frac{\Delta V}{V_o} \Leftrightarrow \alpha = \frac{\Delta P / P_o}{\Delta V / V_o} \quad (9.17)$$

avec un résultat similaire pour la puissance réactive :

$$\frac{\Delta Q}{Q_o} = \beta \frac{\Delta V}{V_o} \Leftrightarrow \beta = \frac{\Delta Q / Q_o}{\Delta V / V_o} \quad (9.18)$$

Comme on le voit, α et β représentent les sensibilités “normalisées” (sans dimension) de la puissance à la tension.

Le modèle peut être amélioré en tenant compte de la sensibilité de la charge à la fréquence f :

$$P = P_o \left(1 + D_p \frac{f - f_N}{f_N} \right) \left(\frac{V}{V_o} \right)^\alpha \quad (9.19)$$

$$Q = Q_o \left(1 + D_q \frac{f - f_N}{f_N} \right) \left(\frac{V}{V_o} \right)^\beta \quad (9.20)$$

où f_N est la fréquence nominale du système. La dépendance linéaire vis-à-vis de $f - f_N$ se justifie par la faible amplitude des variations de fréquence auxquelles les réseaux modernes sont sujets. On vérifie aisément que :

$$D_p = \left. \frac{\Delta P / P_o}{\Delta f / f_N} \right]_{V=V_o} \quad \text{et} \quad D_q = \left. \frac{\Delta Q / Q_o}{\Delta f / f_N} \right]_{V=V_o}$$

Le tableau ci-après présente des valeurs typiques du facteur de puissance (à la tension nominale) et des paramètres α , β , D_p , D_q pour divers types de charges.

composant	$\cos \phi$	α	β	D_p	D_q
conditionnement d'air triphasé central	0.90	0.09	2.5	0.98	-1.3
conditionnement d'air monophasé central	0.96	0.20	2.3	0.90	-2.7
conditionnement d'air en fenêtre	0.82	0.47	2.5	0.56	-2.8
chauffe-eau, cuisinière, four, surgélateur	1.00	2.0	0	0	0
lave-vaisselle	0.99	1.8	3.6	0	-1.4
lessiveuse	0.65	0.08	1.6	3.0	1.8
séchoir électrique	0.99	2.0	3.2	0	-2.5
réfrigérateur	0.8	0.77	2.5	0.53	-1.5
télévision	0.8	2.00	5.1	0	-4.5
lampe à incandescence	1.0	1.55	0	0	0
lampe fluorescente	0.9	0.96	7.4	1	-2.8
moteur industriel	0.88	0.07	0.5	2.5	1.2
moteur de ventilateur	0.87	0.08	1.6	2.9	1.7
pompe agricole	0.85	1.4	1.4	5.0	4.0
four à arc	0.70	2.3	1.6	-1.0	-1.0
transformateur à vide	0.64	3.4	11.5	0	-11.8

9.2.3 Modèle à exposant d'un agrégat de charges

La charge vue du jeu de barres alimentant un réseau de distribution à moyenne tension est un ensemble généralement complexe comprenant de très nombreuses charges de natures diverses et le réseau de distribution lui-même. Une telle charge est difficile à modéliser parce que :

- elle inclut un grand nombre de charges individuelles de natures diverses,
- auxquelles il faut ajouter l'effet du réseau de distribution lui-même ;
- la composition par type de charge n'est pas toujours connue avec précision
- cette composition varie selon l'heure de la journée, selon la saison, etc. . . Par exemple, lorsque l'on effectue une étude à la pointe de consommation, la charge pourra comporter, selon le pays, une grande proportion de chauffage électrique (pointe d'hiver) ou une grande proportion de moteurs provenant de systèmes de conditionnement d'air (pointe d'été) ;
- même si l'on connaissait bien cette composition, il resterait à établir un modèle suffisamment simple de cet ensemble parfois hétérogène.

Il est très courant de recourir également au modèle à exposant pour modéliser les agrégats de charges vus des départs de distribution. Le tableau ci-après donne des exemples de valeurs de α , β , D_p , D_q pour des charges homogènes (c'est-à-dire que les charges individuelles qui la composent appartiennent à une même catégorie de consommateurs).

catégorie de charge	$\cos \phi$	α	β	D_p	D_q
résidentielle, en été	0.9	1.2	2.9	0.8	-2.2
résidentielle, en hiver	0.99	1.5	3.2	1.0	-1.5
commerciale, en été	0.85	1.0	3.5	1.2	-1.6
commerciale, en hiver	0.9	1.3	3.1	1.5	-1.1
industrielle	0.85	0.2	6.0	2.6	1.6
auxiliaires de centrales	0.8	0.1	1.6	2.9	1.8

En pratique, dans les études faisant intervenir un tel modèle de charge, on effectue un calcul de load flow préliminaire pour déterminer le point de fonctionnement initial du système. P_o et Q_o sont alors les puissances spécifiées au noeud PQ où est connectée la charge et V_o la tension de celui-ci, fournie par le calcul de load flow.

On peut tenir compte d'une composition non homogène de la charge en combinant différents modèles, avec une pondération pour chaque type :

$$P = P_o \left(1 + D_p \frac{f - f_N}{f_N} \right) \sum_i a_i \left(\frac{V}{V_o} \right)^{\alpha_i} \quad \text{avec} \quad \sum_i a_i = 1 \quad (9.21)$$

$$Q = Q_o \left(1 + D_q \frac{f - f_N}{f_N} \right) \sum_i b_i \left(\frac{V}{V_o} \right)^{\beta_i} \quad \text{avec} \quad \sum_i b_i = 1 \quad (9.22)$$

où a_i (resp. b_i) est la proportion de la puissance active (resp. réactive) totale consommée par la composante de caractéristique α_i (resp. β_i) lorsque $V = V_o$ et $f = f_N$.

Notons enfin que certains modèles valables aux environs de la tension nominale cessent d'être applicables en cas de déviations importantes et/ou prolongées de la tension. Parmi les phénomènes responsables de ceci, nous avons déjà mentionné le décrochage des moteurs asynchrones. Il y a également l'extinction rapide des lampes fluorescentes lorsque la tension tombe en dessous d'environ 0.7 pu.

9.2.4 Identification des paramètres du modèle

Dans de nombreux réseaux, on cherche actuellement à mieux connaître le comportement des charges et en particulier à identifier les paramètres intervenant dans leurs modèles. Il existe deux approches :

1. la première tente d'identifier la composition de la charge, par exemple à partir d'informations collectées par les services de facturation et consignées dans les systèmes informatiques de gestion du réseau de distribution. A chaque type de charge individuelle est associé un modèle typique et les diverses composantes sont combinées en un modèle unique ;
2. la seconde tente de mesurer sur site la réponse de la charge à des variations de tension ou de fréquence. Evidemment, ces essais sont limités à des variations ne perturbant pas les consommateurs. Des variations permanentes de fréquence sont difficiles à réaliser.

Dans le cas du modèle à exposant, les coefficients α et β peuvent être déterminés en mesurant les variations de puissance ΔP et ΔQ résultant d'une variation de tension ΔV et en introduisant ces données dans (9.17, 9.18).

La variation de tension peut être provoquée en agissant sur le(s) transformateur(s) alimentant le départ de distribution. Ce type de transformateur est souvent équipé d'un régleur en charge permettant d'ajuster le rapport de transformation comme décrit à la section 6.5. La variation de tension peut être obtenue :

- en passant rapidement des prises sur le(s) transformateur(s)
- si l'on dispose de deux transformateurs en parallèle, en déclenchant l'un d'entre eux, à condition qu'il n'y ait pas de risque de surcharge du transformateur restant
- si l'on dispose de deux transformateurs en parallèle, en augmentant momentanément le rapport de l'un et en diminuant le rapport de l'autre avant d'en déclencher un.

En ce qui concerne la réponse du système à des variations importantes de tension ou de fréquence, l'analyse d'enregistrements pris au cours d'incidents sévères est du plus haut intérêt.

En tout état de cause, face à l'incertitude sur le comportement de la charge, il importe de procéder à des études de sensibilité dans lesquelles on fait varier les paramètres des modèles dans les plages de valeurs plausibles.

Chapitre 10

Régulation de la fréquence

Dans tout système électrique de puissance, il importe de maintenir la fréquence dans une plage étroite autour de sa valeur nominale (50 ou 60 Hz). Le respect strict de cette valeur est non seulement nécessaire au fonctionnement correct des charges mais, comme on va le voir, il est également l'indicateur d'un équilibre entre puissances actives produites et consommées.

Le maintien de cet équilibre est essentiel car l'énergie électrique n'est pas emmagasinable, du moins pas dans les quantités suffisantes pour faire face aux fluctuations de la demande ou aux incidents. Elle doit donc être produite au moment où elle est demandée.

Considérons par exemple une augmentation brutale de la demande. Dans les toutes premières secondes, l'énergie correspondante va être prélevée sur l'énergie cinétique stockée dans les masses tournantes des machines synchrones dans les centrales. Ceci va entraîner une diminution de la vitesse de rotation de ces unités, c'est-à-dire de la fréquence du réseau. Cet écart de vitesse est détecté et corrigé automatiquement par les *régulateurs de vitesse*. Dans l'exemple qui nous occupe, ces régulateurs vont augmenter l'admission de fluide (vapeur, gaz ou eau) dans les turbines de manière à ramener les vitesses autour de leurs valeurs nominales, et donc la fréquence du réseau. Une fois le système revenu à l'équilibre, les unités conservent cette admission de fluide plus élevée, donc une production de puissance plus élevée, équilibrant la demande également plus élevée.

Cette régulation en centrale est appelée *régulation primaire*¹. Elle intervient la première sur l'échelle des temps : typiquement, quelques secondes après une perturbation. Comme nous le verrons dans ce chapitre, il existe également une *régulation secondaire*, intervenant typiquement en quelques minutes.

1. dans ce chapitre, "régulation" et "réglage" sont utilisés indistinctement

10.1 Régulateur de vitesse

10.1.1 Description et schéma bloc

Un schéma de principe de la régulation de vitesse est donné à la figure 10.1. Un dispositif mesure la vitesse de rotation de l'ensemble (turbine + générateur). Dans le régulateur, cette mesure est comparée à la vitesse de rotation nominale (correspondant à une fréquence de 50 ou 60 Hz) et l'écart entraîne une augmentation ou une diminution de l'admission de fluide dans la turbine, par action sur ses *soupapes de réglage*. Le régulateur de vitesse utilise un *servomoteur* pour manoeuvrer les soupapes. Ce dispositif (voir figure 10.1) comporte un piston, mû par de l'huile sous pression. Celle-ci est admise dans le cylindre, d'un côté ou de l'autre du piston selon le sens de la correction, par une *vanne pilote*. Le piston se déplace tant que la vanne pilote ne bloque pas l'admission d'huile, c'est-à-dire tant que le signal de correction qui commande la vanne pilote n'est pas revenu à zéro.

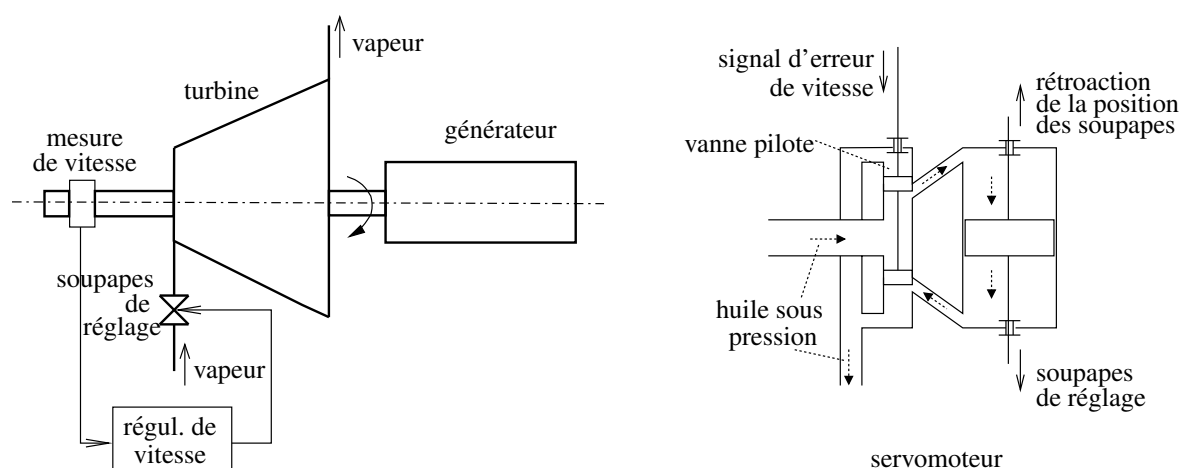


FIGURE 10.1 – schéma de principe de la régulation de vitesse et du servomoteur

On peut établir un schéma bloc comme à la figure 10.2. Le signal d'entrée ω_m est la vitesse de rotation. z représente la fraction d'ouverture des soupapes de la turbine ($0 \leq z \leq 1$). $G(s)$ est la fonction de transfert entre z et la puissance mécanique P_m délivrée par la turbine :

$$P_m = G(s) z \quad (10.1)$$

$F(s)$ est la fonction de transfert entre ω_m et z . Nous allons la détailler quelque peu.

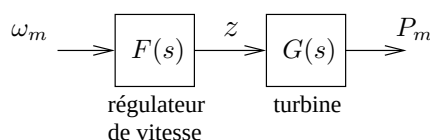


FIGURE 10.2 – schéma bloc du régulateur de vitesse et de la turbine

Un premier type de régulateur de vitesse est décrit par le schéma bloc de la figure 10.3. Il s'agit d'un schéma simplifié ; en particulier, on n'a pas représenté les limites imposées à z et à sa dérivée.

p est le nombre de paires de pôles du générateur ; $p\omega_m$ est donc la vitesse électrique. En régime établi, celle-ci est égale à la pulsation $\omega = 2\pi f$ du système. Le servomoteur est représenté par un gain et un intégrateur. Lorsque le système est en régime établi, l'entrée de l'intégrateur est nécessairement nulle. Le régulateur de vitesse ajuste donc les soupapes de réglage jusqu'à ce que l'erreur de vitesse soit totalement annulée. Un tel régulateur est dit *isochrone*.

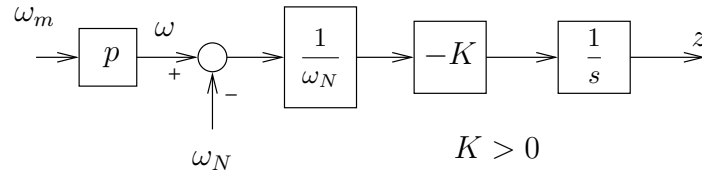


FIGURE 10.3 – schéma bloc d'un régulateur de vitesse isochrone

Dans un réseau comportant plusieurs générateurs, un seul d'entr'eux peut être doté d'un régulateur isochrone. En effet, d'inévitables petites différences entre les consignes de deux régulateurs isochrones les conduiraient à "se disputer" la correction des erreurs de fréquence. Cependant, il n'est pas pensable de faire fonctionner un réseau de grande taille avec un seul générateur doté d'un régulateur isochrone car ce générateur devrait à lui seul assurer l'équilibre entre production et consommation de tout le système.

Pour répartir l'effort sur un certain nombre de générateurs on a recours à un autre type de régulateur, dont le schéma bloc est représenté à la figure 10.4. Ce dernier diffère du régulateur isochrone par la présence d'une rétroaction de la position de la soupape de réglage z . De plus, il fait intervenir z^o la consigne d'ouverture des soupapes, que l'on peut ajuster pour modifier la production de l'unité. Comme on le verra dans les sections suivantes, le paramètre σ joue un rôle important dans la participation de l'unité à l'équilibre production-consommation.

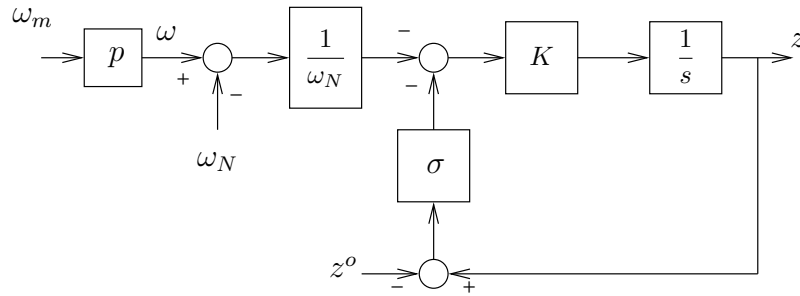


FIGURE 10.4 – schéma bloc d'un régulateur de vitesse avec statisme

Quelques manipulations permettent de passer du schéma de la figure 10.4 à celui de la figure 10.5, plus utilisé en pratique. T_{sm} est relié aux paramètres de la figure 10.4 par :

$$T_{sm} = \frac{1}{K\sigma}$$

On tire aisément :

$$z = \frac{1}{1 + sT_{sm}} \left(z^o - \frac{\omega - \omega_N}{\sigma \omega_N} \right) \quad (10.2)$$

qui fait apparaître T_{sm} comme une constante de temps, relative au servomoteur.

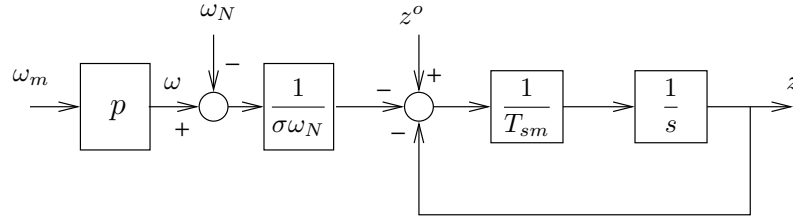


FIGURE 10.5 – schéma bloc d'un régulateur de vitesse avec statisme (seconde version)

10.1.2 Caractéristique statique d'un ensemble turbine-régulateur

La caractéristique statique d'un ensemble turbine - régulateur de vitesse est la relation, en régime établi, entre la puissance mécanique fournie par la turbine et la fréquence du réseau.

En posant $s = 0$ dans (10.1) et (10.2), on trouve aisément :

$$P_m = G(0) \left(z^o - \frac{\omega - \omega_N}{\sigma \omega_N} \right) \quad (10.3)$$

Considérons qu'en régime établi, lorsque les soupapes de réglage sont ouvertes au maximum ($z = 1$), la turbine délivre sa puissance nominale P_N . La relation (10.1) seule donne :

$$P_N = G(0)$$

et en introduisant cette relation dans (10.3) on obtient :

$$P_m = P_N z^o - \frac{P_N}{\sigma} \frac{\omega - \omega_N}{\omega_N} = P^o - \frac{P_N}{\sigma} \frac{f - f_N}{f_N} \quad (10.4)$$

où P^o est la consigne de production de la turbine. Le diagramme f - P_m correspondant se présente donc sous la forme d'une droite inclinée, comme montré à la figure 10.6. Les limites de puissance sont évidemment 0 et P_N , mais la puissance minimale peut être plus élevée (voir ligne en pointillé) pour des raisons de stabilité de combustion dans les unités thermiques ou de vibrations dans les unités hydrauliques.

Le rapport de la variation relative de fréquence à la variation relative de puissance est donné par :

$$\left| \frac{\Delta f / f_N}{\Delta P_m / P_N} \right| = \left| \frac{\Delta \omega / \omega_N}{\Delta P_m / P_N} \right| = \sigma$$

σ est appelé le *statisme* du régulateur de vitesse. En Europe, il vaut typiquement 4 % pour les unités thermiques². Comme le montre la figure 10.6, une variation de fréquence $\Delta f = \sigma f_N =$

2. 5 % en Amérique du nord

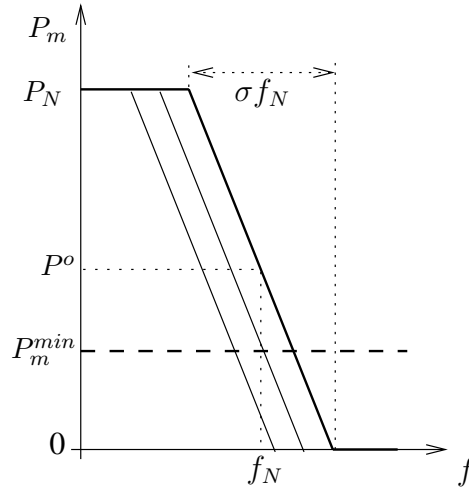


FIGURE 10.6 – caractéristique statique (idéale) de l’ensemble turbine-régulateur

$0.04 \times 50 = 2$ Hz provoquerait une variation de puissance mécanique ΔP_m égale à la puissance nominale P_N . Un statisme infini correspond à une machine fonctionnant à puissance constante et ne participant donc pas à la régulation de la fréquence.

Lorsque l’on ajuste la consigne de puissance P^o , la caractéristique se translate comme représenté à la figure 10.6.

La pente de la caractéristique statique indique clairement que les régulateurs de vitesse sont du type proportionnel. Ils laissent donc une erreur statique sur la fréquence. Comme nous allons le voir à la section suivante, cette propriété permet précisément le partage de l’effort par les différents générateurs interconnectés.

10.2 Régulation primaire

10.2.1 Hypothèses de modélisation

Déterminons à présent les variations de fréquence et de production résultant d’une perturbation du bilan de puissance active. Nous supposons que les transitoires qui suivent cette perturbation sont éteints, auquel cas toutes les machines tournent à la même vitesse électrique et la fréquence est la même dans tout le réseau.

Nous supposons pour simplifier que le réseau est sans pertes et que la puissance de chaque turbine est intégralement transformée en puissance électrique.

Par contre, nous considérons la sensibilité de la charge à la fréquence, en écrivant que la puissance active totale consommée par les charges vaut :

$$P_c = P_c^o p(f) \quad (10.5)$$

où $p(f)$ traduit la dépendance vis-à-vis de la fréquence f . A la fréquence nominale, on a $p(f_N) = 1$ et $P_c = P_c^o$. Dans ce qui suit, nous nous limiterons à de faibles variations de f autour de f_N , ce qui autorise la linéarisation :

$$p(f) \simeq p(f_N) + \left. \frac{dp}{df} \right|_{f=f_N} (f - f_N) = 1 + D(f - f_N)$$

où D est le coefficient de sensibilité de la charge à la fréquence (pu/Hz). L'expression de la puissance consommée par la charge devient donc :

$$P_c = P_c^o (1 + D(f - f_N)) \quad (10.6)$$

10.2.2 Partage des variations de production entre générateurs

La fréquence étant la même partout, les caractéristiques des divers générateurs peuvent être combinées en une caractéristique globale :

$$P_m = \sum_{i=1}^n P_{mi} = \sum_{i=1}^n P_i^o - \frac{f - f_N}{f_N} \sum_{i=1}^n \frac{P_{Ni}}{\sigma_i} \quad (10.7)$$

où n est le nombre de générateurs en service. Sous les hypothèses simplificatrices mentionnées, le bilan de puissance du système s'écrit simplement :

$$\sum_{i=1}^n P_i^o - \frac{f - f_N}{f_N} \sum_{i=1}^n \frac{P_{Ni}}{\sigma_i} = P_c^o (1 + D(f - f_N)) \quad (10.8)$$

Supposons pour simplifier (mais sans perdre en généralité) que le réseau fonctionne initialement à la fréquence nominale f_N . En ce point de fonctionnement, l'équation (10.8) devient simplement :

$$\sum_{i=1}^n P_i^o = P_c^o \quad (10.9)$$

Supposons que survient une perturbation, sous la forme d'une augmentation ΔP_c de la consommation (ou de la perte d'une unité de produisant cette même puissance).

Si l'on suppose que les consignes de production sont inchangées, ce qui est bien le cas après réglage primaire, le bilan de puissance (10.8) devient :

$$\sum_{i=1}^n P_i^o - \frac{f - f_N}{f_N} \sum_{i=1}^n \frac{P_{Ni}}{\sigma_i} = P_c^o (1 + D(f - f_N)) + \Delta P_c \quad (10.10)$$

ou encore, compte-tenu de (10.9) :

$$- \frac{f - f_N}{f_N} \sum_{i=1}^n \frac{P_{Ni}}{\sigma_i} = P_c^o D(f - f_N) + \Delta P_c \quad (10.11)$$

En posant :

$$\beta = DP_c^o + \frac{1}{f_N} \sum_{i=1}^n \frac{P_{Ni}}{\sigma_i} \quad (10.12)$$

et en notant :

$$\Delta f = f - f_N$$

la relation (10.9) s'écrit simplement :

$$-\beta \Delta f = \Delta P_c \quad (10.13)$$

Le paramètre β est appelé *énergie réglante* du système. On notera qu'il a en effet la dimension d'une énergie. Il caractérise la précision de la régulation primaire de fréquence dans le système considéré : pour une même variation de charge, la variation de fréquence est d'autant plus faible que l'énergie réglante est élevée.

On en déduit la variation de fréquence :

$$\Delta f = -\frac{\Delta P_c}{\beta} \quad (10.14)$$

et donc la variation de puissance du j-ème générateur :

$$\Delta P_{mj} = -\frac{\Delta f}{f_N} \frac{P_{Nj}}{\sigma_j} = \frac{\Delta P_c}{f_N} \frac{P_{Nj}}{\beta \sigma_j} \quad (10.15)$$

De ces relations on peut tirer les conclusions suivantes :

- l'existence d'une erreur statique de fréquence permet un partage *prévisible* et *réglable* de la variation de production par les différentes machines ;
- tous les statismes étant fixés, un générateur participe d'autant plus que sa puissance nominale est élevée ;
- toutes les puissances nominales étant fixées, un générateur participe d'autant plus que son statisme est petit ;
- la variation de fréquence est d'autant plus petite que le nombre de générateurs participant à la régulation est élevé.

10.2.3 Réserve primaire

A un instant donné, seule une fraction du nombre total de générateurs participe à la régulation primaire (cf exercice). Pour participer, un producteur doit ménager une réserve, c'est-à-dire placer sa production en dessous de la capacité maximale de la turbine, de façon à être prêt à produire davantage. Ceci est une barrière à l'utilisation de sources d'énergie renouvelable pour le réglage primaire de la fréquence (étant donné que l'on souhaite produire un maximum d'énergie de cette nature).

La participation à la réserve primaire est un service que le producteur offre sur le marché correspondant. S'il est retenu, il est rétribué pour la mise à disposition de sa réserve (même si elle n'est

pas activée) ce qui compense ses pertes de revenus liées à la diminution de la production. Il est également payé en cas d'activation de cette réserve, avec un prix différent pour les augmentations de production (ceci compensant ses coûts de production plus élevés) et pour les diminutions (ceci compensant sa perte de revenus).

10.3 Régulation secondaire

10.3.1 Interconnexion des réseaux

Si l'on excepte les systèmes insulaires autonomes et quelques autres situations³, la plupart des réseaux gérés par des gestionnaires distincts sont regroupés au sein de grandes interconnexions.

Les avantages de celles-ci sont :

- une meilleure régulation de la fréquence, comme montré à la section précédente
- une assistance mutuelle en cas d'incident
- et donc une réduction de la réserve primaire que chaque gestionnaire de réseau doit contracter auprès de producteurs, pour faire face à la perte de générateurs
- la possibilité de vendre et d'acheter de l'énergie à des conditions plus avantageuses
- une meilleure tenue de la tension à l'interface entre systèmes, après interconnexion.

Notons cependant au passage que la constitution de grands systèmes interconnectés n'est pas exempte d'inconvénients :

- des incidents peuvent se propager d'un réseau à un autre, via les lignes d'interconnexion
- des flux de puissance peuvent traverser un réseau situé au sein d'une structure maillée, suite à des modifications topologiques ou des injections inattendues (p.ex. variabilité de la production éolienne) dans les réseaux voisins
- des oscillations électromécaniques lentes (0.1 à 0.5 Hz) et mal amorties peuvent apparaître.

Ces inconvénients peuvent conduire à s'interconnecter via des liaisons à courant continu. Les réseaux ainsi connectés conservent chacun leur fréquence. Dans ce chapitre, toutefois, nous nous intéressons à une interconnexion à courant alternatif dans laquelle la fréquence est unique en régime établi.

10.3.2 Régulation primaire d'une interconnexion : exemple à deux réseaux

Considérons pour simplifier deux réseaux, repérés respectivement 1 et 2, interconnectés via un certain nombre de lignes de transport (cf. figure 10.7). Comme à la section précédente, nous supposons pour simplifier qu'il n'y a pas de pertes dans le réseau, en particulier dans les lignes d'intercon-

3. p.ex. Texas et Québec qui ne sont pas connectés en courant alternatif avec le reste de l'Amérique du Nord

nexion. Soit P_{12} la puissance active totale transitant du réseau 1 vers le réseau 2 via les lignes d'interconnexion.

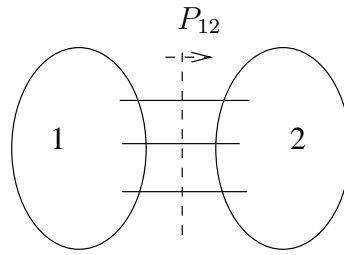


FIGURE 10.7 – interconnexion de deux réseaux

En adaptant les résultats de la section précédente, on établit aisément :

- la caractéristique combinée des générateurs du réseau 1 :

$$P_{m1} = \sum_{i \in 1} P_{mi} = \sum_{i \in 1} P_i^o - \frac{f - f_N}{f_N} \sum_{i \in 1} \frac{P_{Ni}}{\sigma_i}$$

- la caractéristique de la charge du réseau 1 :

$$P_{c1} = P_{c1}^o + D_1 P_{c1}^o (f - f_N)$$

- l'équilibre des puissances dans le réseau 1 :

$$P_{m1} = P_{c1} + P_{12}$$

- la caractéristique combinée des générateurs du réseau 2 :

$$P_{m2} = \sum_{i \in 2} P_{mi} = \sum_{i \in 2} P_i^o - \frac{f - f_N}{f_N} \sum_{i \in 2} \frac{P_{Ni}}{\sigma_i}$$

- la caractéristique de la charge du réseau 2 :

$$P_{c2} = P_{c2}^o + D_2 P_{c2}^o (f - f_N)$$

- l'équilibre des puissances dans le réseau 2 :

$$P_{m2} = P_{c2} + P_{21} = P_{c2} - P_{12}$$

- l'équilibre des puissances dans le système complet :

$$P_{m1} + P_{m2} = P_{c1} + P_{c2}$$

Ces différentes caractéristiques sont représentées à la figure 10.8.

Comme à la section 10.2.2, on suppose que le système fonctionne initialement à la fréquence f_N et on considère l'effet d'une augmentation de la charge. Celle-ci est supposée se produire dans le réseau 1 et est notée ΔP_{c1} .

En appliquant les relations du réglage primaire et en traitant P_{12} comme une charge (resp. une production) vis-à-vis du réseau 1 (resp. réseau 2), on obtient aisément :

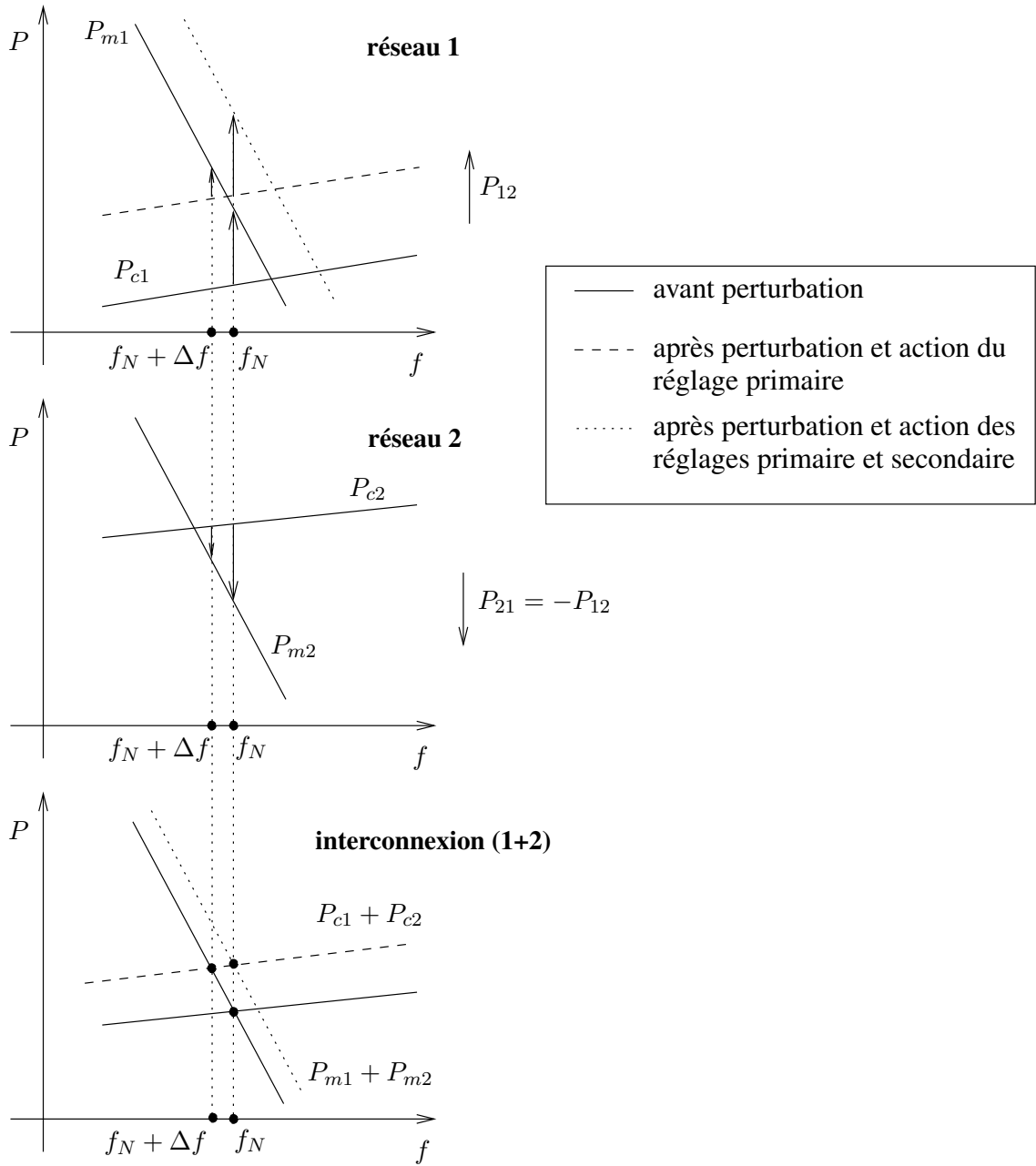


FIGURE 10.8 – caractéristiques des générateurs et des charges dans les deux réseaux interconnectés

- dans le réseau 1 :

$$-\beta_1 \Delta f = \Delta P_{c1} + \Delta P_{12}$$

- dans le réseau 2 :

$$-\beta_2 \Delta f = -\Delta P_{12} \quad (10.16)$$

où β_1 et β_2 sont les énergies réglantes des réseaux 1 et 2, respectivement.

De ces relations, on tire :

- la variation de fréquence :

$$\Delta f = -\frac{\Delta P_{c1}}{\beta_1 + \beta_2}$$

- la variation de l'échange de puissance :

$$\Delta P_{12} = -\frac{\beta_2}{\beta_1 + \beta_2} \Delta P_{c1}$$

Ces variations sont montrées à la figure 10.8. Sous l'effet de l'augmentation de consommation, il y a un déplacement de la caractéristique de la charge dans le réseau 1 et un déplacement correspondant à l'échelle de l'interconnexion. Le diagramme relatif à celle-ci permet de trouver la nouvelle fréquence $f_N + \Delta f$ du système, à l'intersection des courbes de charge et de production. Connaissant cette nouvelle fréquence, on peut remonter dans les diagrammes des réseaux individuels et déterminer la nouvelle puissance échangée.

On voit que le transit de puissance de 1 vers 2 diminue suite à la perturbation. En effet, les machines du réseau 2 participant à la régulation primaire contribuent à équilibrer l'augmentation de la charge dans le réseau 1 : c'est l'assistance mutuelle déjà mentionnée. Il en résulte un flux de puissance de 2 vers 1, qui diminue le transit existant avant perturbation. Cette diminution est d'autant plus forte que β_2 est élevé par rapport à β_1 , ce qui est normalement le cas si le réseau 2 est grand par rapport au réseau 1.

10.3.3 Objectifs de la régulation secondaire

Les deux premiers objectifs de la régulation secondaire sont : (i) d'éliminer l'erreur de fréquence sur laquelle repose le réglage primaire et (ii) de ramener les échanges de puissance entre réseaux interconnectés aux valeurs désirées (spécifiées dans les contrats d'achat d'énergie).

Ce ne sont généralement pas les mêmes générateurs qui participent au réglage primaire et au réglage secondaire. Suite à une perturbation conduisant à un déficit de production, une partie de la réserve primaire est utilisée pour compenser le déficit. Les générateurs participant à la réserve secondaire vont alors prendre la relève de ceux impliqués dans la régulation primaire, dont la réserve va ainsi être reconstituée, afin de faire face à une autre perturbation. Cette reconstitution de la réserve primaire est le troisième objectif de la régulation secondaire.

10.3.4 Principe de la régulation secondaire

Revenons à l'exemple de la figure 10.8. Pour ramener la fréquence f et l'échange de puissance P_{12} à leurs valeurs avant perturbation, il est clair qu'il faut augmenter la production des générateurs du réseau 1. A cette fin, le réglage secondaire va modifier les consignes de production P_i^o de certains générateurs connectés à ce réseau. Comme on l'a vu à la figure 10.6, ceci a pour effet de translater les caractéristiques des générateurs de la région 1. Quand la fréquence revient à la valeur f_N , le transit P_{12} revient à sa valeur avant perturbation. Notons que les consignes des générateurs du réseau 2 ne doivent pas être ajustées ; les productions de ces générateurs se modifient suite à la modification de la fréquence induite par les ajustements dans le réseau 1. Les puissances qu'ils ont apportées momentanément grâce au réglage primaire sont "effacées" par le réglage secondaire.

10.3.5 Mise en oeuvre de la régulation secondaire

La régulation secondaire s'effectue au sein d'une *zone de réglage*. Celle-ci peut coïncider avec un pays ou avec le réseau géré par une compagnie, voire plusieurs compagnies. Dans chaque zone de réglage, on mesure la fréquence ainsi que la *somme* des transits dans les lignes connectant cette zone au reste du système, l'objectif étant de ramener cette somme à une valeur de consigne. En pratique, cette régulation est assurée par un logiciel exécuté dans un centre de conduite recevant les mesures requises à intervalle régulier (de l'ordre de quelques secondes).

Reprenons à notre exemple à deux zones de réglage. Dans chacune, on calcule l'*erreur de réglage de zone*⁴, soit pour la zone 1 :

$$E_1 = P_{12} - P_{12}^0 + \lambda_1(f - f_N) = \Delta P_{12} + \lambda_1 \Delta f$$

et pour la zone 2 :

$$E_2 = P_{21} - P_{21}^0 + \lambda_2(f - f_N) = -\Delta P_{12} + \lambda_2 \Delta f$$

Chacun de ces signaux est traité, dans sa zone, par un régulateur proportionnel-intégral pour obtenir la correction de production dans la zone 1 :

$$\Delta P_1^o = -K_{i1} \int E_1 dt - K_{p1} E_1$$

et dans la zone 2 :

$$\Delta P_2^o = -K_{i2} \int E_2 dt - K_{p2} E_2$$

Les constantes K sont toutes positives.

Ayant défini un certain nombre de générateurs participant au réglage secondaire, on distribue le signal ΔP_1^o (resp. ΔP_2^o) sur ces derniers. La consigne du i -ème générateur de la zone 1 devient donc :

$$P_i^o + \rho_i \Delta P_1^o \quad \text{avec} \quad \sum_i \rho_i = 1$$

et de même dans la zone 2.

A l'issue du réglage, si l'on n'a pas rencontré de limites, le terme intégral impose :

$$\begin{aligned} E_1 = 0 &\Rightarrow \Delta P_{12} + \lambda_1 \Delta f = 0 \\ E_2 = 0 &\Rightarrow -\Delta P_{12} + \lambda_2 \Delta f = 0 \end{aligned}$$

dont la seule solution est :

$$\Delta f = 0 \quad \text{et} \quad \Delta P_{12} = 0 \quad (10.17)$$

qui correspond bien aux deux buts recherchés pour le réglage.

Il faut noter que ce réglage est *entièrement distribué* ; il ne nécessite pas de centraliser les mesures en un point unique, ni d'échanger les mesures avec les autres zones de l'interconnexion.

4. en anglais, *area control error*

Choix des paramètres λ_i

Du point de vue de l'erreur finale, quelles que soient les valeurs de λ_1 et λ_2 , on aboutit à (10.17) après réglage secondaire. Toutefois, le choix de ces paramètres influence la dynamique de la régulation. De ce point de vue, il est judicieux de prendre :

$$\lambda_1 = \beta_1 \quad \text{et} \quad \lambda_2 = \beta_2 \quad (10.18)$$

où β_1 et β_2 sont les énergies réglantes définies antérieurement. En effet, dans ce cas, l'erreur de réglage dans la zone 2 devient :

$$E_2 = -\Delta P_{12} + \beta_2 \Delta f$$

valeur qui est nulle, en vertu de (10.16). En d'autres termes, *les générateurs de la zone 2 ne réagissent pas*, ce qui est souhaitable puisque, comme expliqué plus haut, seuls les générateurs de la zone 1 doivent être ajustés. Plus λ_2 s'écarte de β_2 , plus les générateurs de la zone 2 réagissent *momentanément et inutilement* pendant que la régulation secondaire se met en oeuvre.

En pratique, il n'est pas possible de réaliser exactement la condition (10.18) car les énergies réglantes changent avec la charge du réseau et il peut y avoir des imprécisions au niveau des participations des unités au réglage primaire. On s'efforce toutefois de s'en rapprocher le plus possible.

Choix des paramètres K_i , K_p et ρ_i

Les coefficients K_i et K_p sont choisis de manière à optimiser la dynamique du système. En particulier le réglage secondaire ne doit pas être trop "énergique" pour ne pas interférer avec le réglage primaire, ce qui pourrait générer des oscillations indésirables, voire à la limite une instabilité.

Dans certains réseaux, on ne considère que le terme intégral ($K_p = 0$).

Les coefficients ρ_i distribuent le signal de correction sur un certain nombre de générateurs. Ces générateurs doivent évidemment disposer de réserves de production, dont la somme constitue la *réserve secondaire*.

Tant pour le réglage primaire que pour le réglage secondaire, la variation de production qu'une unité participante s'engage à fournir, dans un temps donné, et donc en particulier son coefficient ρ_i , doit tenir compte du taux de variation maximal permis par la turbine. Ce taux est de l'ordre de :

- quelques pour-cents de la puissance nominale par minute pour une unité thermique ;
- la puissance nominale par minute pour une unité hydraulique.

10.3.6 Extension à plus de deux zones de réglage

Les développements qui précèdent s'appliquent bien entendu à un ensemble quelconque de zones de réglage. L'erreur de réglage de zone fait intervenir la *somme* des puissances que la zone désire

échanger avec toutes les autres, chacune intervenant avec son signe.

Considérons par exemple le cas à trois zones de réglage représenté à la figure 10.9. Supposons que la zone 1 veuille vendre 1000 MW à la zone 3 et que la zone 2 ne désire rien acheter ni vendre. Les consignes du réglage secondaire sont mises à 1000, 0 et -1000 MW respectivement. A l'issue de ce réglage on a :

$$\begin{aligned} E_1 = 0 &\Rightarrow (P_{12} + P_{13}) - 1000 + \lambda_1 \Delta f = 0 \\ E_2 = 0 &\Rightarrow (-P_{12} + P_{23}) - 0 + \lambda_2 \Delta f = 0 \\ E_3 = 0 &\Rightarrow (-P_{13} - P_{23}) + 1000 + \lambda_3 \Delta f = 0 \end{aligned}$$

dont on tire aisément :

$$\begin{aligned} \Delta f &= 0 \\ P_{12} + P_{13} &= 1000 \\ P_{12} &= P_{23} \\ P_{13} + P_{23} &= 1000 \end{aligned}$$

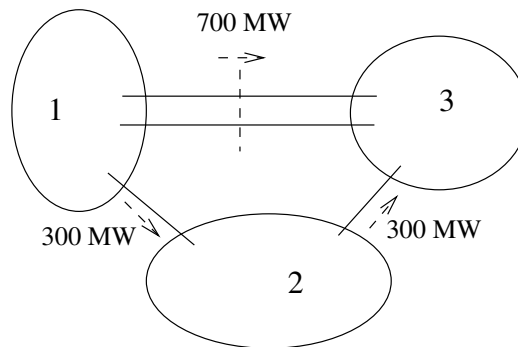


FIGURE 10.9 – exemple à trois zones de réglage

Notons que les transits de puissance individuels ne sont pas contrôlables. Comme illustré à la figure 10.9, une partie du flux de puissance du réseau 1 vers le réseau 3 passe par réseau 2, les électrons n'obéissant qu'aux lois de l'électricité ! Ce transit entraîne des pertes pour le réseau 2, ainsi qu'une mise en charge de ses lignes pouvant diminuer ses marges de sécurité.

Pour se dédommager des pertes, le gestionnaire doit faire payer l'usage de son réseau. Pour faire face à des situations dangereuses, il peut installer un ou plusieurs transformateurs déphaseurs destinés à réduire la puissance active qui le traverse, c'est-à-dire dans l'exemple ci-dessus forcer tout ou partie des 1000 MW à passer par les lignes qui connectent directement les réseaux 1 et 3⁵. Il faut évidemment étudier les emplacements optimaux et l'efficacité de tels déphaseurs en considérant un modèle détaillé du réseau. Dans l'éventualité où plusieurs gestionnaires installeraient de tels dispositifs, se profile le problème de l'interaction entre ces différents dispositifs !

5. ce qui n'est bien entendu possible que parce qu'une liaison directe existe entre 1 et 3

Chapitre 11

Régulation de la tension

Il existe deux différences fondamentales entre la régulation de la fréquence et celle de la tension :

- la fréquence est un “signal” commun à tous les composants d’un même réseau. Aussi grand soit ce dernier, en régime établi, la fréquence est la même partout. Lorsqu’on augmente la production d’une quelconque des centrales, la fréquence est ajustée par les régulateurs de vitesse à une valeur un peu supérieure.

Il n’y a pas de signal ni de comportement équivalent pour la régulation de tension. Les réglages ont une portée locale : lorsque l’on ajuste la tension en un noeud d’un réseau, cela influence la tension des noeuds situés dans un certain voisinage. Au delà, les effets sont négligeables ;

- la fréquence est usuellement tenue près de sa valeur nominale avec grande précision, parce que tout écart de fréquence est révélateur d’un déséquilibre entre puissances actives produite et consommée.

En comparaison, la régulation de la tension est moins précise. Dans un réseau de transport, on admet couramment un écart de $\pm 5\%$ par rapport à la valeur nominale. En fait, on ne peut empêcher de tels écarts, qui proviennent des chutes de tension créées par le passage du courant dans les impédances du réseau.

Il convient toutefois de maintenir la tension dans des limites acceptables :

- elle ne doit pas être trop élevée sous peine d’endommager les isolants, les appareils sensibles, etc. . .
- elle ne peut pas être trop basse, sous peine de perturber, voire interrompre le fonctionnement de certains composants : mise hors service des charges se protégeant contre les sous-tensions, blocage de l’électronique de puissance dans les redresseurs et onduleurs, décrochage des moteurs asynchrones, etc. . .

Les deux manières les plus usuelles d’ajuster les tensions d’un réseau sont :

- produire ou consommer de la puissance réactive en ses noeuds
- ajuster le nombre de spires des transformateurs qui permettent de passer d’un niveau de tension à un autre (dans certains cas, des transformateurs dédiés au réglage de la tension).

11.1 Contrôle de la tension par condensateur ou inductance shunt

Le moyen le plus économique de corriger une chute de tension en un jeu de barres est d'y connecter des bancs de condensateurs shunt, afin d'y produire de la puissance réactive. De même, les augmentations de tension peuvent être corrigées en connectant des selfs shunt¹, afin d'y consommer de la puissance réactive.

Le principe de ce contrôle est illustré à la figure 11.1. On suppose que, sous l'effet d'une perturbation, la caractéristique QV du réseau passe de la droite 1 à la droite 2. En l'absence de compensation shunt au jeu de barres considéré, le point de fonctionnement passe de A en B sous l'effet de la perturbation. La caractéristique QV de la compensation shunt est simplement :

$$Q = B V^2$$

avec $B > 0$ pour un condensateur et $B < 0$ pour une self, soit une parabole dans le plan (V, Q) . Après connexion du condensateur, le nouveau point de fonctionnement est C. La figure 11.1 montre la correction (partielle) de la tension ainsi obtenue. La même figure montre la correction, au moyen d'une inductance shunt, d'une augmentation de tension due au passage de la caractéristique 1 à la caractéristique 3.

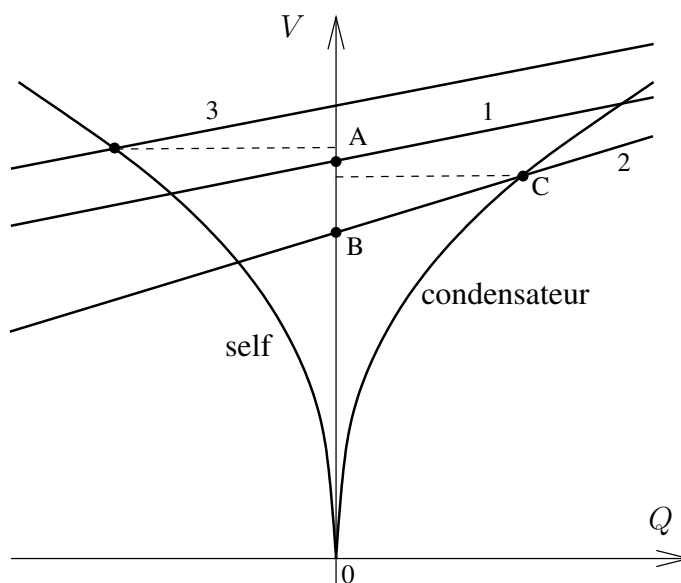


FIGURE 11.1 – caractéristiques QV d'un réseau et de la compensation shunt

Ces enclenchements peuvent être commandés manuellement ou automatiquement. Dans le premier cas, il est souvent réalisé, depuis un centre de conduite, par un opérateur disposant d'une télémessure de la tension. Dans le second cas, un dispositif situé dans le poste enclenche le condensateur automatiquement lorsque la tension passe sous un seuil bas et y reste pendant un délai spécifié et le déclenche lorsque la tension dépasse un seuil haut pendant un temps donné.

1. par facilité, on parle souvent de "capacités shunt" et "d'inductances shunt"

Evidemment, on ne peut pas parler de “régulation” mais plutôt de réglage “par tout ou rien” ou “par paliers” si plusieurs capacités (ou selfs) shunt sont disponibles en parallèle. Par ailleurs, l’élément shunt étant mis en/hors service par fermeture/ouverture de son disjoncteur, des manoeuvres répétées et/ou rapides ne sont pas possibles. A la section 11.4, nous nous intéresserons à un dispositif électronique permettant de faire varier continûment et rapidement la valeur de la susceptance shunt.

11.2 Régulation de tension des machines synchrones

11.2.1 Description

La figure 11.2 donne le schéma de principe du système d’excitation d’une machine synchrone.

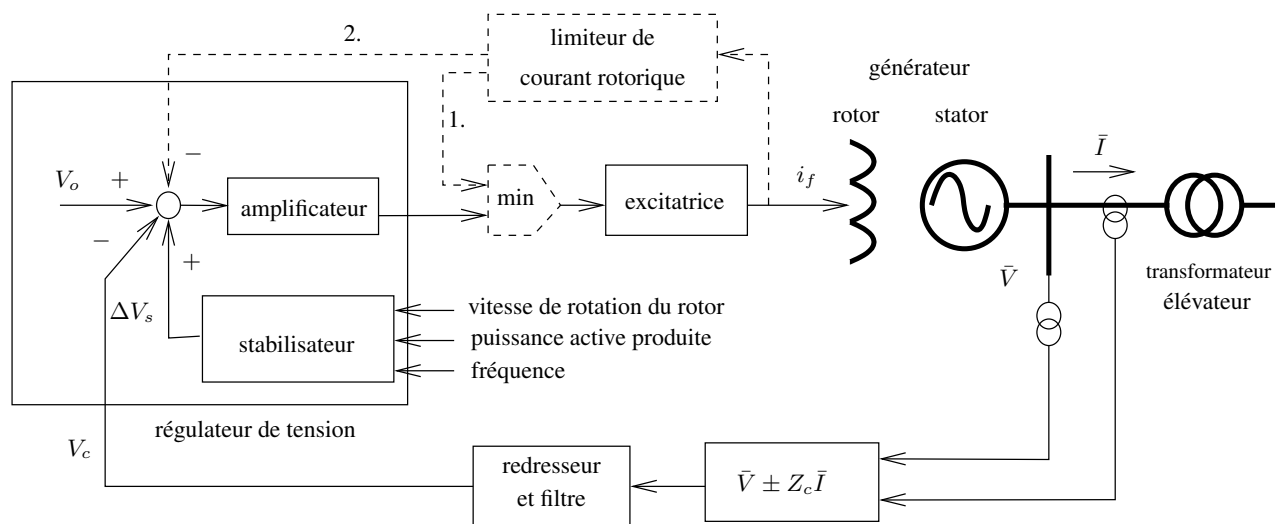


FIGURE 11.2 – schéma de principe du système d’excitation d’une machine synchrone

La tension V au jeu de barres MT du générateur est mesurée au moyen d’un transformateur de potentiel, puis redressée et filtrée pour donner un signal continu V_c , proportionnel à la valeur efficace de la tension alternative.

Le régulateur de tension compare le signal V_c à la consigne de tension V_o , amplifie la différence et met le résultat sous la forme adéquate pour la commande de l’excitatrice (p.ex. impulsions pour l’allumage de thyristors, etc. . .). Le principe général de cette régulation est d’augmenter la tension d’excitation v_f du générateur lorsque la tension terminale V diminue ou lorsque la consigne V_o augmente, et inversement.

Le régulateur est souvent doté de boucles de compensation internes² destinées à procurer une réponse dynamique satisfaisante à l’ensemble régulateur-excitatrice-générateur. Le réglage de cette

2. non représentées à la figure 11.2

compensation se fait généralement en relevant l'évolution de la tension V en réponse à un échelon de consigne V_o , le générateur fonctionnant à vide. Au départ de cette réponse indicielle, on ajuste le temps de réponse, le taux de dépassement, l'erreur statique, etc. . . du système.

Très souvent, le régulateur est aussi doté d'un "stabilisateur", circuit dont le rôle est d'ajouter au signal d'erreur $V_o - V_c$ une composante transitoire ΔV_s améliorant la dynamique de la machine en fonctionnement sur le réseau. Nulle en régime établi, cette composante ΔV_s améliore l'amortissement des oscillations du rotor³ suite à une perturbation. ΔV_s est élaborée au départ de mesures de la vitesse rotorique, de la fréquence, de puissance active, etc. . . passées au travers de fonctions de transfert appropriées⁴.

L'excitatrice est une machine auxiliaire qui procure le niveau de puissance requis par l'enroulement d'excitation du générateur. En régime établi, cette machine fournit une tension et un courant continus mais elle doit également être capable de faire varier rapidement la tension d'excitation v_f en réponse à une perturbation survenant sur le réseau.

Une excitatrice peut se présenter sous forme :

- d'une machine tournante placée sur le même axe que la turbine et le générateur. Cette machine tire donc la puissance d'excitation $v_f i_f$ de la puissance fournie par la turbine. Dans les systèmes anciens, il s'agissait d'une machine à courant continu ; dans les systèmes modernes, il s'agit d'une machine à courant alternatif (en fait, une machine synchrone du type décrit dans ce chapitre, mais de puissance beaucoup plus faible que le générateur qu'elle alimente) dont la sortie est redressée ;
- d'un système "statique" dans lequel la puissance d'excitation $v_f i_f$ est fournie par un transformateur alimenté lui-même par le réseau et dont la sortie est également redressée.

Dans certains cas, le signal V_c utilisé est de la forme :

$$V_c = |\bar{V} \pm Z_c \bar{I}| \quad (11.1)$$

où Z_c est une impédance de compensation et $\bar{I}$ le courant mesuré à la sortie du générateur. L'objectif est le suivant :

- en utilisant le signe moins et en prenant pour Z_c une fraction (entre 50 et 90 %) de l'impédance série du transformateur élévateur, V_c représente la tension en un point fictif situé entre le jeu de barres de la machine et le jeu de barres réseau correspondant. De la sorte, la chute de tension dans le transformateur élévateur est partiellement compensée et la tension du réseau est mieux régulée ;
- le signe plus est utilisé pour réguler la tension en un point fictif situé "à l'intérieur" du générateur. Une telle compensation est utilisée lorsque plusieurs générateurs sont connectés au même jeu de barres MT, une configuration typique des centrales hydrauliques où plusieurs générateurs de petite puissance partagent le même transformateur élévateur. Si on laissait les régulateurs des différentes machines contrôler la même tension, les petites différences inévitables entre générateurs et régulateurs pourraient conduire à un déséquilibre important entre les puissances réactives produites par les diverses machines, tandis que la compensation en question assure un partage équitable.

3. oscillations par rapport au mouvement uniforme correspondant au régime établi parfait

4. la synthèse de celles-ci est étudiée dans le cours ELEC0047

Dans ce qui suit, nous considérons $Z_c = 0$.

11.2.2 Réponse à une perturbation d'une machine régulée en tension

Considérons pour simplifier une machine :

- à rotor lisse ($X_d = X_q = X$) dont on néglige la saturation et la résistance statorique ($R_a = 0$) ;
- dont la régulation de tension serait infiniment précise, de sorte que la tension V aux bornes de la machine peut être supposée constante ($V \simeq V_o$) ;
- dont la puissance active P est constante, étant donné que nous nous intéressons ici à la régulation de tension et à la puissance réactive.

Nous allons considérer la réaction de la machine à une perturbation extérieure.

Considérons d'abord ce qui se passe dans le réseau. La figure 11.3 montre la caractéristique QV du réseau vu des bornes de la machine, approximée par la droite inclinée numérotée 1. Sous l'hypothèse énoncée plus haut, la caractéristique QV de la machine est une horizontale. Le point de fonctionnement du système est l'intersection de la caractéristique de la machine et de celle du réseau. C'est le point A à la figure 11.3.

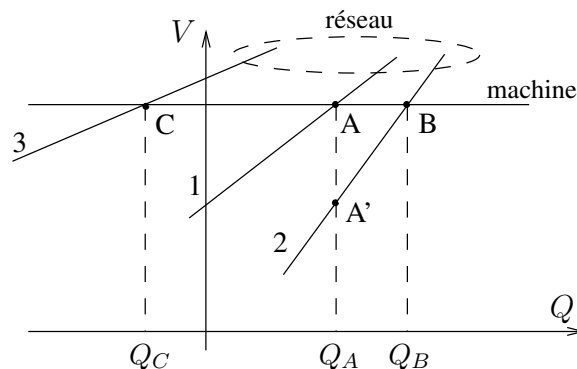


FIGURE 11.3 – caractéristiques QV du réseau et du générateur régulé en tension

Supposons qu'une perturbation survienne dans le réseau modifiant sa caractéristique de 1 en 2. Si la machine produisait une puissance réactive constante Q_A , le nouveau point de fonctionnement serait A' et la tension aux bornes de la machine tomberait. Cependant, sous l'effet de la régulation, le nouveau point de fonctionnement est B. Le maintien de la tension requiert que la machine produise davantage de puissance réactive ($Q_B > Q_A$).

De même, une perturbation qui ferait passer de la caractéristique 1 à la caractéristique 3 causerait un accroissement de tension si la machine produisait une puissance réactive constante mais, sous l'effet de la régulation, le nouveau point de fonctionnement est C, où la machine produit moins de puissance réactive ($Q_C < Q_A$).

Considérons à présent ce qui se passe à l'intérieur de la machine régulée. Le diagramme de phaseur correspondant à (8.52) est donné à la figure 11.4. Sous l'hypothèse d'une régulation de tension

parfaite, nous supposons que $\bar{V}$ est constant.

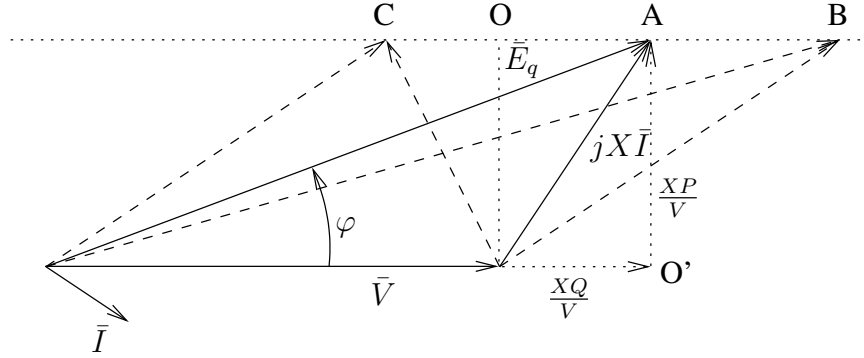


FIGURE 11.4 – diagramme de phaseur d'une machine sous contrôle d'un régulateur de tension parfait ($X_d = X_q = X, R_a = 0$)

Dans ce diagramme, la projection de $jX\bar{I}$ sur l'axe de $\bar{V}$ vaut :

$$XI_Q = \frac{XQ}{V}$$

et sur la perpendiculaire à cet axe :

$$XI_P = \frac{XP}{V}$$

P et V étant constants, l'extrémité du vecteur $jX\bar{I}$ doit donc se déplacer sur une parallèle à $\bar{V}$.

Le point 0 correspond à une production de puissance réactive nulle par la machine. Lorsque le phaseur $\bar{E}_q$ aboutit à droite (resp. à gauche) du point 0, on parle de *fonctionnement en sur-excitation* (resp. *fonctionnement en sous-excitation*).

Les diagrammes de phaseur correspondant aux points de fonctionnement A, B et C de la figure 11.3 sont repérés à la figure 11.4. En réponse à la première perturbation, l'extrémité du vecteur $\bar{E}_q$ se déplace vers la droite. L'amplitude E_q de cette f.e.m. et donc le courant d'excitation i_f augmentent sous l'action du régulateur. En réponse à la seconde perturbation, E_q et donc le courant d'excitation i_f diminuent sous l'action du régulateur.

11.2.3 Prise en compte du gain du régulateur de tension

En régime établi, on peut en première approximation représenter le système d'excitation et le générateur par le schéma bloc statique de la figure 11.5, dans lequel G_1 (resp. G_2) est le gain statique du régulateur de tension (resp. de l'excitatrice) et le dernier bloc tient compte des relations (8.40) et (8.43).

On déduit aisément de cette figure :

$$E_q = G_1 G_2 \frac{\omega_N L_{df}}{\sqrt{3} R_f} (V_o - V) = G(V_o - V) \quad (11.2)$$

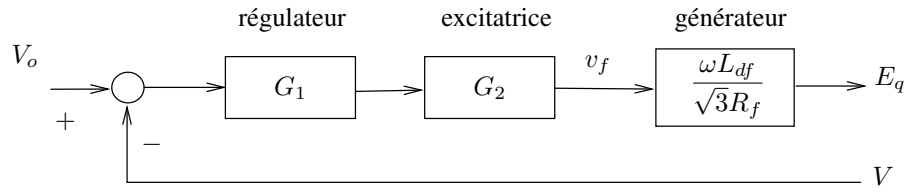


FIGURE 11.5 – schéma bloc de la régulation de tension en régime établi

où G est le gain statique en boucle ouverte de l'ensemble (régulateur + excitatrice + générateur). C'est un nombre sans dimension, que l'on peut encore définir comme la sensibilité :

$$G = \left| \frac{\Delta E_q}{\Delta V} \right|$$

Ce paramètre vaut typiquement entre 20 et 400, les valeurs faibles se rapportant généralement aux systèmes d'excitation plus anciens.

Le diagramme de phaseur de la figure 11.4 donne par ailleurs :

$$E_q^2 = \left(V + X \frac{Q}{V} \right)^2 + \left(X \frac{P}{V} \right)^2 \quad (11.3)$$

Le comportement en régime établi de la machine régulée en tension s'obtient en remplaçant E_q par l'expression (11.2) dans l'équation (11.3). Ceci permet par exemple de déterminer la caractéristique QV de la machine. Un exemple est donné à la figure 11.6 pour une machine avec $X = 2.2$ pu. Deux valeurs du gain G et deux niveaux de production active P ont été considérés. Les puissances active et réactive sont en per unit dans la base de la machine. Pour chaque combinaison, la consigne V_o a été calculée de manière à ce que la machine ne produise pas de puissance réactive sous une tension terminale de 1 pu. Ceci explique pourquoi les quatre courbes passent par le même point, et ce à des fins de comparaison.

Les courbes montrent une légère chute de la tension au fur et à mesure que la puissance réactive produite augmente. Ceci provient de l'erreur statique introduite par le régulateur proportionnel. La chute de tension est évidemment plus prononcée pour des gains G faibles. La pente de la caractéristique QV n'est que faiblement influencée par la puissance active produite.

Notons que l'on rencontre parfois (en France, par exemple) des régulateurs comportant un terme intégral qui annule l'erreur statique de régulation. Dans ce cas on peut supposer en régime établi que $V = V_o$.

11.2.4 Dispositifs limiteurs

A la section 8.8 nous avons énuméré les différentes limites imposées au fonctionnement d'un générateur. Dans ce qui suit, nous nous intéressons aux deux limites affectant la production de puissance réactive.

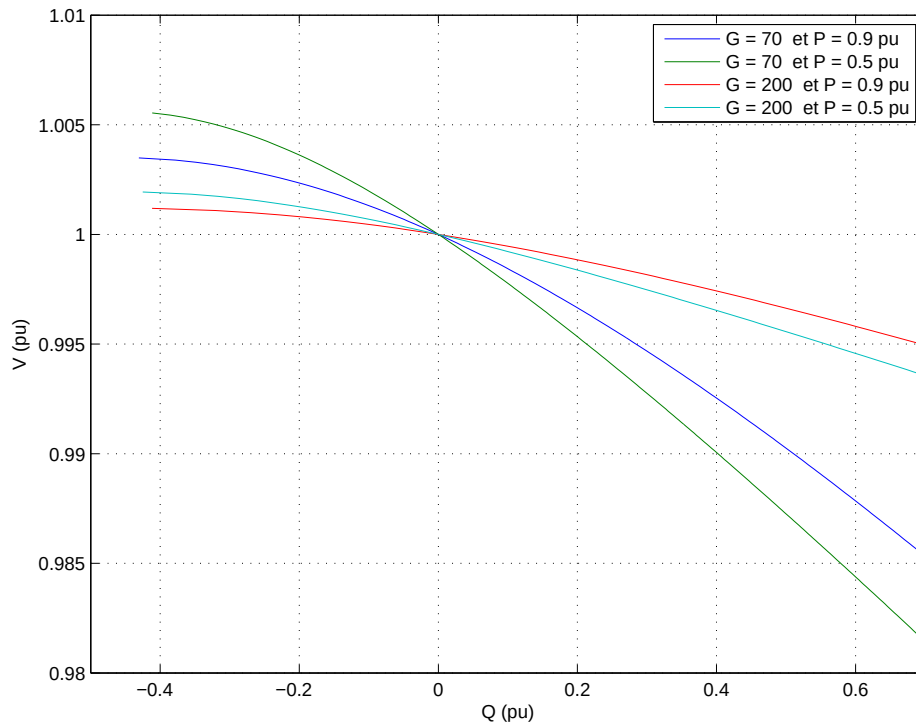


FIGURE 11.6 – caractéristiques QV d’un générateur sous contrôle de son régulateur de tension

Limiteur de courant rotorique (ou de sur-excitation)

En réponse à une perturbation importante, tel un court-circuit, il importe de laisser le régulateur de tension et l’excitatrice fournir un courant d’excitation élevé. Dans de telles circonstances, la tension d’excitation peut croître rapidement jusqu’à une valeur “de plafond” et le courant d’excitation peut atteindre une valeur de l’ordre de $2 I_{fmax}$, où I_{fmax} est le courant maximal admissible en permanence, défini à la section 8.8.

Une telle valeur ne peut être tolérée pendant plus que quelques secondes, sous peine de détériorer l’enroulement d’excitation. Toutefois, étant donné que l’échauffement est proportionnel à $\int i^2 dt$, une surcharge plus petite pourra être tolérée plus longtemps. Cette capacité de surcharge est illustrée par la courbe de la figure 11.7 (norme ANSI), donnant la relation entre le courant et la durée admise pour celui-ci. Une telle caractéristique est dite *à temps inverse*.

Les limiteurs les plus simples (souvent les plus anciens) fonctionnent avec un seuil de courant et un délai de passage en limite fixes ; ils n’exploitent donc pas vraiment la capacité de surcharge thermique décrite plus haut. Par contre, de nombreux limiteurs (souvent de construction plus récente) ont une caractéristique à temps inverse.

Une fois le délai de surcharge écoulé, le courant rotorique doit être diminué. Deux techniques sont utilisées à l’heure actuelle pour transférer le contrôle de l’excitation au limiteur :

1. la première (branche 1 à la figure 11.2) consiste à fournir à l’excitatrice le plus petit des signaux fournis, respectivement, par le régulateur de tension et par le limiteur. Ceci “ouvre”

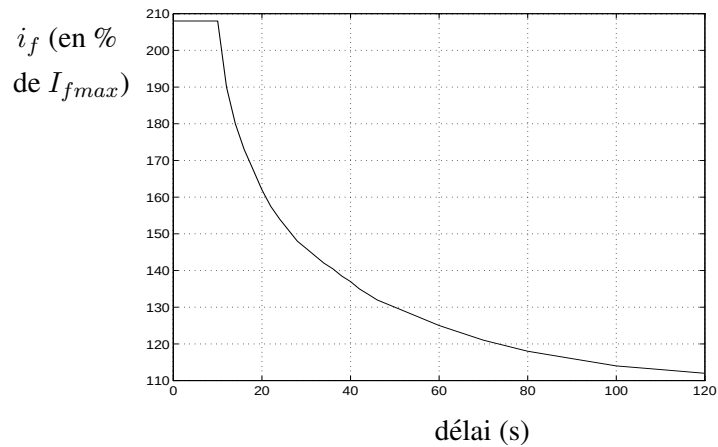


FIGURE 11.7 – relation entre le courant rotorique et la durée admise pour celui-ci

donc la boucle de rétroaction de la tension V . Comme le régulateur est mis hors service, ses boucles internes de compensation (cf section 11.2.1) le sont aussi et le limiteur doit être conçu pour assurer la stabilité du système d'excitation ;

2. dans la seconde technique (branche 2 à la figure 11.2), le limiteur injecte un signal au point d'entrée du régulateur. En temps normal, ce signal est nul tandis que lorsque le limiteur agit, il est tel que le courant d'excitation est ramené à la limite désirée. Ceci peut être vu comme une diminution de la consigne V_o telle que le courant d'excitation reste au niveau désiré. Avec cette technique, la protection de l'enroulement d'excitation repose sur le régulateur de tension. Un système doit donc être prévu pour protéger le générateur en cas de dysfonctionnement du régulateur.

Dans de nombreux cas, le régulateur de tension reprend automatiquement le contrôle de la tension dès que les conditions de fonctionnement le permettent, par exemple suite à une intervention dans le réseau.

Limiteur de courant statorique

Les limiteurs de courant statorique ne sont pas aussi répandus que les limiteurs rotoriques. La raison principale est la plus grande inertie thermique du stator, qui autorise une action plus lente par l'opérateur en centrale. Ce dernier réagira à une alarme de surcharge statorique, soit en diminuant la consigne V_o (ce qui réduit la production de puissance réactive) soit en réduisant la puissance active produite.

Dans certains pays, on rencontre toutefois des limiteurs (automatiques) de courant statorique qui agissent sur le système d'excitation de la façon décrite pour le rotor.

11.2.5 Caractéristique QV d'une machine en limite de courant rotorique et statorique

Les courbes QV d'une machine en limite sont illustrées à la figure 11.8, relative à un turbo-alternateur d'une puissance nominale apparente de 1200 MVA, dont la turbine a une puissance nominale de 1020 MW. On distingue :

- en trait plein, les courbes QV relatives à la limite de courant rotorique, pour trois niveaux de puissance active. On voit qu'en limite de courant rotorique, la production réactive du générateur varie quelque peu avec la tension ;
- en trait pointillé long, les courbes QV correspondant au courant statorique nominal I_N . Elles correspondent à la relation :

$$S = VI_N = \sqrt{P^2 + Q^2} \quad \Rightarrow \quad Q = \sqrt{(VI_N)^2 - P^2}$$

- en trait pointillé court, la courbe QV relative au fonctionnement sous contrôle du régulateur de tension, comme détaillé à la section 11.2.3.

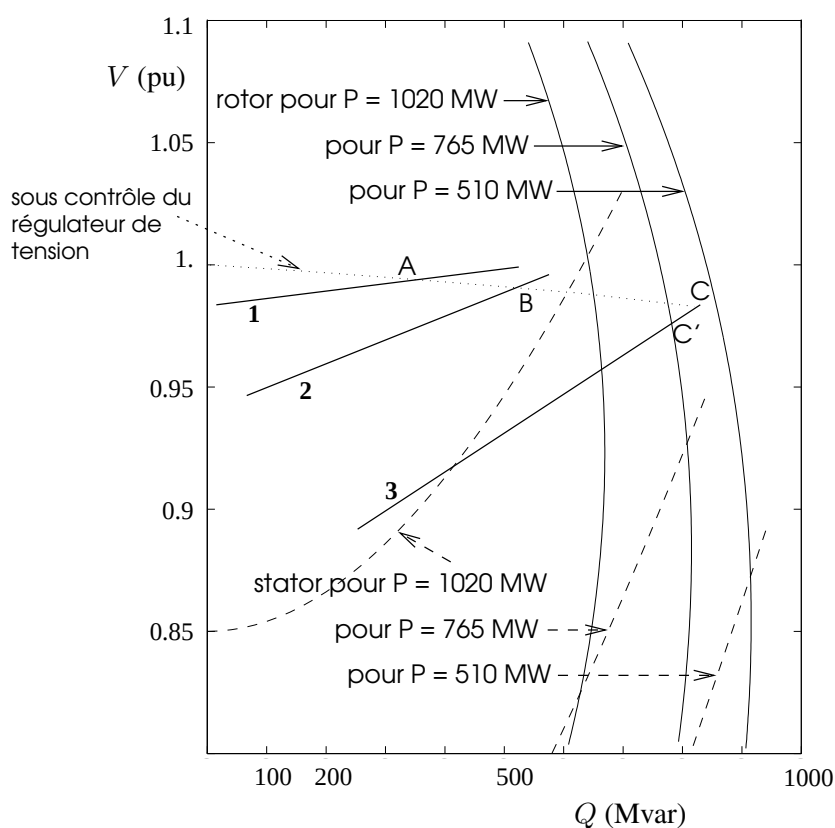


FIGURE 11.8 – caractéristiques QV du réseau et du générateur

Considérons d'abord le cas d'une production de 765 MW et supposons que la caractéristique de réseau soit la courbe numérotée 1. Le point de fonctionnement du système est A, à l'intersection des caractéristiques du réseau et du générateur.

Une première perturbation fait passer la caractéristique du réseau de la courbe 1 à la courbe 2. Le nouveau point de fonctionnement du système est B. En ce point, le générateur est toujours sous contrôle du régulateur de tension. La tension reste très proche de sa valeur avant incident tandis que la production de puissance réactive augmente en réaction à la perturbation.

Une seconde perturbation fait passer la caractéristique du réseau de la courbe 2 à la courbe 3. Le point de fonctionnement se déplace de B en C. Toutefois, en ce point, le courant rotorique est supérieur à la limite permise. Sous l'action du limiteur, la caractéristique du générateur se modifie et le point de fonctionnement qui en résulte est C'. La machine n'est plus contrôlée en tension.

Dans le cas d'une production de 1020 MW, la seconde perturbation entraîne le dépassement de la limite statorique, plus contraignante que la limite rotorique. Si le courant statorique est ramené à la valeur maximale permise (par l'opérateur ou par un dispositif limiteur), la chute de tension est plus sévère que dans le premier cas.

Dans les situations extrêmes où la tension du générateur passé en limite décroît fortement, les auxiliaires de la centrale (p.ex. les moteurs des pompes) risquant de ne plus être alimentés correctement, une protection de sous-tension déclenche le générateur. La perte correspondante des productions active et réactive risque d'aggraver la situation. Une telle protection ne doit donc pas être réglée à un niveau de tension trop élevé sous peine de déclencher la machine dans une situation d'urgence où l'on en a précisément besoin pour soutenir le réseau.

11.3 Compensateurs synchrones

Un compensateur synchrone est une machine synchrone équipée d'un régulateur de tension et utilisée seulement pour réguler la tension en un point d'un réseau.

Une telle machine est capable de produire ou d'absorber de la puissance réactive, selon nécessité. Par contre, elle n'est pas équipée de turbine et ne fournit pas de puissance active. Elle fonctionne en fait comme un moteur synchrone qui n'entraîne aucune charge mécanique. Elle consomme donc une faible puissance active correspondant aux pertes Joule statoriques et aux frottements mécaniques.

Le diagramme de phaseur de la figure 11.4 se simplifie et devient celui de la figure 11.9, qui montre séparément les fonctionnements en sur-excitation et sous-excitation.

Au lieu d'installer des compensateurs synchrones on opte plutôt à l'heure actuelle pour des compensateurs statiques, qui font appel à l'électronique de puissance.

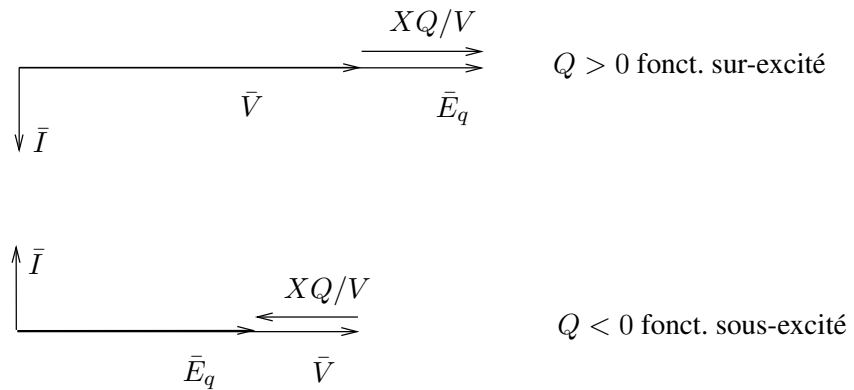


FIGURE 11.9 – diagramme de phaseur d’un compensateur synchrone ($X_d = X_q = X$, $R_a = 0$)

11.4 Compensateurs statiques de puissance réactive

11.4.1 Usage

Les compensateurs statiques de puissance réactive (en abrégé, compensateurs statiques⁵) sont des dispositifs rapides d’injection de puissance réactive faisant appel à l’électronique de puissance.

On les rencontre d’abord comme éléments de compensation dynamique des charges, où ils servent :

- à équilibrer des charges présentant un déséquilibre entre phases
- à stabiliser la tension aux bornes d’une charge variant rapidement (fours à arc, laminoirs, etc. . .). Ces variations peuvent donner lieu au “flicker de tension”, fluctuation d’une fréquence entre 2 et 10 Hz qui cause le “papillotement” des lampes à incandescence et perturbe appareils électroniques, téléviseurs, etc. . .

Dans ce cours, c’est aux applications réseau que nous nous intéressons. Dans ce contexte, les compensateurs statiques sont utilisés :

- pour maintenir quasi constantes les tensions en certains noeuds
- pour améliorer la stabilité de fonctionnement.

Les compensateurs statiques constituent la première génération de dispositifs FACTS⁶, apparus à la fin des années 70.

Les compensateurs statiques font appel au *thyristor*, composant électronique utilisé comme interrupteur. Son symbole est donné à la figure 11.10. Il fonctionne selon le principe suivant :

- le thyristor laisse passer le courant quand l’anode est à un potentiel électrique supérieur à la cathode ($v_A - v_C > 0$) et si une impulsion de tension est envoyée sur la gachette (cette impulsion est donnée par le circuit de commande, indépendant mais synchronisé sur la partie puissance) ;
- lorsque le courant veut changer de sens, le thyristor se bloque et le courant ne peut plus passer.

5. en anglais : Static Var Compensators (SVC)

6. Flexible Alternating Current Transmission Systems

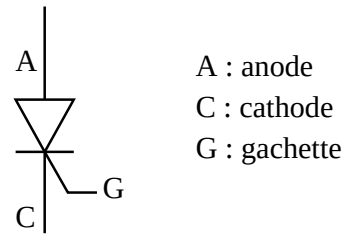


FIGURE 11.10 – thyristor

11.4.2 TSC : principe

Le premier type de compensateur statique est le *Thyristor Switched Capacitor* (TSC) dont le schéma de principe est donné à la figure 11.11.

Le TSC est constitué d'un certain nombre de condensateurs shunt en parallèle, chacun doté d'un interrupteur bidirectionnel à thyristors. Lorsque la tension au jeu de barres HT diminue (resp. augmente), le nombre de condensateurs mis en service augmente (resp. diminue). La variation est donc typiquement par paliers. La logique de contrôle comporte une bande morte dans laquelle il n'y a pas de réaction du dispositif.

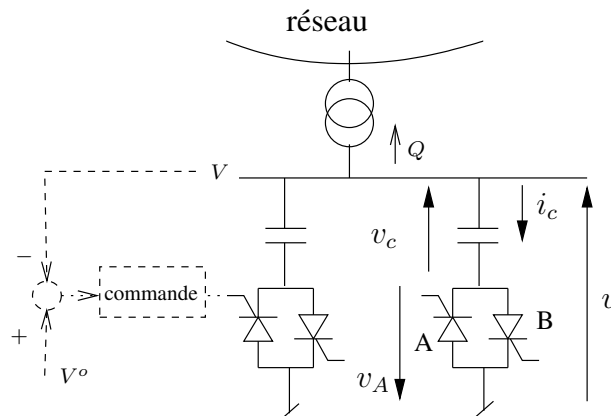


FIGURE 11.11 – schéma de principe d'un TSC

Le processus de commutation d'un TSC est illustré à la figure 11.12.

En $t = 0$, le thyristor B est en train de conduire. Le courant i_c qui traverse le condensateur est en avance de 90° sur la tension v_c à ses bornes, qui est aussi la tension v du réseau (cf fig. 11.11). A l'instant t_1 , le courant s'annule et le thyristor B se bloque. Dans les instants qui suivent immédiatement t_1 , le condensateur reste chargé à la tension de crête V_{max} , tandis que la tension v du réseau diminue. La tension v_A aux bornes du thyristor A (cf fig. 11.11) valant $v_c - v > 0$, ce dernier est polarisé dans le bon sens pour la conduction. L'envoi d'un signal sur sa gachette le fait conduire. Il importe de ne pas attendre pour envoyer ce signal, car la différence de tension aux bornes du thyristor est en train d'augmenter et sa commutation créerait alors un courant transitoire

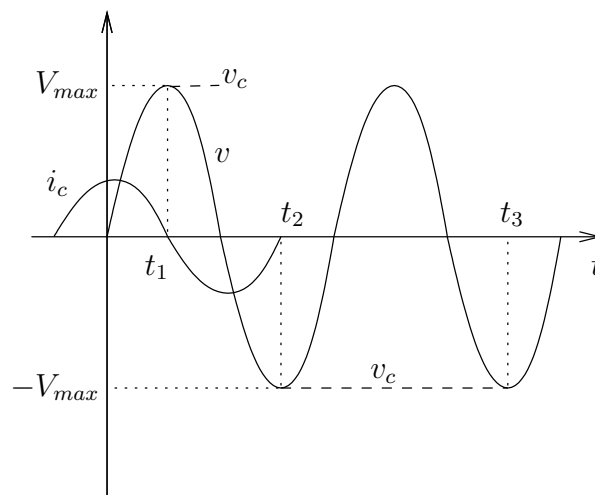


FIGURE 11.12 – commutation dans un TSC

important. En pratique, on ne peut empêcher complètement ce dernier ; c'est la raison pour laquelle on place en série avec le condensateur une faible inductance, qui n'est pas représentée à la figure 11.11.

Le thyristor A se bloque à l'instant t_2 . Si l'on suppose qu'à cet instant on désire mettre le condensateur hors service, on n'envoie pas de commande sur la gachette du thyristor B. Ce faisant, le condensateur reste chargé à la tension $-V_{max}$. On pourra le remettre en service au plus tôt à l'instant t_3 , quand la tension v du réseau sera à nouveau égale à $-V_{max}$ ⁷. En conclusion, la commutation du condensateur ne peut se faire qu'à des multiples entiers de la demi-période.

11.4.3 TCR : principe

Le second type de compensateur statique est le *Thyristor Controlled Reactor* (TCR) dont le schéma de principe est donné à la figure 11.13.

Dans un TCR, on retarde l'instant d'allumage des thyristors placés en série avec l'inductance, comme représenté à la figure 11.14. Dans cette figure, α est l'*angle de retard à l'allumage* mesuré par rapport au zéro de tension, tandis que σ est l'*angle de conduction*. Ce dernier peut varier de 180 à 0 degrés.

Pour différentes valeurs de σ , on obtient les ondes de courant montrées à la figure 11.15. Un développement en série de Fourier de ce signal périodique montre que l'amplitude de la fondamentale (50 ou 60 Hz) vaut :

$$I_{fond} = \frac{V}{\omega L} \frac{\sigma - \sin \sigma}{\pi} \quad (11.4)$$

7. notons qu'en pratique, si l'on attend suffisamment longtemps, le condensateur finit par se décharger, ce qui complique le choix de l'instant de commutation

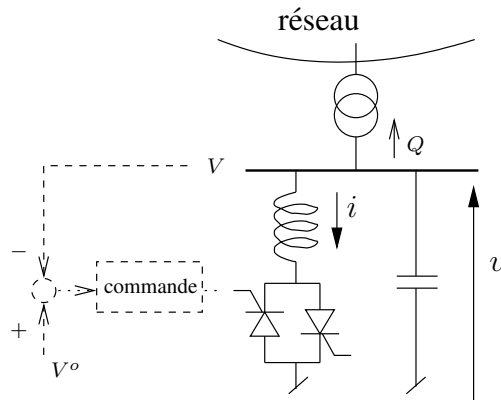


FIGURE 11.13 – schéma de principe d'un TCR

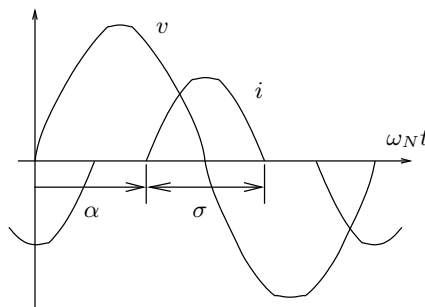


FIGURE 11.14 – retard à l'allumage dans un TCR

où σ est exprimé en radians. Quand on fait varier σ de π à 0, I_{fond} varie de $V/\omega L$ à zéro, ce qui revient à considérer que l'on a une inductance variant entre L et l'infini. Le TCR se comporte donc comme une inductance continûment variable.

Ceci permet de faire varier l'absorption de puissance réactive. Pour obtenir un dispositif pouvant produire de la puissance réactive, on place un condensateur fixe en parallèle avec l'inductance variable. La production réactive de l'ensemble est maximale quand les thyristors ne conduisent pas ; elle est minimale lorsqu'ils conduisent en permanence. En général, la plage de variation va de l'absorption à la production.

Contrairement au TSC, le TCR permet un réglage continu de la susceptance mais il génère des harmoniques, qui doivent être filtrés. L'onde de courant étant symétrique dans le temps, elle ne contient que des harmoniques d'ordre impair. Ceux-ci peuvent être filtrés comme suit :

- pour obtenir un système triphasé, trois TCR monophasés sont montés en triangle, conformément au schéma de la figure 11.16(a). Dans ce montage, les trois phases étant équilibrées, les harmoniques de rang 3, 6, 9, etc ... circulent dans le triangle et les courants de ligne en sont exempts. A titre indicatif, la figure 11.17 montre l'évolution des courants dans deux branches du triangle et dans la ligne incidente à ces deux branches. Les autres harmoniques (de rang 5, 7, 11, etc...) sont éliminés au moyen de filtres (qui représentent une partie importante de l'investissement).
- on peut éliminer les harmoniques de rang 5 et 7 en utilisant deux systèmes triphasés de même

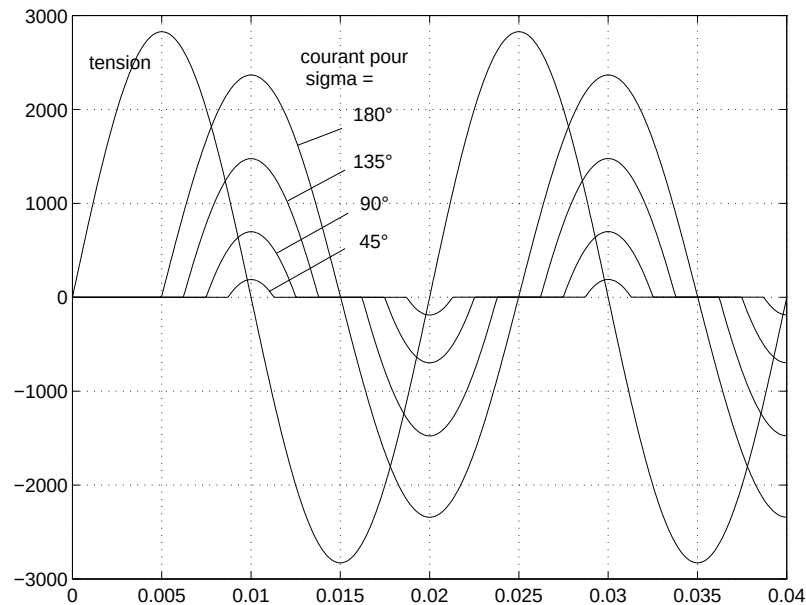


FIGURE 11.15 – ondes de tension et de courant dans un TCR pour différentes valeurs de σ

puissance, connectés aux enroulements secondaires d'un transformateur à trois enroulements, l'un étant monté en triangle et l'autre en étoile (cf figure 11.16(b)). Grâce au déphasage de 30 degrés entre tensions secondaires, les courants de ligne au primaire sont exempts des harmoniques 5 et 7 ; les autres harmoniques sont éliminés avec des filtres plus simples.

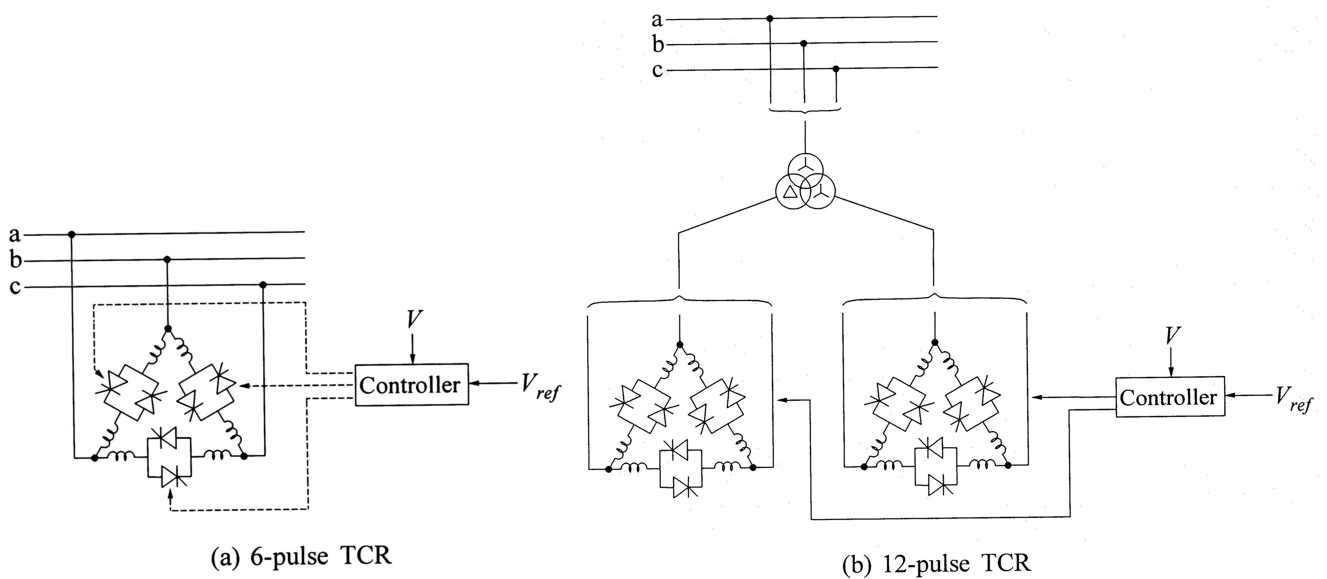


FIGURE 11.16 – montages des TCR pour éliminer les principaux harmoniques

Mentionnons que l'on peut combiner au sein d'un même compensateur le TCR, le TSC, des capacités commutables par disjoncteurs et des capacités fixes. En anglais, un tel ensemble est appelé *Static Var System*.

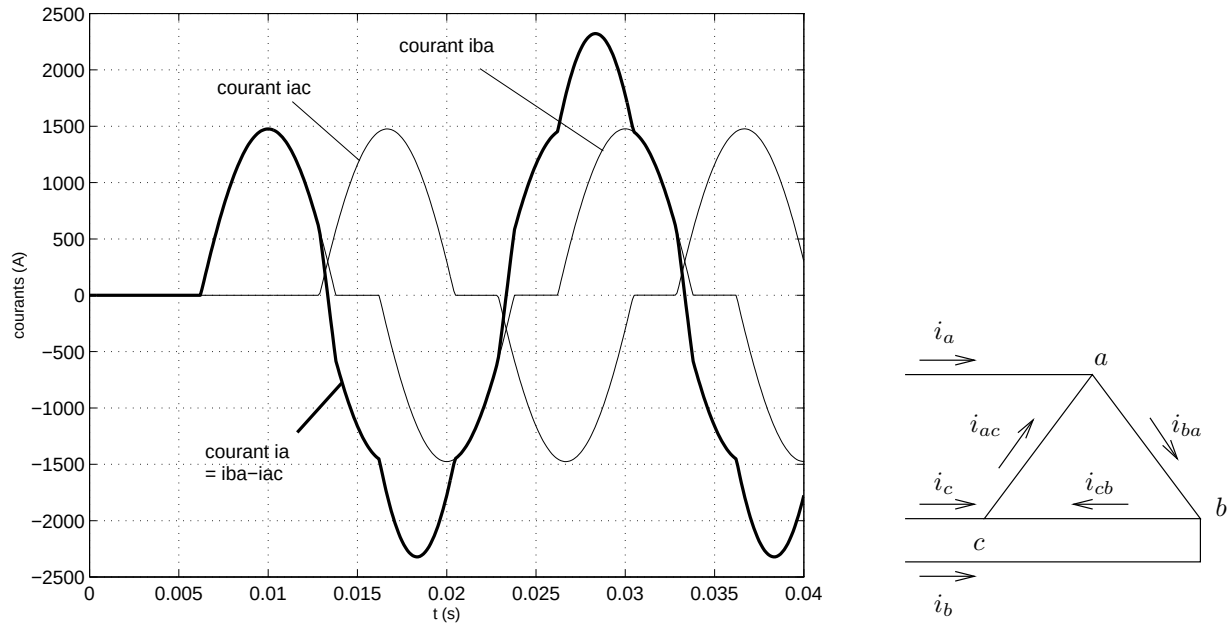


FIGURE 11.17 – ondes de courant dans un TCR

11.4.4 Schéma bloc et puissance nominale

Le schéma-bloc du TCR, relatif au fonctionnement *en régime établi*, est donné à la figure 11.18. La valeur efficace V de la tension au jeu de barres du réseau est mesurée et comparée à la consigne V_o . La différence est amplifiée et sert à commander les thyristors. Les limites $-1/(\omega L)$ et 0 correspondent aux conditions de conduction limites des thyristors, tandis que $B_C = \omega C$ est la susceptance du condensateur en parallèle. B est la susceptance de l'ensemble.

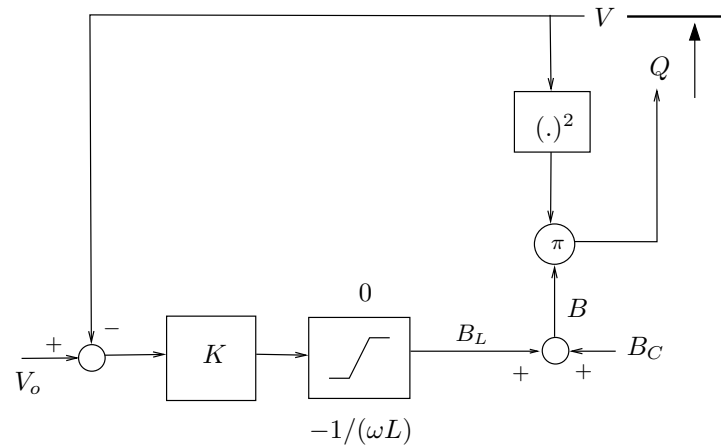


FIGURE 11.18 – schéma-bloc statique du TCR

La puissance nominale du compensateur est donnée par :

$$Q_{nom} = \max \left(\left| B_C - \frac{1}{\omega L} \right|, B_C \right) \cdot U_{nom}^2 \quad (11.5)$$

où U_{nom} est la tension nominale entre phases du jeu de barres. Dans de nombreux cas, l'appareil est destiné à produire plus de puissance réactive qu'à en consommer. Dans ce cas (11.5) devient :

$$Q_{nom} = B_C U_{nom}^2 \quad (11.6)$$

Les paramètres sont usuellement fournis en per unit dans la base $(U_{nom}/\sqrt{3}, Q_{nom})$. Dans cette base et en supposant que (11.6) s'applique, on a simplement : $B_C = 1$ pu. Dans cette même base, K vaut de 25 à 100 pu/pu.

11.4.5 Caractéristique QV et régulation de tension

Un exemple de caractéristique QV est donné à la figure 11.19. Le système per unit précité a été utilisé ; dans l'exemple considéré, on a $B_C = 1$ pu et $B_C - \frac{1}{\omega L} = -0.3$ pu.

La plage de fonctionnement normal est la partie à faible pente, où la tension est contrôlée. Une vue agrandie de cette partie de la caractéristique est donnée à la figure 11.20, pour deux valeurs de K . Dans les deux cas, la consigne V_o a été choisie de manière à ce que le compensateur ne fournisse pas de puissance réactive sous $V = 1$ pu. On voit que la caractéristique peut être assimilée à un segment de droite.

Les autres parties de la caractéristique correspondent aux conditions de conduction limites des thyristors, soit respectivement $Q = (B_C - \frac{1}{\omega L})V^2 = -0.3V^2$ et $Q = B_C V^2 = V^2$.

La caractéristique complète est montrée en trait gras à la figure 11.19.

La figure 11.21, désormais familière, superpose la caractéristique du compensateur à celle du réseau vu du même jeu de barres, dans trois configurations différentes. Lors du passage de la caractéristique 1 à la caractéristique 2, le compensateur maintient la tension (presque) constante en produisant plus de puissance réactive. Si une perturbation plus importante conduit à la caractéristique 3, le compensateur entre en limite et se comporte alors comme un simple condensateur.

En fonctionnement normal, l'opérateur du centre de conduite (ou un système automatique) est usuellement chargé de maintenir la production réactive du compensateur dans un intervalle autour de zéro. L'objectif est de ménager une "marge de manoeuvre" sur le compensateur afin que celui-ci puisse répondre rapidement à des incidents. Pour corriger la production de réactif lorsqu'elle est sortie de l'intervalle spécifié, des condensateurs shunt peuvent être enclenchés ou déclenchés en parallèle avec le TCR. Ces manoeuvres peuvent être effectuées via un dispositif TSC, ce qui conduit à combiner les montages TSC et TCR.

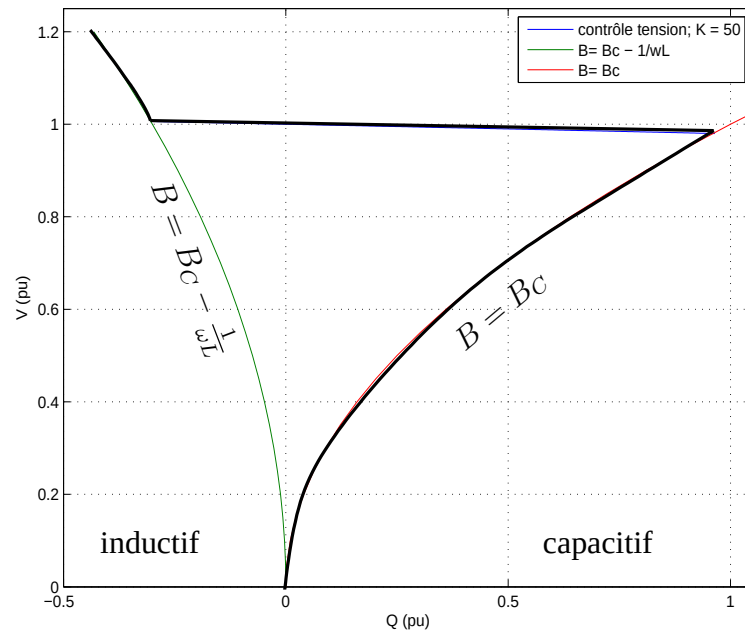


FIGURE 11.19 – caractéristique QV d'un compensateur statique

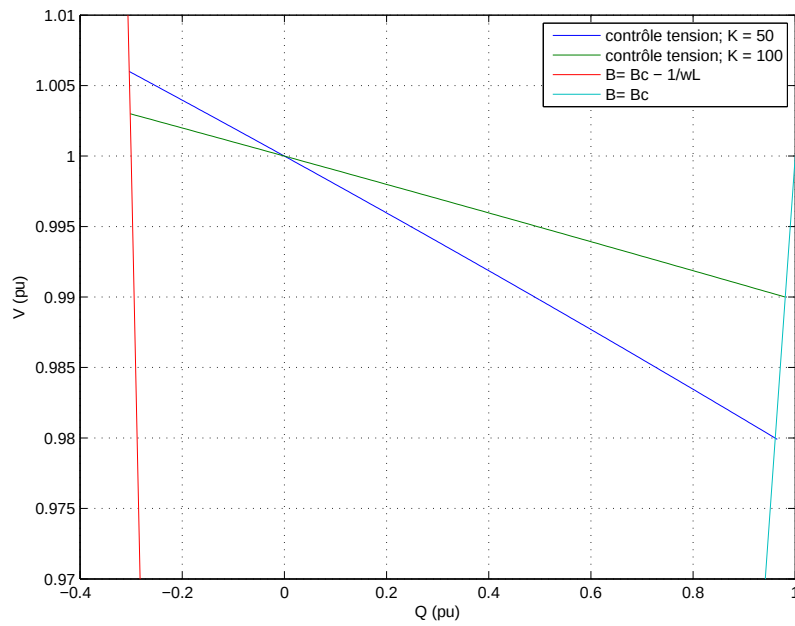


FIGURE 11.20 – vue agrandie de la figure 11.19 pour deux valeurs du gain K

11.4.6 Comparaison avec un compensateur synchrone

Par rapport aux compensateurs synchrones, les compensateurs statiques présentent une plus grande vitesse de réponse, ne contribuent pas aux courants de court-circuit et sont d'un entretien plus aisé. Par contre, par construction, ils ne présentent pas de f.e.m. interne, ce qui diminue leur capacité à

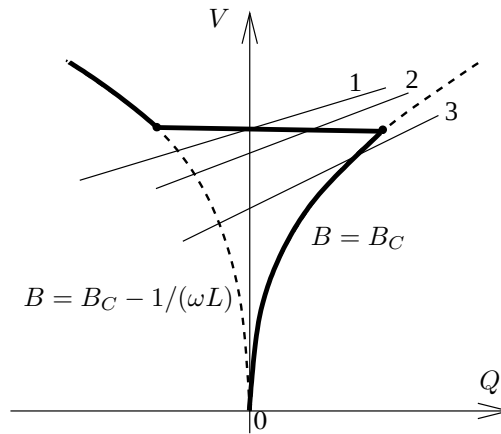


FIGURE 11.21 – principe de la régulation de tension par un compensateur statique

soutenir la tension en régime très perturbé.

Notons toutefois qu'il s'agit de dispositifs relativement coûteux, dont l'usage se justifie dans les cas où l'on a besoin d'une grande rapidité d'action et/ou une régulation précise. Dans les autres cas, il convient d'analyser si des condensateurs ou inductances shunt manoeuvrés par ouverture/fermeture de disjoncteur ne suffisent pas.

11.5 Régulation de tension par les régleurs en charge

11.5.1 Principe

Le procédé le plus couramment utilisé pour contrôler la tension des réseaux de tensions nominales inférieures consiste à doter les transformateurs qui les alimentent de régleurs en charge automatiques. Ces derniers sont pilotés par un asservissement dont le rôle est de maintenir la tension du jeu de barres contrôlé au voisinage d'une consigne, en ajustant le rapport de transformation. De la sorte, les variations de tension au niveau supérieur sont corrigées.

On trouve de tels dispositifs sur les transformateurs qui alimentent les réseaux de distribution à Moyenne Tension (MT), où ils constituent le moyen le plus répandu de régler la tension. Les autres moyens sont les condensateurs shunt et les unités de production distribuées connectées sur le réseau MT. Le réglage de celles-ci n'est pas encore très répandu mais il est appelé à prendre de l'importance dans le futur, si l'on assiste au déploiement attendu des sources d'énergie renouvelable.

On rencontre également des transformateurs avec régleurs en charge automatiques entre les niveaux de transport (THT) et de répartition (HT). Là encore, ils constituent souvent le principal

moyen de régler la tension en l'absence de générateurs ⁸.

La figure 11.22 montre le schéma équivalent simplifié d'un transformateur alimentant le jeu de barres de départ d'un réseau de distribution. Le rapport r est ajusté automatiquement pour maintenir V_2 dans une *bande morte* $[V_2^0 - \epsilon V_2^0 + \epsilon]$.

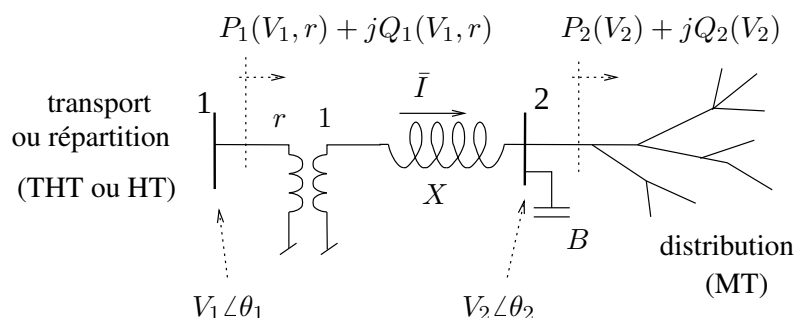


FIGURE 11.22 – réseau de distribution alimenté par un transformateur avec régleur en charge

La consigne V_2^0 est généralement choisie un peu supérieure à la tension nominale, de manière à compenser la chute de tension dans le réseau de distribution et alimenter le consommateur le plus éloigné du départ sous une tension encore correcte, garantie par le contrat de fourniture d'électricité.

Mentionnons que dans certains cas, au lieu du module $|\bar{V}_2|$ de la tension MT, on fournit au régleur en charge le signal $|\bar{V}_2 - Z_c \bar{I}|$ où $\bar{I}$ est le courant entrant dans le réseau de distribution et Z_c une impédance de compensation. De la sorte, la tension n'est pas régulée à la sortie du transformateur mais bien en un point situé "en aval", c'est-à-dire plus près des consommateurs situés en bout de réseau de distribution ⁹.

Les régleurs en charge agissent lentement par rapport aux régulations décrites plus haut dans ce chapitre. Ils passent les prises une par une tant que la tension contrôlée reste en dehors de sa bande morte. Le délai minimum T_m pour passer une prise est d'origine mécanique ; il est de l'ordre de 5 s. Des délais supplémentaires, allant de quelques secondes à une ou deux minutes, sont intentionnellement ajoutés à T_m , de manière à éviter des passages de prises fréquents ou inutiles, cause d'usure de l'équipement. En particulier, suite à une perturbation, il importe de laisser s'éteindre les transitoires sur le réseau avant de corriger, si nécessaire, les tensions MT. Ces délais additionnels peuvent être fixes ou variables. Dans le second cas, on utilise souvent une caractéristique à temps inverse dans laquelle le délai est plus grand pour des erreurs de tension plus petites. Mentionnons également que très souvent le premier passage de prise s'effectue avec un délai plus important (p.ex. de 30 à 60 s) que les passages ultérieurs (p.ex. 10 s par prise). Enfin, dans le cas où il y a plusieurs niveaux de régleurs en charge en cascade, c'est le régleur de niveau de tension le plus élevé qui doit agir le premier, sous peine de créer des oscillations entre régleurs.

La limite inférieure du rapport de transformation r est de l'ordre de 0.85 - 0.90 pu/pu et la limite

8. typiquement, les centrales qui débitaient sur le réseau HT, du temps où ce dernier constituait l'ossature principale du réseau électrique, ont été remplacées par des centrales de plus grande puissance, débitant sur le réseau THT

9. Cette technique s'apparente à celle consistant à compenser une partie de l'impédance du transformateur élévateur d'un générateur synchrone

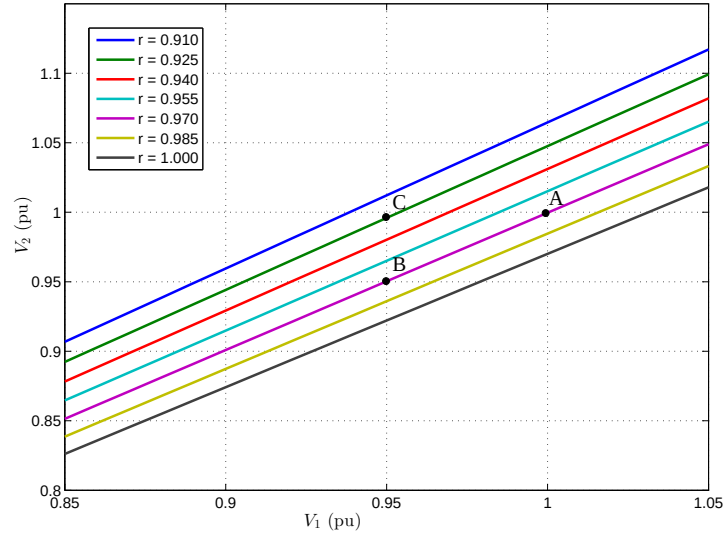


FIGURE 11.23 – variation de V_2 avec V_1 , pour différentes valeurs de r

supérieure de l'ordre de 1.10 - 1.15 pu/pu. Le pas de variation Δr est quant à lui de l'ordre de 0.005 - 0.015 pu/pu. Pour des raisons de stabilité, Δr est généralement inférieur à 2ϵ ; dans ce cas, il y a deux positions possibles du régleur dans la bande morte.

11.5.2 Comportement d'un réseau de distribution contrôlé par un régleur en charge automatique

Revenons à la figure 11.22. Supposons que la puissance complexe consommée par les charges et le réseau en aval du noeud 2 varie avec la tension V_2 selon le modèle (statique) $P_2(V_2) + jQ_2(V_2)$. En fait, les fonctions $P_2(V_2)$ et $Q_2(V_2)$ décrivent le comportement de la charge après extinction de la dynamique décrite par (9.11), c'est-à-dire en pratique quelques secondes après une variation de tension V_2 . Il s'agit de la caractéristique *à court terme* de la charge.

Pour une des valeurs possibles de r , il existe également des caractéristique à court terme $P_1(V_1, r)$ et $Q_1(V_1, r)$ de l'ensemble (charge + condensateur + réseau de distribution + transformateur) vu du jeu de barres de transport. La détermination numérique de ces caractéristiques a fait l'objet de l'exercice 3 du chapitre 9. Les courbes des figures 11.23 à 11.26 se rapportent aux données de l'exercice en question.

Le bilan de puissance du système s'écrit :

$$P_1 = P_2(V_2) \quad (11.7)$$

$$Q_1 = Q_2(V_2) - BV_2^2 + XI^2 = Q_2(V_2) - BV_2^2 + X \frac{P_2^2(V_2) + Q_2^2(V_2)}{V_2^2} \quad (11.8)$$

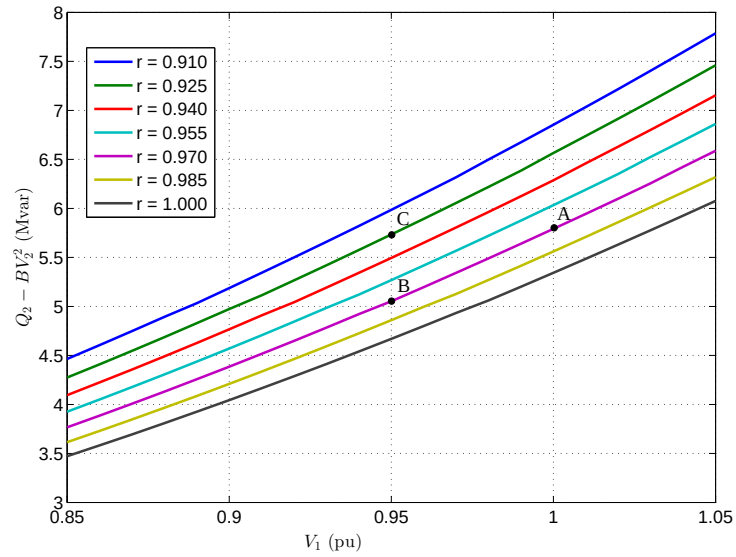


FIGURE 11.24 – variation avec V_1 de la puissance réactive $Q_2 - BV_2^2$ à la sortie MT du transformateur, pour différentes valeurs de r

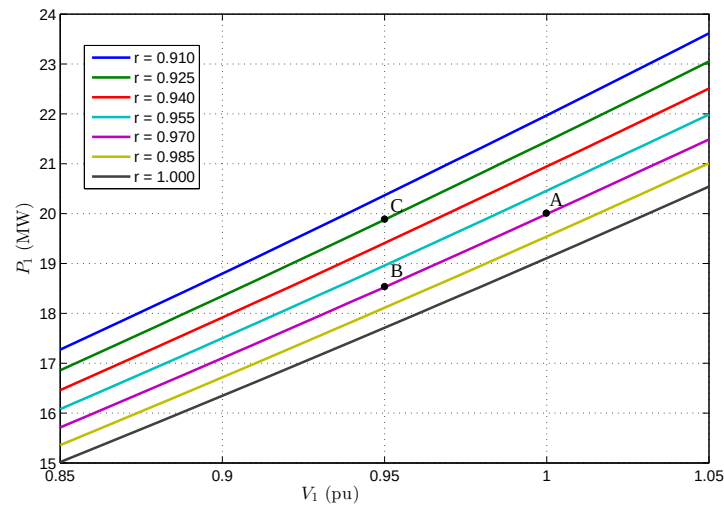


FIGURE 11.25 – variation de la puissance active $P_1 = P_2$ avec V_1 , pour différentes valeurs de r

Considérons le point de fonctionnement initial du système repéré par le point A dans les figures 11.23 à 11.26. Supposons que la tension V_1 du jeu de barres de transport subit une décroissance permanente. Dans les figures en question, on a considéré une chute de 5 %. Après extinction des transitoires, le nouveau point de fonctionnement est B, situé sur la même caractéristique à court terme, et ce aussi longtemps que le rapport r ne se modifie pas.

Sous l'effet de la diminution de V_1 , V_2 diminue également, comme le montre la figure 11.23. Il en résulte une diminution des puissances active et réactive à la sortie du transformateur, comme le

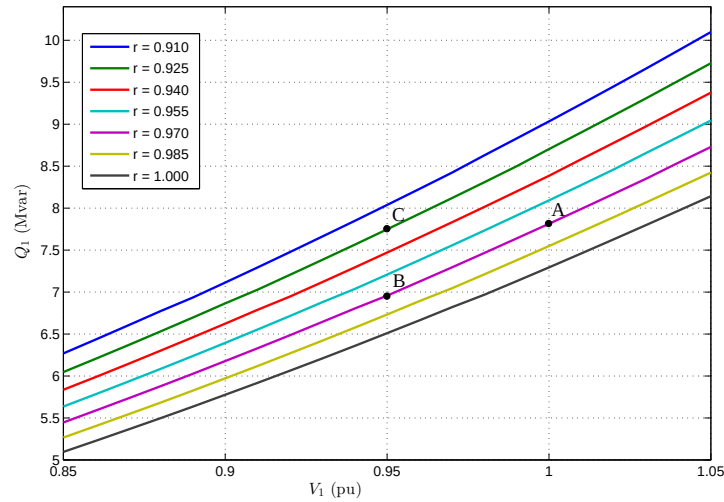


FIGURE 11.26 – variation de la puissance réactive Q_1 avec V_1 , pour différentes valeurs de r

montrent les figures 11.24 et 11.25.

On peut supposer que V_2 est sortie de la bande morte $[V_2^o - \epsilon, V_2^o + \epsilon]$ surveillée par le régulateur en charge. La temporisation initiale une fois écoulée, le régulateur passe donc progressivement des prises. Sous l'effet du changement de r , les caractéristiques à court terme se modifient. Dans l'exemple considéré, il y a trois passages de prise successifs, qui conduisant au point de fonctionnement C où la tension V_2 est revenue dans sa bande morte. Les puissances active et réactive à la sortie MT du transformateur reviennent donc à leurs valeurs initiales, comme le confirment les figures 11.24 et 11.25.

Le transformateur étant supposé sans pertes actives (cf équation (11.7)), la puissance P_1 revient aussi à sa valeur avant perturbation. Il en va de même de la puissance réactive Q_1 car chaque terme du membre de droite de l'équation (11.8) revient également à sa valeur avant perturbation.

En pratique, la tension V_2 n'est pas exactement ramenée à sa valeur initiale, à cause du nombre fini de prises de réglage et de l'insensibilité du régulateur en charge dans sa bande morte. Toutefois, si l'on néglige cette erreur finale, on voit que l'effet du régulateur en charge est de restaurer les puissances soutirées au réseau de transport à leurs valeurs avant perturbation.

Vu les temps de réaction des régulateurs en charge, ce processus de restauration de la charge est un exemple typique de *dynamique à long terme*. On peut dire que, sous l'effet du régulateur en charge, la caractéristique à long terme de la charge est une puissance constante. Ceci néglige l'effet de la bande morte et ne s'applique pas si le régulateur arrive en butée avant de rallier la bande morte.

Chapitre 12

Analyse des défauts équilibrés

Un *défaut* est une perturbation qui empêche le flux normal de puissance dans un réseau d'énergie électrique. Une grande partie des défauts survenant dans les réseaux d'énergie électrique sont causés par la foudre, qui crée un court-circuit entre au moins une des phases et la terre.

Ce chapitre est consacré à l'étude des courts-circuits triphasés symétriques, pour lesquels on peut encore recourir à une analyse par phase, d'où le nom de *défaut équilibré*.

12.1 Phénomènes liés aux défauts

12.1.1 Foudre

Le lecteur trouvera dans l'appendice de ce chapitre une vue d'ensemble des phénomènes physiques (par ailleurs assez complexes) liés à la foudre.

12.1.2 Court-circuit dans le réseau

Comme expliqué dans l'appendice, le rôle du câble de garde est d'attirer sur lui (à la manière d'un paratonnerre) les décharges de foudre les plus importantes. Cependant, il se peut qu'il en remplisse pas son rôle et que la ligne soit touchée au niveau d'un de ses conducteurs de phase.

Quand la foudre touche un conducteur de phase, les charges électriques se déversent dans les deux directions, à partir du point d'impact. Ceci donne naissance à deux ondes de tension se propageant le long de la ligne à la vitesse de la lumière¹. Lorsqu'une telle onde atteint la chaîne d'isolateurs

1. en première approximation la tension maximale de chaque onde vaut $V = Z_c I / 2$. Pour $Z_c \simeq 300 \Omega$ et $I / 2 = 15$ kA, on obtient $V = 4.500.000$ V !

la plus proche, cette dernière est soumise à une différence de potentiel très élevée. S'il y a rupture diélectrique de l'intervalle d'air qui l'entoure, un arc électrique apparaît entre le conducteur et le pylône.

Lors que la foudre touche le câble de garde (ce qui suit s'applique évidemment au cas où la foudre aurait touché directement un pylône), le courant se déverse dans ce conducteur jusqu'à atteindre les pylônes les plus proches. Dans le meilleur des cas, ces courants s'écoulent dans le sol, via les fondations des pylônes. Cependant, le haut du pylône concerné peut monter en tension sous l'effet de l'injection brusque d'un courant élevé dans la structure métallique et dans la prise de terre (qui, toutes deux, présentent une impédance). Cette montée en tension peut conduire à une différence de potentiel élevée entre le haut du pylône et le conducteur de phase, conduisant à une rupture diélectrique de l'intervalle d'air qui entoure la chaîne d'isolateurs et donc à un arc électrique entre le conducteur de phase et le pylône.

Dans ces deux scénarios de contournement d'isolateur, même après que les charges provenant du coup de foudre se soient évacuées dans le sol, l'air ionisé par l'arc reste conducteur et une connexion de faible impédance demeure entre le réseau et la terre, créant ainsi un court-circuit, dont le courant est en fait alimenté par les générateurs.

12.1.3 Protections et disjoncteurs

Les courants circulant dans le réseau en présence du court-circuit ont une amplitude élevée par rapport aux courants existant en fonctionnement normal. Ils doivent être rapidement éliminés sous peine de détériorer les équipements. Par ailleurs, la mise au potentiel nul d'un point du réseau de transport risque de déstabiliser le système (rupture de synchronisme entre générateurs ou instabilité de tension). Enfin, les consommateurs subissent une chute de tension d'autant plus marquée qu'ils sont proches du défaut ; certains processus industriels sont sensibles à de tels creux de tension.

Les protections détectent l'apparition des courants élevés (ou la diminution de l'impédance vue des extrémités de la ligne) et envoient aux disjoncteurs concernés l'ordre d'ouverture. Le délai total d'élimination du défaut se décompose en trois parties :

1. temps pour les circuits de détecter le défaut et d'envoyer l'ordre d'ouverture au disjoncteur
2. temps pour les contacts de ce dernier de se mettre en mouvement
3. temps pour éteindre d'arc électrique qui a pris naissance dès que les contacts électriques se sont écartés.

Pour les disjoncteurs qui équipent les réseaux de transport, on peut considérer que le délai total d'élimination est d'*au plus 5 alternances (0.1 s)*. Les disjoncteurs les plus performants permettent de descendre à 2 alternances. Notons que les disjoncteurs qui équipent les réseaux de répartition ou de distribution sont généralement plus lents (mais moins coûteux !). Ils peuvent prendre 8 alternances, voire davantage, pour éliminer un défaut apparu à ces niveaux de tension inférieurs.

Lorsque les disjoncteurs d'extrémité de la ligne en court-circuit ont déconnecté celle-ci du reste du réseau, l'arc électrique n'est plus alimenté et s'éteint de lui-même.

Le réseau se retrouve privé de la ligne ainsi mise hors service. Dans les grands réseaux de transport, on souhaite généralement la remettre en service le plus rapidement possible. C'est le rôle du dispositif de *ré-enclenchement automatique* de la ligne. Ce dernier doit cependant attendre que l'air ait recouvré ses propriétés d'isolant. Le délai est typiquement de l'ordre de 0.3 seconde.

Le court-circuit causé par la foudre est typiquement un *défaut fugitif* : la mise hors service de la ligne suffit à le faire disparaître. Un *défaut permanent* est causé par le contact de la ligne avec un objet, par la glace accumulée sur les isolateurs, voire dans les cas extrêmes, la chute des pylônes. Dans ce cas, le ré-enclenchement se fait sur défaut et les disjoncteurs doivent être à nouveau ouverts dans les plus brefs délais.

12.1.4 Types de défaut

Les différents défauts qu'un système triphasé peut subir sont repris à la figure 12.1, on n'en considère pas les variantes de courts-circuits avec impédance, pour simplifier.

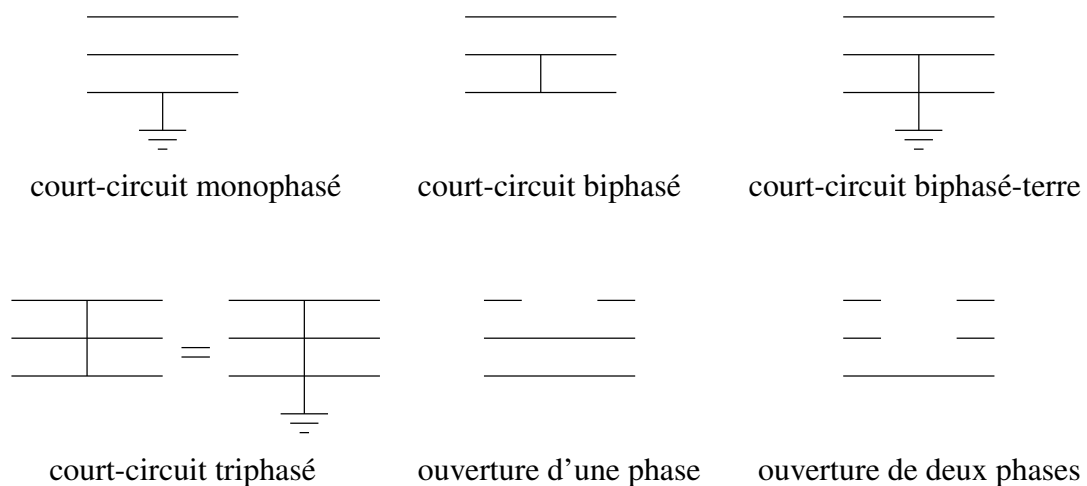


FIGURE 12.1 – les différents types de défaut

De tous les courts-circuits, le monophasé est le plus courant, puisque de 70 à 80 % des défauts sont de ce type. Le court-circuit triphasé ne se produit que dans environ 5 % des cas, mais il est le plus sévère et les équipements doivent pouvoir y faire face. Notons que si les trois phases sont court-circuitées, le système triphasé reste équilibré. Le point commun aux trois phases est virtuellement au potentiel nul et il est équivalent de considérer que le court-circuit s'est produit entre les phases et la terre.

Dans ce chapitre nous nous limitons au court-circuit triphasé, pour lequel une analyse par phase s'applique encore. Les autres types de défauts créent un déséquilibre, dont l'analyse fait intervenir les trois phases ; le chapitre suivant lui est consacré.

12.2 Comportement de la machine synchrone pendant un court-circuit

Les principaux composants responsables de la production des courants de court-circuit sont les générateurs synchrones. Dans cette section, nous considérons comment les représenter dans les études de courts-circuits équilibrés.

12.2.1 Expression du courant de court-circuit d'un générateur fonctionnant initialement à vide

Sur une période d'un ou deux dixièmes de seconde après apparition d'une perturbation, la vitesse de rotation d'une machine synchrone ne peut changer significativement, étant donné l'inertie mécanique des masses tournantes. Dans cet intervalle de temps, les transitoires sont essentiellement de nature électromagnétique ; ils proviennent des variations des flux magnétiques dans les divers enroulements de la machine.

Considérons le cas simple d'un générateur fonctionnant initialement à vide et soumis à l'instant $t = 0$ à un court-circuit triphasé sans impédance. La machine reçoit une tension d'excitation continue $v_f = R_f i_f^o$ et tourne à la vitesse de synchronisme :

$$\theta = \theta^o + \omega_N t$$

où θ^o est la position du rotor au moment où survient le court-circuit. La relation (8.43) donne l'amplitude de la tension aux bornes d'une des phases de la machine :

$$E_q^o = \frac{\omega_N L_{fd} i_f^o}{\sqrt{3}}$$

Considérons d'abord le cas où seul le circuit d'excitation est pris en compte au rotor. On peut établir l'expression analytique du courant de court-circuit en considérant les équations de Park, en leur appliquant la transformée de Laplace², en extrayant les expressions de $I_d(s)$ et $I_q(s)$, en revenant au domaine temporel pour obtenir $i_d(t)$ et $i_q(t)$ et enfin en employant la transformée de Park inverse pour obtenir les courants au stator. Ce développement analytique, assez long, et complété par quelques simplifications justifiées par les ordres de grandeurs des paramètres fournit les expressions suivantes pour le courant dans la phase a :

$$\begin{aligned} i_a(t) = & -\sqrt{2}E_q^o \left[\frac{1}{X_d} + \left(\frac{1}{X_d'} - \frac{1}{X_d} \right) e^{-t/T_d'} \right] \cos(\omega_N t + \theta_o) \\ & + \sqrt{2}E_q^o \frac{1}{2} \left(\frac{1}{X_d'} - \frac{1}{X_q} \right) e^{-t/T_\alpha} \cos(2\omega_N t + \theta_o) + \sqrt{2}E_q^o \frac{1}{2} \left(\frac{1}{X_d'} + \frac{1}{X_q} \right) e^{-t/T_\alpha} \cos \theta_o \end{aligned} \quad (12.1)$$

2. le fait que l'on suppose la vitesse de rotation constante supprime une non-linéarité majeure en présence de laquelle il ne serait pas possible d'utiliser la transformée de Laplace

et pour le courant d'excitation :

$$i_f(t) = i_f^o + \frac{X_d - X'_d}{X'_d} i_f^o e^{-t/T'_d} - \frac{X_d - X'_d}{X'_d} i_f^o e^{-t/T_\alpha} \cos \omega_N t \quad (12.2)$$

Ces expressions font intervenir :

- la constante de temps du circuit d'excitation lorsque le stator est court-circuité :

$$T'_d = \frac{L_{ff} - \frac{L_{fd}^2}{L_{dd}}}{R_f} \quad (12.3)$$

On notera que si le stator était ouvert, la constante de temps du même enroulement (qui n'interagirait alors avec aucun autre circuit) serait :

$$T'_{do} = \frac{L_{ff}}{R_f} \quad (12.4)$$

La constante de temps est donc plus petite lorsque le stator est court-circuité :

$$T'_d < T'_{do} \quad (12.5)$$

- la réactance *transitoire dans l'axe direct* :

$$X'_d = \omega_N L'_d$$

elle-même fonction de l'*inductance transitoire dans l'axe direct* :

$$L'_d = L_{dd} - \frac{L_{fd}^2}{L_{ff}} \quad (12.6)$$

En utilisant (12.3) et (12.6) on établit aisément que :

$$L'_d = L_{dd} \frac{T'_d}{T'_{do}} \quad \text{et donc} \quad X'_d = X_d \frac{T'_d}{T'_{do}} \quad (12.7)$$

On voit aisément que la réactance transitoire est plus petite que la réactance synchrone.

- la constante de temps statorique :

$$T_\alpha = \frac{2}{R_a} \frac{1}{\frac{1}{L'_d} + \frac{1}{L_{qq}}} \quad (12.8)$$

12.2.2 Interprétation physique de l'évolution du courant

Les différentes composantes du courant i_a se justifient comme suit.

Avant apparition du court-circuit, l'enroulement statorique a est le siège d'un flux alternatif $\psi_{af}(t)$ créé par l'enroulement d'excitation en mouvement. Lors de l'application du défaut, ce circuit est refermé sur lui-même et un courant peut y circuler. En vertu de la loi de Lenz, ce courant est tel que,

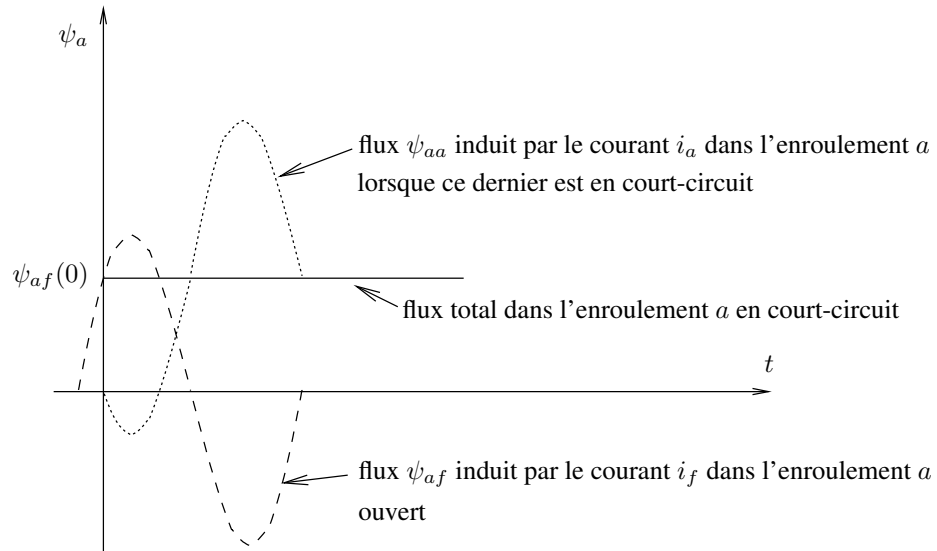


FIGURE 12.2 – flux dans l'enroulement a avec et sans court-circuit

dans les premiers instants, le flux dans l'enroulement reste constant, égal à la valeur $\psi_{af}(0)$ qu'il avait au moment où le court-circuit est apparu. Plus précisément, ce courant produit un flux ψ_{aa} qui s'oppose aux variations de flux que tente d'imposer le circuit d'excitation en mouvement. La situation est représentée à la figure 12.2. Pour produire ce flux ψ_{aa} , le courant induit dans la bobine a doit comporter une composante *unidirectionnelle* et une composante alternative de pulsation ω_N .

Les composantes alternatives des courants induits dans les trois phases sont de même amplitude mais déphasées de 120 degrés électriques les unes par rapport aux autres. Ensemble, elles produisent un champ magnétique H_{ac} tournant à la même vitesse que le rotor. Ce champ est dirigé selon l'axe direct et dans le sens opposé au champ produit par le courant d'excitation i_f^o .

Les composantes unidirectionnelles des courants induits au stator diffèrent d'une phase à l'autre car les trois phases embrassent des flux différents à l'instant $t = 0$. Ensemble, ces composantes créent un champ magnétique H_{dc} fixe par rapport au stator, c'est-à-dire tournant à la vitesse ω_N par rapport au rotor.

Dans un court intervalle de temps après l'apparition du court-circuit, le flux dans l'enroulement d'excitation ne peut pas non plus changer. Un courant unidirectionnel va donc y être induit pour créer un champ qui s'oppose au champ H_{ac} provenant du stator, et un courant alternatif pour s'opposer au champ H_{dc} . On retrouve bien ces deux composantes dans l'expression (12.2).

Suite à la dissipation d'énergie dans les résistances, les flux, tant statoriques que rotorique, ne restent pas constants :

- les composantes unidirectionnelles des courants statoriques décroissent avec la constante de temps T_α . Cette décroissance est également celle de l'enveloppe de la composante alternative du courant d'excitation ;
- la composante unidirectionnelle du courant d'excitation décroît avec la constante de temps T_d' . Cette décroissance est également celle de l'enveloppe de la composante alternative des courants

statoriques.

Comme mentionné plus haut, le champ magnétique H_{ac} associé aux composantes alternatives est dirigé selon l'axe direct de la machine. Ceci explique pourquoi seules les réactances dans l'axe direct interviennent dans les expressions de ces composantes.

Enfin, le chemin offert aux lignes du champ magnétique H_{dc} , fixe par rapport au stator, comporte en fait un entrefer de largeur variable, suivant la position du rotor. Ceci se marque de deux manières :

- la composante unidirectionnelle du courant i_a fait intervenir la moyenne entre $1/X'_d$, valeur correspondant à l'alignement de l'axe direct avec celui de la phase a , et $1/X_q$, valeur correspondant à l'alignement de l'axe en quadrature avec celui de la phase a ;
- l'apparition d'une composante alternative à la pulsation $2\omega_N$. Cette pulsation se justifie par le fait que quand le rotor a fait un *demi* tour, la largeur de l'entrefer est de nouveau la même.

12.2.3 Expressions tenant compte des autres enroulements rotoriques

Lorsque l'on prend en compte les autres enroulements rotoriques, on aboutit à l'expression plus précise du courant statorique que voici :

$$i_a(t) = -\sqrt{2}E_q^o \left[\frac{1}{X_d} + \left(\frac{1}{X'_d} - \frac{1}{X_d} \right) e^{-t/T'_d} + \left(\frac{1}{X''_d} - \frac{1}{X'_d} \right) e^{-t/T''_d} \right] \cos(\omega_N t + \theta_o) \quad (12.9)$$

$$+ \sqrt{2}E_q^o \frac{1}{2} \left(\frac{1}{X''_d} - \frac{1}{X'_d} \right) e^{-t/T_\alpha} \cos(2\omega_N t + \theta_o) + \sqrt{2}E_q^o \frac{1}{2} \left(\frac{1}{X''_d} + \frac{1}{X'_d} \right) e^{-t/T_\alpha} \cos \theta_o$$

dans laquelle :

- X''_d est la *réactance subtransitoire dans l'axe direct*. Cette réactance provient de la réaction de l'amortisseur modélisé par le circuit d_1 . On a nécessairement :

$$X''_d < X'_d < X_d$$

- T''_d est la constante de temps *subtransitoire*, associée elle aussi à l'amortisseur dans l'axe direct. Cette constante de temps est plus petite que T'_d ;
- X''_q est la *réactance subtransitoire dans l'axe en quadrature*. Cette réactance provient de la réaction de l'amortisseur modélisé par le circuit q_2 .

On voit que la présence des amortisseurs modifie l'amplitude de :

- la composante alternative du courant
- la composante unidirectionnelle. En pratique $X''_d \simeq X'_d$ et l'amplitude vaut plus simplement :

$$\sqrt{2} \frac{E_q^o}{X'_d}$$

- la composante à $2\omega_N$. Compte tenu de l'approximation ci-dessus, l'amplitude de cette composante est très faible en pratique et peut être négligée.

Enfin, une expression plus précise pour la constante de temps T_α est :

$$T_\alpha = \frac{X_d''}{\omega_N R_a} \quad (12.10)$$

Le tableau ci-dessous donne l'ordre de grandeur des diverses réactances et constantes de temps apparaissant plus haut.

	machine à			machine à	
	rotor lisse (pu)	pôles saillants (pu)		rotor lisse (s)	pôles saillants (s)
X_d'	0.2-0.4	0.3-0.5	T_{do}'	8.0-12.0	3.0-8.0
X_d''	0.15-0.30	0.25-0.35	T_d'	0.95-1.30	1.0-2.5
X_q''	0.15-0.30	0.25-0.35	T_d''	0.02-0.05	0.02-0.05
			T_α	0.02-0.60	0.02-0.20

12.2.4 Exemple numérique et discussion

A titre d'illustration, considérons une machine caractérisée par :

$$E_q^o = 1, \quad X_d = X_q = 2, \quad X_d' = 0.3, \quad X_d'' = X_q'' = 0.2, \quad R_a = 0.005 \text{ pu}$$

$$T_{do}' = 9, \quad T_d'' = 0.0333 \text{ s}$$

On en déduit :

$$\begin{aligned} T_d' &= T_{do}' \frac{X_d'}{X_d} = 1.35 \text{ s} \\ T_\alpha &= \frac{X_d''}{R_a} = 40 \text{ pu} \\ &= \frac{40}{2\pi 50} = 0.127 \text{ s} \end{aligned} \quad (12.11)$$

On suppose que le court-circuit se produit au moment où l'axe direct coïncide avec l'axe de la phase a , c'est-à-dire que $\theta^o = 0$. Dans ces conditions, la valeur initiale de la composante unidirectionnelle est maximale dans la phase a . Celles dans les phases b et c sont négatives, égales et d'amplitude plus faible.

Il importe de noter que les courbes ci-après se rapportent à un court-circuit permanent et à une vitesse de rotation constante. En pratique, le court-circuit est éliminé par les protections après le délai déjà mentionné, tandis que la vitesse varie sous l'effet du déséquilibre entre couples mécanique et électromagnétique (au point que si le défaut est éliminé trop tard, la machine perd le synchronisme). Les courbes ne peuvent donc être utilisées que sur le court intervalle de temps correspondant au court-circuit.

La figure 12.3 montre l'évolution du courant i_a sur un intervalle de temps de l'ordre de 15 fois T_d'' ou un tiers de T_d' . On voit que la composante unidirectionnelle retarde légèrement le premier passage par zéro du courant, nécessaire à la coupure par le disjoncteur. Elle peut aussi provoquer la saturation du noyau magnétique des transformateurs de mesure utilisés par les protections.

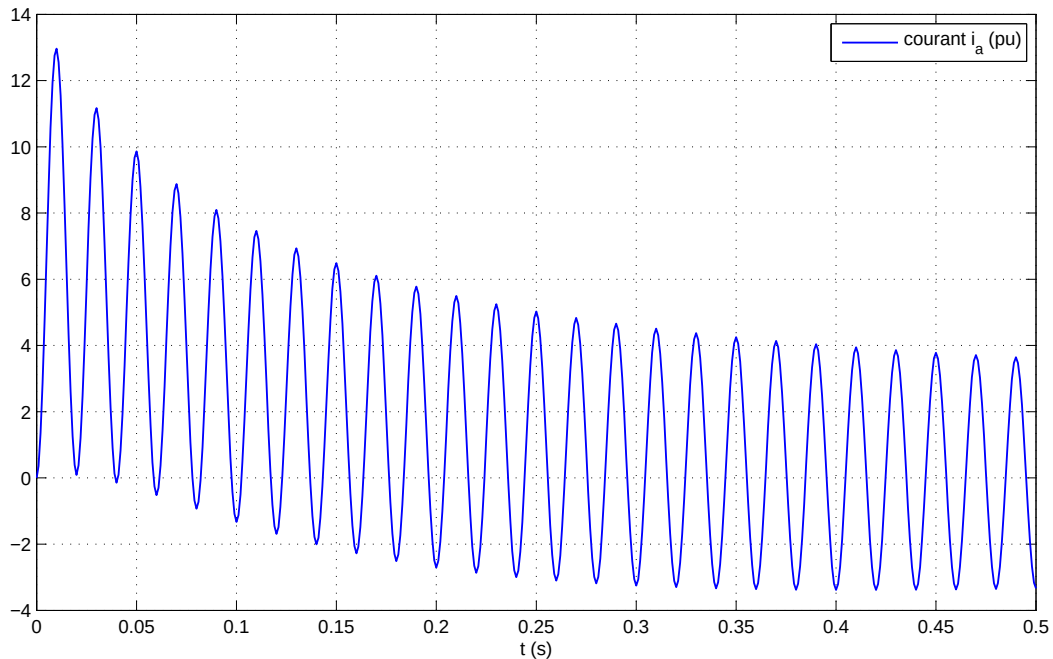


FIGURE 12.3 – évolution du courant de court-circuit dans la phase a

Les courants dans les phases a et b sont comparés à la figure 12.4. On voit qu'une fois éteints les transitoires initiaux, i_b devient égal à i_a déphasé de 120 degrés.

La figure 12.5 montre séparément les composantes alternative et unidirectionnelle du courant i_a . Elles sont initialement de valeurs opposées, ce qui conduit à un courant initialement nul.

Enfin, la figure 12.6 montre l'évolution de l'amplitude de la composante alternative du courant i_a . La ligne horizontale en pointillé donne l'amplitude vers laquelle tend le courant, soit $\sqrt{2}E_q^o/X_d$. On voit qu'avant d'atteindre cette valeur, l'amplitude du courant de défaut est nettement plus élevée, sous l'effet du circuit d'excitation. La relation (12.1) montre en effet qu'en $t = 0$, l'amplitude vaut $\sqrt{2}E_q^o/X_d'$. L'expression plus précise (12.9) montre que sous l'effet supplémentaire des amortisseurs, l'amplitude initiale est $\sqrt{2}E_q^o/X_d''$, soit une valeur encore un peu plus élevée. La figure 12.6 montre l'accroissement correspondant du courant de défaut.

Il résulte de ceci que les disjoncteurs sont appelés à couper un courant nettement plus important que celui qu'on aurait en régime établi, ce qui doit être pris en compte dans leur dimensionnement.

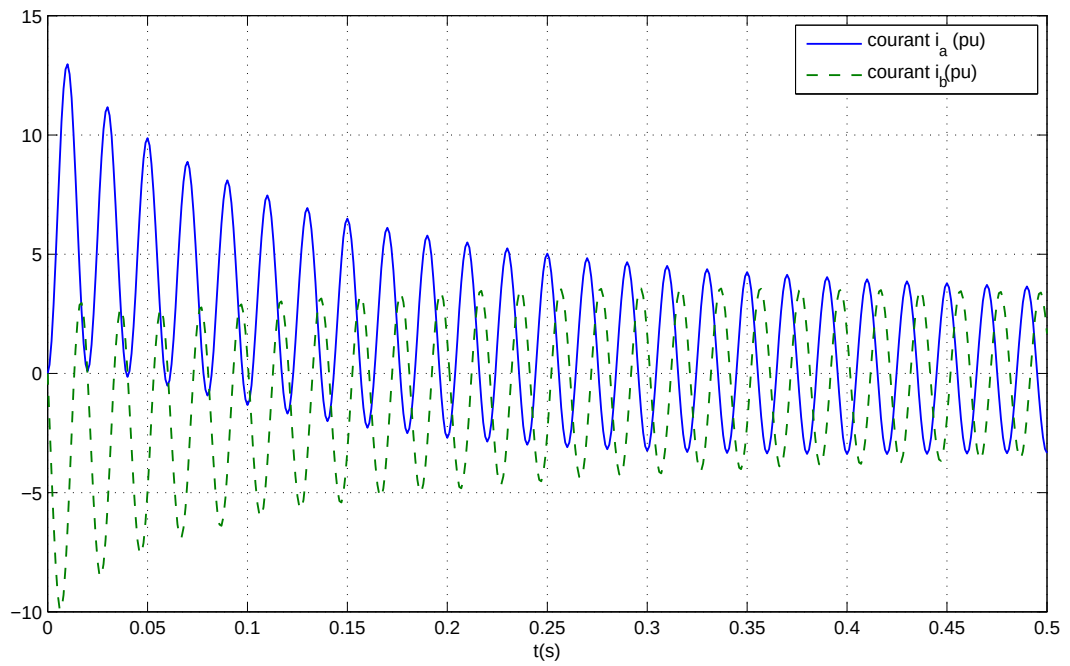


FIGURE 12.4 – évolution des courants de court-circuit dans les phases a et b

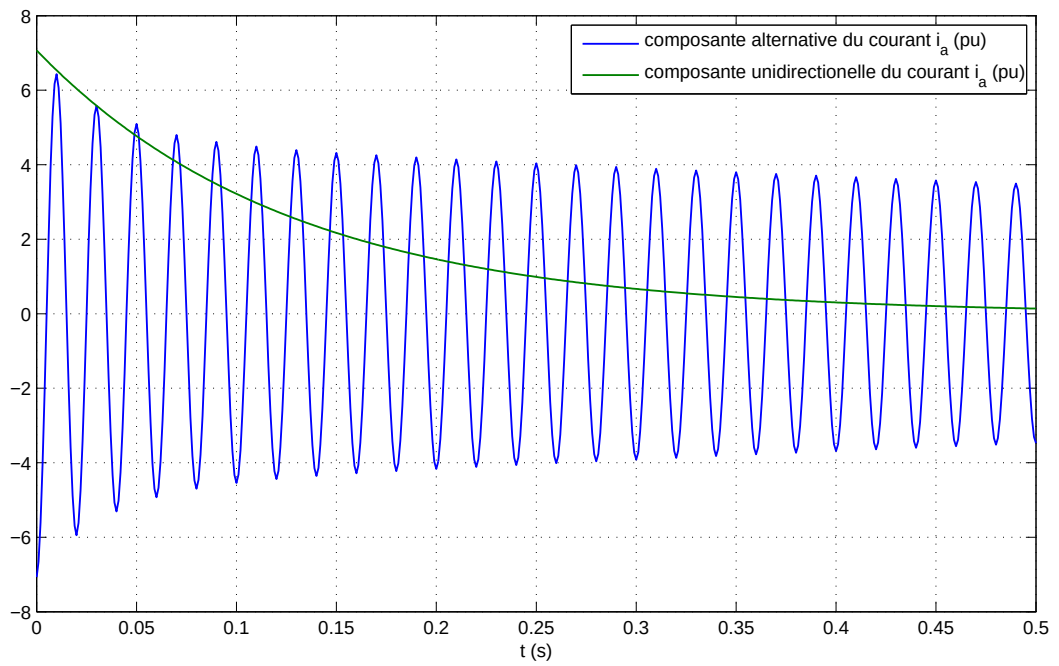


FIGURE 12.5 – composantes alternative et unidirectionnelle du courant dans la phase a

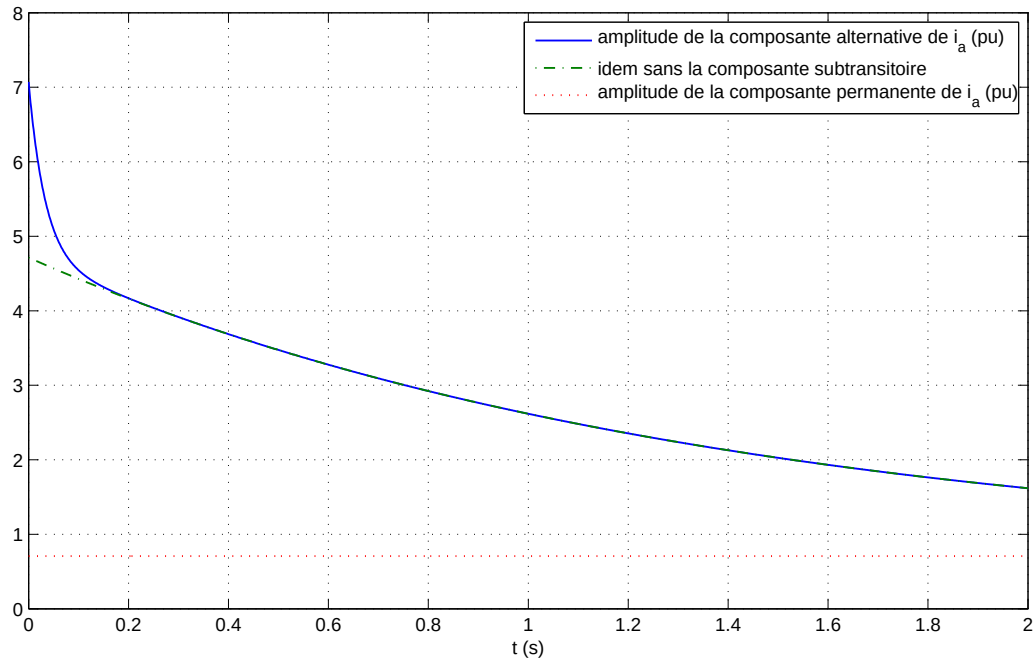


FIGURE 12.6 – évolution de l’amplitude de la composante alternative du courant dans la phase a

12.2.5 Effet de la localisation du défaut

En pratique, il est extrêmement rare d’avoir un court-circuit aux bornes d’un générateur, ce dernier étant abrité dans la centrale. L’endroit le plus proche du générateur où un court-circuit est susceptible de survenir est le poste où se trouve le transformateur élévateur du générateur.

Pour étudier l’effet de la localisation du court-circuit, il convient d’interposer, dans chaque phase, un dipôle (R_e, L_e) entre la machine et le court-circuit. Compte tenu des valeurs typiques du rapport L_e/R_e , on montre que :

- plus le défaut est éloigné de la machine, plus la composante alternative du courant décroît lentement ; la constante de temps se situe entre T'_d (court-circuit aux bornes de la machine) et T'_{do} (court-circuit infiniment éloigné de celle-ci) ;
- plus le défaut est éloigné de la machine, plus la composante unidirectionnelle décroît rapidement.

12.2.6 Simplifications usuelles pour le calcul des courants de court-circuit

Les calculs de courants de court-circuit usuels négligent :

- les composantes unidirectionnelles des courants produits par les machines synchrones. En effet, ces composantes décroissent assez rapidement, d’autant plus que le défaut est éloigné des machines. Toutefois, on peut compenser cette approximation en multipliant le courant calculé sans cette composante par un facteur empirique supérieur à l’unité, afin de se placer en sécurité ;

- les composantes alternatives de pulsation $2\omega_N$ des courants, négligeables pour la raison mentionnée précédemment.

Les calculs portent donc sur les composantes alternatives de pulsation ω_N , ce qui permet de calculer les courants de défaut via les techniques (mais pas les paramètres !) s'appliquant au régime sinusoïdal établi.

Dans les réseaux de transport, compte tenu de la rapidité des disjoncteurs, on considère qu'il faut pouvoir couper la valeur initiale de cette composante, ce qui revient à considérer la réactance subtransitoire des machines dans les calculs. Ceci procure une marge de sécurité puisque ultérieurement l'amplitude du courant de défaut décroît.

Dans les réseaux de distribution, les disjoncteurs sont moins rapides et l'on considère généralement la réactance transitoire dans les calculs. Dans ce cas, il est encore plus légitime de négliger la composante unidirectionnelle du courant de défaut.

12.2.7 Schéma équivalent simplifié d'une machine synchrone

L'analyse du court-circuit d'un générateur initialement à vide a montré que ce dernier se comporte comme une f.e.m. d'amplitude $\sqrt{2}E_q''$ derrière la réactance subtransitoire X_d'' . Qu'en est-il dans le cas usuel où le générateur produit un courant avant apparition du défaut ?

En fait, on peut montrer que la machine synchrone obéit au schéma équivalent de la figure 12.7, dans lequel la f.e.m. $\bar{E}''$ est constante. En effet, on démontre (cf cours ELEC0047) que cette f.e.m. est proportionnelle aux flux dans les enroulements rotoriques de la machine, lesquels ne changent dans les premiers instants qui suivent le court-circuit. En pratique, on appelle $\bar{E}''$ la f.e.m. derrière réactance subtransitoire. La résistance statorique a été négligée.

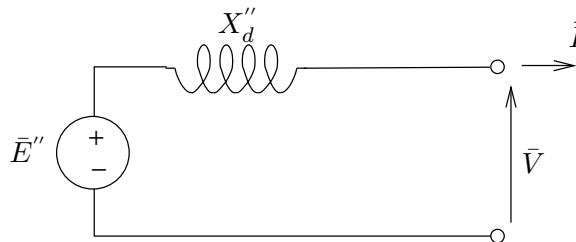


FIGURE 12.7 – schéma équivalent d'une machine synchrone

Supposant le court-circuit appliqué en $t = 0$, la continuité de cette f.e.m. se traduit par :

$$\bar{E}''(0^+) = \bar{E}''(0^-) = \bar{V}(0^-) + jX_d''\bar{I}(0^-) \quad (12.12)$$

On peut donc déterminer cette f.e.m. au départ d'un calcul de load flow pré-incident fournissant $\bar{V}(0^-)$ et $\bar{I}(0^-)$. On retrouve évidemment le cas du générateur initialement à vide en posant $\bar{I}(0^-) = 0$ d'où $\bar{E}''(0^-) = \bar{V}(0^-) = E_q^o$.

12.3 Calcul des courants de court-circuit triphasé

Il importe d'évaluer l'amplitude des courants de défaut pour :

- dimensionner les disjoncteurs, qui doivent avoir un pouvoir de coupure suffisant pour interrompre les courants en question ;
- régler les protections commandant les disjoncteurs. Celles-ci doivent impérativement agir lorsqu'il leur incombe d'éliminer le défaut mais ne doivent pas agir intempestivement lorsque ce n'est pas nécessaire.

12.3.1 Hypothèses de calcul

Considérons un réseau à N noeuds, comportant des lignes, des câbles et des transformateurs. Soit n le nombre de machines synchrones. Sans perte de généralité, nous supposons que les noeuds du réseau sont numérotés en réservant les n premiers numéros aux noeuds de connexion des machines synchrones.

Chaque machine synchrone est représentée par le schéma équivalent de Thévenin de la figure 12.7.

Chaque charge est supposée se comporter à admittance constante. Cette hypothèse est raisonnable pour de nombreux équipements, dans les premiers instants qui suivent un court-circuit. L'admittance équivalente Y_c d'une charge peut se calculer à partir des puissances active $P(0^-)$ et réactive $Q(0^-)$ qu'elle consomme et de la tension $V(0^-)$ à ses bornes, toutes grandeurs relatives à la situation avant court-circuit³ :

$$Y_c = \frac{P(0^-) - jQ(0^-)}{[V(0^-)]^2}$$

Nous supposons qu'un court-circuit d'impédance Z_f se produit au noeud f du réseau. On a donc :

$$\bar{V}_f = Z_f \bar{I}_f \quad (12.13)$$

où $\bar{V}_f$ est la tension au noeud f pendant le défaut et $\bar{I}_f$ est le courant de défaut. Z_f est nul dans le cas d'un court-circuit franc.

Le système se présente comme indiqué à la figure 12.8.a. En fait, la formulation qui suit s'accommode mieux d'un équivalent de Norton pour chaque générateur. Ceci conduit au schéma de la figure 12.8.b, auquel nous nous référerons dans ce qui suit. Le courant $\bar{I}_f$ est compté positivement lorsqu'il sort du réseau.

3. en MW et Mvar, P et Q représentent les puissances consommées par phase ; en per unit, ils représentent la puissance triphasée. La notation (0^-) évoque la situation juste avant l'apparition du défaut en $t = 0$

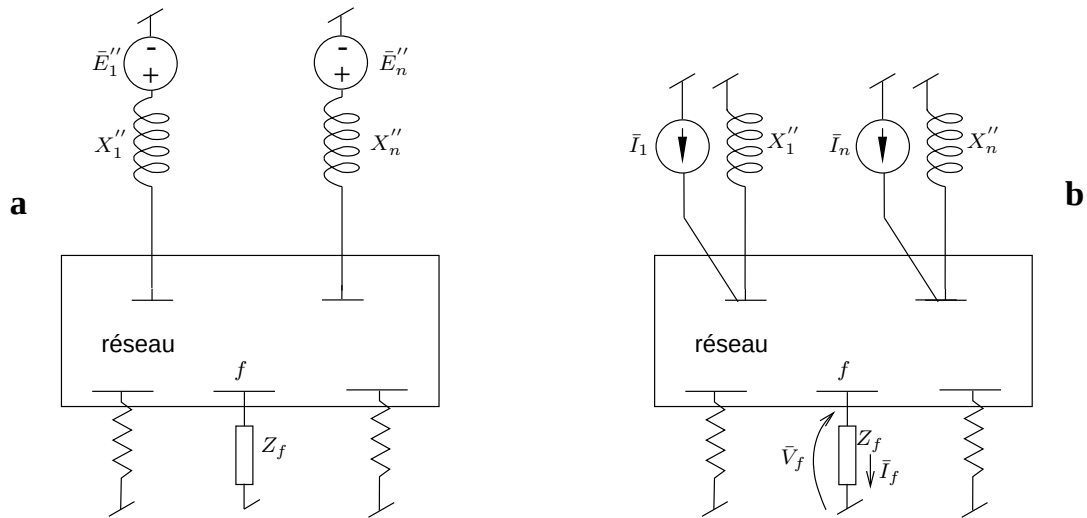


FIGURE 12.8 – réseau soumis à un court-circuit, machines représentées par leur schémas équivalents de Thévenin et de Norton

12.3.2 Matrice d'admittance et équations du réseau

Considérons un réseau à N noeuds comprenant des lignes, des câbles et des transformateurs. Ajoutons-y :

- aux noeuds générateurs : les admittances $1/jX''$
 - aux noeuds charges : les admittances Y_c
 - aux noeuds avec de la compensation shunt : les admittances jB_c correspondantes.
- comme représenté à la figure 12.8.b.

Soit $\mathbf{Y}$ la matrice d'admittance du circuit ainsi obtenu. La terre est prise comme noeud de référence (non compté parmi les N). La méthode la plus générale pour construire $\mathbf{Y}$ est la suivante :

- partir d'une matrice carrée, de dimension N , composée entièrement de zéros
- déterminer les matrices d'admittance des quadripôles décrivant les lignes, les câbles et les transformateurs. Ces matrices de dimension 2 ont été détaillées aux chapitres 4 et 6
- ajouter les termes de chacune de ces matrices aux termes appropriés de $\mathbf{Y}$, en fonction de la numérotation adoptée pour les N noeuds. Un exemple pour un réseau à 4 noeuds est donné à la figure 12.9
- ajouter aux termes diagonaux les contributions des machines, des charges et de la compensation shunt.

Il est très aisé d'implémenter ces règles dans un logiciel de calcul.

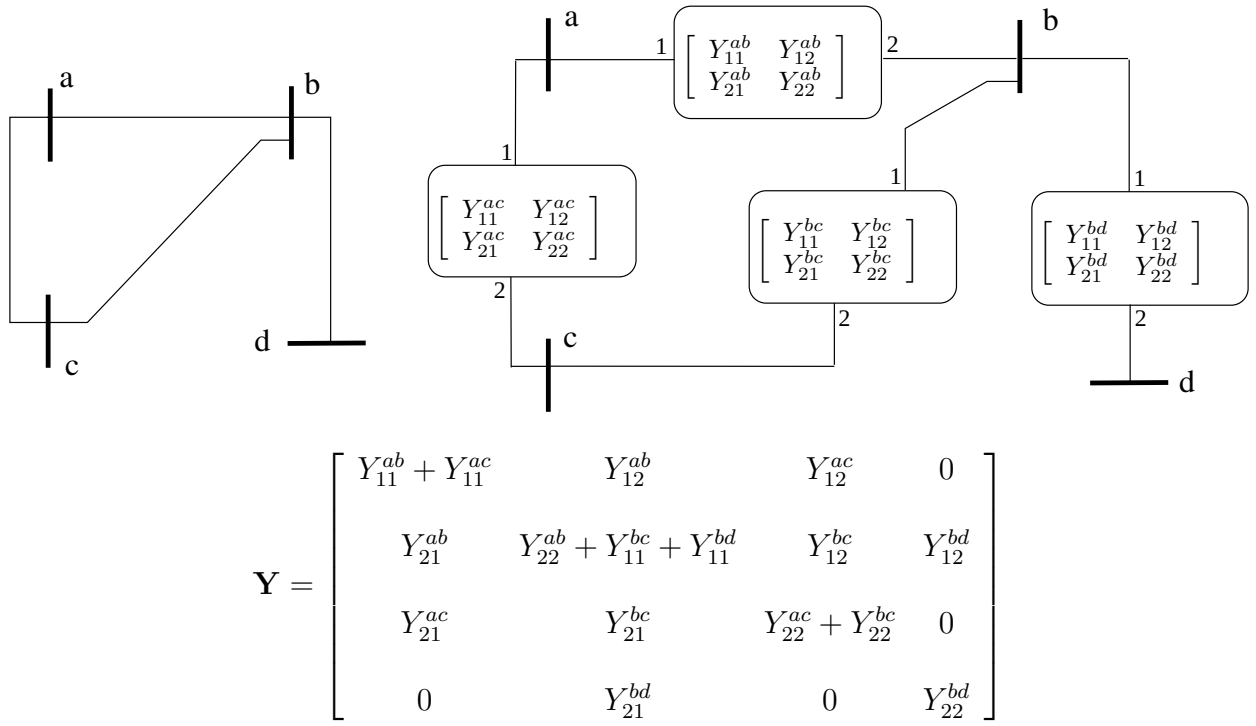


FIGURE 12.9 – Exemple de construction de la matrice d’admittance $\mathbf{Y}$

Dans le cas où il n’y a pas de transformateurs déphaseurs, la matrice $\mathbf{Y}$ peut aussi être construite “par inspection” comme suit :

- les transformateurs sont représentés par le schéma équivalent de la figure 6.5, dans lequel n est réel. Ce schéma peut, à son tour, être remplacé par le schéma équivalent en pi de la figure 6.7
- les lignes et les câbles, quant à eux, sont représentés par le schéma équivalent en pi de la figure 4.6
- ces différents schémas en pi sont assemblés conformément à la topologie du réseau
- à cet ensemble, on ajoute les admittances shunt provenant des machines, des charges et de la compensation shunt
- les termes de $\mathbf{Y}$ s’obtiennent un à un comme suit :
 - le terme non diagonal $[\mathbf{Y}]_{ij} (i \neq j)$ est la somme de toutes les admittances joignant les noeuds i et j , changée de signe ;
 - le terme diagonal $[\mathbf{Y}]_{ii}$ est la somme de toutes les admittances connectées au noeud i .

Notons qu’en l’absence de transformateurs déphaseurs, $\mathbf{Y}$ est symétrique. Elle est non singulière pour autant qu’il existe au moins un élément shunt dans chaque partie connexe du schéma unifilaire du réseau.

Les équations du réseau s’écrivent :

$$\bar{\mathbf{I}} = \mathbf{Y} \bar{\mathbf{V}} \quad (12.14)$$

où $\bar{\mathbf{I}}$ est le vecteur des courants injectés aux N noeuds et $\bar{\mathbf{V}}$ celui des tensions à ces mêmes noeuds. Toutes les composantes de ces vecteurs sont des nombres complexes. Les courants sont considérés comme positifs quand ils *entrent* dans le réseau.

Notons que l'impédance de défaut Z_f a été conservée à l'extérieur du circuit modélisé par $\mathbf{Y}$ et n'intervient donc pas dans cette matrice. On pourrait l'inclure en ajoutant $Y_f = 1/Z_f$ au terme diagonal correspondant de $\mathbf{Y}$. Cependant, cela présente deux inconvénients pratiques :

- on est amené à calculer les courants de défaut en de nombreux noeuds. Incorporer Y_f à la matrice d'admittance $\mathbf{Y}$ requiert de modifier celle-ci pour chaque défaut ;
- on considère fréquemment des défauts francs (comme cas pessimiste). Ceci correspond à une valeur infinie pour Y_f , ce qui n'est pas compatible avec l'utilisation de $\mathbf{Y}$. Notons néanmoins qu'en pratique, on peut donner une très grande valeur à Y_f , ce qui donne une tension quasiment nulle.

La formulation qui suit permet de n'utiliser que la seule matrice $\mathbf{Y}$ relative à la configuration saine (sans défaut) du réseau et s'applique au cas particulier où $Z_f = 0$.

12.3.3 Calcul des tensions pendant défaut par superposition

Nous allons calculer les tensions pendant défaut. A partir de celles-ci il est possible de calculer le courant dans n'importe quelle branche du réseau.

Sous l'effet des courants $\bar{I}_1, \bar{I}_2, \dots, \bar{I}_n$ injectés par les générateurs et du courant $\bar{I}_f$ soutiré au noeud f (cf figure 12.8.b), les tensions prennent une valeur $\bar{\mathbf{V}}$ qui satisfait à :

$$\mathbf{Y} \bar{\mathbf{V}} = \begin{bmatrix} \bar{I}_1 \\ \bar{I}_2 \\ \vdots \\ \bar{I}_n \\ 0 \\ \vdots \\ \vdots \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ \vdots \\ \vdots \\ 0 \\ -\bar{I}_f \\ 0 \\ \vdots \\ 0 \end{bmatrix}$$

Par superposition, la solution $\bar{\mathbf{V}}$ est la somme de deux termes :

$$\bar{\mathbf{V}} = \bar{\mathbf{V}}(0^-) + \Delta \bar{\mathbf{V}} \quad (12.15)$$

avec

$$\mathbf{Y} \bar{\mathbf{V}}(0^-) = \begin{bmatrix} \bar{I}_1 \\ \bar{I}_2 \\ \vdots \\ \bar{I}_n \\ 0 \\ \vdots \\ \vdots \\ 0 \end{bmatrix}$$

et

$$\mathbf{Y} \Delta \bar{\mathbf{V}} = -\bar{I}_f \begin{bmatrix} 0 \\ \vdots \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix} = -\bar{I}_f \mathbf{e}_f \quad (12.16)$$

$\bar{\mathbf{V}}(0^-)$ n'est rien d'autre que le vecteur des tensions aux noeuds avant l'application du défaut (c'est-à-dire en $t = 0^-$). $\Delta \bar{\mathbf{V}}$ apparaît donc comme une correction représentant l'effet du court-circuit. $\mathbf{e}_f$ est un vecteur unitaire dont toutes les composantes sont nulles, à l'exception de la f -ème qui vaut 1.

A ce stade, on ne connaît pas la valeur de $\bar{I}_f$. Provisoirement, résolvons le système linéaire :

$$\mathbf{Y} \Delta \bar{\mathbf{V}}^{(1)} = \mathbf{e}_f \quad (12.17)$$

Les membres de droite des systèmes linéaires (12.16) et (12.17) diffèrent par le facteur $-\bar{I}_f$. On a donc :

$$\Delta \bar{\mathbf{V}} = -\bar{I}_f \Delta \bar{\mathbf{V}}^{(1)}$$

En introduisant ce résultat dans (12.15) on obtient :

$$\bar{\mathbf{V}} = \bar{\mathbf{V}}(0^-) - \bar{I}_f \Delta \bar{\mathbf{V}}^{(1)} \quad (12.18)$$

La f -ème composante de cette relation vectorielle s'écrit :

$$\bar{V}_f = \bar{V}_f(0^-) - \bar{I}_f \Delta \bar{V}_f^{(1)}$$

En combinant cette relation avec (12.13), on obtient enfin la valeur du courant de défaut :

$$\bar{I}_f = \frac{\bar{V}_f(0^-)}{\Delta \bar{V}_f^{(1)} + Z_f} \quad (12.19)$$

et en remplaçant dans (12.18, 12.15), on trouve les tensions recherchées :

$$\bar{\mathbf{V}} = \bar{\mathbf{V}}(0^-) - \frac{\bar{V}_f(0^-)}{\Delta \bar{V}_f^{(1)} + Z_f} \Delta \bar{\mathbf{V}}^{(1)}$$

12.3.4 Relation avec le schéma équivalent de Thévenin

Au chapitre 3, nous avons mentionné le lien entre le courant de court-circuit, l'impédance de Thévenin et la puissance de court-circuit. Considérons le schéma équivalent de Thévenin du réseau, dans sa configuration avant court-circuit et vu du jeu de barres f (cf figure 12.10).

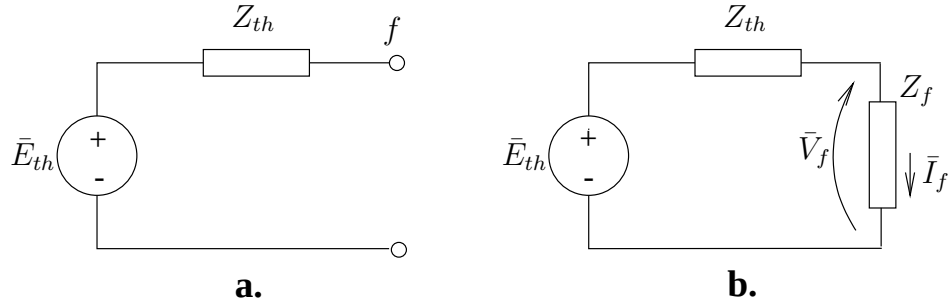


FIGURE 12.10 – schéma équivalent de Thévenin sans et avec défaut

La f.e.m. de Thévenin est donnée par :

$$\bar{E}_{th} = \bar{V}_f(0^-) \quad (12.20)$$

L'impédance de Thévenin peut être obtenue en injectant un courant unitaire au noeud f du réseau passifié et en relevant la tension apparaissant en ce noeud. Dans ces conditions, les tensions aux noeuds valent :

$$\mathbf{Y}^{-1} \mathbf{e}_f = \Delta \bar{\mathbf{V}}^{(1)}$$

et celle au noeud f vaut :

$$Z_{th} = \Delta \bar{V}_f^{(1)} = [\mathbf{Y}^{-1} \mathbf{e}_f]_f = [\mathbf{Y}^{-1}]_{ff} = [\mathbf{Z}]_{ff}$$

où $\mathbf{Z}$ est la *matrice d'impédance* aux jeux de barres. L'impédance de Thévenin vue d'un noeud f est donc égale au terme diagonal (f, f) de la matrice d'impédance.

Le courant de défaut s'obtient très aisément à partir du schéma de la figure 12.10.b :

$$\bar{I}_f = \frac{\bar{E}_{th}}{Z_{th} + Z_f} = \frac{\bar{V}_f(0^-)}{[\mathbf{Z}]_{ff} + Z_f} = \frac{\bar{V}_f(0^-)}{\Delta \bar{V}_f^{(1)} + Z_f} \quad (12.21)$$

On retrouve bien l'expression (12.19).

12.4 Appendice. Phénomènes physiques liés à la foudre

Le texte de cet appendice est tiré de l'ouvrage :

Electric Power System essentials, P. Schavemaker, L. van der Sluis, John Wiley & Sons, 2008, ISBN 978-0470-51027-8

The probability of a direct hit by lightning on an overhead transmission system is high compared with the vulnerability of other parts of the power system to this force of nature. For this reason, transmission line towers are interconnected by ground wires, also named shield wires, that hang well above the phase conductors in the tower and are electrically connected with the tower frame and via the towers to the ground below. The ground wires shield

the phase wires from a direct hit by lightning and reduce the tower ground resistance in dry or rocky soil.

Lightning is a part of the Earth-ionosphere electric system, depicted in Figure 3.43. The voltage between the ionosphere and the Earth surface, in fact an enormous capacitor, is roughly 300000 V. This 'capacitor' has a small leakage current, the fair-weather current, which totals about 1400 A when we take the complete Earth surface into account. Thunderstorms act as the 'battery' in this electric system; the number of thunderstorms, any time of the day, all over the

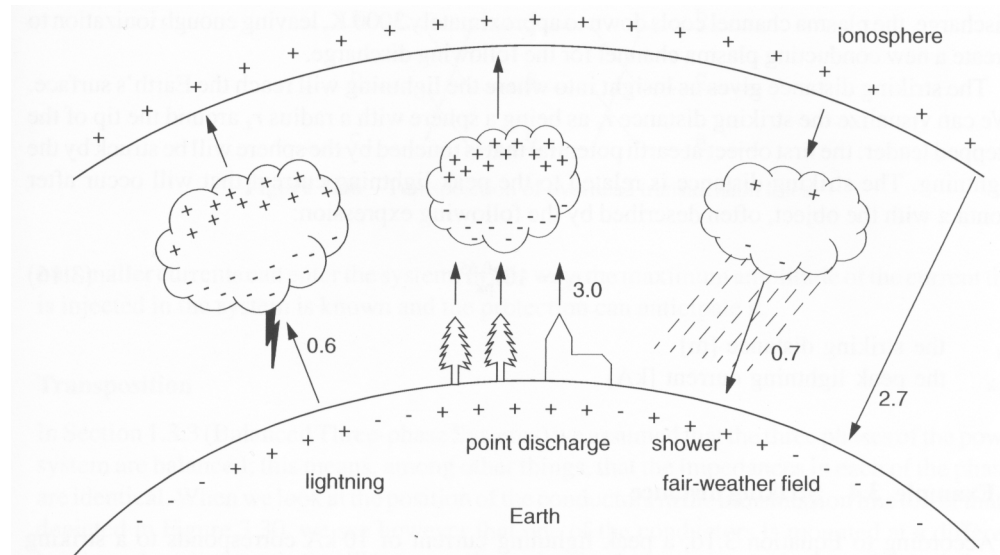


Figure 3.43 The Earth-ionosphere electric system [38]. Currents are in $\mu\text{A}/\text{km}^2$; average of the total Earth, over a long time.

world, is approximately 1500. The power in the Earth-ionosphere electric system is, however, only 500 MW, equivalent to the power of a small generator station.

Lightning mostly occurs on summer days when the ambient temperature is high and the air is humid. Because of the temperature difference, and with that the difference in density, the humid air is lifted to higher altitudes with a considerable lower ambient temperature. Cold air can contain less water than warm air and raindrops are formed. The raindrops have a size of a few millimeters and are polarized by the electric field that is present between the lower part of the ionosphere and the Earth's surface (the strength of this atmospheric field is on summer days in the order of 60 V/m and can reach values of 500 V/m on a dry winter day). The charge separation in the clouds, that you can see in Figure 3.43, is a complex mechanism which is not yet fully understood. Here follows one 'explanation'. Larger raindrops fall down and meet smaller and lighter raindrops that are lifted up by the upward flow of air. These smaller raindrops can be positively or negatively charged. The negative drops are attracted by the positive charge in the lower part of the large raindrop, whereas the positive drops are pushed away. Even though the positive drops are attracted by the negative charge in the upper part of the large raindrop, the small positive raindrops can not attach anymore. This mechanism makes that the larger raindrops, in the lower part of the cloud, are negatively charged and that most of the small raindrops in the top of the cloud have a positive charge. This process is illustrated in Figure 3.44.

The clouds move at great heights and the average field strength is far below the average breakdown strength of air. Inside the thundercloud, the space charge formed by the accumulating negative raindrops create a locally strong electric field in the order of 10000 V/m and this

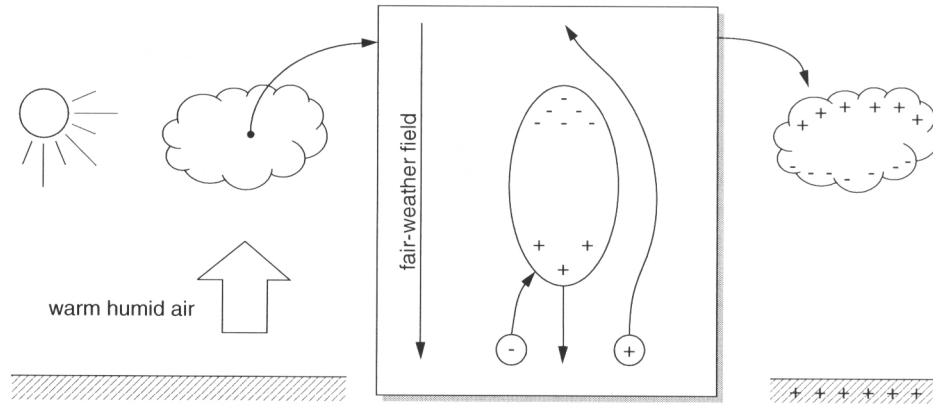


Figure 3.44 Charge separation in clouds. The highlighted box in the middle shows the interaction between a large polarized raindrop falling down and two smaller charged drops that are lifted up.

electric field accelerates the quickly moving ions to considerable velocities. Collision between the accelerated negative ions and air molecules creates new negative ions, which on their part are accelerated, collide with air molecules and free fresh negative ions. An avalanche takes place, and the space charge and the resulting electric field grow in a very short period of time. The strong electric field initiates discharges inside the cloud, and a negative stream of electrons emerges as a dim spark called a stepped leader or a dart leader that jumps in steps of approximately 30 m and reaches the Earth in about 10 ms. When the stepped leader approaches the Earth's surface, the electric field increases so strongly that a positive leader travels upwards, preferably from a pointed object (which gives locally an even stronger electric field). After making contact, the main channel is formed.

Because of the stochastic behavior of the space charge accumulation, the stepped leader also creates branches to the main channel. The main channel carries initially a discharge current of a few hundred Amperes, having a speed of approximately 150 km/s. This discharge current heats up the main channel and the main discharge, the return stroke, is a positive discharge, and it travels at approximately half the speed of light equalizing the charge difference between the thundercloud and the Earth. The main discharge current can be 100000 A or more and the temperature of the plasma in the main channel can reach values as high as 30000 K. The pressure in the main channel is typically 20 bar. The creation of the return stroke takes place between 5 and 10 μ s and is accompanied by a shockwave that we experience as thunder. A lightning stroke consists of several of these discharges, usually three or four, with an interval time of 10–100 μ s. The human eye records this as flickering of the lightning and our ear hears the rolling of the thunder. After each

discharge, the plasma channel cools down to approximately 3000 K, leaving enough ionization to create a new conducting plasma channel for the following discharge.

The striking distance gives us insight into where the lightning will reach the Earth's surface. We can visualize the striking distance r_s as being a sphere with a radius r_s around the tip of the stepped leader: the first object at earth potential that is touched by the sphere will be struck by the lightning. The striking distance is related to the peak lightning current that will occur after contact with the object, often described by the following expression:

$$r_s = 10I_{pk}^{0.65} \quad (3.16)$$

r_s the striking distance [m]
 I_{pk} the peak lightning current [kA]

Example 3.4 Striking distance

According to Equation 3.16, a peak lightning current of 10 kA corresponds to a striking distance of about 45 m, whereas a peak lightning current of 100 kA corresponds to a striking distance of about 200 m.

The ground wire is positioned such that the phase conductors are protected from too large lightning currents. In other words; the ground wires are positioned in such a way that, before the phase conductor is within the striking distance of a large lightning discharge, either the ground wire or the Earth's surface is touched by the sphere. This is illustrated in Figure 3.45. The sphere with radius r_{s1} corresponds to a large lightning current and touches either the ground or the shield wires; the phase conductor is protected from large lightning currents. Smaller currents, such as the sphere with radius r_{s2} , can hit either the ground or the phase conductor. Therefore, the ground or shield wires act as a kind of filter: the system is protected against too large lightning currents, but smaller currents can enter the system. In this way, the maximum amplitude of the current that is injected in the system is known and the protection can anticipate it.

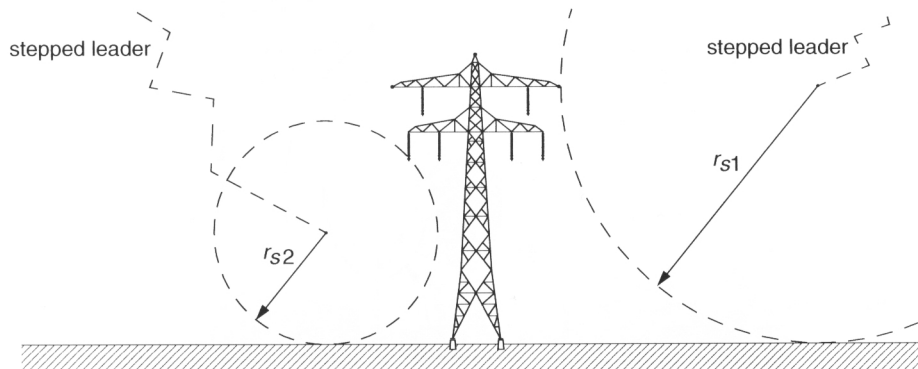


Figure 3.45 Lightning protection by means of shield wires; r_{s1} corresponds to a large lightning current that hits either the ground or the shield wires; r_{s2} corresponds to a smaller lightning current that hits either the ground or the phase conductor.

Chapitre 13

Analyse des systèmes et régimes triphasés déséquilibrés