

Electromagnétisme : PEIP 2 Polytech

8 septembre 2021

Table des matières

1	Systèmes de coordonnées et vecteurs	6
1.1	Systèmes de coordonnées	6
1.1.1	Repère cartésien	6
1.1.2	Repère cylindrique	8
1.1.3	Coordonnées sphériques	12
2	Le champ électrostatique	16
2.1	Notions générales	16
2.1.1	Phénomènes électrostatiques : notion de charge électrique	16
2.1.2	Structure de la matière	18
2.1.3	Matériaux isolants et matériaux conducteurs	18
2.2	Force et champ électrostatiques	19
2.2.1	La force de Coulomb	19
2.2.2	Champ électrostatique créé par une charge ponctuelle	20
2.2.3	Champ créé par un ensemble de charges - principe de superposition	21
2.2.4	Exemple : champ créé par un segment fini de charge	22
2.3	Propriétés de symétrie du champ électrostatique	24
2.3.1	Quelques Compléments	26
3	Lois fondamentales de l'électrostatique	27
3.1	Flux du champ électrostatique	27
3.1.1	Lignes de champ	27
3.1.2	Définition de flux	28
3.1.3	Notion d'angle solide	28
3.1.4	Théorème de Gauss	29
3.1.5	Exemples d'applications	30
3.1.6	Circulation du champ électrostatique	32
3.1.7	Potentiel créé par une charge ponctuelle	34
3.1.8	Potentiel créé par un ensemble de charges	35
3.2	Equations différentielles et intégrales de l'électrostatique	36
3.2.1	Forme locale du théorème de Gauss	36
3.2.2	Dérivation de la forme locale de Gauss (esquisse)	36
3.2.3	Equations fondamentales de l'électrostatique	38
3.2.4	Equation de Poisson	38
4	Conducteurs en équilibre	40
4.1	Conducteurs isolés	40
4.1.1	Notion d'équilibre électrostatique	40

4.1.2	Quelques propriétés des conducteurs en équilibre	40
4.1.3	Pression électrostatique	42
4.1.4	Pouvoir des pointes	43
4.1.5	Capacité d'un conducteur isolé	44
4.1.6	Superposition des états d'équilibre	44
4.2	Systèmes de conducteurs en équilibre	45
4.2.1	Théorème des éléments correspondants	45
4.2.2	Phénomène d'influence électrostatique	46
4.2.3	Coefficients d'influence électrostatique	47
4.3	Le condensateur	50
4.3.1	Condensation de l'électricité	50
4.3.2	Capacités de quelques condensateurs simples	51
4.3.3	Associations de condensateurs	53
4.4	Complément : Relations de continuité du champ électrique	54
5	Dipôle électrique - Energie électrostatique	56
5.1	Le dipôle électrostatique	56
5.1.1	Potentiel électrostatique créé par deux charges électriques	56
5.1.2	Champ électrique à grande distance d'un dipôle : $r \gg d$	58
5.1.3	Développements multipolaires	59
5.2	Diélectriques	60
5.2.1	Champ créé par un diélectrique	61
5.2.2	Le vecteur polarisation	62
5.2.3	Condensateur avec diélectrique	62
5.2.4	Déplacement électrique	63
5.2.5	Condensateur et déplacement électrique	65
5.3	Energie potentielle électrostatique	65
5.3.1	Energie électrostatique d'une charge ponctuelle	65
5.3.2	Energie électrostatique d'un ensemble de charges ponctuelles	65
5.3.3	Energie stockée dans le champ électrique	67
5.3.4	Energie électrostatique de conducteurs en équilibre	67
5.3.5	Energie électrostatique d'une distribution continue de charges	68
5.3.6	Quelques exemples	69
5.4	Actions électrostatiques sur des conducteurs	70
5.4.1	Notions de mécanique du solide	70
5.4.2	Calcul direct des actions électrostatiques sur un conducteur chargé	72
5.4.3	Calcul des actions à partir de l'énergie électrostatique	
	(Méthode des travaux virtuels)	73
5.4.4	Exemple du condensateur	75
5.4.5	Exemple du dipôle	76
6	Courant et champ magnétique	77
6.1	Courant et résistance électriques	77
6.1.1	Le courant électrique	77
6.2	Le champ magnétique	80
6.2.1	Bref aperçu historique	80
6.2.2	Nature des effets magnétiques	81
6.3	Expressions du champ magnétique	83

6.3.1	Champ magnétique créé par une charge en mouvement	83
6.3.2	Champ magnétique créé par un ensemble de charges en mouvement	84
6.3.3	Champ créé par un circuit électrique (formule de Biot et Savart)	84
6.3.4	Propriétés de symétrie du champ magnétique	86
6.4	Calcul du champ dans quelques cas simples	88
6.4.1	Champ créé par un segment de fil rectiligne	88
6.4.2	Champ créé par un fil infini	89
6.4.3	Spire circulaire (sur l'axe)	89
6.4.4	Champ d'un solénoïde fini (sur l'axe)	90
6.4.5	Solénoïde infini	91
7	Lois fondamentales de la magnétostatique	92
7.1	Flux du champ magnétique	92
7.1.1	Conservation du flux magnétique	92
7.1.2	Lignes de champ et tubes de flux	94
7.2	Circulation du champ magnétique	95
7.2.1	Circulation du champ autour d'un fil infini	95
7.3	Le théorème d'Ampère	96
7.3.1	Relations de continuité du champ magnétique	98
7.4	« Potentiel vecteur »	100
7.5	Quatre façons de calculer le champ magnétique	101
7.6	Le dipôle magnétique	101
7.6.1	Champ magnétique créé par une spire	101
8	Champ magnétique en présence de la matière	105
8.1	Le modèle du dipôle en physique	105
8.2	La magnétisation	106
8.3	Le champ « $\mathbf{H}$ »	107
9	Actions et énergie magnétiques	109
9.1	Force magnétique sur une particule chargée	109
9.1.1	La force de Lorentz	109
9.1.2	Trajectoire d'une particule chargée en présence d'un champ magnétique	110
9.1.3	Distinction entre champ électrique et champ électrostatique	111
9.2	Actions magnétiques sur un circuit fermé	112
9.2.1	La force de Laplace	112
9.2.2	Définition légale de l'Ampère	114
9.2.3	Moment de la force magnétique exercée sur un circuit	115
9.2.4	Exemple du dipôle magnétique	115
9.3	Energie potentielle d'interaction magnétique	117
9.3.1	Le théorème de Maxwell	117
9.3.2	Energie potentielle d'interaction magnétique	118
9.3.3	Expressions générales de la force et du couple magnétiques	119
9.3.4	La règle du flux maximum	121
9.4	Laplace et le principe d'Action et de Réaction	121

10 Induction électromagnétique	123
10.1 Les lois de l'induction	123
10.1.1 L'approche de Faraday	123
10.1.2 La loi de Faraday	124
10.1.3 La loi de Lenz	127
10.2 Induction mutuelle et auto-induction	128
10.2.1 Induction mutuelle entre deux circuits fermés	128
10.2.2 Auto-induction	129
10.3 Régimes variables	130
10.3.1 Définition du régime quasi-stationnaire	130
10.3.2 Forces électromotrices (fém) induites	130
10.3.3 Couplage entre bobines : Transformateurs	132
10.4 Retour sur l'énergie magnétique	133
10.4.1 L'énergie magnétique des circuits	133
10.4.2 Forme intégrale de l'énergie magnétique	134
10.4.3 Bilan énergétique d'un circuit électrique	136
11 Annexe des mathématiques	138
11.1 Formules utiles	138
11.2 Analyse vectorielle	138
11.2.1 Le flux d'un champ	138
11.2.2 La divergence du champ	139
11.2.3 Le rotationnel d'un champ	139
11.2.4 L'opérateur « nabla »	139
11.3 Théorème de Green-Ostrogradsky	140
11.3.1 Énoncé général	140
11.3.2 Exemple du théorème d'Ostrogradsky	141
11.4 Théorème de Stokes	141
11.4.1 Énoncé général	141
11.4.2 Exemple	141
11.5 Identités d'analyse vectorielle	142
11.6 Règle BAC - CAB	144
11.7 Exercices d'analyse vectorielle	144

Chapitre 1

Systèmes de coordonnées et vecteurs

1.1 Systèmes de coordonnées

1.1.1 Repère cartésien

Repérage d'un point en coordonnées cartésiennes

Un repère cartésien est défini par un point origine O et trois axes (Ox, Oy, Oz) perpendiculaires entre eux (voir figure 1.1a)). Les vecteurs unitaires portés par les axes sont : $\vec{u}_x, \vec{u}_y, \vec{u}_z$ que nous allons le plus souvent dénoter simplement par $\hat{x}, \hat{y}$, et $\hat{z}$.

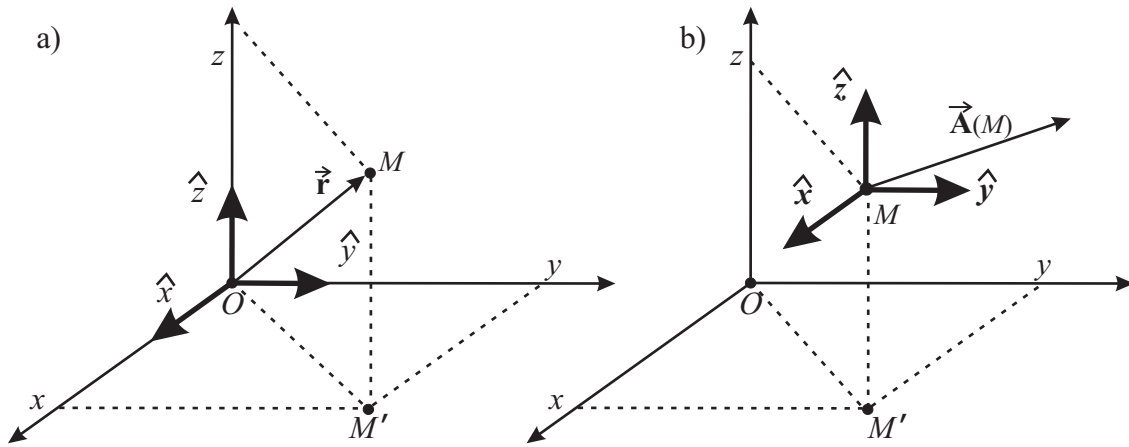


FIGURE 1.1 – Repère cartésien

On doit bien noter la disposition relative des directions (Ox, Oy, Oz) . Telles qu'elles sont placées, elles définissent un trièdre direct. Dans un tel trièdre, un bonhomme transpercé des pieds à la tête par Oy , regardant la direction Oz , a la direction Ox à sa gauche. On peut noter aussi que Ox, Oy , et Oz sont respectivement orientés selon les directions de l'index, du majeur, et du pouce de la main droite. Un point M de l'espace est repéré par les trois composantes du vecteur $\vec{r}$ joignant O à M (voir fig. 1.1a) :

$$\vec{r}(x, y, z) = \overrightarrow{OM} = x\hat{x} + y\hat{y} + z\hat{z} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} = (x, y, z) . \quad (1.1)$$

M' est la projection de M dans le plan (xOy) . Les composantes x et y de $\vec{r}$ sont les coordonnées du point M' dans ce plan. La composante z est obtenue en traçant la parallèle à OM' passant par M . On

dira indistinctement qu'un objet se trouve au point M ou en $\vec{r}$. Les composants des vecteurs, x, y, z , sont des nombres réels et elles peuvent être positives, négatives ou nulles.

Repérage d'un vecteur en coordonnées cartésiennes

Quand il s'agit de repérer un vecteur $\vec{A}$ (M) dont le point d'application est situé au point $M(x, y, z)$, ou $\vec{r}(x, y, z)$, on peut décrire ce vecteur avec le même base de vecteurs unitaires $\hat{x}, \hat{y}, \hat{z}$ (voir fig. 1.1 p)). Nous appelons donc $\hat{x}, \hat{y}, \hat{z}$, un repère orthonormé **global** parce qu'on peut l'utiliser à décrire un vecteur ayant n'importe lequel point d'application.

Déplacement (différentielle) en coordonnées cartésiennes

La différentielle de la position se trouve facilement à partir de sa définition :

$$d\vec{OM} \equiv \frac{\partial \vec{OM}}{\partial x} dx + \frac{\partial \vec{OM}}{\partial y} dy + \frac{\partial \vec{OM}}{\partial z} dz = \hat{x} dx + \hat{y} dy + \hat{z} dz . \quad (1.2)$$

La différentielle volume autour d'un point, $d\mathcal{V}$, correspond au produit des trois différentielles de déplacement

$$d\mathcal{V} = dx dy dz . \quad (1.3)$$

On en déduit également la différentielle de surface orientée $d\vec{S}$

$$d\vec{S} = \hat{x} dy dz + \hat{y} dx dz + \hat{z} dx dy . \quad (1.4)$$

Exemple - Densité de charge :

La plus souvent en physique, on ne parle pas de charges ponctuelles, mais de densité volumique de charge, $\rho(\vec{r})$, distribuée dans un gaz, liquide, plasma ou objet solide. La densité de charge, $\rho_v(\vec{r})$, est analogue à la densité de masse étudiée en cours de mécanique : notamment, si l'on considère un différentielle de volume, $d\mathcal{V}$ autour du point $\vec{r}$ qui enferme une quantité charge appelée dq , la densité volumique de charge en ce point s'écrit par définition :

$$\rho_v(\vec{r}) \equiv \frac{dq}{d\mathcal{V}} . \quad (1.5)$$

La charge totale, Q_{tot} dans un volume $\mathcal{V}$ quelconque s'obtient en intégrant ρ_v sur ce volume,

$$Q_{\text{tot}} = \iiint_{\mathcal{V}} \rho_v(\vec{r}) d\mathcal{V} . \quad (1.6)$$

Exercices résolus :

1. On considère une densité volumique de charge donnée par $\rho_v(x, y, z) = \frac{\rho_0}{a^6} xy^2 z^3$ à l'intérieur d'un cube de côté, a (le cube occupe la région $a > x > 0$, $a > y > 0$, et $a > z > 0$ et ρ_0 et a sont des constantes). La charge totale contenue dans le cube est obtenue en intégrant sur le volume :

$$\begin{aligned} Q_{\text{cube}} &= \iiint_{\text{cube}} \rho(x, y, z) d\mathcal{V} = \int_0^a dx \int_0^a dy \int_0^a dz \frac{\rho_0}{a^6} xy^2 z^3 \\ &= \frac{\rho_0}{a^6} \times \int_0^a x dx \int_0^a y^2 dy \int_0^a z^3 dz \\ &= \frac{\rho_0}{a^6} \times \frac{a^2}{2} \times \frac{a^3}{3} \times \frac{a^4}{4} = \frac{\rho_0}{24} a^3 . \end{aligned}$$

2. On considère un rectangle défini dans un plan $z = cte$ avec une largeur a selon l'axe Ox et une longueur b selon l'axe Oy . Sa densité surfacique est donnée par $\sigma(x, y) = \frac{\sigma_0}{ab^3}xy^3$. Trouver la charge totale de ce rectangle. Ici, le fait que z est constant nous dicte que $dz = 0$ et l'éq. (1.4) nous donne par conséquent que $\vec{dS} \equiv \hat{n}dS = \hat{z}dxdy$ comme il se doit. On ne s'intéresse ici qu'à l'amplitude $dS = dxdy$ de la différentielle de surface. On obtient la charge totale du rectangle en intégrant σ sur sa surface :

$$\begin{aligned} Q_{\text{rect}} &= \iint_{\text{rect}} \sigma(x, y) dS = \frac{\sigma_0}{ab^3} \int_0^a x dx \int_0^b y^3 dy \\ &= \frac{\sigma_0}{ab^3} \frac{a^2}{2} \frac{b^4}{4} = \frac{ab\sigma_0}{8} . \end{aligned}$$

Gradient en coordonnées cartésiennes

La différentielle en coordonnées cartésiennes d'un champ scalaire Φ s'exprime :

$$d\Phi = \frac{\partial \Phi}{\partial x} dx + \frac{\partial \Phi}{\partial y} dy + \frac{\partial \Phi}{\partial z} dz . \quad (1.7)$$

Le gradient en coordonnées cartésiennes est défini telle que :

$$d\Phi = \vec{\text{grad}} \Phi \cdot d\vec{OM} . \quad (1.8)$$

En insérant les expressions de l'éq. (1.2) et (1.7) dans (1.8) on en déduit qu'en coordonnées cartésiennes que l'opérateur gradient s'exprime :

$$\vec{\text{grad}} = \hat{x} \frac{\partial}{\partial x} + \hat{y} \frac{\partial}{\partial y} + \hat{z} \frac{\partial}{\partial z} . \quad (1.9)$$

1.1.2 Repère cylindrique

Repérage d'un point en coordonnées cylindriques

En coordonnées cylindriques, un point M de l'espace est repéré comme un point de cylindre (droit, à base circulaire) dont l'axe Oz est généralement confondu avec l'axe Oz du repère cartésien.

Le point M (ou $\vec{r}$) est repéré par

- le rayon ρ du cylindre sur lequel il s'appuie
- z sa distance par rapport au plan de référence xOy
- ϕ l'angle (Ox, OM') où M' est la projection de M sur le plan xOy .

La notation $\vec{r}(\rho, \phi, z)$ vient se substituer à $\vec{r}(x, y, z)$ du repère cartésien. Vous pouvez facilement vérifier que, pour un point donné, les composantes cartésiennes et cylindriques sont liées par :

$$x = \rho \cos \phi \quad y = \rho \sin \phi \quad z = z . \quad (1.10)$$

Repérage d'un vecteur en coordonnées cylindriques

Nous nous posons la question de repérer un vecteur dont le point d'application est situé au point $M(\rho, \phi, z)$, ou $\vec{r}(\rho, \phi, z)$. Pour cela nous attachons à M un repère orthonormé local $(\hat{\rho}, \hat{\phi}, \hat{z})$. Nous l'appelons **local** parce qu'il n'est pas le même pour tous les points M de l'espace. Ce repère local est fait de 3 vecteurs unitaires de base orthogonaux $(\hat{\rho}, \hat{\phi}, \hat{z})$:

- $\hat{\rho}$ (ou $\vec{u}_\rho$) est un vecteur parallèle à $\vec{OM}'$.
- $\hat{\phi}$ (ou $\vec{u}_\phi$) est dans le de croissance de ϕ , c.-à-d. un vecteur contenu dans le plan xOy et perpendiculaire au vecteur $\hat{\rho}$.

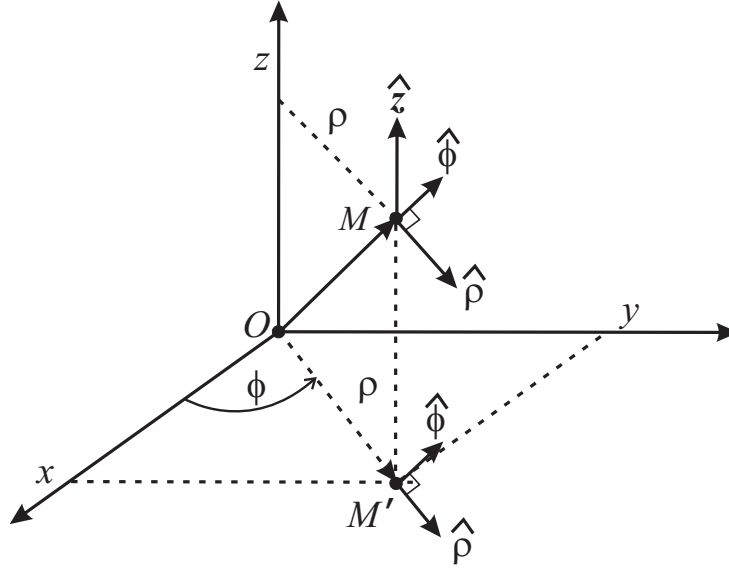


FIGURE 1.2 – coordonnées cylindriques.

— $\hat{z}$ (ou $\vec{u}_z$) est parallèle à l'axe Oz .

En coordonnées cylindriques, un vecteur $\vec{E}(M)$ (ou simplement $\vec{E}(\vec{r})$) attaché au point $M(\rho, \phi, z)$ est repéré par trois composantes (E_ρ, E_ϕ, E_z) dans un repère orthonormé **local** $(\hat{\rho}, \hat{\phi}, \hat{z})$: Dans ce repère, le vecteur champ électrique a 3 composantes et s'écrit

$$\vec{E}(M) = E_\rho \hat{\rho} + E_\phi \hat{\phi} + E_z \hat{z} \quad \text{ou} \quad \vec{E}(M) = \begin{pmatrix} E_\rho \\ E_\phi \\ E_z \end{pmatrix}.$$

Au point M , la relation entre les vecteurs unitaires $(\hat{\rho}, \hat{\phi}, \hat{z})$ et les vecteurs unitaires cartésiennes $(\hat{x}, \hat{y}, \hat{z})$ s'écrivent :

$$\begin{aligned} \hat{\rho} &\equiv \frac{\frac{\partial \vec{OM}}{\partial \rho}}{\left| \frac{\partial \vec{OM}}{\partial \rho} \right|} = \cos \phi \hat{x} + \sin \phi \hat{y} \\ \hat{\phi} &\equiv \frac{\frac{\partial \vec{OM}}{\partial \phi}}{\left| \frac{\partial \vec{OM}}{\partial \phi} \right|} = -\sin \phi \hat{x} + \cos \phi \hat{y} \\ \hat{z} &\equiv \frac{\frac{\partial \vec{OM}}{\partial z}}{\left| \frac{\partial \vec{OM}}{\partial z} \right|} = \hat{z}, \end{aligned} \tag{1.11}$$

où nous avons utilisé la relation

$$\vec{OM} = \rho \cos \phi \hat{x} + \rho \sin \phi \hat{y} + z \hat{z}, \tag{1.12}$$

que l'on a obtenue en insérant les relations de l'éq. (1.10) dans l'éq. (1.1).

On peut voir les relations de l'éq. (1.11) comme une relation matricielle (tensorielle) :

$$\begin{pmatrix} \hat{\rho} \\ \hat{\phi} \\ \hat{z} \end{pmatrix} = \begin{pmatrix} \cos \phi & \sin \phi & 0 \\ -\sin \phi & \cos \phi & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} \hat{x} \\ \hat{y} \\ \hat{z} \end{pmatrix} = T \begin{pmatrix} \hat{x} \\ \hat{y} \\ \hat{z} \end{pmatrix}.$$

Les relations inverses sont obtenues en prenant l'inverse de la matrice T . Puisque les deux bases sont orthonormées, on a $T^{-1} = T^t$ où T^t est la transpose de la matrice T . On obtient de cette manière les vecteurs unitaires $(\hat{x}, \hat{y}, \hat{z})$ en fonction des $(\hat{\rho}, \hat{\phi}, \hat{z})$:

$$\begin{pmatrix} \hat{x} \\ \hat{y} \\ \hat{z} \end{pmatrix} = T^t \begin{pmatrix} \hat{\rho} \\ \hat{\phi} \\ \hat{z} \end{pmatrix} = \begin{pmatrix} \cos \phi & -\sin \phi & 0 \\ \sin \phi & \cos \phi & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} \hat{\rho} \\ \hat{\phi} \\ \hat{z} \end{pmatrix},$$

c'est-à-dire :

$$\begin{aligned} \hat{x} &= \cos \phi \hat{\rho} - \sin \phi \hat{\phi} \\ \hat{y} &= \sin \phi \hat{\rho} + \cos \phi \hat{\phi} \\ \hat{z} &= \hat{z}. \end{aligned} \tag{1.13}$$

On peut également vérifier ces relations avec de la géométrie.

Déplacement (différentielle) en coordonnées cylindriques

En coordonnées cylindriques le vecteur position s'exprime

$$\overrightarrow{OM} = \rho \hat{\rho} + z \hat{z},$$

et la différentielle de déplacement est donc :

$$d\overrightarrow{OM} = \frac{\partial \overrightarrow{OM}}{\partial \rho} d\rho + \frac{\partial \overrightarrow{OM}}{\partial \phi} d\phi + \frac{\partial \overrightarrow{OM}}{\partial z} dz.$$

Si l'on veut exprimer $d\overrightarrow{OM}$ en coordonnées cylindriques, il faut tenir compte du fait que le vecteur unitaire local $\hat{\rho}$ dépend de la coordonnée ϕ (voir eq. (1.11)) :

$$\begin{aligned} \frac{\partial \overrightarrow{OM}}{\partial \rho} &= \hat{\rho} + \rho \frac{\partial \hat{\rho}}{\partial \rho} = \hat{\rho} \quad \left(\text{puisque } \frac{\partial \hat{\rho}}{\partial \rho} = \mathbf{0} \right) \\ \frac{\partial \overrightarrow{OM}}{\partial \phi} &= \rho \frac{\partial \hat{\rho}}{\partial \phi} = \rho \frac{\partial}{\partial \phi} (\cos \phi \hat{x} + \sin \phi \hat{y}) = \rho (-\sin \phi \hat{x} + \cos \phi \hat{y}) = \rho \hat{\phi}. \end{aligned}$$

Un déplacement en coordonnées cylindriques s'exprime donc

$$d\overrightarrow{OM} = \hat{\rho} d\rho + \hat{\phi} \rho d\phi + \hat{z} dz. \tag{1.14}$$

Cette formule est très utile afin d'en déduire des volumes et des surfaces élémentaires. Par exemple, un élément de volume élémentaire en coordonnées cylindriques (plus précisément la différentielle du volume) s'exprime :

$$dV = (d\rho) (\rho d\phi) (dz) = \rho d\rho d\phi dz. \tag{1.15}$$

On peut également servir de l'éq. (1.14) afin d'exprimer la différentielle de surface orientée :

$$d\vec{S} = \hat{\rho} \rho d\phi dz + \hat{\phi} \rho d\rho dz + \hat{z} \rho d\rho d\phi. \tag{1.16}$$

Exemples :

1. On peut utiliser l'éq. (1.15) afin de dériver la formule pour un cylindre de rayon R et de cote L :

$$\begin{aligned} \text{Volume}_{\text{cylindre } R,L} &= \iiint_{\text{cylindre}} dV = \int_0^R d\rho \int_0^{2\pi} \rho d\phi \int_0^L dz = L \int_0^R \rho d\rho \int_0^{2\pi} d\phi \\ &= 2\pi L \int_0^R \rho d\rho = \pi R^2 L . \end{aligned}$$

2. On peut utiliser l'éq. (1.16) afin de calculer la charge totale d'un disque de rayon a et de charge surfacique $\sigma(\rho) = \sigma_0 \frac{\rho^2}{a^2}$. Puisqu'il s'agit d'un disque, on a $dz = 0$, et $d\vec{S} = dS\hat{n} = \hat{z}\rho d\rho d\phi$ avec $\hat{n} = \hat{z}$ comme le vecteur directeur de la surface. Pour cette utilisation, on n'a besoin que de l'amplitude, $dS = \rho d\rho d\phi$, de la différentielle de surface. On calcul la charge totale avec l'intégrale suivante :

$$\begin{aligned} Q_{\text{disque}} &= \iint_{\text{disque}} \sigma(\rho) dS = \int_0^a \rho d\rho \int_0^{2\pi} d\phi \sigma_0 \frac{\rho^2}{a^2} \\ &= \frac{2\pi\sigma_0}{a^2} \int_0^a \rho^3 d\rho = \frac{2\pi\sigma_0}{a^2} \left. \frac{\rho^4}{4} \right|_0^a = \frac{\pi\sigma_0 a^2}{2} . \end{aligned}$$

Gradient en coordonnées cylindriques

La différentielle en coordonnées cylindriques d'un champ scalaire Φ s'exprime :

$$d\Phi = \frac{\partial\Phi}{\partial\rho} d\rho + \frac{\partial\Phi}{\partial\phi} d\phi + \frac{\partial\Phi}{\partial z} dz . \quad (1.17)$$

Le gradient en coordonnées cylindriques est défini telle que :

$$d\Phi = \overrightarrow{\text{grad}} \Phi \cdot d\overrightarrow{OM} . \quad (1.18)$$

Une comparaison entre (1.14), (1.17) et (1.18) montre que l'expression du gradient en coordonnées cylindriques s'écrit :

$$\overrightarrow{\text{grad}} \Phi = \frac{\partial\Phi}{\partial\rho} \hat{\rho} + \frac{1}{\rho} \frac{\partial\Phi}{\partial\phi} \hat{\phi} + \frac{\partial\Phi}{\partial z} \hat{z} \quad (1.19)$$

Exemple : Lorsque le potentiel électrique $V(M)$ est exprimé en coordonnées cylindriques (ρ, ϕ, z) , les composantes du champ électrique dans le repère cylindrique attaché au point M sont données par :

$$\vec{E}(\rho, \phi, z) = -\overrightarrow{\text{grad}} V(\rho, \phi, z)$$

$$\begin{aligned} \vec{E} &= E_\rho \hat{\rho} + E_\phi \hat{\phi} + E_z \hat{z} \\ E_\rho &= -\frac{\partial V}{\partial \rho} \\ E_\phi &= -\frac{1}{\rho} \frac{\partial V}{\partial \phi} \\ E_z &= -\frac{\partial V}{\partial z} \end{aligned} .$$

Le potentiel créé par une distribution linéique de charge avec une densité par unité de longueur λ est donné par $V(\rho) = -\frac{\lambda}{2\pi\epsilon_0} \ln(\rho) + Cte$. On obtient immédiatement le champ électrique par

$$\vec{E}(\rho) = -\overrightarrow{\text{grad}} V(\rho) = \frac{\lambda}{2\pi\epsilon_0\rho} \hat{\rho} .$$

1.1.3 Coordonnées sphériques

Repérage d'un point en coordonnées sphériques

En coordonnées sphériques, un point $M(r, \theta, \phi)$ est considéré comme un point d'une sphère centrée sur O . Le point M est repéré

- par le rayon r de la sphère à laquelle il appartient
- L'angle θ entre la direction $\vec{Oz}$ et la direction $\vec{OM}$. $\theta = (\vec{Oz}, \vec{OM})$
- l'angle ϕ entre la direction $\vec{Ox}$ et la direction $\vec{OM'}$ où M' est la projection de M dans le plan xOy . : $\phi = (\vec{Ox}, \vec{OM'})$

Un point $M(r, \theta, \phi)$ étant donné, on trouve que ses coordonnées cartésiennes s'écrivent en fonction des coordonnées sphériques ; ainsi :

$$x = r \sin \theta \cos \phi \quad y = r \sin \theta \sin \phi \quad z = r \cos \theta \quad (1.20)$$

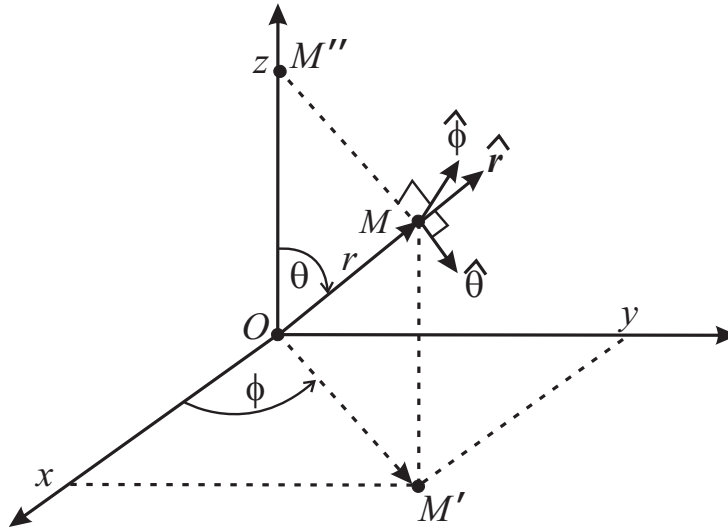


FIGURE 1.3 – Coordonnées sphériques

En géographie, où on est amené à repérer un point sur la sphère terrestre, l'angle θ indiquerait la latitude par rapport au pôle nord et l'angle ϕ , la longitude est par rapport au méridien de référence.

Repérage d'un vecteur en coordonnées sphériques

En coordonnées sphériques, un vecteur $\vec{E}(M)$ (ou simplement $\vec{E}(\vec{r})$) attaché au point $M(r, \theta, \phi)$ est repéré par trois composantes (E_r, E_θ, E_ϕ) dans un repère orthonormé **local** $(\hat{r}, \hat{\theta}, \hat{\phi})$:

$$\vec{E}(M) = E_r \hat{r} + E_\theta \hat{\theta} + E_\phi \hat{\phi} ,$$

avec

- $\hat{r}$ (ou $\vec{u}_r$) est un vecteur parallèle à $\vec{OM}$.
- $\hat{\theta}$ (ou $\vec{u}_\theta$) est parallèle au vecteur tangent en M au cercle de rayon r décrit dans le plan qui contient à la fois les directions $\vec{Oz}$, $\vec{OM}$ et $\vec{OM'}$.
- $\hat{\phi}$ (ou $\vec{u}_\phi$) est tangent en M au cercle de centre M'' et de rayon $M''M = OM'$, contenu dans le plan perpendiculaire à $\vec{Oz}$.

Au point M , la relation entre les vecteurs unitaires $(\hat{r}, \hat{\theta}, \hat{\phi})$ et les vecteurs unitaires cartésiennes $(\hat{x}, \hat{y}, \hat{z})$ s'écrivent :

$$\begin{aligned}\hat{r} &\equiv \frac{\frac{\partial \vec{OM}}{\partial r}}{\left| \frac{\partial \vec{OM}}{\partial r} \right|} = \sin \theta \cos \phi \hat{x} + \sin \theta \sin \phi \hat{y} + \cos \theta \hat{z} \\ \hat{\theta} &\equiv \frac{\frac{\partial \vec{OM}}{\partial \theta}}{\left| \frac{\partial \vec{OM}}{\partial \theta} \right|} = \cos \theta \cos \phi \hat{x} + \cos \theta \sin \phi \hat{y} - \sin \theta \hat{z} \\ \hat{\phi} &\equiv \frac{\frac{\partial \vec{OM}}{\partial \phi}}{\left| \frac{\partial \vec{OM}}{\partial \phi} \right|} = -\sin \phi \hat{x} + \cos \phi \hat{y} \quad ,\end{aligned}\tag{1.21}$$

où nous avons utilisé la relation,

$$\vec{OM} = r \sin \theta \cos \phi \hat{x} + r \sin \theta \sin \phi \hat{y} + r \cos \theta \hat{z} \quad ,\tag{1.22}$$

que l'on a obtenue en insérant les relations de l'éq. (1.20) dans l'éq. (1.1).

On peut voir les relations de l'éq. (1.21) comme une relation matricielle (tensorielle)

$$\begin{pmatrix} \hat{r} \\ \hat{\theta} \\ \hat{\phi} \end{pmatrix} = \begin{pmatrix} \sin \theta \cos \phi & \sin \theta \sin \phi & \cos \theta \\ \cos \theta \cos \phi & \cos \theta \sin \phi & -\sin \theta \\ -\sin \phi & \cos \phi & 0 \end{pmatrix} \begin{pmatrix} \hat{x} \\ \hat{y} \\ \hat{z} \end{pmatrix} = T \begin{pmatrix} \hat{x} \\ \hat{y} \\ \hat{z} \end{pmatrix} \quad ,\tag{1.23}$$

ainsi que les relation inverses,

$$\begin{aligned}\begin{pmatrix} \hat{x} \\ \hat{y} \\ \hat{z} \end{pmatrix} &= T^{-1} \begin{pmatrix} \hat{r} \\ \hat{\theta} \\ \hat{\phi} \end{pmatrix} = T^t \begin{pmatrix} \hat{r} \\ \hat{\theta} \\ \hat{\phi} \end{pmatrix} \\ &= \begin{pmatrix} \sin \theta \cos \phi & \cos \theta \cos \phi & -\sin \phi \\ \sin \theta \sin \phi & \cos \theta \sin \phi & \cos \phi \\ \cos \theta & -\sin \theta & 0 \end{pmatrix} \begin{pmatrix} \hat{r} \\ \hat{\theta} \\ \hat{\phi} \end{pmatrix} \quad ,\end{aligned}\tag{1.24}$$

où nous avons encore utilisé le fait que les deux bases sont orthonomées implique que $T^{-1} = T^t$.

Exemple de coordonnées sphériques : Considérons le potentiel et le champ électriques créés par une charge ponctuelle q placée à l'origine O . En coordonnées sphériques, ceux-ci s'expriment entièrement en fonction du vecteur radial $\vec{r}$ et la coordonnée radiale $r = |\vec{r}|$:

$$V(r) = \frac{q}{4\pi\epsilon_0} \frac{1}{r} \quad \vec{E}(\vec{r}) = \frac{q}{4\pi\epsilon_0} \frac{\hat{r}}{r^2} = \frac{q}{4\pi\epsilon_0} \frac{\vec{r}}{r^3} \quad (\text{avec } \vec{r} = r\hat{r}) \quad ,$$

ce qui est plus simple et "naturel" que les expressions en coordonnées cartésiennes :

$$V(x, y, z) = \frac{q}{4\pi\epsilon_0} \frac{1}{\sqrt{x^2 + y^2 + z^2}} \quad \vec{E}(x, y, z) = \frac{q}{4\pi\epsilon_0} \frac{x\hat{x} + y\hat{y} + z\hat{z}}{(x^2 + y^2 + z^2)^{3/2}} \quad .$$

Position et déplacement (différentielle) en coordonnées sphériques

En coordonnées sphériques, le vecteur position s'écrit simplement

$$\vec{OM} = r\hat{r} \quad .$$

La différentielle, $d\vec{OM}$, en coordonnées sphériques s'exprime :

$$d\vec{OM} = \frac{\partial \vec{OM}}{\partial r} dr + \frac{\partial \vec{OM}}{\partial \theta} d\theta + \frac{\partial \vec{OM}}{\partial \phi} d\phi .$$

Afin d'exprimer $d\vec{OM}$ en coordonnées sphériques, il faut tenir compte du fait que le vecteur unitaire local $\hat{r}$ dépend des coordonnées θ , et ϕ (mais pas sur r) :

$$\begin{aligned} \frac{\partial \vec{OM}}{\partial r} &= \hat{r} + r \frac{\partial \hat{r}}{\partial r} = \hat{r} \\ \frac{\partial \vec{OM}}{\partial \theta} &= r \frac{\partial \hat{r}}{\partial \theta} = r \hat{\theta} \\ \frac{\partial \vec{OM}}{\partial \phi} &= r \frac{\partial \hat{r}}{\partial \phi} = r \sin \theta \hat{\phi} . \end{aligned}$$

Un déplacement en coordonnées sphériques s'exprime donc

$$d\vec{OM} = dr \hat{r} + r d\theta \hat{\theta} + r \sin \theta d\phi \hat{\phi} . \quad (1.25)$$

Comme pour les autres systèmes de coordonnées, on peut utiliser la différentielle de position afin de trouver les différentielles de volume et de surface. La différentielle d'un élément de volume élémentaire en coordonnées sphériques est donc :

$$dV = (dr) (r d\theta) (r \sin \theta d\phi) = r^2 dr \sin \theta d\theta d\phi . \quad (1.26)$$

On peut également servir de l'éq. (1.25) afin d'exprimer la différentielle de surface orientée :

$$d\vec{S} = \hat{r} r^2 \sin \theta d\theta d\phi + \hat{\theta} r \sin \theta dr d\phi + \hat{\phi} r dr d\theta . \quad (1.27)$$

Exemples :

1. On peut utiliser l'éq. (1.26) afin de dériver la formule pour le volume d'une sphère de rayon R :

$$\begin{aligned} \text{Volume}_{\text{sphère de rayon } R} &= \iiint_{\text{sphère}} dV = \int_0^R dr \int_0^\pi d\theta \int_0^{2\pi} r^2 \sin \theta d\phi = \int_0^R r^2 dr \int_0^\pi \sin \theta d\theta \int_0^{2\pi} d\phi \\ &= 2\pi \int_0^R r^2 dr \int_{-1}^1 du = 4\pi \int_0^R r^2 dr = \frac{4\pi}{3} R^3 . \end{aligned}$$

où nous avons fait le changement de variable $u = \cos \theta$ et la relation différentielle associée $du = -\sin \theta d\theta$ afin de trouver :

$$\int_0^\pi \sin \theta d\theta \Rightarrow \int_{-1}^1 du = 2 .$$

2. Quelle est la charge totale d'une sphère de rayon R dont la charge volumique s'exprime $\rho_v(r) = \rho_0 \frac{r}{R}$? (**Attn !** : Il ne faut pas confondre la densité de charge volumique ρ_v avec la coordonnée cylindrique ρ). Cet accident de notation ne pose pas trop de difficultés ici puisque il s'agit dans ce problème de coordonnées sphériques, r , θ , ϕ . La charge totale se trouve par une intégration de la densité volumique :

$$\begin{aligned} Q_{\text{sphère}} &= \iiint_{\text{sphère}} \rho_v(r) dV = \int_0^R dr \int_0^\pi d\theta \int_0^{2\pi} \rho_0 \frac{r}{R} r^2 \sin \theta d\phi \\ &= 4\pi \rho_0 \int_0^R \frac{r^3}{R} dr = \frac{4\pi \rho_0}{R} \frac{r^4}{4} \Big|_0^R = \pi \rho_0 R^3 . \end{aligned}$$

3. Quelle est la charge totale sur une **surface** sphérique définie par $r = a$ quand la densité de charge surfacique s'exprime $\sigma(\theta) = \sigma_0 \sin^2 \theta$? Puisque il s'agit d'une surface à rayon constant, on a $dr = 0$ ce qui donne dans l'éq. (1.27) que $\vec{dS} \equiv \hat{n}dS = \hat{r}a^2 \sin \theta d\theta d\phi$. Ici, on ne s'intéresse qu'à l'amplitude de la différentielle de surface, $\vec{dS}$, c.-à-d. $dS = a^2 \sin \theta d\theta d\phi$, et on obtient la charge totale en intégrant la densité surfacique :

$$\begin{aligned} Q_{\text{surface}} &= \iint_{\text{surface}} \sigma(\theta) dS = \int_0^\pi d\theta \int_0^{2\pi} d\phi a^2 \sigma_0 \sin^2 \theta \sin \theta \\ &= 2\pi a^2 \sigma_0 \int_0^\pi \sin^2 \theta \sin \theta d\theta = 2\pi a^2 \sigma_0 \int_{-1}^1 (1 - \cos^2 \theta) d(\cos \theta) \\ &= 2\pi a^2 \sigma_0 \int_{-1}^1 (1 - u^2) du = 2\pi a^2 \sigma_0 \left(u - \frac{u^3}{3} \right) \Big|_{-1}^1 = \frac{8\pi a^2 \sigma_0}{3}. \end{aligned}$$

Gradient en coordonnées sphériques

La différentielle en coordonnées sphériques s'écrit :

$$d\Phi = \frac{\partial \Phi}{\partial r} dr + \frac{\partial \Phi}{\partial \theta} d\theta + \frac{\partial \Phi}{\partial \phi} d\phi \equiv \overrightarrow{\text{grad}} \Phi \cdot d\overrightarrow{OM} \quad (1.28)$$

Une comparaison entre cette équation et l'éq. (1.25) montre que l'expression du gradient en coordonnées sphériques est donnée par :

$$\overrightarrow{\text{grad}} \Phi = \hat{r} \frac{\partial \Phi}{\partial r} + \hat{\theta} \frac{1}{r} \frac{\partial \Phi}{\partial \theta} + \hat{\phi} \frac{1}{r \sin \theta} \frac{\partial \Phi}{\partial \phi} \quad (1.29)$$

Exemple : Pour une charge ponctuelle située à l'origine par exemple, si on se rappelle que son potentiel électrique, s'écrit $V(r) = q/(4\pi\epsilon_0 r)$, on obtient toute suite son champ électrique en coordonnées sphériques :

$$\begin{aligned} \vec{E}(\vec{r}) &= -\overrightarrow{\text{grad}} V(r) = -\hat{r} \frac{\partial V}{\partial r} = -\frac{q}{4\pi\epsilon_0} \hat{r} \frac{\partial}{\partial r} \frac{1}{r} \\ &= \frac{q}{4\pi\epsilon_0} \frac{\hat{r}}{r^2} \end{aligned}$$

alors que le calcul est plus onéreux en coordonnées cartésiennes

$$\begin{aligned} V(x, y, z) &= \frac{q}{4\pi\epsilon_0} \frac{1}{(x^2 + y^2 + z^2)^{1/2}} \\ \vec{E}(x, y, z) &= -\overrightarrow{\text{grad}} V(x, y, z) = -\frac{q}{4\pi\epsilon_0} \left(\hat{x} \frac{\partial V}{\partial x} + \hat{y} \frac{\partial V}{\partial y} + \hat{z} \frac{\partial V}{\partial z} \right) \\ &= \frac{q}{4\pi\epsilon_0} \left(\frac{x}{(x^2 + y^2 + z^2)^{3/2}} \hat{x} + \frac{y}{(x^2 + y^2 + z^2)^{3/2}} \hat{y} + \frac{z}{(x^2 + y^2 + z^2)^{3/2}} \hat{z} \right) \\ &= \frac{q}{4\pi\epsilon_0} \frac{x\hat{x} + y\hat{y} + z\hat{z}}{(x^2 + y^2 + z^2)^{3/2}} = \frac{q}{4\pi\epsilon_0} \frac{\vec{r}}{r^3} \end{aligned}$$

Chapitre 2

Le champ électrostatique

2.1 Notions générales

2.1.1 Phénomènes électrostatiques : notion de charge électrique

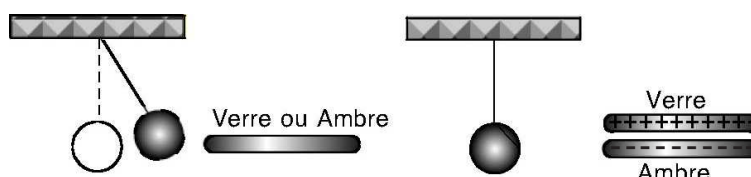
Quiconque a déjà vécu l'expérience désagréable d'une « décharge électrique » lors d'un contact avec un corps étranger connaît un effet provoqué par une charge électrostatique. Une autre manifestation de l'électricité statique consiste en l'attraction de petits corps légers (bouts de papier par ex.) avec des corps frottés (règles, pour continuer sur le même ex.). Ce type de phénomène est même rapporté par Thalès de Milet, aux alentours de 600 av. J.-C. : il avait observé l'attraction de brindilles de paille par de l'ambre jaune frotté... Le mot électricité, qui désigne l'ensemble de ces manifestations, provient de « elektron », qui signifie ambre en grec.

L'étude des phénomènes électriques a continué jusqu'au XIX^{ème} siècle, où s'est élaborée la théorie unifiée des phénomènes électriques et magnétiques, appelée électromagnétisme. C'est à cette époque que le mot « statique » est apparu pour désigner les phénomènes faisant l'objet de ce cours. Nous verrons plus loin, lors du cours sur le champ magnétique, pourquoi il en est ainsi. On se contentera pour l'instant de prendre l'habitude de parler de phénomènes électrostatiques.

Pour les mettre en évidence et pour apporter une interprétation cohérente, regardons deux expériences simples.

— Expérience 1 :

Prenons une boule (faite de sureau ou de polystyrène, par ex.) et suspendons-la par un fil. Ensuite on approche une tige, de verre ou d'ambre, après l'avoir frottée préalablement : la tige attire la boule. Par contre, si l'on approche simultanément deux tiges (ambre et verre) côte à côte, on peut arriver à une situation où rien ne se passe.

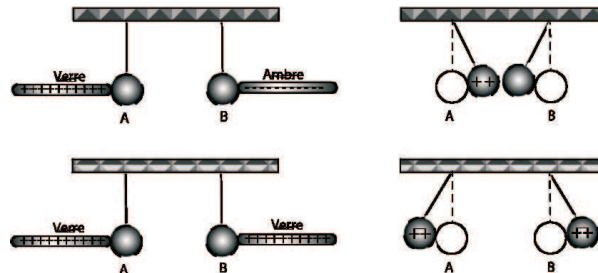


Tout se passe donc comme si chacune des tiges était, depuis son frottement, porteuse d'électricité, mais celle-ci se manifeste en deux états contraires (car capables d'annuler les effets de

l'autre). On a ainsi qualifié arbitrairement de positive l'électricité contenue dans le verre (frotté avec de la soie), et de négative celle portée par l'ambre (idem, ou encore du plastique frotté avec de la fourrure).

— Expérience 2 :

Prenons maintenant deux boules A et B, préalablement mises en contact avec une tige frottée (elles sont « électrisées »), et suspendons-les côte à côte. Si elles ont été mises en contact toutes deux avec une tige de même matériau, elles se *repoussent*.



Par contre, si elles ont été mises en contact avec des tiges de matériau différent (ex. A avec du verre frotté et B avec de l'ambre frotté), alors elles s'attirent. Si, du fait de leur attraction, elles viennent à se toucher, on observe qu'elles perdent alors toute électrisation : elles prennent une position d'équilibre vis-à-vis de leur poids.

Cette expérience est assez riche. On peut tout d'abord en conclure que deux corps portant une électricité de même nature (soit positive, soit négative) se repoussent, tandis qu'ils s'attirent s'ils portent des charges opposées.

Mais cette expérience nous montre également que cette électricité est capable, non seulement d'agir à distance (répulsion ou attraction), mais également de se déplacer d'un corps à un autre. Mais alors qu'est-ce qui se déplace ? Si l'on suspend les boules à une balance, même très précise, nous sommes incapables de détecter la moindre variation de poids entre le début de l'expérience et le moment où elles sont électrisées. Pourtant, le fait qu'il soit nécessaire qu'il y ait un contact entre deux matériaux pour que l'électricité puisse passer de l'un à l'autre, semble indiquer que cette électricité est portée par de la matière.

On explique l'ensemble des effets d'électricité statique par l'existence, au sein de la matière, de particules portant une **charge électrique** q , positive ou négative, et libres de se déplacer. C'est Robert A. Millikan qui a vérifié pour la première fois en 1909, grâce à une expérience mettant en jeu des gouttes d'huile, le fait que toute charge électrique Q est quantifiée, c'est-à-dire qu'elle existe seulement sous forme de multiples d'une charge élémentaire e , **indivisible** ($Q = Ne$). La particule portant cette charge élémentaire est appelée l'**électron**.

Dans le système d'unités international, l'**unité de la charge électrique est le Coulomb (symbole C)** (et **A.s en unités SI de base**). Des phénomènes d'électricité statique mettent en jeu des nanocoulombs (nC) voire des microcoulombs (μC), tandis que l'on peut rencontrer des charges de l'ordre du Coulomb en électrocinétique.

L'ensemble des expériences de la physique (et en particulier celles décrites plus haut) ne peuvent s'expliquer que si la charge électrique élémentaire est invariante : on ne peut ni la détruire ni l'engendrer, et ceci est valable quel que soit le référentiel. C'est ce que l'on décrit par la notion d'invariance relativiste de la charge électrique.

2.1.2 Structure de la matière

La vision moderne de la matière décrit celle-ci comme étant constituée d'atomes. Ceux-ci sont eux-mêmes constitués d'un noyau (découvert en 1911 par Rutherford) autour duquel « gravite » une sorte de nuage composé d'électrons et portant l'essentiel de la masse. Ces électrons se repoussent les uns les autres mais restent confinés autour du noyau car celui-ci possède une charge électrique positive qui les attire. On attribue cette charge positive à des particules appelées **protons**. Cependant, le noyau atomique ne pourrait rester stable s'il n'était composé que de protons : ceux-ci ont en effet tendance à se repousser mutuellement. Il existe donc une autre sorte de particules, les **neutrons** (découverts en 1932 par Chadwick) portant une charge électrique nulle. Les particules constituant le noyau atomique sont appelées les **nucléons**.

Dans le tableau de Mendeleïev tout élément chimique X est représenté par la notation A_ZX . Le nombre A est appelé le nombre de masse : c'est le nombre total de nucléons (protons et neutrons). Le nombre Z est appelé le nombre atomique et est le nombre total de protons constituant le noyau. La charge électrique nucléaire totale est donc $Q = +Ze$, le cortège électronique possédant alors une charge totale $Q = -Ze$, assurant ainsi la neutralité électrique d'un atome. Exemple : le Carbone ${}^{12}_6\text{C}$ possède 12 nucléons, dont 6 protons (donc 6 électrons) et 6 neutrons, le Cuivre ${}^{63}_{29}\text{Cu}$ 63 nucléons dont 29 protons (donc 29 électrons) et 34 neutrons. L'atome de cuivre existe aussi sous la forme ${}^{64}_{29}\text{Cu}$, c'est-à-dire avec 35 neutrons au lieu de 34 : c'est ce qu'on appelle un **isotope**.

Valeurs des charges électriques et des masses des constituants atomiques dans le Système International :

Electron : $q = -e = -1,605 \cdot 10^{-19}\text{C}$: $m_e = 9,109 \cdot 10^{-31}\text{kg}$

Proton : $q = +e = 1,605 \cdot 10^{-19}\text{C}$: $m_p = 1,672 \cdot 10^{-27}\text{kg}$

Neutron : $q = 0\text{C}$: $m_n = 1,674 \cdot 10^{-27}\text{kg}$

Comme on peut le remarquer, même une charge de l'ordre du Coulomb (ce qui est énorme), correspondant à environ 10^{18} électrons, ne produit qu'un accroissement de poids de l'ordre de 10^{-12}kg : c'est effectivement imperceptible.

Si les électrons sont bien des particules quasi-ponctuelles, les neutrons et les protons en revanche ont une taille non nulle (inférieure à 10^{-15}m). Il s'avère qu'ils sont eux-mêmes constitués de **quarks**, qui sont aujourd'hui, avec les électrons, les vraies briques élémentaires de la matière. Les protons ainsi que les neutrons forment ainsi une classe de particules appelée les **baryons**.

A l'heure actuelle, l'univers (ou plutôt l'ensemble reconnu de ses manifestations) est descriptible à l'aide de quatre forces fondamentales :

1. La force électromagnétique, responsable de la cohésion de l'atome (électrons-nucléons) et qui joue un rôle primordiale dans la vaste majorité des phénomènes naturelles et technologiques qui nous entoure.
2. La force nucléaire faible, responsable pour certains décroissances nucléaires exotiques. (mais qui a joué un rôle important dans les premiers instants de l'univers, et à la mort des étoiles)
3. La force nucléaire forte, responsable de la cohésion des baryons (quarks-quarks) et du noyau (protons-neutrons) ;
4. La force gravitationnelle, qui garde nos pieds sur terre et qui est responsable de la structure à grande échelle de l'univers (cohésion des corps astrophysiques, cohésion des systèmes planétaires, des galaxies, des amas galactiques, moteur de la cosmologie).

2.1.3 Matériaux isolants et matériaux conducteurs

Un matériau est ainsi constitué d'un grand nombre de charges électriques, mais celles-ci sont toutes compensées (même nombre d'électrons et de protons). Aux températures usuelles, la matière est élec-

triquement neutre. En conséquence, lorsque des effets d'électricité statique se produisent, cela signifie qu'il y a eu un déplacement de charges, d'un matériau vers un autre : c'est ce que l'on appelle l'**électrisation** d'un corps. Ce sont ces charges, en excès ou en manque, en tout cas non compensées, qui sont responsables des effets électriques sur ce corps (ex : baguette frottée).

Un matériau est dit **conducteur parfait** si, lorsqu'il devient électrisé, les porteurs de charge non compensés peuvent se déplacer librement dans tout le volume occupé par le matériau.

Ce sera un **isolant (ou diélectrique) parfait** si les porteurs de charge non compensés ne peuvent se déplacer librement et restent localisés à l'endroit où ils ont été déposés. Un matériau quelconque se situe évidemment quelque part entre ces deux états extrêmes

Refaisons une expérience d'électricité statique : prenons une baguette métallique par la main et frottons-la avec un chiffon. Cela ne marchera pas, la baguette ne sera pas électrisée. Pourquoi ? Etant nous-mêmes d'assez bons conducteurs, les charges électriques arrachées au chiffon et transférées à la baguette sont ensuite transférées sur nous et l'on ne verra plus d'effet électrique particulier au niveau de la baguette. Pour que cette expérience marche, il est nécessaire d'isoler électriquement la baguette (en la tenant avec un matériau diélectrique).

2.2 Force et champ électrostatiques

2.2.1 La force de Coulomb

Charles Auguste de Coulomb (1736-1806) a effectué une série de mesures (à l'aide d'une balance de torsion) qui lui ont permis de déterminer avec un certain degré de précision les propriétés de la force électrostatique exercée par une charge « ponctuelle » q_1 sur une autre charge « ponctuelle » q_2 :

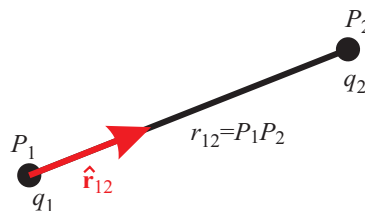
1. La force est radiale, c'est-à-dire dirigée selon la droite qui joint les deux charges ;
2. Elle est proportionnelle au produit des charges : soit attractive si elles sont de signe opposé, soit répulsive sinon ;
3. Enfin, elle varie comme l'inverse du carré de la distance entre les deux charges.

L'expression mathématique moderne de la force de Coulomb et traduisant les propriétés ci-dessus est la suivante

$$\vec{F}_{1 \rightarrow 2} = K \frac{q_1 q_2}{r_{12}^2} \hat{r}_{12} \equiv \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{r_{12}^2} \hat{r}_{12} , \quad (2.1)$$

avec les définitions :

$$\vec{r}_{12} \equiv \overrightarrow{P_1 P_2} \quad r_{12} \equiv |\vec{r}_{12}| \quad \hat{r}_{12} \equiv \frac{\overrightarrow{P_1 P_2}}{|\overrightarrow{P_1 P_2}|} = \frac{\vec{r}_{12}}{r_{12}} , \quad (2.2)$$



La constante multiplicative, K , vaut,

$$K \equiv \frac{1}{4\pi\epsilon_0} \approx 9 \cdot 10^9 \text{ SI } (\text{N.m}^2.\text{C}^{-2}) \text{ ou en unités fondamentales : } (\text{kg.m}^3.\text{s}^{-4}.\text{A}^{-2}) . \quad (2.3)$$

La constante, ϵ_0 , joue un rôle particulier et est appelée la **permittivité électrique du vide** (unités : Farad/m ; et en unités de base : $\text{kg}^{-1} \cdot \text{m}^{-3} \cdot \text{s}^4 \cdot \text{A}^2$)

Remarques :

1. Cette expression n'est valable que pour des charges **immobiles** (approximation de l'électrostatique) et **dans le vide**. Cette loi est la base même de toute l'électrostatique.
2. Cette force obéit au principe d'Action et de Réaction de la mécanique classique. (c.-à-d. si un système (dénote a) exerce une force électrostatique, $\vec{F}_{a \rightarrow b}$, sur un système b , le système b exerce une force, $\vec{F}_{b \rightarrow a} = -\vec{F}_{a \rightarrow b}$ sur le système a).
3. A part la valeur numérique de la constante K , ainsi que l'existence de charges positives et négatives, cette loi a exactement les mêmes propriétés vectorielles que la force de la gravitation (loi de Newton). Il ne sera donc pas étonnant de trouver des similitudes entre ces deux lois.
4. Le chapeau sur $\hat{r}_{12}$ indique qu'il s'agit d'un vecteur **unitaire** (c.-à-d. $|\hat{r}_{12}| = 1$).

Ordres de grandeur

- Quel est le rapport entre la force d'attraction gravitationnelle et la répulsion coulombienne entre deux électrons ?

$$\frac{F_e}{F_g} = \frac{e^2}{4\pi\epsilon_0 G m_e^2} \approx 4 \cdot 10^{42} .$$

La force électrostatique apparaît donc largement dominante vis-à-vis de l'attraction gravitationnelle. Cela implique donc que tous les corps célestes sont exactement électriquement neutres.

- Quelle est la force de répulsion coulombienne entre deux charges de 1 C situées à 1 km ?

$$\frac{F_e}{g} = \frac{1}{4\pi\epsilon_0} \frac{1}{(10^3)^2} \frac{1}{10} \approx 10^3 \text{kg} .$$

C'est une force équivalente au poids exercé par une tonne !

2.2.2 Champ électrostatique créé par une charge ponctuelle

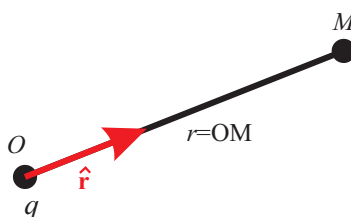
Soit une charge q_1 située en un point O de l'espace, exerçant une force électrostatique sur une autre charge q_2 située en un point M . L'expression de cette force est donnée par la loi de Coulomb ci-dessus (eq. (2.1)). Mais comme pour l'attraction gravitationnelle, on peut la mettre sous une forme plus intéressante,

$$\vec{F}_{1 \rightarrow 2} = q_2 \vec{E}_1(M) ,$$

où

$$\vec{E}_1(M) = \frac{1}{4\pi\epsilon_0} \frac{q_1}{r^2} \hat{r} \quad r \equiv |\vec{OM}| \equiv OM \quad \hat{r} = \frac{\vec{OM}}{OM} .$$

L'intérêt de cette séparation vient du fait que l'on distingue clairement ce qui dépend uniquement de la particule qui subit la force (ici, c'est sa charge q_2 , pour la gravité c'est sa masse), de ce qui ne dépend que d'une source extérieure, ici le vecteur $\vec{E}_1(M)$.



Définition 1 Une particule de charge q située en O crée en tout point M de l'espace distinct de O un champ vectoriel :

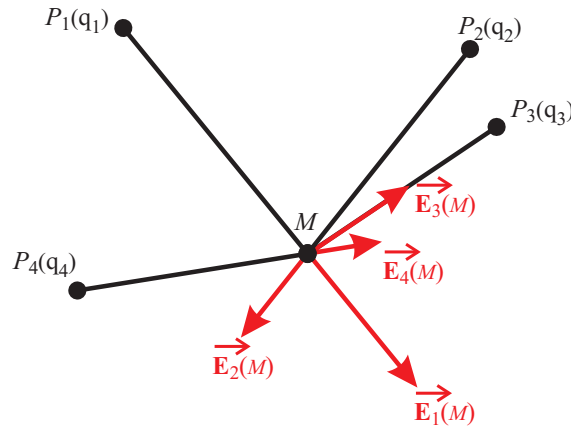
$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \frac{q}{r^2} \hat{r} \quad r \equiv |\vec{OM}| \equiv OM \quad \hat{r} = \frac{\vec{OM}}{OM} \quad (2.4)$$

appelé *champ électrostatique*. L'unité usuel du champ électrique est le **Volt/mètre** (symbole V.m^{-1}). En unités SI de base, le champ électrique est : $\text{kg.m.s}^{-3} \cdot \text{A}^{-1}$

Cette façon de procéder découle de (ou implique) une nouvelle vision de l'espace : les particules chargées se déplacent maintenant dans un espace où existe (se trouve défini) un champ vectoriel. Elles subissent alors une force en fonction de la valeur du champ au lieu où elle se trouve.

2.2.3 Champ créé par un ensemble de charges - principe de superposition

On considère maintenant N particules de charges électriques q_i , situées en des points P_i : quel est le champ électrostatique créé par cet ensemble de charges en un point M ?



La réponse n'est absolument pas évidente car l'on pourrait penser que la présence du champ créé par des particules voisines modifie celui créé par une particule. En fait, il n'en est rien et l'expérience montre que la force totale subie par une charge q située en M est simplement la **superposition** des forces élémentaires,

$$\vec{F}(M) = \sum_{i=1}^N \vec{F}_i(M) = \sum_{i=1}^N \frac{q}{4\pi\epsilon_0} \frac{q_i}{P_i M^2} \hat{u}_i = q \sum_{i=1}^N \frac{1}{4\pi\epsilon_0} \frac{q_i}{P_i M^2} \hat{u}_i = q \vec{E}(M) ,$$

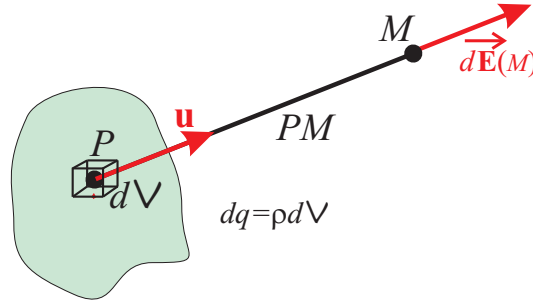
où $\hat{u}_i \equiv \frac{\vec{P_i M}}{P_i M}$, et il en résulte donc que,

$$\vec{E}(M) = \sum_{i=1}^N \frac{1}{4\pi\epsilon_0} \frac{q_i}{P_i M^2} \hat{u}_i , \quad (2.5)$$

est le champ électrostatique créé par un ensemble discret de charges.

Cette propriété de superposition des effets électrostatiques est un fait d'expérience et énoncé comme le principe de superposition (comme tout principe, il n'est pas démontré).

En pratique, cette expression est rarement utilisable puisque nous sommes la plupart du temps amenés à considérer des matériaux comportant un nombre gigantesque de particules. C'est simplement dû au fait que l'on ne considère que des échelles spatiales très grandes devant les distances inter-particulaires, perdant ainsi toute possibilité de distinguer une particule de l'autre. Il est dans ce cas plus habile d'utiliser des distributions continues de charges.



Soit P un point quelconque d'une distribution de charges (solide, gaz, ou plasma) et $dq(P)$ la charge élémentaire contenue en ce point. Le champ électrostatique total créé en un point M par cette distribution de charges est

$$\vec{E}(M) = \int_{\text{distribution}} d\vec{E}(M) \quad \text{avec} \quad d\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \frac{dq}{PM^2} \hat{u} \quad (2.6)$$

Mathématiquement, tout se passe donc comme une charge ponctuelle dq était située en un point P de la distribution, créant au point M un champ électrostatique $d\vec{E}(M)$, avec $\vec{PM} = PM\hat{u}$. Il s'agit évidemment d'une approximation, permettant de remplacer une somme presque infinie par une intégrale.

En désignant par dV le volume infinitésimal autour du point P , on peut définir $\rho(P) \equiv \frac{dq}{dV}$ comme étant la densité volumique de charges (unité : Cm^{-3}). Le champ électrostatique créé par une telle distribution est donc :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \iiint \frac{\rho(P) \vec{PM}}{PM^3} dV = \frac{1}{4\pi\epsilon_0} \iiint \frac{\rho(P)}{PM^2} \hat{u} dV. \quad (2.7)$$

Lorsque l'une des dimensions de la distribution de charges est beaucoup plus petite que les deux autres (ex : un plan ou une sphère creuse), on peut généralement faire une intégration sur cette dimension. On définit alors la densité surfacique de charges $\sigma(P) = \frac{dq}{dS}$ (unité : Cm^{-2}) produisant un champ total

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \iint \frac{\sigma(P)}{PM^2} \hat{u} dS. \quad (2.8)$$

Enfin, si deux des dimensions de la distribution sont négligeables devant la troisième (ex : un fil), on peut définir une densité linéique de charges $\lambda(P) = \frac{dq}{dl}$ (unité : Cm^{-1}) produisant un champ total

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \int \frac{\lambda(P)}{PM^2} \hat{u} dl. \quad (2.9)$$

L'utilisation de l'une ou l'autre de ces trois expressions dépend de la géométrie de la distribution de charges considérée. L'expression générale à retenir est celle de l'éq. (2.6) ci-dessus.

2.2.4 Exemple : champ créé par un segment fini de charge

On considère un segment rectiligne P_1P_2 de densité linéique homogène λ_0 . On veut calculer le champ électrique $\vec{E}$ à un point M à une distance ρ du segment. Compte tenu des symétries, on travaille en coordonnées cylindriques avec l'axe z confondu avec l'axe du segment. Les bouts du segment sont respectivement z_1 et z_2 . On obtient le champ $\vec{E}$ en appliquant l'expression intégrale de l'éq. (2.9) :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \int_{P_1}^{P_2} \lambda_0 \frac{\vec{PM}}{|\vec{PM}|^3} dl. \quad (2.10)$$

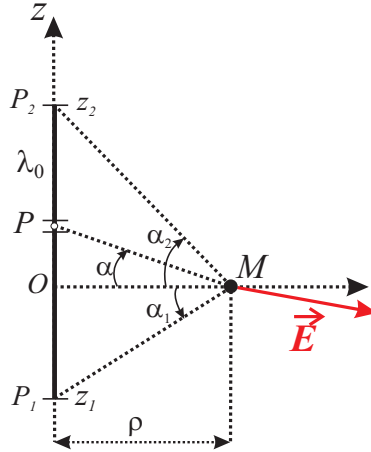


FIGURE 2.1 – Segment rectiligne de charge

Dans ce système de coordonnées cylindriques on a $\overrightarrow{OP} = z\hat{z}$ et :

$$d\ell = dz \quad \overrightarrow{PM} = -z\hat{z} + \rho\hat{\rho} ,$$

et l'intégrale s'écrit :

$$\begin{aligned} \vec{E}(M) &= \frac{\lambda_0}{4\pi\epsilon_0} \int_{z_1}^{z_2} \frac{\overrightarrow{PM}}{|\overrightarrow{PM}|^3} dz = \frac{\lambda_0}{4\pi\epsilon_0} \int_{z_1}^{z_2} \frac{-z\hat{z} + \rho\hat{\rho}}{(\rho^2 + z^2)^{3/2}} dz \\ &= \hat{\rho} \frac{\lambda_0 \rho}{4\pi\epsilon_0} \int_{z_1}^{z_2} \frac{dz}{(\rho^2 + z^2)^{3/2}} - \hat{z} \frac{\lambda_0}{4\pi\epsilon_0} \int_{z_1}^{z_2} \frac{z}{(\rho^2 + z^2)^{3/2}} dz . \end{aligned} \quad (2.11)$$

La deuxième intégrale de l'éq. (2.11) s'effectue avec un changement de variable

$$u^2 = \rho^2 + z^2 \Rightarrow u du = z dz , \quad (2.12)$$

ce qui amène à :

$$\begin{aligned} - \int_{z_1}^{z_2} \frac{z}{(\rho^2 + z^2)^{3/2}} dz &= - \int_{(\rho^2 + z_1^2)^{1/2}}^{(\rho^2 + z_2^2)^{1/2}} \frac{du}{u^2} = \frac{1}{u} \Big|_{(\rho^2 + z_1^2)^{1/2}}^{(\rho^2 + z_2^2)^{1/2}} \\ &= \left[\frac{1}{(\rho^2 + z_2^2)^{1/2}} - \frac{1}{(\rho^2 + z_1^2)^{1/2}} \right] = \frac{1}{\rho} [\cos \alpha_2 - \cos \alpha_1] . \end{aligned} \quad (2.13)$$

Pour le composant du champ dans la direction, $\hat{\rho}$, il faut évaluer l'intégrale :

$$\int_{z_1}^{z_2} \frac{dz}{(\rho^2 + z^2)^{3/2}} = \int_{z_1}^{z_2} \frac{dz}{|\overrightarrow{PM}|^3} . \quad (2.14)$$

Pour cette intégrale, il convient de faire un changement de variables vers l'angle $\alpha = (\overrightarrow{MO}, \overrightarrow{MP})$ et on peut ainsi écrire,

$$|\overrightarrow{OP}| = z = \rho \tan \alpha, \quad PM \equiv |\overrightarrow{PM}| = \frac{\rho}{\cos \alpha} . \quad (2.15)$$

On trouve dz en prenant la différentielle de $z = \rho \tan \alpha$:

$$dz = \rho d(\tan \alpha) = \rho \left(1 + \frac{\sin^2 \alpha}{\cos^2 \alpha} \right) d\alpha = \frac{\rho}{\cos^2 \alpha} d\alpha . \quad (2.16)$$

Mettant les relations de (2.15) et (2.16) dans l'éq. (2.14) on obtient :

$$\begin{aligned} \int_{z_1}^{z_2} \frac{dz}{|\vec{PM}|^3} &= \int_{\alpha_1}^{\alpha_2} \frac{\cos^3 \alpha}{\rho^3} \frac{\rho}{\cos^2 \alpha} d\alpha = \frac{1}{\rho^2} \int_{\alpha_1}^{\alpha_2} \cos \alpha d\alpha \\ &= \frac{1}{\rho^2} (\sin \alpha_2 - \sin \alpha_1) = \frac{1}{\rho^2} \left(\frac{z_2}{(\rho^2 + z_2^2)^{1/2}} - \frac{z_1}{(\rho^2 + z_1^2)^{1/2}} \right). \end{aligned} \quad (2.17)$$

Utilisant les résultats des intégrations des équations (2.13) et (2.17) dans l'éq. (2.11), on obtient enfin $\vec{E}(M)$ en fonction de ρ et des angles α_1 et α_2 :

$$\vec{E}(M) = \frac{\lambda_0}{4\pi\epsilon_0\rho} [\hat{\rho}(\sin \alpha_2 - \sin \alpha_1) + \hat{z}(\cos \alpha_2 - \cos \alpha_1)] , \quad (2.18)$$

ou encore en fonction de ρ et les positions, z_1 et z_2 :

$$\vec{E}(M) = \frac{\lambda_0}{4\pi\epsilon_0\rho} \left[\hat{\rho} \left(\frac{z_2/\rho}{(1 + \frac{z_2^2}{\rho^2})^{1/2}} - \frac{z_1/\rho}{(1 + \frac{z_1^2}{\rho^2})^{1/2}} \right) + \hat{z} \left(\frac{1}{(1 + \frac{z_2^2}{\rho^2})^{1/2}} - \frac{1}{(1 + \frac{z_1^2}{\rho^2})^{1/2}} \right) \right]. \quad (2.19)$$

Les expressions (2.18) et (2.19) sont un peu compliquées, mais il convient de remarquer que si l'on s'intéresse qu'au champ près du segment tel que $z_2/\rho \rightarrow \infty$ et $z_1/\rho \rightarrow -\infty$ (c.-à-d. $\alpha_1 \rightarrow -\frac{\pi}{2}$ et $\alpha_2 \rightarrow \frac{\pi}{2}$), l'expression du champ est assez simple :

$$\vec{E}(\rho) \rightarrow \frac{\lambda_0}{2\pi\epsilon_0\rho} \hat{\rho}. \quad (2.20)$$

Nous verrons dans le prochain chapitre comment retrouver ce résultat avec beaucoup plus de facilité.

2.3 Propriétés de symétrie du champ électrostatique

Principe de Curie : « Lorsque certaines causes produisent certains effets, les éléments de symétrie des causes doivent se retrouver dans les effets produits. »

Du fait que le champ soit un effet créé par une distribution de charges, il contient des informations sur les causes qui lui ont donné origine. Ainsi, si l'on connaît les propriétés de symétrie d'une distribution de charges, on pourra connaître celles du champ électrostatique associé. Ces propriétés sont fondamentales car elles permettent de simplifier considérablement le calcul du champ électrostatique.

Dans un espace homogène et isotrope, si l'on fait subir une transformation géométrique à un système physique (ex : ensemble de particules, distribution de charges) susceptible de créer certains effets (forces, champs), alors ces effets subissent les mêmes transformations. Si un système physique S possède un certain degré de symétrie, on pourra alors déduire les effets créés par ce système en un point à partir des effets en un autre point.

Le principe de Curie nous permet d'arriver à certaines conclusions concernant le comportement du champ produit par des systèmes possédant une ou plusieurs symétries

Règles de symétrie - invariances

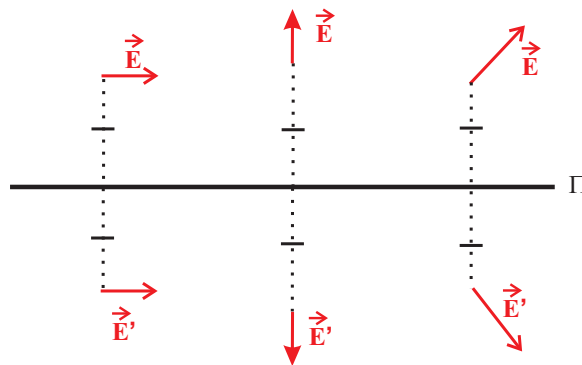
- **Invariance par translation :** si S est invariant dans toute translation parallèle à un axe Oz , le champ de vecteur associé ne dépend pas de z .
- **Symétrie axiale :** si S est invariant dans toute rotation ϕ autour d'un axe Oz , alors un champ de vecteur associé exprimé en coordonnées cylindriques (ρ, ϕ, z) ne dépend pas de ϕ .

- **Symétrie cylindrique** : si S est invariant par translation le long de l'axe Oz et rotation autour de ce même axe, alors un champ de vecteur associé exprimé en coordonnées cylindriques (ρ, ϕ, z) ne dépend que de la distance à l'axe ρ .
- **Symétrie sphérique** : si S est invariant dans toute rotation autour d'un point fixe O , alors un champ de vecteur associé en coordonnées sphériques (r, θ, ϕ) ne dépend que de la distance au centre r .

Pour des systèmes possédant des plans de symétrie ou d'antisymétrie, le principe de Curie nous permet de connaître la direction et l'amplitude d'un vecteur si nous connaissons préalablement son vecteur symétrique.

Transformations géométriques d'un vecteur

Considérons un système possédant une symétrie par rapport à un plan Π .



Soit $\vec{A}'(M')$ le vecteur obtenu par symétrie par rapport à un plan Π à partir de $\vec{A}(M)$. D'après la figure ci-dessus, on voit que :

1. $\vec{A}'(M') = \vec{A}(M)$ si $\vec{A}(M)$ est engendré par les mêmes vecteurs de base que Π .
2. $\vec{A}'(M') = -\vec{A}(M)$ si $\vec{A}(M)$ est perpendiculaire à Π .

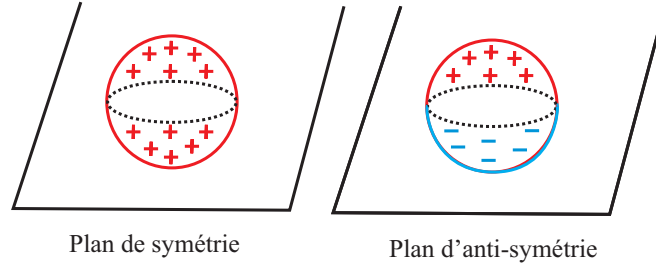
Les signes des règles ci-dessus sont inversés s'il s'agit d'un système antisymétrique.

Les deux règles de transformation combinées avec le principe de Curie nous permet d'imposer les contraintes suivantes concernant les directions d'un champ de vecteurs aux points contenus dans le plan de symétrie (ou d'antisymétrie) :

- **Plan de symétrie Π** : si un système S admet un plan de symétrie Π , alors en tout point de ce plan, le champ électrostatique, $\vec{E}$, est contenu dans ce plan.
- **Plan d'antisymétrie Π'** : si, par symétrie par rapport à un plan Π' , un système S est transformé en $-S$, alors en tout point de ce plan, le champ électrostatique, $\vec{E}$, lui est perpendiculaire.

Remarque importante

Nous verrons en magnétostatique qu'il convient de faire la distinction entre vrais vecteurs (ou vecteurs axiaux) et pseudo-vecteurs (ou vecteurs polaires), ces derniers étant définis à partir du produit vectoriel de deux vecteurs vrais. Ainsi, le champ électrostatique est un vrai vecteur tandis que le champ magnétique est un pseudo-vecteur. Tout ce qui a été dit ci-dessus n'est valable que pour les vrais vecteurs.



2.3.1 Quelques Compléments

- **Pourquoi un vrai vecteur $\vec{A}(x_1, x_2, x_3)$ est indépendant de la variable x_1 si le système S n'en dépend pas ?**

Soit un point $M(x_1, x_2, x_3)$ dont les coordonnées sont exprimé dans un système quelconque. Soit un point $M' \equiv M(x_1 + dx_1, x_2, x_3)$ lui étant infiniment proche. On a alors

$$\vec{A}(M') = \begin{cases} A_1(M') = A_1(x_1 + dx_1, x_2, x_3) \simeq A_1(x_1 + dx_1, x_2, x_3) + \frac{\partial A_1}{\partial x_1} dx_1 \\ A_2(M') = A_2(x_1 + dx_1, x_2, x_3) \simeq A_2(x_1 + dx_1, x_2, x_3) + \frac{\partial A_2}{\partial x_1} dx_1 \\ A_3(M') = A_3(x_1 + dx_1, x_2, x_3) \simeq A_3(x_1 + dx_1, x_2, x_3) + \frac{\partial A_3}{\partial x_1} dx_1 \end{cases}$$

c'est-à-dire, de façon plus compacte $\vec{A}(M') = \vec{A}(M) + \frac{\partial \vec{A}(M)}{\partial x_1} dx_1$. Si le système physique S reste invariant lors d'un changement de M en M' , alors (Principe de Curie) $\vec{A}(M') = \vec{A}(M)$. On a donc $\frac{\partial \vec{A}(M)}{\partial x_1} = \vec{0}$ en tout point M , ce qui signifie que $\vec{A}(x_2, x_3)$ ne dépend pas de x_1 . On peut suivre le même raisonnement pour chacune des autres coordonnées.

- **Pourquoi un vrai vecteur appartient nécessairement à un plan de symétrie ?**

Si M appartient à un plan de symétrie Π du système, son symétrique par rapport à Π est par définition le même point (c.-à.-d. $M' = M$) ce qui entraîne par le principe de Curie que n'importe quel vecteur $\vec{A}(M)$ s'appuyant sur ce point doit satisfaire :

$$\vec{A}'(M) = \vec{A}(M) . \quad (2.21)$$

Si un vecteur $\vec{A}(M)$ est engendré par les mêmes vecteurs de base que Π , la règle 1 entraîne $\vec{A}'(M') = \vec{A}'(M) = \vec{A}(M)$, ce qui est consistant avec (2.21). Par contre, si on veut considérer que $\vec{A}(M)$ soit perpendiculaire à Π , la règle 2 ci-dessus impose que $\vec{A}'(M') = \vec{A}'(M) = -\vec{A}(M)$ ce qui n'est consistant avec (2.21) que si $\vec{A}(M) = \vec{0}$.

- **Pourquoi un vrai vecteur est nécessairement perpendiculaire à un plan Π' d'anti-symétrie ?**

Si M appartient à un plan d'antisymétrie Π' du système, son symétrique par rapport à Π' est par définition le même point (c.-à.-d. $M' = M$) ce qui entraîne par le principe de Curie que n'importe quel vecteur $\vec{A}(M)$ s'appuyant sur ce point doit satisfaire :

$$\vec{A}'(M) = -\vec{A}(M) . \quad (2.22)$$

Si un vecteur $\vec{A}(M)$ soit engendré par les mêmes vecteurs de base que Π' , la règle 1 entraîne $\vec{A}'(M') = \vec{A}'(M) = \vec{A}(M)$ ce qui n'est consistant avec (2.22) que si $\vec{A}(M) = \vec{0}$. Par contre, si le vecteur $\vec{A}(M)$ est perpendiculaire à Π' , la règle 2 ci-dessus impose que $\vec{A}'(M) = -\vec{A}(M)$ ce qui est consistant avec (2.22).

Chapitre 3

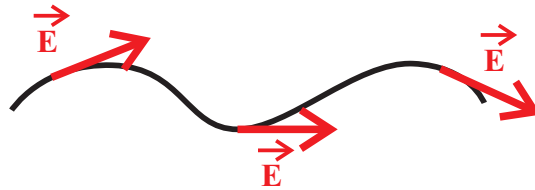
Lois fondamentales de l'électrostatique

3.1 Flux du champ électrostatique

3.1.1 Lignes de champ

Le concept de *lignes de champ* (également appelées lignes de force) est très utile pour se faire une représentation spatiale d'un champ de vecteurs.

Définition : Une ligne de champ d'un champ de vecteur quelconque est une courbe C définie dans l'espace telle que, en chacun de ses points le vecteur $\vec{E}$ soit tangent. Considérons un déplacement



élémentaire $d\vec{\ell}$ le long d'une ligne de champ électrostatique C . Le fait que le champ $\vec{E}$ soit en tout point de C parallèle à $d\vec{\ell}$ s'écrit :

$$\vec{E} \wedge d\vec{\ell} = \vec{0} . \quad (3.1)$$

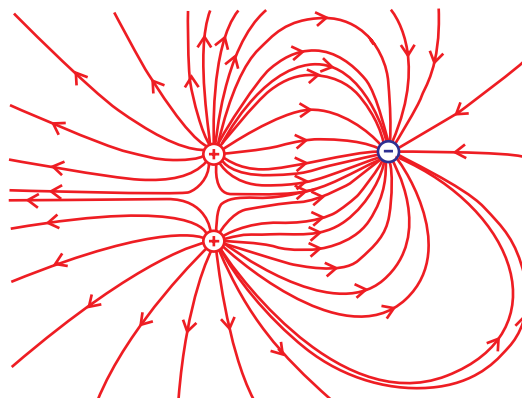


FIGURE 3.1 – Lignes de champ électrique associées avec deux charges positives et une charge négative.

Une façon de penser aux lignes de champ est la suivante : si l'on place une charge positive sur une ligne de champ, la force sera parallèle à la ligne de champ avec l'orientation de la ligne du champ donnée par la direction de la force. Un exemple de cette démarche se trouve en Figure [3.1.1](#) ci-dessus où on a représenté les lignes de champ de deux charges positives et une charge négative.

Il est important de remarquer qu'avec la représentation par lignes du champ, nous n'avons pas d'information directe sur l'intensité du champ. L'intensité du champ correspond à la densité des lignes du champ (une densité de lignes champ élevée correspond à une région de forte amplitude du champ).

3.1.2 Définition de flux

La notion de flux est toujours associée avec une surface S (de façon explicite ou implicite).

Définition : Le flux, Φ_v , d'un champ vectoriel, $\vec{V}$, à travers une surface orientée, $\vec{S}$, est par définition :

$$\Phi_v \equiv \iint_S \vec{V} \cdot d\vec{S}. \quad (3.2)$$

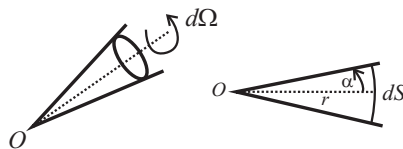
C'est-à-dire, le flux d'un champ vectoriel à travers une surface est la mesure du nombre de lignes du dit champ traversant cette surface.

Dans l'Eq. (3.2), le double intégrale, $\iint_S$, indique une intégration sur la surface S . On se rappelle que la différentielle de surface orientée s'écrit, $d\vec{S} = dS\hat{n}$ ou $\hat{n}$ est le vecteur unitaire normale à la surface. Pour une surface, S , quelconque, il va de soit qu'il y a deux sens possibles pour le vecteur $\hat{n}$, et il faut en générale imposer un choix (souvent arbitraire) sur le sens de $\hat{n}$. Quand il s'agit d'une surface fermée par contre, l'intégrale est dénoté, $\oint_S$, et la convention habituelle est d'orienter $\hat{n}$ de l'intérieur vers l'extérieur de la surface.

Notre notion intuitive de flux correspond au transport « de quelque chose ». En effet, si le vecteur, $\vec{V}$, représente la vitesse d'un fluide de densité homogène, le flux est proportionnel à la quantité de fluide qui traverse la surface par unité de temps. Néanmoins, en physique, on peut parler de flux de n'importe quel champ vectoriel, comme le champ électrique $\vec{E}$, même en l'absence d'une quelconque quantité transportée à travers la surface.

3.1.3 Notion d'angle solide

La notion d'angle solide est l'extension naturelle dans l'espace de l'angle défini dans un plan. Par exemple, le cône de lumière construit par l'ensemble des rayons lumineux issus d'une lampe torche est entièrement décrit par la donnée de deux grandeurs : la direction (une droite) et l'angle maximal d'ouverture des rayons autour de cette droite. On appelle cette droite la génitrice du cône et l'angle en question, l'angle au sommet.



Définition : l'angle solide élémentaire $d\Omega$, délimité par un cône coupant un élément de surface élémentaire dS située sur une sphère de rayon r centrée sur le sommet en O du cône vaut :

$$d\Omega = \frac{dS}{r^2}. \quad (3.3)$$

Cet angle solide est toujours positif et indépendant de la distance r . Son unité est le « stéradian » (symbole sr).

En coordonnées sphériques, la surface élémentaire à r constant vaut $dS = r^2 \sin\theta d\theta d\phi$. L'angle solide élémentaire s'écrit alors $d\Omega = \sin\theta d\theta d\phi$. Ainsi, l'angle solide délimité par un cône de révolution,

d'angle au sommet α vaut

$$\Omega = \int d\Omega = \int_0^{2\pi} d\phi \int_0^\alpha \sin \theta d\theta = 2\pi (1 - \cos \alpha) . \quad (3.4)$$

Le demi-espace, engendré avec $\alpha = \pi/2$ (radians), correspond donc à un angle solide de 2π stéradians, tandis que l'espace entier correspond à un angle solide de 4π stéradians ($\alpha = \pi$).



D'une façon générale, le cône (ou le faisceau lumineux de l'exemple ci-dessus) peut intercepter une surface quelconque, dont la normale $\hat{n}$ fait un angle θ avec la génératrice de vecteur directeur $\hat{r}$. Avec la surface orientée $\vec{dS} \equiv dS \hat{n}$, l'angle solide élémentaire est défini par :

$$d\Omega \equiv \frac{\vec{dS} \cdot \hat{r}}{r^2} = \frac{dS \hat{n} \cdot \hat{r}}{r^2} = \frac{dS \cos \theta}{r^2} = \frac{dS'}{r^2} , \quad (3.5)$$

où dS' est la surface effective (qui, par exemple, serait « vue » par un observateur situé en O).

3.1.4 Théorème de Gauss

On considère maintenant une charge ponctuelle q située en un point O de l'espace. Le flux du champ électrostatique $\vec{E}$, créé par cette charge, à travers une surface élémentaire quelconque **orientée** est par définition :

$$d\Phi_e = \vec{E} \cdot \vec{dS} \equiv \vec{E} \cdot \hat{n} dS . \quad (3.6)$$

Par convention, on oriente le vecteur unitaire $\hat{n}$, normal à la surface dS , vers l'extérieur, c'est-à-dire dans la direction qui s'éloigne de la charge q . Ainsi, pour $q > 0$, le champ $\vec{E}$ est dirigé dans le même sens que $\hat{n}$ et l'on obtient un flux positif.

A partir du champ créé par une charge ponctuelle, on obtient alors :

$$d\Phi_e = \frac{q}{4\pi\epsilon_0} \frac{\hat{r} \cdot \hat{n}}{r^2} dS = \frac{q}{4\pi\epsilon_0} d\Omega , \quad (3.7)$$

c'est-à-dire un flux dépendant directement de l'angle solide sous lequel est vue la surface et non de sa distance r (notez bien que $d\Omega > 0$, q pouvant être positif ou négatif). Ce résultat est une simple conséquence de la décroissance du champ électrostatique en $1/r^2$: on aurait le même genre de résultat avec le champ gravitationnel.

Que se passe-t-il lorsqu'on s'intéresse au flux total à travers une surface (quelconque) fermée ? Prenons le cas illustré dans la figure 3.1.4 ci-dessus. On a une charge q située à l'intérieur de la surface S (enfermant ainsi un volume V), surface orientée (en chaque point de S , le vecteur $\hat{n}$ est dirigé vers l'extérieur). Pour le rayon 1, on a simplement :

$$d\Phi_1 = \frac{q}{4\pi\epsilon_0} d\Omega , \quad (3.8)$$

mais le rayon 2 traverse plusieurs fois la surface, avec des directions différentes. On aura alors une contribution au flux :

$$\begin{aligned} d\Phi_2 &= \frac{q}{4\pi\epsilon_0} \left(\frac{\hat{r} \cdot \hat{n}_1}{r_1^2} dS_1 + \frac{\hat{r} \cdot \hat{n}_2}{r_2^2} dS_2 + \frac{\hat{r} \cdot \hat{n}_3}{r_3^2} dS_3 \right) \\ &= \frac{q}{4\pi\epsilon_0} (d\Omega - d\Omega + d\Omega) \\ &= \frac{q}{4\pi\epsilon_0} d\Omega . \end{aligned} \quad (3.9)$$

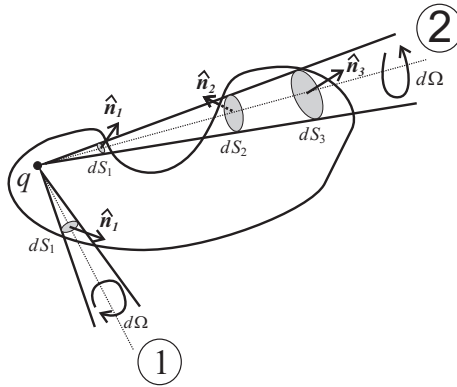


FIGURE 3.2 – Flux électrique à travers une surface fermée

Ce résultat est général, puisque pour une charge se trouvant à l'intérieur de S , un rayon dans une direction donnée va toujours traverser S un nombre impair de fois. En intégrant alors sur toutes les directions (c'est-à-dire sur les 4π stéradians), on obtient un flux électrique total

$$\Phi_e = \oiint_S \vec{E} \cdot d\vec{S} = \frac{q}{\epsilon_0}. \quad (3.10)$$

En vertu du principe de **superposition**, ce résultat se généralise aisément à un ensemble quelconque de charges. Donc nous avons obtenu le :

Théorème 1 - Théorème de Gauss : le flux du champ électrique à travers une surface fermée quelconque est égal, dans le vide, à $1/\epsilon_0$ fois la charge électrique contenue à l'intérieur de cette surface (**N.B.** vecteur normale locale de la surface orientée vers l'extérieur) :

$$\Phi_e \equiv \oiint_S \vec{E} \cdot d\vec{S} = \frac{Q_{\text{int}}}{\epsilon_0}. \quad (3.11)$$

Remarques :

1. Du point de vue physique, le théorème de Gauss fournit le lien entre le flux du champ électrostatique et les sources du champ, à savoir les charges électriques.
2. La démonstration précédente utilise la loi de Coulomb qui, elle, est un fait expérimental et n'est pas démontrée. Inversement, on peut retrouver la loi de Coulomb à partir du théorème de Gauss : c'est ce qui est fait dans l'électromagnétisme, dans lequel le théorème de Gauss constitue une loi fondamentale, non démontrable (l'une des quatre équations de Maxwell).

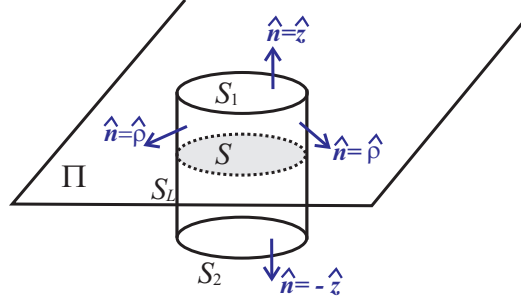
3.1.5 Exemples d'applications

Le théorème de Gauss fournit une méthode très utile pour calculer le champ $\vec{E}$ lorsque celui-ci possède des propriétés de symétrie particulières. Celles-ci doivent en effet permettre de calculer facilement le flux Φ_e . Comme le théorème de Gauss est valable pour une surface quelconque, il nous suffit de trouver une surface S adaptée, c'est-à-dire respectant les propriétés de symétrie du champ, appelée « surface de Gauss ».

Champ électrostatique créé par un plan infini uniformément chargé

On considère un plan infini Π portant une charge électrique σ uniforme par unité de surface. Pour utiliser Gauss, il nous faut d'abord connaître les propriétés de symétrie du champ $\vec{E}$. Tous les plans perpendiculaires au plan infini Π sont des plans de symétrie de celui-ci : $\vec{E}$ appartient aux plans de symétrie, il est donc perpendiculaire à Π . Si ce plan est engendré par les vecteurs

$(\hat{x}, \hat{y})$ alors $\vec{E} = E_z(x, y, z) \hat{z}$. Par ailleurs, l'invariance par translation selon x et y nous fournit $\vec{E}(x, y, z) = E_z(z) \hat{z}$. Le plan Π (c.-à-d. le plan $z = 0$) est lui-même plan de symétrie, donc $E_z(z)$ est impaire.



Etant donné ces propriétés de symétrie, la surface de Gauss la plus adaptée est un cylindre de sections parallèles au plan et situées à des hauteurs symétriques (dénoté par $\pm z$). Le flux à travers cette surface est alors :

$$\begin{aligned} \Phi_e &= \oiint_S \vec{E} \cdot d\vec{S} = \oiint_{S_1} \vec{E} \cdot d\vec{S}_1 + \oiint_{S_2} \vec{E} \cdot d\vec{S}_2 + \oiint_{S_L} \vec{E} \cdot d\vec{S}_L \\ &= E_z(z) S - E_z(-z) S + 0 = 2E_z S \end{aligned} \quad (3.12)$$

où nous avons utilisé le fait que $S_1 = S_2 = S$ et que $\vec{E} \cdot d\vec{S}_L = 0$. Mettant ce résultat pour Φ_e dans le théorème de Gauss nous donne :

$$\Phi_e = 2E_z S = \frac{Q_{\text{int}}}{\epsilon_0} = \frac{1}{\epsilon_0} \iint_S \sigma dS = \frac{\sigma S}{\epsilon_0}. \quad (3.13)$$

Il s'ensuit que le champ électrostatique créé par un plan infini uniformément chargé vaut :

$$\vec{E} = E_z \hat{z} = \frac{\sigma}{2\epsilon_0} \hat{z}. \quad (3.14)$$

Ce résultat n'est valable que pour $z > 0$. En se rappelant que le champ E_z est impair en z , on peut écrire ce résultat sous une forme valable pour $z \leq 0$:

$$\vec{E} = \text{sgn}(z) \frac{\sigma}{2\epsilon_0} \hat{z} \equiv \frac{z}{|z|} \frac{\sigma}{2\epsilon_0} \hat{z}. \quad (3.15)$$

Remarques :

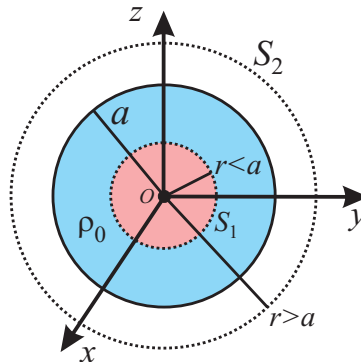
1. Le champ **ne varie pas avec la distance**. Ceci est une conséquence du fait que le plan est supposé infini.
2. On peut encore **appliquer ce résultat pour une surface quelconque chargée uniformément**. Il suffit alors d'interpréter $\vec{E}$ comme le champ au voisinage immédiat de la surface : suffisamment près, celle-ci peut être assimilée à un plan infini.
3. Le champ subit une variation de direction brusque en traversant le plan de charge surfacique. Ce « saut » du champ est caractéristique des problèmes concernant des densités de charge surfacique. Nous le reverrons plus en détail dans le chapitre suivant.

Champ créé par une boule uniformément chargée

On considère une boule (sphère pleine) de centre O et rayon a , chargée avec une distribution volumique homogène de charges ρ_0 . Cette distribution possédant une symétrie sphérique, le champ électrostatique qui en résulte aura la même symétrie, donc $\vec{E}(\vec{r}) = E_r(r) \hat{r}$.

La surface de Gauss adaptée à la symétrie du problème est une sphère de rayon r centrée sur l'origine et l'évaluation de l'intégrale surfacique dans le théorème de Gauss donne :

$$\begin{aligned}\Phi_e &= \oiint_S \vec{E} \cdot d\vec{S} = \oiint_S E(r) dS = E_r(r) 4\pi r^2 \\ &= \frac{Q_{\text{int}}}{\epsilon_0} .\end{aligned}\quad (3.16)$$



Avec ce résultat en main, il nous suffit à déterminer le champ dans les deux régions distinctes, $r < a$ et $r > a$:

- Lorsque $r < a$, il n'y a qu'une partie de la charge totale de la sphère à l'intérieur de la surface de Gauss, S_1 ($Q_{\text{int}} = \frac{4}{3}\pi r^3 \rho_0$), et la formule de l'éq. (3.16) nous donne :

$$E_r(r) = \frac{Q_{\text{int}}}{4\pi\epsilon_0 r^2} = \frac{\frac{4}{3}\pi r^3 \rho_0}{4\pi r^2 \epsilon_0} = \frac{\rho_0}{3\epsilon_0} r \quad r < a \quad (3.17)$$

- Lorsque $r > a$, la sphère de Gauss enferme un volume V supérieur à celui de la boule, mais la distribution de charges n'est non nulle que jusqu'en $r = a$. Donc on a simplement que la charge à l'intérieur de S_2 est $Q_{\text{int}} = Q_{\text{tot}} = \frac{4}{3}\pi a^3 \rho_0$ ce qui fournit dans l'éq. (3.16) que :

$$E_r(r) = \frac{Q_{\text{tot}}}{4\pi\epsilon_0 r^2} = \frac{\rho_0}{3\epsilon_0} \frac{a^3}{r^2} \quad r > a \quad (3.18)$$

On vient ainsi de démontrer, sur un cas simple, qu'une distribution de charges à symétrie sphérique produit à l'extérieur de la distribution, le même champ qu'une charge **ponctuelle** égale, située en O .

3.1.6 Circulation du champ électrostatique

On va démontrer ci-dessous qu'il existe un champ scalaire V , appelé potentiel électrostatique, défini dans tout l'espace et qui permet de reconstruire le champ électrostatique $\vec{E}$. Outre une commodité de calcul (il est plus facile d'additionner deux scalaires que deux vecteurs), l'existence d'un tel scalaire traduit des propriétés importantes du champ électrostatique. Mais tout d'abord, est-il possible d'obtenir un champ de vecteurs à partir d'un champ scalaire ?

Prenons un scalaire $V(M)$ défini en tout point M de l'espace (on dit un champ scalaire). Une variation dV de ce champ lorsqu'on passe d'un point M à un point M' infiniment proche est alors fourni par la différentielle totale :

$$dV(M) = \sum_{i=1}^3 \frac{\partial V}{\partial x_i} dx_i = \overrightarrow{\text{grad}V} \cdot d\overrightarrow{OM} , \quad (3.19)$$

où le vecteur $\overrightarrow{\text{grad}}V$, est le gradient du champ scalaire V et constitue un champ de vecteurs défini partout. Ses composantes dans un système de coordonnées donné sont obtenues très simplement. Par exemple, en coordonnées cartésiennes, on a $d\overrightarrow{OM} = dx\hat{x} + dy\hat{y} + dz\hat{z}$ et

$$dV = \frac{\partial V}{\partial x}dx + \frac{\partial V}{\partial y}dy + \frac{\partial V}{\partial z}dz, \quad (3.20)$$

d'où l'expression suivante pour le gradient en coordonnées cartésiennes

$$\overrightarrow{\text{grad}}V = \frac{\partial V}{\partial x}\hat{x} + \frac{\partial V}{\partial y}\hat{y} + \frac{\partial V}{\partial z}\hat{z} = \left(\begin{array}{c} \frac{\partial V}{\partial x} \\ \frac{\partial V}{\partial y} \\ \frac{\partial V}{\partial z} \end{array} \right). \quad (3.21)$$

En faisant de même en coordonnées cylindriques et sphériques on trouve respectivement $d\overrightarrow{OM} = d\rho\hat{\rho} + \rho d\phi\hat{\phi} + dz\hat{z}$ et $d\overrightarrow{OM} = dr\hat{r} + r d\theta\hat{\theta} + r \sin\theta d\phi\hat{\phi}$ ce qui amènent aux expressions respectives pour le gradient en coordonnées cylindriques et sphériques :

$$\begin{aligned} \overrightarrow{\text{grad}}V &= \hat{\rho}\frac{\partial V}{\partial \rho} + \hat{\phi}\frac{1}{\rho}\frac{\partial V}{\partial \phi} + \hat{z}\frac{\partial V}{\partial z} = \left(\begin{array}{c} \frac{\partial V}{\partial \rho} \\ \frac{1}{\rho}\frac{\partial V}{\partial \phi} \\ \frac{\partial V}{\partial z} \end{array} \right) \\ \overrightarrow{\text{grad}}V &= \hat{r}\frac{\partial V}{\partial r} + \hat{\theta}\frac{1}{r}\frac{\partial V}{\partial \theta} + \hat{\phi}\frac{1}{r \sin\theta}\frac{\partial V}{\partial \phi} = \left(\begin{array}{c} \frac{\partial V}{\partial r} \\ \frac{1}{r}\frac{\partial V}{\partial \theta} \\ \frac{1}{r \sin\theta}\frac{\partial V}{\partial \phi} \end{array} \right). \end{aligned} \quad (3.22)$$

Un déplacement $d\overrightarrow{OM} = \overrightarrow{MM'}$ le long d'une courbe (ou surface) définie par $V = \text{Constante}$ correspond à $dV = 0$, ce qui signifie que $\overrightarrow{\text{grad}}V$ est un vecteur qui est perpendiculaire en tout point à cette courbe (ou surface).

Par ailleurs, plus les composantes du gradient sont élevées et plus il y a une variation rapide de V . Or, c'est bien ce qui semble se produire, par exemple, au voisinage d'une charge électrique q : les lignes de champ électrostatique sont des droites qui convergent ($q < 0$) ou divergent ($q > 0$) toutes vers la charge. Il est donc tentant d'associer le champ $\vec{E}$ (vecteur) au gradient d'une fonction scalaire V .

En fait, depuis Newton (1687) et sa loi de gravitation universelle, de nombreux physiciens et mathématiciens s'étaient penché sur les propriétés de cette force radiale en $1/r^2$. En particulier Lagrange avait ainsi introduit en 1777 une fonction scalaire appelée potentiel, plus « fondamentale » puisque la force en dérive. C'est Poisson qui a introduit le potentiel électrostatique en 1813, par analogie avec la loi de Newton.

Définition : le potentiel électrostatique V est relié au champ électrostatique $\vec{E}$ par

$$\boxed{\vec{E} = -\overrightarrow{\text{grad}}V}. \quad (3.23)$$

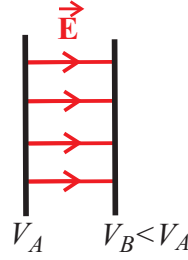
Remarques :

1. Le signe moins est une convention liée à celle adoptée pour l'énergie électrostatique (cf. chapitre 5).
2. La conséquence de cette définition du potentiel est $dV(M) = -\vec{E} \cdot d\overrightarrow{OM}$ pour un déplacement infinitésimal quelconque.
3. Les lignes de champ électrostatique sont perpendiculaires aux courbes équipotentiellles.

Définition : la circulation du champ électrostatique le long d'une courbe allant de A vers B est

$$\boxed{\int_A^B \vec{E} \cdot d\vec{\ell} = - \int_A^B dV = V(A) - V(B)}. \quad (3.24)$$

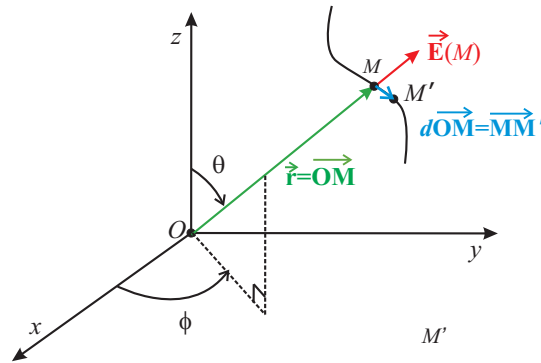
Remarques :



1. Cette circulation est conservatrice : elle ne dépend pas du chemin suivi. Du coup, on dit qu'un champ est **conservateur** du moment qu'on peut l'exprimer partout comme un gradient d'un champ scalaire.
2. La circulation du champ électrostatique sur une courbe fermée (on retourne en A) est nulle. On verra plus loin que ceci est d'une grande importance en électrocinétique.
3. D'après la relation ci-dessus, le long d'une ligne de champ, c'est-à-dire pour $\vec{E} \cdot d\vec{\ell} > 0$ on a $V(A) > V(B)$. Les lignes de champ électrostatiques vont dans le sens des potentiels décroissants.

3.1.7 Potentiel créé par une charge ponctuelle

Nous venons de voir l'interprétation géométrique du gradient d'une fonction scalaire et le lien avec la notion de circulation. Mais nous n'avons pas encore prouvé que le champ électrostatique pouvait effectivement se déduire d'un potentiel V !



Considérons donc une charge ponctuelle q située en un point O pris comme l'origine d'un système de coordonnées sphériques. En un point M de l'espace, cette charge crée un champ électrostatique $\vec{E}$. Le potentiel électrostatique est alors donné par :

$$dV(M) = -\vec{E} \cdot d\vec{OM} = -\frac{q}{4\pi\epsilon_0} \frac{\hat{r} \cdot d\vec{r}}{r^2} = -\frac{q}{4\pi\epsilon_0} \frac{dr}{r^2}, \quad (3.25)$$

c'est-à-dire, après intégration suivant r ,

$$V(r) - V(\infty) = \frac{q}{4\pi\epsilon_0} \int_r^\infty \frac{dr}{r^2} = -\frac{q}{4\pi\epsilon_0} \frac{1}{r} \Big|_r^\infty = \frac{q}{4\pi\epsilon_0} \frac{1}{r}, \quad (3.26)$$

ce qu'on écrit souvent dans la forme

$$V(r) = \frac{q}{4\pi\epsilon_0} \frac{1}{r} + V(\infty) = \frac{q}{4\pi\epsilon_0} \frac{1}{r} + V_0, \quad (3.27)$$

où la constante d'intégration V_0 est le potentiel à l'infini.

Remarques :

1. La constante d'intégration V_0 correspond à la valeur absolue du potentiel électrostatique à l'infini. Ceci est arbitraire puisque seulement des différences de potentiel sont mesurables. C'est une convention universelle de prendre le potentiel nul à l'infini (c.-à-d. $V_0 = 0$).
2. L'unité du potentiel est le Volt . En unités du système international (SI) le Volt vaut :

$$[V] = [E.m] = \text{N m C}^{-1} = \text{kg.m}^2.\text{s}^{-3}.\text{A}^{-1} (\text{unités SI de base}) .$$

3. Si l'on veut se donner une représentation du potentiel, on peut remarquer qu'il mesure le degré d'électrification d'un conducteur (voir Chapitre 4). Il y a en fait une analogie formelle entre d'un côté, le potentiel V et la température T d'un corps, et de l'autre, entre la charge Q et la chaleur déposée dans ce corps.

3.1.8 Potentiel créé par un ensemble de charges

Considérons maintenant un ensemble de N charges ponctuelles q_i distribuées aux points P_i dans tout l'espace. En vertu du principe de superposition, le champ électrostatique total $\vec{E} = \sum_{i=1}^N \vec{E}_i$ est la somme vectorielle des champs $\vec{E}_i$ créés par chaque charge q_i . On peut donc définir un potentiel électrostatique total $V(M) = \sum_{i=1}^N V_i(M)$ tel que $\vec{E} = -\text{grad}V$ soit encore vérifié. En utilisant l'expression du potentiel créé par une charge unique, on obtient :

$$V(M) = \frac{1}{4\pi\epsilon_0} \sum_{i=1}^N \frac{q_i}{r_i} + V_0 , \quad (3.28)$$

où $r_i = P_i M = |\vec{P_i M}|$ est la distance entre la charge q_i et le point M .

Lorsqu'on s'intéresse à des échelles spatiales qui sont très grandes par rapport aux distances entre les charges q_i , on peut faire un passage à la limite continue et remplacer la somme discrète par une intégrale $\sum_i q_i(P_i) \rightarrow \int dq(P)$ où P est un point courant autour duquel se trouve une charge « élémentaire » dq . Le potentiel électrostatique créé par une distribution de charges continue est alors

$$V(M) = \frac{1}{4\pi\epsilon_0} \int \frac{dq}{r} + V_0 \quad r \equiv PM . \quad (3.29)$$

Remarques :

1. Pour une distribution de charges finie, V_0 est simplement le potentiel à l'infini (examiner la limite $PM \rightarrow \infty$). Dans ce cas, la convention est de prendre $V_0 = 0$.
2. Pour des distributions de charges linéique λ , surfacique σ et volumique ρ , on obtient respectivement

$$V(M) = \frac{1}{4\pi\epsilon_0} \int \frac{\lambda d\ell}{r}, \quad V(M) = \frac{1}{4\pi\epsilon_0} \iint \frac{\sigma dS}{r}, \quad V(M) = \frac{1}{4\pi\epsilon_0} \iiint \frac{\rho dV}{r} . \quad (3.30)$$

où $r \equiv PM$.

3. Noter que l'on ne peut pas évaluer le potentiel (ni le champ d'ailleurs) d'une particule en utilisant l'expression discrète (c'est-à-dire pour $r_i = 0$). Par contre, on peut le faire avec une distribution continue (surfacique ou volumique) : c'est dû au fait que dq/r converge lorsque r tend vers zéro pour une distribution surfacique ou volumique.

3.2 Equations différentielles et intégrales de l'électrostatique

3.2.1 Forme locale du théorème de Gauss

Le théorème de Gauss, tel que nous l'avons présenté dans l'éq. (3.11), est apparu sous une forme intégrale. Le flux à travers une surface fermée est relié à la quantité de charges intérieures à ce volume, sans dépendre de leur distribution. Puisque cette formule marche pour n'importe quel volume, elle implique l'existence d'une loi locale ou différentielle qui relie deux grandeurs en un point $\vec{r}$. (Comme par exemple la relation $\vec{E} = -\vec{\text{grad}}V$ qui est locale dans le sens que le champ en $\vec{r}$ est lié à la dérivée du potentiel en ce même point $\vec{r}$.)

Si on laisse le volume dans la forme intégrale de Gauss, l'éq. (3.11), devient infinitésimal, ($\delta\mathcal{V} \rightarrow 0$) on démontre dans 3.2.2 ci-dessous que la loi de Gauss intégrale prend la forme :

$$\left[\frac{\partial E_x}{\partial x} + \frac{\partial E_y}{\partial y} + \frac{\partial E_z}{\partial z} \right] \delta\mathcal{V} = \frac{\rho}{\epsilon_0} \delta\mathcal{V} . \quad (3.31)$$

La forme locale (ou différentielle) de l'éq. (3.11) est donc :

$$\frac{\partial E_x}{\partial x} + \frac{\partial E_y}{\partial y} + \frac{\partial E_z}{\partial z} = \frac{\rho}{\epsilon_0} . \quad (3.32)$$

Elle relie, en chaque point de l'espace, la somme des trois dérivées partielles écrites ci-dessus à la densité de charge volumique en ce même point. Les charges surfaciques, linéiques et ponctuelles sont des idéalizations de la distribution des charges qui doivent être traitées avec précaution (La théorie des distributions est notamment bien adapté à ce genre de situation).

Pour un champ vecteur $\vec{A}$ donné, les physiciens ont pris l'habitude de noter :

$$\boxed{\text{div } \vec{A} = \frac{\partial A_x}{\partial x} + \frac{\partial A_y}{\partial y} + \frac{\partial A_z}{\partial z}} . \quad (3.33)$$

La forme locale du théorème de Gauss s'écrit alors :

$$\boxed{\text{div } \vec{E} = \frac{\rho}{\epsilon_0}} . \quad (3.34)$$

La divergence (div) d'un vecteur est un scalaire. Cet être mathématique vient rejoindre le gradient et le rotationnel dans ce que l'on appelle l'analyse vectorielle. **N.B.**

$$\text{Dans le vide } \rho = 0 \text{ et } \text{div } \vec{E} = 0 . \quad (3.35)$$

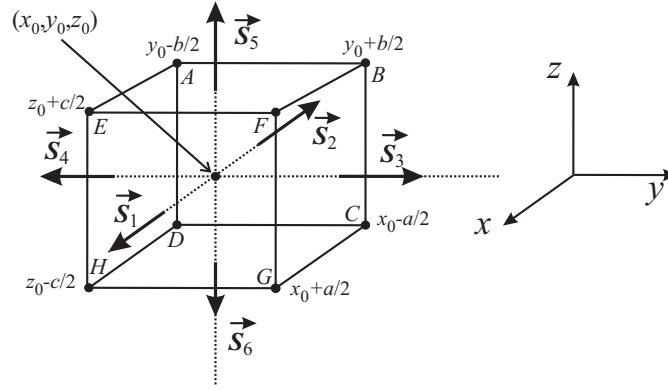
Comme le gradient et le rotationnel, la divergence présente des expressions distinctes en coordonnées cartésiennes, cylindriques et sphériques.

3.2.2 Dérivation de la forme locale de Gauss (esquisse)

Considérons un petit parallélépipède rectangle centré en un point d'abscisse x_0, y_0, z_0 et de côtés a, b et c (voir figure 3.3). Les deux faces perpendiculaires à la direction Ox sont situées en $x = x_0 - a/2$ et $x = x_0 + a/2$. Les côtés de ces faces sont b (parallèlement à Oy) et c (parallèlement à Oz). L'aire de ces deux faces est égale au produit bc .

Reproduire le même raisonnement pour les autres faces.

L'ensemble des 6 rectangles (tels que $ABCD$) forme une surface fermée entourant le volume $\delta\mathcal{V} = abc$ du parallélépipède rectangle


 FIGURE 3.3 – Parallélépipède de volume $\delta V = abc$ centré sur le point (x_0, y_0, z_0)

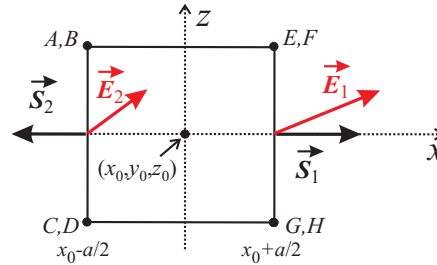
Les vecteurs $\vec{S}_1, \vec{S}_2, \vec{S}_3, \vec{S}_4, \vec{S}_5, \vec{S}_6$ représentant les surfaces des rectangles sont orientées vers l'extérieur du volume et s'écrivent :

$$\begin{aligned} \vec{S}_1 &= \begin{bmatrix} bc \\ 0 \\ 0 \end{bmatrix} & \vec{S}_2 &= \begin{bmatrix} -bc \\ 0 \\ 0 \end{bmatrix} & \vec{S}_3 &= \begin{bmatrix} 0 \\ ac \\ 0 \end{bmatrix} & \vec{S}_4 &= \begin{bmatrix} 0 \\ -ac \\ 0 \end{bmatrix} \\ \vec{S}_5 &= \begin{bmatrix} 0 \\ 0 \\ ab \end{bmatrix} & \vec{S}_6 &= \begin{bmatrix} 0 \\ 0 \\ -ab \end{bmatrix} . \end{aligned} \quad (3.36)$$

Le flux total de $\vec{E}$ à travers la surface parallélépipédique fermée s'écrit :

$$\begin{aligned} \Phi_e &= (\vec{E}_1 \cdot \vec{S}_1 + \vec{E}_2 \cdot \vec{S}_2) + (\vec{E}_3 \cdot \vec{S}_3 + \vec{E}_4 \cdot \vec{S}_4) + (\vec{E}_5 \cdot \vec{S}_5 + \vec{E}_6 \cdot \vec{S}_6) \\ &= \Phi_{12} + \Phi_{34} + \Phi_{56} . \end{aligned} \quad (3.37)$$

Déterminons Φ_{12} , c'est-à-dire la somme des flux du champ électrique à travers les surfaces $\vec{S}_1$ et $\vec{S}_2$. Pour cela, faisons figurer ces deux surfaces de profil. $\vec{E}_1$ est le champ électrique $\vec{E}$ en $x = x_0 + a/2$. $\vec{E}_2 = \vec{E}(x_0 + a/2, y_0, z_0)$. (et $\vec{E}_2 = \vec{E}(x_0 - a/2, y_0, z_0)$).



Puisque bc est l'aire du rectangle $ABCD$, et au vu de l'orientation de $\vec{S}_1$ qui n'a de composante que suivant l'axe Ox , le flux de $\vec{E}$ à travers la surface $\vec{S}_1$ s'écrit : $(bc)E_x(x_0 + a/2, y_0, z_0)$. Faisons de même avec le flux travers $\vec{S}_2$, nous avons donc :

$$\begin{aligned} \Phi_{12} &= \vec{E}_1 \cdot \vec{S}_1 + \vec{E}_2 \cdot \vec{S}_2 \\ &= (bc)E_x(x_0 + a/2, y_0, z_0) + (-bc)E_x(x_0 - a/2, y_0, z_0) \\ &= (bc) [E_x(x_0 + a/2, y_0, z_0) - E_x(x_0 - a/2, y_0, z_0)] . \end{aligned} \quad (3.38)$$

Dans la limite où a tend vers 0, on peut faire un développement limité autour de x_0, y_0, z_0 ,

$$\begin{aligned} E_x(x_0 + a/2, y_0, z_0) &= E_x(x_0, y_0, z_0) + \left(\frac{a}{2}\right) \frac{\partial E_x}{\partial x}(x_0, y_0, z_0) \\ E_x(x_0 - a/2, y_0, z_0) &= E_x(x_0, y_0, z_0) - \left(\frac{a}{2}\right) \frac{\partial E_x}{\partial x}(x_0, y_0, z_0) , \end{aligned} \quad (3.39)$$

d'où

$$\Phi_{12} = abc \frac{\partial E_x}{\partial x}(x_0, y_0, z_0) , \quad (3.40)$$

où le produit abc n'est autre que le volume $\delta\mathcal{V}$ du parallélépipède rectangle. Les flux Φ_{34} et Φ_{56} peuvent être calculés de la même façon, ce qui conduit à :

$$\Phi_e \rightarrow \delta\mathcal{V} \times \left(\frac{\partial E_x}{\partial x} + \frac{\partial E_y}{\partial y} + \frac{\partial E_z}{\partial z} \right) \equiv \delta\mathcal{V} \operatorname{div} \vec{E} . \quad (3.41)$$

La charge à l'intérieur du volume est la densité locale de charge $\rho(x_0, y_0, z_0)$ multipliée par le volume du parallélépipède, c.-à-d. $Q_{\text{int}} = \rho\delta\mathcal{V}$, et la loi de Gauss pour un volume infinitésimal prend donc la forme, $\Phi_e = \delta\mathcal{V} \operatorname{div} \vec{E} = \delta\mathcal{V} \frac{\rho}{\epsilon_0}$.

L'application du théorème de Gauss à un volume infinitésimal donne donc une forme locale (différentielle) de cette loi :

$$\operatorname{div} \vec{E} = \frac{\rho}{\epsilon_0} . \quad (3.42)$$

Nous pouvons également utiliser le théorème de Green-Ostrogradsky de l'éq. (11.15) afin d'aller dans l'autre sens et voir que la forme intégrale du théorème de Gauss est une conséquence directe de sa formulation différentielle.

$$\operatorname{div} \vec{E} = \frac{\rho}{\epsilon_0} \quad \Leftrightarrow \quad \oiint_S \vec{E} \cdot d\vec{S} = \frac{Q_{\text{int}}}{\epsilon_0} . \quad (3.43)$$

Autrement dit : les formes différentielle et intégrale du théorème de Gauss sont deux expressions mathématiques de la même relation physique fondamentale.

3.2.3 Equations fondamentales de l'électrostatique

Il s'avère que toute l'électrostatique du vide peut être formulé en termes d'équations différentielles. En effet, on peut montrer que l'autre loi locale de l'électrostatique, le fait qu'on peut exprimer $\vec{E}$ comme un gradient d'un champ scalaire, est équivalent au fait que la rotationnelle de $\vec{E}$ soit nulle partout :

$$\vec{E} = -\overrightarrow{\operatorname{grad}} V \Leftrightarrow \overrightarrow{\operatorname{rot}} \vec{E} = \vec{0} . \quad (3.44)$$

(Voir 11.17 de l'appendice des maths pour une démonstration de $\overrightarrow{\operatorname{rot}} \overrightarrow{\operatorname{grad}} f = \vec{0}$) **N.B.** Un champ conservateur est caractérisé par $\overrightarrow{\operatorname{rot}} \vec{A} = 0$ dans toute l'espace.

Donc en résumé, on peut formuler toute l'électrostatique du vide avec deux équations différentielles de premier ordre :

$$\boxed{\operatorname{div} \vec{E} = \frac{\rho}{\epsilon_0} \quad \text{et} \quad \overrightarrow{\operatorname{rot}} \vec{E} = \vec{0}} . \quad (3.45)$$

3.2.4 Equation de Poisson

On peut également formuler l'électrostatique en termes d'une seule équation différentielle (de deuxième ordre) appelée équation de Poisson. La combinaison de la forme locale du théorème de Gauss $\operatorname{div} \vec{E} = \rho/\epsilon_0$ et de la relation $\vec{E} = -\overrightarrow{\operatorname{grad}} V$ conduit à l'équation de Poisson. En effet :

$$\operatorname{div} \vec{E} = \frac{\partial E_x}{\partial x} + \frac{\partial E_y}{\partial y} + \frac{\partial E_z}{\partial z} = -\frac{\partial}{\partial x} \left(\frac{\partial V}{\partial x} \right) - \frac{\partial}{\partial y} \left(\frac{\partial V}{\partial y} \right) - \frac{\partial}{\partial z} \left(\frac{\partial V}{\partial z} \right) = \frac{\rho}{\epsilon_0} ,$$

entraîne

$$\frac{\partial^2 V}{\partial x^2} + \frac{\partial^2 V}{\partial y^2} + \frac{\partial^2 V}{\partial z^2} = -\frac{\rho}{\epsilon_0} ,$$

ce qui est l'équation de Poisson. Elle se synthétise en :

$$\boxed{\Delta V = -\frac{\rho}{\epsilon_0}} \quad (3.46)$$

où :

$$\boxed{\Delta \equiv \text{div } \overrightarrow{\text{grad}}}$$

définie un opérateur d'analyse vectorielle appelé *Laplacien*. Nous avons démontré qu'en **coordonnées cartésiennes**, elle s'exprime simplement comme

$$\Delta = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} . \quad (3.48)$$

On peut démontrer mathématiquement (mais nous ne le ferons pas ici) que la solution générale de l'équation de Poisson, l'éq. (3.46), peut s'écrire sous forme intégrale :

$$V(\vec{x}) = \frac{1}{4\pi\epsilon_0} \iiint \frac{\rho(\vec{x}')}{|\vec{x} - \vec{x}'|} d^3x' + V_0 , \quad (3.49)$$

qui n'est rien d'autre que l'expression pour le potentiel trouvé en (3.29) et (3.30) (avec une notation un peu plus sophistiquée). Prenant le gradient de cette expression, on obtient

$$\vec{E}(\vec{x}) = -\overrightarrow{\text{grad}} V(\vec{x}) = \frac{1}{4\pi\epsilon_0} \iiint \frac{\vec{x} - \vec{x}'}{|\vec{x} - \vec{x}'|^3} \rho(\vec{x}') d^3x' , \quad (3.50)$$

qui est de nouveau l'expression intégrale de l'équation du champ électrique déduite de la force de Coulomb dans le chapitre 2. Donc la boucle est bouclée, et on retrouve la physique de Coulomb en prenant les équations différentielles de l'éq. (3.45) comme point de départ. Bien qu'on puisse en déduire une formulation de l'autre, l'histoire (et la pratique) de la physique démontrent que la formulation de l'électromagnétisme en termes de champs est plus porteur pour les développements qui vont suivre que le point de vu de Coulomb en termes de forces agissant entre des charges à distance.

Chapitre 4

Conducteurs en équilibre

4.1 Conducteurs isolés

4.1.1 Notion d'équilibre électrostatique

Jusqu'à présent, nous nous sommes intéressés uniquement aux charges électriques et à leurs effets. Que se passe-t-il pour un corps **conducteur** dans lequel les charges sont libres de se déplacer ?

Prenons une baguette en plastique et frottons-la. On sait qu'elle devient électrisée parce qu'elle devient alors capable d'attirer des petits bouts de papier. Si on la met en contact avec une autre baguette, alors cette deuxième devient également électrisée, c'est-à-dire atteint un certain degré d'électrisation. Au moment du contact des deux baguettes, des charges électriques passent de l'une à l'autre, modifiant ainsi le nombre de charges contenues dans chacune des baguettes, jusqu'à ce qu'un équilibre soit atteint. Comment définir un tel équilibre ?

Définition : *l'équilibre électrostatique d'un conducteur est atteint lorsqu'aucune charge électrique ne se déplace plus à l'intérieur du conducteur.*

Du point de vue des charges élémentaire, cela signifie que **le champ électrostatique total à l'intérieur du conducteur est nul.**

Comme le champ dérive d'un potentiel, cela implique qu'un **conducteur à l'équilibre électrostatique est équipotentiel.**

Remarques :

1. Si le conducteur est chargé, le champ électrostatique total est (principe de superposition) la somme du champ extérieur et du champ créé par la distribution de charges contenues dans le conducteur. Cela signifie que les charges s'arrangent (se déplacent) de telle sorte que le champ qu'elles créent compense exactement, en tout point du conducteur, le champ extérieur.
2. Nous voyons apparaître ici une analogie possible avec la thermodynamique :

Equilibre électrostatique $\Longleftrightarrow$ Equilibre thermodynamique

Potentiel électrostatique $\Longleftrightarrow$ Température

Courant (flux de charges électriques) $\Longleftrightarrow$ flux de chaleur

En effet, à l'équilibre thermodynamique, deux corps de températures initialement différentes mis en contact, acquièrent la même température finale en échangeant de la chaleur (du plus chaud vers le plus froid).

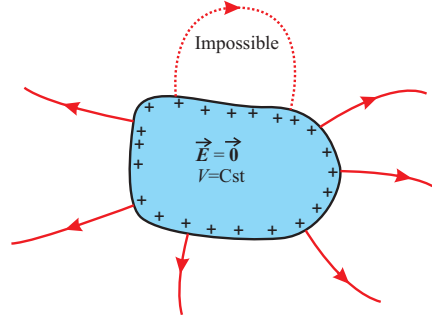
4.1.2 Quelques propriétés des conducteurs en équilibre

- (a) **Lignes de champ** Nous avons vu que, à l'intérieur d'un conducteur (chargé ou non) le champ électrostatique total est nul. Mais ce n'est pas forcément le cas à l'extérieur, en particulier si le

conducteur est chargé. Puisqu'un conducteur à l'équilibre est équipotentiel, cela entraîne alors que, sa surface étant au même potentiel, le champ électrostatique est normal à la surface d'un conducteur. Par ailleurs, aucune ligne de champ ne peut « revenir » vers le conducteur. En effet, la circulation du champ le long de cette ligne impose

$$V(A) - V(B) = \int_A^B \vec{E} \cdot d\vec{\ell}.$$

Si les points A et B appartiennent au même conducteur, alors la circulation doit être nulle, ce qui est impossible le long d'une ligne de champ (où, par définition $\vec{E}$ est parallèle à $d\vec{\ell}$).



- (b) **Distribution des charges** Si un conducteur est chargé, où se trouvent les charges non compensées? Supposons qu'elles soient distribuées avec une distribution volumique ρ . Prenons un volume quelconque $\mathcal{V}$ situé à l'intérieur d'un conducteur à l'équilibre électrostatique. En vertu du théorème de Gauss, on a

$$\oiint_S \vec{E} \cdot d\vec{S} = \iiint_{\mathcal{V}} \frac{\rho}{\epsilon_0} dV = 0,$$

puisque le champ $\vec{E}$ est nul partout. Cela signifie que $\rho = 0$ (autant de charges + que de charges -) et donc, qu'à l'équilibre, aucune charge non compensée ne peut se trouver dans le volume occupé par le conducteur. **Toutes les charges non compensées se trouvent donc nécessairement localisées à la surface du conducteur.**

Ce résultat peut se comprendre par l'effet de répulsion que celles-ci exercent les unes sur les autres. A l'équilibre, les charges tendent donc à se trouver aussi éloignées les unes des autres qu'il est possible de le faire.

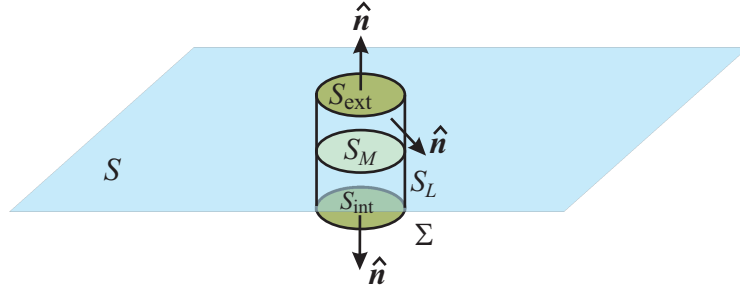
- (c) **Théorème de Coulomb** En un point M infiniment voisin de la surface S d'un conducteur, le champ électrostatique $\vec{E}$ est normal à S . Considérons une petite surface S_{ext} parallèle à la surface S du conducteur. On peut ensuite construire une surface fermée Σ en y adjoignant une surface rentrant à l'intérieur du conducteur S_{int} ainsi qu'une surface latérale S_L . En appliquant le théorème de Gauss sur cette surface fermée, on obtient

$$\begin{aligned} \Phi &= \oiint_{\Sigma} \vec{E} \cdot d\vec{S} = \iint_{S_L} \vec{E} \cdot d\vec{S}_L + \iint_{S_{\text{int}}} \vec{E} \cdot d\vec{S}_{\text{int}} + \iint_{S_{\text{ext}}} \vec{E} \cdot d\vec{S}_{\text{ext}} \\ &= \iint_{S_{\text{ext}}} \vec{E} \cdot d\vec{S}_{\text{ext}} = ES_{\text{ext}} = \frac{Q_{\text{int}}}{\epsilon_0} = \frac{1}{\epsilon_0} \iint_{S_M} \sigma dS = \frac{\sigma S_M}{\epsilon_0}, \end{aligned} \quad (4.1)$$

où S_M est la surface dessinée par le *tube de flux* passant par S_{ext} , donc $S_M = S_{\text{ext}}$ (on peut choisir ces surfaces aussi petites que l'on veut). Mettant ce résultat dans la relation $\frac{\sigma S_M}{\epsilon_0} = ES_{\text{ext}}$ obtenue dans l'éq. (4.1) on obtient le :

Théorème de Coulomb : le champ électrostatique à proximité immédiate d'un conducteur de densité surfacique σ vaut

$$\boxed{\vec{E} = \frac{\sigma}{\epsilon_0} \hat{n}}, \quad (4.2)$$

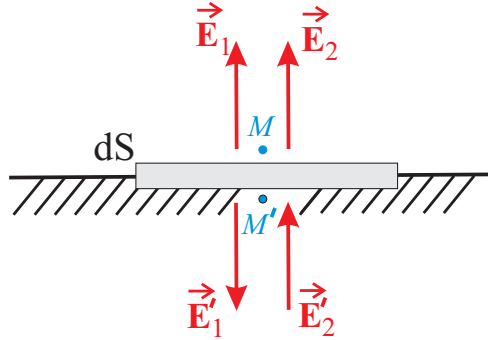


où $\hat{n}$ est un vecteur unitaire normal au conducteur et dirigé vers l'extérieur.

Lorsque le champ au voisinage d'un conducteur dépasse une certaine limite, une étincelle est observée : le milieu entourant le conducteur devient alors conducteur. Ce champ maximal, de l'ordre de 3 Méga V/m dans l'air, est appelé champ disruptif. Il correspond à l'ionisation des particules du milieu (molécules dans le cas de l'air).

4.1.3 Pression électrostatique

Soient deux points M et M' infiniment proches de la surface d'un conducteur de densité surfacique σ , M situé à l'extérieur tandis que M' est situé à l'intérieur. Considérons maintenant une surface élémentaire dS située entre ces deux points. Soit $\vec{E}_1$ le champ créé en M par les charges situées sur dS et $\vec{E}_2$ le champ créé en M par toutes les autres charges situées à la surface du conducteur. Soient $\vec{E}'_1$ et $\vec{E}'_2$ les champs respectifs en M' .



On a alors les trois propriétés suivantes :

1. $\vec{E}_2 = \vec{E}'_2$ car M et M' sont infiniment proches.
2. $\vec{E}'_2 = -\vec{E}'_1$ car le champ électrostatique à l'intérieur du conducteur est nul.
3. $\vec{E}_1 = -\vec{E}'_1$ car $\vec{E}_1$ est symétrique par rapport à dS , considérée comme un plan puisque M et M' peuvent être infiniment rapprochés.

Grâce à ces trois propriétés, on en déduit que $\vec{E}_2 = \vec{E}_1$, c'est-à-dire que la contribution de l'ensemble du conducteur est égale à celle de la charge située à proximité immédiate. Comme le champ total vaut $\vec{E} = \vec{E}_1 + \vec{E}_2 = \frac{\sigma}{\epsilon_0} \hat{n}$ (théorème de Coulomb), on en déduit que le champ créé par l'ensemble du conducteur (à l'exclusion des charges situées en dS) au voisinage du point M est $\vec{E}_2 = \frac{\sigma}{2\epsilon_0} \hat{n}$.

Autrement dit, la force électrostatique $d^2\vec{F}$ subie par cette charge $dq = \sigma dS$ de la part de l'ensemble des autres charges du conducteur vaut

$$d^2\vec{F} = dq\vec{E}_2 = \sigma dS \frac{\sigma}{2\epsilon_0} \hat{n} = \frac{\sigma^2}{2\epsilon_0} \hat{n} dS . \quad (4.3)$$

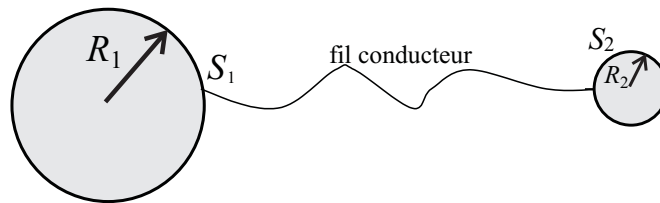
Quel que soit le signe de σ , la force est toujours normale et dirigée vers l'extérieur du conducteur. Cette propriété est caractéristique d'une pression, force par unité de surface. Ainsi, la pression électrostatique

subie en tout point d'un conducteur vaut :

$$P = \frac{\sigma^2}{2\epsilon_0} . \quad (4.4)$$

Cette pression est en général trop faible pour arracher les charges de la surface du conducteur. Mais elle peut déformer ou déplacer la surface, les charges communiquant au solide la force électrostatique qu'elles subissent.

4.1.4 Pouvoir des pointes



L'expression « Pouvoir des pointes » décrit le fait expérimental que, à proximité d'une pointe, le champ électrostatique est toujours très intense. En vertu du théorème de Coulomb, cela signifie que la densité surfacique de charges est, au voisinage d'une pointe, très élevée. On peut aborder ce phénomène avec deux sphères chargées de rayons différents, reliées par un fil conducteur et placées loin l'une de l'autre. On peut donc considérer que chaque sphère est isolée mais qu'elle partage le même potentiel V . Cela implique alors :

$$\begin{aligned} V_1 = V_2 &\implies \frac{1}{4\pi\epsilon_0} \iint_{S_1} \frac{\sigma_1 dS_1}{R_1} = \frac{1}{4\pi\epsilon_0} \iint_{S_2} \frac{\sigma_2 dS_2}{R_2} \implies \frac{\sigma_1 R_1}{\epsilon_0} = \frac{\sigma_2 R_2}{\epsilon_0} \\ &\implies \frac{\sigma_2}{\sigma_1} = \frac{R_1}{R_2} \quad \text{Th. de Coulomb} \implies \frac{|\vec{E}_2|}{|\vec{E}_1|} = \frac{R_1}{R_2} . \end{aligned} \quad (4.5)$$

Donc, plus le rayon de la sphère sera petit et plus sa densité de charges sera élevée et par conséquence du théorème de Coulomb de l'éq. (4.2), plus son champ électrique proche sera grand aussi. Tout se passe comme si les charges « préféraient » les zones à forte courbure. A priori, cela semble en contradiction avec l'idée naïve que les charges non compensées ont tendance à se repousser mutuellement. Le résultat ci-dessus nous montre l'effet d'une pointe (accumulation de charges), mais ne nous offre aucune explication de ce phénomène. Qu'est ce qui, physiquement, a permis une « accumulation » de charges sur une pointe ?

Prenons une sphère chargée placée seule dans l'espace. Se repoussant mutuellement, les charges vont produire une distribution surfacique uniforme. Maintenant, si l'on fait une pointe (zone convexe) les charges situées en haut de la pointe « voient » non seulement le champ électrostatique créé par les charges immédiatement voisines, mais également celui créé par les charges situées sur les bords de la pointe. Quand une charge se retrouve, sous l'effet répulsif des autres charges, repoussée vers la pointe, le champ qu'elle-même crée devient moins important (puisqu'elle est éloignée des autres charges) vis-à-vis des charges restées sur la partie uniforme de la sphère. Cela permet ainsi à une autre charge de prendre sa place : cette nouvelle charge se déplace donc et se retrouve elle-même repoussée sur la pointe. Le conducteur atteint l'équilibre électrostatique lorsque le champ répulsif créé par toutes les charges accumulées au niveau de la pointe compense celui créé par les charges restées sur le « corps » du conducteur.

4.1.5 Capacité d'un conducteur isolé

Nous avons vu qu'il était possible de faire une analogie entre la température d'un corps et le potentiel électrostatique. Or, pour une quantité de chaleur donnée, la température d'un corps dépend en fait de sa capacité calorifique. Il en va de même pour le potentiel électrostatique : il dépend de la capacité du corps à « absorber » les charges électriques qu'il reçoit. On peut donc suivre cette analogie et définir une nouvelle notion, la capacité électrostatique :

$$\text{Capacité électrostatique} \iff \text{Capacité calorifique} .$$

Soit un conducteur à l'équilibre électrostatique isolé dans l'espace, chargé avec une distribution surfacique σ et porté au potentiel V . Celui-ci s'écrit :

$$V(M) = \frac{1}{4\pi\epsilon_0} \iint_{\text{Surface}} \frac{\sigma(P) dS}{PM} ,$$

en tout point M du conducteur, le point P étant un point quelconque de sa surface. Par ailleurs, la charge électrique totale portée par ce conducteur s'écrit

$$Q = \iint_{\text{Surface}} \sigma(P) dS .$$

Si on multiplie la densité surfacique par un coefficient constant a , on obtient une nouvelle charge totale $Q' = aQ$ et un nouveau potentiel $V' = aV$. On a ainsi un nouvel état d'équilibre électrostatique, parfaitement défini. On voit donc que, quoi qu'on fasse, tout état d'équilibre d'un conducteur isolé (caractérisé par Q et V) est tel que le rapport Q/V reste constant (cela résulte de la linéarité de Q et V en fonction de σ).

Définition : La capacité électrostatique d'un conducteur à l'équilibre est définie par

$$\boxed{C = \frac{Q}{V}} , \quad (4.6)$$

où Q est la charge électrique totale du conducteur porté au potentiel V . L'unité de la capacité est le Farad (symbole F) - unité fondamentale ($\text{A}^2 \cdot \text{s}^4 \cdot \text{m}^{-2} \cdot \text{kg}^{-1}$).

Remarques :

1. La capacité C d'un conducteur est une grandeur toujours **positive**. Elle ne dépend que des caractéristiques géométriques et du matériau dont est fait le conducteur.
2. Les unités couramment utilisées en électrocinétique sont le nF ou pF.
3. La permittivité du vide, ϵ_0 , a la dimension de $\text{F} \cdot \text{m}^{-1}$ (unités de base : $\text{m}^{-3} \cdot \text{kg}^{-1} \cdot \text{s}^4 \cdot \text{A}^2$).
4. Exemple : capacité d'une sphère de rayon R , chargée avec une densité surfacique σ

$$\begin{aligned} V = V(O) &= \frac{1}{4\pi\epsilon_0} \iint_{\text{Surface}} \frac{\sigma(P) dS}{OP} = \frac{1}{4\pi\epsilon_0} \iint_{\text{Surface}} \frac{\sigma dS}{R} = \frac{1}{4\pi\epsilon_0 R} \iint_{\text{Surface}} \sigma dS = \frac{Q}{4\pi\epsilon_0 R} \\ C &= \frac{Q}{V} = 4\pi\epsilon_0 R . \end{aligned} \quad (4.7)$$

4.1.6 Superposition des états d'équilibre

Nous avons vu qu'un conducteur isolé, à l'équilibre électrostatique, est caractérisé par sa charge Q et son potentiel V , qui sont reliés entre eux par la capacité C du conducteur. Inversement, étant

donné un conducteur de capacité C , la donnée de sa distribution surfacique σ détermine complètement son état d'équilibre, puisque $Q = \int \int_{\text{Surface}} \sigma dS$.

Soit maintenant un autre état d'équilibre du même conducteur défini par une densité surfacique σ' . Le conducteur porte alors une charge Q' et a un potentiel V' . Du fait de la linéarité de Q et V avec σ , toute combinaison linéaire de σ et σ' est encore un état d'équilibre :

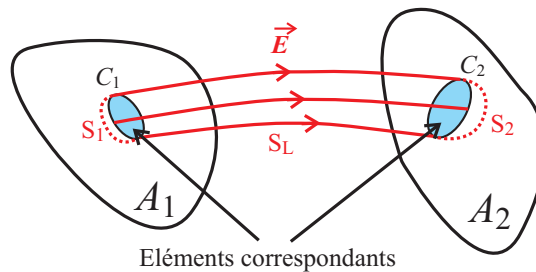
$$\sigma'' = a\sigma + b\sigma' \Leftrightarrow \begin{cases} Q'' = aQ + bQ' \\ V'' = \frac{Q''}{C} = aV + bV' \end{cases} .$$

On a donc ici un résultat qui nous sera utile plus tard : **toute superposition d'états d'équilibre (d'un conducteur ou d'un ensemble de conducteurs) est également un état d'équilibre.**

4.2 Systèmes de conducteurs en équilibre

4.2.1 Théorème des éléments correspondants

Soit deux conducteurs (A_1) et (A_2), placés l'un à coté de l'autre et portant des densités surfaciques σ_1 et σ_2 à l'équilibre. S'ils ne sont pas au même potentiel, des lignes de champ électrostatique relient (A_1) à (A_2). Soit un petit contour fermé C_1 situé sur la surface de (A_1) tel que l'ensemble des lignes de champ s'appuyant sur C_1 rejoignent (A_2) et y dessinent un contour fermé C_2 (on en déduit par construction que toutes les lignes de champ s'appuyant sur la surface A_1 bornée par C_1 se terminent sur la surface A_2 bornée par C_2). L'ensemble de ces lignes de champ constitue ce qu'on appelle un



tube de flux : le flux du champ électrostatique à travers la surface latérale S_L dessinée par ce tube est nul par construction ($\vec{E} \cdot d\vec{S} = 0$). Soit une surface fermée produite $S = S_L + S_1 + S_2$ où S_1 est une surface qui s'appuie sur C_1 et plonge à l'intérieur de conducteur A_1 et S_2 une surface analogue pour le conducteur A_2 .

En vertu du théorème de Gauss, on a

$$\begin{aligned} \Phi &= \oint_S \vec{E} \cdot d\vec{S} = \iint_{S_L} \vec{E} \cdot d\vec{S}_L + \iint_{S_1} \vec{E} \cdot d\vec{S}_1 + \iint_{S_2} \vec{E} \cdot d\vec{S}_2 = 0 \\ &= \frac{Q_{\text{int}}}{\epsilon_0} = \frac{q_1}{\epsilon_0} + \frac{q_2}{\epsilon_0} , \end{aligned}$$

où q_1 est la charge totale contenue sur la surface de (A_1) embrassée par C_1 tandis que q_2 est la charge contenue sur la surface correspondante de (A_2). Du coup $q_2 = -q_1$ nécessairement.

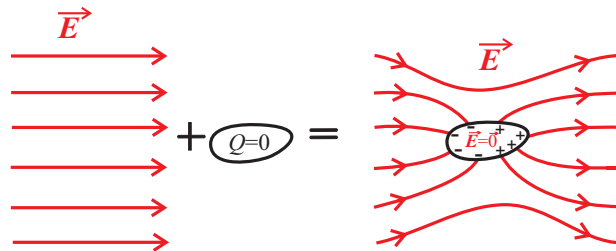
Théorème : les charges électriques portées par deux éléments correspondants sont opposées.

Cette démonstration a fait appel à un concept important, le **tube de flux**, qui relayait dans cet exemple les deux éléments correspondants, C_1 et C_2 . En générale, le tube de flux est définie comme suit :

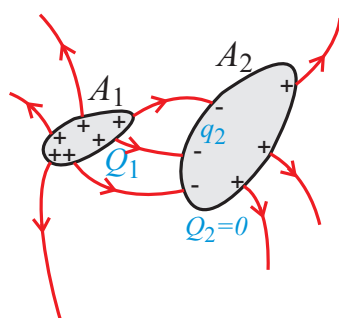
- Soit un contour **fermé** C tel que le champ électrostatique lui est perpendiculaire, c'est-à-dire tel que $\vec{E} \perp d\vec{\ell}$ où $d\vec{\ell}$ est un vecteur élémentaire de C . En chaque point de C passe donc une ligne de champ particulière. L'ensemble de toutes les lignes de champ passant par C dessine alors une surface dans l'espace, une sorte de tube. Par construction, le flux du champ électrostatique est nul à travers la surface latérale du tube, de telle sorte que le flux est conservé : ce qui rentre à la base du tube ressort de l'autre côté. On appelle un tel « rassemblement » de lignes de champ un tube de flux.

4.2.2 Phénomène d'influence électrostatique

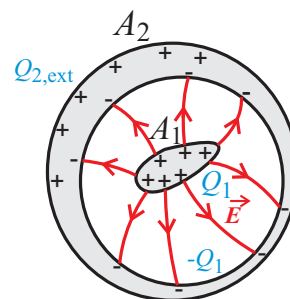
Jusqu'à présent nous n'avons abordé que les conducteurs chargés, isolés dans l'espace. Que se passe-t-il lorsque, par exemple, on place un conducteur neutre dans un champ électrostatique uniforme ? Étant neutre, sa charge $Q = \iint \sigma dS$ doit rester nulle. Mais étant un conducteur, les charges sont libres de se déplacer : on va donc assister à un déplacement de charges positives dans la direction de $\vec{E}$ et de charges négatives dans la direction opposée. On obtient alors une polarisation du conducteur (création de pôles + et -), se traduisant par une distribution surfacique σ non-uniforme (mais telle que $Q = 0$).



Considérons maintenant le cas plus compliqué d'un conducteur (A_1) de charge Q_1 avec une densité surfacique σ_1 , placé à proximité d'un conducteur neutre (A_2). En vertu de ce qui a été dit précédemment, on voit apparaître une densité surfacique σ_2 non-uniforme sur (A_2) due au champ électrostatique de (A_1). Mais, en retour, la présence de charges σ_2 situées à proximité de (A_1) modifie la distribution de charges σ_1 ! A l'équilibre électrostatique, les deux distributions de charges σ_1 et σ_2 dépendent l'une de l'autre. On appelle cette action réciproque, l'influence électrostatique. Dans cet exemple, l'influence est dite partielle, car l'ensemble des lignes de champ électrostatique issues de (A_1) n'aboutissent pas sur (A_2). Soit q_2 la charge portée par la région de (A_2) reliée à (A_1). En vertu du théorème des éléments correspondants, on a $|q_2| < |Q_1|$.



Influence partielle



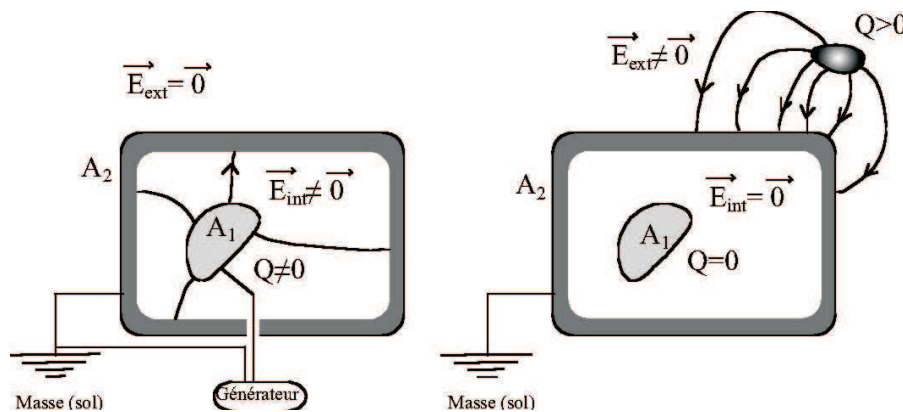
Influence totale

On peut créer des conditions d'influence électrostatique totale en plaçant (A_1) à l'intérieur de (A_2). Puisque l'ensemble des lignes de champ issues de (A_1) aboutit sur l'intérieur de (A_2), on voit apparaître la charge $Q_2^{\text{int}} = -Q_1$ sur la face correspondante interne de (A_2), et ceci quelle que soit la position de (A_1). Cette propriété (démontrée à partir du théorème des éléments correspondants)

est connue sous le nom de théorème de Faraday. La charge électrique totale sur (A_2) est simplement $Q_2 = Q_2^{\text{int}} + Q_2^{\text{ext}} = -Q_1 + Q_2^{\text{ext}}$.

Notion d'écran ou de blindage électrostatique - la cage de Faraday : Un conducteur à l'équilibre a un champ nul : de ce fait, s'il possède une cavité, celle-ci se trouve automatiquement isolée (du point de vue électrostatique) du monde extérieur. On définit par écran électrostatique parfait tout conducteur creux maintenu à un potentiel constant.

Lorsqu'on relie (A_2) au sol, on a $Q_2^{\text{ext}} = 0$ (les charges s'écoulent vers la Terre ou proviennent de celle-ci). Dans ce cas, le champ électrostatique mesuré à l'extérieur de (A_2) est nul, malgré la présence de (A_1) chargé à l'intérieur de (A_2). Ainsi, l'espace extérieur à (A_2) est protégé de toute influence électrostatique provenant de la cavité. L'inverse est également vrai.



Prenons maintenant le cas où (A_1) porte une charge nulle et où (A_2) est placé à proximité d'autres conducteurs chargés. A l'équilibre, on aura $Q_2^{\text{int}} = 0$ mais aussi un champ électrostatique non nul mesuré à l'extérieur de (A_2), dépendant de la distribution surfacique externe de (A_2). Ainsi, malgré la charge portée par la surface extérieure de (A_2), la cavité interne possède un champ électrostatique nul. Nous voyons donc que le champ électrostatique régnant à l'intérieur de (A_2) est parfaitement indépendant de celui à l'extérieur. Noter que ceci reste vrai même si (A_2) n'est pas maintenu à potentiel constant.

Une combinaison linéaire de ces deux situations permettant de décrire tous les cas possibles, nous venons de démontrer que tout conducteur creux maintenu à potentiel constant constitue bien un écran électrostatique dans les deux sens. Un tel dispositif est appelé **cage de Faraday**.

Alors que la distribution des charges Q_2^{int} dépend de la position de (A_1), celle des charges Q_2^{ext} portées par la surface externe de (A_2) dépend, elle, uniquement de ce qui se passe à l'extérieur.

Applications :

1. Protection contre la foudre : un paratonnerre est en général complété par un réseau de câbles entourant l'édifice à protéger, reliés à la Terre.
2. Tout conducteur transportant un courant faible est entouré d'une gaine métallique (appelée blindage) reliée au sol. Cette gaine est parfois simplement le châssis de l'appareil.

4.2.3 Coefficients d'influence électrostatique

Nous avons vu que lorsque plusieurs conducteurs sont mis en présence les uns des autres, ils exercent une influence électrostatique réciproque. A l'équilibre (mécanique et électrostatique), les densités surfaciques de chaque conducteur dépendent des charges qu'ils portent, de leur capacité et de leurs positions relatives. Si l'on cherche à calculer, par exemple, le potentiel pris par l'un des conducteurs, alors il nous faut résoudre le problème complet : calculer le potentiel de chaque conducteur.

Soit un ensemble de N conducteurs (A_i) de charge électrique totale Q_i et potentiel V_i , en équilibre électrostatique. Prenons (A_1) et appliquons la notion vue précédemment de superposition des états

d'équilibre. On peut toujours décomposer la distribution surfacique sur (A_1) de la forme $\sigma_1 = \sum_{j=1}^N \sigma_{1j}$ où σ_{11} est la densité surfacique de charges apparaissant sur (A_1) si tous les autres conducteurs étaient portés au potentiel nul (mais présents) et σ_{1j} celle apparaissant lorsque tous (y compris A_1) sont portés au potentiel nul, sauf (A_j) . On peut alors écrire que la charge totale sur (A_1) est :

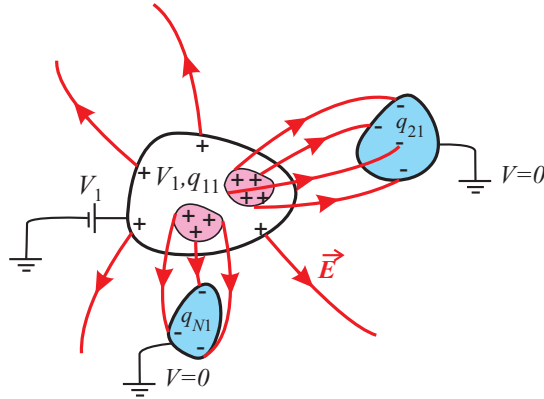
$$\begin{aligned} Q_1 &= \iint \sigma_1 dS = \sum_{j=1}^N \iint \sigma_{1j} dS = q_{11} + q_{12} + \dots + q_{1N} \\ &= C_{11}V_1 + C_{12}V_2 + \dots + C_{1N}V_N . \end{aligned} \quad (4.8)$$

Pour connaître Q_1 il faut donc connaître les N états d'équilibre électrostatique.

Un moyen de mesurer tous les coefficients $C_{i,1}$ est de considérer un premier état d'équilibre, celui où on garde le premier armature à potentiel V_1 et tous les autres armatures sont mis au potentiel nul. Utilisant un exposant ⁽¹⁾ pour indiquer qu'il s'agit du premier état d'équilibre, les charges, $Q_i^{(1)}$, sur chacun des armatures s'écrit :

$$\begin{aligned} Q_1^{(1)} &\equiv q_{11} = C_{11}V_1 \\ Q_2^{(1)} &\equiv q_{21} = C_{21}V_1 \\ &\vdots \equiv \vdots = \vdots \\ Q_N^{(1)} &\equiv q_{N1} = C_{N1}V_1 . \end{aligned} \quad (4.9)$$

En effet, la charge apparaissant sur (A_1) ne peut être due qu'à V_1 , C_{11} étant la capacité du conducteur



(A_1) en présence des autres conducteurs. Mais par influence, une distribution σ_{j1} apparaît sur tous les autres conducteurs (A_j) . Celle-ci dépend du nombre de lignes de champ qui joignent (A_1) à chaque conducteur (A_j) . En vertu du théorème des éléments correspondants, la charge qui « apparaît » est de signe opposé à celle sur (A_1) , elle-même proportionnelle à q_{11} donc à V_1 : les coefficients d'influence $C_{j1} (j \neq 1)$ sont donc négatifs.

Considérons maintenant le deuxième état d'équilibre, où tous les conducteurs sauf (A_2) sont mis au potentiel nul. On a alors dans ce cas

$$\begin{aligned} Q_1^{(2)} &\equiv q_{12} = C_{12}V_2 \\ Q_2^{(2)} &\equiv q_{22} = C_{22}V_2 \\ &\vdots \equiv \vdots = \vdots \\ Q_N^{(2)} &\equiv q_{N2} = C_{N2}V_2 . \end{aligned} \quad (4.10)$$

Bien évidemment, en reproduisant cette opération, on obtient que l'état d'équilibre le plus général est décrit par :

$$Q_i = \sum_{j=1}^N Q_i^{(j)} = \sum_{j=1}^N q_{ij} = \sum_{j=1}^N C_{ij} V_j, \quad (4.11)$$

ou sous forme matricielle :

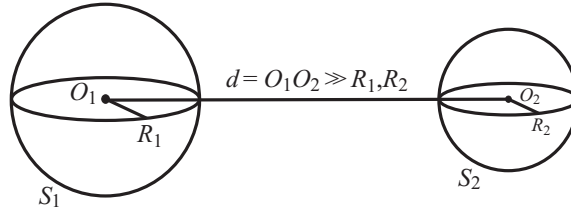
$$\begin{bmatrix} Q_1 \\ \vdots \\ Q_N \end{bmatrix} = \begin{bmatrix} C_{11} & \cdots & C_{1N} \\ \vdots & \ddots & \vdots \\ C_{N1} & \cdots & C_{NN} \end{bmatrix} \begin{bmatrix} V_1 \\ \vdots \\ V_N \end{bmatrix}. \quad (4.12)$$

Les coefficients C_{ij} sont appelés **coefficients d'influence**. Les coefficients C_{ii} sont parfois appelés **coefficients de capacité ou capacités des conducteurs en présence des autres**. Il ne faut pas les confondre avec les capacités propres C_i des conducteurs isolés, seuls dans l'espace. D'une façon générale, on a les propriétés suivantes :

1. Les C_{ii} sont toujours positifs.
2. Les C_{ij} sont toujours négatifs et $C_{ij} = C_{ji}$ (matrice symétrique).
3. $C_{ii} \geq -\sum_{j \neq i} C_{ji}$, l'égalité n'étant possible que dans le cas d'une influence totale.

La dernière inégalité est une conséquence du théorème des éléments correspondants. En effet, prenons le conducteur (A_1) porté au potentiel V_1 alors que les autres sont mis au potentiel nul. Tous les tubes de flux partant de (A_1) n'aboutissent pas nécessairement à un autre conducteur (ils ne le feraient que pour une influence totale). Donc, cela signifie que la charge totale située sur (A_1) est (en valeur absolue) supérieure à l'ensemble des charges situées sur les autres conducteurs, c'est-à-dire $q_{11} = C_{11}V_1 \geq |q_{21}| + \dots + |q_{N1}| = \sum_{j \neq 1} |C_{j1}| V_1$.

Exemple : Soient deux conducteurs sphériques, (A_1) et (A_2) , de rayons R_1 et R_2 portant une charge Q_1 et Q_2 , situés à une distance d l'un de l'autre.



1. *A quels potentiels se trouvent ces deux conducteurs ?*

En vertu du principe de superposition, le potentiel de (A_1) , pris en son centre O_1 est

$$V_1 = \frac{1}{4\pi\epsilon_0} \iint_{S_1} \frac{\sigma_1 dS_1}{P_1 O_1} + \frac{1}{4\pi\epsilon_0} \iint_{S_2} \frac{\sigma_2 dS_2}{P_2 O_1},$$

où le premier terme est dû aux charges Q_1 et le second à celles situées sur (A_2) . Lorsque la distance d est beaucoup plus grande que les rayons, on peut assimiler $P_2 O_1 \simeq O_2 O_1 = d$ pour tout point P_2 de la surface de (A_2) et l'on obtient

$$V_1 \simeq \frac{Q_1}{4\pi\epsilon_0 R_1} + \frac{Q_2}{4\pi\epsilon_0 d} = \frac{Q_1}{C_1} + \frac{Q_2}{C_d},$$

où l'on reconnaît en C_1 la capacité d'une sphère isolée et en $C_d = 4\pi\epsilon_0 d$, un coefficient qui dépend à la fois de la géométrie des deux conducteurs et de leur distance. En faisant de même pour (A_2) , on obtient

$$V_2 \simeq \frac{Q_2}{4\pi\epsilon_0 R_2} + \frac{Q_1}{4\pi\epsilon_0 d} = \frac{Q_2}{C_2} + \frac{Q_1}{C_d},$$

où C_2 est la capacité de (A_2) isolée.

2. Trouver les coefficients de capacité de ce système.

A partir des deux expressions approximatives pour V_1 et V_2 précédentes, on obtient donc un problème linéaire qui peut se mettre sous la forme matricielle suivante :

$$\begin{bmatrix} \frac{1}{C_1} & \frac{1}{C_d} \\ \frac{1}{C_d} & \frac{1}{C_2} \end{bmatrix} \begin{bmatrix} Q_1 \\ Q_2 \end{bmatrix} \simeq \begin{bmatrix} V_1 \\ V_2 \end{bmatrix}, \quad (4.13)$$

c'est-à-dire $V_i = D_{ij}Q_j$ où la matrice D_{ij} est connue à partir de l'inverse des diverses capacités. Si l'on veut se ramener au problème précédent (calcul des charges connaissant les potentiels), c'est-à-dire à la résolution de $Q_i = C_{ij}V_j$, où C_{ij} est la matrice des coefficients d'influence, il faut inverser la matrice D_{ij} .

On se rappelle comment inverser une matrice 2×2 :

$$A = \begin{bmatrix} a & c \\ d & b \end{bmatrix}, \quad A^{-1} = \frac{1}{|A|} \begin{bmatrix} b & -c \\ -d & a \end{bmatrix} = \frac{1}{ab - cd} \begin{bmatrix} b & -c \\ -d & a \end{bmatrix} \quad (4.14)$$

Avec l'éq. (4.13) et l'éq. (4.14), on obtiendra en effet $Q_i = D_{ij}^{-1}V_j$, donc $C_{ij} = D_{ij}^{-1}$ avec :

$$C_{11} = \frac{C_1}{1 - \frac{C_1 C_2}{C_d^2}}, \quad C_{22} = \frac{C_2}{1 - \frac{C_1 C_2}{C_d^2}}, \quad C_{12} = C_{21} = -\frac{\frac{C_1 C_2}{C_d}}{1 - \frac{C_1 C_2}{C_d^2}}. \quad (4.15)$$

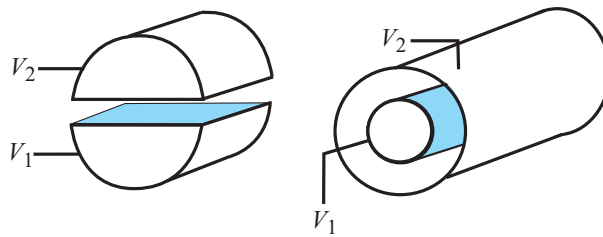
On voit clairement sur cet exemple que, premièrement, les capacités en présence des autres conducteurs C_{ii} ne sont pas identifiables aux capacités propres C_i des conducteurs isolés dans l'espace et deuxièmement, que les coefficients d'influence C_{ij} sont bien négatifs.

4.3 Le condensateur

4.3.1 Condensation de l'électricité

Définition : On appelle condensateur tout système de deux conducteurs en influence électrostatique. Il y a deux sortes de condensateurs :

- à armatures rapprochées (à influence quasi totale)
- à influence totale



En général, les deux armatures sont séparées par un matériau isolant (un diélectrique), ce qui a pour effet d'accroître la capacité du condensateur. Dans ce qui suit on suppose qu'il n'y a que du vide. Soient donc deux conducteurs (A_1) et (A_2) portant une charge totale Q_1 et Q_2 et de potentiels V_1 et V_2 . D'après la section précédente, on a $Q_1 = C_{11}V_1 + C_{12}V_2$ et $Q_2 = C_{21}V_1 + C_{22}V_2$, c.-à.-d. :

$$\begin{bmatrix} Q_1 \\ Q_2 \end{bmatrix} = \begin{bmatrix} C_{11} & C_{12} \\ C_{21} & C_{22} \end{bmatrix} \begin{bmatrix} V_1 \\ V_2 \end{bmatrix}. \quad (4.16)$$

Les coefficients C_{ij} étant indépendants des valeurs de Q et de V , il suffit, pour les trouver, de considérer des cas particuliers simples (formellement on a ici 2 équations à 4 inconnues).

Regardons ce qui se passe dans le cas d'un condensateur à influence totale, c'est-à-dire un condensateur pour lequel on a :

$$Q_2 = Q_2^{\text{ext}} + Q_2^{\text{int}} = Q_2^{\text{ext}} - Q_1 . \quad (4.17)$$

Si on relie (A_2) à la masse ($V_2 = 0$, $Q_2^{\text{ext}} = 0$ car on néglige toute influence extérieure), alors on obtient

$$\begin{cases} Q_2 = -Q_1 \\ C_{21} = -C_{11} . \end{cases} \quad (4.18)$$

La première relation n'est vraie que si (A_2) est à la masse, mais la seconde est générale. Par ailleurs, on sait par l'analyse dans la section précédent que $C_{12} = C_{21}$. Par convention, la capacité C du condensateur, sa charge Q et sa tension entre armatures sont alors définies de la façon suivante,

$$\begin{aligned} C &= C_{11} \\ Q &= Q_1 \\ U &= V_1 - V_2 . \end{aligned} \quad (4.19)$$

Mettant ces relations dans l'éq. (4.16) on voit que la première ligne nous donne la relation des condensateurs

$$Q = CU . \quad (4.20)$$

Remarques :

1. Pourquoi appelle-t-on ces dispositifs des condensateurs? Parce qu'ils permettent de mettre en évidence le phénomène de « condensation de l'électricité », à savoir l'accumulation de charges électriques dans une petite zone de l'espace. Ainsi, en construisant des condensateurs de capacité C élevée, on obtient des charges électriques Q élevées avec des tensions U faibles.
2. La charge située sur l'armature (A_2) est $Q_2 = Q_2^{\text{ext}} - Q$ (pour un condensateur à influence totale) et, en toute rigueur, ne vaut $-Q$ que lorsque (A_2) est mise à la masse. En général, elle reste cependant négligeable devant Q dans les cas considérés dans ce cours et on n'en tiendra donc pas compte.

Pour un condensateur à armatures rapprochées, on obtient le même résultat, moyennant une séparation faible (devant leur taille) des conducteurs. Dans ce type de condensateur, les charges Q_1 et Q_2 correspondent à celles qui se trouvent réparties sur l'ensemble de la surface de chaque conducteur. Mais si la distance est faible, l'influence électrostatique va condenser les charges sur les surfaces en regard, de telle sorte que l'on peut faire l'hypothèse suivante :

$$\begin{aligned} Q_1 &= Q_1^{\text{ext}} + Q_1^S \simeq Q_1^S \\ Q_2 &= Q_2^{\text{ext}} + Q_2^S = Q_2^{\text{ext}} - Q_1^S \simeq Q_2^S - Q_1 , \end{aligned} \quad (4.21)$$

ce qui nous ramène à une expression identique à celle d'un condensateur à influence totale.

4.3.2 Capacités de quelques condensateurs simples

Dans ce qui suit, nous allons voir plusieurs exemples de calculs de capacités. Pour obtenir la capacité C d'un condensateur, il faut calculer la relation entre sa charge Q et sa tension U , c'est-à-dire :

$$U = V_1 - V_2 = \int_1^2 \vec{E} \cdot d\vec{\ell} = \frac{Q}{C} . \quad (4.22)$$

Autrement dit, il faut être capable de calculer la circulation du champ électrostatique entre les deux armatures ainsi que la charge Q .

- (a) **Condensateur sphérique** : Soit un condensateur constitué de deux armatures sphériques de même centre O , de rayons respectifs R_1 et R_2 , séparées par un vide. ($R_2 > R_1$). D'après le théorème de Gauss, le champ électrostatique en un point M situé à un rayon r entre les deux armatures vaut :

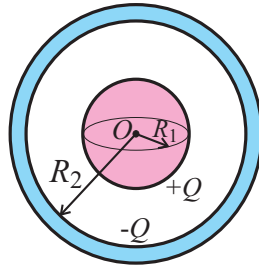
$$\vec{E}(r) = \frac{Q}{4\pi\epsilon_0 r^2} \hat{r},$$

en coordonnées sphériques, ce qui donne une tension

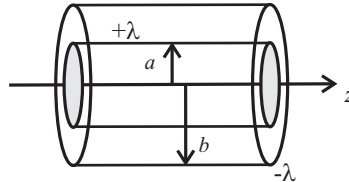
$$U = V_1 - V_2 = \int_{R_1}^{R_2} \vec{E} \cdot d\vec{\ell} = \frac{Q}{4\pi\epsilon_0} \int_{R_1}^{R_2} \frac{dr}{r^2} = \frac{Q}{4\pi\epsilon_0} \left(\frac{1}{R_1} - \frac{1}{R_2} \right) = \frac{Q}{4\pi\epsilon_0} \frac{R_2 - R_1}{R_1 R_2},$$

et fournit donc une capacité totale :

$$C = \frac{Q}{U} = 4\pi\epsilon_0 \frac{R_1 R_2}{R_2 - R_1}. \quad (4.23)$$



- (b) **Condensateur cylindrique** : Soit un condensateur constitué de deux armatures cylindriques coaxiales de rayons a et b séparées par un vide ($b > a$) et de longueurs ℓ quasi-infinie ($\ell \gg b$). Soit $\lambda = Q/\ell$ la charge par unité de longueur du cylindre intérieur. D'après le théorème de



Gauss, le champ électrostatique entre les deux armatures s'écrit :

$$\vec{E}(\rho) = \frac{\lambda}{2\pi\epsilon_0 \rho} \hat{\rho},$$

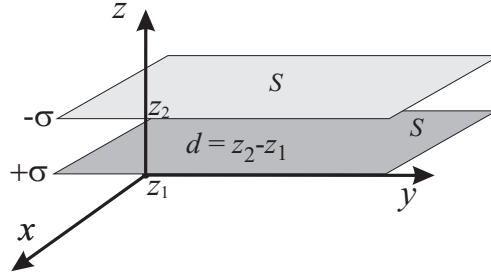
en coordonnées cylindriques, ce qui donne une différence de tension :

$$U = V_1 - V_2 = \int_a^b \vec{E} \cdot d\vec{\ell} = \frac{\lambda}{2\pi\epsilon_0} \int_a^b \frac{d\rho}{\rho} = \frac{Q}{2\pi\epsilon_0 \ell} (\ln b - \ln a),$$

et une capacité par unité de longueur,

$$C/\ell = \frac{Q}{\ell U} = \frac{2\pi\epsilon_0}{\ln \frac{b}{a}}. \quad (4.24)$$

- (c) **Condensateur plan** Soient deux armatures (A_1) et (A_2) planes parallèles, orthogonales à un même axe Oz de vecteur unitaire $\hat{z}$, de surface S et situées à une distance $d = z_2 - z_1$ l'une de l'autre. On utilise l'approximation de planes infinies, c.-à-d. $S \gg d$. L'armature (A_1) porte une densité surfacique de charges σ et (A_2), en vertu du théorème des éléments correspondants,



porte une densité $-\sigma$. Entre les deux armatures, le champ électrostatique est la superposition des champs créés par ces deux plans infinis, c'est-à-dire

$$\vec{E} = \vec{E}_1 + \vec{E}_2 = \frac{\sigma}{2\epsilon_0} \hat{z} + \frac{-\sigma}{2\epsilon_0} (-\hat{z}) = \frac{\sigma}{\epsilon_0} \hat{z}.$$

La différence de potentiel entre les deux armatures est alors

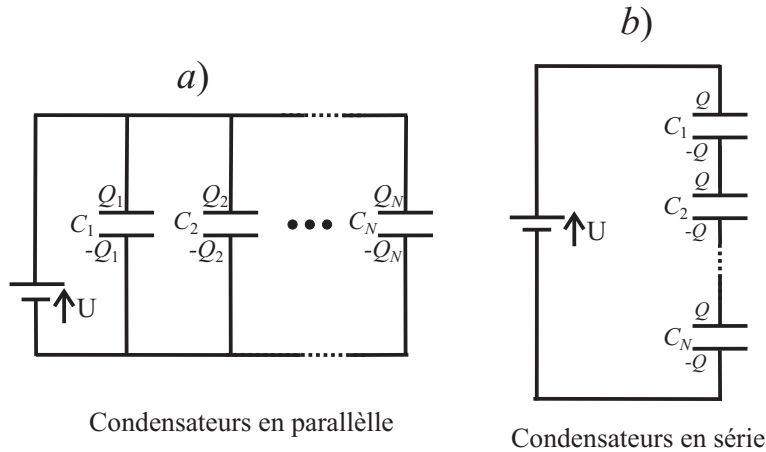
$$U = V_1 - V_2 = \int_{z_1}^{z_2} \vec{E} \cdot d\vec{\ell} = \frac{\sigma}{\epsilon_0} \int_{z_1}^{z_2} dz = \frac{\sigma}{\epsilon_0} d,$$

d'où une capacité

$$C = \frac{Q}{U} = \frac{\sigma S}{U} = \frac{S\epsilon_0}{d}. \quad (4.25)$$

Dans la pratique, cette relation s'applique à tous les condensateurs dans le vide (de façon approximative) à condition que les dimensions de la surface S sont très grandes devant la distance de séparation d des armatures.

4.3.3 Associations de condensateurs



- a) **Condensateurs en parallèle** Soient N condensateurs de capacités C_i mis en parallèle avec la même tension $U = V_1 - V_2$. La charge électrique de chacun d'entre eux est donnée par $Q_i = C_i U$. La charge électrique totale est simplement :

$$Q = \sum_i^N Q_i = \left(\sum_i^N C_i \right) U,$$

ce qui correspond à une capacité équivalente :

$$C_{eq} = \sum_i^N C_i, \quad (4.26)$$

qui est la somme des capacités individuelles.

- b) **Condensateurs en série** Soient N condensateurs de capacités C_i mis en série les uns derrière les autres. On porte aux potentiels V_0 et V_N les deux extrémités de la chaîne et on apporte la charge Q sur le premier condensateur. En supposant que tous les condensateurs sont initialement neutres, il s'établit la charge $\pm Q$ (par influence) sur les armatures des condensateurs adjacents. La tension totale aux bornes de la chaîne de condensateurs s'écrit alors simplement

$$\begin{aligned} U = V_0 - V_N &= (V_0 - V_1) + (V_1 - V_2) + \dots + (V_{N-1} - V_N) \\ &= \frac{Q}{C_1} + \frac{Q}{C_2} + \dots + \frac{Q}{C_N} = \left(\sum_i^N \frac{1}{C_i} \right) Q, \end{aligned}$$

et correspond à celle d'une capacité unique C de capacité équivalente

$$\frac{1}{C_{\text{eq}}} = \sum_i^N \frac{1}{C_i}. \quad (4.27)$$

4.4 Complément : Relations de continuité du champ électrique

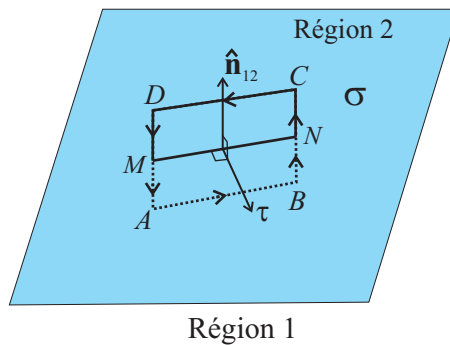
Dans ce chapitre, nous avons vu qu'il y a une discontinuité du champ électrique en présence de charges surfaciques d'un conducteur (théorème de Coulomb). Ici, on fait appel aux lois fondamentales afin d'exprimer les contraintes générales sur le champ à la traversée de n'importe quelle nappe de charge surfacique.

Soit une distribution surfacique de charge σ sur une surface Σ séparant l'espace en deux régions 1 et 2. Pour la composante tangentielle du champ, nous remarquons que puisque le champ électrostatique $\vec{E}$ s'exprime comme un gradient du potentiel, on a :

$$\oint_C \vec{E} \cdot d\vec{\ell} = - \oint_C dV = 0.$$

sur n'importe quel contour $\mathcal{C}$.

Considérons maintenant le contour $ABCD$ traversant la surface de charge : La circulation du



champ sur ce contour s'écrit alors :

$$\int_{AB} \vec{E} \cdot d\vec{\ell} + \int_{BC} \vec{E} \cdot d\vec{\ell} + \int_{CD} \vec{E} \cdot d\vec{\ell} + \int_{DA} \vec{E} \cdot d\vec{\ell} = 0.$$

Laissant les cotés DA et BC tendre vers zéro, on obtient

$$\int_{MN} (\vec{E}_1 - \vec{E}_2) \cdot d\vec{\ell} = 0.$$

Puisque MN est quelconque (sur la surface), on doit avoir :

$$(\vec{E}_1 - \vec{E}_2) \cdot d\vec{\ell} = 0,$$

ce qui veut dire que la composante tangentielle du champ à une surface doit être continue même en présence de charges surfaciques.

Une façon courante d'exprimer cette propriété de continuité est de remarquer qu'on peut écrire

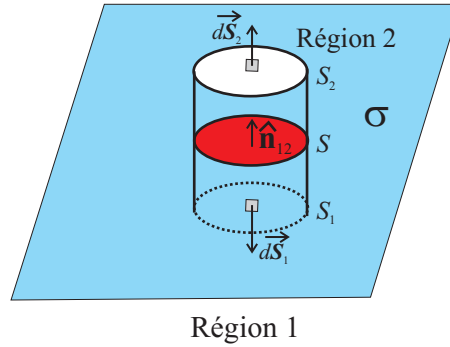
$$\begin{aligned} (\vec{E}_1 - \vec{E}_2) \cdot d\vec{\ell} &= (\vec{E}_1 - \vec{E}_2) \cdot (\hat{n}_{12} \wedge \hat{\tau}) d\ell \\ &= [(\vec{E}_1 - \vec{E}_2) \wedge \hat{n}_{12}] \cdot \hat{\tau} d\ell, \end{aligned}$$

c'est-à-dire (puisque la direction de $\hat{\tau}$ est arbitraire)

$$\boxed{\hat{n}_{12} \wedge (\vec{E}_2 - \vec{E}_1) = \vec{0}}, \quad (4.28)$$

sur n'importe quelle surface.

Considérons maintenant une surface fermée fictive en forme de cylindre fermé contenant une nappe de charge surfacique séparant les deux régions :



On applique le théorème de Gauss à cette surface qui s'écrit :

$$\iint_{S_1} \vec{E} \cdot d\vec{S}_1 + \iint_{S_2} \vec{E} \cdot d\vec{S}_2 + \iint_{S_L} \vec{E} \cdot d\vec{S}_L = \frac{Q_{\text{int}}}{\epsilon_0} = \frac{1}{\epsilon_0} \iint_S \sigma dS,$$

où S_L est la surface latérale. Lorsqu'on fait tendre cette surface vers zéro (S_1 tend vers S_2), on obtient

$$\iint_{S_2} \vec{E}_2 \cdot d\vec{S}_2 + \iint_{S_1} \vec{E}_1 \cdot d\vec{S}_1 = \frac{1}{\epsilon_0} \iint_S \sigma dS \Rightarrow \iint_S (\vec{E}_2 - \vec{E}_1) \cdot \hat{n}_{12} dS = \frac{1}{\epsilon_0} \iint_S \sigma dS,$$

puisque $d\vec{S}_2 = -d\vec{S}_1 = dS \hat{n}_{12}$ dans cette limite. Ce résultat étant valable quelque soit la surface S choisie, on vient donc de démontrer que :

$$\boxed{\hat{n}_{12} \cdot (\vec{E}_2 - \vec{E}_1) = \frac{\sigma}{\epsilon_0}} \quad (4.29)$$

En résumé avec une densité surfacique de charge σ sur une surface Σ séparant deux milieux 1 et 2 :

- la composante tangentielle (à la surface) du champ électrique est continue.
- la composante normale du champ électrique est discontinue en présence d'une charge surfacique σ , et cette discontinuité s'écrit comme $E_{2,\perp} - E_{1,\perp} = \frac{\sigma}{\epsilon_0}$ sur toute la surface.

Chapitre 5

Dipôle électrique - Energie électrostatique

Nous avons étudié jusqu'ici les lois de l'électrostatique valables dans le vide et dans les métaux parfaitement conducteurs. Nous nous proposons maintenant de commencer d'étendre ces lois à des matériaux autres que les métaux, les diélectriques.

En contraste avec les conducteurs qui ont une grande quantité de charges libres se déplaçant à l'intérieur du matériau, la grande majorité des charges dans les diélectriques peuvent difficilement se déplacer et sont liées aux atomes ou molécules du matériel. Ce qu'il faut comprendre est que même si des systèmes comme des atomes ou molécules sont globalement neutres, ça ne veut pas dire qu'ils ont un comportement électrique nul et qu'il faut en tenir compte dans la description des matériaux. Le modèle qui va nous permettre de comprendre le comportement des constituants des diélectriques est celui du dipôle électrique.

5.1 Le dipôle électrostatique

5.1.1 Potentiel électrostatique créé par deux charges électriques



Il existe dans la nature des systèmes électriquement neutres mais dont le centre de gravité des charges négatives n'est pas confondu avec celui des charges positives. Un tel système peut souvent être décrit (on dit modélisé) en première approximation par deux charges électriques ponctuelles, $+q$ et $-q$ situées à une distance $d = 2a$ l'une de l'autre. On appelle un tel système de charges un dipôle électrostatique.

Définition : Quand deux charges égales et opposées sont séparées par une distance $\overrightarrow{BA}$, (vecteur de position relative de la charge positive par rapport à la charge négative) on définit le moment dipolaire électrique, $\vec{p}$, tel que :

$$\vec{p} \equiv q\overrightarrow{BA} = q|\overrightarrow{BA}|\hat{r}_{BA} \equiv qd\hat{r}_{BA} \equiv p\hat{r}_{BA} . \quad (5.1)$$

- Dans la nature, certains molécules sont caractérisés par un moment dipolaire permanent, tels que l'eau, le chlorure d'hydrogène, et monoxyde de carbone ; alors que d'autres, comme le dioxyde de carbone sont dépourvus de moment dipolaire permanent. (voir la figure 5.1.1)
- Les unités SI (de C.m) ne sont pas commodes pour exprimer le moment dipolaire des molécules. Par conséquent, une unité du système CGS est toujours assez utilisé pour mesurer le moment dipolaire des molécules. Il s'agit du debye (symbole D), avec $1 \text{ D} = 3,33564 \times 10^{-30} \text{ C.m}$. Le moment dipolaire de l'eau par exemple est $1,8546 \pm 0,0040 \text{ D}$.

Connaître l'effet (la force) électrostatique que ces deux charges créent autour d'eux nécessite le calcul du champ électrostatique. Nous aurions pu appliquer le principe de superposition aux champs électriques et calculer ainsi la somme vectorielle des champs électriques créé par chacune des charges ($\pm q$). Il s'avère plus simple de calculer le potentiel créé produit par le dipôle et calculer le champ électrique en appliquant la formule $\mathbf{E} = -\text{grad}V$. C'est ce que nous allons faire avec l'aide de la figure 5.1.1 ci-dessous.

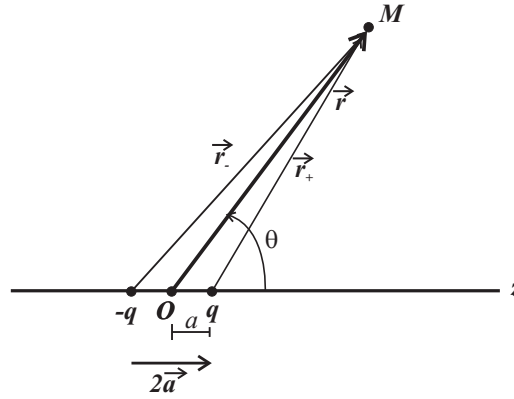


FIGURE 5.1 – Repérage du point M par rapport à un dipôle idéal, $\mathbf{p} = p\hat{\mathbf{z}}$ avec $p = 2aq$.

Le système est invariant par rapport à des rotations autour de l'axe de symétrie, z , donc le potentiel ne dépend pas de la coordonnée ϕ . D'après la figure 5.1.1 et le chapitre précédent, le potentiel, $V(M)$, créé en un point M repéré par ses coordonnées sphériques (r, θ, ϕ) est simplement :

$$\begin{aligned} V(M) &= V_{+q}(M) + V_{-q}(M) \\ &= \frac{q}{4\pi\epsilon_0} \left(\frac{1}{r_+} - \frac{1}{r_-} \right) = \frac{q}{4\pi\epsilon_0} \left(\frac{r_- - r_+}{r_+ r_-} \right), \end{aligned}$$

où l'on a choisi arbitrairement $V = 0$ à l'infini et $r_{\pm} \equiv |\vec{r}_{\pm}|$ avec $\vec{r}_{\pm} = \vec{r} \mp a\hat{\mathbf{z}}$. Lorsqu'on ne s'intéresse qu'à l'action électrostatique à grande distance, c'est-à-dire à des distances $r \gg a$, on peut faire un développement limité de $V(M)$. Au premier ordre en a/r , on obtient :

$$r_{\pm} = (\vec{r}_{\pm} \cdot \vec{r}_{\pm})^{1/2} = r \left(1 \mp 2\frac{a}{r} \hat{\mathbf{r}} \cdot \hat{\mathbf{z}} + \frac{a^2}{r^2} \right)^{1/2} \simeq r \left(1 \mp 2\frac{a}{r} \hat{\mathbf{r}} \cdot \hat{\mathbf{z}} \right)^{1/2} \simeq r \left(1 \mp \frac{a}{r} \cos \theta \right),$$

c'est-à-dire : $r_- - r_+ = 2a \cos \theta + O((a/r)^2)$, et $r_- r_+ = r^2 + O((a/r)^2)$. Le potentiel créé à grande distance par un dipôle électrostatique est donc donné par une expression facile à retenir,

$$V(\vec{r}) = V(r, \theta) = \frac{2aq \cos \theta}{4\pi\epsilon_0 r^2} = \frac{\vec{p} \cdot \hat{\mathbf{r}}}{4\pi\epsilon_0 r^2}. \quad (5.2)$$

où $p = qd = 2aq$.

5.1.2 Champ électrique à grande distance d'un dipôle : $r \gg d$

Pour calculer le champ électrostatique, il nous suffit maintenant d'utiliser $\mathbf{E} = -\overrightarrow{\text{grad}}V$ en coordonnées sphériques. On obtient ainsi :

$$\vec{\mathbf{E}}(r, \theta, \phi) = \begin{pmatrix} E_r = -\frac{\partial V}{\partial r} = \frac{2p \cos \theta}{4\pi\epsilon_0 r^3} \\ E_\theta = -\frac{1}{r} \frac{\partial V}{\partial \theta} = \frac{p \sin \theta}{4\pi\epsilon_0 r^3} \\ E_\phi = -\frac{1}{r \sin \theta} \frac{\partial V}{\partial \phi} = 0 \end{pmatrix}. \quad (5.3)$$

En coordonnées sphériques, le champ électrique s'exprime donc (voir l'éq. (1.24)) :

$$\vec{\mathbf{E}}(\vec{\mathbf{r}}) = \frac{1}{4\pi\epsilon_0} \frac{p}{r^3} \left(2 \cos \theta \hat{\mathbf{r}} + \sin \theta \hat{\boldsymbol{\theta}} \right), \quad (5.4)$$

où il ne faut pas oublier que θ est l'angle entre $\vec{\mathbf{p}}$ et $\vec{\mathbf{r}}$.

Dans l'éq. (5.4), on est parfois gêné par la dépendance explicite en θ par rapport à l'axe du dipôle. On peut mettre ce résultat sous une forme plus facile à retenir en remarquant que nous avons dérivé le champ électrique en prenant $\vec{\mathbf{p}} = 2a\hat{\mathbf{z}}$ et l'expression du vecteur unitaire $\hat{\mathbf{z}}$ en coordonnées sphériques s'écrit donc :

$$\hat{\mathbf{z}} = \cos \theta \hat{\mathbf{r}} - \sin \theta \hat{\boldsymbol{\theta}},$$

ce qui permet d'écrire le champ électrique sous la forme compacte :

$$\vec{\mathbf{E}}(M) = \vec{\mathbf{E}}(\vec{\mathbf{r}}) = \frac{1}{4\pi\epsilon_0} \frac{3(\vec{\mathbf{p}} \cdot \hat{\mathbf{r}}) \hat{\mathbf{r}} - \vec{\mathbf{p}}}{r^3}. \quad (5.5)$$

On aurait pu aussi bien se rappeler que $\cos \theta = z/r$ où $r = (x^2 + y^2 + z^2)^{1/2}$, afin de réécrire l'expression de l'éq. (5.2) en coordonnées cartésiennes,

$$V(x, y, z) = \frac{p}{4\pi\epsilon_0} \frac{\cos \theta}{r^2} = \frac{p}{4\pi\epsilon_0} \frac{z}{(x^2 + y^2 + z^2)^{3/2}},$$

ce qui nous permettrait de calculer le champ $\vec{\mathbf{E}}$ du dipôle directement en coordonnées cartésiennes,

$$\vec{\mathbf{E}}(x, y, z) = - \begin{pmatrix} \frac{\partial V}{\partial x} \\ \frac{\partial V}{\partial y} \\ \frac{\partial V}{\partial z} \end{pmatrix} = \frac{p}{4\pi\epsilon_0} \begin{pmatrix} \frac{3zx}{r^5} \\ \frac{3zy}{r^5} \\ \frac{2z^2 - x^2 - y^2}{r^5} \end{pmatrix}, \quad (5.6)$$

en se rappelant que cette formule s'applique seulement quand le dipôle est orienté sur l'axe Oz où $\vec{\mathbf{p}} = p\hat{\mathbf{z}}$, alors que l'expression de l'éq. (5.5) s'applique à n'importe quelle orientation du dipôle.

Par construction, notre dipôle modèle possède une symétrie de révolution autour de l'axe qui le porte (ici l'axe Oz) : le potentiel ainsi que le champ électrostatique possèdent donc également cette symétrie. Cela va nous aider à visualiser les lignes de champ ainsi que les équipotentiels. Par exemple, le plan médiateur défini par $\theta = \pi/2$ ($z = 0$) est une surface équipotentielle $V = 0$. Les équipotentiels sont des surfaces (dans l'espace 3D ; dans le plan ce sont des courbes) définies par $V = \text{Constante} = V_0$, c'est-à-dire :

$$r(\theta) = \sqrt{\frac{p \cos \theta}{4\pi\epsilon_0 V_0}}. \quad (5.7)$$

L'équation des lignes de champ est obtenue à partir de l'équation $\vec{\mathbf{E}} \wedge d\vec{\boldsymbol{\ell}} = \vec{\mathbf{0}}$ et on tenant compte du fait que la composant E_ϕ du champ dipolaire est nulle, on trouve :

$$\begin{pmatrix} E_r \\ E_\theta \\ 0 \end{pmatrix} \wedge \begin{pmatrix} dr \\ r d\theta \\ r \sin \theta d\phi \end{pmatrix} = \vec{\mathbf{0}} \Rightarrow d\phi = 0 \text{ et } E_r r d\theta = E_\theta dr \Rightarrow \frac{dr}{r} = \frac{E_r}{E_\theta} d\theta = \frac{2 \cos \theta d\theta}{\sin \theta}. \quad (5.8)$$

L'équation des lignes de champs, $\vec{E}$, est ensuite obtenue par intégration :

$$\int \frac{dr}{r} = 2 \int \frac{d(\sin \theta)}{\sin \theta} \Rightarrow \ln r = \ln \sin^2 \theta + C \Rightarrow r(\theta) = K \sin^2 \theta . \quad (5.9)$$

où K est une constante d'intégration dont la valeur (arbitraire) définit la ligne de champ.

Il ne faut pas oublier que les équations des surfaces équipotentielles, eq. (5.7) et les lignes du champ, ont été dérivées dans l'approximation (5.9). Si on veut connaître les lignes du champ près de la distribution de charge, il faut se renseigner sur les détails précis sur la distribution des charges du dipôle. Les équipotentielles et les lignes du champ du dipôle modèle sont illustré dans la figure 5.2. On peut constater dans cette figure que, *loin du dipôle*, les équipotentielles et les lignes du champ correspondent aux fonctions, de $r(\theta)$, décrites respectivement dans les l'éqs. (5.7) et (5.9).

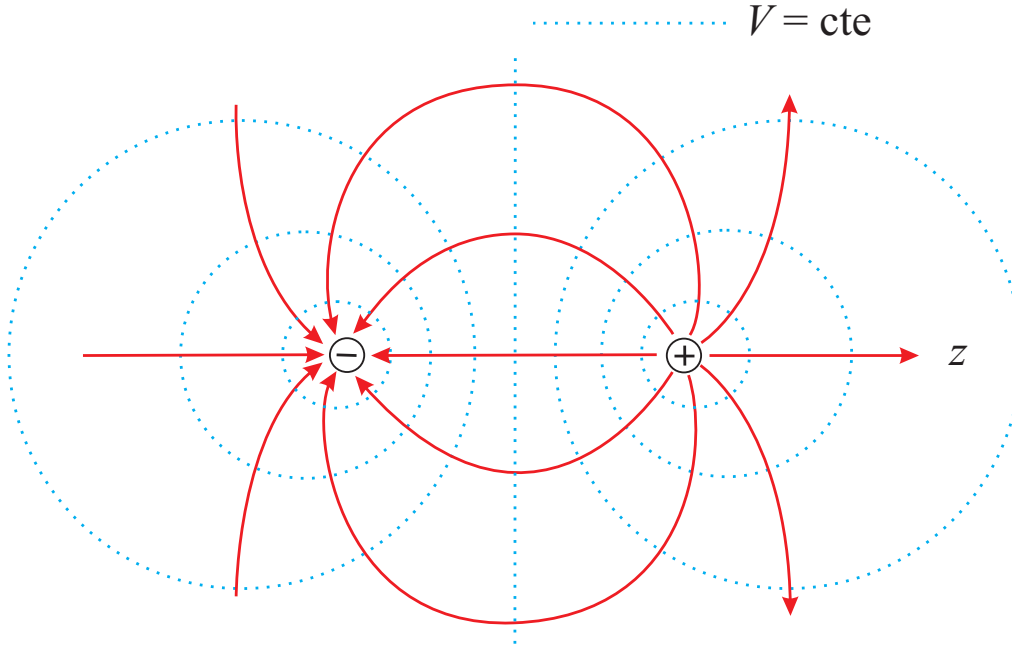


FIGURE 5.2 – Lignes de champ (en rouge - trait plein) et équipotentielles (en bleu - en traits pointillés) d'un modèle de dipôle électrique à une échelle où la séparation des charges + et - n'est pas négligeable.

5.1.3 Développements multipolaires

Lorsqu'on a affaire à une distribution de charges électriques et qu'on ne s'intéresse qu'au champ créé à une grande distance devant les dimensions de cette distribution, on peut également utiliser une méthode de calcul approchée du potentiel. Le degré de validité de ce calcul dépend directement de l'ordre du développement limité utilisé : plus on va à un ordre élevé et meilleure sera notre approximation. Par exemple, l'expression du dipôle ci-dessus n'est valable que pour $r \gg a$, mais lorsque r tend vers a , il faut prendre en compte les ordres supérieurs, les termes dits multipolaires.

Prenons le cas d'une distribution de N charges ponctuelles q_i situées en $\vec{r}_i = \overrightarrow{OP}_i$. Le potentiel créé en un point M repéré par le vecteur position $\vec{r} = \overrightarrow{OM}$ (coordonnées sphériques) est

$$V(\vec{r}) = \sum_{i=1}^N \frac{1}{4\pi\epsilon_0} \frac{q_i}{|\vec{r} - \vec{r}_i|} . \quad (5.10)$$

En supposant $r \gg r_i$, on peut montrer que ce potentiel admet le développement suivant

$$V(\vec{r}) = \frac{1}{4\pi\epsilon_0} \sum_{i=1}^N \left(\frac{q_i}{r} + \frac{q_i r_i \cos \theta_i}{r^2} + \frac{q_i r_i^2}{r^3} (3 \cos^2 \theta_i - 1) + \dots \right) .$$

où θ_i est l'angle entre $\vec{r}$ et $\vec{r}_i$. Faire un développement multipolaire d'une distribution quelconque de charges consiste à arrêter le développement limité à un ordre donné, dépendant du degré de précision souhaité. Dans le développement ci-dessus, le premier terme (ordre zéro ou monopolaire) correspond à assimiler la distribution à une charge totale placée en O . Cela peut être suffisant vu de très loin, si cette charge totale est non nulle. Dans le cas contraire (ou si l'on souhaite plus de précision) on obtient le deuxième terme qui peut se mettre sous la forme

$$\frac{\vec{p} \cdot \hat{r}}{4\pi\epsilon_0 r^2},$$

où le vecteur

$$\vec{p} = \sum_{i=1}^N q_i \vec{r}_i, \quad (5.11)$$

est le **moment dipolaire** du système de charges. Si la charge est décrite par une distribution volumique de charge, le moment dipolaire s'écrit respectivement :

$$\vec{p} = \iiint \rho(P) \vec{OP} dV \quad \vec{p} = \iint \sigma(P) \vec{OP} dS. \quad (5.12)$$

En la définition du moment dipolaire ci-dessus, on aurait pu s'inquiéter du fait qu'il semble que le moment dipolaire dépend du choix de l'origine, O . En réalité, le moment dipolaire ne dépend pas du choix d'origine à condition que la charge totale soit nulle. Comme démonstration, imaginons qu'on calcule le moment dipolaire autour d'une autre origine O' . Le moment dipolaire dans ce nouveau système est :

$$\begin{aligned} \vec{p}' &= \iiint \rho(P) \vec{O'P} dV \\ &= \iiint \rho(P) (\vec{O'O} + \vec{OP}) dV \\ &= \vec{O'O} \iiint \rho(P) dV + \iiint \rho(P) \vec{OP} dV \\ &= \vec{O'O} Q_T + \vec{p}. \end{aligned}$$

Donc $\vec{p}' = \vec{p}$ à condition que la charge totale du système, Q_T , soit nulle.

5.2 Diélectriques

L'expérience démontre que le modèle du dipôle électrique est souvent suffisant pour décrire le comportement des constituants fondamentaux des objets diélectriques. Il faut se rappeler une fois encore que les milieux matériels ou liquides macroscopiques sont composés d'un nombre gigantesque d'atomes et/ou molécules.

Certains constituants comme les molécules d'eau ont un moment dipolaire mais en absence de champs électriques, ces dipôles sont orientés aléatoirement afin que le moment dipolaire global de l'eau soit quasi nul. D'autres constituants (comme le CO_2) n'ont pas de moment dipolaire propre mais peuvent acquérir un moment dipolaire sous l'effet d'un champ électrique $\vec{E}$. Dans ces deux cas, on parle de **moment dipolaires induits** (en présence de champ électriques). Quelques rares matériaux comme les ferroélectriques, ont des moments dipolaires qui ont tendance à s'aligner les uns avec les autres. On parle dans ces cas de **moments dipolaires permanents** qui peuvent persister même en absence de champ électriques. On appelle tout ces types de matériaux des diélectriques et plus précisément :

Un milieu diélectrique est, par définition, une substance possédant l'une des propriétés suivantes

- tout volume $d\mathcal{V}$ de la substance possède un moment électrique dipolaire $\vec{d}\vec{p}$, **on dit dans ce cas que la polarisation est permanente** ;
- tout volume $d\mathcal{V}$ de la substance **est susceptible d'acquérir** sous l'action d'un champ électrique extérieur $\vec{E}$ un moment dipolaire électrique $\vec{d}\vec{p}$: **la polarisation est dite induite par le champ $\vec{E}$** .

Puisque une petite quantité macroscopique d'un diélectrique contient un grand nombre d'atomes et/ou molécules, on décrit généralement les diélectriques par un **vecteur polarisation**, $\vec{P}$, qui peut être **interprété** comme une **densité volumique du moment dipolaire** telle que $\vec{d}\vec{p} = \vec{P}d\mathcal{V}$ (ceci est analogue avec ce que nous avons fait pour la charge $dq = \rho d\mathcal{V}$).

5.2.1 Champ créé par un diélectrique

Pour simplicité, prenons le cas (assez courant) où la charge totale diélectrique soit négligeable. Il y a plusieurs façons (équivalentes) de procéder à calculer le champ produit par un diélectrique, mais ils prennent des points de vu physique différents sur le problème. La méthode que l'on choisira sera déterminée par la nature de la question posée.

- (1) **Méthode directe** : Imaginons pour l'instant, qu'on connaît le **vecteur polarisation**, $\vec{P}$. On peut calculer la contribution au potentiel du diélectrique en mettant $\vec{d}\vec{p} = \vec{P}d\mathcal{V}$ dans l'équation (5.2) et intégrant (comme nous avons calculé le champ électrique à partir de la loi de Coulomb). Le potentiel créé par le diélectrique est donc :

$$V(M) = \frac{1}{4\pi\epsilon_0} \iiint \frac{\vec{P}(\mathcal{Q}) \cdot \hat{u}}{r^2} d\mathcal{V}, \quad (5.13)$$

où $r = \mathcal{Q}M$ est la distance entre le point d'intégration, $\mathcal{Q}$, et M ainsi que $d\mathcal{V}$ est maintenant un volume d'intégration autour du point $\mathcal{Q}$ ($\hat{u} \equiv \frac{\vec{\mathcal{Q}M}}{r}$). Bien-entendu, avec $V(M)$ en main on obtient le champ électrique produit par le diélectrique avec $\vec{E} = -\vec{\text{grad}}V$.

Remarque : Dans ce chapitre, nous remplaçons le point d'intégration, P , utilisé dans les chapitres précédents par $\mathcal{Q}$, afin d'éviter confusion avec le vecteur polarisation, $\vec{P}$

(2) Méthode des charges de polarisation

On peut transformer l'éq. (5.13) en effectuant une intégration par parties tri-dimensionnelle en remarquant que $\vec{P}(\mathcal{Q}) \cdot \hat{u}/r^2 = \vec{P} \cdot \vec{\text{grad}}_{\mathcal{Q}} \frac{1}{r}$ (où $\vec{\text{grad}}_{\mathcal{Q}}$ est le gradient par rapport à la position $\mathcal{Q}$). En faisant appel à l'identité de l'éq. (5.33) on trouve,

$$\vec{P} \cdot \vec{\text{grad}}_{\mathcal{Q}} \frac{1}{r} = \text{div}_{\mathcal{Q}} \left(\frac{1}{r} \vec{P} \right) - \frac{1}{r} \text{div}_{\mathcal{Q}} \vec{P}. \quad (5.14)$$

Insérant cette relation dans l'intégrale de l'éq. (5.13) et en se servant du théorème de Green-Ostrogradsky (cf. 11.15), on voit que l'expression de l'éq. (5.13) est équivalente à une expression en termes d'intégrales de la surface, $\mathcal{S}$, et du volume $\mathcal{V}$ du diélectrique :

$$V(M) = \frac{1}{4\pi\epsilon_0} \iint_{\mathcal{S}} \frac{\vec{P}(\mathcal{Q}) \cdot \hat{n} d\mathcal{S}}{r} - \frac{1}{4\pi\epsilon_0} \iiint_{\mathcal{V}} \frac{\text{div}_{\mathcal{Q}} P(\mathcal{Q})}{r} d\mathcal{V}. \quad (5.15)$$

Le sens physique de cette formule devient plus claire si on définit une **densité de charge de polarisation surfacique**, σ_{pol} , ainsi qu'une **densité volumique de charges de polarisation**, ρ_{pol} , telles que :

$$\sigma_{\text{pol}}(\mathcal{Q}) \equiv \vec{P}(\mathcal{Q}) \cdot \hat{n} \quad \rho_{\text{pol}}(\mathcal{Q}) \equiv -\text{div}_{\mathcal{Q}} \vec{P}(\mathcal{Q}) \quad (5.16)$$

Avec ces définitions, la formule de l'éq. (5.15) du potentiel s'écrit en termes de «charges de polarisation» :

$$V(M) = \frac{1}{4\pi\epsilon_0} \iint_S \frac{\sigma_{\text{pol}}}{r} dS + \frac{1}{4\pi\epsilon_0} \iiint_V \frac{\rho_{\text{pol}}}{r} dV. \quad (5.17)$$

Avec la formule de l'éq. (5.17), on peut calculer le potentiel $V(M)$ et déterminer le champ électromagnétique avec $\vec{E} = -\text{grad}V$, comme avec la méthode directe, mais on peut également amener le gradient sous l'intégrale afin d'obtenir $\vec{E}$ une formule directe pour $\vec{E}$ analogue à la formule de Coulomb :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \iint_S \frac{\sigma_{\text{pol}}}{r^2} \hat{u} dS + \frac{1}{4\pi\epsilon_0} \iiint_V \frac{\rho_{\text{pol}}}{r^2} \hat{u} dV, \quad (5.18)$$

avec $r = QM$, et $\hat{u} = \vec{QM}/QM$ comme d'habitude.

Remarque : L'expression (5.18) est un exemple du **théorème de compensation** qui dit que l'effet électromagnétique de tout objet matériel peut être reproduit par des charges et courants agissant dans le vide. Enfin compte, nous avons déjà rencontré un autre exemple de ce théorème avec les conducteurs parfait quand leurs effets étaient traités comme des charges surfaciques. Vous pouvez également regarder le chapitre 10 du *Cours de Physique de Feynman - Électromagnétisme 1* pour une présentation didactique du champ $\vec{P}$.

5.2.2 Le vecteur polarisation

Afin qu'on puisse appliquer les formules dérivées dans 5.2.1 il nous faut connaître le vecteur polarisation $\vec{P}$ partout. Pour un vaste nombre de milieux dans des conditions ordinaires, $\vec{P}$ est simplement proportionnel au champ électrique $\vec{E}$ présent à cet endroit (Appelés **milieux linéaires**). La constante de proportionnalité entre $\vec{P}$ et $\epsilon_0 \vec{E}$, est dénoté χ_e et on l'appelle la **susceptibilité électrique** du matériel,

$$\vec{P} \equiv \chi_e \epsilon_0 \vec{E}, \quad (5.19)$$

où ϵ_0 est mis dans la définition afin que χ_e puisse être un nombre sans dimension. Bien qu'en principe, on doive pouvoir calculer χ_e théoriquement, ceci nécessiterait un calcul en mécanique quantique de grande envergure et en pratique on prend des valeurs de χ_e mesuré dans le laboratoire.

Si on met l'expression de l'éq. (5.19) dans les équations pour calculer le champ V ou $\vec{E}$ de la section précédente (les éq. (5.15) ou (5.18)) on s'en aperçoit qu'afin de calculer le champ $\vec{E}$ à un point donné M , il faut déjà connaître le champ $\vec{E}$ dans tout le matériel. Une telle situation, où la connaissance d'une quantité (ici le champ $\vec{E}$) est requise afin de la calculer ce même quantité est souvent rencontré en physique et dénommé calcul auto-consistant. Nous verrons plus loin comment la linéarité de $\vec{P}$ par rapport à $\vec{E}$ nous permettra de contourner ce problème d'auto-consistance. Pour l'instant, voyons comment le problème d'auto-consistance se manifeste dans le problème simple d'un condensateur avec milieu diélectrique entre ces armatures.

5.2.3 Condensateur avec diélectrique

On considère un condensateur plan avec un diélectrique de susceptibilité χ_e remplissant le volume $V = Sd$ entre les deux armatures (de surface S) d'un condensateur (voir la figure 5.3 ci-dessous). La charge sur l'armature A_1 est Q ($\sigma = Q/S$). On dénote par $\vec{E}_0$ le champ qui serait présent entre les armatures si le diélectrique n'était pas présent, $\vec{E}_0 = Q/(\epsilon_0 S) \hat{z}$. Afin de résoudre le problème auto-consistant, on fait l'hypothèse que le vecteur $\vec{P}$ entre les armatures soit constant et dirigé selon la direction $\hat{z}$ et voir si on trouve une solution.

Avec l'hypothèse de $\vec{P}$ constant entre les armatures on voit qu'il n'y a pas de charges volumiques dans le matériel $\rho_{\text{pol}} = -\text{div}_Q \vec{P}(Q) = 0$, il suffit de calculer le champ créé par les charges de

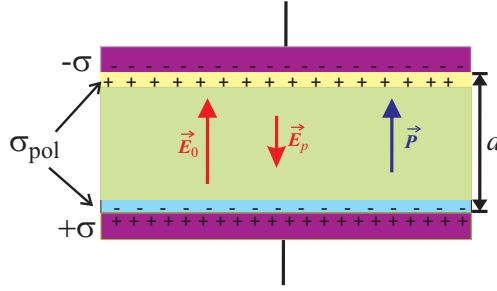


FIGURE 5.3 – Condensateur avec un diélectrique entre les armatures.

polarisation surfaciques, σ_{pol} . On se rappelle que $\sigma_{\text{pol}} = \vec{P} \cdot \hat{n}$ où $\hat{n}$ est le vecteur normale à la surface du diélectrique et dirigé vers l'extérieur. Il y a donc une charge surfacique de $\sigma_{\text{pol}} = P$ sur la face supérieure du diélectrique ainsi qu'une charge $\sigma_{\text{pol}} = -P$ sur la face inférieure (voir la figure 5.3). Le diélectrique doit donc avoir le même comportement qu'un condensateur dont le champ électrique à l'intérieur est en opposition au champ électrique imposée $\vec{E}_0$. Donc le champ de polarisation est $\vec{E}_{\text{pol}} = -\sigma_{\text{pol}}/\epsilon_0 \hat{z} = -\vec{P}/\epsilon_0 = -\chi_e \vec{E}$ à l'intérieur du diélectrique et nul ailleurs. Le champ électrique, $\vec{E}$, réellement présent entre les armatures est alors la superposition de $\vec{E}_0$ et de $\vec{E}_{\text{pol}}$

$$\begin{aligned}\vec{E} &= \vec{E}_0 + \vec{E}_{\text{pol}} = \vec{E}_0 - \chi_e \vec{E} \\ \vec{E} &= \frac{1}{(1 + \chi_e)} \vec{E}_0.\end{aligned}$$

Donc, on voit que l'effet de placer un diélectrique entre les armatures est de diminuer le champ électrique par un facteur $1 + \chi_e$ qu'on appelle la **constante diélectrique relatif** du matériel

$$\boxed{\epsilon_r = (1 + \chi_e)} . \quad (5.20)$$

Remarque : On constate que ϵ_r est un nombre sans dimension. On l'appelle constante diélectrique relatif afin d'éviter confusion avec le fait que certains préfèrent parler de **constant diélectrique (ou permittivité) du matériel**, la quantité $\epsilon_d = \epsilon_r \epsilon_0$ qui a évidemment les mêmes dimensions que la permittivité du vide ϵ_0 .

Le fait d'avoir diminué $\vec{E}$ entre les armatures diminue également leur différence de potentiel.

$$U = V_1 - V_2 = \int_{A_1}^{A_2} \vec{E} \cdot d\vec{\ell} = \frac{d}{\epsilon_r} E_0 = \frac{d}{\epsilon_r} \frac{Q}{\epsilon_0 S} \equiv \frac{Q}{C} .$$

Donc on peut en conclure que la capacité d'un condensateur avec un diélectrique entre les armatures est augmentée par le facteur ϵ_r

$$\boxed{C = \frac{\epsilon_0 \epsilon_r S}{d}} . \quad (5.21)$$

5.2.4 Déplacement électrique

C'est souvent assez pénible de formuler un problème d'électrostatique en termes charges de polarisation. Après tout, ces « charges » n'existent que grâce à un champ électrique appliqué au système, et disparaissent aussitôt qu'on enlève le champ appliqué. Il est difficile de mesurer les charges de polarisation et on ne peut pas les manipuler directement lors d'une expérience. On peut se demander donc, pourquoi en parler de tout. Y'a-t-il un moyen de formuler l'électrostatique de façon de prendre ces charges de polarisation en compte automatiquement afin d'oublier ou presque leur existence ? La réponse est oui, et c'est la raison pour laquelle les physiciens choisissent souvent d'utiliser un champ « auxiliaire », $\vec{D}$

qui intègre le vecteur polarisation dans sa définition dès le départ, le **déplacement diélectrique** $\vec{D}$, définit tel que :

$$\vec{D}(\vec{r}) \equiv \epsilon_0 \vec{E}(\vec{r}) + \vec{P}(\vec{r}) . \quad (5.22)$$

L'équation différentielle de $\vec{D}$ est :

$$\text{div } \vec{D} = \epsilon_0 \text{div } \vec{E} + \text{div } \vec{P} = \rho - \rho_{\text{pol}} \equiv \rho_{\text{libre}} , \quad (5.23)$$

où ρ_{libre} correspond aux charges réellement manipulées dans une expérience. Avec ce champ, on ne parle pas des charges de polarisations mais seulement des charges libres comme celles qu'on place sur les armatures d'un condensateur.

Si les symétries du problème sont suffisant, on peut résoudre $\vec{D}$ en faisant appel à la forme intégrale de l'éq. (5.23) :

$$\oint_S \vec{D} \cdot d\vec{S} = Q_{\text{libre,int}} , \quad (5.24)$$

ce qui est un analogue du théorème de Gauss. Cette expression nous dicte également les unités du champ $\vec{D}$ comme C.m⁻².

Il ne faut pas oublier, que c'est toujours le champ $\vec{E}$ qui est relié aux quantités mesurables (différences de potentiel) et que le champ « auxiliaire », $\vec{D}$ sert surtout comme un moyen de trouver $\vec{E}$. Le lien entre les deux quantités est directe pour des matériaux **linéaires** (dont le vecteur polarisation satisfait l'éq. (5.19)). Mettant la relation (5.19) dans (5.22), on obtient une relation linéaire entre $\vec{D}$ et $\vec{E}$ (appelée **relation constitutive**) :

$$\vec{D} \equiv \epsilon_0 \vec{E} + \vec{P} = \epsilon_0 (1 + \chi_e) \vec{E} \equiv \epsilon_0 \epsilon_r \vec{E} \equiv \epsilon \vec{E} , \quad (5.25)$$

où ϵ est la permittivité du diélectrique et $\epsilon_r = \epsilon/\epsilon_0 = (1 + \chi_e)$ est la constante diélectrique (**relative**) du diélectrique.

En résumé, en présence de diélectriques, les lois de l'électrostatique peuvent être formulées à partir de deux formules :

$$\left\{ \begin{array}{l} \text{div } \vec{D} = \rho_{\text{libre}} \\ \text{rot } \vec{E} = \vec{0} \end{array} \right. , \quad (5.26)$$

associées avec la relation constitutive :

$$\vec{D} = \epsilon_0 \epsilon_r \vec{E} . \quad (5.27)$$

Utilisant les mêmes techniques que dans le complément du chapitre 4, les conditions limites associés avec les formules de l'éq. (5.26) sont :

$$\hat{n}_{12} \wedge (\vec{E}_2 - \vec{E}_1) = \vec{0} \quad \hat{n}_{12} \cdot (\vec{D}_2 - \vec{D}_1) = \sigma_{\text{libre}} . \quad (5.28)$$

On verra plus tard que ces conditions limites sont toujours valables même dans le régime de champs variables dans le temps. Dans des applications plus avancées, on exploite souvent ces conditions limites dans la recherche de solutions.

5.2.5 Condensateur et déplacement électrique

Afin d'apprécier l'utilité du champ auxiliaire $\vec{D}$, reprenons le condensateur plane traité en 5.2.3. La charge Q sur les l'armature A_1 est simplement la charge « libre » et la densité surfacique libre est simplement $\sigma = Q/S$. Grâce à la symétrie du problème, $\vec{D}$ est dans la direction de séparation des armatures et on obtient $\vec{D} = \sigma \hat{z}$ en utilisant l'éq. (5.24) et les mêmes techniques de surfaces de Gauss que celles employées pour le condensateur du vide (voir chapitre 4). On en déduit que le champ entre les armatures est :

$$\vec{E} = \frac{\vec{D}}{\epsilon_0 \epsilon_r} = \frac{\sigma}{\epsilon_0 \epsilon_r} \hat{z} , \quad (5.29)$$

ce qui amène à $U = \frac{Qd}{S\epsilon_0\epsilon_r}$ et $C = \frac{S\epsilon_0\epsilon_r}{d}$ exactement comme nous avons trouvé dans la section 5.2.3 mais avec nettement moins d'effort.

5.3 Energie potentielle électrostatique

5.3.1 Energie électrostatique d'une charge ponctuelle

Comment mesure-t-on l'énergie potentielle gravitationnelle d'un corps de masse m ? On le déplace d'une position initiale jusqu'à une position finale (on exerce donc une force) puis on le lâche sans vitesse initiale. S'il acquiert une vitesse, c'est qu'il développe de l'énergie cinétique. Or, en vertu du principe de conservation de l'énergie, cette énergie ne peut provenir que d'un autre réservoir énergétique, appelé énergie potentielle. Comment s'est constituée cette énergie potentielle gravitationnelle ? Grâce au déplacement du corps par l'opérateur. Ainsi, le travail effectué par celui-ci est une mesure directe de l'énergie potentielle. On va suivre le même raisonnement pour l'énergie électrostatique.

Définition 2 *l'énergie potentielle électrostatique d'une particule chargée placée dans un champ électrostatique, $\mathcal{E}_e$, est égale au travail qu'il faut fournir pour amener de façon quasi-statique cette particule de l'infini à sa position actuelle.*

Prenons une particule de charge q placée dans un champ $\vec{E}$. Pour la déplacer de l'infini vers un point M , un opérateur doit fournir une force qui s'oppose à la force de Coulomb. Si ce déplacement est fait suffisamment lentement, la particule n'acquiert aucune énergie cinétique. Cela n'est possible que si, à tout instant, $\vec{F}_{\text{ext}} = -\vec{F} = -q\vec{E}$. Le travail fourni par l'opérateur sera donc

$$\mathcal{E}_e(M) \equiv \int_{\infty}^M dW = \int_{\infty}^M \vec{F}_{\text{ext}} \cdot d\vec{r} = - \int_{\infty}^M q\vec{E} \cdot d\vec{r} = q[V(M) - V(\infty)] .$$

Puisqu'on peut toujours définir le potentiel nul à l'infini, on obtient l'expression suivante pour l'énergie électrostatique d'une charge ponctuelle située en M

$$\mathcal{E}_e(M) = qV(M) .$$

On voit donc que le potentiel électrostatique est une mesure (à un facteur q près) de l'énergie électrostatique : c'est dû au fait que V est lié à la circulation du champ. Autre remarque importante : l'énergie est indépendante du chemin suivi.

5.3.2 Energie électrostatique d'un ensemble de charges ponctuelles

Dans la section précédente, nous avons considéré une charge q placée dans un champ $\vec{E}$ extérieur et nous avons ainsi négligé le champ créé par la charge elle-même. Mais lorsqu'on a affaire à un ensemble de N charges ponctuelles q_i , chacune d'entre elles va créer sur les autres un champ électrostatique

et ainsi mettre en jeu une énergie d'interaction électrostatique. Quelle sera alors l'énergie potentielle électrostatique de cet ensemble de charges ?

Soit la charge ponctuelle q_1 placée en P_1 . On amène alors une charge q_2 de l'infini jusqu'en P_2 , c'est-à-dire que l'on fournit un travail $W_2 = q_2 V_1(P_2) = q_1 V_2(P_1) = W_1$ identique à celui qu'il aurait fallu fournir pour amener q_1 de l'infini en P_1 en présence de q_2 déjà située en P_2 . Cela signifie que ce système constitué de 2 charges possède une énergie électrostatique :

$$\mathcal{E}_e = \frac{q_1 q_2}{4\pi\epsilon_0 r_{12}} = W_1 = W_2 = \frac{1}{2} (W_1 + W_2) ,$$

où $r_{12} = P_1 P_2$.

Remarque : Dans cette approche, nous avons considéré q_2 immobile alors que l'on rapprochait q_1 . En pratique évidemment, c'est la distance entre les deux charges qui diminue du fait de l'action de l'opérateur extérieur à la fois sur q_1 et q_2 (avec $\vec{F}_{\text{ext} \rightarrow 1} = -\vec{F}_{\text{ext} \rightarrow 2}$ puisque $\vec{F}_{1 \rightarrow 2} = -\vec{F}_{2 \rightarrow 1}$). On aurait aussi bien pu calculer le travail total fourni par l'opérateur en évaluant le déplacement de q_1 et de q_2 de l'infini à la distance intermédiaire (« $M/2$ »). Une autre façon de comprendre cela, c'est de réaliser que nous avons évalué le travail fourni par l'opérateur dans le référentiel lié à q_2 (immobile). Celui-ci est identique au travail évalué dans un référentiel fixe (où q_1 et q_2 se déplacent) car le déplacement des charges s'effectue de manière quasi-statique (aucune énergie n'a été communiquée au centre de masse).

Si maintenant on amène une 3^{ème} charge q_3 de l'infini jusqu'en P_3 (q_1 et q_2 fixes), il faut fournir un travail supplémentaire,

$$\begin{aligned} W_3 &= q_3 V_{1+2}(P_3) = q_3 (V_1(P_3) + V_2(P_3)) \\ &= \frac{1}{4\pi\epsilon_0} \left(\frac{q_1 q_3}{r_{13}} + \frac{q_2 q_3}{r_{23}} \right) , \end{aligned}$$

correspondant à une énergie électrostatique de ce système de 3 charges :

$$\mathcal{E}_e = \frac{1}{4\pi\epsilon_0} \left(\frac{q_1 q_2}{r_{12}} + \frac{q_1 q_3}{r_{13}} + \frac{q_2 q_3}{r_{23}} \right) .$$

Ainsi, on voit qu'à chaque couple $q_i q_j$ est associée une énergie potentielle d'interaction. Pour un système de N charges, on aura alors :

$$\boxed{\mathcal{E}_e = \frac{1}{4\pi\epsilon_0} \sum_{\text{couples}} \frac{q_i q_j}{r_{ij}}} . \quad (5.30)$$

Cette formule est assez pratique d'utilisation pour des systèmes composés de quelques charges ponctuelles.

Il existe néanmoins, une forme équivalente à l'éq. (5.30) qui s'avère parfois plus pratique en présence de symétries, et qui est surtout indispensable quand on veut généraliser l'énergie électrostatique à des distributions de charge. La dérivation de cette forme alternative est la suivante :

$$\mathcal{E}_e = \frac{1}{4\pi\epsilon_0} \sum_{\text{couples}} \frac{q_i q_j}{r_{ij}} = \frac{1}{4\pi\epsilon_0} \sum_{i=1}^N \sum_{j>i}^N \frac{q_i q_j}{r_{ij}} = \frac{1}{2} \sum_{i=1}^N \frac{q_i}{4\pi\epsilon_0} \sum_{j \neq i} \frac{q_j}{r_{ij}} = \frac{1}{2} \sum_{i=1}^N q_i V_i ,$$

où le facteur 1/2 apparaît parce que chaque couple est compté deux fois. On remarque qu'ici V_i est le potentiel créé en P_i par toutes **les autres** charges du système (le potentiel dû à la charge q_i étant exclu).

L'énergie électrostatique d'un ensemble de N charges ponctuelles peut donc s'écrire de façon alternative (mais équivalente à celle de (5.30)) comme suit :

$$\boxed{\mathcal{E}_e = \frac{1}{2} \sum_{i=1}^N q_i V_i(P_i) \quad \text{avec} \quad V_i(P_i) = \frac{1}{4\pi\epsilon_0} \sum_{j \neq i} \frac{q_j}{r_{ij}}} . \quad (5.31)$$

5.3.3 Energie stockée dans le champ électrique

Les équations (5.38), (5.36) et (5.37) nous disent qu'une distribution de charges peut emmagasiner de l'énergie électrostatique. Mais où est-elle stockée? Sous quelle forme? Une réponse possible à cette question est de voir l'énergie électrostatique comme étant stocké dans le champ électrique lui-même. On peut dériver une telle formulation en se rappelant qu'au niveau fondamentale, toute distribution de charge peut être décrite comme une densité volumique du champ $\rho(\vec{r})$:

$$\mathcal{E}_e = \frac{1}{2} \int_V \rho(P) V(P) dV_P = \frac{1}{2} \int_V \rho V dV, \quad (5.32)$$

où nous avons supprimé la dépendance spatiale sur P afin de mieux voir le contenu physique de cette équation.

L'équation (5.32) est sans ambiguïté dans le vide, mais en présence de diélectriques, c'est la charge libre qu'il faut lire dans cette équation. En invoquant l'équation $\text{div } \vec{D} = \rho_{\text{libre}}$ on peut réécrire cette équation comme :

$$\mathcal{E}_e = \frac{1}{2} \iiint_V V \text{div } \vec{D} dV = \frac{1}{2} \iiint_V \text{div} (V \vec{D}) dV - \frac{1}{2} \iiint_V \vec{D} \cdot \overrightarrow{\text{grad}} V dV,$$

où nous avons effectué une intégration par parties tridimensionnelle en utilisant l'identité mathématique (démontrée dans l'eq. (11.19) de l'annexe des mathématiques de chapitre 11) :

$$\text{div} (f \vec{A}) \equiv \vec{A} \cdot \overrightarrow{\text{grad}} f + f \text{div } \vec{A}. \quad (5.33)$$

En invoquant le théorème d'Ostrogradsky (eq. (11.15)), et $\vec{E} = -\overrightarrow{\text{grad}} V$ on obtient

$$\mathcal{E}_e = \frac{1}{2} \oint_S V \vec{D} \cdot \hat{n} dS + \frac{1}{2} \iiint_V \vec{D} \cdot \vec{E} dV,$$

où il faut se rappeler que la surface d'intégration S dans cette équation est rejetée à l'infini afin d'entourer tout l'espace. De ce fait, l'intégrale surfacique est nulle puisque à des grandes distances, V décroît comme $1/r$ et $\vec{D}$ décroît comme $1/r^2$, tandis que la surface S augmente seulement comme r^2 . Nous avons enfin une expression de l'énergie électrostatique d'une distribution de charges exprimée entièrement en termes du champ $\vec{D}$ et $\vec{E}$:

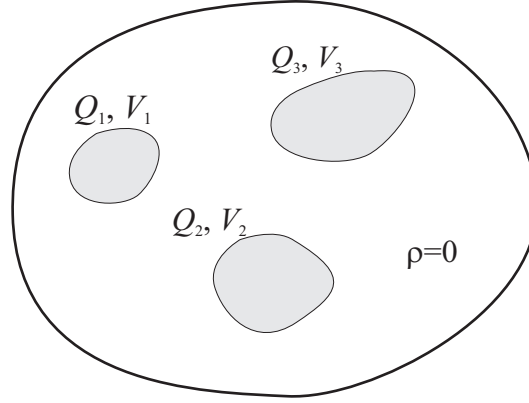
$$\mathcal{E}_e = \frac{1}{2} \iiint_V \vec{D} \cdot \vec{E} dV. \quad (5.34)$$

Pour des charges dans le vide, $\vec{D} = \epsilon_0 \vec{E}$ (puisque $\vec{P}$ y est nul), on peut interpréter la quantité $\frac{\epsilon_0}{2} \vec{E}^2$, comme une densité volumique d'énergie électrique. En présence de diélectriques par contre, une partie de l'énergie est contenue dans les degrés de liberté du diélectrique, et même si on peut toujours dans une certaine mesure interpréter $\frac{1}{2} \vec{D} \cdot \vec{E} = \frac{\epsilon_0 \epsilon_r}{2} \vec{E}^2$ comme une densité d'énergie locale, il devient difficile à distinguer la partie de l'énergie stockée dans le champ électrique de celle qui est stockée dans le matériel sans savoir les détails sur les propriétés microscopiques du diélectrique.

5.3.4 Energie électrostatique de conducteurs en équilibre

Soit un conducteur isolé, de charge Q distribuée sur sa surface S . L'énergie potentielle électrostatique de ce conducteur est alors,

$$\mathcal{E}_e = \frac{1}{2} \int dq V(P) = \frac{V}{2} \int dq = \frac{QV}{2},$$



puisque'il est équipotentiel, c'est-à-dire :

$$\mathcal{E}_e = \frac{1}{2} QV = \frac{1}{2} CV^2 = \frac{1}{2} \frac{Q^2}{C} .$$

Ceci est l'énergie nécessaire pour amener un conducteur de capacité C au potentiel V . Puisque cette énergie est toujours positive cela signifie que, quel que soit V (et donc sa charge Q), cela coûte toujours de l'énergie.

Soit un ensemble de N conducteurs chargés placés dans un volume $\mathcal{V}$. A l'équilibre, ils ont une charge Q_i et un potentiel V_i . En dehors du volume occupé par chaque conducteur, il n'y a pas de charge donc $dq = 0$. L'énergie électrostatique de cette distribution de charges est alors simplement,

$$\mathcal{E}_e = \frac{1}{2} \int_{\mathcal{V}} dq V(P) = \frac{1}{2} \sum_{i=1}^N \int_{S_i} V_i \sigma_i(P) dS_i = \frac{1}{2} \sum_{i=1}^N V_i \int_{S_i} \sigma_i(P) dS_i ,$$

c'est-à-dire :

$$\mathcal{E}_e = \frac{1}{2} \sum_{i=1}^N Q_i V_i . \quad (5.35)$$

5.3.5 Energie électrostatique d'une distribution continue de charges

Pour une distribution continue de charges, la généralisation de la formule précédente est immédiate. Soit $dq = \rho d\mathcal{V}$ la charge située autour d'un point P quelconque de la distribution. L'énergie électrostatique de cette distribution s'écrit

$$\boxed{\mathcal{E}_e = \frac{1}{2} \iiint \rho(P) V(P) d\mathcal{V}_P = \frac{1}{2} \iiint \rho V d\mathcal{V}} , \quad (5.36)$$

où nous avons de nouveau supprimé la dépendance explicite en P dans la deuxième égalité.

Remarque : Nous n'avons pas spécifié une limite à l'intégration puisque $\rho = 0$ en dehors de la distribution des charges. Donc en principe, l'intégrale est effectuée sur tout l'espace, mais en pratique on n'effectue l'intégrale que sur les régions contenant des distributions de charges non-nulles.

Il est important de remarquer qu'en passant à une distribution continue dans l'éq. (5.36), le potentiel à un point P ,

$$V(P) = \frac{1}{4\pi\epsilon_0} \iiint \frac{\rho(P') d\mathcal{V}_{P'}}{PP'} = \frac{1}{4\pi\epsilon_0} \iiint \frac{\rho d\mathcal{V}}{PP'} ,$$

est maintenant le potentiel créé par **toute la distribution de charge**. En effet, pour un volume infinitésimal $d\mathcal{V}$ de taille caractéristique a , c'est-à-dire, $d\mathcal{V} \propto a^3$, le potentiel au centre de ce volume

créé par charges à l'intérieur de dV est forcément proportionnel à a^{-1} . De ce fait, la région autour du point P fera une contribution $\rho dV/PP' \propto \rho(P)a^2$ à l'intégrale et celle-ci sera nulle dans le passage à la limite infinitésimale de $a \rightarrow 0$.

De la même manière, on obtient un résultat analogue pour une distribution surfacique :

$$\boxed{\mathcal{E}_e = \frac{1}{2} \iint_S \sigma V dS}, \quad (5.37)$$

où S est la surface contenant la charge. Par contre, on ne peut pas généraliser ces arguments aux fils infiniment minces en interaction. Pour de tels cas, on est obligé de procéder comme nous avons fait pour des charges ponctuelles et exclure le potentiel créé par le fil lui-même ou d'abandonner le concept d'un fil infiniment mince et revenir à une distribution de charges surfacique ou volumique.

Un effet subtil mais important est que le $\mathcal{E}_e$ obtenu par les éqs. (5.36) ou (5.37) seront non-négatives, c.-à-d. $\mathcal{E}_e \geq 0$. Ceci vient du passage de charges ponctuelles à la distribution continue de charge qui a eu un effet subtil mais important. Pour une distribution volumique (ou surfacique) de charge, il n'est pas nécessaire d'exclure explicitement la charge située en P puisque $dq(P)$ tend vers zéro avec l'élément infinitésimal (contribution nulle à l'intégrale, absence de divergence). Du coup, les expressions intégrales de $\mathcal{E}_e$ des équations, (5.36) et (5.37) tiennent compte du travail nécessaire d'établir une distribution de charge alors que les expressions de (5.30) et (5.31) ne tiennent pas compte de l'énergie nécessaire d'établir les charges ponctuelles que nous avons pris comme existant ex nihilo (car l'énergie d'une distribution de charge vraiment « ponctuelle » tendrait vers infini quand sa taille tend vers zéro).

5.3.6 Quelques exemples

— Exemple 1 : Le condensateur

Soit un condensateur constitué de deux armatures. L'énergie électrostatique de ce système de deux conducteurs est

$$\mathcal{E}_e = \frac{1}{2} (Q_1 V_1 + Q_2 V_2) = \frac{1}{2} Q_1 (V_1 - V_2) = \frac{1}{2} Q U,$$

c'est-à-dire

$$\mathcal{E}_e = \frac{1}{2} Q U = \frac{1}{2} C U^2 = \frac{1}{2} \frac{Q^2}{C}. \quad (5.38)$$

Energie emmagasinée par un condensateur plan - On se rappelle du chapitre 4 qu'un condensateur plan est constitué de deux armatures de surfaces, S , en influence quasi-totale et séparées d'une distance $d \ll S$ (remplit éventuellement par un diélectrique). Nous avons vu également que la capacité de ce condensateur est $C = S\epsilon_r\epsilon_0/d$, et donc l'énergie stockée dans le condensateur est donnée par l'éq. (5.38). On se rappelle que $\vec{E}$ et $\vec{D}$ sont parallèles entre les armatures et nul ailleurs. Entre les armatures on a $|\vec{E}| = \sigma/(\epsilon_r\epsilon_0)$ et $|\vec{D}| = \sigma$ et on peut ainsi vérifier que le résultat de l'éq. (5.38) est consistant avec l'expression de l'éq. (5.34) pour de l'énergie stockée dans le champ électrique :

$$\begin{aligned} \mathcal{E}_e &= \frac{1}{2} \iiint_V \vec{D} \cdot \vec{E} dV = \frac{Sd}{2} \sigma \left(\frac{\sigma}{\epsilon_r\epsilon_0} \right) = \frac{1}{2\epsilon_r\epsilon_0} Sd \left(\frac{Q}{S} \right)^2 \\ &= \frac{1}{2} \frac{Q^2}{\frac{S\epsilon_r\epsilon_0}{d}} = \frac{1}{2} \frac{Q^2}{C}. \end{aligned}$$

— Exemple 2 : Le dipôle

Soit un dipôle électrostatique placé dans un champ électrostatique $\vec{E}_{\text{ext}}$. On s'intéresse au potentiel d'interaction électrostatique entre ce dipôle et le champ et non pas à celle qui existe

entre la charge $+q$ et $-q$ du dipôle lui-même. On considère donc le dipôle comme un système de deux charges, $-q$ placée en un point B et $+q$ en A , n'interagissant pas entre elles. L'énergie électrostatique de ce système de charges est simplement

$$\mathcal{E}_{e,\text{ext}} = qV_{\text{ext}}(A) - qV_{\text{ext}}(B) = q \int_B^A dV = -q \int_B^A \vec{E}_{\text{ext}} \cdot d\vec{r} \simeq -q \vec{E}_{\text{ext}} \cdot \vec{BA},$$

ce qui donne

$$\mathcal{E}_{e,\text{ext}} = -\vec{p} \cdot \vec{E}_{\text{ext}}, \quad (5.39)$$

où $\vec{p} = q\vec{BA}$ est le moment dipolaire électrique.

Remarque : L'énergie électrostatique entre la charge $+q$ et $-q$ du dipôle lui-même est $\mathcal{E}_{e,\text{int}} = -\frac{q^2}{4\pi\epsilon_0 BA} = \vec{p} \cdot \vec{E}_{\text{int}}(A)$ où $\vec{E}_{\text{int}}(A)$ est le champ créé par la charge $-q$ en A . Si le champ extérieur $\vec{E}_{\text{ext}}$ est bien supérieur à $\vec{E}_{\text{int}}(A)$ cela signifie que le dipôle est profondément modifié (voire brisé) par le champ : l'énergie d'interaction est supérieure à l'énergie interne de liaison. Cependant, la distance AB étant en général très petite, cela ne se produit pas et le dipôle se comporte comme un système lié, sans modification de son énergie interne (ceci n'est pas tout à fait exact : un champ extérieur peut faire osciller les deux charges autour de leur position d'équilibre, induisant ainsi une variation de leur énergie de liaison).

— Exemple 3 : Un conducteur chargé placé dans un champ extérieur

Soit un conducteur portant une charge Q et mis au potentiel V en l'absence de champ extérieur. Il possède donc une énergie électrostatique interne $\mathcal{E}_{e,\text{int}} = \frac{QV}{2}$ correspondant à l'énergie qu'il a fallu fournir pour déposer les Q charges au potentiel V sur le conducteur. Si maintenant il existe un champ extérieur $\vec{E}_{\text{ext}}$, alors le conducteur prend un nouveau potentiel V' et son énergie peut s'écrire $\mathcal{E}_{e,\text{int}} = \frac{QV'}{2}$. Comment calculer V' ?

La méthode directe consiste à prendre en compte la polarisation du conducteur sous l'effet du champ extérieur et calculer ainsi la nouvelle distribution surfacique σ (avec $Q = \iint_S \sigma dS$).

Une autre méthode consiste à considérer la conservation de l'énergie : en plaçant le conducteur dans un champ extérieur, on lui fournit une énergie potentielle d'interaction électrostatique qui s'ajoute à son énergie électrostatique « interne ». Supposons (pour simplifier) que le champ extérieur $\vec{E}_{\text{ext}}$ est constant à l'échelle du conducteur. Alors ce dernier se comporte comme une charge ponctuelle placée dans un champ et possède donc une énergie potentielle d'interaction électrostatique $\mathcal{E}_{e,\text{ext}} = QV_{\text{ext}}$. L'énergie électrostatique totale sera alors

$$\mathcal{E}_e = QV_{\text{ext}} + Q\frac{V}{2} = Q\frac{V'}{2},$$

c'est-à-dire $V' = V + 2V_{\text{ext}}$

5.4 Actions électrostatiques sur des conducteurs

5.4.1 Notions de mécanique du solide

- a) Calcul direct des actions (force et moment d'une force)

Un conducteur étant un solide, il faut faire appel à la mécanique du solide. Tout d'abord, on choisit un point de référence O , des axes et un système de coordonnées respectant le plus possible

la symétrie du solide. La force et le moment de cette force par rapport au point O sont alors

$$\begin{aligned}\vec{F} &= \iiint_{\text{solide}} d^3\vec{F} \\ \vec{\Gamma}_O &= \iiint_{\text{solide}} d^3\vec{\Gamma}_O = \iiint_{\text{solide}} \vec{OP} \wedge d^3\vec{F},\end{aligned}$$

où $d^3\vec{F}$ est la force s'exerçant sur un élément infinitésimal centré autour d'un point P quelconque du solide et où l'intégrale porte sur tous les points du solide. Le formalisme de la mécanique du solide considère ensuite que la force totale ou résultante $\vec{F}$ s'applique au barycentre G du solide.

- b) Liens entre travail d'une action (force ou moment) et l'action elle-même.

Lors d'une *translation pure* du solide, considéré comme indéformable, tout point P du solide subit une translation d'une quantité fixe : $d\vec{r} = \vec{r}' - \vec{r} = \vec{\varepsilon}$. La force totale responsable de ce déplacement doit fournir un travail

$$\begin{aligned}dW &= \iiint_{\text{solide}} d^3\vec{F} \cdot d\vec{r} = \left(\iiint_{\text{solide}} d^3\vec{F} \right) \cdot \vec{\varepsilon} = \vec{F} \cdot \vec{\varepsilon} \\ &= \sum_{i=1}^3 F_i dx_i,\end{aligned}$$

où $\vec{F}$ est la résultante de la force s'exerçant sur le solide et les x_i les coordonnées du centre de masse du solide.

Dans le cas de *rotations pures*, on ne s'intéresse qu'au moment des forces responsables de ces rotations. Celles-ci sont décrites par trois angles infinitésimaux $d\alpha$ autour de trois axes Δ_i , passant par le centre d'inertie G du solide et engendrés par les vecteurs unitaires $\hat{e}_i$. L'expression générale du moment d'une force (ou couple) par rapport à G est alors

$$\vec{\Gamma}_O = \sum_{i=1}^3 \Gamma_i \hat{e}_i$$

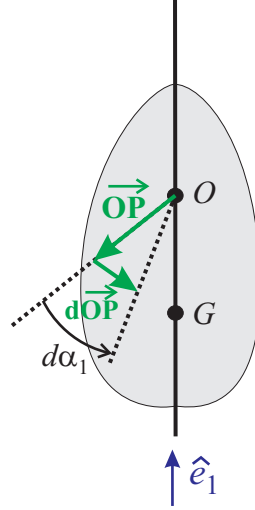
Lors de rotations du solide, le vecteur repérant la position d'un de ses points P quelconque varie suivant la règle

$$d(\vec{OP}) = \sum_{i=1}^3 d\alpha_i \vec{e}_i \wedge \vec{OP}.$$

Le travail fourni par le moment de la force est

$$\begin{aligned}dW &= \iiint_{\text{solide}} d^3\vec{F} \cdot d\vec{OP} = \iiint_{\text{solide}} d^3\vec{F} \cdot \left(\sum_{i=1}^3 d\alpha_i \hat{e}_i \wedge \vec{OP} \right) \\ &= \sum_{i=1}^3 d\alpha_i \hat{e}_i \cdot \left(\iiint_{\text{solide}} \vec{OP} \wedge d^3\vec{F} \right) = \sum_{i=1}^3 d\alpha_i \hat{e}_i \cdot \vec{\Gamma} \\ &= \sum_{i=1}^3 d\alpha_i \Gamma_i.\end{aligned}$$

Dans le cas général d'une translation accompagnée de rotations, chaque effet produit une contribution au travail fourni lors de l'interaction.



c) *Calcul des actions à partir de l'énergie potentielle (méthode des travaux virtuels)*

Si l'on a cherché le lien entre travail de l'action et les composantes de celle-ci, c'est qu'il est possible de calculer ces dernières en appliquant le principe de conservation de l'énergie. En effet une force produit un mouvement de translation de l'ensemble du solide tandis que le moment de la force produit un mouvement de rotation. Ces deux actions correspondent à un travail, donc à une modification de l'énergie d'interaction.

L'énergie totale $\mathcal{E}_{\text{tot}}$ d'un solide en électrostatique s'écrit $\mathcal{E}_{\text{tot}} = \mathcal{E}_c + \mathcal{E}_e$ où $\mathcal{E}_c$ est son énergie cinétique et $\mathcal{E}_e$ son énergie électrostatique (aussi appelée : potentielle d'interaction). Si le solide est *isolé*, son énergie totale reste constante, c'est-à-dire $d\mathcal{E}_{\text{tot}} = 0$, et l'on obtient ainsi le théorème de l'énergie cinétique

$$d\mathcal{E}_c = dW = -d\mathcal{E}_e, \quad (5.40)$$

où nous avons vu que au dessus que $dW = \sum_{i=1}^3 F_i dx_i$. Si l'on a par ailleurs l'expression de l'énergie électrostatique $\mathcal{E}_e$, alors on peut directement exprimer la force ou son moment (exprimés dans dW) en fonction de $d\mathcal{E}_e$.

Si, lors de l'évolution du solide, celui-ci n'est pas isolé et reçoit ou perd de l'énergie, on a $d\mathcal{E}_{\text{tot}} \neq 0$, c'est-à-dire :

$$d\mathcal{E}_c = dW = d\mathcal{E}_{\text{tot}} - d\mathcal{E}_e. \quad (5.41)$$

On voit donc que dans ce cas, le lien entre la force (ou son moment) et l'énergie potentielle n'est plus directe. Si l'on veut faire un tel lien, il faudra alors retrancher au travail la partie due à cet apport (ou perte) d'énergie mécanique. Il faudra alors considérer chaque cas particulier. Nous allons illustrer cette approche dans la section 5.4.3 ci-dessous.

5.4.2 Calcul direct des actions électrostatiques sur un conducteur chargé

Considérons maintenant le cas d'un conducteur chargé placé dans un champ électrostatique $\vec{E}_{\text{ext}}$. Celui-ci produit une force de Coulomb sur chaque charge électrique distribuée sur la surface S du conducteur. Nous permettons que les conducteurs baignent éventuellement dans un milieu diélectrique (liquide de préférence) caractérisé par ϵ_r . D'après ce que nous avons vu précédemment, la force totale s'écrit

$$\vec{F} = \iint_S d^2 \vec{F} = \iint_S \sigma \vec{E}_{\text{ext}} dS = \iint_S P \hat{n} dS = \iint_S \frac{\sigma^2}{2\epsilon_r \epsilon_0} \hat{n} dS, \quad (5.42)$$

où $P = \sigma^2 / (2\varepsilon_r \varepsilon_0)$ est la pression électrostatique dans un milieu diélectrique. Le moment de la force électrostatique s'écrit

$$\begin{aligned}\vec{\Gamma}_O &= \iint_S \vec{OP} \wedge d^2\vec{F} = \iint_S \vec{OP} \wedge \vec{E}_{\text{ext}} \sigma dS \\ &= \iint_S \vec{OP} \wedge \hat{n} P dS = \iint_S \vec{OP} \wedge \hat{n} \frac{\sigma^2}{2\varepsilon_r \varepsilon_0} dS.\end{aligned}\quad (5.43)$$

Mais ces expressions ne sont utilisables que si l'on peut calculer la densité surfacique σ . Lorsque ce n'est pas le cas, il faut utiliser la méthode ci-dessous.

5.4.3 Calcul des actions à partir de l'énergie électrostatique (Méthode des travaux virtuels)

Soit un système de deux conducteurs chargés (A1) et (A2). Pour connaître la force $\vec{F}_{1 \rightarrow 2}$ exercée par (A1) sur (A2), on imagine une petite modification des paramètres de A2 (position et/ou orientation) et on examine le changement apporté à l'énergie électrostatique. Ceci nous permettra d'en déduire les forces (moments) associées avec ces modifications. Cette méthode s'appelle méthode des travaux virtuels. Un conducteur en équilibre électrostatique étant caractérisé par un potentiel V et une charge Q , il y a deux cas extrêmes qu'il faut considérer séparément.

a) *Système isolé : charges constantes*



On se place à l'équilibre mécanique, où les forces externes empêchent des déplacements des armatures (A1) et (A2) (c.-à-d. $\vec{F}_{\text{ext}} = -\vec{F}$). Imaginons maintenant un qu'on laisse le système faire un déplacement élémentaire autour de cette position. Nous sommes dans le cas d'un système isolé où $d\mathcal{E}_{\text{tot}} = 0$ ce qui implique par conservation d'énergie que $dW = -d\mathcal{E}_e$ (voir eq. (5.40)). L'énergie de ce travail, dW est donc fournie par l'énergie électrostatique de (A2) et nous pouvons écrire :

$$dW = \sum_{i=1}^3 F_i dx_i = \vec{F} \cdot d\vec{r} = -d\mathcal{E}_e.$$

Or, l'énergie électrostatique est une fonction de la position de (A2), donc $d\mathcal{E}_e = \sum_{i=1}^3 \left(\frac{\partial \mathcal{E}_e}{\partial x_i} \right)_Q dx_i$. Autrement dit, dans le cas d'un déplacement d'un conducteur *isolé* on doit avoir à tout moment :

$$dW = \sum_{i=1}^3 F_i dx_i = -d\mathcal{E}_e = - \sum_{i=1}^3 \left(\frac{\partial \mathcal{E}_e}{\partial x_i} \right)_Q dx_i,$$

c'est-à-dire une force électrostatique,

$$F_i = - \left(\frac{\partial \mathcal{E}_e}{\partial x_i} \right)_Q, \quad (5.44)$$

exercée par (A1) sur (A2). Notez que les variables x_i décrivent la *distance* entre (A1) et (A2). Cette force peut aussi s'interpréter comme une force interne exercée par un conducteur sur une partie de lui-même. Ainsi, cette expression est également valable dans le cas d'un conducteur qui serait soumis à une déformation : ce serait la force exercée par le conducteur sur une partie de lui-même lors d'une modification de son énergie d'interaction électrostatique, $\mathcal{E}_e$.

Dans le cas de rotations pures, l'énergie dépend des différents angles et l'on va plutôt écrire pour un conducteur *isolé* (Q constant)

$$dW = \sum_{i=1}^3 \Gamma_i d\alpha_i = -d\mathcal{E}_e = - \sum_{i=1}^3 \left(\frac{\partial \mathcal{E}_e}{\partial \alpha_i} \right)_Q d\alpha_i ,$$

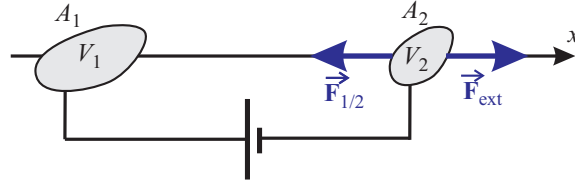
c'est-à-dire un moment des forces électrostatiques $\vec{\Gamma} = \sum_{i=1}^3 \Gamma_i \hat{e}_i$ dont les composantes vérifient :

$$\Gamma_i = - \left(\frac{\partial \mathcal{E}_e}{\partial \alpha_i} \right)_Q . \quad (5.45)$$

L'utilisation des expressions eq.(5.44) et (5.45) nécessite de calculer l'énergie électrostatique $\mathcal{E}_e$ et sa dépendance en fonction de la position du (ou des) conducteur(s).

La présence du signe moins indique que les actions électrostatiques (forces et moments) tendent toujours à ramener le conducteur vers une position d'énergie minimale.

b) *Système relié à un générateur : potentiels constants*



A proximité de l'équilibre mécanique ($\vec{F}_{\text{ext}} = -\vec{F}$) on effectue un petit déplacement autour de cette position. Les forces internes effectuent toujours un travail $dW = \vec{F} \cdot d\vec{r}$, mais il existe une deuxième source d'énergie, le générateur. Lors du déplacement, celui-ci maintient les potentiels V_1 et V_2 constants. Cela ne peut se faire qu'en modifiant la charge Q_1 et Q_2 de chaque conducteur. Ainsi, le générateur fournit un travail permettant d'amener des charges dQ_1 au potentiel V_1 et dQ_2 au potentiel V_2 , c'est-à-dire une énergie fournie $d\mathcal{E}_{\text{Gen}} = V_1 dQ_1 + V_2 dQ_2$.

Nous sommes donc dans le cas traité en eq.(5.41) où le changement d'énergie totale, $d\mathcal{E}_{\text{tot}}$ est celui fournie par le générateur (c.-à-d. $d\mathcal{E}_{\text{tot}} = d\mathcal{E}_{\text{Gen}}$). La conservation d'énergie s'écrit maintenant :

$$\begin{aligned} dW &= d\mathcal{E}_{\text{Gen}} - d\mathcal{E}_e \\ \vec{F} \cdot d\vec{r} &= V_1 dQ_1 + V_2 dQ_2 - \frac{1}{2} (V_1 dQ_1 + V_2 dQ_2) \\ \sum_{i=1}^3 F_i dx_i &= \frac{1}{2} (V_1 dQ_1 + V_2 dQ_2) = d\mathcal{E}_e = \sum_{i=1}^3 \left(\frac{\partial \mathcal{E}_e}{\partial x_i} \right)_V dx_i , \end{aligned}$$

où nous avons utilisé le fait qu'à potentiel constant, la variation d'énergie électrostatique s'écrit : $d\mathcal{E}_e = \frac{1}{2} (V_1 dQ_1 + V_2 dQ_2)$. Autrement dit, dans le cas d'un déplacement d'un conducteur relié à un générateur (V maintenu constant), la force électrostatique vaut :

$$F_i = + \left(\frac{\partial \mathcal{E}_e}{\partial x_i} \right)_V . \quad (5.46)$$

Dans le cas de rotations pures, l'énergie dépend des différents angles et l'on obtient un moment des forces électrostatiques égal à :

$$\Gamma_i = + \left(\frac{\partial \mathcal{E}_e}{\partial \alpha_i} \right)_V . \quad (5.47)$$

Les expressions obtenues dans les deux cas considérés sont générales et indépendantes du déplacement élémentaire. En fait celui-ci ne constitue qu'un artifice de calcul, connu sous le nom de méthode des travaux virtuels. Notez qu'une telle méthode s'appuie sur le principe de conservation de l'énergie et donc, nécessite l'identification de l'ensemble des sources d'énergie présentes.

5.4.4 Exemple du condensateur

L'énergie électrostatique du condensateur s'écrit $\mathcal{E}_e = \frac{1}{2}QU = \frac{1}{2}CU^2 = \frac{1}{2}\frac{Q^2}{C}$. D'après la section précédente, lorsque le condensateur est isolé, la force électrostatique entre les deux armatures s'écrit :

$$F_i = - \left(\frac{\partial \mathcal{E}_e}{\partial x_i} \right)_Q = \frac{Q^2}{2C^2} \left(\frac{\partial C}{\partial x_i} \right) = \frac{U^2}{2} \left(\frac{\partial C}{\partial x_i} \right) .$$

Par contre, lorsque le condensateur est relié à un générateur, on a

$$F_i = + \left(\frac{\partial \mathcal{E}_e}{\partial x_i} \right)_U = \frac{U^2}{2} \left(\frac{\partial C}{\partial x_i} \right) .$$

Ainsi, on vient de démontrer que, **dans tous les cas**, la force électrostatique existant entre les deux armatures d'un condensateur s'écrit :

$$\vec{F} = \frac{U^2}{2} \vec{\text{grad}} C . \quad (5.48)$$

On obtient de même que le moment par rapport à l'axe Δ_i de la force électrostatique s'écrit dans tous les cas :

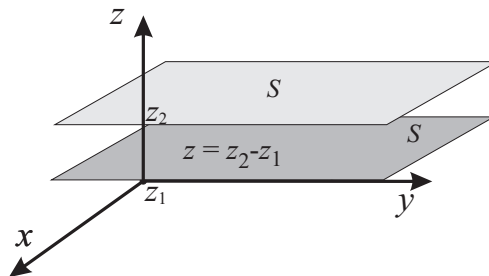
$$\Gamma_i = \frac{U^2}{2} \left(\frac{\partial C}{\partial \alpha_i} \right) . \quad (5.49)$$

Remarques :

1. On aurait pu s'attendre que les forces et moments soient les mêmes dans les deux formulations parce que peu importe qu'on soit en train de garder Q ou U constant lors d'un déplacement, les charges et champs sont les mêmes dans les deux cas à un instant donné.
2. Les actions électrostatiques tendent toujours à augmenter la capacité C d'un condensateur.
3. La force obtenue par la méthode des travaux virtuels est la même que celle donnée par l'expression $\vec{F} = \iint d^2\vec{F} = \iint P \hat{n} dS$, ce qui signifie que la distribution de charges σ doit s'arranger de telle sorte que ce soit effectivement le cas.

Exemple : le condensateur plan

Soit un condensateur plan de capacité $C(z) = \varepsilon_r \varepsilon_0 S / z$, où S est la surface d'influence mutuelle commune aux deux armatures et $z = z_2 - z_1$ la distance entre celles-ci.



La force exercée par l'armature 1 sur l'armature 2 est :

$$\begin{aligned}\vec{F}_{1 \rightarrow 2} &= \frac{U^2}{2} \overrightarrow{\text{grad}}_2 C = -\vec{F}_{2 \rightarrow 1} \\ \vec{F}_{1 \rightarrow 2} &= \hat{z} \frac{U^2}{2} \frac{\partial C}{\partial z_2} = -\hat{z} \frac{U^2}{2} \frac{\varepsilon_r \varepsilon_0 S}{(z_2 - z_1)^2} \\ \vec{F}_{2 \rightarrow 1} &= \hat{z} \frac{U^2}{2} \frac{\partial C}{\partial z_1} = \hat{z} \frac{U^2}{2} \frac{\varepsilon_r \varepsilon_0 S}{(z_2 - z_1)^2} .\end{aligned}$$

Notez que la bonne utilisation de la formule générale (portant sur le gradient de C) nécessite la compréhension de sa démonstration (ce que signifie les variables x_i).

5.4.5 Exemple du dipôle

Soit un dipôle électrostatique de moment dipolaire $\vec{p}$ placé dans un champ extérieur $\vec{E}_{\text{ext}}$. On cherche dans un premier temps à calculer la force électrostatique exercée par ce champ sur le dipôle. Celui-ci restant à charge constante, on va donc utiliser l'expression obtenue pour un système isolé. En se rappelant que l'énergie potentielle d'interaction électrique est $\mathcal{E}_e = -\vec{p} \cdot \vec{E}_{\text{ext}}$, nous avons :

$$F_i = - \left(\frac{\partial \mathcal{E}_e}{\partial x_i} \right)_Q = \frac{\partial (\vec{p} \cdot \vec{E}_{\text{ext}})}{\partial x_i} ,$$

c'est-à-dire une expression vectorielle

$$\vec{F} = \overrightarrow{\text{grad}} (\vec{p} \cdot \vec{E}_{\text{ext}}) . \quad (5.50)$$

Sous l'effet de cette force, un dipôle aura tendance à se déplacer vers les régions où le champ électrostatique est le plus fort.

Le moment de la force électrostatique est donné par

$$\Gamma_i = - \left(\frac{\partial \mathcal{E}_e}{\partial \alpha_i} \right)_Q = \frac{\partial (\vec{p} \cdot \vec{E}_{\text{ext}})}{\partial \alpha_i} ,$$

avec $\vec{\Gamma} = \sum_i \Gamma_i \hat{e}_i$. On peut cependant clarifier considérablement cette expression. Il suffit en effet de remarquer que lors d'une rotation pure, le vecteur moment dipolaire varie comme

$$d\vec{p} = \sum_{i=1}^3 d\alpha_i \hat{e}_i \wedge \vec{p} = \sum_{i=1}^3 \frac{\partial \vec{p}}{\partial \alpha_i} d\alpha_i ,$$

puisqu'il dépend a priori de la position du point considéré, donc des angles α_i . En supposant alors que le champ $\vec{E}_{\text{ext}}$ est constant à l'échelle du dipôle, on obtient,

$$\Gamma_i = \frac{\partial (\vec{p} \cdot \vec{E}_{\text{ext}})}{\partial \alpha_i} = \frac{\partial \vec{p}}{\partial \alpha_i} \cdot \vec{E}_{\text{ext}} = (\hat{e}_i \wedge \vec{p}) \cdot \vec{E}_{\text{ext}} = \hat{e}_i \cdot (\vec{p} \wedge \vec{E}_{\text{ext}}) ,$$

c'est-à-dire l'expression vectorielle suivante :

$$\vec{\Gamma} = \vec{p} \wedge \vec{E}_{\text{ext}} . \quad (5.51)$$

Le moment des forces électrostatiques a donc tendance à aligner le dipôle dans la direction du champ extérieur.

Chapitre 6

Courant et champ magnétique

6.1 Courant et résistance électriques

6.1.1 Le courant électrique

Nous avons vu qu'il était possible d'électriser un matériau conducteur, par exemple par frottements. Si l'on met ensuite ce conducteur en contact avec un autre, le deuxième devient à son tour électrisé, c'est-à-dire qu'il a acquis une certaine charge Q . Cela signifie que lors du contact des charges se sont déplacées de l'un vers l'autre. On définit alors le courant par,

$$I = \frac{dQ}{dt}, \quad (6.1)$$

où les unités sont les Ampères (symbole A). Dans le système international, l'Ampère est l'une des 4 unités fondamentales (avec le mètre, le kilogramme et la seconde), de telle sorte que

$$1\text{C} = 1\text{A s} \quad (\text{Ampère seconde}). \quad (6.2)$$

La définition de l'éq. (6.1) de I ne nous renseigne pas sur son signe, il faut choisir une convention. Par exemple, soit $Q > 0$ la charge du conducteur initialement chargé (A_1). On a affaire ici à une décharge de (A_1) vers (A_2). Si l'on désire compter positivement le courant de (A_1) vers (A_2), alors il faut mettre un signe moins à l'expression de l'éq. (6.1).

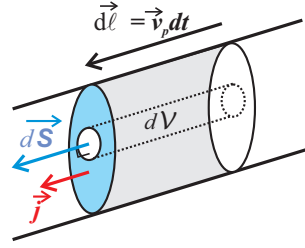
La densité de courant électrique

La raison physique du courant est un déplacement de charges, c'est-à-dire l'existence d'une vitesse organisée (par opposition à la vitesse d'agitation thermique) de celles-ci. Considérons donc un fil conducteur de section S , dans lequel se trouvent n_p porteurs de charge q_p par unité de volume, animés d'une vitesse $\vec{v}_p$ dans le référentiel du laboratoire. Pendant un instant dt , ces charges parcourent une distance $d\vec{\ell} = \vec{v}_p dt$. Soit $dS\hat{n}$ un élément infinitésimal de surface mesuré sur la section du fil, orienté dans une direction arbitraire. La quantité de charge électrique, dQ , qui traverse cette surface pendant dt est celle contenue dans le volume élémentaire $dV = dS d\ell$ associé :

$$dQ = n_p q_p dV = n_p q_p \vec{v}_p dt \cdot dS\hat{n} \equiv dt \vec{j} \cdot d\vec{S},$$

où on a définie une *densité de courant*, $\vec{j}$, qui décrit les caractéristiques locales du transport de charge par les porteurs de courant :

$$\vec{j} = n_p q_p \vec{v}_p = \rho_p \vec{v}_p, \quad (6.3)$$



exprimée en Ampères par mètre carré (A m^{-2}). La charge totale dQ traversant une section du fil est l'intégrale des charges dQ sur la section du fil. Finalement donc, le courant I circulant dans le fil est relié à la densité par

$$I = \frac{dQ}{dt} = \frac{1}{dt} \iint_{\text{section}} dQ = \frac{1}{dt} \iint_{\text{section}} \vec{j} \cdot \hat{n} dS dt ,$$

c'est-à-dire :

$$I = \iint_{\text{section}} \vec{j} \cdot d\vec{S} . \quad (6.4)$$

On dit que le courant dans un circuit est le flux à travers la section du fil de la densité de courant. Le sens du courant (grandeur algébrique) est alors donné par le sens du vecteur densité de courant.

Un conducteur est un cristal (ex, cuivre) dans lequel se déplacent des particules chargées (ex, électrons). Suivant le matériau, les porteurs de charges responsables du courant peuvent être différents. Dans un métal, ce sont des électrons, dits de conduction (la nature et le signe des porteurs de charge peuvent être déterminés grâce à l'effet Hall –voir chapitre 9). Dans un gaz constitué de particules ionisées, un plasma, ou bien dans un électrolyte, il peut y avoir plusieurs espèces chargées en présence. En toute généralité, on doit donc définir la densité locale de courant de la forme

$$\vec{j} = \sum_{\alpha} n_{\alpha} q_{\alpha} \vec{v}_{\alpha} , \quad (6.5)$$

où l'on fait une sommation sur toutes les espèces (électrons et ions) en présence. Dans le cas particulier d'un cristal composé d'ions immobiles (dans le référentiel du laboratoire) et d'électrons en mouvement, on a

$$\vec{j} = -n_e e \vec{v}_e , \quad (6.6)$$

où $e \simeq 1,6 \times 10^{-19} \text{C}$ est la charge élémentaire et n_e la densité locale d'électrons libres. **La densité de courant (donc le sens attribué à I) est ainsi dans le sens contraire du déplacement réel des électrons.**

Loi d'Ohm microscopique (ou locale)

Dans la plupart des conducteurs, on observe une proportionnalité entre la densité de courant et le champ électrostatique local,

$$\vec{j} = \gamma \vec{E} , \quad (6.7)$$

où le coefficient de proportionnalité γ (souvent dénoté par σ) est appelé la **conductivité** du milieu (unités : voir plus bas). On définit également $\eta = 1/\gamma$, la **résistivité** du milieu. La conductivité est une grandeur locale positive, dépendant uniquement des propriétés du matériau. Ainsi, le cuivre possède une conductivité $\gamma_{CU} = 58 \times 10^6 \text{ S/m}$, tandis que celle du verre (isolant) vaut $\gamma_{\text{verre}} = 10^{-11} \text{ S/m}$.

Une telle loi implique que les lignes de champ quasi-électrostatique sont également des lignes de courant, indiquant donc le chemin pris par les charges électriques. Par ailleurs, comme γ est positif, cela implique que **le courant s'écoule dans la direction des potentiels décroissants.**

D'où peut provenir cette loi ? Prenons le cas simple d'une charge électrique q_p soumise à la force de Coulomb mais aussi à des collisions (modèle de Drude). Ces collisions peuvent se décrire comme une force de frottement proportionnelle à la vitesse (moyenne) $\vec{v}_p$ de la charge. La relation fondamentale de la dynamique s'écrit :

$$m_p \frac{d\vec{v}_p}{dt} = q_p \vec{E} - k \vec{v}_p, \quad (6.8)$$

où m_p est la masse des porteurs. Cette équation montre qu'en régime permanent (stationnaire, mais non statique), la charge q_p atteint une vitesse limite $\vec{v}_l = \mu \vec{E}$ où $\mu = q_p/k$ est appelé la **mobilité** des charges. Ce régime est atteint en un temps caractéristique $\tau \simeq m_p/k$, appelé temps de relaxation.

Insérant ces valeurs dans l'expression de la densité de courant $\vec{j} = n_p q_p \vec{v}_l$, on obtient la loi d'Ohm avec :

$$\vec{j} = \frac{n_p q_p^2}{k} \vec{E} \equiv \gamma \vec{E} \quad \Rightarrow \quad \gamma = \frac{n_p q_p^2}{k} \Rightarrow k = \frac{n_p q_p^2}{\gamma}. \quad (6.9)$$

Cette expression microscopique pour γ nous permet d'en déduire le temps de temps de relaxation des porteurs de courant :

$$\tau = \frac{m_p}{k} = \frac{\gamma m_p}{n_p q_p^2}. \quad (6.10)$$

Ainsi, la loi d'Ohm microscopique (ou locale) s'explique bien par ce modèle simple de collisions des porteurs de charge. Mais collisions avec quoi ? On a longtemps cru que c'étaient des collisions avec les ions du réseau cristallin du conducteur, mais il s'avère qu'il s'agit en fait de collisions avec les impuretés contenues dans celui-ci.

Prenons le cas du cuivre, métal conducteur au sein duquel existe une densité numérique d'électrons de conduction de l'ordre de $n_e \simeq 8 \times 10^{28} \text{ m}^{-3}$. Le temps de relaxation est alors de $\tau = \frac{\gamma_{CU} m_e}{e^2 n_e} \approx 2 \cdot 10^{-14} \text{ s}$. C'est le temps typique entre deux collisions. Quelle est la distance maximale parcourue par les électrons pendant ce temps (libre parcours moyen) ? Elle dépend de leur vitesse réelle : celle-ci est la somme de la vitesse moyenne $\vec{v}_e$ (le courant) et de la vitesse d'agitation thermique de norme $v_{th} = \sqrt{kT/m_e} \approx 10^5 \text{ m/s}$ à température ambiante, mais dont la valeur moyenne (vectorielle) est nulle. La vitesse du au courant est la plupart du temps, complètement négligeable par rapport à cette très grande vitesse thermique (par exemple, un fil de Cuivre d'une section de 1 mm^2 parcouru par un courant de 1 A, possède une densité de courant de $j_e = 10^6 \text{ Am}^{-2}$ et une vitesse moyenne de $v_e = j/n_e e \simeq 0,07 \text{ mm/s}$). Le libre parcours moyen d'un électron serait alors de :

$$l = v_{th} \tau \approx 2 \cdot 10^{-9} \text{ m} = 20 \text{ Å},$$

ce qui est un ordre de grandeur supérieur à la distance inter-atomique (de l'ordre de l'Angström). Ce ne sont donc pas les collisions avec les ions du réseau du cuivre qui sont la cause de la loi d'Ohm.

Résistance d'un conducteur : loi d'Ohm macroscopique

Considérons maintenant une portion AB d'un conducteur parcouru par un courant I . S'il existe un courant, cela signifie qu'il y a une chute de potentiel entre A et B, $U = V_A - V_B = \int_A^B \vec{E} \cdot d\vec{\ell}$. On définit alors la résistance de cette portion par

$$R = \frac{U}{I} = \frac{\int_A^B \vec{E} \cdot d\vec{\ell}}{\iint_S \gamma \vec{E} \cdot d\vec{S}}, \quad (6.11)$$

où l'unité est l'Ohm (symbole Ω). Dans le cas simple d'un conducteur filiforme de section S où, sur une longueur L , le champ électrostatique est uniforme, on obtient le lien entre la résistance d'un conducteur

(propriété macroscopique) et sa résistivité, η (propriété microscopique) :

$$R = \frac{L}{\gamma S} = \eta \frac{L}{S}, \quad (6.12)$$

qui montre que les unités de la résistivité sont le Ωm (Ohm mètre).

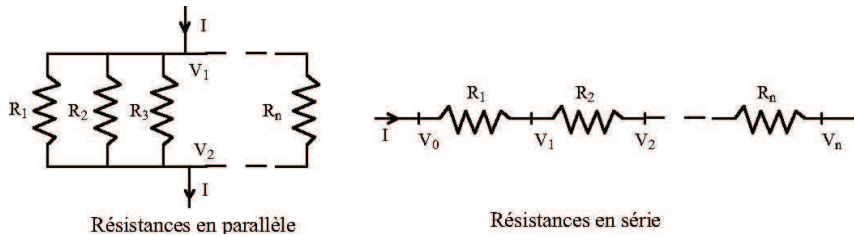
Associations de résistances

(a) Résistances en série Soient N résistances R_i mises bout à bout dans un circuit et parcourues par un courant I . La tension aux bornes de la chaîne est simplement :

$$U = (V_0 - V_1) + (V_1 - V_2) + \cdots + (V_{N-1} - V_N) = R_1 I + R_2 I + \cdots + R_N I$$

c'est-à-dire analogue à celle obtenue par une résistance unique dont la valeur est :

$$R_{\text{eq}} = \sum_{i=1}^N R_i. \quad (6.13)$$



(b) Résistances en parallèle Soient N résistances R_i mises en parallèle sous une tension $U = V_1 - V_2$ et alimentées par un courant I . Le courant se sépare alors en n courants :

$$I_i = \frac{U}{R_i},$$

dans chacune des N branches. En vertu de la conservation du courant (voir ci-dessous), on a :

$$I = \sum_{i=1}^N I_i = \sum_{i=1}^N \frac{U_i}{R_i} = \frac{U}{R_{\text{eq}}},$$

c'est-à-dire que l'ensemble des N branches est analogue à une résistance équivalente en série :

$$\frac{1}{R_{\text{eq}}} = \sum_{i=1}^N \frac{1}{R_i}. \quad (6.14)$$

6.2 Le champ magnétique

6.2.1 Bref aperçu historique

Les aimants sont connus depuis l'Antiquité, sous le nom de magnétite, pierre trouvée à proximité de la ville de Magnésie (Turquie). C'est de cette pierre que provient le nom actuel de champ magnétique.

Les chinois furent les premiers à utiliser les propriétés des aimants, il y a plus de 1000 ans, pour faire des boussoles. Elles étaient constituées d'une aiguille de magnétite posée sur de la paille flottant sur de l'eau contenue dans un récipient gradué.

Au XVIII^{ème} siècle, Franklin découvre la nature électrique de la foudre (1752). Or, il y avait déjà à cette époque de nombreux témoignages de marins attirant l'attention sur des faits étranges :

- Les orages perturbent les boussoles
 - La foudre frappant un navire aimante tous les objets métalliques.
- Franklin en déduisit « *la possibilité d'une communauté de nature entre les phénomènes électriques et magnétiques* » .

Coulomb (1785) montre la décroissance en $1/r^2$ des deux forces.

Mais il faut attendre la fin du XIX^{ème} siècle pour qu'une théorie complète apparaisse, la théorie de l'électromagnétisme. Tout commença avec l'expérience d'Oersted en 1820. Il plaça un fil conducteur au dessus d'une boussole et y fit passer un courant. En présence d'un courant l'aiguille de la boussole est effectivement déviée, prouvant sans ambiguïté un lien entre le courant électrique et le champ magnétique. Par ailleurs, il observa :

- Si on inverse le sens du courant, la déviation change de sens.
- La force qui dévie l'aiguille est non radiale

L'étude quantitative des interactions entre aimants et courants fut faite par les physiciens Biot et Savart (1820). Ils mesurèrent la durée des oscillations d'une aiguille aimantée en fonction de sa distance à un courant rectiligne. Ils trouvèrent que la force agissant sur un pôle est dirigée perpendiculairement à la direction reliant ce pôle au conducteur et qu'elle varie en raison inverse de la distance. De ces expériences, Laplace déduisit ce qu'on appelle aujourd'hui la loi de Biot et Savart. Une question qui s'est ensuite immédiatement posée fut : si un courant dévie un aimant, alors est-ce qu'un aimant peut faire dévier un courant ? Ceci fut effectivement prouvé par Davy en 1821 dans une expérience où il montra qu'un arc électrique était dévié dans l'entrefer d'un gros aimant.

L'élaboration de la théorie électromagnétique mit en jeu un grand nombre de physiciens de renom : Oersted, Ampère, Arago, Faraday, Foucault, Henry, Lenz, Maxwell, Weber, Helmholtz, Hertz, Lorentz et bien d'autres. Si elle débuta en 1820 avec Oersted, elle ne fut mise en équations par Maxwell qu'en 1873 et ne trouva d'explication satisfaisante qu'en 1905, dans le cadre de la théorie de la relativité d'Einstein.

Dans ce cours, nous traiterons dans les chapitres 6 à 10 de la question suivante : comment produire un champ magnétique à partir de courants permanents ? Dans le chapitre 10 nous aborderons le problème inverse : comment produire de l'électricité à partir d'un champ magnétique ?

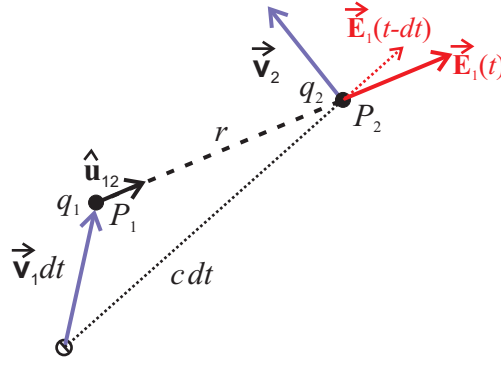
6.2.2 Nature des effets magnétiques

Jusqu'à présent nous n'avons abordé que des particules chargées immobiles, ou encore des conducteurs (ensembles de particules) en équilibre. Que se passe-t-il lorsqu'on considère enfin le mouvement des particules ? Soient deux particules q_1 et q_2 situées à un instant t aux points P_1 et P_2 . En l'absence de mouvement, la particule q_1 crée au point P_2 un champ électrostatique $\vec{E}_1(P_2)$ et la particule q_2 subit une force dont l'expression est donnée par la loi de Coulomb :

$$\vec{F}_{1 \rightarrow 2} = q_2 \vec{E}_1(P_2) .$$

Qui dit force, dit modification de la quantité de mouvement de q_2 puisque $\vec{F}_{1 \rightarrow 2} = \frac{d\vec{p}_2}{dt} \simeq \frac{\Delta \vec{p}_2}{\Delta t}$. Autrement dit la force électrostatique due à q_1 crée une modification $\Delta \vec{p}_2$ pendant un temps Δt . Une force correspond en fait à un transfert d'information (ici de q_1 vers q_2) pendant un court laps de temps. Or, rien ne peut se propager plus vite que la vitesse c de la lumière. Cette vitesse étant grande mais finie, tout transfert d'information d'un point de l'espace à un autre prend nécessairement un temps fini. Ce temps pris par la propagation de l'information introduit donc un retard, comme nous allons le voir.

On peut considérer l'exemple ci-dessus comme se qui se passe effectivement dans le référentiel propre de q_1 . Dans un référentiel fixe, q_1 est animée d'une vitesse $\vec{v}_1$. Quelle serait alors l'action de



q_1 sur une particule q_2 animée d'une vitesse $\vec{v}_2$? Soit dt le temps qu'il faut à l'information (le champ électrostatique créé par q_1) pour se propager de q_1 vers q_2 . Pendant ce temps, q_1 parcourt une distance $v_1 dt$ et q_2 parcourt la distance $v_2 dt$. Autrement dit, lorsque q_2 ressent les effets électrostatiques dus à q_1 , ceux-ci ne sont plus radiaux : le champ $\vec{E}_1(t - dt)$ « vu » par q_2 est dirigé vers l'ancienne position de q_1 et dépend de la distance $c dt$ et non pas de la distance r . On voit ici qu'il faut corriger la loi de Coulomb qui nous aurait donné le champ $\vec{E}_1(t)$ qui est faux (suppose propagation instantanée de l'information *i.e.* une vitesse infinie). Les effets électriques ne peuvent se résumer au champ électrostatique.

Cependant, l'expérience montre que la prise en compte d'une correction tenant compte de la vitesse finie du transport d'information ne suffirait pas à expliquer la trajectoire de q_2 : une force supplémentaire apparaît, plus importante d'ailleurs que celle associée avec la vitesse finie d'information ! En première approximation, la force totale exercée par q_1 sur q_2 s'écrit en fait :

$$\vec{F}_{1 \rightarrow 2} \simeq \frac{q_1 q_2}{4\pi\epsilon_0 r_{12}^2} \left[\hat{u}_{12} + \frac{\vec{v}_2}{c} \wedge \left(\frac{\vec{v}_1}{c} \wedge \hat{u}_{12} \right) \right]. \quad (6.15)$$

Dans cette expression, on voit donc apparaître un deuxième terme qui dépend des vitesses des deux particules ainsi que la vitesse de propagation de la lumière. Ce deuxième terme s'interprète comme la contribution d'un champ magnétique créé par q_1 . Autrement dit :

$$\vec{F}_{1 \rightarrow 2} \simeq q_2 \left(\vec{E}_1(P_2) + \vec{v}_2 \wedge \vec{B}_1(P_2) \right), \quad (6.16)$$

où nous avons introduit un champ magnétique, $\vec{B}_1$, créé par la charge q_1 où :

$$\vec{B}_1(P_2) = \frac{q_1}{4\pi\epsilon_0 c^2} \frac{(\vec{v}_1 \wedge \hat{u}_{12})}{r_{12}^2} = \frac{q_1}{4\pi\epsilon_0 c^2} \frac{\vec{v}_1 \wedge \overrightarrow{P_1 P_2}}{|\overrightarrow{P_1 P_2}|^3}. \quad (6.17)$$

(nous avons utilisé le fait que $r_{12} = |\overrightarrow{P_1 P_2}|$ et $\hat{u}_{12} = \overrightarrow{P_1 P_2} / |\overrightarrow{P_1 P_2}|$). L'unité SI du champ magnétique est le Tesla. A partir de l'équation (6.16) on en déduit qu'en unités fondamentales 1T correspond à $1\text{kg.s}^{-2}.\text{A}^{-1}$.

L'inspection des équations (6.15) à (6.17) montre que la force magnétique a un rapport de $(v/c)^2$ à la force de Coulomb (donc normalement très faible). Les expressions (6.17) et (6.16) sont le fruit de décennies de travail par nombreux physiciens de renom. Nous renonçons donc à un développement « historique » du sujet et on va utiliser ces expressions, que l'on admettra, afin de dériver les lois fondamentales de la magnétostatique (obtenues historiquement à travers des expériences).

Nous reviendrons plus tard (chapitre 9) sur les propriétés de la force magnétique exprimée dans l'éq. (6.16). Cette expression n'est valable que pour des particules se déplaçant à des vitesses beaucoup plus petites que celle de la lumière (approximation de la magnétostatique). Dernière remarque : la force magnétique dépend de la vitesse de la particule, ce qui implique que le champ magnétique dépend du choix du système référentiel (voir discussion chapitre 9) !

6.3 Expressions du champ magnétique

6.3.1 Champ magnétique créé par une charge en mouvement

D'après l'éq. (6.17) ci-dessus, le champ magnétique créé en un point M par une particule de charge q située en un point P et animée d'une vitesse $\vec{v}$ dans un référentiel galiléen est :



$$\vec{B}(M) = \frac{\mu_0 q \vec{v} \wedge \overrightarrow{PM}}{4\pi |\overrightarrow{PM}|^3},$$

où $\mu_0 = 1/(\epsilon_0 c^2)$. L'unité du champ magnétique dans le système international est le **Tesla (T)** ($1 \text{ kg.s}^{-2}.\text{A}^{-1}$). Une autre unité appartenant au système CGS, le **Gauss (G)**, est également très souvent utilisée :

$$1 \text{ Gauss} = 10^{-4} \text{ Tesla} \quad . \quad (6.18)$$

Le facteur μ_0 est la perméabilité du vide : il décrit la capacité du vide à « laisser passer » le champ magnétique. Sa valeur dans le système d'unités international MKSA est :

$$\mu_0 \equiv 4\pi 10^{-7} \text{ H.m}^{-1} \text{ (H pour Henry)} \quad (6.19)$$

$$(1 \text{ Henry} = 1 \text{ V.A}^{-1}\text{s} = 1 \text{ N.m.A}^{-2} = 1 \text{ m}^2.\text{kg.s}^{-2}.\text{A}^{-2}) .$$

Remarques :

- La valeur de μ_0 donné en l'éq. (6.19) est exacte, directement liée à la définition de l'Ampère (voir Chapitre 9). Le facteur 4π a été introduit pour simplifier les équations de Maxwell.
- Nous avons vus que les phénomènes électriques et magnétiques sont intimement liés. Les expériences de l'époque montrèrent que la vitesse de propagation était toujours la même, à savoir c , la vitesse de la lumière. Cela signifiait qu'il y avait donc un lien secret entre le magnétisme, l'électricité et la lumière, et plongeait les physiciens dans la plus grande perplexité. On remarque que :

$$\mu_0 \epsilon_0 c^2 = 1 \quad , \quad (6.20)$$

ce qui permet de définir la valeur de la permittivité du vide (caractéristique décrivant sa capacité à affaiblir les forces électrostatiques) :

$$\epsilon_0 = \frac{1}{c^2 \mu_0} \simeq \frac{10^7}{4\pi . 9.10^{16}} = \frac{10^{-9}}{36\pi} \text{ F.m}^{-1} \quad (\text{F pour Farad}) .$$

Deux propriétés importantes du champ magnétique :

- De même que pour le champ électrostatique, **le principe de superposition** s'applique au champ magnétique. Si on considère deux particules 1 et 2 alors le champ magnétique créé en un point M quelconque de l'espace sera la somme vectorielle des champs créés par chaque particule.
- Du fait du produit vectoriel, le champ magnétique est ce qu'on appelle un pseudo-vecteur (voir plus bas).

6.3.2 Champ magnétique créé par un ensemble de charges en mouvement

Considérons N particules de charges q_i situées en des points P_i et de vitesse $\vec{v}_i$. En vertu du principe de superposition, le champ magnétique créé en un point M est la somme vectorielle des champs créés par chaque particule et vaut

$$\vec{B}(M) = \frac{\mu_0}{4\pi} \sum_{i=1}^N \frac{q_i \vec{v}_i \wedge \overrightarrow{P_i M}}{|\overrightarrow{P_i M}|^3}.$$

Si le nombre de particules est très grand dans un volume $\mathcal{V}$ donné et qu'on s'intéresse à des échelles spatiales bien plus grandes que la distance entre ces particules, il est avantageux d'utiliser une description continue. Il faut donc définir des distributions continues comme nous l'avons fait en électrostatique. Mais des distributions continues de quoi ?

Le passage à la limite continue consiste à assimiler tout volume élémentaire $d\mathcal{V}$, situé autour d'un point P' quelconque de la distribution de charges en mouvement, à une charge dq animée d'une vitesse moyenne $\vec{v}$. Le champ magnétique résultant s'écrit alors

$$\vec{B}(M) = \frac{\mu_0}{4\pi} \int_V \frac{dq \vec{v} \wedge \overrightarrow{PM}}{|\overrightarrow{PM}|^3},$$

où l'intégrale porte sur le volume V total embrassé par ces charges. En toute généralité, considérons α espèces différents de particules (ex : électrons, ions) chacune animée d'une vitesse $\vec{v}_\alpha$ de charge q_α et d'une densité numérique n . On peut alors écrire $dq \vec{v} = \sum_\alpha n_\alpha q_\alpha \vec{v}_\alpha d\mathcal{V}$, où la somme porte sur le nombre d'espèces différentes et non sur le nombre de particules. On reconnaît ainsi l'expression générale du vecteur densité locale de courant $\vec{j} = \sum_\alpha n_\alpha q_\alpha \vec{v}_\alpha$.

L'expression du champ magnétique créé par une distribution volumique de charges quelconque est donc :

$$\boxed{\vec{B}(M) = \frac{\mu_0}{4\pi} \iiint \frac{\vec{j}(P) \wedge \overrightarrow{PM}}{|\overrightarrow{PM}|^3} d\mathcal{V}}. \quad (6.21)$$

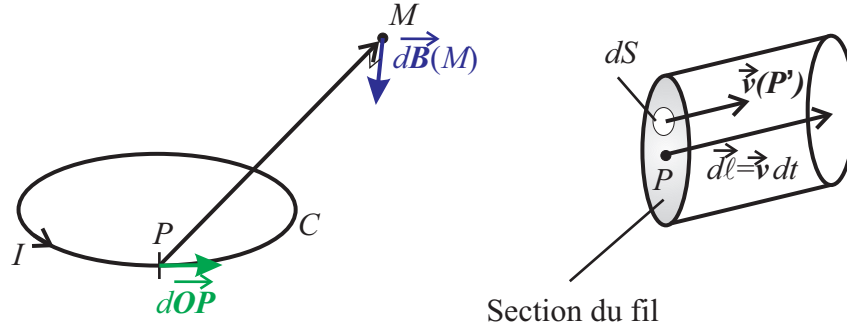
Ce résultat est général et valable quelle que soit la forme du conducteur. On peut l'appliquer, par exemple, à l'intérieur d'un métal de volume V quelconque.

Quelques ordres de grandeur :

- Champ magnétique interstellaire moyen : $B \approx \mu G$
- Champ magnétique terrestre : $B_\perp \approx 0,4 G$, $B_{\text{horizontal}} \approx 0,3 G$
- Un aimant courant $B \approx 10 \text{ mT}$
- Champ magnétique dans une tache solaire $B \approx kG \approx 0.1 \text{ Tesla}$
- Un électroaimant ordinaire $B \approx \text{Tesla}$
- Une bobine supraconductrice $B \approx 20 \text{ Tesla}$
- Champ magnétique d'une étoile à neutrons 10^8 Tesla

6.3.3 Champ créé par un circuit électrique (formule de Biot et Savart)

Dans le cas particulier d'un circuit filiforme fermé, parcouru par un courant permanent I , la formule précédente va nous fournir la loi de Biot et Savart. Dans ce cas, le volume élémentaire s'écrit $d\mathcal{V} = dS dl$ où dS est un élément de surface transverse situé en P et dl un élément de longueur du fil. Or, on considère seulement les cas où le point M est situé à une distance telle du fil qu'on peut considérer



celui-ci comme très mince. Plus précisément, le vecteur vitesse (ou densité de courant) a la même orientation sur toute la section du fil ($\vec{j}$ parallèle à $\vec{d\ell}$ et à $\vec{dS}$). Ainsi, on écrit,

$$\begin{aligned}\vec{B}(M) &= \frac{\mu_0}{4\pi} \oint_{\text{circuit}} d\ell \iint_{\text{section}} \frac{\vec{j}(P') \wedge \overrightarrow{P'M}}{|\overrightarrow{P'M}|^3} dS = \frac{\mu_0}{4\pi} \oint_{\text{circuit}} \frac{\left[\iint_{\text{section}} \vec{j}(P') dS \right] d\ell \wedge \overrightarrow{PM}}{|\overrightarrow{PM}|^3} \\ &= \frac{\mu_0}{4\pi} \oint_{\text{circuit}} \frac{\left[\iint_{\text{section}} \vec{j}(P) \cdot \vec{dS} \right] \vec{d\ell} \wedge \overrightarrow{PM}}{|\overrightarrow{PM}|^3} = \frac{\mu_0}{4\pi} \oint_{\text{circuit}} \frac{I \vec{d\ell} \wedge \overrightarrow{PM}}{|\overrightarrow{PM}|^3},\end{aligned}$$

où l'on a utilisé $P'M \gg PP'$ (donc $\overrightarrow{P'M} \simeq \overrightarrow{PM}$), P étant un point sur le fil (centre de la section). Par ailleurs, nous avons utilisé le fait que la normale à la section ainsi que $\vec{d\ell}$ étaient orientés dans le sens du courant ($\vec{j} dS d\ell = (\vec{j} \cdot \vec{dS}) \vec{d\ell}$).

Formule de Biot et Savart : en un point M quelconque de l'espace, le champ magnétique créé par un circuit parcouru par un courant permanent I est :

$$\boxed{\vec{B}(M) = \frac{\mu_0 I}{4\pi} \oint_{\text{circuit}} \frac{\vec{d\ell} \wedge \overrightarrow{PM}}{|\overrightarrow{PM}|^3}}, \quad (6.22)$$

où P est un point quelconque le long du circuit et $\vec{d\ell} = \overrightarrow{dOP}$ (O étant l'origine, **quelconque**, du système).

La formule de Biot et Savart (1820) a été établie expérimentalement et fournit un lien explicite entre le champ magnétique et le courant. Mais ce n'est que plus tard (1880+) que les physiciens ont réalisé que le courant était dû au déplacement de particules dans un conducteur.

Règles mnémotechniques :

Dans l'utilisation de la formule de Biot et Savart, il faut faire attention au fait que le champ magnétique créé par un circuit est la *somme vectorielle* de tous les $\vec{dB}$, engendrées par un élément de circuit, dont le sens est donné par celui du courant I ,

$$\boxed{d\vec{B}(M) = \frac{\mu_0 I}{4\pi} \frac{\vec{d\ell} \wedge \overrightarrow{PM}}{|\overrightarrow{PM}|^3}}.$$

Or chaque $\vec{dB}$ est défini par un produit vectoriel. Il faut donc faire extrêmement attention à l'orientation des circuits. Voici quelques règles mnémotechniques :

- Règle des trois doigts de la main droite
- Règle du bonhomme d'Ampère

6.3.4 Propriétés de symétrie du champ magnétique

Ces propriétés sont fondamentales car elles permettent de simplifier considérablement le calcul du champ magnétique. Du fait que le champ soit un effet créé par un courant, il contient des informations sur les causes qui lui ont donné origine. Cette règle se traduit par la présence de certaines symétries et invariances si les sources de courant en possèdent également. Ainsi, si l'on connaît les propriétés de symétrie de la densité de courant, on pourra connaître celles du champ magnétique.

Vecteurs et pseudo-vecteurs

Un vecteur polaire, ou vrai vecteur, est un vecteur dont la direction, le module et le sens sont parfaitement déterminés. Exemples : vitesse d'une particule, champ électrostatique, densité de courant. Un vecteur axial, ou pseudo-vecteur, est un vecteur dont le sens est défini à partir d'une convention d'orientation d'espace et dépend donc de cette convention. Exemples : le vecteur rotation instantanée, le champ magnétique, la normale à une surface.

Cette différence provient du produit vectoriel : le sens du produit vectoriel dépend de la convention d'orientation de l'espace. Le produit vectoriel de deux vrais vecteurs (respectivement pseudo-vecteurs) est un pseudo-vecteur (resp. vrai vecteur), tandis que celui d'un vrai vecteur par un pseudo-vecteur est un pseudo-vecteur.

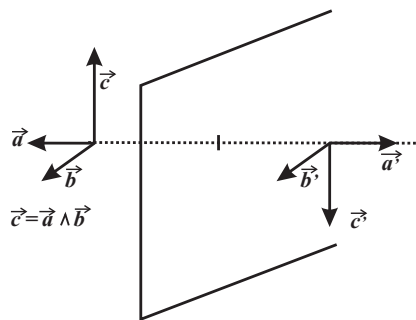


FIGURE 6.1 – Transformation par rapport à un plan de symétrie

Orienter l'espace revient à associer à un axe orienté un sens de rotation dans un plan perpendiculaire à cet axe. Le sens conventionnellement choisi est déterminé par la règle du tire-bouchon de Maxwell ou la règle du bonhomme d'Ampère (pour le champ magnétique mais aussi pour le vecteur rotation instantanée).

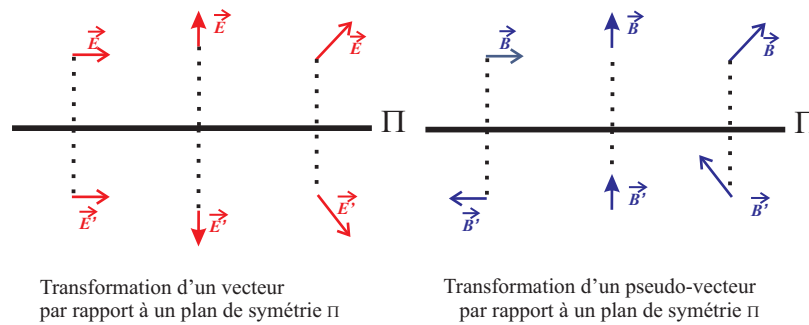
Transformations géométriques d'un vecteur ou d'un pseudo-vecteur

Vecteurs et pseudo-vecteurs se transforment de la même manière sous une rotation ou une translation. Il n'en est pas de même dans la symétrie par rapport à un plan ou à un point. Dans ces transformations

- un vecteur est transformé en son symétrique,
- un pseudo-vecteur est transformé en l'opposé du symétrique.

Soit $\vec{A}'(M')$ le vecteur obtenu par symétrie par rapport à un plan Π à partir de $\vec{A}(M)$. D'après la figure ci-dessus, on voit que

1. Si $\vec{A}(M)$ est un vrai vecteur
 - $\vec{A}'(M') = \vec{A}(M)$ si $\vec{A}(M)$ est engendré par les mêmes vecteurs de base que S ;
 - $\vec{A}'(M') = -\vec{A}(M)$ si $\vec{A}(M)$ est perpendiculaire à S .
2. au contraire, si $\vec{A}(M)$ est un pseudo-vecteur
 - $\vec{A}'(M') = \vec{A}(M)$ si $\vec{A}(M)$ est perpendiculaire à S .
 - $\vec{A}'(M') = -\vec{A}(M)$ si $\vec{A}(M)$ est engendré par les mêmes vecteurs de base que S ;



Ces deux règles de transformation vont nous permettre de déterminer des règles de symétrie utiles.

Dans un espace homogène et isotrope, si l'on fait subir une transformation géométrique à un système physique (ex : ensemble de particules, distribution de charges et/ou de courants) susceptible de créer certains effets (forces, champs), alors ces effets subissent les mêmes transformations. Si un système physique S possède un certain degré de symétrie, on pourra alors déduire les effets créés par ce système en un point à partir des effets en un autre point.

Règles de symétrie

- **Invariance par translation** : si S est invariant dans toute translation parallèle à un axe Oz , les effets ne dépendent pas de z .
- **Symétrie axiale** : si S est invariant dans toute rotation ϕ autour d'un axe Oz , alors ses effets exprimés en coordonnées cylindriques (ρ, ϕ, z) ne dépendent pas de ϕ .
- **Symétrie cylindrique** : si S est invariant par translation le long de l'axe Oz et rotation autour de ce même axe, alors ses effets exprimés en coordonnées cylindriques (ρ, ϕ, z) ne dépendent que de la distance à l'axe ρ .
- **Symétrie sphérique** : si S est invariant dans toute rotation autour d'un point fixe O , alors ses effets exprimés en coordonnées sphériques (r, θ, ϕ) ne dépendent que de la distance au centre r .
- **Plan de symétrie Π** : Si S admet un plan de symétrie Π , alors en tout point de ce plan,
 - un effet à caractère vectoriel est contenu dans le plan
 - un effet à caractère pseudo-vectoriel lui est perpendiculaire.
- **Plan d'antisymétrie Π'** : si, par symétrie par rapport à un plan Π' , S est transformé en $-S$ alors en tout point de ce plan,
 - un effet à caractère vectoriel est perpendiculaire au plan
 - un effet à caractère pseudo-vectoriel est contenu dans ce plan.

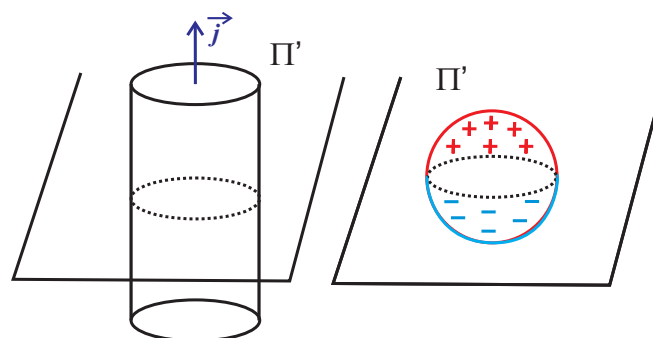


FIGURE 6.2 – Exemples de plans d'anti-symétrie

Voici quelques règles simples et très utiles, directement issues de la liste ci-dessus :

- Si $\vec{j}$ est invariant par rotation autour d'un axe, $\vec{B}$ l'est aussi.

- Si $\vec{j}$ est poloïdal (porté par $\hat{\rho}$ et/ou $\hat{z}$), alors $\vec{B}$ est toroïdal (porté par $\hat{\phi}$).
- Si $\vec{j}$ est toroïdal, alors $\vec{B}$ est poloïdal.
- Si le système de courants possède un plan de symétrie, alors $\vec{j}$ est contenu dans ce plan et donc $\vec{B}$ lui est perpendiculaire.

6.4 Calcul du champ dans quelques cas simples

6.4.1 Champ créé par un segment de fil rectiligne

On considère le champ magnétique $\delta\vec{B}$ créé par un segment rectiligne P_1P_2 parcouru par un courant d'intensité I en un point M situé à une distance ρ du fil. On écrit $\delta\vec{B}$ parce que ce segment ne peut être qu'une partie d'un circuit complet, et il faut en général calculer le champ $\vec{B}$ produit par le circuit en entier (par superposition).

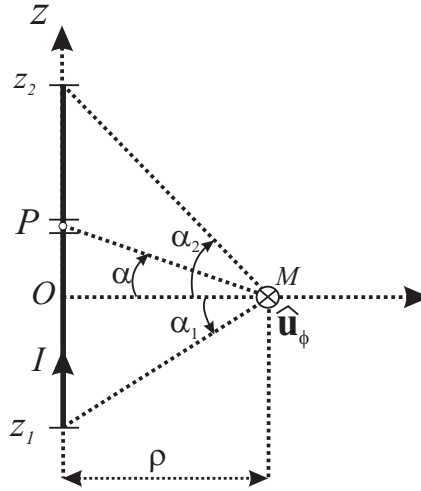


FIGURE 6.3 – Segment rectiligne d'un circuit

A partir de la loi Biot-Savart on écrit :

$$\delta\vec{B}(M) = \frac{\mu_0}{4\pi} I \int_{P_1}^{P_2} d\vec{\ell} \wedge \frac{\vec{PM}}{|\vec{PM}|^3}. \quad (6.23)$$

Compte tenu de la symétrie, on utilise des coordonnées cylindriques avec $\vec{OP} = z\hat{z}$ et on a :

$$d\vec{\ell} = \hat{z}dz \quad \vec{PM} = -z\hat{z} + \rho\hat{\rho}.$$

Avec ces coordonnées on trouve :

$$d\vec{\ell} \wedge \vec{PM} = dz\hat{z} \wedge (-z\hat{z} + \rho\hat{\rho}) = dz\rho\hat{\phi}. \quad (6.24)$$

Il convient dans ce problème de faire un changement de variables vers l'angle α ($\vec{MO}, \vec{MP}$) et on peut ainsi écrire .

$$|\vec{OP}| = z = \rho \tan \alpha, \quad PM \equiv |\vec{PM}| = \frac{\rho}{\cos \alpha}.$$

On trouve dz en prenant la différentielle de $z = \rho \tan \alpha$:

$$dz = \rho d(\tan \alpha) = \rho \left(1 + \frac{\sin^2 \alpha}{\cos^2 \alpha} \right) d\alpha = \frac{\rho}{\cos^2 \alpha} d\alpha. \quad (6.25)$$

Mettant les relations expressions de (6.24) et (6.25) dans (6.23), on obtient :

$$\begin{aligned}\vec{\delta B}(\rho) &= \frac{\mu_0}{4\pi} I \int_{z_1}^{z_2} \frac{d\vec{\ell} \wedge \vec{PM}}{PM^3} = \frac{\mu_0}{4\pi} I \hat{\phi} \int_{z_1}^{z_2} \frac{dz\rho}{PM^3} \\ &= \frac{\mu_0}{4\pi\rho} I \hat{\phi} \int_{\alpha_1}^{\alpha_2} \cos\alpha d\alpha \\ &= \frac{\mu_0}{4\pi\rho} I \hat{\phi} (\sin\alpha_2 - \sin\alpha_1) .\end{aligned}\quad (6.26)$$

On peut également exprimer le résultat en fonction de z en utilisant $\sin\alpha = z/(z^2 + \rho^2)^{1/2}$ (voir figure 6.3) :

$$\vec{\delta B}(\rho) = \frac{\mu_0}{4\pi\rho} I \hat{\phi} \left(\frac{z_2}{(z_2^2 + \rho^2)^{1/2}} - \frac{z_1}{(z_1^2 + \rho^2)^{1/2}} \right) . \quad (6.27)$$

6.4.2 Champ créé par un fil infini

Dans l'approximation, où le point M est suffisamment près du fil pour utiliser l'approximation du fil « infini » ($\rho \ll z_1$ et $\rho \ll z_2$), on peut ignorer les contributions des autres portions du circuit et on prend la limite $z_1 \rightarrow -\infty$, $z_2 \rightarrow \infty$ ($\alpha_1 \rightarrow -\frac{\pi}{2}$, $\alpha_2 \rightarrow \frac{\pi}{2}$) en équations (6.27) ou (6.26) ce qui donne :

$$\boxed{\vec{B}(\rho) \rightarrow \frac{\mu_0}{2\pi\rho} I \hat{\phi}} . \quad (6.28)$$

On constate que dans la limite du fil infini, la nature toroïdale du champ aurait pu être déduite directement des symétries du problème. Nous verrons dans le chapitre 7 comment obtenir (6.28) facilement en utilisant le théorème d'Ampère.

6.4.3 Spire circulaire (sur l'axe)

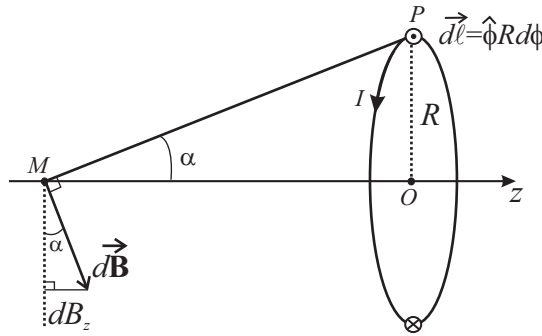


FIGURE 6.4 – Spire de rayon R parcourue par un courant

Considérons maintenant le cas d'une spire circulaire de centre O et de rayon R , parcourue par un courant permanent I (voir la figure 6.4). La densité de courant étant toroïdale et invariante par rotation autour de l'axe z , c'est-à-dire :

$$\vec{j}(\rho, \phi, z) = j(\rho, z) \hat{\phi} .$$

Le champ magnétique sera poloïdal en conséquence (c.-à-d.) :

$$\vec{B}(\rho, \phi, z) = B_\rho(\rho, z) \hat{\rho} + B_z(\rho, z) \hat{z} .$$

Comme le champ pour une spire de taille finie est assez complexe, on ne s'intéresse ici qu'au champ magnétique sur l'axe Oz de la spire. Sur l'axe z , la composante radiale du champ doit s'annuler à cause de la symétrie du système et nous n'avons qu'à calculer la composante selon $\hat{z}$ (la superposition des composantes de $\vec{B}$ selon $\hat{\rho}$ sera nulle - on peut aussi remarquer que tout plan qui contient l'axe Oz est un plan d'antisymétrie du système et que le champ $\vec{B}$ est contraint à appartenir à chacun de ses plans, donc il est obligé d'être orienté selon $\hat{z}$).

En projetant la loi de Biot et Savart de l'éq. (6.22) sur l'axe Oz , on obtient la composante dB_z produit par un élément $d\vec{\ell} = d\vec{OP} = \hat{\phi}dOP$ du circuit :

$$dB_z = d\vec{B} \cdot \hat{z} = |d\vec{B}| \sin \alpha = \frac{\mu_0 I}{4\pi} \frac{dOP \sin \alpha}{PM^2} = \frac{\mu_0 I}{4\pi} \frac{dOP \sin^3 \alpha}{R^2} = \frac{\mu_0 I}{4\pi} \frac{\sin^3 \alpha d\phi}{R},$$

où nous avons utilisé : $d\ell = dOP = R d\phi$, $PM = R/\sin \alpha$, et la figure 6.4 pour évaluer $d\vec{B} \cdot \hat{z}$.

En intégrant les composants, dB_z , sur tout le circuit, on obtient :

$$B_z = \hat{z} \cdot \frac{\mu_0 I}{4\pi} \oint_{\text{spire}} \frac{d\vec{\ell} \wedge \vec{PM}}{|\vec{PM}|^3} = \oint_{\text{spire}} dB_z = \frac{\mu_0 I \sin^3 \alpha}{4\pi R} \int_0^{2\pi} d\phi = \frac{\mu_0 I \sin^3 \alpha}{2 R}. \quad (6.29)$$

Quand il faut exprimer l'éq. (6.29) en fonction de la position, z , on utilise le fait que $\sin \alpha = R/(z^2 + R^2)^{1/2}$ dans l'éq. (6.29) pour écrire :

$$B_z(z) = \frac{\mu_0 I}{2R} \sin^3 \alpha = \frac{\mu_0 I}{2R} \frac{1}{\left(1 + \frac{z^2}{R^2}\right)^{3/2}}. \quad (6.30)$$

6.4.4 Champ d'un solénoïde fini (sur l'axe)

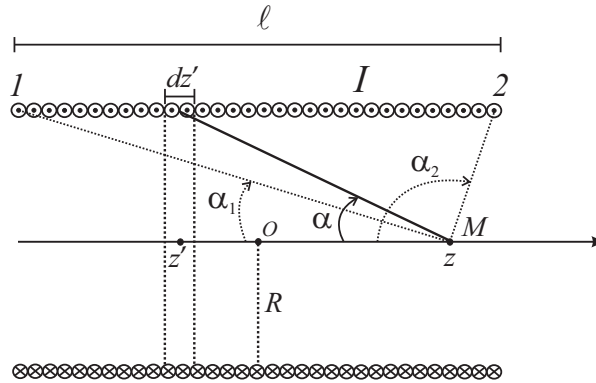


FIGURE 6.5 – Modèle d'un Solénoïde fini : Calcul de champ magnétique sur l'axe

Un solénoïde est constitué d'un enroulement d'un fil conducteur autour d'un cylindre de longueur ℓ et de rayon R . On suppose que ce fil est suffisamment mince pour pouvoir modéliser ce solénoïde comme une juxtaposition de N spires coaxiales, avec $n = N/\ell$ spires par unité de longueur. Chaque spire est alors parcourue par un courant permanent I . Comme pour la spire simple vue plus haut, les propriétés de symétrie du courant montrent que le champ magnétique du solénoïde, qui est la somme vectorielle du champ créé par chaque spire, est suivant z uniquement. Autour d'un point P situé en z' , sur une épaisseur $dOP = dz'$, il y a ndz' spires. Ces spires créent donc un champ en un point M ($OM = z$) quelconque de l'axe,

$$dB_z(z) = \frac{\mu_0 I n dz'}{2R} \sin^3 \alpha,$$

où nous avons utilisé (6.29) et l'angle α est reliée à z et z' par :

$$\cot \alpha = \frac{z - z'}{R} .$$

Si on prend la différentielle de cette expression, on obtient :

$$dz' = -Rd\left(\frac{\cos \alpha}{\sin \alpha}\right) = -R\left(-1 - \frac{\cos^2 \alpha}{\sin^2 \alpha}\right) d\alpha = \frac{R}{\sin^2 \alpha} d\alpha .$$

Le champ magnétique total s'écrit donc,

$$\begin{aligned} B_z &= \int_1^2 dB_z = \frac{\mu_0 I n}{2} \int_{\alpha_1}^{\alpha_2} \sin \alpha d\alpha \\ &= \frac{\mu_0 I n}{2} (\cos \alpha_1 - \cos \alpha_2) . \end{aligned} \quad (6.31)$$

6.4.5 Solénoïde infini

Dans la limite du solénoïde infini $R \ll \ell$, on a $\alpha_1 \rightarrow 0$ et $\alpha_2 \rightarrow \pi$ et l'expression de l'éq. (6.31) pour le champ sur l'axe devient

$$\boxed{\vec{B} \rightarrow \mu_0 I n \hat{z}} , \quad (6.32)$$

où on se rappelle que $n = N/\ell$ correspond au nombre de spires par unité de longueur. Avec le théorème d'Ampère du chapitre 7 nous pourrions montrer que ceci est la valeur partout à l'intérieur du solénoïde (dans la limite de $R \ll \ell$) et que $\vec{B} = \vec{0}$ à l'extérieur du solénoïde. On remarquera les analogies du solénoïde vis-à-vis du champ $\vec{B}$ par rapport au condensateur pour le champ $\vec{E}$.

Chapitre 7

Lois fondamentales de la magnétostatique

7.1 Flux du champ magnétique

7.1.1 Conservation du flux magnétique

Nous avons vu dans le chapitre précédent que la forme microscopique de la loi de Biot-Savart permet de calculer le champ magnétique créé par n'importe quelle distribution de courant volumique $\vec{j}$ (voir l'éq. (6.21)). Cette loi est l'analogue de la loi de Coulomb en électrostatique qui nous a permis d'en déduire les deux lois fondamentales de l'électrostatique : $\text{div} \vec{E} = \rho/\epsilon_0$ et $\text{rot} \vec{E} = \vec{0}$.

De façon analogue, nous verrons dans ce chapitre, que la loi de Biot et Savart nous permet d'en déduire deux lois fondamentales de la magnétostatique. Pour ce faire, nous allons d'abord écrire eq. (6.21) avec une notation un peu plus sophistiquée où $\vec{r}$ indique la position du point de mesure du champ M et $\vec{r}'$ indique la position de la source du champ P . L'élément de volume, dV autour du point P sera ici dénoté d^3r' et la forme microscopique de la loi de Biot-Savart avec cette notation s'écrit :

$$\vec{B}(\vec{r}) = \frac{\mu_0}{4\pi} \iiint \left[\vec{j}(\vec{r}') \wedge \frac{\vec{r} - \vec{r}'}{|\vec{r} - \vec{r}'|^3} \right] d^3r'. \quad (7.1)$$

Commençons par prendre la divergence du champ $\vec{B}$ en magnétostatique :

$$\begin{aligned} \text{div}_r \vec{B}(\vec{r}) &= \frac{\mu_0}{4\pi} \iiint \text{div}_r \left[\vec{j}(\vec{r}') \wedge \frac{\vec{r} - \vec{r}'}{|\vec{r} - \vec{r}'|^3} \right] d^3r' \\ &= \frac{\mu_0}{4\pi} \iiint \left\{ \frac{\vec{r} - \vec{r}'}{|\vec{r} - \vec{r}'|^3} \cdot \overrightarrow{\text{rot}}_r \vec{j}(\vec{r}') - \vec{j}(\vec{r}') \cdot \overrightarrow{\text{rot}}_r \overrightarrow{\text{grad}}_r \frac{1}{|\vec{r} - \vec{r}'|} \right\} d^3r', \end{aligned} \quad (7.2)$$

où l'indice r sur div_r , $\overrightarrow{\text{rot}}_r$, et $\overrightarrow{\text{grad}}_r$ nous rappelle que ces opérateurs agissent sur la variable $\vec{r}$ (et pas sur $\vec{r}'$). Dans la deuxième ligne, nous avons utilisé l'identité :

$$\text{div}(\vec{a} \wedge \vec{b}) = \vec{b} \cdot \overrightarrow{\text{rot}} \vec{a} - \vec{a} \cdot \overrightarrow{\text{rot}} \vec{b}, \quad (7.3)$$

(dérivé dans l'éq. (11.22) de l'Annexe des mathématiques) et le fait que :

$$-\overrightarrow{\text{grad}}_r \frac{1}{|\vec{r} - \vec{r}'|} = \frac{\vec{r} - \vec{r}'}{|\vec{r} - \vec{r}'|^3}. \quad (7.4)$$

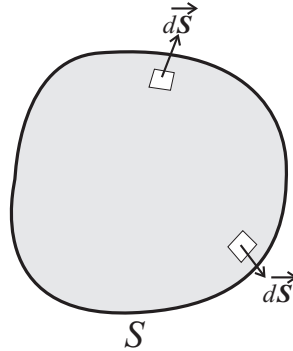
Puisque $\vec{j}(\vec{r}')$ est indépendant de $\vec{r}$, on a $\overrightarrow{\text{rot}}_r \vec{j}(\vec{r}') = \vec{0}$. Insérant cette relation et l'identité,

$$\overrightarrow{\text{rot}} \overrightarrow{\text{grad}} f \equiv 0, \quad (7.5)$$

dans l'éq. (7.2) nous donne une loi fondamentale de la magnétostatique :

$$\boxed{\operatorname{div} \vec{B} = 0} . \quad (7.6)$$

Considérons maintenant une surface fermée S quelconque, c'est-à-dire pour laquelle on peut définir localement un élément de surface $\vec{dS} = dS \hat{n}$ dont le vecteur normal est orienté vers l'extérieur (par convention). On peut maintenant appliquer le théorème de Green-Ostrogradsky sur le volume V à



l'intérieur de cette surface :

$$\oiint_S \vec{B} \cdot \vec{dS} = \iiint_V \operatorname{div} \vec{B} dV = 0 .$$

où dans le dernier membre de droite nous avons utilisé le fait que $\operatorname{div} \vec{B} = 0$.

Ceci prend donc la forme intégrale d'une loi générale appelé la conservation du flux magnétique qui dit que le flux du champ magnétique, Φ_m , à travers une surface fermée est nul, c.-à-d. :

$$\boxed{\Phi_m = \oiint_S \vec{B} \cdot \vec{dS} = 0} . \quad (7.7)$$

Bien que nous avons obtenu cette loi dans le contexte du magnétostatique, on peut faire l'hypothèse, qui s'avère être vraie, que le champ magnétique reste conservé même si les champs et les courants varient avec le temps. Les équations (7.6) et (7.7) expriment donc deux formes d'une même loi fondamentale de la physique « la conservation du flux magnétique » .

La conservation du flux magnétique est une propriété très importante et montre une différence fondamentale entre le champ magnétique et le champ électrostatique. Nous avons vu, avec le théorème de Gauss, que le flux du champ électrostatique dépend des charges électriques contenues à l'intérieur de la surface :

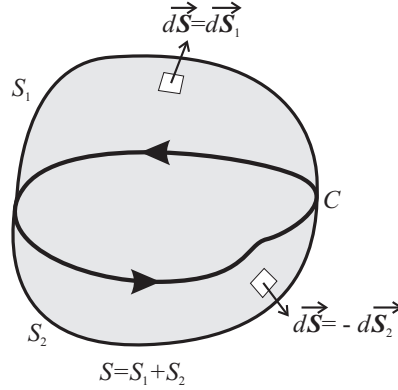
$$\Phi_e = \oiint_S \vec{E} \cdot \vec{dS} = \frac{Q_{\text{int}}}{\epsilon_0} .$$

Si la charge totale est positive, le flux est positif et il « sort » de cette surface un champ électrostatique (source). Si la charge est négative, le flux est négatif et le champ « rentre », converge vers la surface (puits). Cette propriété reste d'ailleurs également valable en régime variable. Rien de tel n'a jamais été observé pour le champ magnétique. On ne connaît pas de charge magnétique analogue à la charge électrique (ce serait un « monopôle magnétique ») : donc tout le champ qui rentre dans une surface fermée doit également en ressortir. La source la plus élémentaire de champ magnétique est un dipôle (deux polarités), comme l'aimant dont on ne peut dissocier le pôle nord du pôle sud.

A partir de l'équation (7.7), on montre que le flux à travers une surface ouverte s'appuyant sur un contour fermé C est indépendant du choix de cette surface. Prenons deux surfaces distinctes S_1 et S_2 s'appuyant sur un même contour C . La réunion de ces surfaces, $S = S_1 + S_2$, forme une surface fermée. En orientant la surface S de l'intérieur vers l'extérieur, la conservation du flux magnétique impose

$$\Phi_S = \Phi_{S_1} - \Phi_{S_2} = 0 ,$$

donc $\Phi_{S_1} = \Phi_{S_2}$, c'est-à-dire le flux magnétique à travers n'importe quelle surface s'appuyant sur le même contour C (et utilisant la même convention de normale) est indépendant du choix de la surface.



7.1.2 Lignes de champ et tubes de flux

Le concept de lignes de champ (également appelées lignes de force) est très utile pour se faire une représentation spatiale d'un champ de vecteurs. Ce sont ces lignes de champ qui sont tracées par la matière sensible au champ magnétique, telle que la limaille de fer au voisinage d'un aimant.

Définition 3 Une ligne de champ d'un champ de vecteur quelconque est une courbe C dans l'espace telle qu'en chacun de ses points le vecteur $\vec{y}$ soit tangent.

Considérons un déplacement élémentaire $\vec{d\ell}$ le long d'une ligne de champ magnétique C . Le fait que le champ magnétique $\vec{B}$ soit en tout point de C parallèle à $\vec{d\ell}$ s'écrit :

$$\vec{B} \wedge \vec{d\ell} = \vec{0}.$$

En coordonnées cartésiennes $\vec{d\ell} = dx\hat{x} + dy\hat{y} + dz\hat{z}$ et les lignes de champ sont calculées en résolvant :

$$\frac{dx}{B_x} = \frac{dy}{B_y} = \frac{dz}{B_z}.$$

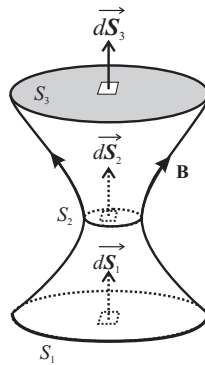
En coordonnées sphériques $\vec{d\ell} = dr\hat{r} + r d\theta\hat{\theta} + r \sin\theta d\phi\hat{\phi}$ et l'équation des lignes devient

$$\frac{dr}{B_r} = \frac{r d\theta}{B_\theta} = \frac{r \sin\theta d\phi}{B_\phi}. \quad (7.8)$$

La conservation du flux magnétique implique que les lignes de champ magnétique se referment sur elles-mêmes.

Un tube de flux est une sorte de «rassemblement» de lignes de champ. Soit une surface S_1 s'appuyant sur une courbe fermée C telle que le champ magnétique y soit tangent (c'est-à-dire $\vec{B} \perp \vec{d\ell}$ où $\vec{d\ell}$ est un vecteur élémentaire de C). En chaque point de C passe donc une ligne de champ particulière. En prolongeant ces lignes de champ on construit ainsi un tube de flux.

Tout au long de ce tube, le flux magnétique est conservé. En effet, considérons une portion de tube cylindrique entre S_1 et S_3 , ayant un rétrécissement en une surface S_2 . La surface $S = S_1 + S_3 + S_L$, où S_L est la surface latérale du tube, constitue une surface fermée. Donc le flux à travers S est nul. Par ailleurs, le flux à travers la surface latérale est également nul, par définition des lignes de champ ($\vec{B} \cdot \vec{d\vec{S}} = 0$ sur S_L). Donc, le flux en S_1 est le même qu'en S_3 . On peut faire le même raisonnement



pour S_2 . Cependant puisque $S_1 > S_2$ pour un flux identique, cela signifie que le champ magnétique est plus concentré en S_2 . D'une manière générale, plus les lignes de champ sont rapprochées et plus le champ magnétique est localement élevé. Les exemples les plus célèbres de tubes de flux rencontrés dans la nature sont les taches solaires.

7.2 Circulation du champ magnétique

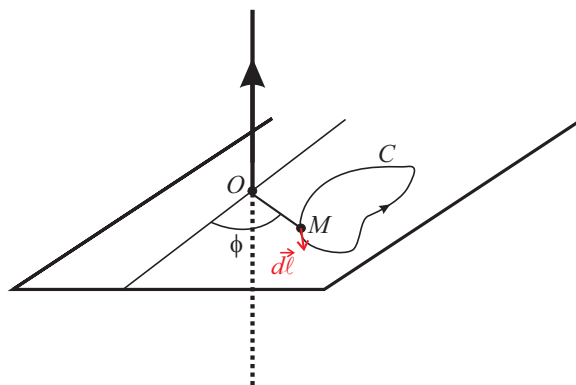
7.2.1 Circulation du champ autour d'un fil infini

Nous avons vu dans le chapitre précédent que la loi Biot et Savart prédit que le champ $\vec{B}$ créé par un fil infini en un point $M(r, \phi, z)$ s'écrit en coordonnées cylindriques :

$$\vec{B} = \frac{\mu_0 I}{2\pi\rho} \hat{\phi}.$$

Considérons maintenant une courbe fermée quelconque, C . Un déplacement élémentaire le long de cette courbe s'écrit $d\vec{\ell} = d\rho\hat{\rho} + \rho d\phi\hat{\phi} + dz\hat{z}$. La circulation de $\vec{B}$ sur la courbe fermée C vaut alors :

$$\oint_C \vec{B} \cdot d\vec{\ell} = \mu_0 I \oint \frac{d\phi}{2\pi}.$$



Plusieurs cas de figures peuvent se présenter :

- Si C n'enlace pas le fil, $\oint d\phi = 0$
- Si C enlace le fil, $\oint d\phi = 2\pi$
- Si C enlace le fil N fois, $\oint d\phi = N2\pi$

La circulation de $\vec{B}$ sur une courbe fermée est donc directement liée au courant qui traverse la surface délimitée par cette courbe. C'est Ampère qui, en recherchant une explication du magnétisme

dans une théorie de la dynamique des courants, découvrit cette propriété du champ magnétique. Elle est démontrée ici sur un cas particulier à partir de la loi de Biot et Savart mais elle traduit une loi magnétostatique fondamentale connu sous le nom de *théorème d'Ampère*.

7.3 Le théorème d'Ampère

Théorème 2 *La circulation de $\vec{B}$ le long d'une courbe C quelconque, orientée et fermée, appelée contour d'Ampère, est égale à μ_0 fois la somme algébrique des courants enlacés par le contour (c.-à-d. le flux du courant traversent une surface ouverte délimitée par C)*

$$\oint_C \vec{B} \cdot d\vec{\ell} = \mu_0 I_{\text{enl}} . \quad (7.9)$$

Cette relation fondamentale est l'équivalent du théorème de Gauss pour le champ électrostatique : elle relie le champ ($\vec{B}$ ou $\vec{E}$) à ses sources (le courant I ou la charge Q) *dans le vide* (à l'intérieur d'un matériau il faut les corriger). Cependant, à la différence du théorème de Gauss, elle n'est valable qu'en régime permanent (courants continus).

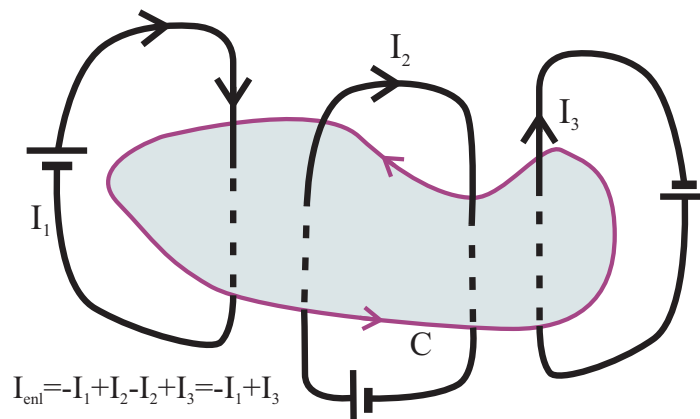


FIGURE 7.1 – Courants enlacés par un contour, C , dans le contexte du théorème d'Ampère

Remarques :

- Le théorème d'Ampère et la loi de Biot et Savart ont la même cause originelle.
- Le choix du sens de la circulation sur le contour d'Ampère choisi est purement arbitraire. Une fois ce choix fait, la règle du bonhomme d'Ampère permet d'attribuer un signe aux courants qui traversent la surface ainsi délimitée.
- Comme pour le théorème de Gauss, ce qui compte c'est la somme algébrique des sources : par exemple, si deux courants de même amplitude mais de sens différents traversent la surface, le courant total sera nul (voir figure ci-dessus).
- Bien qu'il exprime une loi fondamentale, le théorème de d'Ampère n'est exacte qu'en présence de courants stationnaires. C'est Maxwell qui a été le premier à comprendre qu'il fallait modifier cette loi en présence de champs et courants non-stationnaires (c.-à-d. qui varient dans le temps).

Comme pour la forme intégrale du théorème de Gauss, le théorème d'Ampère est une forme intégrale d'une loi fondamentale. On peut dériver la forme différentielle de cette loi en laissant le contour dans l'éq. (7.9) devenir le contour autour d'une différentielle de surface, $d\vec{S}$. Dans la limite infinitesimale de $dS \rightarrow 0$, un calcul similaire à celui que nous avons fait pour le théorème de Gauss dans la section l'éq. (3.2.2) nous montrerait que l'éq. (7.9) prendrait alors la forme :

$$\text{rot } \vec{B} \cdot d\vec{S} = \mu_0 \vec{j} \cdot d\vec{S} .$$

Puisque cette relation doit tenir pour n'importe quelle position ou orientation de la surface $\vec{dS}$, on obtient donc une forme différentielle du théorème d'Ampère :

$$\boxed{\vec{\text{rot}} \vec{B} = \mu_0 \vec{j}} . \quad (7.10)$$

De nos jours, on préfère souvent énoncer la forme « différentielle » de théorème d'Ampère de l'éq. (7.10) et effectuer le passage vers sa forme « intégrale » de l'éq. (7.9) en se servant du théorème de Stokes :

Théorème 3 « Théorème de Stokes » : La circulation d'un champ vectoriel $\vec{A}$ le long d'un contour C quelconque est égal au flux du rotationnel de $\vec{A}$ à travers toute surface s'appuyant sur C :

$$\boxed{\oint_C \vec{A} \cdot d\vec{\ell} = \iint_S \vec{\text{rot}} \vec{A} \cdot d\vec{S}} . \quad (7.11)$$

En appliquant ce théorème au champ $\vec{B}$ et en se servant de l'équation (7.10) on obtient donc la forme intégrale du théorème d'Ampère :

$$\oint_C \vec{B} \cdot d\vec{\ell} = \iint_S (\vec{\text{rot}} \vec{B}) \cdot d\vec{S} = \mu_0 \iint_S \vec{j} \cdot d\vec{S} \equiv \mu_0 I_{\text{enl}} .$$

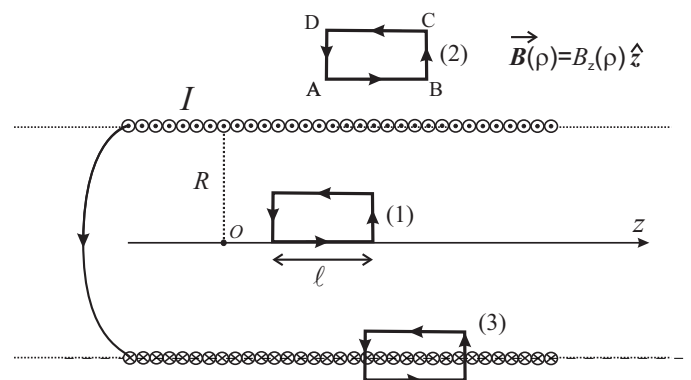
Exemple d'utilisation : le solénoïde infini

Considérons un solénoïde infini, comportant n spires par unité de longueur, chacune parcourue par un courant I permanent. Etant donné la géométrie cylindrique du solénoïde, on se place en coordonnées cylindriques, l'axe z étant l'axe du solénoïde. La densité de courant est toroïdal et s'écrit $\vec{j}(\rho, \phi, z) = j_\phi(\rho) \hat{\phi}$ puisqu'il y a invariance par rotation autour de l'axe z et translation le long de ce même axe. Donc, le champ magnétique est poloïdal et s'écrit :

$$\vec{B}(\rho, \phi, z) = B_z(\rho) \hat{z} .$$

(Il n'y a pas de champ dans la direction $\hat{\rho}$ à cause du fait que $z = \text{Cste}$ est un plan de symétrie pour un solénoïde *infini*.)

On choisit trois contours d'Ampère, chacun en forme de rectangle de longueur arbitraire ℓ et aux arrêts (A,B,C,D) (voir figure) :



— Contour (1) :

$$\oint_{(1)} \vec{B} \cdot d\vec{\ell} = \int_{AB} \vec{B} \cdot d\vec{\ell} + \int_{BC} \vec{B} \cdot d\vec{\ell} + \int_{CD} \vec{B} \cdot d\vec{\ell} + \int_{DA} \vec{B} \cdot d\vec{\ell} = 0 ,$$

$$\Rightarrow \int_0^l \vec{B}(\rho_{AB}) \cdot \hat{z} dz + \int_l^0 \vec{B}(\rho_{CD}) \cdot \hat{z} dz = 0 ,$$

$$\Rightarrow B_z(\rho_{AB}) \ell = B_z(\rho_{CD}) \ell .$$

Donc, le champ magnétique est uniforme à l'intérieur du solénoïde (infini) (pour $\rho < R$ $B_z(\rho) = \text{Cste} \equiv B_{\text{int}}$).

- Contour (2) : on obtient le même résultat, c'est-à-dire un champ uniforme à l'extérieur. Mais comme ce champ doit être nul à l'infini, on en déduit qu'il est nul partout (pour $\rho > R$, $B_z(\rho) = B_{\text{ext}} = 0$).
- Contour (3) :

$$\oint_{(3)} \vec{B} \cdot d\vec{\ell} = \int_{AB} \vec{B} \cdot d\vec{\ell} + \int_{BC} \vec{B} \cdot d\vec{\ell} + \int_{CD} \vec{B} \cdot d\vec{\ell} + \int_{DA} \vec{B} \cdot d\vec{\ell} = -n\ell\mu_0 I$$

$$\Rightarrow - \int_0^l \vec{B}(\rho_{CD}) \cdot \hat{z} dz = -n\ell\mu_0 I$$

$$\Rightarrow B_z(\rho_{CD}) = B = \mu_0 n I .$$

En résumé : le champ magnétique est faible l'extérieur d'un solénoïde et le champ à l'intérieur (orienté le long de l'axe du solénoïde) est approximativement uniforme avec :

$$\vec{B}_{\text{int}} = \mu_0 n I \hat{z} , \quad (7.12)$$

où n est le nombre de spires par unité de longueur et I est le courant dans chaque spire.

7.3.1 Relations de continuité du champ magnétique

Puisque le courant est la source du champ magnétique, on peut se demander ce qui se passe à la traversée d'une nappe de courant. Comme pour le champ électrostatique, va-t-on voir une discontinuité dans le champ ?

Soit une distribution surfacique de courant $\vec{j}_s$ séparant l'espace en deux régions 1 et 2. Considérons

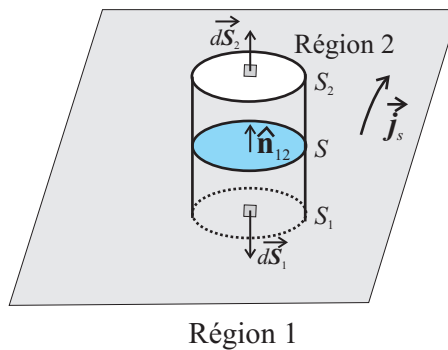


FIGURE 7.2 – Nappe de courant surfacique : continuité de la composante normale de $\vec{B}$

une surface fictive fermée et infinitésimale (illustrée en figure 7.2), traversant la nappe de courant. La conservation du flux magnétique à travers cette surface s'écrit :

$$\oiint_S \vec{B} \cdot d\vec{S} = \iint_{S_1} \vec{B} \cdot d\vec{S}_1 + \iint_{S_2} \vec{B} \cdot d\vec{S}_2 + \iint_{S_L} \vec{B} \cdot d\vec{S}_L = 0 ,$$

où S_L est la surface latérale. Lorsqu'on fait tendre le volume contenu par cette surface à zéro (c.-à-d. S_1 tend vers S_2), on obtient

$$\iint_{S_2} \vec{B}_2 \cdot d\vec{S}_2 + \iint_{S_1} \vec{B}_1 \cdot d\vec{S}_1 = 0 \Rightarrow \iint_S (\vec{B}_2 - \vec{B}_1) \cdot \hat{n}_{12} dS = 0 ,$$

puisque $d\vec{S}_2 = -d\vec{S}_1 = dS\hat{n}_{12}$ dans cette limite. Ce résultat étant valable quelque soit la surface S choisie, on vient donc de démontrer que :

$$\boxed{\hat{n}_{12} \cdot (\vec{B}_2 - \vec{B}_1) = 0} . \quad (7.13)$$

Pour la composante tangentielle, nous allons utiliser le théorème d'Ampère. Considérons le contour d'Ampère infinitésimal illustré en figure 7.3 :

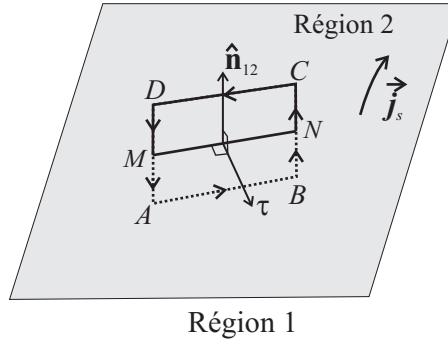


FIGURE 7.3 – Nappe de courant surfacique : discontinuité de la composante tangentielle de $\mathbf{B}$

Le théorème d'Ampère pour ce contour s'écrit alors :

$$\oint \vec{B} \cdot d\vec{\ell} = \int_{AB} \vec{B} \cdot d\vec{\ell} + \int_{BC} \vec{B} \cdot d\vec{\ell} + \int_{CD} \vec{B} \cdot d\vec{\ell} + \int_{DA} \vec{B} \cdot d\vec{\ell} = \mu_0 I_{\text{enl}} .$$

Le courant I_{enl} est celui qui circule sur la nappe, autrement dit, il est défini par la densité de courant surfacique :

$$I_{\text{enl}} = \iint_{ABCD} \vec{j} \cdot d\vec{S} = \int_{MN} (\vec{j}_s \cdot \hat{\tau}) d\ell ,$$

où $(\overrightarrow{MN}, \hat{n}_{12}, \hat{\tau})$ est un trièdre direct. Dans la limite $DA \rightarrow 0$, le théorème d'Ampère fournit :

$$\int_{MN} (\vec{B}_1 - \vec{B}_2) \cdot d\vec{\ell} = \mu_0 \int_{MN} (\vec{j}_s \cdot \hat{\tau}) d\ell .$$

Puisque MN est quelconque (sur la surface), on doit avoir :

$$(\vec{B}_1 - \vec{B}_2) \cdot d\vec{\ell} = \mu_0 \vec{j}_s \cdot \hat{\tau} d\ell ,$$

mais

$$\begin{aligned} (\vec{B}_1 - \vec{B}_2) \cdot d\vec{\ell} &= (\vec{B}_1 - \vec{B}_2) \cdot (\hat{\tau} \wedge -\hat{n}_{12}) d\ell \\ &= [(\vec{B}_1 - \vec{B}_2) \wedge \hat{n}_{12}] \cdot \hat{\tau} d\ell , \end{aligned}$$

c'est-à-dire (puisque la direction de $\hat{\tau}$ est arbitraire)

$$\boxed{\hat{n}_{12} \wedge (\vec{B}_2 - \vec{B}_1) = \mu_0 \vec{j}_s} . \quad (7.14)$$

En résumé, à la traversée d'une nappe de courant,

- la composante normale du champ magnétique reste continue,
- la composante tangentielle du champ magnétique est discontinue.

7.4 « Potentiel vecteur »

Une conséquence mathématique de la loi $\text{div } \vec{B} = 0$ (cf. l'éq. (7.6)), est que l'on peut toujours définir un champ vectoriel $\vec{A}$ tel que $\vec{B} = \text{rot } \vec{A}$. La démonstration se repose sur l'éq. (11.18) de l'appendice sur les mathématiques. On appelle $\vec{A}$ le « potentiel vecteur » même s'il n'a pas les propriétés d'un potentiel.

De plus est, le champ $\vec{A}$ n'est pas bien définie puisqu'on peut toujours ajouter le gradient d'un champ scalaire f à $\vec{A}$ sans changer sa rotationnelle

$$\begin{aligned}\vec{A}' &= \vec{A} + \text{grad } f \\ \text{rot } \vec{A}' &= \text{rot } \vec{A} + \text{rot } \text{grad } f = \text{rot } \vec{A} = \vec{B},\end{aligned}$$

puisque $\text{rot } \text{grad } f = \vec{0}$ pour n'importe quelle fonction scalaire, f comme démontré dans l'éq. (11.17) de l'Annexe des mathématiques.

Insérant $\vec{B} = \text{rot } \vec{A}$ dans $\text{rot } \vec{B} = \mu_0 \vec{j}$, on obtient une équation différentielle pour $\vec{A}$:

$$\text{rot } \text{rot } \vec{A} \equiv \text{grad } \text{div } \vec{A} - \Delta \vec{A} = \mu_0 \vec{j}, \quad (7.15)$$

où nous avons utilisé encore une autre identité mathématique (voir la démonstration (11.20) de l'Annexe des mathématiques) $\text{rot } \text{rot} \equiv \text{grad } \text{div} - \Delta$. On peut enlever une partie de la liberté dans la définition de $\vec{A}$ en imposant la contrainte de la « gauge de Coulomb », c.-à.-d. on impose la condition :

$$\text{div } \vec{A} = 0. \quad (7.16)$$

Ainsi l'équation (7.15) dans cette gauge devient :

$$\Delta \vec{A} = -\mu_0 \vec{j}. \quad (7.17)$$

On remarque qu'en dehors de sa nature vectoriel cette équation à la même forme que l'équation de Poisson ($\Delta V = -\rho/\epsilon_0$) que nous avons déjà rencontré dans l'électrostatique.

La solution pour $\vec{A}$ de l'éq. (7.17) se trouve par analogie directe avec celle de V en électrostatique :

$$\boxed{\vec{A}(M) = \frac{\mu_0}{4\pi} \iiint \frac{\vec{j}(P) dV}{|\vec{PM}|}}. \quad (7.18)$$

On obtient la formule pour $\vec{A}$ produit par un circuit filiforme, en suivant la même logique qui nous a permis de passer de l'éq. (6.21) pour le champ $\vec{B}$ en termes de $\vec{j}$ à la loi de Biot et Savart (l'éq. (6.22)) dans le chapitre 6. On obtient ainsi :

$$\vec{A}(M) = \frac{\mu_0 I}{4\pi} \oint_{\text{circuit}} \frac{d\vec{\ell}_P}{|\vec{PM}|}. \quad (7.19)$$

Même si les équations (7.18) et (7.19) sont analogues à l'expression intégrale d'un potentiel scalaire (d'où le nom potentiel vecteur), elles sont moins pratiques à l'utilisation, puisque qu'il s'agit d'effectuer des intégrales d'une quantité vectorielle.

Les physiciens se méfiaient au départ du potentiel vecteur à cause de la liberté du choix de gauge et du fait qu'il ne s'agit pas d'une quantité directement mesurable (en contraste avec les champs $\vec{E}$, $\vec{B}$, et V). Ce point de vue était radicalement modifié à l'arrivée de la mécanique quantique dans laquelle la force joue un rôle moins important que l'énergie et par conséquent les champs $\vec{A}$ et V se sont trouvés dans le rôle de vedette tandis que les champs $\vec{E}$ et $\vec{B}$ ont été relégués à des rôles secondaires.

7.5 Quatre façons de calculer le champ magnétique

En guise de résumé voici des conseils sur les méthodes à employer pour calculer le champ magnétique.

- **La formule de Biot et Savart** : elle n'est pratique que lorsqu'on sait calculer l'addition vectorielle des champs $d\vec{B}$ créés par tous les éléments du circuit (souvent des circuits filiformes).
- **Le théorème d'Ampère** : il faut être capable de calculer la circulation du champ sur un contour choisi. Cela nécessite donc une symétrie relativement simple des courants.
- **La conservation du flux** : à n'utiliser que si l'on connaît déjà son expression dans une autre région de l'espace.
- **Le potentiel vecteur** : On calcule le potentiel vecteur $\vec{A}$ par une méthode qui ressemble à celle du calcul du potentiel scalaire en électrostatique. Néanmoins, il faut calculer ses trois composantes dans une région donnée et ensuite pouvoir calculer $\text{rot}\vec{A}$ afin d'obtenir $\vec{B}$.

Dans tous les cas, il faut prendre en compte les propriétés de symétrie de la densité de courant.

7.6 Le dipôle magnétique

7.6.1 Champ magnétique créé par une spire

Soit une spire plane, de forme quelconque, de centre d'inertie O , parcourue par un courant permanent I . Nous allons calculer le champ magnétique créé par cette spire en tout point M de l'espace, situé à grande distance de la spire (précisément, à des distances grandes comparées à la taille de la spire). On pose (cf. figure ci-dessous)

$$\vec{r} = \overrightarrow{OM} \quad \vec{r}' = \overrightarrow{PM} \quad \vec{\rho} = \overrightarrow{OP} = \vec{r} - \vec{r}' \quad \hat{r} = \frac{\vec{r}}{r}.$$

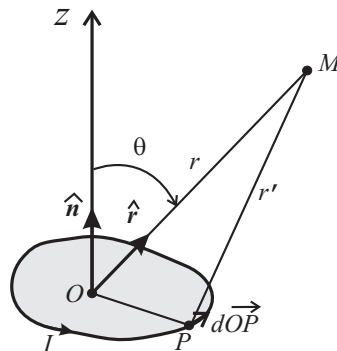


FIGURE 7.4 – Dipôle magnétique : champ loin d'une spire de courant

Bien qu'on pourrait trouver le champ magnétique créé par le circuit directement à partir de la formule de Biot et Savart, il est instructif d'arriver à ce résultat en employant le potentiel vecteur $\vec{A}$. Le potentiel vecteur associé avec la spire de courant s'écrit :

$$\vec{A}(M) = \frac{\mu_0 I}{4\pi} \oint_{\text{spire}} \frac{d\vec{\ell}}{PM} = \frac{\mu_0 I}{4\pi} \oint_{\text{spire}} \frac{d\vec{\rho}}{r'}$$

avec $r' = PM = \left| \overrightarrow{OM} - \overrightarrow{OP} \right| = \left(r^2 + \rho^2 - 2\overrightarrow{OP} \cdot \overrightarrow{OM} \right)^{1/2}.$

où nous avons utilisé le fait que $\vec{d\ell} \equiv d\vec{OP} = d\vec{\rho}$ dans le plan de la spire. On se rappelle qu'on emploie la limite $r \gg \rho$, pour tout point P appartenant à la spire, donc :

$$\frac{1}{r'} = \frac{1}{r} \left(1 + \frac{\rho^2}{r^2} - 2 \frac{\vec{OP} \cdot \vec{r}}{r^2} \right)^{-1/2} = \frac{1}{r} + \frac{\vec{r} \cdot \vec{\rho}}{r^3} + O\left(\frac{\rho^2}{r^2}\right).$$

et l'on obtient :

$$\vec{A}(M) = \vec{A}(\vec{r}) = \frac{\mu_0 I}{4\pi} \oint_{\text{spire}} \frac{d\vec{\rho}}{r'} \simeq \frac{\mu_0 I}{4\pi r} \oint_{\text{spire}} d\vec{\rho} + \frac{\mu_0 I}{4\pi r^3} \oint_{\text{spire}} d\vec{\rho} (\vec{r} \cdot \vec{\rho}). \quad (7.20)$$

La première intégrale de l'éq.(7.20) se fait rapidement. Si on décompose le vecteur $\vec{\rho}$ dans une base $(\hat{e}_1, \hat{e}_2)$ engendrant le plan de la spire, on a :

$$\oint_{\text{spire}} d\vec{\rho} = \vec{0} = \hat{e}_1 \oint_{\text{spire}} d\rho_1 + \hat{e}_2 \oint_{\text{spire}} d\rho_2 = \vec{0}.$$

puisque chaque intégrale est nulle :

$$\oint_{\text{spire}} d\rho_1 = [\rho_1(P_0) - \rho_1(P_0)] = 0 = \oint_{\text{spire}} d\rho_2.$$

Il faut maintenant évaluer la deuxième intégrale de l'éq.(7.20). Si on décompose les vecteurs $\vec{\rho}$ et $\vec{r}$ dans une base $(\hat{e}_1, \hat{e}_2)$ engendrant le plan de la spire, on a :

$$\oint_{\text{spire}} d\vec{\rho} (\vec{\rho} \cdot \vec{r}) = \oint_{\text{spire}} \{d\rho_1 (\rho_1 r_1 + \rho_2 r_2) \hat{e}_1 + d\rho_2 (\rho_1 r_1 + \rho_2 r_2) \hat{e}_2\} = 0. \quad (7.21)$$

On remarque d'abord que :

$$\begin{aligned} \oint_{\text{spire}} d(\vec{\rho} \cdot \vec{\rho}) &= [\rho^2]_{P_0}^{P_0} = 0 \\ \Rightarrow \oint_{\text{spire}} d(\rho_1^2 + \rho_2^2) &= 2 \oint_{\text{spire}} (\rho_1 d\rho_1 + \rho_2 d\rho_2) = 0, \end{aligned} \quad (7.22)$$

puisqu'on revient au même point départ de P_0 . De la même manière :

$$\oint_{\text{spire}} d(\rho_1 \rho_2) = \oint_{\text{spire}} \rho_1 d\rho_2 + \oint_{\text{spire}} \rho_2 d\rho_1 = [\rho_1 \rho_2]_{P_0}^{P_0} = 0,$$

et on a donc l'égalité :

$$\oint_{\text{spire}} \rho_2 d\rho_1 = - \oint_{\text{spire}} \rho_1 d\rho_2. \quad (7.23)$$

Utilisant les relations (7.22) et (7.23) dans l'éq.(7.21) donne :

$$\begin{aligned} \oint_{\text{spire}} d\vec{\rho} (\vec{\rho} \cdot \vec{r}) &= r_2 \oint_{\text{spire}} \rho_2 d\rho_1 \hat{e}_1 + r_1 \oint_{\text{spire}} \rho_1 d\rho_2 \hat{e}_2 \\ &= (-r_2 \hat{e}_1 + r_1 \hat{e}_2) \frac{1}{2} \oint_{\text{spire}} (\rho_1 d\rho_2 - \rho_2 d\rho_1) \\ &= \frac{\hat{n} \wedge \vec{r}}{2} \oint_{\text{spire}} (\rho_1 d\rho_2 - \rho_2 d\rho_1) \\ &= S \hat{n} \wedge \vec{r}, \end{aligned} \quad (7.24)$$

où $\hat{\mathbf{n}} = \hat{\mathbf{e}}_3$ est le vecteur normal au plan de la spire (vecteur de base de l'axe 3) et S sa surface. Dans la dernière ligne de l'éq. (7.24) nous avons utilisé l'intégrale :

$$\frac{1}{2} \oint_{\text{spire}} \vec{\rho} \wedge d\vec{\rho} = \hat{\mathbf{n}} \left[\frac{1}{2} \oint_{\text{spire}} (\rho_1 d\rho_2 - \rho_2 d\rho_1) \right] = S\hat{\mathbf{n}}. \quad (7.25)$$

Ce calcul est général, valable quelle que soit la surface. En effet, une surface élémentaire dS , telle que :

$$\boxed{\frac{1}{2} \vec{\rho} \wedge d\vec{\rho} = dS\hat{\mathbf{n}}},$$

est toujours engendrée lors d'un petit déplacement du vecteur $\vec{\rho}$ (voir la figure 7.5)

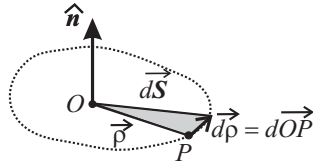


FIGURE 7.5 – Surface élémentaire d'une spire

On obtient donc une expression relativement simple pour le potentiel $\vec{A}$ loin de la spire :

$$\vec{A}(\vec{r}) \simeq \frac{\mu_0 I}{4\pi r^2} S\hat{\mathbf{n}} \wedge \hat{\mathbf{r}} = \frac{\mu_0}{4\pi} \left(\frac{\vec{\mathbf{m}} \wedge \hat{\mathbf{r}}}{r^2} \right), \quad (7.26)$$

où on voit apparaître une grandeur importante car décrivant complètement la spire « vue » depuis une grande distance, à savoir le **moment magnétique dipolaire** :

$$\boxed{\vec{\mathbf{m}} = IS\hat{\mathbf{n}}}. \quad (7.27)$$

Le champ magnétique se déduit de $\vec{B} = \text{rot } \vec{A}$:

$$\vec{B}(\vec{r}) = \frac{\mu_0}{4\pi} \text{rot} \left(\frac{\vec{\mathbf{m}} \wedge \hat{\mathbf{r}}}{r^2} \right) = \frac{\mu_0}{4\pi} \text{rot} \left(\frac{\vec{\mathbf{m}} \wedge \vec{r}}{r^3} \right).$$

On invoque maintenant une relation d'analyse vectorielle (cf. l'éq. (11.21) de l'annexe des mathématiques) :

$$\text{rot} (f \vec{V}) = f \text{rot } \vec{V} + \overrightarrow{\text{grad}} f \wedge \vec{V} \quad \text{avec} \quad \vec{V} = \vec{\mathbf{m}} \wedge \vec{r} \quad \text{et} \quad f = \frac{1}{r^3}.$$

Pour le deuxième terme du second membre ci-dessus, on se rappelle que $\overrightarrow{\text{grad}} (1/r^3) = -3\vec{r}/r^5$, alors que pour le premier terme du second membre, on utilise le fait que :

$$\begin{aligned} \text{rot} (\vec{\mathbf{m}} \wedge \vec{r}) &= \text{rot} \left(\begin{bmatrix} m_x \\ m_y \\ m_z \end{bmatrix} \wedge \begin{bmatrix} x \\ y \\ z \end{bmatrix} \right) = \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ z m_y - y m_z & x m_z - z m_x & y m_x - x m_y \end{vmatrix} \\ &= 2m_x \hat{x} + 2m_y \hat{y} + 2m_z \hat{z} = 2\vec{\mathbf{m}}, \end{aligned}$$

où nous avons utilisé la formulation « déterminant » afin de calculer le rotationnel. Mettant tout ceci dans l'Eq. (7.6.1), on obtient le champ, $\vec{B}$, du dipôle magnétique :

$$\vec{B}(\vec{r}) \simeq \frac{\mu_0}{4\pi r^3} [2\vec{\mathbf{m}} - 3\hat{\mathbf{r}} \wedge (\vec{\mathbf{m}} \wedge \hat{\mathbf{r}})]. \quad (7.28)$$

On obtient une expression équivalente (et un peu plus simple) pour $\vec{B}$ en faisant appel à l'égalité $\hat{r} \wedge (\vec{m} \wedge \hat{r}) = \vec{m} (\hat{r} \cdot \hat{r}) - \hat{r} (\hat{r} \cdot \vec{m})$:

$$\vec{B}(\vec{r}) = \frac{\mu_0}{4\pi r^3} [3\hat{r}(\vec{m} \cdot \hat{r}) - \vec{m}] , \quad (7.29)$$

ce qui peut également s'écrire :

$$\vec{B}(\vec{r}) = -\frac{\mu_0}{4\pi} \overrightarrow{\text{grad}} \left[\frac{\vec{m} \cdot \hat{r}}{r^2} \right] . \quad (7.30)$$

En coordonnées sphériques, $\hat{r} \cdot \vec{m} = m \cos \theta$ et les composantes poloïdales du champ s'écrivent :

$$B_r = \frac{\mu_0}{4\pi r^3} 2m \cos \theta \quad B_\theta = \frac{\mu_0}{4\pi r^3} m \sin \theta .$$

Remarque : On constate que le champ $\vec{B}$ d'un dipôle magnétique est analogue au champ $\vec{E}$ produit par un dipôle électrique :

$$\vec{E}(\vec{r}) \rightarrow -\frac{1}{4\pi\epsilon_0} \overrightarrow{\text{grad}} \left[\frac{\vec{p} \cdot \hat{r}}{r^2} \right] .$$

(voir chapitre 5 et l'éq. (5.5) en particulier)

Cette analogie est assez surprenante compte tenu du fait que les équations et les sources de $\vec{B}$ sont bien différentes que celles de $\vec{E}$.

Lignes de champ d'un dipôle magnétique

Pour le dipôle magnétique « ponctuelle » les équations pour les lignes de champ, $\vec{B}$, sont entièrement analogues à celles du champ $\vec{E}$ pour le dipôle électrique. Ici, l'équation des lignes de champ en coordonnées sphériques fournit (voir (7.8)) :

$$\vec{d\ell} \wedge \vec{B} = \vec{0} \Rightarrow d\phi = 0 \quad \text{et} \quad \frac{dr}{\frac{\mu_0}{4\pi r^3} 2m \cos \theta} = \frac{r d\theta}{\frac{\mu_0}{4\pi r^3} 2m \sin \theta} \Rightarrow \frac{dr}{r} = \frac{\cos \theta d\theta}{\sin \theta} ,$$

donc les lignes de champ magnétique sont obtenues en intégrant les deux côtés :

$$\int \frac{dr}{r} = \int \frac{\cos \theta d\theta}{\sin \theta} \Rightarrow \ln r = 2 \ln (\sin \theta) + C \Rightarrow r(\theta) = K \sin^2 \theta ,$$

où K est une constante.

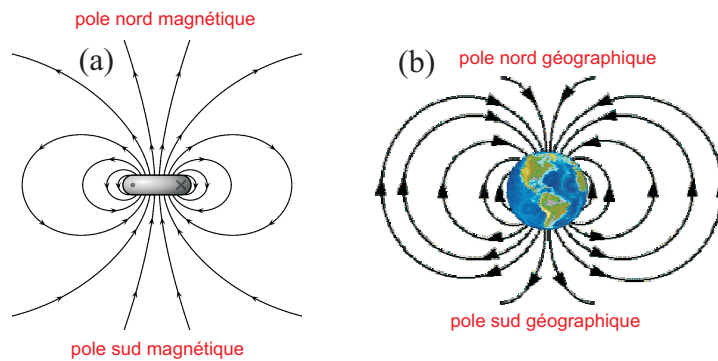


FIGURE 7.6 – Lignes de champ magnétiques produites par des dipôles magnétiques. Exemples : (a) une spire de courant, (b) le champ magnétique terrestre.

Chapitre 8

Champ magnétique en présence de la matière

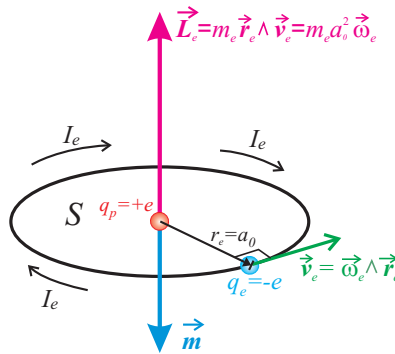
8.1 Le modèle du dipôle en physique

Comme nous avons remarqué dans le chapitre précédent, l'expression du champ magnétique créé par une spire de courant (dipôle magnétique $\vec{m} = IS\hat{n}$) est formellement équivalente à celle du champ électrostatique créé par un système de deux charges opposées (dipôle électrique $\vec{p} = q\vec{d}$)

$$\vec{B}(\vec{r}) \rightarrow -\frac{\mu_0}{4\pi} \vec{\text{grad}} \left[\frac{\vec{m} \cdot \hat{r}}{r^2} \right] \quad \vec{E}(\vec{r}) \rightarrow -\frac{1}{4\pi\epsilon_0} \vec{\text{grad}} \left[\frac{\vec{p} \cdot \hat{r}}{r^2} \right] .$$

Cependant, pour le champ magnétique, il s'avère impossible de séparer le dipôle en une charge magnétique « + » et une autre « - ». Le dipôle est la première source de champ magnétique. C'est la raison pour laquelle il joue un si grand rôle dans la modélisation des effets magnétiques observés dans la nature, au niveau microscopique comme macroscopique.

L'origine du champ magnétique d'un matériau quelconque (ex : aimant) doit être microscopique. En utilisant le modèle atomique de Bohr, on peut se convaincre que les atomes (du moins certains) ont un **moment magnétique dipolaire intrinsèque**. Le modèle de Bohr de l'atome d'Hydrogène consiste en un électron de charge $q_e = -e$ en mouvement circulaire uniforme autour d'un noyau central (un proton) avec une période $T_e = \frac{2\pi}{\omega}$.



Si on regarde sur des échelles de temps longues par rapport à T_e , tout se passe comme s'il y avait un courant

$$I_e = \frac{q_e}{T_e} = \frac{q_e \omega}{2\pi} .$$

On a donc une sorte de spire circulaire, de rayon moyen la distance moyenne au proton, c'est-à-dire le

rayon de Bohr, a_0 . L'atome d'Hydrogène aurait donc un moment magnétique intrinsèque :

$$\begin{aligned}\vec{m} &= I_e S \hat{n} = \frac{q_e \omega}{2} a_0^2 \hat{n} = \frac{q_e}{2m_e} (m_e \omega a_0^2 \hat{n}) \\ &= \frac{q_e}{2m_e} \vec{L}_e,\end{aligned}\tag{8.1}$$

où $\vec{L}_e$ est le moment cinétique de l'électron (voir l'éq. (9.8)) et le facteur $\gamma \equiv q/2m$ est appelé **le rapport gyromagnétique**.

Ce raisonnement peut se généraliser aux autres atomes. En effet, un ensemble de charges en rotation autour d'un axe produit un moment magnétique proportionnel au moment cinétique total. Cela se produit même si la charge totale est nulle (matériau ou atome neutre) : ce qui compte c'est l'existence d'un courant. Il suffit donc d'avoir un décalage, même léger, entre les vitesses des charges « + » et celles des charges « - ».

Du coup, on peut expliquer qualitativement les propriétés magnétiques des matériaux en fonction de l'orientation des moments magnétiques des atomes qui les composent :

- **Matériaux diamagnétiques** : produisent un moment magnétique induit, proportionnel au champ magnétique appliqué, qui s'oppose à ce dernier. Le champ $\vec{B}$ résultant est d'intensité inférieure au champ appliqué.
- **Matériaux paramagnétiques** : Leurs constituants ont des moments dipolaires magnétiques intrinsèques qui peuvent s'aligner avec un champ magnétique externe. Elles peuvent ainsi être aimantées momentanément et le $\vec{B}$ résultant est d'intensité supérieure au champ appliqué.
- **Matériaux ferromagnétiques** : ceux dont les moments sont déjà orientés dans une direction particulière, de façon permanente (aimants naturels).

La Terre est connue pour avoir un champ magnétique dipolaire, où le pôle Nord magnétique correspond approximativement au pôle Sud géographique. Au niveau macroscopique, l'explication de l'existence du champ magnétique observé sur les planètes et sur les étoiles est encore aujourd'hui loin d'être satisfaisante. La théorie de l'effet dynamo essaye de rendre compte des champs observés par la présence de courants, essentiellement azimutaux, dans le cœur des astres.

Plusieurs faits connus restent partiellement inexpliqués :

- **Les cycles magnétiques** : le Soleil a un champ magnétique à grande échelle qui ressemble à celui de la Terre, approximativement dipolaire. Cependant, il y a une inversion de polarité tous les 11 ans. Pour la Terre, on a pu mettre en évidence que le cycle est de l'ordre de 700.000 ans en moyenne et que la dernière inversion remonte à . . . ~ 700.000 ans!! Par ailleurs, on observe actuellement des fluctuations du champ terrestre assez importantes, mais de là à prédire exactement la prochaine inversion, le mystère reste entier.
- **Non-alignement avec le moment cinétique de l'astre** : s'il est de l'ordre d'une dizaine de degrés pour la Terre (avec une modification de la direction de l'axe magnétique d'environ 15° par an), il est de 90° pour celui de Neptune!

8.2 La magnétisation

On vient de voir dans la section précédente que les constituants atomiques de la matière peuvent agir comme de petites boucles de courants et donnent naissance à des champs magnétiques dipolaires. Pour les milieux diamagnétiques et paramagnétiques le moment de magnétique est proportionnelle

au champ magnétique incident sur l'atome. Donc, par analogie avec le traitement des diélectriques, on définit une densité volumique de moment dipolaire $\vec{M}$ (appelé magnétisation) tel que le moment magnétique d'un volume infinitésimal dV est donné par $d\vec{m} = \vec{M}dV$.

On définit également une **susceptibilité magnétique** χ_m qui donne la proportionnalité entre $\vec{M}$ et $\vec{B}$, caractéristique du matériau en question :

$$\vec{M} = \frac{\chi_m}{\mu_0} \vec{B}, \quad (8.2)$$

où χ_m est un nombre sans dimension qui est typiquement de l'ordre de $\chi_m \sim -10^{-5}$ pour les matériaux diamagnétiques et $\chi_m \sim 10^{-3}$ pour des matériaux paramagnétiques. A cause de la petite taille de χ_m on peut dans beaucoup de situations ignorer la présence de la matière sur les effets magnétiques et calculer le champ magnétique comme s'il s'agissait du vide.

Les matériaux ferromagnétiques font une grande exception à la règle ci-dessus. Pour ces matériaux, les moments dipolaires s'agissent entre eux fortement et tendent à tous s'aligner dans le même sens (au moins à l'intérieur d'un domaine cristallin).

La densité de courant, $\vec{j}_m$, (de nature atomique) associée avec l'existence de $\vec{M}$, se trouve avec la relation :

$$\vec{j}_m = \text{rot} \vec{M}. \quad (8.3)$$

L'équation d'Ampère s'écrit donc,

$$\text{rot} \vec{B} = \mu_0 (\vec{j}_m + \vec{j}_{\text{libre}}), \quad (8.4)$$

où $\vec{j}_{\text{libre}}$ sont des courants manipulés dans une expérience (usuellement dans des fils électriques)

8.3 Le champ « $\mathbf{H}$ »

Puisque nous n'avons pas de contrôle direct de $\vec{j}_m$ (et ses moments dipolaires magnétiques microscopiques associés), il s'avère pratique en présence de la matière de définir un champ auxiliaire $\vec{H}$ de façon analogue avec le champ auxiliaire $\vec{D}$ en électrostatique (voir le section 5.2). Le champ $\vec{H}$ regroupe le $\vec{M}$ avec le champ $\vec{B}$ de la façon suivante :

$$\vec{H} \equiv \frac{\vec{B}}{\mu_0} - \vec{M}. \quad (8.5)$$

On obtient l'équation différentielle de $\vec{H}$ en prenant la rotationnelle de l'éq. (8.5) et utilisant ensuite les équations (8.4) et (8.3). On obtient ainsi **L'équation différentielle de $\vec{H}$ en magnétostatique** :

$$\text{rot} \vec{H} = \vec{j}_{\text{libre}}. \quad (8.6)$$

Pour des matériaux diamagnétiques et paramagnétiques, $\vec{M}$ est proportionnelle à $\vec{B}$, et on peut écrire

$$\begin{aligned} \vec{H} &= \frac{\vec{B}}{\mu_0} - \frac{\chi_m \vec{B}}{\mu_0} = \frac{1}{\mu_0} (1 - \chi_m) \vec{B} \\ &\equiv \frac{\vec{B}}{\mu_r \mu_0} \end{aligned} \quad (8.7)$$

$$\Rightarrow \mu_r = \frac{1}{(1 - \chi_m)}.$$

On appelle $\mu_r = 1/(1 - \chi_m)$ la perméabilité magnétique relative du matériel. La relation entre $\vec{H}$ et $\vec{B}$ pour les milieux linéaires est :

$$\boxed{\vec{H} = \frac{\vec{B}}{\mu_0 \mu_r}} . \quad (8.8)$$

Si la symétrie du problème est suffisamment élevée, on peut obtenir $\vec{H}$ en faisant appel à la forme intégrale de l'éq. (8.6) :

$$\boxed{\oint_C \vec{H} \cdot d\vec{\ell} = I_{\text{enl, libre}}} . \quad (8.9)$$

Cette expression nous dicte les unités du champ $\vec{H}$ comme étant A.m^{-1} .

Remarque : Si on compare attentivement les définitions des champs auxiliaires $\vec{H}$ et $\vec{D}$, ainsi que les paramètres constitutifs associés, ε_r et μ_r , on remarquera quelques différences un peu troublantes (de signe etc.). Ces différences regrettables ne viennent pas de la physique elle-même mais plutôt d'un accident de parcours historique. Elles proviennent du fait qu'au début, les physiciens pensaient que $\vec{H}$ était le champ fondamental et $\vec{B}$ le champ auxiliaire. Ainsi, autrefois (et parfois encore), on appelait $\vec{H}$ le champ magnétique et $\vec{B}$ le champ « d'induction magnétique ». De nos jours, on préfère appeler le champ fondamentale $\vec{B}$ le **champ magnétique**, et le champ $\vec{H}$ simplement le champ « H ».

Chapitre 9

Actions et énergie magnétiques

9.1 Force magnétique sur une particule chargée

Ce qui a été dit aux chapitres précédents concerne plus particulièrement les aspects macroscopiques, par exemple, le champ magnétique mesurable créé par un circuit électrique. Or, le courant circulant dans un circuit est dû au déplacement de particules chargées. Nous prendrons donc le parti ici de poser l'expression de la force magnétique s'exerçant sur une particule (sans la démontrer) puis de montrer comment s'exprime cette force sur un circuit. Historiquement bien sûr, c'est la force de Laplace qui a été mise en évidence la première, la force de Lorentz n'est venue que bien plus tard...

9.1.1 La force de Lorentz

La force totale, électrique et magnétique (on dit électromagnétique) subie par une particule de charge q et de vitesse $\vec{v}$ mesurée dans un référentiel galiléen est *l'équation de Lorentz* :

$$\boxed{\vec{F} = q \left(\vec{E} + \vec{v} \wedge \vec{B} \right)} . \quad (9.1)$$

On appelle cette force **la force de Lorentz**. On peut la mettre sous la forme :

$$\vec{F} = \vec{F}_e + \vec{F}_m \quad \text{où} \quad \begin{cases} \vec{F}_e = q\vec{E} \\ \vec{F}_m = q\vec{v} \wedge \vec{B} \end{cases} ,$$

où $\vec{F}_e$ est la composante *électrique* et $\vec{F}_m$ la composante *magnétique*. La composante magnétique de la force de Lorentz (parfois appelée force magnétique) possède un ensemble de propriétés remarquables :

1. **La force magnétique ne fournit pas de travail.** Si on applique la relation fondamentale de la dynamique pour une particule de masse m et charge q , on obtient :

$$\vec{F}_m = q\vec{v} \wedge \vec{B} = m \frac{d\vec{v}}{dt} ,$$

donc

$$\frac{d}{dt} \left(\frac{1}{2} m v^2 \right) = \frac{d}{dt} \left(\frac{1}{2} m \vec{v} \cdot \vec{v} \right) = m \vec{v} \cdot \frac{d\vec{v}}{dt} = q \vec{v} \cdot (\vec{v} \wedge \vec{B}) = 0 .$$

L'énergie cinétique de la particule est donc bien conservée.

2. **La force magnétique est une correction en $(v/c)^2$ à la force de Coulomb**, où c est la vitesse de la lumière. Ceci se voit en regardant la force d'interaction entre deux particules chargées :

$$\vec{F}_{1 \rightarrow 2} \simeq \frac{q_1 q_2}{4\pi\epsilon_0 r_{12}^2} \left[\hat{u}_{12} + \frac{\vec{v}_2}{c} \wedge \left(\frac{\vec{v}_1}{c} \wedge \hat{u}_{12} \right) \right] . \quad (9.2)$$

3. **Violation du principe d'action et de réaction.** On peut aisément vérifier sur un cas particulier simple que la force magnétique ne satisfait pas le 3^{ème} principe de Newton. Pour cela, il suffit de prendre une particule 1 se dirigeant vers une particule 2. Le champ magnétique créé par 1 sera alors nul à l'emplacement de la particule 2 :

$$\vec{B}_1 = \frac{\mu_0 q_1}{4\pi r^2} \vec{v}_1 \wedge \hat{u}_{12} = 0 ,$$

et donc la force $\vec{F}_{1 \rightarrow 2}$ sera nulle. Mais si la deuxième particule ne se dirige pas vers la première, son champ magnétique sera non nul en 1 et il y aura une force $\vec{F}_{2 \rightarrow 1}$ non nulle... (N.B. Il ne faut pas penser pour autant que l'électromagnétisme viole la conservation de la quantité de mouvement. La théorie est incomplète à ce stade puisque on ne parle que de courants stationnaires. La conservation de la quantité de mouvement sera restaurée une fois que nous aurons les équations complètes de l'électrodynamique.)

9.1.2 Trajectoire d'une particule chargée en présence d'un champ magnétique

Considérons une particule de masse m placée dans un champ magnétique uniforme avec une vitesse initiale $\vec{v}(t=0) = \vec{v}_0$. La relation fondamentale de la dynamique s'écrit

$$\frac{d}{dt} \vec{v} = \frac{q}{m} (\vec{v} \wedge \vec{B}) . \quad (9.3)$$

Puisque la force magnétique est nulle dans la direction du champ, cette direction est privilégiée. On va donc tirer parti de cette information et décomposer la vitesse en deux composantes, l'une parallèle et l'autre perpendiculaire au champ, $\vec{v} = \vec{v}_{\parallel} + \vec{v}_{\perp}$. L'équation du mouvement s'écrit alors

$$\left\{ \begin{array}{l} \frac{d}{dt} \vec{v}_{\parallel} = \vec{0} \\ \frac{d}{dt} \vec{v}_{\perp} = \frac{q}{m} (\vec{v}_{\perp} \wedge \vec{B}) \end{array} \right. .$$

La trajectoire reste donc rectiligne uniforme dans la direction du champ. Prenons un repère cartésien dont l'axe z est donné par le champ $\vec{B} = B\hat{z}$. Dans ce repère, l'éq. (9.3) s'écrit,

$$\frac{d}{dt} \begin{bmatrix} v_x \\ v_y \\ v_{\parallel} \end{bmatrix} = \frac{q}{m} \begin{bmatrix} v_x \\ v_y \\ v_{\parallel} \end{bmatrix} \wedge \begin{bmatrix} 0 \\ 0 \\ B \end{bmatrix} = \frac{qB}{m} \begin{bmatrix} v_y \\ -v_x \\ 0 \end{bmatrix} .$$

et l'équation portant sur la composante perpendiculaire pour une charge, q_e , négative se décompose alors en deux équations :

$$\begin{aligned} \frac{d}{dt} v_x &= -\omega v_y \\ \frac{d}{dt} v_y &= \omega v_x \end{aligned} \quad \text{où} \quad \omega = \left| \frac{qB}{m} \right| .$$

Ce système se ramène à deux équations de la forme $\frac{d^2 v_i}{dt^2} = -\omega^2 v_i$ (pour $i = x, y$) et on a donc pour solution,

$$\begin{aligned} \frac{dx}{dt} &= v_x = v_{\perp 0} \cos \omega t \\ \frac{dy}{dt} &= v_y = v_{\perp 0} \sin \omega t , \end{aligned}$$

où l'on a choisi une vitesse initiale suivant x , $v_{\perp 0} = v_{\perp 0} \hat{x}$. En intégrant une deuxième fois ce système on obtient

$$\left\{ \begin{array}{l} x = \frac{v_{\perp 0}}{\omega} \sin \omega t \\ y = -\frac{v_{\perp 0}}{\omega} \cos \omega t \end{array} \right. .$$

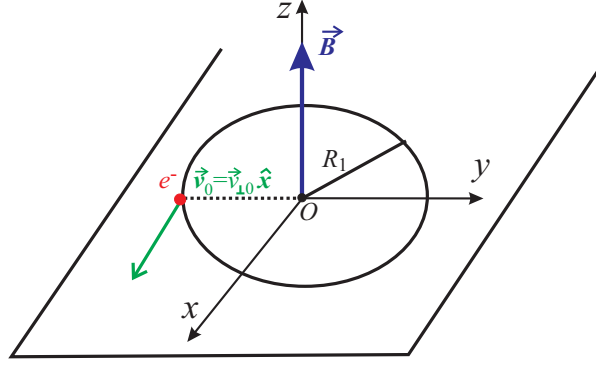


FIGURE 9.1 – Cas particulier d’une particule de charge négative (rotation dans le sens direct)

où les constantes d’intégration ont été choisies nulles (choix arbitraire). La trajectoire est donc un cercle de rayon $R_L = \left| \frac{v_{\perp 0}}{\omega} \right| = \left| \frac{mv_{\perp 0}}{qB} \right|$, le **rayon de Larmor**, décrit avec la pulsation $\omega = \frac{|q|B}{m}$, dite **pulsation gyro-synchrotron**. Ce cercle est parcouru dans le sens conventionnel positif pour des charges négatives.

Le rayon de Larmor correspond à la « distance » la plus grande que peut parcourir une particule dans la direction transverse avant d’être déviée de sa trajectoire. Cela correspond donc à une sorte de distance de piégeage. A moins de recevoir de l’énergie cinétique supplémentaire, une particule chargée est ainsi piégée dans un champ magnétique.

Il est intéressant de noter que plus l’énergie cinétique transverse d’une particule est élevée (grande masse ou grande vitesse transverse) et plus le rayon de Larmor est grand. Inversement, plus le champ magnétique est élevé et plus ce rayon est petit.

Remarque : Nous avons vu au chapitre 7 qu’une charge en mouvement crée un champ magnétique. Donc, une particule mise en rotation par l’effet d’un champ magnétique extérieur va créer son propre champ. Il n’en a pas été tenu compte dans le calcul précédent, celui-ci étant la plupart du temps négligeable.

9.1.3 Distinction entre champ électrique et champ électrostatique

Nous allons traiter ici un problème un peu subtil. En mécanique classique, il y a trois principes fondamentaux : le principe d’inertie, la relation fondamentale de la dynamique et le principe d’action et de réaction. Nous avons déjà vu que la force magnétique $\vec{F}_m = q(\vec{v} \wedge \vec{B})$ ne satisfaisait pas au 3^{ème} principe. Mais il y a pire. Pour pouvoir appliquer la relation fondamentale de la dynamique, il faut se choisir un référentiel galiléen. Ce choix étant arbitraire, les lois de la physique doivent être indépendantes de ce choix (invariance galiléenne). Autrement dit, les véritables forces doivent être indépendantes du référentiel. Il est clair que ce n’est pas le cas de la force magnétique $\vec{F}_m$. En effet, considérons une particule q se déplaçant dans un champ magnétique avec une vitesse constante dans le référentiel du laboratoire. Dans ce référentiel, elle va subir une force magnétique qui va dévier sa trajectoire. Mais si on se place dans le référentiel propre de la particule (en translation uniforme par rapport au laboratoire, donc galiléen), sa vitesse est nulle. Il n’y a donc pas de force et elle ne devrait pas être déviée ! Comment résoudre ce paradoxe ? C’est Lorentz qui a donné une solution formelle à ce problème, mais c’est Einstein qui lui a donné un sens grâce à la théorie de la relativité. La véritable force, électromagnétique, est la force de Lorentz :

$$\vec{F} = q(\vec{E} + \vec{v} \wedge \vec{B}) .$$

Supposons que cette particule soit soumise à un champ électrostatique $\vec{E}$ et un champ magnétique $\vec{B}$, mesurés dans le référentiel $\mathbb{R}$ du laboratoire. Dans un référentiel $\mathbb{R}'$ où la particule est au repos, le terme magnétique sera nul. Si on exige alors l'invariance de la force, on doit écrire

$$\vec{F}' = q\vec{E}' = \vec{F} = q\left(\vec{E}_s + \vec{v} \wedge \vec{B}\right) .$$

Le champ $\vec{E}'$ « vu » dans le référentiel $\mathbb{R}'$ est donc la somme du champ électrostatique $\vec{E}_s$ et d'un autre champ, appelé champ électromoteur $\vec{E}_m = \vec{v} \wedge \vec{B}$. Ainsi, on a bien conservé l'invariance de la force lors d'un changement de référentiel, mais au prix d'une complexification du champ électrique qui n'est plus simplement un champ électrostatique !

Deux conséquences importantes :

1. La circulation d'un champ électrique $\vec{E} = \vec{E}_s + \vec{E}_m$ est en générale non nulle à cause du terme électromoteur (d'où son nom d'ailleurs) : celui-ci peut donc créer une différence de potentiel qui va engendrer un courant, ce qui n'est pas possible avec un champ purement électrostatique.
2. Le champ électrique « vu » dans $\mathbb{R}'$ dépend du champ magnétique « vu » dans $\mathbb{R}$. On ne peut donc pas appliquer la règle de changement de référentiel classique (transformation galiléenne) mais une autre, plus complexe (transformation de Lorentz). Champs électrique et magnétique dépendent tous deux du référentiel, la compréhension de ce phénomène électromagnétique ne pouvant se faire que dans le contexte de la relativité.

Ceci dit, nous utiliserons tout de même l'expression de la force magnétique ou de Lorentz pour calculer, par exemple, des trajectoires de particules dans le formalisme de la mécanique classique. On ne devrait pas obtenir des résultats trop aberrants *tant que leurs vitesses restent très inférieures à celle de la lumière*.

Une dernière remarque : nous avons implicitement supposé que la charge q de la particule était la même dans les deux référentiels. Cela n'est a priori pas une évidence. Nous pouvons en effet imposer que toute propriété fondamentale de la matière soit effectivement invariante par changement de référentiel. Le concept de masse, par exemple, nécessite une attention particulière. En effet, tout corps massif possède un invariant appelé « masse au repos ». Cependant sa « masse dynamique » (impulsion divisée par sa vitesse) sera d'autant plus élevée que ce corps aura une vitesse s'approchant de celle de la lumière. *Nous admettrons donc que la charge électrique est bien un invariant (dit relativiste)*.

9.2 Actions magnétiques sur un circuit fermé

9.2.1 La force de Laplace

Nous avons vu que la force subie par une particule chargée en mouvement dans un champ magnétique, la force Lorentz, s'écrit ; $\vec{F} = q\left(\vec{E} + \vec{v} \wedge \vec{B}\right)$. Considérons un milieu comportant α espèces différentes de particules chargées, chaque espèce ayant une densité volumique n_α , et une vitesse $\vec{v}_\alpha$. Ces divers porteurs de charges sont donc responsables d'une densité locale de courant

$$\vec{j} = \sum_{\alpha} n_{\alpha} q_{\alpha} \vec{v}_{\alpha} .$$

Par ailleurs, chaque particule étant soumise à la force de Lorentz, la force s'exerçant sur un élément de volume $d\mathcal{V}$ comportant $\sum_{\alpha} n_{\alpha} \vec{v}_{\alpha}$ particules s'écrit,

$$d^3\vec{F} = \sum_{\alpha} n_{\alpha} q_{\alpha} \left(\vec{E} + \vec{v}_{\alpha} \wedge \vec{B}\right) d\mathcal{V} .$$

On voit donc apparaître une force due au champ électrique. Cependant, si le volume élémentaire que l'on considère est suffisamment grand pour que s'y trouve un grand nombre de particules et si le conducteur est électriquement neutre, on doit avoir

$$\sum_{\alpha} n_{\alpha} q_{\alpha} = 0 ,$$

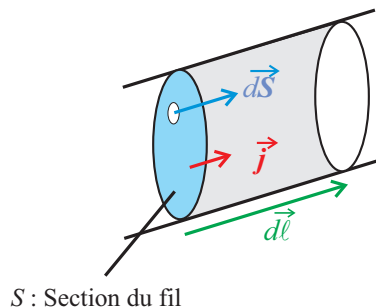
ce qui annule la force électrique.

On obtient alors $\overrightarrow{d^3\mathbf{F}} = \sum_{\alpha} n_{\alpha} q_{\alpha} (\vec{v}_{\alpha} \wedge \vec{B}) dV = (\sum_{\alpha} n_{\alpha} q_{\alpha} \vec{v}_{\alpha}) \wedge \vec{B} dV$, c'est-à-dire :

$$\boxed{\overrightarrow{d^3\mathbf{F}} = \vec{j} \wedge \vec{B} dV} . \quad (9.4)$$

Nous avons donc ci-dessus l'expression générale de la force créée par un champ magnétique extérieur sur une densité de courant quelconque circulant dans un conducteur neutre (la résultante est évidemment donnée par l'intégrale sur le volume).

Dans le cas particulier d'un conducteur filiforme, l'élément de volume s'écrit $dV = \vec{dS} \cdot \vec{d\ell}$, où $\vec{d\ell}$ est un élément de longueur infinitésimal orienté dans la direction de $\vec{j}$ et $\vec{dS}$ une surface infinitésimale (voir figure ci-dessous). Dans le cas d'un circuit filiforme (très mince donc où l'on peut considérer que



le champ $\vec{B}$ est constant), la force qui s'exerce sur une longueur $d\ell$ du fil s'écrit

$$\begin{aligned} \vec{dF}_L &= \iint (\vec{j} \wedge \vec{B}) \vec{dS} \cdot \vec{d\ell} = \iint (\vec{j} \cdot \vec{dS}) \vec{d\ell} \wedge \vec{B} \\ &= I \vec{d\ell} \wedge \vec{B} . \end{aligned}$$

La force qui s'exerce sur un conducteur fermé, parcouru par un courant permanent I , appelée **force de Laplace**, vaut

$$\boxed{\mathbf{F}_L = \oint_{\text{circuit}} \vec{dF}_L = I \oint_{\text{circuit}} \vec{d\ell} \wedge \vec{B}} . \quad (9.5)$$

Cette force s'applique sur un circuit qui est un solide. Dans ce cours, on ne considèrera que des circuits pour lesquels on pourra appliquer le principe fondamental de la mécanique, en assimilant ceux-ci à des points matériels (leur centre d'inertie). Aucun élément de longueur ne sera privilégié : la force $\vec{dF}_L = I \vec{d\ell} \wedge \vec{B}$ s'applique au milieu de chaque portion $\vec{d\ell}$.

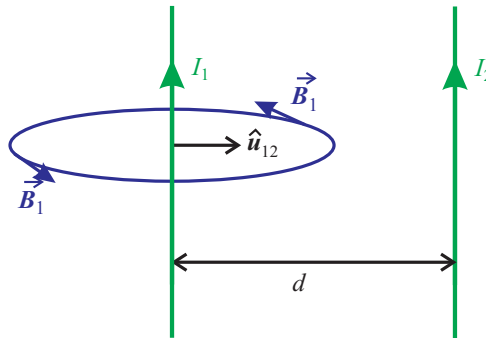
Remarques :

1. Ayant été établie à partir d'équations valables uniquement en régime permanent, cette expression n'est vraie que pour un courant permanent. Il faut en particulier faire attention à intégrer la force sur le circuit **fermé**.

2. **Pour des circuits fermés de forme complexe**, il devient difficile de calculer la force magnétique à partir de l'expression de la force de Laplace. Dans ce cas, *il vaut mieux utiliser une méthode énergétique (travaux virtuels, voir 9.3.3 plus bas)*.
3. A partir de la force de Lorentz, qui est une force microscopique agissant sur des particules individuelles et qui ne travaille pas, nous avons obtenu une force macroscopique agissant sur un solide. Cette force est capable de déplacer le solide et donc d'exercer un travail non nul. Comment comprendre ce résultat ? Il faut interpréter la force de Laplace comme la résultante de l'action des particules sur le réseau cristallin du conducteur. C'est donc une sorte de réaction du support à la force de Lorentz agissant sur ses constituants chargés. Au niveau microscopique cela se traduit par la présence d'un champ électrostatique, **le champ de Hall**(voir TD).
4. **Bien que la force de Lorentz ne satisfasse pas le principe d'Action et de Réaction, la force de Laplace entre deux circuits, elle, le satisfait !** (voir complément 9.4) La raison profonde réside dans l'hypothèse du courant permanent parcourant les circuits (I le même, partout dans chaque circuit) : en régime permanent, il n'y a plus de problème de délai lié à la vitesse de propagation finie de la lumière.

9.2.2 Définition légale de l'Ampère

Considérons le cas de deux fils infinis parcourus par un courant I_1 et I_2 , situés à une distance d l'un de l'autre. Grâce au théorème d'Ampère, il est alors facile de calculer le champ magnétique créé



par chaque fil. La force par unité de longueur subie par le fil 2 à cause du champ $\vec{B}_1$ vaut

$$\begin{aligned} \frac{d\vec{F}_{1 \rightarrow 2}}{d\ell_2} &= \frac{I_2 d\vec{\ell}_2 \wedge \vec{B}_1}{d\ell_2} = -\frac{\mu_0 I_1 I_2}{2\pi d} \hat{u}_{12} \\ &= -\frac{I_1 d\vec{\ell}_1 \wedge \vec{B}_2}{d\ell_1} = -\frac{d\vec{F}_{2 \rightarrow 1}}{d\ell_1}. \end{aligned}$$

Cette force est attractive si les deux courants sont dans le même sens, répulsive sinon. Puisqu'il y a une force magnétique agissant sur des circuits parcourus par un courant, on peut mesurer l'intensité de celui-ci. C'est par la mesure de cette force qu'a été établie la définition légale de l'Ampère (A) :

L'Ampère est l'intensité de courant passant dans deux fils parallèles, situés à 1 mètre l'un de l'autre, et produisant une attraction réciproque de $2 \cdot 10^{-7}$ Newtons par unité de longueur de fil.

C'est cette définition qui nous a donnée la valeur par définition de la perméabilité magnétique du vide, $\mu_0 = 4\pi 10^{-7}$ S.I. (unité H.m^{-1} ou N.A^{-2}) puisque avec $I_1 = I_2 = 1\text{A}$ et $d = 1\text{m}$, on trouve :

$$\left| \frac{d\vec{F}_{1 \rightarrow 2}}{d\ell_2} \right| = \frac{\mu_0}{2\pi} = \frac{4\pi 10^{-7}}{2\pi} = 2 \cdot 10^{-7} \text{ N.m}^{-1} \quad (9.6)$$

9.2.3 Moment de la force magnétique exercée sur un circuit

Puisqu'un circuit électrique est un solide, il faut utiliser le formalisme de la mécanique du solide. On va introduire ici les concepts minimaux requis.

Soit un point P quelconque appartenant à un circuit électrique et le point O , le centre d'inertie de ce circuit. Si ce circuit est parcouru par un courant permanent I et plongé dans un champ magnétique $\vec{B}$, alors chaque élément de circuit $d\vec{\ell} = d\vec{OP}$, situé autour de P , subit une force de Laplace $d\vec{F}_L = I d\vec{\ell} \wedge \vec{B}$. Le moment par rapport à O de la force de Laplace sur l'ensemble du circuit est alors

$$\vec{\Gamma} \equiv \oint_{\text{circuit}} \vec{OP} \wedge d\vec{F}_L . \quad (9.7)$$

Soient trois axes Δ_i , passant par le centre d'inertie O du circuit et engendrés par les vecteurs unitaires $\hat{e}_i$. Le moment des forces s'écrit alors $\vec{\Gamma} = \sum_i \Gamma_i \hat{e}_i$. L'existence d'un moment non nul se traduit par la mise en rotation du circuit autour d'un ou plusieurs axes Δ_i . Autrement dit, par une modification de la « quantité de rotation » du solide, c'est-à-dire son moment cinétique. Le moment cinétique du solide par rapport à O est

$$\vec{L} \equiv \oint_{\text{circuit}} \vec{OP} \wedge \vec{v} dm , \quad (9.8)$$

où dm est la masse élémentaire située sur l'élément $d\vec{OP}$, et $\vec{v}$ sa vitesse. Dans tous les cas de figure étudiés dans ce cours, on admettra qu'on peut choisir les axes Δ_i tel que le moment cinétique d'un circuit peut se mettre sous la forme,

$$\vec{L} = \sum_{i=1}^{i=3} \mathcal{I}_i \Omega_i \hat{e}_i ,$$

où Ω est le vecteur instantané de rotation et $\mathcal{I}_i$ les 3 moments d'inertie du circuit par rapport aux 3 axes Δ_i ($[\mathcal{I}]$ est la matrice d'inertie, ici diagonale). Le moment d'inertie par rapport à l'axe Δ_i est défini par,

$$\mathcal{I}_i = \oint_{\text{circuit}} dm r_i^2 ,$$

où r_i est la distance d'un point P quelconque du circuit à l'axe Δ_i . La dynamique d'un circuit soumis à plusieurs moments de forces extérieures est donnée par le théorème du moment cinétique (pour un solide)

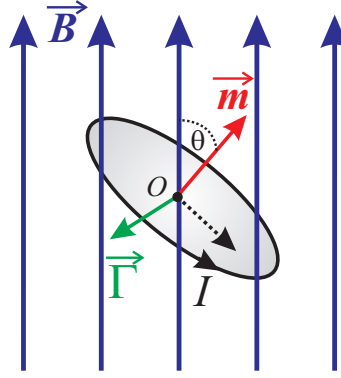
$$\frac{d\vec{L}}{dt} = \vec{\Gamma}_{\text{ext}} . \quad (9.9)$$

Dans le cas d'un circuit tournant autour d'un seul axe Oz , avec une vitesse angulaire $\vec{\Omega} = \Omega \hat{z} = \dot{\theta} \hat{z}$ et un moment d'inertie $\mathcal{I}$ constant, on obtient l'équation simplifiée suivante :

$$\mathcal{I}_z \ddot{\theta} = \Gamma_z . \quad (9.10)$$

9.2.4 Exemple du dipôle magnétique

Considérons le cas simple d'un dipôle magnétique, c'est-à-dire d'une spire parcourue par un courant I permanent, plongé dans un champ magnétique extérieur $\vec{B}$ constant (soit dans tout l'espace, soit ayant une variation spatiale sur une échelle bien plus grande que la taille de la spire). La force de



Laplace s'écrit alors,

$$\vec{F}_L = \oint_{\text{spire}} I \vec{d\ell} \wedge \vec{B} = \left(\oint_{\text{spire}} I \vec{d\ell} \right) \wedge \vec{B} = \vec{0} ,$$

(puisue dans le plan xOy du circuit, le déplacement vectoriel s'écrit : $\vec{d\ell} = \overrightarrow{dOP} = dx\hat{x} + dy\hat{y}$ comme on peut voir dans la figure 9.2 ci-dessous.)

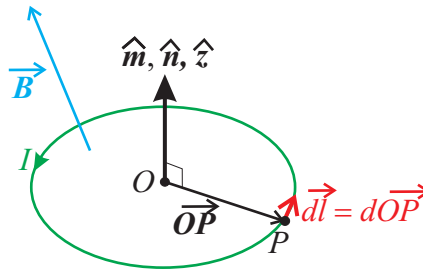


FIGURE 9.2 – Circuit dans un champ magnétique

Un champ magnétique constant ne va donc engendrer aucun mouvement de translation de la spire. Le moment de la force de Laplace par rapport au centre d'inertie O de la spire s'écrit :

$$\begin{aligned} \vec{\Gamma} &= \oint_{\text{spire}} \overrightarrow{OP} \wedge d\vec{F}_L = \oint_{\text{spire}} \overrightarrow{OP} \wedge (I d\overrightarrow{OP} \wedge \vec{B}) \\ &= I \oint_{\text{spire}} d\overrightarrow{OP} (\overrightarrow{OP} \cdot \vec{B}) - I \vec{B} \oint_{\text{spire}} (\overrightarrow{OP} \cdot d\overrightarrow{OP}) = I \oint_{\text{spire}} d\overrightarrow{OP} (\overrightarrow{OP} \cdot \vec{B}) - 0 \\ &= I \oint_{\text{spire}} (dx\hat{x} + dy\hat{y}) (xB_x + yB_y) = IB_y\hat{x} \oint_{\text{spire}} ydx + IB_x\hat{y} \oint_{\text{spire}} xdy \\ &= (-IB_y\hat{x} + IB_x\hat{y}) \oint_{\text{spire}} xdy = (IB_x\hat{y} - IB_y\hat{x}) \frac{1}{2} \oint_{\text{spire}} (xdy - ydx) = (IB_x\hat{y} - IB_y\hat{x}) S \\ &= IS\hat{n} \wedge \vec{B} , \end{aligned}$$

où $\hat{n} = \vec{z}$, ($\hat{x}$ et $\hat{y}$ sont dans le plan du dipôle magnétique : voir la figure 9.2). En se rappelant que $\vec{m} = IS\hat{n}$, on a

$$\boxed{\vec{\Gamma} = \vec{m} \wedge \vec{B}} . \quad (9.11)$$

Malgré une résultante des forces nulle, le champ magnétique exerce un moment qui va avoir tendance à faire tourner la spire sur elle-même, de telle sorte que son moment magnétique dipolaire $\vec{m}$ s'aligne dans la direction de $\vec{B}$.

Remarques :

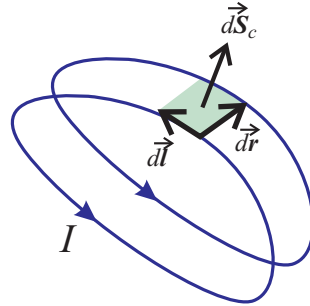
1. L'expression (9.11) ci-dessus n'est valable que dans le cas d'un dipôle.
2. On utilise souvent le terme « *couple magnétique* » pour décrire le moment des forces magnétiques sur un circuit, ceci pour éviter de confondre avec le moment magnétique dipolaire.
3. Les matériaux ferromagnétiques sont ceux pour lesquels on peut assimiler leurs atomes à des dipôles alignés dans le même sens. Mis en présence d'un champ magnétique externe, ils auront tendance à se mettre dans la direction du champ, ce qui va produire un mouvement macroscopique d'ensemble.

9.3 Energie potentielle d'interaction magnétique

9.3.1 Le théorème de Maxwell

Un circuit électrique parcouru par un courant produit un champ magnétique engendrant une force de Laplace sur un deuxième circuit, si celui-ci est lui-même parcouru par un courant. Chaque circuit agit sur l'autre, ce qui signifie qu'il y a une énergie d'origine magnétique mise en jeu lors de cette interaction. D'une façon générale, un circuit parcouru par un courant permanent placé dans un champ magnétique ambiant possède une énergie potentielle d'interaction magnétique.

Pour la calculer, il suffit d'évaluer le travail de la force de Laplace lors d'un déplacement virtuel de ce circuit (méthode des travaux virtuels, comme en électrostatique).



Considérons un élément $d\vec{\ell}$ d'un circuit filiforme, orienté dans la direction du courant I . Cet élément subit une force de Laplace $d\vec{F}_L$. Pour déplacer le circuit d'une quantité $d\vec{r}$, cette force doit fournir un travail

$$\begin{aligned} d^2W_L &= d\vec{F}_L \cdot d\vec{r} = I \left(d\vec{\ell} \wedge \vec{B} \right) \cdot d\vec{r} = I \left(d\vec{r} \wedge d\vec{\ell} \right) \cdot \vec{B} \\ &= IdS_c \hat{n} \cdot \vec{B} \equiv Id^2\Phi_c, \end{aligned}$$

où $dS_c \hat{n}$ est la surface élémentaire décrite lors du déplacement de l'élément de circuit (les trois vecteurs forment un trièdre direct). On reconnaît alors l'expression du flux magnétique à travers cette surface balayée, appelé flux coupé. Pour l'ensemble du circuit, le travail dû à un déplacement élémentaire $d\vec{r}$ est

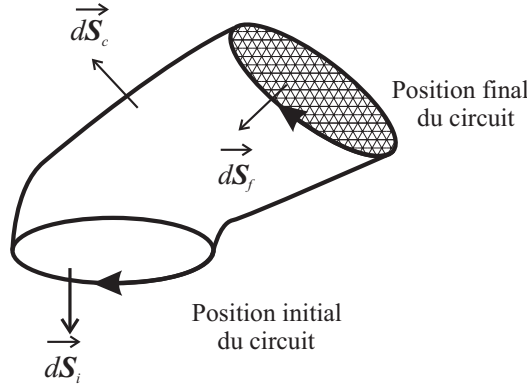
$$dW_L = \oint_{\text{circuit}} d^2W_L = \oint_{\text{circuit}} Id^2\Phi_c = Id\Phi_c.$$

Théorème de Maxwell :

Le déplacement d'un circuit électrique fermé dans un champ magnétique extérieur engendre un travail des forces magnétiques égal au produit du courant traversant le circuit par le flux coupé par celui-ci lors de son déplacement.

Commentaires sur la notion de Flux coupé

Le nom de flux coupé provient de notre représentation du champ magnétique sous forme de lignes de champ. Lors du déplacement du circuit, celui-ci va en effet passer à travers ces lignes, donc les « couper ». La notion de flux coupé est très importante car elle permet parfois de simplifier les calculs considérablement. Par ailleurs, *dans le cas d'un champ magnétique constant dans le temps*, nous allons démontrer que le flux coupé par le circuit Φ_c lors de son déplacement est exactement égal à la variation du flux total $\Delta\Phi$.



Soit un circuit C orienté, parcouru par un courant I et déplacé dans un champ magnétique extérieur (voir figure 9.3.1 ci-dessus). Ce circuit définit à tout instant une surface S s'appuyant sur C . Lors du déplacement de sa position initiale vers sa position finale, une surface fermée $\Sigma = S_i + S_f + S_c$ est ainsi décrite, où S_c est la surface balayée lors du déplacement. On choisit d'orienter les normales à chaque surface vers l'extérieur. La conservation du flux magnétique impose alors

$$\Phi_\Sigma = \Phi_{S_i} - \Phi_{S_f} + \Phi_c = 0 ,$$

c'est-à-dire

$$\Phi_c = \Phi_{S_f} - \Phi_{S_i} .$$

On a donc bien $\Phi_c = \Delta\Phi$ et le travail fait par la force de Laplace est :

$$\boxed{W_L = I\Delta\Phi} , \quad (9.12)$$

qui est vérifié algébriquement. *Ne pas oublier que ce raisonnement n'est valable que pour un champ magnétique extérieur statique* (pas de variation temporelle du champ au cours du déplacement du circuit).

9.3.2 Énergie potentielle d'interaction magnétique

Considérons un circuit électrique parcouru par un courant permanent I et placé dans un champ magnétique statique. Le circuit est donc soumis à la force de Laplace : cela signifie qu'il est susceptible de se déplacer et donc de développer une vitesse. On pourra calculer cette vitesse en appliquant, par exemple, le théorème de l'énergie cinétique $\Delta\mathcal{E}_c = W_L = I\Delta\Phi$. Mais d'où provient cette énergie ?

Si l'on en croit le principe de conservation de l'énergie, cela signifie que le circuit possède un réservoir d'énergie potentielle $\mathcal{U}_m$, lié à la présence du champ magnétique extérieur. L'énergie mécanique du circuit étant $\mathcal{E} = \mathcal{E}_c + \mathcal{U}_m$, on obtient $d\mathcal{U}_m = -dW_L$.

L'énergie potentielle magnétique d'un circuit parcouru par un courant permanent I et placé dans une champ magnétique extérieur est donc

$$\boxed{\mathcal{U}_m = -I\Phi + \text{Constante}} . \quad (9.13)$$

La valeur de la constante, comme pour toute énergie potentielle d'interaction, est souvent choisie arbitrairement nulle.

9.3.3 Expressions générales de la force et du couple magnétiques

L'existence d'une énergie potentielle se traduit par une action possible (reconversion de cette énergie). Ainsi la résultante $\vec{F}_L = \oint d\vec{F}_L$ des forces magnétiques exercées sur le circuit est donnée par

$$d\mathcal{U}_m = \sum_{i=1}^3 \frac{\partial \mathcal{U}_m}{\partial x_i} dx_i = -dW_L = -\vec{F}_L \cdot d\vec{r} = -\sum_{i=1}^3 F_i dx_i ,$$

où les dx_i mesurent les déplacements (translations) dans les trois directions de l'espace par rapport au centre d'inertie du circuit (là où s'applique la force magnétique). On obtient ainsi l'expression générale de la force de Laplace agissant sur un circuit parcouru par un courant permanent, c'est-à-dire,

$$F_i = -\frac{\partial \mathcal{U}_m}{\partial x_i} ,$$

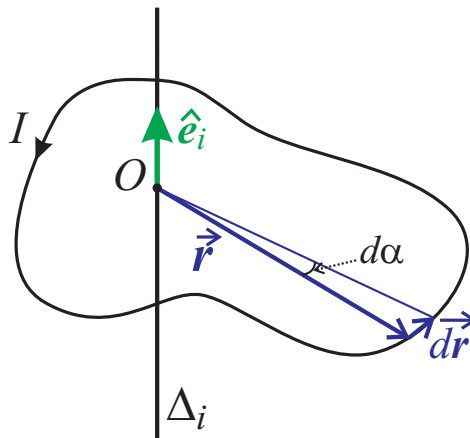
ou sous forme vectorielle,

$$\boxed{\vec{F}_L = -\vec{\text{grad}} \mathcal{U}_m = I \vec{\text{grad}} \Phi} . \quad (9.14)$$

Remarques :

1. La force totale (s'exerçant donc sur le centre d'inertie du circuit) a tendance à pousser le circuit vers les régions où le flux sera maximal.
2. Cette expression est valable uniquement pour des courants permanents. Noter qu'elle s'applique néanmoins pour des circuits déformés et donc pour lesquels il y aura aussi une modification du flux sans réel déplacement du circuit.

On peut faire le même raisonnement dans le cas d'un mouvement de rotation pure du circuit. Prenons le cas général de rotations d'angles infinitésimaux $d\alpha_i$ autour de trois axes Δ_i , passant par le centre d'inertie O du circuit et engendrés par les vecteurs unitaires $\hat{e}_i$.



Soit le vecteur $\vec{r} = \overrightarrow{OP} = r\hat{r}$ reliant un point P quelconque d'un circuit et le point O . La vitesse du point P s'écrit en toute généralité (voir un cours de mécanique)

$$\frac{d\vec{r}}{dt} = \frac{dr}{dt} \hat{r} + \vec{\Omega} \wedge \vec{r} ,$$

où le premier terme correspond à une translation pure et le second à une rotation pure, décrite par le vecteur instantané de rotation,

$$\vec{\Omega} = \begin{pmatrix} \dot{\alpha}_1 \\ \dot{\alpha}_2 \\ \dot{\alpha}_3 \end{pmatrix} = \sum_{i=1}^3 \frac{d\alpha_i}{dt} \hat{e}_i .$$

L'expression générale du moment de la force magnétique par rapport à O est $\vec{\Gamma} = \sum_{i=1}^3 \Gamma_i \hat{e}_i$. Le travail dû à la force de Laplace lors d'une rotation pure ($r = OP$ reste constant)

$$\begin{aligned} dW_L &= \oint_{\text{circuit}} \vec{dF}_L \cdot \vec{dr} = \oint_{\text{circuit}} \vec{dF}_L \cdot \left(\sum_{i=1}^3 d\alpha_i \hat{e}_i \wedge \vec{r} \right) \\ &= \sum_{i=1}^3 d\alpha_i \hat{e}_i \cdot \left(\oint_{\text{circuit}} \vec{r} \wedge \vec{dF}_L \right) = \sum_{i=1}^3 d\alpha_i \hat{e}_i \cdot \vec{\Gamma}_L = \sum_{i=1}^3 d\alpha_i \Gamma_i \\ &= Id\Phi = \sum_{i=1}^3 I \frac{\partial \Phi}{\partial \alpha_i} d\alpha_i , \end{aligned}$$

d'où

$$\Gamma_i = I \frac{\partial \Phi}{\partial \alpha_i} .$$

Autrement dit, le moment de la force magnétique par rapport à un axe Δ_i passant par le centre d'inertie O du circuit, dépend de la variation de flux lors d'une rotation du circuit autour de cet axe.

Exemple : Le dipôle : En supposant que le champ magnétique extérieur est constant à l'échelle d'un dipôle de moment magnétique dipolaire $\vec{m} = IS\hat{n}$, on obtient un flux $\Phi = \vec{B} \cdot S\hat{n}$.

La force magnétique totale s'écrit alors,

$$\vec{F}_L = I \overrightarrow{\text{grad}} \Phi = \overrightarrow{\text{grad}} (IS\hat{n} \cdot \vec{B}) ,$$

c'est-à-dire

$$\boxed{\vec{F}_L = \overrightarrow{\text{grad}} (\vec{m} \cdot \vec{B})} . \quad (9.15)$$

Le moment de la force magnétique (couple magnétique) s'écrit

$$\Gamma_i = I \frac{\partial \Phi}{\partial \alpha_i} = I \frac{\partial}{\partial \alpha_i} (\vec{B} \cdot S\hat{n}) = \vec{B} \cdot \frac{\partial (IS\hat{n})}{\partial \alpha_i} = \vec{B} \cdot \frac{\partial \vec{m}}{\partial \alpha_i} .$$

Or le moment magnétique dipolaire varie de la façon suivante lors d'une rotation,

$$d\vec{m} = \sum_{i=1}^3 d\alpha_i \hat{e}_i \wedge \vec{m} = \sum_{i=1}^3 \frac{\partial \vec{m}}{\partial \alpha_i} d\alpha_i ,$$

et on obtient donc,

$$\Gamma_i = \vec{B} \cdot (\hat{e}_i \wedge \vec{m}) = \hat{e}_i \cdot (\vec{m} \wedge \vec{B}) ,$$

c'est-à-dire l'expression vectorielle

$$\boxed{\vec{\Gamma}_L = \vec{m} \wedge \vec{B}} . \quad (9.16)$$

Remarquer que ce calcul est bien plus facile que le calcul direct effectué à la section 9.2.4.

9.3.4 La règle du flux maximum

Un solide est dans une position d'équilibre stable si les forces et les moments auxquels il est soumis tendent à le ramener vers cette position s'il en est écarté. D'après le théorème de Maxwell on a

$$dW_L = Id\Phi = I(\Phi_f - \Phi_i) = \vec{F}_L \cdot \vec{dr}.$$

Si la position est stable, cela signifie que l'opérateur doit fournir un travail, autrement dit un déplacement $\vec{dr}$ dans le sens contraire de la force (qui sera une force de rappel), donc $dW < 0$ ou $\Phi_f < \Phi_i$.

Un circuit tend toujours à se placer dans des conditions d'équilibre stable, où le flux du champ est maximum.

Cette règle est très utile pour se forger une intuition des actions magnétiques.

9.4 Laplace et le principe d'Action et de Réaction

On démontre ici que le principe d'Action et de Réaction est bien vérifié pour la force de Laplace s'exerçant entre deux circuits C_1 et C_2 quelconques, parcourus par des courants permanents I_1 et I_2 . La force exercée par C_1 sur C_2 s'écrit

$$\vec{F}_{1 \rightarrow 2} = \oint_{C_2} I_2 d\vec{P}_2 \wedge \vec{B}_1 = \oint_{C_2} I_2 d\vec{P}_2 \wedge \left[\oint_{C_1} \frac{\mu_0 I_1}{4\pi} \frac{d\vec{P}_1 \wedge \hat{u}_{12}}{(P_1 P_2)^2} \right] = \frac{\mu_0 I_1 I_2}{4\pi} \oint_{C_2} \oint_{C_1} d\vec{P}_2 \wedge \frac{d\vec{P}_1 \wedge \hat{u}_{12}}{(P_1 P_2)^2},$$

où P_1 (resp. P_2) est un point quelconque de C_1 (resp. C_2) et $\vec{P}_1 \vec{P}_2 = P_1 P_2 \hat{u}_{12}$. La force exercée par C_2 sur C_1 vaut

$$\vec{F}_{2 \rightarrow 1} = \oint_{C_1} I_1 d\vec{P}_1 \wedge \vec{B}_2 = \oint_{C_1} I_1 d\vec{P}_1 \wedge \left[\oint_{C_2} \frac{\mu_0 I_2}{4\pi} \frac{d\vec{P}_2 \wedge \hat{u}_{21}}{(P_2 P_1)^2} \right] = -\frac{\mu_0 I_1 I_2}{4\pi} \oint_{C_2} \oint_{C_1} d\vec{P}_1 \wedge \frac{d\vec{P}_2 \wedge \hat{u}_{12}}{(P_1 P_2)^2},$$

puisque $\hat{u}_{21} = -\hat{u}_{12}$. Par ailleurs, on a

$$\begin{aligned} d\vec{P}_2 \wedge (d\vec{P}_1 \wedge \hat{u}_{12}) &= d\vec{P}_1 (d\vec{P}_2 \cdot \hat{u}_{12}) - \hat{u}_{12} (d\vec{P}_2 \cdot d\vec{P}_1) \\ d\vec{P}_1 \wedge (d\vec{P}_2 \wedge \hat{u}_{12}) &= d\vec{P}_2 (d\vec{P}_1 \cdot \hat{u}_{12}) - \hat{u}_{12} (d\vec{P}_1 \cdot d\vec{P}_2). \end{aligned}$$

Il nous suffit donc de calculer le premier terme puisque le second est identique dans les deux expressions des forces. Mathématiquement, les expressions

$$\oint_{C_1} \oint_{C_2} [] = \oint_{C_2} \oint_{C_1} [] ,$$

sont effectivement équivalentes si ce qui se trouve dans le crochet (la fonction à intégrer) est symétrique par rapport aux variables de chacune des deux intégrales. Mais dans celle de gauche, le point P_1 reste d'abord constant lors de l'intégrale portant sur C_2 , tandis que dans celle de droite, c'est le point P_2 qui est maintenu constant lors de l'intégration sur C_1 . Ainsi, on peut écrire

$$\begin{aligned} \oint_{C_2} \frac{d\vec{P}_1 (d\vec{P}_2 \cdot \hat{u}_{12})}{(P_1 P_2)^2} &= d\vec{P}_1 \oint_{C_2} \frac{d\vec{P}_2 \cdot \hat{u}_{12}}{(P_1 P_2)^2} \\ \oint_{C_1} \frac{d\vec{P}_2 (d\vec{P}_1 \cdot \hat{u}_{12})}{(P_1 P_2)^2} &= d\vec{P}_2 \oint_{C_1} \frac{d\vec{P}_1 \cdot \hat{u}_{12}}{(P_1 P_2)^2}. \end{aligned}$$

Posant $\vec{r} = \overrightarrow{P_1 P_2}$ et remarquant que $\overrightarrow{dP_2} = d\overrightarrow{P_1 P_2} + d\overrightarrow{P_1} = d\overrightarrow{P_1 P_2}$ (puisque $d\overrightarrow{P_1 P_2} = d(\overrightarrow{P_2 O} - \overrightarrow{P_1 O}) = \overrightarrow{dP_2} - \overrightarrow{dP_1}$ et $\overrightarrow{dP_1} = 0$ pendant l'intégration sur C_2), on obtient :

$$\oint_{C_2} \frac{(\overrightarrow{dP_2} \cdot \hat{u}_{12})}{(P_1 P_2)^2} = \oint_{C_2} \frac{\overrightarrow{dr} \cdot \vec{r}}{r^3} = \oint_{C_2} \frac{d(r^2)}{2r^3} \stackrel{u \rightarrow r^2}{=} \frac{1}{2} \oint_{C_2} \frac{du}{u^{3/2}} = - \oint_{C_2} d(u^{-1/2}) = 0 .$$

puisque l'on fait un tour complet sur C_2 et l'on revient donc au point de départ. Le résultat est évidemment le même pour l'intégrale portant sur C_1 . En résumé, on obtient,

$$\begin{aligned} \vec{F}_{1 \rightarrow 2} &= \frac{\mu_0 I_1 I_2}{4\pi} \oint_{C_2} \oint_{C_1} \overrightarrow{dP_2} \wedge \frac{\overrightarrow{dP_1} \wedge \hat{u}_{12}}{(P_1 P_2)^2} = -\frac{\mu_0 I_1 I_2}{4\pi} \oint_{C_2} \oint_{C_1} \frac{\hat{u}_{12} \cdot (\overrightarrow{dP_2} \cdot \overrightarrow{dP_1})}{(P_1 P_2)^2} \\ &= -\frac{\mu_0 I_1 I_2}{4\pi} \oint_{C_2} \oint_{C_1} \frac{\hat{u}_{12} \cdot (\overrightarrow{dP_1} \cdot \overrightarrow{dP_2})}{(P_1 P_2)^2} = -\vec{F}_{2 \rightarrow 1} , \end{aligned}$$

Ceci achève la démonstration.

Chapitre 10

Induction électromagnétique

10.1 Les lois de l'induction

10.1.1 L'approche de Faraday

Jusqu'à maintenant, nous nous sommes intéressés essentiellement à la création d'un champ magnétique à partir d'un courant permanent. Ceci fut motivé par l'expérience de Oersted. A la même époque, le physicien anglais Faraday était préoccupé par la question inverse : puisque ces deux phénomènes sont liés, comment produire un courant à partir d'un champ magnétique ? Il fit un certain nombre d'expériences qui échouèrent car il essayait de produire un courant permanent. En fait, il s'aperçut bien de certains effets troublants, mais ils étaient toujours transitoires.

Exemple d'expérience : on enroule sur un même cylindre deux fils électriques. L'un est relié à une pile et possède un interrupteur, l'autre est seulement relié à un galvanomètre, permettant ainsi de mesurer tout courant qui serait engendré dans ce second circuit. En effet, Faraday savait que lorsqu'un courant permanent circule dans le premier circuit, un champ magnétique serait engendré et il s'attendait donc à voir apparaître un courant dans le deuxième circuit. En fait rien de tel n'était observé : lorsque l'interrupteur était fermé ou ouvert, rien ne se passait. Par contre, lors de son ouverture ou de sa fermeture, une déviation fugace de l'aiguille du galvanomètre pouvait être observée (cela n'a pas été perçu immédiatement). Une telle déviation pouvait également s'observer lorsque, un courant circulant dans le premier circuit, on déplaçait le deuxième circuit.

Autre expérience : prenons un aimant permanent et plaçons le à proximité d'une boucle constituée d'un fil conducteur relié à un galvanomètre. Lorsque l'aimant est immobile, il n'y a pas de courant mesurable dans le fil. Par contre, lorsqu'on déplace l'aimant, on voit apparaître un courant dont le signe varie selon qu'on approche ou qu'on éloigne l'aimant. De plus, ce courant est d'autant plus important que le déplacement est rapide.

Ces deux types d'expériences ont amené Faraday à écrire ceci : « *Quand le flux du champ magnétique à travers un circuit fermé change, il apparaît un courant électrique.* »

Dans les deux expériences, si on change la résistance R du circuit, alors le courant I apparaissant est également modifié, de telle sorte que $e = RI$ reste constant. Tous les faits expérimentaux mis en évidence par Faraday peuvent alors se résumer ainsi :

Loi de Faraday : la variation temporelle du flux magnétique à travers un circuit fermé y engendre une fém induite

$$e = - \frac{d\Phi}{dt} \quad (\text{ expression 1 }) . \quad (10.1)$$

L'induction électromagnétique est donc un phénomène qui dépend intrinsèquement du temps et sort du cadre de la magnétostatique (étude des phénomènes magnétiques stationnaires). Nous allons l'étudier

dans un premier temps dans le contexte d'une approximation dite « quasi-stationnaire » (voir section 10.3.1). L'importance de l'induction provient du fait qu'il s'agit d'un phénomène magnétique produisant un effet électrique.

10.1.2 La loi de Faraday

Posons-nous la question de Faraday. Comment crée-t-on un courant ? Un courant est un déplacement de charges dans un matériau conducteur. Ces charges sont mises en mouvement grâce une différence de potentiel (ddp) qui est maintenue par une force électromotrice ou fém (elle s'exprime donc en Volts). Une pile, en convertissant son énergie chimique pendant un instant dt , fournit donc une puissance P (travail W par unité de temps) modifiant l'énergie cinétique des dQ porteurs de charge et produisant ainsi un courant I .

Soit P_p la puissance communiquée à une particule de charge q_p se déplaçant à une vitesse $\vec{v}_p$. Sachant que dans un conducteur il y a n_p porteurs de charge par unité de volume, la puissance totale P que doit fournir le générateur (par ex. une pile) est

$$\begin{aligned} P &\equiv \iiint_V n_p P_p dV = \oint_{\text{circuit}} d\ell \iint_{\text{section}} n_p P_p dS = \oint_{\text{circuit}} d\ell \iint_{\text{section}} n_p \vec{F}_p \cdot \vec{v}_p dS \\ &= \oint_{\text{circuit}} \iint_{\text{section}} (n_p q_p \vec{v}_p \cdot d\vec{S}) \frac{\vec{F}_p \cdot d\vec{\ell}}{q_p} = \oint_{\text{circuit}} \frac{\vec{F}_p \cdot d\vec{\ell}}{q_p} \iint_{\text{section}} (\vec{j} \cdot d\vec{S}) \\ &= I \oint_{\text{circuit}} \frac{\vec{F}_p \cdot d\vec{\ell}}{q_p} = I e . \end{aligned}$$

On pose donc que la fém d'un circuit est

$$e \equiv \frac{P}{I} = \oint_{\text{circuit}} \frac{\vec{F}_p \cdot d\vec{\ell}}{q_p} , \quad (10.2)$$

où $\vec{F}_p$ est la force qui s'exerce sur les charges mobiles q_p (c.-à-d. les porteurs du courant). Or la force de Coulomb est incapable de produire une fém, puisque la circulation du champ **électrostatique** (donc le travail) est nulle sur un circuit fermé,

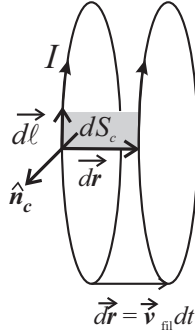
$$e = \oint_{\text{circuit}} \vec{E}_s \cdot d\vec{\ell} = V(A) - V(A) = 0 . \quad (10.3)$$

Il faut quand même se rappeler que l'induction est toujours associée avec une variation des paramètres physique dans le temps, (soit une variation temporelle du champ $\vec{B}$, soit le déplacement d'un circuit) et rien n'empêche que la circulation du champ électrique dans de telles circonstances ne soit pas nulle.

La force $\vec{F}_p$ responsable de l'apparition d'une fém sur les porteurs n'est rien d'autre que la force de Lorentz, c'est-à-dire

$$e = \oint_{\text{circuit}} (\vec{E} + \vec{v}_{\text{fil}} \wedge \vec{B}) \cdot d\vec{\ell} \quad (\text{ expression 2 }) , \quad (10.4)$$

où $\vec{v}_{\text{fil}}$ est la vitesse du fil électrique qui contient les porteurs du courant. **N.B.** Ce n'était pas nécessaire de tenir compte de la vitesse des porteurs du courant par rapport au fil électrique dans l'expression 2, puisque ce composant de leur vitesse est toujours parallèle à $d\vec{\ell}$.



Reprenons maintenant l'expérience qui consiste à déplacer un circuit fermé avec une vitesse $\vec{v}$ dans un champ magnétique $\vec{B}_s$ et un champ électrique $\vec{E}_s$ statiques. Que se passe-t-il pendant un instant dt ? La force de Lorentz (due à ce mouvement d'ensemble) agissant sur chaque particule q du conducteur s'écrit $\vec{F} = q (\vec{E}_s + \vec{v}_{\text{fil}} \wedge \vec{B}_s)$, fournissant ainsi une f.é.m.,

$$\begin{aligned} e &= \oint_{\text{circuit}} (\vec{E}_s + \vec{v}_{\text{fil}} \wedge \vec{B}_s) \cdot d\vec{\ell} = -\frac{1}{dt} \oint_{\text{circuit}} (\vec{v}_{\text{fil}} dt \wedge d\vec{\ell}) \cdot \vec{B}_s \\ &= -\frac{1}{dt} \oint_{\text{circuit}} dS_c \hat{n} \cdot \vec{B}_s, \end{aligned}$$

où $dS_c \hat{n}$ est la surface orientée élémentaire, décrite lors du déplacement du circuit. On reconnaît alors l'expression du flux coupé à travers cette surface élémentaire. On a donc

$$e = -\frac{1}{dt} \oint_{\text{circuit}} d^2\Phi_c = -\frac{d\Phi_c}{dt} = -\frac{d\Phi}{dt},$$

puisque la variation du flux coupé est égale à celle du flux total à travers le circuit (conservation du flux magnétique, cf. théorème de Maxwell). Attention au sens de $d\vec{\ell}$: il doit être cohérent avec $d\Phi_c = d\Phi$.

Nous venons de démontrer la loi de Faraday dans le cas d'un circuit rigide, déplacé dans un champ électromagnétique statique. Nous avons vu apparaître naturellement l'expression du flux coupé. En fait, la seule chose qui compte, c'est l'existence d'un mouvement d'ensemble du tout ou d'une partie du circuit (revoir démonstration pour s'en convaincre). Ainsi, l'expression de la fém induite

$$\boxed{e = -\frac{d\Phi_c}{dt} \quad (\text{expression 3})}, \quad (10.5)$$

reste valable pour un circuit déformé et/ou déplacé dans un champ magnétique statique.

Première difficulté

Prenons l'expérience de la roue de Barlow. L'appareil consiste en un disque métallique mobile autour d'un axe fixe, plongeant dans un champ magnétique et touchant par son bord extérieur une cuve de mercure. Un circuit électrique est ainsi établi entre la cuve et l'axe et on ferme ce circuit sur un galvanomètre permettant de mesurer tout courant. Lorsqu'on fait tourner le disque, un courant électrique est bien détecté, en cohérence avec la formule de l'expression 2. Cependant, il n'y a pas de variation du flux total à travers la roue ! Ce résultat expérimental semble donc contradictoire avec $e = -\frac{d\Phi_c}{dt}$!

Comment comprendre cela ? Même si, globalement, il n'y a pas de variation du flux total, il n'en reste pas moins que les charges du disque conducteur se déplacent dans un champ magnétique. On pourrait donc faire fi de l'égalité $d\Phi_c = d\Phi$ et calculer ainsi une f.é.m. non nulle. Cependant, **la cause**

physique fondamentale de l'induction réside dans l'expression 2. Il faut donc utiliser *les expressions 1 et 3 uniquement comme des moyens parfois habiles de calculer la fém.*

Deuxième difficulté

Si on se place maintenant dans le référentiel du circuit rigide, on verra un champ magnétique variable (c'est le cas, par ex, lorsqu'on approche un aimant d'un circuit immobile). Dans ce cas, le flux coupé est nul et on devrait donc avoir une fém nulle, ce qui n'est pas le cas d'après l'expérience de Faraday. Ce résultat expérimental semble cette fois-ci en contradiction avec $e = -\frac{d\Phi_c}{dt}$!

Résolution de ce paradoxe

Puisque, dans ce dernier cas, le champ magnétique dépend du temps, il n'y a plus de lien direct entre le flux coupé et le flux total à travers le circuit. Si on revient à l'expression 2, on voit que dans le référentiel du circuit la force « magnétique » est nulle et il ne reste plus que le terme « électrique ». Or, lorsque nous avons considéré les changements de repère dans section III.1.3, nous avons vu que le champ électrique dépend du champ magnétique (de rotation non nul) dans un autre repère. Afin d'être en accord avec l'expérience de Faraday, on doit en conclure que la variation d'un champ magnétique doit produire un champ électrique de rotationnel non nul $\vec{E}_m$ (dit champ électromoteur) qui va s'ajouter au champ électrostatique, $\vec{E}_s$. Avec cette supposition, on obtient :

$$\begin{aligned} e &= -\frac{d\Phi}{dt} \quad (\text{expérience de Faraday}) \\ &= \oint_{\text{circuit}} \vec{E}_m \cdot d\vec{\ell} = \oint_{\text{circuit}} \vec{E} \cdot d\vec{\ell} \quad (\text{circulation du champ électromoteur}), \end{aligned}$$

où nous avons utilisé $\vec{E} = \vec{E}_s + \vec{E}_m$ et l'équation 10.3 dans la dernière égalité. En se rappelant de la définition du flux magnétique, cette relation s'écrit,

$$\frac{d\Phi}{dt} = \frac{d}{dt} \iint_{\text{circuit}} \vec{B} \cdot d\vec{S} = \iint_{\text{circuit}} \frac{\partial \vec{B}}{\partial t} \cdot d\vec{S} = - \oint_{\text{circuit}} \vec{E} \cdot d\vec{\ell}.$$

Cette équation intégrale décrit ainsi un nouvel effet physique, totalement indépendant de tout ce que nous avons vu jusqu'à présent : l'induction. Comme, les champs sont présents même en absence d'un circuit, on en déduit la validité de cette relation pour un contour arbitraire C et elle devient ainsi *une des quatre équations fondamentales du champ électromagnétique* :

$$\boxed{\oint_C \vec{E} \cdot d\vec{\ell} = - \iint_S \frac{\partial \vec{B}}{\partial t} \cdot d\vec{S}}. \quad (10.6)$$

Comme les autres lois fondamentales, on peut exprimer l'éq. 10.6 sous forme différentielle. Le théorème de Stokes nous dicte que

$$\oint_C \vec{E} \cdot d\vec{\ell} = \iint_S \text{rot } \vec{E} \cdot d\vec{S}, \quad (10.7)$$

pour n'importe quelle champ vectoriel et puisque le contour C est arbitraire, une comparaison entre l'éq. 10.7 et l'éq. 10.6, nous dicte *l'expression différentielle du rotationnel de $\vec{E}$* :

$$\boxed{\text{rot } \vec{E} = -\frac{\partial \vec{B}}{\partial t}}. \quad (10.8)$$

Résumé/Bilan

Que se passe-t-il si on *déplace un circuit (rigide ou non) dans un champ magnétique variable* ? Quelle expression faut-il utiliser ? En fait, il faut revenir à la force de Lorentz dans le cas général de champs variables. On aura alors une fém induite

$$\begin{aligned}
 e &= \oint_{\text{circuit}} \left(\vec{E} + \vec{v}_{\text{fil}} \wedge \vec{B} \right) \cdot d\vec{\ell} = \oint_{\text{circuit}} \vec{E} \cdot d\vec{\ell} - \frac{d\Phi_c}{dt} \\
 &= - \iint_S \frac{\partial \vec{B}}{\partial t} \cdot d\vec{S} - \frac{d\Phi_c}{dt} = - \frac{d\Phi}{dt}
 \end{aligned} \tag{10.9}$$

Le premier terme décrit la circulation non nulle d'un champ électromoteur, associé à la variation temporelle du champ magnétique, tandis que le deuxième terme décrit la présence d'un flux coupé dû au déplacement du circuit et/ou à sa déformation.

Note Bene : La convention est d'appeler « *e* » la force électromotrice (f.é.m.) du circuit même si on voit qu'elle est en réalité *l'intégrale d'une force par unité de charge sur un tour complet du circuit*. De ce fait « *e* » a les unités de Volts même si l'idée de potentiel électrique n'est plus applicable en présence d'induction. Puisque « *e* » est mesuré en Volts on dit parfois « différence de potentiel » dû à la force électromotrice mais ceci est un léger abus de langage.

Quelques subtilités :

1. Puisque les porteurs de courant subissent une f.é.m., on aurait pu se demander pourquoi cette force n'entraîne pas une force mécanique sur le fil. Il faut se rappeler que le fil est neutre et qu'il y a autant de charges immobiles (de charge opposée) qu'il y a de charges porteurs de courant. Les forces sur ces charges sont égales et opposées aux forces sur les porteurs, et par conséquent, la force d'induction ne produit pas de force mécanique sur un fil neutre.
2. On pouvait également se demander pourquoi nous avons invoqué seulement la vitesse du fil, $\vec{v}_{\text{fil}}$, et pas la vitesse moyenne des porteurs de courant, $\vec{v}_p$, qui est responsable pour le courant ($I = Sq n_p |\vec{v}_p|$). La véritable vitesse moyenne des charges de conduction est ainsi $\vec{v} = \vec{v}_{\text{fil}} + \vec{v}_p$. Les raisons pour cet apparent oubli sont les suivantes :
 - (a) La vitesse $\vec{v}_p$ associé avec le courant est dirigé selon la direction $d\vec{\ell}$ du fil. La force électromotrice, $\vec{v}_p \wedge \vec{B}$ qu'elle génère sera donc perpendiculaire au fil. Cette force va rapidement redistribuer les porteurs de courant afin d'établir un **champ de Hall** qui produira une force qui lui est égale et opposée (voir TD). Donc, la plupart du temps on se soucie gère de cette force et elle ne contribue pas dans le calcul de la f.é.m. du circuit (parce qu'elle est perpendiculaire à $d\vec{\ell}$).
 - (b) D'autre part, le fait que les charges sont en mouvement avec une vitesse $\vec{v}_p$ par rapport au fil, ils subissent une force mécanique $\vec{v}_p \wedge \vec{B}$ que les charges immobiles ne subissent pas. Cette force mécanique est précisément la force de Laplace que nous avons déjà étudié dans le chapitre 9. Nous constatons donc, que **la force de Lorentz est à la fois responsable de la force de Laplace et au terme $\vec{v}_{\text{fil}} \wedge \vec{B}$ dans la force électromotrice**.

10.1.3 La loi de Lenz

Enoncé : *l'induction produit des effets qui s'opposent aux causes qui lui ont donné naissance.*

Cette loi est, comme la règle du flux maximum, déjà contenue dans les équations et donc n'apporte rien de plus, hormis une intuition des phénomènes physiques. En l'occurrence, la « loi de Lenz » n'est que l'expression du signe « $-$ » contenu dans la loi de Faraday.

Exemple : si on approche un circuit du pôle nord d'un aimant, le flux augmente et donc la fém induite est négative. Le courant induit sera alors négatif et produira lui-même un champ magnétique induit opposé à celui de l'aimant. **Deux conséquences :**

1. L'augmentation du flux à travers le circuit est amoindrie.
2. Il apparaît une force de Laplace $\vec{F} = I \vec{\text{grad}} \Phi$ négative, s'opposant à l'approche de l'aimant.

Le signe « - » dans la loi de Faraday (la loi de Lenz) décrit le fait que dans des conditions normales, il n'y a pas d'emballement possible (ex. courant ne faisant qu'augmenter).

Remarque sur la convention de signe

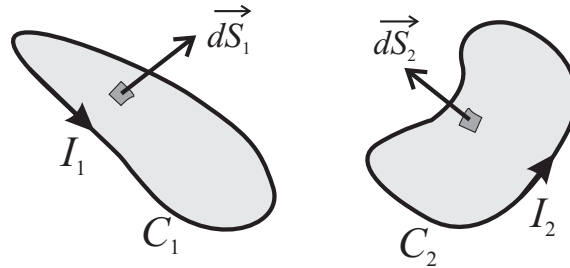
La détermination du sens du courant induit se fait de la façon suivante :

1. On se choisit arbitrairement un sens de circulation le long du circuit.
2. Ce sens définit, grâce à la règle du bonhomme d'Ampère, une normale au circuit.
3. Le signe du flux est alors déterminé en faisant le produit scalaire du champ magnétique par cette normale.
4. En utilisant ensuite la loi de Faraday, on obtient la valeur et le signe de la fém.
5. Enfin, le courant est obtenu à partir de la loi d'Ohm (son signe peut aussi être directement connu en utilisant la loi de Lenz).

10.2 Induction mutuelle et auto-induction

10.2.1 Induction mutuelle entre deux circuits fermés

Soient deux circuits fermés, orientés, traversés par des courants I_1 et I_2 .



Le premier crée un champ magnétique $\vec{B}_1$ dont on peut calculer le flux Φ_{12} à travers le deuxième circuit,

$$\Phi_{12} = \iint_{S_2} \vec{B}_1 \cdot d\vec{S}_2 = \left[\frac{\mu_0}{4\pi} \iint_{S_2} \oint_{C_1} \frac{d\vec{\ell}_1 \wedge \vec{PM}}{|\vec{PM}|^3} \cdot d\vec{S}_2 \right] I_1, \quad (10.10)$$

où P est un point quelconque du circuit C_1 (l'élément de longueur valant $d\vec{\ell}_1 = d\vec{OP}$) et M un point quelconque de la surface délimitée par C_2 , à travers laquelle le flux est calculé. De même, on a pour le flux créé par le circuit C_2 sur le circuit C_1 :

$$\Phi_{21} = \iint_{S_1} \vec{B}_2 \cdot d\vec{S}_1 = \left[\frac{\mu_0}{4\pi} \iint_{S_1} \oint_{C_2} \frac{d\vec{\ell}_2 \wedge \vec{PM}}{|\vec{PM}|^3} \cdot d\vec{S}_1 \right] I_2, \quad (10.11)$$

où P est cette fois-ci un point du circuit C_2 et M un point de la surface délimitée par C_1 , à travers laquelle le flux est calculé. Les termes entre crochets dépendent de la distance entre les deux circuits

et de facteurs uniquement géométriques liés à la forme de chaque circuit. Comme, dans le cas général, ils sont difficiles voire impossible à calculer, il est commode de poser

$$\begin{aligned}\Phi_{12} &= M_{12}I_1 \\ \Phi_{21} &= M_{21}I_2 .\end{aligned}$$

Le signe des coefficients dépend de l'orientation respective des circuits et suit la même logique que pour le courant induit. D'après les choix pris pour le sens de circulation le long de chaque circuit (voir figure), les flux sont négatifs pour des courants I_1 et I_2 positifs. Donc les coefficients sont négatifs dans cet exemple.

Théorème : *Le coefficient d'induction mutuelle ou inductance mutuelle (unités : Henry)*

$$\boxed{M = M_{12} = M_{21}} . \quad (10.12)$$

Il met en jeu une énergie potentielle d'interaction magnétique entre les deux circuits

$$\boxed{\mathcal{U}_m = -MI_1I_2 + \text{Constante}} .$$

Il nous faut démontrer que les inductances sont bien les mêmes pour chaque circuit. La raison profonde réside dans le fait qu'ils sont en interaction, donc possèdent chacun la même énergie potentielle d'interaction. Si on déplace C_2 , il faut fournir un travail

$$dW_2 = I_2 d\Phi_{12} = I_1 I_2 dM_{12}$$

Mais ce faisant, on engendre une variation du flux à travers C_1 et donc un travail

$$dW_1 = I_1 d\Phi_{21} = I_2 I_1 dM_{21}$$

Puisqu'ils partagent la même énergie d'interaction (chaque travail correspond au mouvement relatif de C_1 par rapport à C_2), on a $dW_1 = dW_2$ et donc

$$dM_{12} = dM_{21} \Rightarrow M_{12} = M_{21} + \text{Constante}$$

Cette constante d'intégration doit être nulle puisque, si on éloigne les circuits l'un de l'autre à l'infini, l'interaction tend vers zéro et donc les inductances s'annulent.

N.B. Quand nous disons ci-dessus que les deux circuits possèdent le même potentiel d'interaction, ça revient à dire que le principe d'action et réaction est satisfait pour les circuits. Nous avons démontré dans la section [9.4](#) que ceci est bien vérifié pour des circuits avec courants stationnaires.

10.2.2 Auto-induction

Si on considère un circuit isolé, parcouru par un courant I , on s'aperçoit qu'on peut produire le même raisonnement que ci-dessus. En effet, le courant I engendre un champ magnétique dans tout l'espace et il existe donc un flux de ce champ à travers le circuit lui-même,

$$\Phi = \iint_S \vec{B} \cdot d\vec{S} = \left[\frac{\mu_0}{4\pi} \iint_S \oint_C \frac{d\vec{\ell} \wedge \vec{PM}}{|\vec{PM}|^3} \cdot d\vec{S} \right] I ,$$

qu'on peut simplement écrire

$$\boxed{\Phi = LI} , \quad (10.13)$$

où L est le coefficient d'auto-induction ou auto-inductance (ou self), exprimé en Henry. Il ne dépend que des propriétés géométriques du circuit et est nécessairement positif (alors que le signe de l'inductance mutuelle dépend de l'orientation d'un circuit par rapport à l'autre).

10.3 Régimes variables

10.3.1 Définition du régime quasi-stationnaire

Avec les lois que nous avons énoncé jusqu'à présent, nous sommes en mesure d'étudier certains régimes variables. En effet, tous les raisonnements basés sur la notion d'un champ (électrique ou magnétique) constant au cours du temps peuvent aisément être appliqués à des systèmes physiques variables (champs dépendant du temps), pourvu que cette variabilité s'effectue sur des échelles de temps longues par rapport au temps caractéristique d'ajustement du champ. Voici tout de suite un exemple concret.

La plupart des lois de la magnétostatique supposent un courant permanent, c'est-à-dire le même dans le tout le circuit. Lorsqu'on ferme un interrupteur, un signal électromagnétique se propage dans tout le circuit et c'est ainsi que peut s'établir un courant permanent : cela prend un temps de l'ordre de ℓ/c , où ℓ est la taille du circuit et c la vitesse de la lumière. Si l'on a maintenant un générateur de tension sinusoïdale de période T , alors on pourra malgré tout utiliser les relations déduites de la magnétostatique si

$$\boxed{T \gg \ell/c} . \quad (10.14)$$

Ainsi, bien que le courant soit variable, la création d'un champ magnétique obéira à la loi de Biot et Savart tant que le critère ci-dessus reste satisfait. Ce type de régime variable est également appelé régime quasi-stationnaire.

10.3.2 Forces électromotrices (fém) induites

Considérons tout d'abord le cas d'un circuit isolé rigide (non déformable). Nous avons vu qu'une fém induite apparaissait dès lors que le flux variait. D'après la loi de Faraday et l'expression ci-dessus, cette fém vaudra :

$$\boxed{e = -L \frac{dI}{dt}} , \quad (10.15)$$

(L étant constant pour un circuit rigide). En régime variable, si le courant diminue, on verra donc apparaître une fém positive engendrant un courant induit qui va s'opposer à la décroissance du courant dans le circuit. La self d'un circuit tend donc à atténuer les variations de courant.

Dans les schémas électriques la self est symbolisée par une bobine. C'est en effet la façon la plus commode de produire une self : plus le nombre de spires est élevé et plus grande sera l'auto-inductance L (le cylindre sur lequel on fait l'enroulement est d'ailleurs souvent constitué de fer doux, matériau ferromagnétique, pour amplifier le champ, donc L). Puisque le même flux magnétique traverse les N spires de la bobine le flux à travers la bobine s'écrit,

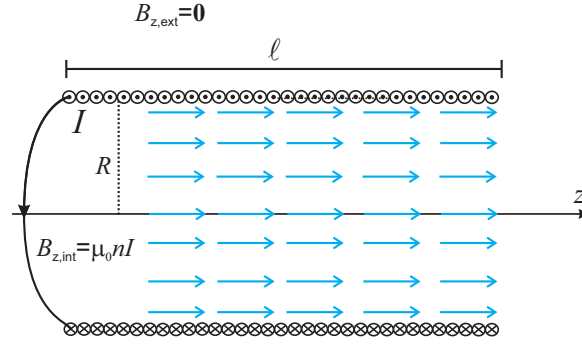
$$\Phi_{\text{bobine}} \equiv \iint_{\text{bobine}} \vec{B} \cdot d\vec{S} \simeq N \iint_{\text{spire}} \vec{B} \cdot d\vec{S} ,$$

ce qui donne la règle

$$\boxed{\Phi_{\text{bobine}} \simeq N \Phi_{\text{spire}}} . \quad (10.16)$$

Exemple : On peut appliquer cette règle afin de trouver l'auto induction d'une bobine circulaire avec un coeur vide (pas de fer doux) avec N spires circulaires de rayon R , distribués uniformément sur une longueur ℓ dans l'approximation d'une bobine infini, c'est-à-dire $R \ll \ell$.

On se rappelle qu'en chapitre 7, que nous avons montré que le champ $\vec{B}$ est constant (dirigé selon l'axe de la bobine) et que $|\vec{B}| = \mu_0 n I = (\mu_0 N I) / \ell$. Le flux magnétique à travers une spire circulaire



s'écrit donc :

$$\Phi_{\text{spire}} = \iint_{\text{spire}} \vec{B} \cdot d\vec{S} = \pi R^2 (\mu_0 N I) / \ell .$$

Avec la règle de l'éq. (10.16), nous obtenons $\Phi_{\text{bobine}} = N \Phi_{\text{spire}} = (\mu_0 \pi R^2 N^2 I) / \ell$ et par la définition de L de l'éq. (10.13) on obtient :

$$L_{\text{self, vide}} \simeq (\mu_0 \pi R^2 N^2) / \ell , \quad (10.17)$$

comme coefficient d'auto-induction de la bobine.

Si l'on considère maintenant deux bobines couplées C_1 et C_2 , alors l'expression des flux totaux à travers ces circuits s'écrit :

$$\begin{cases} \Phi_1 = \Phi_{11} + \Phi_{21} = L_1 I_1 + M I_2 \\ \Phi_2 = \Phi_{22} + \Phi_{12} = L_2 I_2 + M I_1 \end{cases} .$$

On aura donc en régime variable des fém induites dans chaque circuit

$$e_1 = -L_1 \frac{dI_1}{dt} - M \frac{dI_2}{dt} \quad (10.18a)$$

$$e_2 = -M \frac{dI_1}{dt} - L_2 \frac{dI_2}{dt} . \quad (10.18b)$$

Ce couplage entre deux circuits électriquement indépendants peut avoir des conséquences importantes (parfois désastreuses), comme l'apparition soudaine d'un courant dans un circuit fermé non alimenté. En effet, supposons que I_2 soit nul à un instant et qu'il y ait à ce moment là une variation de courant I_1 . L'induction mutuelle va alors engendrer un courant I_2 induit, qui va à son tour modifier I_1 .

On utilise le système de deux équations couplées dans l'éq. (10.18) afin d'éliminer certaines variables. Par exemple, s'intéresse souvent aux forces électromotrices dans les deux circuits. On peut donc utiliser l'éq. (10.18b) afin d'exprimer :

$$\frac{dI_2}{dt} = -\frac{e_2}{L_2} - \frac{M}{L_2} \frac{dI_1}{dt} . \quad (10.19)$$

Mettant ceci dans l'éq. (10.18a), on élimine la variation $\frac{dI_2}{dt}$ afin d'obtenir une équation reliant e_1 , e_2 , et $\frac{dI_1}{dt}$:

$$e_1 = \frac{M}{L_2} e_2 - \left(\frac{L_1 L_2 - M^2}{L_2} \right) \frac{dI_1}{dt} . \quad (10.20)$$

Dans le cas générale, il nous faudrait d'autres informations reliant dI_1/dt à e_1 et e_2 afin d'établir complètement la relation entre e_1 et e_2 , mais cette relation suffit pour traiter le cas d'un transformateur en couplage quasi-parfait comme nous verrons dans la section 10.3.3 ci-dessous.

10.3.3 Couplage entre bobines : Transformateurs

Puisque le champ magnétique est généralement très faible à l'extérieur des bobines, l'induction mutuelle entre deux bobines est quasi-nulle sauf si une bobine est à l'intérieur de l'autre, ou s'il y a un guide du flux magnétique d'une bobine vers l'autre comme c'est le cas dans la plupart des transformateurs (voir figure 10.2).

Considérons d'abord le cas simple d'une bobine de longueur ℓ_2 avec N_2 spires de rayon R_2 à l'intérieur d'une bobine de longueur $\ell_1 \geq \ell_2$ et de rayon $R_1 > R_2$. On laissera l'axe de la bobine de 2 faire un angle θ avec l'axe de la bobine à l'extérieur comme illustré dans la figure 10.1 (a)).

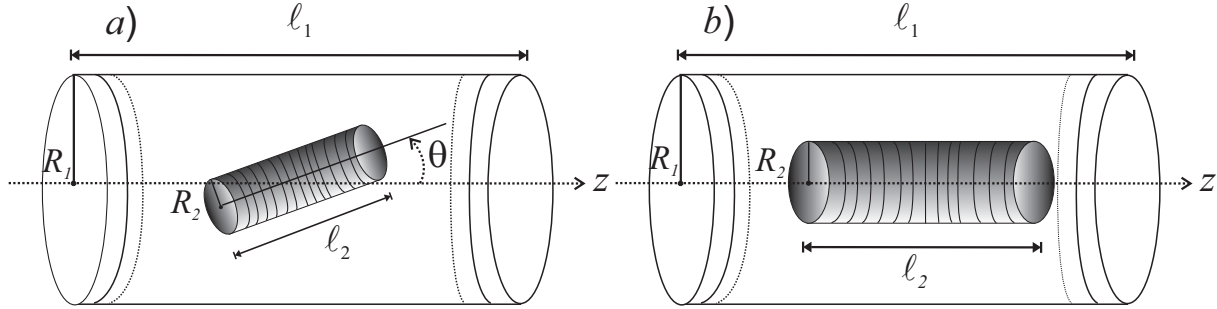


FIGURE 10.1 – Bobines en couplage magnétique

En prenant l'approximation simplificatrice de bobine 1 infini, le champ magnétique produit par la bobine 1 est homogène sur l'ensemble de la bobine 2 et donné par $\vec{B}_1 = \mu_0 n_1 I_1 \hat{z}$, où l'axe de la bobine 1 est pris selon $\hat{z}$ et I_1 est le courant qui parcourt cette bobine. Par conséquent, on peut aisément calculer le flux magnétique produit par I_1 et traversant la bobine 2 comme $\Phi_{12} = \mu_0 n_1 N_2 \pi R_2^2 \cos \theta I_1 \equiv M_{12} I_1$, donc l'influence mutuelle est :

$$M = M_{12} = \mu_0 N_1 N_2 \pi R_2^2 \cos \theta / \ell_1 . \quad (10.21)$$

Le calcul de $\Phi_{21} = M_{21} I_1$ serait moins évident, mais il n'est pas nécessaire car on sait que les inductions mutuelles sont égaux, $M_{21} = M_{12} = M$. Il est pratique à définir un facteur de couplage magnétique, k , entre les deux circuits de la façon suivante :

$$k \equiv \frac{M}{(L_1 L_2)^{1/2}} . \quad (10.22)$$

Dans le cas d'un transformateur, on cherche à maximiser l'induction mutuelle et on prend généralement l'axe de bobine 2 comme étant orientée selon le même axe que la bobine 1, donc $\theta = 0$ et $\cos \theta = 1$ comme illustré dans la figure 10.1 b) ce qui nous donne :

$$M = \mu_0 N_1 N_2 \pi R_2^2 / \ell_1 \quad \Rightarrow \quad k = \frac{R_2}{R_1} \left(\frac{\ell_2}{\ell_1} \right)^{1/2} . \quad (10.23)$$

Dans le calcul de k ci-dessus, nous avons utilisé les formules de bobines infinies pour les auto-inductances, c.-à-d. $L_1 = \mu_0 N_1^2 \pi R_1^2 / \ell_1$ et $L_2 = \mu_0 N_2^2 \pi R_2^2 / \ell_2$. Puisque $\ell_1 > \ell_2$ et $R_1 > R_2$, on constate par inspection de l'expression pour k dans l'éq. (10.23) que $k \leq 1$, mais que $k \rightarrow 1$ dans la limite où les dimensions de la bobine 2 tendent vers celles de la bobine 1, c.-à-d. quand $\ell_2 \rightarrow \ell_1 = \ell$ et $R_2 \rightarrow R_1 = R$.

Dans cette limite de dimensions égales, l'inductance mutuelle devient $M \simeq \mu_0 N_1 N_2 \pi R^2 / \ell \simeq \frac{N_2}{N_1} L_1 \simeq \frac{N_1}{N_2} L_2$, ce qui amène à un facteur de couplage quasi-parfait de $k \simeq 1$. La condition de facteur de couplage parfait, $k \simeq 1$ correspond donc à $M^2 = L_1 L_2$ et si l'on injecte cette condition dans l'expression reliant les forces électromotrices e_2 à e_1 des deux circuits donné dans l'éq. (10.20) on obtient :

$$e_1 = \frac{M}{L_2} e_2 = \frac{N_1}{N_2} e_2 , \quad (10.24)$$

ce qui nous donne le principe de base des transformateurs où le rapport des forces électromotrices de deux bobines en couplage quasi-parfait ne dépend que du rapport du nombre de spires dans les bobines, indépendamment de la fréquence des oscillations de courants dans les circuits.

Dans la pratique, les deux bobines sont rarement l'une à l'intérieur de l'autre, mais plutôt chacune enroulée séparément autour d'un même « noyau » de fer doux comme illustré dans la figure 10.2. Ce noyau a un double emploi. D'abord et principalement, il guide le flux magnétique entre les deux bobines afin d'assurer que le même flux magnétique traverse chacun des deux bobines, assurant ainsi un couplage magnétique optimal de $k \approx 1$ entre les deux bobines. Deuxièmement, les moments dipolaires des atomes de fer doux s'alignent avec le champ magnétique produit par les bobines et renforcent ce champ magnétique de façon à augmenter l'inductance mutuelle et les auto-inductances des bobines, tout en gardant un couplage optimal de $M^2 \simeq L_1 L_2$. Ces commentaires sont destinés à une simple ébauche de principe et nous référons le lecteur intéressé à la vaste littérature sur ce sujet d'importance majeur dans l'acheminement de courant électrique.

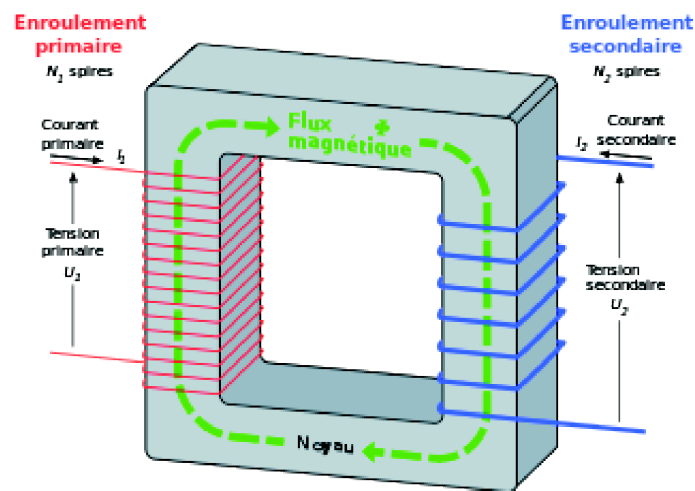


FIGURE 10.2 – Schématique d'un transformateur de tension monophasique.

10.4 Retour sur l'énergie magnétique

10.4.1 L'énergie magnétique des circuits

Dans le chapitre 9 nous avons vu que l'énergie magnétique d'un circuit parcouru par un courant permanent I placé dans un champ magnétique extérieur $\vec{B}_{\text{ext}}$ s'écrit $\mathcal{U}_m = -I\Phi$. Or, un tel circuit produit un flux $\Phi = LI$ à travers lui-même, ce qui semblerait impliquer une énergie magnétique... négative ? Étrange. D'autant plus que nous avons interprété cette énergie comme une énergie potentielle d'interaction entre le circuit et le champ extérieur.

Afin de comprendre ce résultat, il faut se rappeler que ce qu'on appelle l'énergie magnétique potentielle n'est pas l'énergie totale du système. L'énergie magnétique potentielle était dérivée sous la condition qu'un générateur de courant fournissait du travail afin de garder le courant constant dans les circuits. Nous avons déjà rencontré une situation analogue en électrostatique quand il s'agissait de calculer la force sur les armatures d'un condensateur quand ceux-ci sont tenues à potentiel électrique constant par un générateur.

Il nous faut de nouveau raisonner par la méthode des travaux virtuels. Tout effet a nécessité un travail et est donc porteur d'énergie. Un conducteur portant une charge électrique Q (et potentiel V)

possède une énergie électrostatique :

$$\mathcal{E}_e = \frac{Q^2}{2C} = \frac{1}{2}CV^2 = \frac{1}{2}QV ,$$

où C est la capacité du conducteur. Cette énergie est stockée dans (portée par) le champ électrostatique. Nous avons calculé cette énergie en évaluant le travail fourni pour constituer ce réservoir de charges.

Il nous faut faire un raisonnement similaire pour le cas magnétique ici. S'il existe un courant I , c'est qu'un générateur a fourni une puissance $P(t) = ei(t)$ pendant un certain temps. Cela signifie que le circuit (décrit par une self L) a reçu une puissance

$$P_m(t) = -ei(t) = Li \frac{di}{dt} = \frac{d}{dt} \left(\frac{Li^2}{2} \right) ,$$

puisque celui-ci crée un champ magnétique (on néglige ici toute dissipation). Partant d'un courant nul à $t = 0$, on obtient après un temps Δt un courant $I = i(\Delta t)$ et une énergie emmagasinée

$$\boxed{\mathcal{E}_m = \int_0^{\Delta t} P_m dt = \frac{1}{2}LI^2} . \quad (10.25)$$

Cette énergie est stockée dans le champ magnétique qui est créé par un courant d'amplitude I , circulant dans un circuit de self L . Le facteur $1/2$ provient de l'action du circuit sur lui-même. Si l'on prend en compte la dissipation (voir plus bas), on obtient que l'énergie nécessaire à la création d'un courant I (ou la génération du champ $\vec{B}$ associé) doit être supérieure.

Prenons maintenant le cas de deux circuits en interaction. Chacun est parcouru par un courant permanent et engendre ainsi un champ magnétique. L'énergie magnétique totale emmagasinée est alors

$$\begin{aligned} \mathcal{E}_m &= - \int_0^t e_1 I_1 dt - \int_0^t e_2 I_2 dt \\ &= \int_0^t \left[L_1 I_1 \frac{dI_1}{dt} + M I_1 \frac{dI_2}{dt} \right] dt + \int_0^t \left[L_2 I_2 \frac{dI_2}{dt} + M I_2 \frac{dI_1}{dt} \right] dt \\ &= \frac{1}{2} (L_1 I_1^2 + L_2 I_2^2) + M I_1 I_2 . \end{aligned}$$

On voit donc que $\mathcal{E}_m \neq \mathcal{E}_{m,1} + \mathcal{E}_{m,2}$: il y a un troisième terme, correspondant à l'interaction entre les deux circuits.

10.4.2 Forme intégrale de l'énergie magnétique

On se rappelle que l'énergie électrique $\mathcal{E}_e$ peut être écrite dans la forme d'une intégrale volumique :

$$\mathcal{E}_e = \frac{1}{2} \iiint \rho V d\mathcal{V} .$$

Nous avons vu également que $\mathcal{E}_e$ peut être interprétée comme étant stockée par le champ électrique et qu'on peut le calculer via l'intégrale :

$$\mathcal{E}_e = \frac{1}{2} \iiint \vec{D} \cdot \vec{E} d\mathcal{V} = \frac{\epsilon_0}{2} \iiint \epsilon_r \vec{E}^2 d\mathcal{V} .$$

Notre objectif dans cette section est d'obtenir des expressions analogues pour l'énergie magnétique. A partir de l'éq. (10.25) on peut en déduire

$$\mathcal{E}_m = \frac{1}{2}LI^2 = \frac{1}{2}I\Phi_{\text{circ}} .$$

où Φ_{circ} est le flux magnétique à travers le circuit. Il s'avère utile dans ce contexte de faire appel au potentiel vecteur $\vec{A}$ dans l'expression du flux magnétique à travers une surface S s'appuyant sur le circuit :

$$\Phi_{\text{circ}} \equiv \iint_S \vec{B} \cdot d\vec{S} = \iint_S \text{rot} \vec{A} \cdot d\vec{S} = \oint_{\text{circ}} \vec{A} \cdot d\vec{\ell} = \oint_{\text{circ}} \vec{A}(M) \cdot d\vec{OM},$$

où dans la dernière expression, nous avons choisi de dénoter par M le point d'intégration sur le circuit. Mettant ce résultat dans l'expression $\mathcal{E}_m = \frac{1}{2} I \Phi$ et en utilisant l'expression intégrale de l'éq. (7.18) pour le potentiel vecteur $\vec{A}(M)$, on obtient :

$$\begin{aligned} \mathcal{E}_m &= \frac{1}{2} I \oint_{\text{circ}} \iiint \frac{\mu_0}{4\pi} \frac{\vec{j}(P)}{PM} d\mathcal{V}_P \cdot d\vec{OM} = \frac{1}{2} \iiint \oint_{\text{circ}} \frac{\mu_0}{4\pi} \frac{I d\vec{OM}}{PM} \cdot \vec{j}(P) d\mathcal{V}_P \\ &= \frac{1}{2} \iiint \vec{A}(P) \cdot \vec{j}(P) d\mathcal{V}_P, \end{aligned}$$

où nous avons aussi fait appel à l'expression de l'éq. (7.19) pour le champ $\vec{A}$ généré par un circuit filiforme. Ces manipulations nous indiquent donc, que l'énergie magnétique emmagasinée, $\mathcal{E}_m$, peut s'exprimer en générale via l'intégrale volumique :

$$\boxed{\mathcal{E}_m = \frac{1}{2} \iiint \vec{A} \cdot \vec{j} d\mathcal{V}}, \quad (10.26)$$

où l'intégrale s'effectue sur tout l'espace.

Insérant la relation $\text{rot} \vec{H} = \vec{j}$ dans cette équation et en faisant appel à l'identité 11.22 de section 11.5 sur les mathématiques, on obtient :

$$\begin{aligned} \iiint \vec{A} \cdot (\text{rot} \vec{H}) d\mathcal{V} &= \iiint \text{div} (\vec{A} \wedge \vec{H}) d\mathcal{V} + \iiint \vec{H} \cdot (\text{rot} \vec{A}) d\mathcal{V} \\ &= \oint_{\infty} (\vec{A} \wedge \vec{H}) \cdot \hat{n} dS + \iiint \vec{H} \cdot \vec{B} d\mathcal{V} \\ &= \iiint \vec{H} \cdot \vec{B} d\mathcal{V}. \end{aligned} \quad (10.27)$$

où dans la dernière ligne nous avons utilisé le fait que l'intégrale surfacique du terme $\vec{A} \wedge \vec{H} \cdot \hat{n} dS$, évaluée à l'infini soit nulle (puisque pour grand r , $|\vec{A}| \propto 1/r$ et $|\vec{H}| \propto 1/r^2$). Finalement, en se rappelant que pour un milieu linéaire $\vec{H} = \frac{\vec{B}}{\mu_0 \mu_r}$, on obtient une expression locale pour W_m écrite entièrement en termes du champ magnétique :

$$\boxed{\mathcal{E}_m = \frac{1}{2} \iiint \vec{H} \cdot \vec{B} d\mathcal{V} = \frac{1}{2\mu_0} \iiint \frac{\vec{B}^2}{\mu_r} d\mathcal{V}}. \quad (10.28)$$

Un exemple de l'utilité de cette formule se trouve avec l'énergie emmagasinée dans une bobine. En prenant les mêmes conditions que pour la bobine évalué dans la section précédente on a $|\vec{B}_{\text{int}}| = \mu_0 N I / \ell$ à l'intérieur du cylindre de la bobine (de rayon R et de longueur ℓ), et nul partout ailleurs. La formule de l'éq. (10.28) nous donne :

$$\begin{aligned} \mathcal{E}_m &= \frac{1}{2\mu_0} \iiint \vec{B}^2(\vec{r}) d\mathcal{V} = \frac{\vec{B}_{\text{int}}^2}{2\mu_0} \iiint_{\text{cylindre}} d\mathcal{V} = \frac{(\mu_0^2 N^2 I^2 / \ell^2)}{2\mu_0} (\pi R^2 \ell) \\ &= \frac{1}{2} \frac{\mu_0 \pi R^2 N^2}{\ell} I^2. \end{aligned}$$

Le comparaison avec l'équation (10.25) (c.-à-d. $\mathcal{E}_m = \frac{1}{2} L I^2$) permet d'en déduire que $L = (\mu_0 \pi R^2 N^2) / \ell$ en accord avec ce que nous avons trouvé dans l'éq. (10.17) ci-dessus en faisant appel à des considérations de flux magnétique.

10.4.3 Bilan énergétique d'un circuit électrique

D'après la relation établie en électrocinétique, la tension entre deux points A et B d'un circuit vaut

$$V_A - V_B = RI - e ,$$

où e est la fém située entre A et B , R la résistance totale et le courant I circulant de A vers B (n.b. V_A , V_B , I , et e sont tous des fonctions du temps).

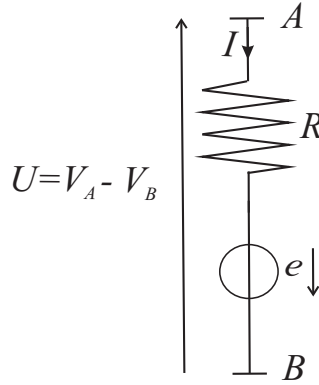


FIGURE 10.3 – Branche de circuit avec résistance, R et force électromotrice, e .

Etant parcouru par un courant, cette branche du circuit va engendrer un champ magnétique, donc produire un flux à travers lui-même qui, si le champ varie, va engendrer une fém (loi de Faraday) $e = -LdI/dt$, et on aura alors,

$$V_A - V_B = RI + L \frac{dI}{dt} .$$

En ajoutant un condensateur de capacité C en série à ces éléments, on obtient un circuit RLC et la différence de potentiel s'écrit :

$$V_A - V_B = RI + L \frac{dI}{dt} + \frac{Q}{C} .$$

On obtient un circuit fermé complet en mettant les extrémités B et A en contact avec $V_A - V_B = 0$. Si aux moins une partie de ce circuit est placée dans un champ magnétique extérieur $\vec{B}_{\text{ext}}$, le champ total sera la somme du champ induit et du champ $\vec{B}_{\text{ext}}$ et l'équation sera :

$$V_A - V_B = 0 = RI + L \frac{dI}{dt} + \frac{Q}{C} + \frac{d\Phi_{\text{ext}}}{dt} = RI + L \frac{dI}{dt} - U_{\text{gen}} ,$$

où $U_{\text{gen}} = -\frac{d\Phi_{\text{ext}}}{dt}$ est la force électromotrice venant de l'extérieur (toute autre source de force électromotrice pourrait en principe être assimilée dans U_{gen} , mais les sources de force électromotrice alternatives sont le plus souvent de type induction).

Ainsi, un circuit composé d'un générateur délivrant une tension U_{gen} , d'une résistance R , d'une bobine de self L et d'un condensateur de capacité C (circuit RLC) aura pour équation

$$\boxed{U_{\text{gen}} = RI + L \frac{dI}{dt} + \frac{Q}{C}} , \quad (10.29)$$

où $I = \frac{dQ}{dt}$ est le courant circulant dans le circuit et Q la charge sur l'une des armatures du condensateur.

La puissance fournie par le générateur se transmet au circuit qui l'utilise alors de la façon suivante :

$$\begin{aligned} P &= U_{\text{gen}} I \\ &= I \left(RI + L \frac{dI}{dt} + \frac{Q}{C} \right) = RI^2 + \frac{d}{dt} \left(\frac{1}{2} LI^2 \right) + \frac{d}{dt} \left(\frac{1}{2} \frac{Q^2}{C} \right) , \end{aligned}$$

c'est-à-dire

$$P = P_J + \frac{d}{dt} (\mathcal{E}_m + \mathcal{E}_e) . \quad (10.30)$$

Une partie de la puissance disponible est donc convertie en chaleur (dissipation par effet Joule), tandis que le reste sert à produire des variations de l'énergie électromagnétique totale du circuit. Dans un circuit « libre » (où $U_{\text{gen}} = 0$), on voit que cette énergie totale diminue au cours du temps, entièrement reconvertie en chaleur.

Chapitre 11

Annexe des mathématiques

11.1 Formules utiles

$$\overrightarrow{\text{grad}}f = \vec{\nabla}f = \hat{x}\frac{\partial}{\partial x}f + \hat{y}\frac{\partial}{\partial y}f + \hat{z}\frac{\partial}{\partial z}f \quad (11.1)$$

$$\text{div} \vec{A} = \vec{\nabla} \cdot \vec{A} = \frac{\partial A_x}{\partial x} + \frac{\partial A_y}{\partial y} + \frac{\partial A_z}{\partial z} \quad (11.2)$$

$$\begin{aligned} \vec{a} \wedge \vec{b} &= \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ a_x & a_y & a_z \\ b_x & b_y & b_z \end{vmatrix} = \begin{bmatrix} a_x \\ a_y \\ a_z \end{bmatrix} \wedge \begin{bmatrix} b_x \\ b_y \\ b_z \end{bmatrix} \\ &= \hat{x}(a_y b_z - a_z b_y) + \hat{y}(a_z b_x - a_x b_z) + \hat{z}(a_x b_y - a_y b_x) \end{aligned} \quad (11.3)$$

$$\begin{aligned} \overrightarrow{\text{rot}} \vec{A} = \vec{\nabla} \wedge \vec{A} &= \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ A_x & A_y & A_z \end{vmatrix} \\ &= \hat{x} \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) + \hat{y} \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) + \hat{z} \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \end{aligned} \quad (11.4)$$

$$\Delta f \equiv \text{div} \overrightarrow{\text{grad}}f = \vec{\nabla} \cdot (\vec{\nabla}f) = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} + \frac{\partial^2 f}{\partial z^2} \quad (11.5)$$

11.2 Analyse vectorielle

11.2.1 Le flux d'un champ

Si dans la notion du flux on ressent communément une idée de mouvement, c'est sans doute parce que l'on parle du flux et du reflux de l'eau. C'est aussi parce que les flux que l'on introduit en physique sont souvent ceux de vecteurs auxquels sont associées des déplacements d'objets ou de fluides.

En fait, le flux est une définition mathématique n'impliquant a priori aucun mouvement. Pour un champ de vecteur $\vec{A}(r)$ et un élément de surface orienté $d\vec{S}$ placé en $\vec{r}$, un élément de flux $d\Phi$ est défini par : $d\Phi = \vec{A}(\vec{r}) \cdot d\vec{S}$

Le flux total à travers une surface S est égal à la somme de ces éléments de flux. L'expression intégrale du flux à travers cette surface est donc

$$\Phi_S = \int d\Phi = \iint_S \vec{A}(\vec{r}) \cdot d\vec{S}.$$

Le champ électrique est un champ de vecteur et, en tant que tel, des éléments de flux lui sont associés.

11.2.2 La divergence du champ

Prenons un petit volume $d\mathcal{V}$ autour d'un point P . On appelle $d\Phi_f$ le flux d'un champ de vecteurs $\vec{A}$ à travers la surface fermée $d\mathcal{S}$ délimitant le volume $d\mathcal{V}$. On appelle la divergence du champ $\vec{A}$, la limite suivante lorsque $d\mathcal{S}$ et $d\mathcal{V}$ tendent vers 0

$$\operatorname{div} \vec{A} \equiv \lim_{d\mathcal{V} \rightarrow 0} \frac{d\Phi_f}{d\mathcal{V}} .$$

On trouvera qu'en coordonnées cartésiennes, la divergence s'exprime

$$\operatorname{div} \vec{A} = \frac{\partial A_x}{\partial x} + \frac{\partial A_y}{\partial y} + \frac{\partial A_z}{\partial z} . \quad (11.6)$$

Un théorème important pour l'électrostatique est le **théorème de d'Ostrogradsky** qui nous apprend que le flux d'un champ $\vec{A}$ à travers une surface fermée S est égale à l'intégrale du divergence du champ $\vec{A}$ sur le volume, $\mathcal{V}$, délimitée par la surface :

$$\oint_S \vec{A} \cdot d\vec{S} = \iiint_{\mathcal{V}} \operatorname{div} \vec{A} d\mathcal{V} . \quad (11.7)$$

11.2.3 Le rotationnel d'un champ

Considérons un champ vectoriel $\vec{A}(x, y, z)$ et ses composantes $A_x(x, y, z)$, $A_y(x, y, z)$ et $A_z(x, y, z)$. On appelle rotationnel de $\vec{A}$, le vecteur **rot** $\vec{A}$ dont les composantes en coordonnées cartésiennes sont

$$\vec{\operatorname{rot} \vec{A}} = \begin{bmatrix} \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) \\ \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) \\ \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \end{bmatrix} .$$

Une façon simple de se rappeler cette formule est l'expression « déterminant » :

$$\begin{aligned} \vec{\operatorname{rot} \vec{A}} &= \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ A_x & A_y & A_z \end{vmatrix} \\ &= \hat{x} \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) + \hat{y} \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) + \hat{z} \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) . \end{aligned}$$

Un théorème important pour le magnétostatique est un des « **théorèmes de Stokes** » qui nous apprend que la circulation de champ $\vec{A}$ le long d'un contour fermée C est égale au flux du rotationnel de $\vec{A}$ à travers toute surface S s'appuyant sur C :

$$\oint_C \vec{A} \cdot d\vec{\ell} = \iint_S \vec{\operatorname{rot} \vec{A}} \cdot d\vec{S} . \quad (11.8)$$

11.2.4 L'opérateur « nabra »

L'opérateur « nabra » noté $\vec{\nabla}$ est largement utilisé dans les ouvrages anglo-saxon. Il s'écrit

$$\vec{\nabla} = \vec{x} \frac{\partial}{\partial x} + \vec{y} \frac{\partial}{\partial y} + \vec{z} \frac{\partial}{\partial z} , \quad (11.9)$$

donc en coordonnées cartésiennes, ses composantes sont

$$\vec{\nabla} = \begin{bmatrix} \frac{\partial}{\partial x} \\ \frac{\partial}{\partial y} \\ \frac{\partial}{\partial z} \end{bmatrix} . \quad (11.10)$$

Avec cette notation, le gradient d'un champ scalaire V s'écrit (en coordonnées cartésiennes)

$$\overrightarrow{\text{grad}} V \equiv \vec{\nabla} V . \quad (11.11)$$

La divergence d'un champ vecteur $\vec{A}$ s'écrit,

$$\text{div } \vec{A} \equiv \vec{\nabla} \cdot \vec{A} , \quad (11.12)$$

et finalement le rotationnel d'un champ vecteur $\vec{A}$ s'écrit

$$\overrightarrow{\text{rot}} \vec{A} \equiv \vec{\nabla} \wedge \vec{A} . \quad (11.13)$$

Notamment,

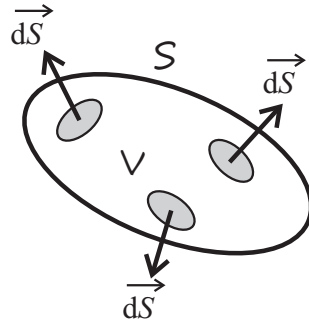
$$\overrightarrow{\text{rot}} \vec{A} = \begin{bmatrix} \frac{\partial}{\partial x} \\ \frac{\partial}{\partial y} \\ \frac{\partial}{\partial z} \end{bmatrix} \wedge \begin{bmatrix} A_x \\ A_y \\ A_z \end{bmatrix} = \begin{bmatrix} \frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \\ \frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \\ \frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \end{bmatrix} . \quad (11.14)$$

Attention : Cette notation peut servir comme aide-mémoire aux expressions de div , $\overrightarrow{\text{grad}}$, et de $\overrightarrow{\text{rot}}$ en coordonnées cartésiennes, mais il ne faut pas oublier qu'elle ne marche qu'en coordonnées cartésiennes. En coordonnées cylindriques et sphériques, il faut revenir aux définitions mathématiques de ces opérateurs afin de retrouver leurs expressions correctes.

11.3 Théorème de Green-Ostrogradsky

11.3.1 Enoncé général

Nous avons montré ci-dessus que le flux de $\vec{E}$ à travers la surface (orientée normale sortante) qui délimite un petit cube est égale à la divergence de $\vec{E}$ multipliée par le volume de ce cube. Ceci se généralise à toute surface fermée et à tout vecteur $\vec{A}$. Soit $\vec{A}$ un champ de vecteur, S une surface fermée dont les éléments de surface $d\vec{S}$ sont orientés dans le sens de la normale sortante, et V le volume délimité par la surface fermée.



On a de façon générale :

$$\oint_S \vec{A} \cdot d\vec{S} = \iiint_V \text{div } \vec{A} dV . \quad (11.15)$$

Avec ce théorème, on peut remonter au théorème intégrale de Gauss, (l'éq. (3.11)) à partir de sa forme local : On remplace $\vec{A}$ par $\vec{E}$ et on invoque la forme locale du théorème de Gauss, $\text{div } \vec{E} = \frac{\rho}{\epsilon_0}$ dans le second membre du théorème. Ensuite, puisque,

$$\iiint_V \rho dV = Q_{\text{int}}$$

ce qui donne,

$$\int_S \vec{E} \cdot d\vec{S} = \frac{Q_{\text{int}}}{\epsilon_0} ,$$

où l'on retrouve l'expression intégrale du théorème de Gauss.

11.3.2 Exemple du théorème d'Ostrogradsky

L'application de ce théorème sur un volume quelconque et pour un champ de vecteur non trivial est très vite compliquée. Nous nous contenterons de le vérifier sur un champ radial de composante :

$$A_x = \frac{x}{a} \quad A_y = \frac{y}{a} \quad A_z = \frac{z}{a} ,$$

et pour un volume délimité par une sphère de rayon R centrée à l'origine. $\vec{A} = \vec{r}/a$ est un vecteur radial de norme r . Le flux de $\vec{A}$ à travers une sphère de rayon R est donc,

$$\oiint_S \vec{A} \cdot d\vec{S} = \frac{R}{a} \times 4\pi R^2 = \frac{4\pi R^3}{a} .$$

On obtient directement que $\text{div } \vec{A} = 3/a$ en tout point de l'espace, et l'intégrale de la divergence de $\vec{A}$ sur le volume de la sphère est :

$$\iiint_V \text{div } \vec{A} dV = \frac{3}{a} \times \frac{4}{3}\pi R^3 = \frac{4\pi R^3}{a} ,$$

ce qui est en accord avec le théorème de Green-Ostrogradsky.

11.4 Théorème de Stokes

11.4.1 Enoncé général

Théorème 4 *Le « Théorème de Stokes » s'écrit : La circulation d'un champ vectoriel $\vec{A}$ le long d'un contour C quelconque est égal au flux du rotationnel de $\vec{A}$ à travers toute surface s'appuyant sur C :*

$$\oint_C \vec{A} \cdot d\vec{\ell} = \iint_S \text{rot } \vec{A} \cdot d\vec{S} \quad (11.16)$$

11.4.2 Exemple

Vérifier le théorème de Stokes pour $\vec{E} = (z^3, 2x, y^2)$ et pour un disque dans un plan $z = Cte$ de rayon R centré au dessus de l'origine.

Réponse : En général

$$\text{rot } \vec{E} = \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ E_x & E_y & E_z \end{vmatrix} = \hat{x} \left(\frac{\partial E_z}{\partial y} - \frac{\partial E_y}{\partial z} \right) + \hat{y} \left(\frac{\partial E_x}{\partial z} - \frac{\partial E_z}{\partial x} \right) + \hat{z} \left(\frac{\partial E_y}{\partial x} - \frac{\partial E_x}{\partial y} \right) .$$

Pour ce problème,

$$\begin{aligned} \text{rot } \vec{E} &= \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ z^3 & 2x & y^2 \end{vmatrix} = \hat{x} \left(\frac{\partial y^2}{\partial y} \right) + \hat{y} \left(\frac{\partial z^3}{\partial z} \right) + \hat{z} \left(\frac{\partial x}{\partial x} \right) \\ &= 2y\hat{x} + 3z^2\hat{y} + 2\hat{z} . \end{aligned}$$

Dans le plan $z = 1$, $d\vec{S} = dS\hat{z} = dx dy \hat{z} = d\rho d\phi \hat{z}$ et $\text{rot } \vec{E} = 2y\hat{x} + 3z^2\hat{y} + 2\hat{z}$. On obtient donc :

$$\iint_{\text{disque}} \text{rot } \vec{E} \cdot d\vec{S} = 2 \iint_{\text{disque}} dS = 2\pi R^2 .$$

Pour calculer la circulation sur le périmètre du disque, il est pratique de paramétrer l'intégrale en termes de l'angle ϕ :

$$\vec{d\ell} = \hat{\phi} R d\phi = R d\phi \left(-\frac{y}{R} \hat{x} + \frac{x}{R} \hat{y} \right) \quad x = R \cos \phi \quad y = R \sin \phi ,$$

et la circulation de $\vec{E}$ sur le périmètre du disque s'exprime donc :

$$\begin{aligned} \oint \vec{E} \cdot \vec{d\ell} &= \int_0^{2\pi} (z^3 \hat{x} + 2x \hat{y} + y^2 \hat{z}) \cdot R d\phi \left(-\frac{y}{R} \hat{x} + \frac{x}{R} \hat{y} \right) \\ &= R \int_0^{2\pi} (z^3 \hat{x} + 2R \cos \phi \hat{y} + R^2 \sin^2 \phi \hat{z}) \cdot (-\sin \phi \hat{x} + \cos \phi \hat{y}) d\phi \\ &= R \int_0^{2\pi} (-z^3 \sin \phi + 2R \cos^2 \phi) d\phi = 2R^2 \int_0^{2\pi} \cos^2 \phi d\phi \\ &= 2\pi R^2 , \end{aligned}$$

où nous avons utilisé la relation :

$$\cos(2\alpha) = \cos^2 \alpha - \sin^2 \alpha \Rightarrow \cos^2 \alpha = \frac{1 + \cos(2\alpha)}{2} .$$

Donc nous avons montré dans notre exemple que :

$$\oint_{\text{cercle}} \vec{E} \cdot \vec{d\ell} = \iint_{\text{disque}} \text{rot} \vec{E} \cdot \vec{dS} = 2\pi R^2 .$$

ce qui vérifie le théorème de Stokes.

11.5 Identités d'analyse vectorielle

1. Montrer que

$$\text{rot}(\overrightarrow{\text{grad}} f) \equiv \vec{0} . \quad (11.17)$$

Démonstration :

$$\begin{aligned} \overrightarrow{\text{grad}} f &= \hat{x} \frac{\partial}{\partial x} f + \hat{y} \frac{\partial}{\partial y} f + \hat{z} \frac{\partial}{\partial z} f \\ \overrightarrow{\text{rot grad}} f &= \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ \frac{\partial f}{\partial x} & \frac{\partial f}{\partial y} & \frac{\partial f}{\partial z} \end{vmatrix} = \hat{x} \left(\frac{\partial^2 f}{\partial y \partial z} - \frac{\partial^2 f}{\partial z \partial y} \right) + \hat{y} \left(\frac{\partial^2 f}{\partial z \partial x} - \frac{\partial^2 f}{\partial x \partial z} \right) + \hat{z} \left(\frac{\partial^2 f}{\partial x \partial y} - \frac{\partial^2 f}{\partial y \partial x} \right) \\ &= \vec{0} \end{aligned}$$

2. Montrer que

$$\text{div}(\overrightarrow{\text{rot A}}) \equiv 0 . \quad (11.18)$$

Démonstration :

$$\begin{aligned} \text{div}(\overrightarrow{\text{rot A}}) &= \frac{\partial}{\partial x} \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) + \frac{\partial}{\partial y} \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) + \frac{\partial}{\partial z} \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \\ &= \left(\frac{\partial^2 A_z}{\partial x \partial y} - \frac{\partial^2 A_y}{\partial x \partial z} \right) + \left(\frac{\partial^2 A_x}{\partial y \partial z} - \frac{\partial^2 A_z}{\partial y \partial x} \right) + \left(\frac{\partial^2 A_y}{\partial z \partial x} - \frac{\partial^2 A_x}{\partial z \partial y} \right) \\ &= 0 \end{aligned}$$

3. Montrer que

$$\operatorname{div}(f \vec{A}) = \overrightarrow{\operatorname{grad} f} \cdot \vec{A} + f \operatorname{div} \vec{A} . \quad (11.19)$$

Démonstration :

$$\begin{aligned} & \frac{\partial}{\partial x} f A_x + \frac{\partial}{\partial y} f A_y + \frac{\partial}{\partial z} f A_z \\ &= A_x \left(\frac{\partial}{\partial x} f \right) + f \left(\frac{\partial}{\partial x} A_x \right) + A_y \left(\frac{\partial}{\partial y} f \right) + f \left(\frac{\partial}{\partial y} A_y \right) + A_z \left(\frac{\partial}{\partial z} f \right) + f \left(\frac{\partial}{\partial z} A_z \right) \\ &= f \left[\left(\frac{\partial}{\partial x} A_x \right) + \left(\frac{\partial}{\partial y} A_y \right) + \left(\frac{\partial}{\partial z} A_z \right) \right] + A_x \left(\frac{\partial}{\partial x} f \right) + A_y \left(\frac{\partial}{\partial y} f \right) + A_z \left(\frac{\partial}{\partial z} f \right) \\ &= f \operatorname{div} \vec{A} + (\overrightarrow{\operatorname{grad} f}) \cdot \vec{A} \end{aligned}$$

4. Montrer que

$$(\overrightarrow{\operatorname{rot} \operatorname{rot}}) \vec{A} = (\overrightarrow{\operatorname{grad} \operatorname{div}} - \Delta) \vec{A} , \quad (11.20)$$

où Δ est le Laplacien.

Démonstration :

$$\begin{aligned} \overrightarrow{\operatorname{rot}} \vec{A} &= \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ A_x & A_y & A_z \end{vmatrix} = \hat{x} \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) + \hat{y} \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) + \hat{z} \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \\ \overrightarrow{\operatorname{rot} \operatorname{rot}} \vec{A} &= \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) & \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) & \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \end{vmatrix} \\ &= \hat{x} \left[\frac{\partial}{\partial y} \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) - \frac{\partial}{\partial z} \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) \right] \\ &\quad + \hat{y} \left[\frac{\partial}{\partial z} \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) - \frac{\partial}{\partial x} \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \right] \\ &\quad + \hat{z} \left[\frac{\partial}{\partial x} \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) - \frac{\partial}{\partial y} \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) \right] \\ &= -\hat{x} \left(\frac{\partial^2 A_x}{\partial y^2} + \frac{\partial^2 A_x}{\partial z^2} \right) - \hat{y} \left(\frac{\partial^2 A_y}{\partial x^2} + \frac{\partial^2 A_y}{\partial z^2} \right) - \hat{z} \left(\frac{\partial^2 A_z}{\partial x^2} + \frac{\partial^2 A_z}{\partial y^2} \right) \\ &\quad + \hat{x} \frac{\partial}{\partial x} \left(\frac{\partial A_y}{\partial y} + \frac{\partial A_z}{\partial z} \right) + \hat{y} \frac{\partial}{\partial y} \left(\frac{\partial A_z}{\partial z} + \frac{\partial A_x}{\partial x} \right) + \hat{z} \frac{\partial}{\partial z} \left(\frac{\partial A_x}{\partial x} + \frac{\partial A_y}{\partial y} \right) \\ &= -\hat{x} \left(\frac{\partial^2 A_x}{\partial x^2} + \frac{\partial^2 A_x}{\partial y^2} + \frac{\partial^2 A_x}{\partial z^2} \right) - \hat{y} \left(\frac{\partial^2 A_y}{\partial x^2} + \frac{\partial^2 A_y}{\partial y^2} + \frac{\partial^2 A_y}{\partial z^2} \right) - \hat{z} \left(\frac{\partial^2 A_z}{\partial x^2} + \frac{\partial^2 A_z}{\partial y^2} + \frac{\partial^2 A_z}{\partial z^2} \right) \\ &\quad + \hat{x} \frac{\partial}{\partial x} \left(\frac{\partial A_x}{\partial x} + \frac{\partial A_y}{\partial y} + \frac{\partial A_z}{\partial z} \right) + \hat{y} \frac{\partial}{\partial y} \left(\frac{\partial A_x}{\partial x} + \frac{\partial A_y}{\partial y} + \frac{\partial A_z}{\partial z} \right) + \hat{z} \frac{\partial}{\partial z} \left(\frac{\partial A_x}{\partial x} + \frac{\partial A_y}{\partial y} + \frac{\partial A_z}{\partial z} \right) , \end{aligned}$$

ce qui s'écrit :

$$(\overrightarrow{\operatorname{rot} \operatorname{rot}}) \vec{A} = (\overrightarrow{\operatorname{grad} \operatorname{div}} - \Delta) \vec{A} .$$

5. Montrer que

$$\overrightarrow{\text{rot}}(f\vec{A}) = (\overrightarrow{\text{grad}}f) \wedge \vec{A} + f\overrightarrow{\text{rot}}\vec{A} \quad (11.21)$$

Démonstration :

$$\begin{aligned} \overrightarrow{\text{rot}}(f\vec{A}) &= \hat{x} \left(\frac{\partial(fA_z)}{\partial y} - \frac{\partial(fA_y)}{\partial z} \right) + \hat{y} \left(\frac{\partial(fA_x)}{\partial z} - \frac{\partial(fA_z)}{\partial x} \right) + \hat{z} \left(\frac{\partial(fA_y)}{\partial x} - \frac{\partial(fA_x)}{\partial y} \right) \\ &= \hat{x} \left(A_z \frac{\partial f}{\partial y} - A_y \frac{\partial f}{\partial z} \right) + \hat{y} \left(A_x \frac{\partial f}{\partial z} - A_z \frac{\partial f}{\partial x} \right) + \hat{z} \left(A_y \frac{\partial f}{\partial x} - A_x \frac{\partial f}{\partial y} \right) \\ &\quad + \hat{x} f \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) + \hat{y} f \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) + \hat{z} f \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \\ &= \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ \frac{\partial f}{\partial x} & \frac{\partial f}{\partial y} & \frac{\partial f}{\partial z} \\ A_x & A_y & A_z \end{vmatrix} + f\overrightarrow{\text{rot}}\vec{A} \\ &= (\overrightarrow{\text{grad}}f) \wedge \vec{A} + f\overrightarrow{\text{rot}}\vec{A} . \end{aligned}$$

6. Montrer que

$$\text{div}(\vec{A} \wedge \vec{B}) = \vec{B} \cdot (\overrightarrow{\text{rot}}\vec{A}) - \vec{A} \cdot (\overrightarrow{\text{rot}}\vec{B}) \quad (11.22)$$

Démonstration :

$$\begin{aligned} \vec{A} \wedge \vec{B} &= \hat{x}(A_y B_z - A_z B_y) + \hat{y}(A_z B_x - A_x B_z) + \hat{z}(A_x B_y - A_y B_x) \\ \text{div}(\vec{A} \wedge \vec{B}) &= \frac{\partial}{\partial x}(A_y B_z - A_z B_y) + \frac{\partial}{\partial y}(A_z B_x - A_x B_z) + \frac{\partial}{\partial z}(A_x B_y - A_y B_x) \\ &= \left(B_z \frac{\partial A_y}{\partial x} + A_y \frac{\partial B_z}{\partial x} - B_y \frac{\partial A_z}{\partial x} - A_z \frac{\partial B_y}{\partial x} \right) \\ &\quad + B_x \frac{\partial A_z}{\partial y} + A_z \frac{\partial B_x}{\partial y} - B_z \frac{\partial A_x}{\partial y} - A_x \frac{\partial B_z}{\partial y} \\ &\quad + B_y \frac{\partial A_x}{\partial z} + A_x \frac{\partial B_y}{\partial z} - B_x \frac{\partial A_y}{\partial z} - A_y \frac{\partial B_x}{\partial z} \\ &= B_x \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) + B_y \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) + B_z \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \\ &\quad - A_x \left(\frac{\partial B_z}{\partial y} - \frac{\partial B_y}{\partial z} \right) - A_y \left(\frac{\partial B_x}{\partial z} - \frac{\partial B_z}{\partial x} \right) - A_z \left(\frac{\partial B_y}{\partial x} - \frac{\partial B_x}{\partial y} \right) \\ &= \vec{B} \cdot (\overrightarrow{\text{rot}}\vec{A}) - \vec{A} \cdot (\overrightarrow{\text{rot}}\vec{B}) . \end{aligned}$$

11.6 Règle BAC - CAB

Il y a une relation célèbre et bien utile quand on rencontre un double produit vectoriel appelé la règle « BAC - CAB ». Il s'écrit :

$$\vec{A} \wedge (\vec{B} \wedge \vec{C}) = \vec{B}(\vec{A} \cdot \vec{C}) - \vec{C}(\vec{A} \cdot \vec{B}) . \quad (11.23)$$

11.7 Exercices d'analyse vectorielle

1. Calculer $\overrightarrow{\text{rot}} \frac{\vec{r}}{r^2}$.

Réponse : On se rappelle que

$$\hat{r} \equiv \frac{\vec{r}}{r} = \frac{x\hat{x} + y\hat{y} + z\hat{z}}{(x^2 + y^2 + z^2)^{1/2}} ,$$

donc on a

$$\begin{aligned}\overrightarrow{\text{rot}} \frac{\hat{\mathbf{r}}}{r^2} &= \overrightarrow{\text{rot}} \frac{\vec{\mathbf{r}}}{r^3} = \begin{vmatrix} \hat{\mathbf{x}} & \hat{\mathbf{y}} & \hat{\mathbf{z}} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ \frac{x}{r^3} & \frac{y}{r^3} & \frac{z}{r^3} \end{vmatrix} \\ &= \hat{\mathbf{x}} \left(\frac{\partial}{\partial y} \frac{z}{r^3} - \frac{\partial}{\partial z} \frac{y}{r^3} \right) + \hat{\mathbf{y}} \left(\frac{\partial}{\partial z} \frac{x}{r^3} - \frac{\partial}{\partial x} \frac{z}{r^3} \right) + \hat{\mathbf{z}} \left(\frac{\partial}{\partial x} \frac{y}{r^3} - \frac{\partial}{\partial y} \frac{x}{r^3} \right) \\ &= -\hat{\mathbf{x}} \frac{3}{r^5} (zy - yz) - \hat{\mathbf{y}} \frac{3}{r^5} (xz - zx) + \hat{\mathbf{z}} \frac{3}{r^5} (yx - xy) = \vec{\mathbf{0}} .\end{aligned}$$

2. Considérer le champ de vecteurs $\vec{\mathbf{E}} = (-y, x, 0) / (x^2 + y^2)$. Calculer son rotationnel et sa circulation sur un cercle de rayon R dans le plan xy centré à l'origine. $\vec{\mathbf{E}}$ dérive-t-il d'un potentiel ?

Réponse : Le rotationnel de $\vec{\mathbf{E}}$

$$\begin{aligned}\overrightarrow{\text{rot}} \vec{\mathbf{E}} &= \begin{vmatrix} \hat{\mathbf{x}} & \hat{\mathbf{y}} & \hat{\mathbf{z}} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ -\frac{y}{(x^2+y^2)} & \frac{x}{(x^2+y^2)} & 0 \end{vmatrix} \\ &= \left(\frac{\partial}{\partial x} \frac{x}{(x^2+y^2)} + \frac{\partial}{\partial y} \frac{y}{(x^2+y^2)} \right) \hat{\mathbf{z}} \\ &= \left(\frac{2}{(x^2+y^2)} - \frac{2x^2}{(x^2+y^2)^2} - \frac{2y^2}{(x^2+y^2)^2} \right) \hat{\mathbf{z}} \\ &= \vec{\mathbf{0}} ,\end{aligned}$$

sauf sur l'axe z ($x \rightarrow y \rightarrow 0$) où il est mal défini. A cause de problème sur l'axe, on a pas le droit d'affirmer que $\vec{\mathbf{E}}$ dérive d'un potentiel, parceque pour cela il faut démontrer que $\overrightarrow{\text{rot}} \vec{\mathbf{E}} = \vec{\mathbf{0}}$ **partout**.

On calcul maintenant sa circulation sur un cercle de rayon R dans le plan xy centré à l'origine. Le déplacement le long d'un cercle de rayon R est

$$d\vec{\ell} = R d\phi \hat{\phi} = R d\phi \left(-\frac{y}{R} \hat{\mathbf{x}} + \frac{x}{R} \hat{\mathbf{y}} \right) ,$$

avec la contrainte

$$x^2 + y^2 = R^2 .$$

La circulation sur ce cercle est donc

$$\begin{aligned}\oint_{\text{cercle}} \vec{\mathbf{E}} \cdot d\vec{\ell} &= \int_0^{2\pi} \frac{-y\hat{\mathbf{x}} + x\hat{\mathbf{y}}}{R^2} \cdot \left(-\frac{y}{R} \hat{\mathbf{x}} + \frac{x}{R} \hat{\mathbf{y}} \right) R d\phi \\ &= \frac{1}{R^2} \int_0^{2\pi} (-y\hat{\mathbf{x}} + x\hat{\mathbf{y}}) \cdot (-y\hat{\mathbf{x}} + x\hat{\mathbf{y}}) d\phi \\ &= \frac{1}{R^2} \int_0^{2\pi} (x^2 + y^2) d\phi = \int_0^{2\pi} d\phi \\ &= 2\pi .\end{aligned}$$

Le théorème de Stokes nous dicte que

$$\iint_S \overrightarrow{\text{rot}} \vec{\mathbf{E}} \cdot d\vec{\mathbf{S}} = \oint_{\text{cercle}} \vec{\mathbf{E}} \cdot d\vec{\ell} ,$$

et on peut donc conclure que,

$$\iint_S \overrightarrow{\text{rot}} \vec{E} \cdot d\vec{S} = 2\pi ,$$

pour n'importe quelle surface S qui s'appuie sur le cercle. On peut en conclure que le rotationnel de $\vec{E}$ n'est pas nul sur l'axe z . et par conséquent le champ $\vec{E}$ ne dérive pas d'un potentiel.