

APPUNTI DI FISICA 2

CdS IN MATEMATICA

E. Smargiassi

Dipartimento di Fisica dell'Università di Trieste

Telefono 040.558.5414

E-mail: smargiassi@units.it

A.A. 2025–2026

Indice

1	Elettrostatica	6
1.1	Carica elettrica e campo elettrico	6
1.1.1	Le leggi di Gauss e di Coulomb	7
1.1.2	Filo carico	8
1.1.3	Dipoli elettrici	9
1.2	Il potenziale elettrostatico	10
1.3	La corrente elettrica	11
1.3.1	La legge di Ohm	12
1.4	Conservazione della carica elettrica	12
1.5	Il campo elettrico nella materia	13
1.5.1	Induzione elettrostatica. Il condensatore	13
1.5.2	I dielettrici	14
2	Magnetostatica	16
2.1	Il campo magnetico	16
2.2	La forza di Lorentz	17
2.2.1	Lo spettrometro di massa	18
2.3	Solenoidalità di B	18
2.4	La legge di Ampère	20
2.4.1	Filo infinito	21
2.4.2	Solenoide	21
2.5	Forze magnetiche su correnti	22
2.5.1	Forza tra fili paralleli	23
2.5.2	Forze e momenti su spire quadrate	23
2.6	Il potenziale vettore	25
2.6.1	L'equazione per il potenziale vettore	26
2.7	La legge di Biot e Savart	27
2.7.1	Filo infinito	27
2.7.2	Spira	28
2.7.3	Solenoide	29
2.8	Cenni sul magnetismo nella materia	30

2.8.1	I tre vettori magnetici	30
2.8.2	Il meccanismo della magnetizzazione	31
3	Elettrodinamica	34
3.1	L'induzione elettromagnetica	34
3.1.1	Campo variabile, circuito fisso	34
3.1.2	Campo fisso, circuito in moto	35
3.1.3	La conservazione dell'energia	36
3.1.4	L'alternatore	37
3.2	Induttanza	38
3.2.1	Autoinduzione	38
3.2.2	Mutua induzione. Il trasformatore	39
3.3	Circuiti oscillanti	40
3.3.1	Oscillazioni libere e risonanza	42
3.3.2	Oscillazioni forzate. Impedenza	42
3.4	La legge di Ampère-Maxwell	44
3.4.1	La corrente di spostamento	45
3.5	Le equazioni di Maxwell	46
3.5.1	L'equazione delle onde	47
3.6	Potenziali elettrodinamici	48
3.6.1	L'equazione d'onda per i potenziali	50
3.6.2	I potenziali ritardati	50
3.7	Lagrangiana ed Hamiltoniana di una particella in campo elettromagnetico	51
3.7.1	Formulazione lagrangiana	51
3.7.2	Formulazione hamiltoniana	53
3.8	Leggi di conservazione per il campo elettromagnetico	54
3.8.1	Energia	54
3.8.2	Quantità di moto	56
4	Onde	59
4.1	Propagazione per onde	59
4.1.1	Onde piane e onde sferiche	59
4.1.2	Onde monocromatiche	61
4.2	Velocità di fase e velocità di gruppo	62
4.2.1	Onde di compressione nei solidi	63
4.3	Effetto Doppler	64
4.4	Le onde elettromagnetiche	66
4.4.1	Onde elettromagnetiche piane	66
4.4.2	Polarizzazione delle onde elettromagnetiche	68
4.5	Energia e quantità di moto di un'onda	70

4.5.1	La pressione di radiazione	71
4.6	La luce come onda elettromagnetica	72
4.6.1	Quantizzazione del campo elettromagnetico: il fotone	72
4.6.2	Lo spettro elettromagnetico	74
4.7	Cenni di ottica geometrica	76
4.7.1	Le onde elettromagnetiche nei mezzi materiali	76
4.7.2	Le leggi dell'ottica geometrica	77
4.7.3	Il principio di Fermat	80
4.8	Interferenza e diffrazione	81
4.8.1	Interferenza	81
4.8.2	Diffrazione	83
5	La teoria della Relatività	85
5.1	Contrasto tra Meccanica Classica ed elettromagnetismo	85
5.1.1	Tentativi di misurare la velocità rispetto all'etere	86
5.1.2	Reazioni ai risultati sperimentali	89
5.2	I postulati della teoria della Relatività Ristretta	90
5.2.1	Velocità della luce e simultaneità	91
5.2.2	Misure di lunghezza trasversali	92
5.2.3	Misure di tempi	93
5.2.4	Misure di lunghezza longitudinali	94
5.2.5	Sincronizzazione degli orologi	95
5.3	Le trasformazioni di Lorentz	95
5.3.1	Le trasformazioni della velocità	97
5.4	Spazio di Minkowski e quadrivettori	98
5.4.1	Gli intervalli quadridimensionali	99
5.4.2	Le trasformazioni generali tra quadrivettori	101
5.5	Dinamica relativistica	103
5.5.1	La lagrangiana di una particella	103
5.5.2	Il quadrimomento. Massa ed energia	104
5.6	Elettrodinamica covariante	109
5.6.1	Il campo magnetico come effetto relativistico	109
5.6.2	Trasformazioni di cariche e correnti	111
5.6.3	Il quadripotenziale	112
5.6.4	Trasformazioni dei campi	112
6	Termodinamica	115
6.1	I sistemi termodinamici	115
6.1.1	Grandezze estensive ed intensive	116
6.1.2	Equilibrio, trasformazioni, cicli. Funzioni di stato	117
6.1.3	Lavoro	118

6.1.4	Principio Zero e temperatura	118
6.2	Il Primo Principio	120
6.2.1	La conservazione dell'energia e il calore	120
6.2.2	Calori specifici e latenti	121
6.2.3	Leggi e trasformazioni dei gas perfetti	123
6.2.4	Gas reali	124
6.2.5	Termodinamica della radiazione	126
6.2.6	Motori termici, frigoriferi, pompe di calore	127
6.3	Il Secondo Principio	129
6.3.1	Formulazioni di Clausius e Kelvin-Planck	129
6.3.2	Entropia	131
6.3.3	Il Terzo Principio	132
6.3.4	Trasformazioni irreversibili e aumento dell'entropia	132
6.3.5	Fasi e passaggi di stato	134
6.4	Cenni all' interpretazione statistica	135
6.4.1	I gas perfetti	135
6.4.2	Equipartizione dell'energia	136
6.4.3	Entropia e disordine	137
7	Elementi di Meccanica Quantistica	140
7.1	Origine	140
7.1.1	Fotoni e onde di materia	142
7.2	L'equazione di Schrödinger	143
7.3	Stati e osservabili	144
7.4	Formulazione generale	145
7.5	Operatori con analogo classico	147
7.6	Commutatori	149
7.6.1	Il Principio di Indeterminazione	150
7.7	L'oscillatore armonico	151
7.8	Il momento angolare	153
7.9	Il modello dell'atomo	155
7.9.1	Atomi idrogenoidi	155
7.9.2	Atomi a molti elettroni	156
7.10	Dinamica quantistica	158
7.10.1	Il propagatore	158
7.10.2	L'equazione di Schrödinger dipendente dal tempo	159
7.10.3	Leggi di conservazione	160
7.11	Spin e sistemi a molte particelle	161
7.11.1	Lo spin	161
7.11.2	Particelle identiche	162
7.11.3	Statistiche quantistiche	164

7.11.4 Fermioni	165
7.11.5 Bosoni	167
A I sistemi di unità di misura dell'elettromagnetismo	168
A.1 Il Sistema Internazionale	168
A.2 Il sistema di Gauss	169
B L'analisi di Fourier	171
B.1 La serie di Fourier	171
B.2 La trasformata di Fourier	173
B.3 Applicazioni	174
B.3.1 Il circuito RLC	174
B.3.2 I potenziali ritardati	175
C La "funzione" δ di Dirac	177
D Relazioni vettoriali	180
D.1 Definizioni e relazioni algebriche	180
D.2 Relazioni differenziali	181
D.3 Relazioni integrali	181

Capitolo 1

Elettrostatica

1.1 Carica elettrica e campo elettrico

A tutti i corpi può essere assegnata una grandezza, eventualmente nulla, chiamata *carica elettrica*, in funzione della quale essi si attraggono o si respingono. Esistono due tipi di carica elettrica, la positiva e la negativa; cariche di segno uguale si respingono, cariche di segno opposto si attraggono. Nel Sistema Internazionale (SI) di misura l'unità di misura della carica è il *Coulomb* (C), la cui definizione precisa è data nel par.2.4. La forza elettrica, nel vuoto, risulta dipendere, tra cariche in quiete relativa, esclusivamente dalla posizione delle cariche e dal loro valore.

La carica elettrica è *quantizzata*, ovvero si presenta sotto forma di “granuli” non ulteriormente divisibili. Tutte i valori delle cariche sono dunque multipli della carica elettrica $-e = -1.6022 \cdot 10^{-19}$ C dell'elettrone¹, od e del protone. Data però la piccolezza di questa carica elementare, possiamo, in ogni problema macroscopico, trattare la carica come fosse una distribuzione continua senza praticamente commettere errori. Per questo definiamo anche una *densità di carica* $\rho(\mathbf{r}, t)$, ovvero la carica per unità di volume, che supporremo una funzione regolare delle coordinate e del tempo.

Gli elettroni sono solitamente i portatori della carica elettrica, in quanto sono le particelle più mobili. Va tenuto presente che in taluni casi, come nelle soluzioni elettrolitiche, anche gli atomi dotati di un elettrone in più od in meno del normale, e dunque carichi (*ioni*), fungono da portatori di carica.

¹ Escludiamo qui la carica dei costituenti di protone e neutrone, i *quark*, che in modulo vale $e/3$, in quanto tali particelle non sono osservabili singolarmente e non sono quindi di interesse per questo corso.

Il problema dello studio dell' interazione elettrica e dei suoi effetti può essere diviso in due parti distinte:

- Si definisce in ogni punto dello spazio e ad ogni tempo un campo, detto *campo elettrico* ed indicato con $\mathbf{E}$, generato dalle cariche elettriche;
- Si studia poi l'effetto di questo campo sul moto delle cariche.

Siccome la forza dipende solo dalla posizione della carica ed è proporzionale al valore della carica stessa, si può *definire* il campo come la forza sentita da una carica di prova q divisa per la q stessa, ovvero

$$\mathbf{E} = \mathbf{F}/q \Leftrightarrow \mathbf{F} = q\mathbf{E} \quad (1.1)$$

dove $\mathbf{E}$ dipende solo dalle altre cariche presenti, supponendo che q sia abbastanza piccola da non influenzarne la disposizione. Il campo $\mathbf{E}$ si misura in N C^{-1} , un' unità detta anche, più spesso, Volt/metro (V m^{-1}). Unendo la 1.1 con la seconda legge di Newton il problema di trovare il moto di cariche soggette ad un campo elettrico dato è in linea di principio risolto.

1.1.1 Le leggi di Gauss e di Coulomb

Rimane da stabilire come si determina il campo. Nel caso di due punti materiali reciprocamente a riposo, dotati di cariche di valore q_1 e q_2 , la forza elettrica che si esercita tra di essi è data dall'espressione²

$$\mathbf{F} = \frac{q_1 q_2}{4\pi\epsilon_0 r^2} \hat{\mathbf{r}} \quad (1.2)$$

detta *legge di Coulomb*. Ne segue che il campo generato da una carica puntiforme q è dato da

$$\mathbf{E} = \frac{q}{4\pi\epsilon_0 r^2} \hat{\mathbf{r}}. \quad (1.3)$$

In queste equazioni compare la costante empirica ϵ_0 , detta *costante dielettrica del vuoto*, che vale $8.854 \cdot 10^{-12} \text{ N}^{-1} \text{ m}^{-2} \text{ C}^2$. Il problema della generazione del campo si riformula esprimendo la legge di Coulomb (e si vedrà che questa espressione è anche più generale) nella forma di

²In queste note usiamo la diffusa convenzione secondo la quale il "cappelletto" posto su di un vettore ($\hat{\mathbf{v}}$) indica che si tratta di un versore, cioè di un vettore di modulo unitario. È chiaro che $\mathbf{v} = v\hat{\mathbf{v}}$ e $\mathbf{v}/v = \hat{\mathbf{v}}$.

Gauss tramite il *flusso del campo elettrico* Φ_E calcolato su di una superficie chiusa:

$$\Phi_E = \oint \mathbf{E} \cdot d\mathbf{S} = q/\epsilon_0 \quad (1.4)$$

dove la carica q è tutta e sola quella contenuta nella regione di spazio delimitata dalla superficie sulla quale si calcola il flusso di $\mathbf{E}$. In forma differenziale l'equazione 1.4 diventa, usando la D.17

$$\nabla \cdot \mathbf{E} = \rho/\epsilon_0. \quad (1.5)$$

Si vede subito che dalla legge di Gauss si ottiene la legge di Coulomb per una carica puntiforme o comunque a simmetria sferica. Si può anche far vedere che se vale la legge di Coulomb vale anche la legge di Gauss; le due quindi, nel caso statico, sono equivalenti.

Si definiscono *linee di campo elettrico* le linee che ad ogni punto hanno come tangente il vettore $\mathbf{E}$. Si tratta di linee orientate: il loro verso è dato dal verso del campo elettrico nel punto tangente. Si vede dalla legge di Gauss che le linee di campo iniziano sulle cariche positive e terminano su quelle negative (si ricordi che le superfici su cui si calcolano i flussi, se chiuse, sono sempre orientate verso l'esterno). In altre parole, iniziano e terminano solo nei punti in cui la divergenza di $\mathbf{E}$ non è nulla. Il termine stesso "divergenza" nasce da qui.

1.1.2 Filo carico

Un caso importante è quello di un filo rettilineo infinito uniformemente carico. Sia λ la densità di carica per unità di lunghezza, allora in una sezione lunga L vi sarà una carica $q = \lambda L$. Il campo, per simmetria, sarà radiale. Applicando la legge di Gauss ad un cilindretto con l'asse principale parallelo al filo, di lunghezza L e avente raggio di base r , si ha allora

$$2\pi r L E(r) = \oint \mathbf{E}(r) \cdot d\mathbf{S} = q/\epsilon_0 = \lambda L/\epsilon_0 \Rightarrow E(r) = \frac{\lambda}{2\pi\epsilon_0 r}. \quad (1.6)$$

La forza per unità di lunghezza tra due fili identici paralleli distanti s vale allora

$$F/L = \int_0^{\lambda L} E(s) dq/L = \int_0^L \frac{\lambda}{2\pi\epsilon_0 s} \lambda dl/L = \frac{\lambda^2}{2\pi\epsilon_0 s}. \quad (1.7)$$

Immaginiamo ora che i fili, che sono composti da circa 10^{25} atomi per metro, presentino uno sbilanciamento di carica elettrica — per esempio un eccesso di elettroni — di una parte su 10^9 . Ovvero, dei circa 10^{25} elettroni che trasportano la corrente, ce ne siano 10^{16} in più di quelli necessari a controbilanciare la carica positiva dei nuclei degli atomi che compongono il filo. Questo sbilanciamento implica un eccesso di carica negativa pari a $\lambda = 1.6 \cdot 10^{-19} \times 10^{16} = 1.6 \cdot 10^{-3}$ C/m. In questo caso la forza per unità di lunghezza tra fili distanti 1 m varrebbe circa $5 \cdot 10^5$ N/m, corrispondente circa al peso di un corpo di massa 50 tonnellate per metro di filo. Basta quindi uno squilibrio piccolissimo tra cariche positive e negative per generare forze enormi.

1.1.3 Dipoli elettrici

Un sistema semplice ma importante è quello detto *dipolo elettrico*. Si definisce *momento di dipolo elettrico* di un sistema di N cariche di carica q_i e poste in $\mathbf{r}_i$, oppure di una distribuzione continua di densità di carica $\rho(\mathbf{r})$, la grandezza

$$\mathbf{p} = \sum_{i=1}^N q_i \mathbf{r}_i = \int \mathbf{r} \rho(\mathbf{r}) d\Omega \quad (1.8)$$

che per due cariche uguali ed opposte distanti $\mathbf{d}$ (vettore che punta dalla carica negativa alla positiva) diventa $\mathbf{p} = q\mathbf{d}$.

Il momento di dipolo è semplicemente il secondo termine di un particolare sviluppo in serie del campo elettrico, detto *sviluppo in multipoli*, in cui ogni termine decade via via più velocemente con la distanza. Il termine di ordine più basso è il termine di “monopolo elettrico”, proporzionale a Q/r^2 , dove Q è la carica totale. Si dimostra che il campo di dipolo vale $[3(\mathbf{p} \cdot \hat{\mathbf{r}})\hat{\mathbf{r}} - \mathbf{p}]/4\pi\epsilon_0 r^3$. Il campo elettrostatico di un sistema neutro pertanto, a grandi distanze, ha sempre quest’ultima forma.

Si dimostra che il momento torcente che agisce su di un dipolo immerso in un campo elettrico uniforme $\mathbf{E}$ vale

$$\boldsymbol{\tau} = \mathbf{p} \times \mathbf{E} \quad (1.9)$$

e l’energia è

$$E = -\mathbf{p} \cdot \mathbf{E} \quad (1.10)$$

e risulta ovviamente minima quando $\mathbf{p}$ ed $\mathbf{E}$ sono paralleli.

1.2 Il potenziale elettrostatico

La legge di Gauss non è in generale sufficiente per determinare univocamente il campo elettrico (nei calcoli precedenti si sono dovute invocare altre considerazioni, principalmente di simmetria). Ci vuole una seconda equazione, e risulta che è sufficiente un'equazione che determini la circuitazione di $\mathbf{E}$ o, equivalentemente tramite la D.18, il suo rotore:

$$\oint \mathbf{E} \cdot d\mathbf{l} = 0 \quad (1.11)$$

$$\nabla \times \mathbf{E} = 0 \quad (1.12)$$

La prima di queste equazioni è equivalente all'affermazione che l'integrale di linea di $\mathbf{E}$, a estremi fissati, è indipendente dal cammino. Queste equazioni dicono che $\mathbf{E}$ è un campo *conservativo*, ovvero derivabile da una funzione scalare detta *potenziale elettrico* ϕ tramite la

$$\phi = - \int_O^P \mathbf{E} \cdot d\mathbf{l} \quad (1.13)$$

$$\mathbf{E} = -\nabla\phi. \quad (1.14)$$

Il potenziale, che si misura in Volt (V), è definito, come l'energia potenziale, a meno di una costante arbitraria, che può essere fissata scegliendo il punto di riferimento O . Solo le differenze di potenziale hanno significato fisico. L'energia potenziale elettrostatica U di una carica q è

$$U = - \int \mathbf{F} \cdot d\mathbf{l} = -q \int \mathbf{E} \cdot d\mathbf{l} = q\phi. \quad (1.15)$$

Potenziale elettrostatico ed energia potenziale elettrostatica sono dunque grandezze strettamente imparentate, ma non sono la stessa cosa. L'energia guadagnata da una carica elementare positiva e che risale la differenza di potenziale di 1 V è detta *elettronvolt* (eV) e vale $1.6022 \cdot 10^{-9}$ J.

Inserendo la definizione del potenziale elettrostatico nella legge di Gauss si ottiene l'equazione per il potenziale elettrostatico:

$$\nabla^2 \phi = -\rho/\epsilon_0 \quad (1.16)$$

L'equazione 1.16 prende il nome di *equazione di Poisson*. La sua versione omogenea, ovvero in assenza di cariche, è detta *equazione di Laplace*, ed è una delle equazioni fondamentali della Fisica Matematica.

Per trovare la soluzione formale della 1.16 partiamo dal fatto che la sua soluzione, nel caso di una carica puntiforme infinitesima dq posta in $\mathbf{r}'$, è la legge di Coulomb:

$$d\phi(\mathbf{r}) = \frac{dq(\mathbf{r}')}{4\pi\epsilon_0|\mathbf{r} - \mathbf{r}'|}.$$

Se allora immaginiamo la distribuzione di carica in questione come la sovrapposizione di tante cariche approssimativamente puntiformi $dq(\mathbf{r}') = \rho(\mathbf{r}')d\Omega'$ possiamo scrivere la soluzione generale dell'equazione 1.16, in virtù della sua linearità, come somma di tutte le soluzioni dei problemi in cui è presente solo la carica $\rho(\mathbf{r}')d\Omega'$:

$$\phi(\mathbf{r}) = \frac{1}{4\pi\epsilon_0} \int \frac{\rho(\mathbf{r}')d\Omega'}{|\mathbf{r} - \mathbf{r}'|}. \quad (1.17)$$

1.3 La corrente elettrica

Un insieme di cariche in moto costituisce un'entità chiamata *corrente elettrica*, la cui intensità, che denota il numero di cariche che passano nell'unità di tempo attraverso una superficie data, è misurata in Ampère (A) = C s⁻¹ e solitamente è indicata con la lettera I . Tale corrente può essere anche descritta in termini di un vettore detto *densità di corrente*

$$\mathbf{j} = \rho \mathbf{v} \quad (1.18)$$

dove ρ è la densità di carica e $\mathbf{v}$ è la velocità delle cariche. Entrambe queste quantità saranno in generale funzione della posizione; la loro dipendenza dal tempo, finché tratteremo di campi elettrici e magnetici statici, sarà ignorata. Sfortunatamente risulta che nei fili elettrici a muoversi sono gli elettroni, carichi negativamente, per cui la velocità effettiva delle cariche è opposta al verso della corrente. Peraltro esistono molti casi in cui si muovono cariche positive o di entrambi i tipi; la conduzione nelle soluzioni elettrolitiche per esempio è di questo tipo. In questi casi la (densità di) corrente totale è la somma delle (densità di) corrente dei diversi tipi di cariche.

Le unità di misura di $\mathbf{j}$ nel SI sono chiaramente A m⁻². La relazione tra corrente e densità di corrente è data dal fatto che la prima è il flusso della seconda:

$$I = \int \mathbf{j} \cdot d\mathbf{S}.$$

1.3.1 La legge di Ohm

Perché si abbia una corrente dal punto 1 al punto 2 bisogna che le cariche siano sottoposte ad una forza che le trascini da 1 a 2. Quasi sempre si tratta di una forza elettrica, dunque deve essere presente una differenza di potenziale $V = \Delta\phi$.

La *legge di Ohm* indica la relazione tra differenza di potenziale e corrente elettrica ai capi di una vasta classe di dispositivi, principalmente fili metallici:

$$V = RI \quad (1.19)$$

dove la costante positiva R è detta *resistenza* e si misura in Ohm (Ω) = $V/A = \text{kg m}^2 \text{C}^{-2} \text{s}^{-1}$. Un dispositivo descritto dalla legge di Ohm con un valore di R apprezzabile è pure detto resistenza o *resistore*. Si tenga però presente che tale legge vale anche per i conduttori.

La legge di Ohm si esprime in forma vettoriale come

$$\mathbf{j} = \sigma \mathbf{E} \quad (1.20)$$

e la *conduttività* (o *conducibilità*) σ , che è una grandezza caratteristica del materiale in cui scorre la corrente, si misura in $\Omega^{-1} \text{m}^{-1} = \text{kg}^{-1} \text{m}^{-3} \text{C}^2 \text{s}$. Il suo reciproco è detto *resistività* e si indica normalmente col simbolo ρ . Le tipiche conducibilità metalliche sono dell'ordine di $10^7 \div 10^8 \Omega^{-1} \text{m}^{-1}$.

La potenza (eventualmente negativa) trasferita dal campo elettrico alle cariche che compongono una corrente I , soggetta ad una differenza di potenziale V , vale dunque

$$W = \frac{dU}{dt} = \frac{dq}{dt} V = VI \quad (1.21)$$

e se vale la legge di Ohm la potenza dissipata P è espressa dalla *legge di Joule*

$$P = VI = RI^2. \quad (1.22)$$

1.4 Conservazione della carica elettrica

Come sappiamo dall'esempio dei due fili esposto in precedenza, basta uno sbilanciamento di carica veramente minimo per generare forze elettrostatiche enormi. Il fatto che queste forze non siano mai osservate è un forte argomento in favore del fatto che la carica totale non cambia mai in

nessuna circostanza. Tutto quello che può fare è di spostarsi da un punto all'altro. In termini quantitativi possiamo scrivere che la carica inclusa in una regione di spazio Ω varia col tempo solo se c'è un flusso di carica in entrata o in uscita attraverso il bordo Σ di tale regione che compensa esattamente questa variazione. In formule, $dq/dt = -I$ (la corrente infatti è positiva quando esce e fa calare q), vale a dire:

$$\frac{d}{dt} \int_{\Omega} \rho \, d\Omega + \oint_{\Sigma} \mathbf{j} \cdot d\mathbf{S} = 0. \quad (1.23)$$

Usando il teorema di Gauss D.17 e sfruttando il fatto che l'eq. 1.23 deve valere per ogni volume Ω ne otteniamo la forma differenziale:

$$\frac{\partial \rho}{\partial t} + \nabla \cdot \mathbf{j} = 0. \quad (1.24)$$

L'equazione 1.23 o 1.24 è detta *equazione di continuità della carica*. Si può osservare che non abbiamo usato alcuna proprietà della carica tranne la sua conservazione (locale); pertanto queste equazioni, previa reinterpretazione dei termini, descrivono la conservazione di qualunque quantità. Per esempio, se ρ = densità di massa e ricordiamo che $\mathbf{j} = \rho \mathbf{v}$, queste equazioni rappresentano la legge di conservazione della massa, e come tali sono fondamentali, per esempio, in fluidodinamica.

1.5 Il campo elettrico nella materia

1.5.1 Induzione elettrostatica. Il condensatore

Quando un materiale è soggetto ad un campo elettrico esterno, le cariche che in esso sono presenti si spostano. Se si tratta di un conduttore lo spostamento è macroscopico, un fenomeno detto *induzione elettrostatica*. Le cariche positive si accumulano sul lato in cui il potenziale è minore e quelle negative sul lato dove è maggiore; all'interno del conduttore il potenziale è uniforme.

Siccome un insieme di cariche genera anche un campo elettrico, oltre che esserne influenzato, ne segue che se abbiamo due conduttori A e B , affacciati l'uno verso l'altro ma isolati tra loro, e vi è per una qualche ragione uno squilibrio di cariche per esempio positive sul lato di A affacciato su B , allora sul lato di B affacciato su A compaiono cariche negative per induzione. Queste cariche — positive su A , negative su B — si attraggono a vicenda e pertanto costituiscono un sistema stabile.

Un dispositivo del genere viene chiamato *condensatore* o *capacitore* e i due corpi A e B sono detti *armature*. Le armature possono poi essere connesse ad altri corpi, dai lati opposti a quelli con cui si affacciano vicendevolmente, in modo tale da far fluire via le altre cariche e da far sì che ciascuna armatura possa rimanere con una carica netta q e $-q$.

Un condensatore è caratterizzato dalla sua *capacità*, ovvero dal rapporto (in modulo) tra la carica q su ciascuna armatura e la differenza di potenziale tra di esse:

$$C = q/V \quad (1.25)$$

capacità che si misura in Farad (F) = $\text{m}^{-2} \text{kg}^{-1} \text{s}^2 \text{C}^2$. L'energia totale accumulata sul condensatore vale

$$U = \frac{1}{2}CV^2 = \frac{q^2}{2C}. \quad (1.26)$$

Il prototipo del condensatore è il *condensatore piano*, composto da due lastre metalliche di area S separate da una distanza d . Se $S \gg d^2$ il campo elettrico può essere preso approssimativamente uniforme e pari a V/d , e la capacità è data da

$$C = \frac{\epsilon_0 S}{d}. \quad (1.27)$$

Per questo condensatore l'energia accumulata U è esprimibile direttamente in funzione del campo elettrico $E = V/d$:

$$U = \frac{1}{2}CV^2 = \frac{1}{2} \frac{\epsilon_0 S}{d} (Ed)^2 = \Omega \frac{\epsilon_0 E^2}{2} \quad (1.28)$$

dove $\Omega = Sd$ è il volume del condensatore.

1.5.2 I dielettrici

Se il corpo soggetto al campo esterno è isolante lo spostamento di cariche al suo interno è microscopico, ovvero su scala atomica, e si rivela nella comparsa di un campo elettrico, sempre diretto in direzione opposta a quello esterno, che si aggiunge a quello esterno. Questo effetto viene detto *polarizzazione dei dielettrici*.

La polarizzazione si può descrivere tramite un vettore $\mathbf{P}$ definito come il momento di dipolo elettrico del materiale per unità di volume e si misura in C m^{-2} . A partire da $\mathbf{P}$ e da $\mathbf{E}$ si definisce un terzo vettore $\mathbf{D} = \epsilon_0 \mathbf{E} + \mathbf{P}$, detto *spostamento elettrico*. L'importanza di questo terzo vettore sta

nel fatto che si può dimostrare che in funzione di esso la legge di Gauss assume una forma più semplice:

$$\oint \mathbf{D} \cdot d\mathbf{S} = q_{libera} \Leftrightarrow \nabla \cdot \mathbf{D} = \rho_{libera} \quad (1.29)$$

dove il pedice “libera” indica che non si considerano le cariche di polarizzazione all’interno del materiale; in pratica, si considerano solo le cariche accessibili direttamente allo sperimentatore.

Si tratta di una relazione del tutto generale; si osservi però che non si può dare una relazione altrettanto generale per il rotore di $\mathbf{D}$ — che non è un campo conservativo, se non nel vuoto — per cui è necessario conoscere la *relazione costitutiva* che lo collega con $\mathbf{E}$.

A questo fine, se si assume che $\mathbf{P}$ ed $\mathbf{E}$ siano paralleli e proporzionali tra loro, si pone

$$\mathbf{D} = \epsilon \mathbf{E} \Leftrightarrow \mathbf{P} = (\epsilon - \epsilon_0) \mathbf{E}$$

dove il nuovo coefficiente ϵ è detto *costante dielettrica* del mezzo e indica di quanto viene ridotto il campo elettrico in presenza di un materiale. Infatti

$$\oint \epsilon \mathbf{E} \cdot d\mathbf{S} = \oint \mathbf{D} \cdot d\mathbf{S} = q_{libera} \Rightarrow \oint \mathbf{E} \cdot d\mathbf{S} = q_{libera} / \epsilon$$

e confrontandola con la 1.4 si vede che l’effetto del materiale sul campo generato da una carica data è di sostituire ϵ_0 con ϵ . Si definisce infine una *costante dielettrica relativa* $\epsilon_r = \epsilon / \epsilon_0$, adimensionale. Si vede immediatamente che, essendo $\epsilon > \epsilon_0$, l’interazione tra due cariche immerse in un dielettrico è diminuita di un fattore ϵ_r , e così pure le differenze di potenziale. Ne segue che la capacità di un condensatore viene *aumentata* di un fattore ϵ_r dalla presenza di un dielettrico.

Capitolo 2

Magnetostatica

2.1 Il campo magnetico

Esiste una seconda interazione tra cariche che dipende dal loro stato di moto, ed in particolare dalla loro velocità. Questa interazione è detta *interazione magnetica*. In completa analogia con il caso elettrostatico, si può dividere il problema dello studio di questa interazione e dei suoi effetti in due parti distinte:

- Si definisce un nuovo campo, detto *campo magnetico* ed indicato con B^1 , generato dalle cariche in moto, ovvero dalle correnti elettriche;
- Si studia poi l'effetto di questo campo sul moto delle cariche.

Una volta risolti entrambi i problemi, e conoscendo il campo elettrico e le altre forze eventualmente presenti, sarà possibile, a meno di difficoltà tecniche a volte notevoli, risolvere completamente il problema del moto di un qualunque sistema di cariche interagenti. In particolare, il problema delle interazioni magnetiche è in generale riconducibile ad un problema di interazione tra correnti.

Il campo magnetico, così come quello elettrico, compare quindi come un intermediario metodologico che non figura nel risultato finale dello studio di un sistema; vedremo però in seguito, nella sezione 3.8, come però questi campi possano vantare un “grado di realtà” più elevato di quanto questa presentazione possa far pensare.

¹Molti testi, specie quelli più vecchi, chiamano B *induzione magnetica*.

2.2 La forza di Lorentz

Il secondo problema posto sopra è in realtà il più semplice dei due. L'esperienza insegna che la forza su di una carica dovuta ad una interazione magnetica è sempre perpendicolare alla velocità della carica. È ragionevole pertanto supporre che tale forza si possa esprimere nel modo seguente:

$$\mathbf{F} = q\mathbf{v} \times \mathbf{B} \quad (2.1)$$

visto che il risultato di un prodotto vettore è un vettore perpendicolare ad entrambi i fattori. In generale, se è presente anche un campo elettrico, la forza su di una carica è data da

$$\mathbf{F} = q(\mathbf{E} + \mathbf{v} \times \mathbf{B}). \quad (2.2)$$

Dalla 2.1 si vede che le unità di $\mathbf{B}$ sono $\text{kg s}^{-1} \text{C}^{-1}$; questa unità è chiamata Tesla (T), o a volte Weber/ m^2 (il Weber è l'unità di misura del flusso magnetico $\Phi_B = \oint \mathbf{B} \cdot d\mathbf{S}$). Il campo magnetico terrestre al suolo ha un'intensità di circa $0.4 \div 0.5 \cdot 10^{-4} \text{ T}$. L'unità 10^{-4} T è chiamata Gauss, ma non fa parte del SI.

È immediato rendersi conto che il campo magnetico, qualunque sia, non trasferisce energia alla carica. Infatti l'energia trasferita per unità di tempo (ovvero la potenza) è

$$W = dE/dt = \mathbf{F} \cdot d\mathbf{r}/dt = \mathbf{F} \cdot \mathbf{v} = 0$$

siccome $\mathbf{F}$ è perpendicolare a $\mathbf{v}$. Pertanto l'energia si conserva anche qualora $\mathbf{B}$ non sia conservativo nel senso dell'analisi vettoriale (ovvero abbia rotazionale nullo od obbedisca ad altre condizioni equivalenti; vedremo subito che in generale non ne soddisfa nessuna).

Infine si noti che si conserva anche la componente del momento angolare $\mathbf{L} = \mathbf{r} \times \mathbf{p}$ lungo il campo $\mathbf{B}$. Infatti, ricordando che $\mathbf{p} = m\mathbf{v}$ e $\mathbf{F} = d\mathbf{p}/dt$,

$$\frac{d\mathbf{L}}{dt} = \frac{d\mathbf{r}}{dt} \times \mathbf{p} + \mathbf{r} \times \frac{d\mathbf{p}}{dt} = m \overbrace{\mathbf{v} \times \mathbf{v}}^{=0} + q \mathbf{r} \times (\mathbf{v} \times \mathbf{B}) \quad (2.3)$$

e l'ultimo termine è perpendicolare a $\mathbf{B}$. Infatti dalla D.3 segue

$$\mathbf{B} \cdot \mathbf{r} \times (\mathbf{v} \times \mathbf{B}) = \mathbf{r} \cdot [(\mathbf{v} \times \mathbf{B}) \times \mathbf{B}] = \mathbf{r} \cdot \mathbf{0} = 0.$$

2.2.1 Lo spettrometro di massa

È elementare risolvere le equazioni del moto per una particella per cui $\mathbf{E} = 0$ e $\mathbf{B} = B\hat{\mathbf{z}}$ costante e uniforme². Posto $\mathbf{r}_0 = 0$, $\mathbf{v}_0 = (v_0, 0, 0)$, l'equazione del moto $m\mathbf{a} = \mathbf{F} = q(\mathbf{v} \times \mathbf{B}) = q(v_x\hat{\mathbf{x}} + v_y\hat{\mathbf{y}} + v_z\hat{\mathbf{z}}) \times B\hat{\mathbf{z}} = qB(v_y\hat{\mathbf{x}} - v_x\hat{\mathbf{y}})$ diventa, scritta per componenti,

$$\begin{aligned}\ddot{x} &= qB\dot{y}/m \\ \ddot{y} &= -qB\dot{x}/m \\ \ddot{z} &= 0\end{aligned}$$

e la sua soluzione è semplice da trovare, usando le normali tecniche di integrazione dei sistemi di equazioni lineari a coefficienti costanti:

$$\begin{aligned}x &= v_0/\omega \sin(\omega t) \\ y &= v_0/\omega [\cos(\omega t) - 1] \\ z &= 0\end{aligned}$$

dove si è introdotta la *frequenza di ciclotrone* $\omega = qB/m$. Si vede che la traiettoria ha equazione cartesiana $x^2 + y^2 + 2v_0y/\omega = 0$ ed è perciò un arco di cerchio di raggio v_0/ω , giacente sul piano xy e con centro in $(0, -v_0/\omega, 0)$. La deflessione delle particelle così lanciate dipende quindi, oltre che dai valori di B e di v_0 , controllati dallo sperimentatore, dal loro rapporto carica/massa q/m . Particelle con valori di q/m diversi vengono deflesse in maniera diversa.

Questo fatto è alla base dello strumento di analisi chimica detto *spettrometro di massa*. Per determinare la composizione chimica quantitativa di un campione basta infatti vaporizzarlo, ionizzare in qualche modo controllato le molecole che si producono, ed inviarle a velocità nota (per esempio facendole passare attraverso una differenza di potenziale elettrico data) in un campo magnetico noto. Le molecole si raggrupperanno sul fianco dello strumento ed in base alle posizioni in cui sono raggruppate si determina q/m e da lì il tipo di molecola in questione. In questo modo si ottengono sensibilità estremamente elevate.

2.3 Solenoidaltà di B

Veniamo al problema di determinare il campo magnetico dal moto delle cariche. Come per il caso del campo elettrico, è sufficiente trovare due

²Si tenga presente sempre la differenza tra i due termini: *costante* si riferisce all'indipendenza dal tempo, *uniforme* all'indipendenza dalla posizione.

equazioni, una per la divergenza di $\mathbf{B}$ ed una per il suo rotore. L'equazione alla divergenza risulta essere semplicissima:

$$\nabla \cdot \mathbf{B} = 0. \quad (2.4)$$

La corrispondente equazione in forma integrale, ovviamente, è

$$\oint \mathbf{B} \cdot d\mathbf{S} = 0 \quad (2.5)$$

ed è equivalente all'affermazione che *il flusso di $\mathbf{B}$ è indipendente dalla superficie su cui lo si calcola*, purché il bordo di tale superficie sia fissato.

Prendiamo infatti due superfici A e B aventi lo stesso bordo λ ; λ potrebbe essere, per fissare le idee, una circonferenza nel piano xy , orientata in senso antiorario per chi vede l'asse z puntare verso l'alto, A la superficie del cerchio incluso e B la semicirconferenza che "sporge" nella direzione dell'asse z positivo. Sia A che B sono, in questo caso, orientate nel senso positivo delle z . Il flusso di $\mathbf{B}$ su queste superfici sia rispettivamente Φ_A e Φ_B .

Consideriamo la superficie $C = A + B$: è una superficie chiusa; però bisogna notare che per convenzione le superfici chiuse sono orientate verso l'esterno. Perciò per calcolare Φ_C bisogna invertire l'orientazione di una delle due superfici, nel nostro esempio la superficie A . Pertanto il flusso su A cambia segno e $\Phi_C = \Phi_B - \Phi_A$, ma siccome il flusso di $\mathbf{B}$ attraverso una qualunque superficie chiusa è nullo, abbiamo $\Phi_C = 0$, dunque $\Phi_B = \Phi_A$, CVD.

Il significato delle equazioni 2.4, 2.5 è chiaro: il campo magnetico non ha sorgenti nel senso in cui le ha il campo elettrico; alternativamente, si può dire che *non esistono cariche magnetiche*, altrimenti dette *monopoli magnetici*. Si tratta di un puro dato sperimentale: non esistono ragioni teoriche o di principio che vietino l'esistenza di monopoli magnetici. Di fatto, però, essi non sono mai stati osservati ed assumeremo dunque la validità delle equazioni 2.4, 2.5.

Le linee di campo magnetico quindi non hanno inizio né fine (si pensi, per confronto, a quelle di campo elettrico, che iniziano sulle cariche positive e terminano su quelle negative, nei punti in cui $\nabla \cdot \mathbf{E} \neq 0$) e possono solo essere o infinite o chiuse su sé stesse. Questo si esprime dicendo che il campo magnetico è *solenoidale*.

Il significato geometrico della solenoidalità rende evidente il fatto esposto più sopra dell'indipendenza dalla superficie del flusso di $\mathbf{B}$: il numero delle linee che passano attraverso una superficie deve essere lo stesso di quelle che escono da un'altra avente lo stesso bordo, e viceversa.

2.4 La legge di Ampère

Siccome l'equazione alla divergenza per il campo magnetico non contiene le sorgenti del campo, che sono le correnti, queste devono essere considerate nell'equazione al rotore, che prende la seguente forma (*equazione di Ampère*):

$$\nabla \times \mathbf{B} = \mu_0 \mathbf{j} \quad (2.6)$$

e la corrispondente equazione integrale è, applicando il teorema di Stokes,

$$\oint \mathbf{B} \cdot d\mathbf{l} = \mu_0 I \quad (2.7)$$

visto che il flusso di $\mathbf{j}$ è I . Si intende che la corrente che compare a secondo membro dell'equazione 2.7 è la corrente totale — ovvero la somma algebrica di tutte le correnti — su di una qualunque superficie che abbia come bordo la curva su cui viene calcolata la circuitazione a primo membro. Si tratta naturalmente di curve e superfici orientate; convenzionalmente si intende orientata la superficie nel senso definito dalla regola della mano destra.

Questa regola dice che, se si osserva il bordo della superficie in modo tale che sia orientato in senso antiorario, allora la superficie è orientata verso l'alto. Così, per esempio, la superficie di un cerchio giacente nel piano xy il cui bordo è percorso in modo tale che un punto vada dall'intercetta con l'asse positivo delle x all'intercetta con l'asse positivo delle y ecc., è orientata verso l'asse z positivo. Un altro modo di definirlo è questo: se si stringe il pugno destro con le dita nel verso in cui è percorso il bordo della superficie, il pollice indica il verso di $\mathbf{S}$.

La costante μ_0 è detta *permeabilità magnetica del vuoto* ed è introdotta per rendere coerenti le unità di misura ai due membri dell'equazione di Ampère. È definita come *esattamente* $4\pi \cdot 10^{-7} \text{ m kg C}^{-2}$. Questa nuova unità di misura è anche nota come Henry/m. Ritroveremo l'Henry (H) più avanti. Si noti come μ_0 non sia una costante empirica come ϵ_0 , che è invece *misurata*. La ragione sta nel fatto che l'unità di corrente, l'Ampère, è appositamente definita in modo tale che μ_0 abbia il valore summenzionato. Infatti nella versione attuale del SI l'Ampère è un'unità fondamentale ed il Coulomb è derivata, vale a dire che 1 C è definito come la quantità di carica che passa in un secondo in un filo percorso da una corrente di 1 A. In altri termini, per definizione 1 A è la corrente che, percorrendo due fili paralleli molto lunghi posti a distanza di 1 m, genera una forza tra di essi pari a $2 \cdot 10^{-7} \text{ N/m}$ (si veda l'Appendice A).

2.4.1 Filo infinito

Applichiamo ora la legge di Ampère al caso di un filo rettilineo infinito percorso da una corrente I . Poniamo l'asse z del nostro sistema di riferimento lungo l'asse del filo. Per simmetria, ci aspettiamo che il campo dipenda solo dalla distanza r dal filo e non dalla coordinata z né dall'angolo di rotazione attorno al filo. Siccome le linee di campo devono essere chiuse, la simmetria ci dice anche che esse devono essere cerchi aventi il centro lungo il filo e giacenti su piani perpendicolari ad esso. Prendendo quindi come percorso d'integrazione uno di questi cerchi, di raggio r , abbiamo che ivi $\mathbf{B}$ è costante in modulo e sempre tangente al cerchio. Perciò $\oint \mathbf{B} \cdot d\mathbf{l} = \int_0^{2\pi r} B \, dl = 2\pi r B$ e la 2.7 diventa

$$B = \frac{\mu_0 I}{2\pi r}. \quad (2.8)$$

Il verso del campo è stabilito, anche in questo caso, dalla regola della mano destra: se si stringe il pugno destro puntando il pollice nel verso della corrente, le altre dita indicano il verso di $\mathbf{B}$. Pertanto, in forma vettoriale, la 2.8 diventa

$$\mathbf{B} = \frac{\mu_0 I}{2\pi r} \hat{\mathbf{z}} \times \hat{\mathbf{r}} = \frac{\mu_0 I}{2\pi r^2} (-y\hat{\mathbf{x}} + x\hat{\mathbf{y}}) \quad (2.9)$$

se si suppone positiva la corrente che scorre nel verso positivo dell'asse z (*Esercizio*: verificare l'identità delle due forme dell'equazione 2.9).

2.4.2 Solenoide

Un solenoide è un avvolgimento di filo sagomato ad elica molto stretta in modo da formare un cilindro, e può essere pensato, con ottima approssimazione, come un insieme di spire circolari (si chiama spira un filo che forma un circuito chiuso) collegate in serie. Tutte queste spire quindi sono percorse dalla stessa corrente I . Il caso più semplice da trattare, e che nel magnetismo svolge lo stesso ruolo illustrativo che il condensatore infinito a facce piane e parallele svolge in elettrostatica, è quello del solenoide infinito i cui avvolgimenti sono uniformemente distribuiti lungo la superficie del cilindro con densità lungo l'asse di $n = N/L$ spire per unità di lunghezza.

È semplice calcolare il campo all'interno di un tale solenoide ideale. Per simmetria il campo $\mathbf{B}$ sarà parallelo all'asse del cilindro, e siccome le linee di campo non possono uscire possiamo supporre che $\mathbf{B} = 0$ al di fuori del solenoide. Prendiamo allora un circuito $ABCD$ tale che due lati — AB e

CD — siano lunghi L e paralleli all'asse del solenoide, AB all'interno e CD all'esterno, e gli altri due — BC e DA — perpendicolari. Allora

$$\oint \mathbf{B} \cdot d\mathbf{l} = \left(\int_A^B + \int_B^C + \int_C^D + \int_D^A \right) \mathbf{B} \cdot d\mathbf{l} =$$

$$= \int_0^L B dl + \int_B^C B \cos(\pi/2) dl + \int_C^D 0 dl + \int_D^A B \cos(\pi/2) dl = BL$$

per cui dalla 2.7, tenendo conto che in un tratto lungo L ci sono $N = nL$ spire ciascuna percorsa dalla corrente I , segue

$$B = \mu_0 I_{totale} / L = \mu_0 NI / L = \mu_0 nI. \quad (2.10)$$

Un solenoide quindi è un buon modo per creare campi magnetici intensi senza dover ricorrere a correnti eccessivamente elevate: è sufficiente aumentare la densità degli avvolgimenti. Questo è il prototipo degli *elettromagneti*.

2.5 Forze magnetiche su correnti

Siccome la forza magnetica si esercita su cariche in moto, e la corrente non è altro che un insieme di cariche in moto, possiamo concludere che un filo percorso da corrente ed immerso in un campo magnetico è sottoposto ad una forza. Si noti che questa forza esiste anche nel caso, quello più comune, in cui il filo è elettricamente neutro.

Per calcolare tale forza partiamo, com'è ovvio, dalla forza di Lorentz 2.1. Per applicarla, dobbiamo in qualche modo collegare i parametri macroscopici del filo — la sua lunghezza e la corrente I che lo percorre — con la velocità e la carica che figurano nelle equazioni 2.1 e 1.18.

Notiamo che $\mathbf{j}$ ha in ogni punto la stessa direzione del filo, e che questa direzione (e verso) è data dal versore tangente al filo $\hat{\mathbf{l}} = d\mathbf{l}/dl$. Inoltre possiamo scrivere per il modulo di $\mathbf{j} = I/S$ se S è la sezione trasversale del filo e $\mathbf{j}$ è uniforme in esso, visto che I è il flusso di $\mathbf{j}$. Dunque $\mathbf{j} d\Omega = \hat{\mathbf{l}} I d\Omega / S = I d\mathbf{l}$, visto che $d\Omega = S dl$ è il volume di un pezzo di filo lungo dl . Applicando la 1.18 troviamo quindi che la forza è

$$\mathbf{F} = \sum_i q_i \mathbf{v}_i \times \mathbf{B} = \int \rho \mathbf{v} \times \mathbf{B} d\Omega = \int \mathbf{j} \times \mathbf{B} d\Omega = I \int d\mathbf{l} \times \mathbf{B}. \quad (2.11)$$

2.5.1 Forza tra fili paralleli

Combinando l'equazione 2.11 con la 2.9 otteniamo la forza per unità di lunghezza tra due fili infiniti paralleli, posti a distanza s e giacenti in direzione z , percorsi da correnti I_1 e I_2 . Per la forza sul filo 1 si ha:

$$\mathbf{F}_1/L = I_1/L \int d\mathbf{l} \times \mathbf{B}_2 = \frac{\mu_0 I_1 I_2}{2\pi s L} \int d\mathbf{l} \times (\hat{\mathbf{z}} \times \hat{\mathbf{s}})$$

dove l'integrale viene calcolato su di un tratto del filo 1 di lunghezza L e $\hat{\mathbf{s}}$ è un versore che punta dal filo 2 al filo 1. È immediato rendersi conto che il doppio prodotto vettore, nelle condizioni date, genera un vettore di modulo dl , la cui direzione è perpendicolare ad entrambi i fili e che passa per essi. Il verso dipende dalla direzione delle correnti: se sono concordi punta verso il filo 2, se sono discordi punta nel verso opposto.

È chiaro che lo stesso risultato si ottiene scambiando i ruoli dei due fili, cioè gli indici 1 e 2. Possiamo perciò concludere che *due fili percorsi da corrente si attraggono se le correnti sono concordi, si respingono se sono discordi*. Il modulo della forza per unità di lunghezza, uguale sia per il filo 1 che per il 2, è

$$F/L = \frac{\mu_0 I_1 I_2}{2\pi s}. \quad (2.12)$$

Prendiamo un caso abbastanza tipico: $I_1 = I_2 = 1$ A e $s = 0.1$ m. In questo caso abbiamo che $F/L = 2 \cdot 10^{-6}$ N/m, una forza molto piccola: infatti la forza su di un metro di filo in questo caso corrisponde al peso di un corpo di circa 0.2 milligrammi. Il filo stesso pesa ben di più. Nella sezione 1.1.2 si è visto che la forza elettrica per unità di lunghezza tra fili in cui si abbia un minimo sbilanciamento di carica (una parte su un miliardo) vale $5 \cdot 10^5$ N/m. Una differenza di undici ordini di grandezza rispetto al caso magnetico presentato qui!

Possiamo concluderne che le interazioni magnetiche sono in realtà piuttosto deboli, e che sono osservabili più facilmente rispetto alle interazioni elettriche per il solo fatto che queste ultime, essendo generate da cariche di due segni opposti, sono schermate facilmente (ovvero, è difficile trovare una carica positiva/negativa libera perchè viene quasi sempre neutralizzata dall' accorrere di cariche di segno opposto). Il campo magnetico invece, essendo generato da correnti, non permette una schermatura così facile.

2.5.2 Forze e momenti su spire quadrate

Consideriamo una spira rigida in foggia di quadrato $ABCD$ di lato a , e percorsa da una corrente I . Sia la spira immersa in un campo magnetico

uniforme $\mathbf{B} = B\hat{x}$ e supponiamola orientata in modo tale che i lati AB e CD siano paralleli all'asse z mentre i lati BC e DA vi siano perpendicolari e formino un angolo φ con l'asse x , e pertanto uno di $\theta = \pi/2 - \varphi$ con l'asse y . Il vettore $\mathbf{S} = a^2\hat{n}$ che descrive la superficie orientata della spira è quindi anch'esso perpendicolare all'asse z e forma un angolo $-\pi/2 + \varphi = -\theta$ (o $\pi/2 + \varphi = \pi - \theta$, a seconda del senso di circolazione della corrente: qui si considererà $-\theta$) con l'asse x .

È facile vedere che la forza sul lato BC è perpendicolare al lato ed è uguale ed opposta a quella sul lato DA ; queste forze pertanto tendono solo a deformare la spira, che però abbiamo supposto rigida, e pertanto si annullano semplicemente tra loro. Diversa è la situazione per le forze sui lati AB e CD : queste forze formano una coppia di braccio $a \sin \theta$ e di intensità IaB . Sulla spira pertanto agisce una forza di risultante nulla ma un momento pari a

$$\tau = a^2 IB \sin \theta \hat{z} = \mathbf{IS} \times \mathbf{B} \quad (2.13)$$

e che pertanto tende a rendere paralleli $\mathbf{S}$ e $\mathbf{B}$ (il momento della forza è nullo quando il piano della spira è perpendicolare a $\mathbf{B}$). Se la spira è composta da N avvolgimenti invece che da uno, il valore di τ naturalmente deve essere moltiplicato per N .

Se all'istante in cui $\mathbf{S}$ e $\mathbf{B}$ sono all'incirca paralleli invertiamo il senso della corrente o il verso del campo magnetico, la spira viene sospinta a fare un altro mezzo giro. Ripetendo indefinitamente il procedimento abbiamo un dispositivo che permette di fornire un momento torcente in maniera continua. Questo dispositivo è alla base di ogni *motore elettrico*.

Il momento magnetico

Definiamo *momento di dipolo magnetico*, o semplicemente *momento magnetico* il vettore

$$\mathbf{m} = \frac{1}{2} \int \mathbf{r} \times \mathbf{j} \, d\Omega. \quad (2.14)$$

Il momento magnetico, che si misura in $\text{m}^2 \text{C s}^{-1}$, nel caso di spire piane si riduce ad un vettore diretto lungo l'asse $\hat{n}$ della spira, con verso dato dalla regola della mano destra, e di modulo pari al prodotto dell'area S per la corrente; infatti

$$\mathbf{m} = \frac{I}{2} \int \mathbf{r} \times d\mathbf{l} = IS\hat{n} = \mathbf{IS} \quad (2.15)$$

visto che il modulo di $\frac{1}{2}\mathbf{r} \times d\mathbf{l}$ è l'areola spazzata dal vettore $\mathbf{r}$ quando la sua estremità si sposta di $d\mathbf{l}$ (*Esercizio: verificarlo*) e la direzione e il verso sono stabiliti dalla regola della mano destra.

Utilizzando questa definizione possiamo riscrivere la 2.13 come

$$\boldsymbol{\tau} = \mathbf{m} \times \mathbf{B}. \quad (2.16)$$

Integrando la 2.13 rispetto a θ possiamo anche scrivere che l'energia di una spira in un campo magnetico uniforme è

$$E = \int \tau_z d\theta = -mB \cos \theta = -\mathbf{m} \cdot \mathbf{B} \quad (2.17)$$

e risulta ovviamente minima quando $\mathbf{m}$ e $\mathbf{B}$ sono paralleli. Si confrontino le 2.16 e 2.17 con le corrispondenti, ed analoghe, relazioni valide per dipoli elettrici in un campo elettrico.

Il concetto di momento magnetico risulta utile in molti casi; lo vedremo, per esempio, nei paragrafi 2.7 e 2.8.

2.6 Il potenziale vettore

Dall'equazione di Ampère segue che il campo magnetico non può essere derivato da un potenziale nello stesso senso in cui il campo elettrico è derivato dal potenziale scalare. Questo sarebbe infatti possibile solo se il rotore di $\mathbf{B}$ fosse ovunque nullo. L'equazione 2.4 però ci dice che è sempre possibile trovare un campo vettoriale $\mathbf{A}$, detto *potenziale vettore*, e misurato in T m, tale che

$$\nabla \times \mathbf{A} = \mathbf{B}. \quad (2.18)$$

Infatti la divergenza di un rotore è sempre nulla. Si noti che l'equazione 2.18, essendo un'equazione differenziale a derivate parziali, determina $\mathbf{A}$ soltanto a meno di una funzione arbitraria, così come il potenziale elettrostatico ϕ è determinato a meno di una costante arbitraria. Tale funzione arbitraria, in considerazione del fatto che $\nabla \times \nabla \Lambda = 0$, può essere scritta come $\nabla \Lambda$, dove Λ è una qualunque funzione scalare delle coordinate e del tempo. Sfrutteremo spesso questa possibilità.

A volte è possibile trovare $\mathbf{A}$, noto $\mathbf{B}$, mediante una semplice ispezione della forma di quest'ultimo. Per esempio si verifica immediatamente che un campo magnetico uniforme diretto lungo l'asse z , $\mathbf{B} = B\hat{\mathbf{z}}$, può essere derivato da un potenziale vettore $\mathbf{A} = B/2(-y\hat{\mathbf{x}} + x\hat{\mathbf{y}})$, ma anche da $\mathbf{A}' = B(x\hat{\mathbf{y}} + z\hat{\mathbf{z}}/2)$ (*Esercizio: verificare che $\mathbf{A} - \mathbf{A}'$ è effettivamente*

il gradiente di una funzione Λ e dare l'espressione esplicita di Λ). Si noti che in questo caso $\nabla \cdot \mathbf{A} = 0$, ma che $\nabla \cdot \mathbf{A}' \neq 0$.

2.6.1 L'equazione per il potenziale vettore

È naturalmente importante trovare un'equazione che permetta di determinare $\mathbf{A}$, note le correnti, nel caso generale. Questo lo si ottiene partendo dall'equazione di Ampère in forma differenziale. Inserendo l'equazione 2.18 nella 2.6 si ottiene

$$\nabla \times \mathbf{B} = \nabla \times \nabla \times \mathbf{A} = \mu_0 \mathbf{j}$$

Utilizzando l'identità vettoriale D.13 e sfruttando l'arbitrarietà di $\mathbf{A}$ con l'imporre – si dimostra che è sempre possibile – che $\nabla \cdot \mathbf{A} = 0$, otteniamo infine l'equazione cercata:

$$\nabla^2 \mathbf{A} = -\mu_0 \mathbf{j} \quad (2.19)$$

La soluzione formale dell'equazione 2.19 può essere calcolata esplicitamente e rigorosamente ma qui ci limitiamo ad un semplice argomento per analogia. Possiamo infatti notare che l'equazione 1.16 è assai simile alla 2.19 a parte il fatto di essere una relazione vettoriale e non scalare. Confrontando le equazioni 1.16, 1.17 e 2.19 possiamo quindi immaginare che la soluzione di quest'ultima sia:

$$\mathbf{A}(\mathbf{r}) = \frac{\mu_0}{4\pi} \int \frac{\mathbf{j}(\mathbf{r}') d\Omega'}{|\mathbf{r} - \mathbf{r}'|} \quad (2.20)$$

che in effetti risulta essere l'espressione corretta.

L'equazione 2.20 è del tutto generale ma, come capita spesso in questi casi, risulta scomoda da applicare direttamente. È in generale opportuno trovare delle formule più pratiche da applicare in casi particolari. Uno di questi, particolarmente importante, è quello in cui la densità di corrente è concentrata entro fili conduttori, in quanto è questo il caso di interesse per la teoria dei circuiti elettrici. Se ricordiamo che per un filo $\mathbf{j} d\Omega = I d\mathbf{l}$, dalla 2.20 si ha

$$\mathbf{A}(\mathbf{r}) = \frac{\mu_0 I}{4\pi} \int \frac{d\mathbf{l}}{|\mathbf{r} - \mathbf{r}'|}. \quad (2.21)$$

Non ci si lasci confondere dalla forma dell'equazione 2.21: l'integrale è calcolato rispetto ad $\mathbf{r}'$, che indica punti sul filo, ed $\mathbf{l} = \mathbf{l}(\mathbf{r}')$. Inoltre, siccome stiamo trattando casi stazionari, l'equazione è valida solo per circuiti

chiusi: se il circuito fosse aperto, ad una estremità avremmo uno svuotamento di carica ed all'altra un accumulo, contrariamente all'ipotesi di stazionarietà.

2.7 La legge di Biot e Savart

Già l'equazione 2.21 è una semplificazione importante, ma è possibile ottenere direttamente un'espressione per $\mathbf{B}$ che spesso risulta utile. Prendendo il rotore della 2.21 ed usando la 2.18 si ha infatti

$$\mathbf{B}(\mathbf{r}) = \frac{\mu_0 I}{4\pi} \int \nabla \times \frac{d\mathbf{l}}{|\mathbf{r} - \mathbf{r}'|}$$

visto che le derivate che compongono il rotore agiscono su funzioni del "punto campo" $\mathbf{r}$ e non del "punto sorgente" $\mathbf{r}'$ (pertanto non agiscono nemmeno su $d\mathbf{l}$) rispetto a cui è calcolato l'integrale.

Per calcolare il secondo membro osserviamo che, posto $\mathbf{s} = \mathbf{r} - \mathbf{r}'$, si ha $\partial(1/s)/\partial x = \partial s/\partial x \partial(1/s)/\partial s = -s_x/s^3$ ed analogamente per y e z . Allora la componente x dell'integrando diventa

$$\nabla \times d\mathbf{l}/s|_x = \begin{vmatrix} \hat{\mathbf{x}} & \hat{\mathbf{y}} & \hat{\mathbf{z}} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ dl_x/s & dl_y/s & dl_z/s \end{vmatrix}_x = dl_y s_z/s^3 - dl_z s_y/s^3 = d\mathbf{l} \times \mathbf{s}/s^3|_x.$$

Al posto del rotore rimane dunque un semplice prodotto vettore tra $d\mathbf{l}$ e $\mathbf{s}$:

$$\mathbf{B}(\mathbf{r}) = \frac{\mu_0 I}{4\pi} \int \frac{d\mathbf{l} \times \mathbf{s}}{s^3}. \quad (2.22)$$

Questa equazione è nota come *legge di Biot e Savart*³.

2.7.1 Filo infinito

Utilizziamo la legge di Biot e Savart per ricalcolare il campo magnetico generato dal solito filo rettilineo infinito percorso da corrente. Data la simmetria del sistema, se prendiamo il filo come giacente lungo l'asse z , possiamo senza perdere di generalità considerare un punto campo di coordinate $(x, y, 0) = x\hat{\mathbf{x}} + y\hat{\mathbf{y}}$ e ad una distanza $r = \sqrt{x^2 + y^2}$ dal filo.

³A volte l'espressione 2.22 è indicata col nome di *formula di (Ampère-)Laplace*, o anche come *prima formula di Laplace*, ed il nome di Biot e Savart è riservato alla 2.8 o alla 2.9.

Allora la distanza s tra il punto campo ed il punto sorgente si esprime in termini di r e dell'angolo θ fra s e $\hat{z}$ come $s \sin \theta = r$, e la posizione lungo il filo come $l = -s \cos \theta = -r \cot \theta$ (il segno negativo viene dal fatto di prendere negative le coordinate al di sotto del piano $z = 0$). Perciò $d\mathbf{l} = \hat{z} r d\theta / \sin^2 \theta$. Si può dare anche l'espressione del vettore $s = x\hat{x} + y\hat{y} - \hat{z}l$, da cui $d\mathbf{l} \times s = \hat{z} \times s r d\theta / \sin^2 \theta = (-y\hat{x} + x\hat{y}) r d\theta / \sin^2 \theta$. L'integrale 2.22 diventa dunque

$$\mathbf{B}(\mathbf{r}) = (-y\hat{x} + x\hat{y}) \frac{\mu_0 I}{4\pi r^2} \int_0^\pi \sin \theta d\theta = (-y\hat{x} + x\hat{y}) \frac{\mu_0 I}{2\pi r^2}$$

che coincide con la 2.9.

2.7.2 Spira

Consideriamo una spira circolare di raggio a che giaccia nel piano yz , percorsa da una corrente I . Il calcolo del campo in un punto generico è possibile ma complicato; ci limitiamo al calcolo lungo l'asse della spira, l'asse x , ipotizzando che la corrente scorra in senso orario se vista da un osservatore con la testa in direzione dell'asse x positivo.

La simmetria indica subito che il campo $\mathbf{B}$ su tale asse deve avere direzione $\hat{x}$; infatti ogni contributo perpendicolare ad essa proveniente da un tratto di spira viene cancellato da un contributo uguale ed opposto da un tratto di spira diametralmente opposto.

Per calcolare il campo parallelo ad $\hat{x}$ nel punto x , vediamo per prima cosa che $|d\mathbf{l} \times s| = s dl = \sqrt{a^2 + x^2} dl$, e con direzione angolata di un angolo tale che $s \cos \theta = a$, dove θ è anche l'angolazione rispetto all'asse x . Perciò $\mathbf{B}$ vale

$$\mathbf{B}(x) = \frac{\mu_0 I}{4\pi} \int \frac{\cos \theta dl}{s^2} \hat{x} = \frac{\mu_0 I a}{4\pi s^3} \int dl \hat{x} = \frac{\mu_0 I a^2}{2(a^2 + x^2)^{3/2}} \hat{x} = \frac{\mu_0 \mathbf{m}}{2\pi(a^2 + x^2)^{3/2}} \quad (2.23)$$

dove $\mathbf{m}$ è il momento magnetico della spira. A grande distanza, a è trascurabile rispetto a x e si ottiene

$$\mathbf{B}(x) = \frac{\mu_0 \mathbf{m}}{2\pi x^3}. \quad (2.24)$$

Si vede che il campo a grandi distanze dipende esclusivamente dal valore di $\mathbf{m}$ e non separatamente dai dettagli della geometria o dal valore di I (si dimostra che in una direzione generica $\hat{r}$ il campo vale $\mu_0[3(\mathbf{m} \cdot \hat{r})\hat{r} - \mathbf{m}]/4\pi r^3$, in perfetta analogia col caso elettrostatico).

Questo è solo un esempio dello sviluppo del campo magnetico in multipoli (in completa analogia col caso elettrico), uno sviluppo in serie in cui ogni termine dello sviluppo decade via via più velocemente con la distanza ed in cui il termine di dipolo è semplicemente il primo termine. A differenza del caso elettrico, infatti, in quello magnetico il termine di monopolo non esiste. Pertanto, eccettuato il caso in cui $\mathbf{m} = 0$, il campo magnetico di un sistema di correnti stazionarie a grandi distanze tende asintoticamente al campo di dipolo. Questo si può verificare, per esempio, calcolando i campi generati da spire piane di altra forma (*Esercizio*: spira quadrata) e verificando che a grandi distanze hanno tutti esattamente la forma 2.24.

2.7.3 Solenoide

Come abbiamo già visto, un solenoide può essere visto come una serie di spire identiche coassiali collegate in parallelo. Il campo al suo interno è quindi calcolabile come somma dei campi generati da tutte le spire. Posto parallelo all'asse x l'asse del solenoide e presa $x = 0$ la coordinata di un punto lungo l'asse posto a distanza L dall'estremità sinistra ed L' da quella destra, si vede subito che il campo in quel punto è la somma di tanti campi paralleli all'asse, dunque è anch'esso parallelo all'asse. Il suo modulo si calcola sommando i contributi di tutte le spire. Ricordando che ci sono $n = N/L$ spire per unità di lunghezza, vediamo che in un tratto dx ci sono $dN = ndx$ spire, ognuna delle quali fornisce un contributo dato dalla 2.23, per cui:

$$B = \int_{-L}^{L'} \frac{\mu_0 I a^2}{2(a^2 + x^2)^{3/2}} n dx.$$

L'integrale si calcola con la sostituzione $x = a \tan t$:

$$\begin{aligned} B &= \frac{\mu_0 n I}{2} \int_{-\arctan(L/a)}^{\arctan(L'/a)} \frac{1 + \tan^2 t}{(1 + \tan^2 t)^{3/2}} dt = \frac{\mu_0 n I}{2} \int_{-\arctan(L/a)}^{\arctan(L'/a)} \cos t dt \\ &= \frac{\mu_0 n I}{2} \sin t \Big|_{-\arctan(L/a)}^{\arctan(L'/a)} = \frac{\mu_0 n I}{2} \left[\frac{L/a}{\sqrt{1 + (L/a)^2}} + \frac{L'/a}{\sqrt{1 + (L'/a)^2}} \right]. \end{aligned}$$

Per $L, L' \rightarrow \infty$ il risultato coincide, come deve, con la 2.10. Si vede anche che il campo all'estremità di un solenoide semiinfinito (es. $L \rightarrow \infty, L' = 0$) è la metà di quello all'interno di uno infinito. (*Esercizio*: trovare il termine correttivo principale, nel caso di un solenoide finito con $L = L'$, al campo di un solenoide infinito).

Si sarà notato che laddove si è abbastanza fortunati da poter invocare considerazioni di simmetria l'utilizzo della legge di Ampère permette calcoli più rapidi e semplici.

2.8 Cenni sul magnetismo nella materia

2.8.1 I tre vettori magnetici

Quando un materiale è soggetto ad un campo magnetico esterno — e, come vedremo, a volte anche quando non lo è — si ha un effetto simile a quello della polarizzazione dei dielettrici nel caso elettrostatico, ovvero il materiale stesso genera un campo magnetico che si aggiunge a quello esterno. Questo effetto viene detto *magnetizzazione*. A differenza però del caso elettrostatico, in cui il campo addizionale è sempre diretto in direzione opposta a quello esterno, nel caso magnetostatico la magnetizzazione può creare un campo sia opposto a quello esterno, ed in questo caso si parla di *diamagnetismo*, sia concorde ad esso, e si parla allora di *paramagnetismo* o *ferromagnetismo*.

La magnetizzazione si può descrivere tramite il *vettore magnetizzazione* $\mathbf{M}$, ovvero il momento magnetico del materiale per unità di volume, che si misura in A m^{-1} . A partire da $\mathbf{M}$ e da $\mathbf{B}$ si definisce un terzo vettore $\mathbf{H} = \mathbf{B}/\mu_0 - \mathbf{M}$, detto *campo magnetizzante*⁴. L'importanza di questo terzo vettore sta nel fatto che si può dimostrare che in funzione di esso la legge di Ampère assume una forma più semplice:

$$\nabla \times \mathbf{H} = \mathbf{j}_{libera} \quad (2.25)$$

e in forma integrale

$$\oint \mathbf{H} \cdot d\mathbf{l} = I_{libera} \quad (2.26)$$

dove il pedice “libera” indica che non si considerano le correnti di magnetizzazione all'interno del materiale; in pratica, si considerano solo le correnti accessibili direttamente allo sperimentatore tramite fili ecc.. (Si confronti con la legge di Gauss espressa come $\oint \mathbf{D} \cdot d\mathbf{S} = q_{libera}$.) Si tratta di una relazione del tutto generale; si osservi però che non si può dare una relazione altrettanto generale per la divergenza di $\mathbf{H}$ — non è un

⁴O semplicemente “campo H ”. Tradizionalmente era il vettore $\mathbf{H}$ ad essere chiamato campo magnetico, e questo uso lo si trova ancora su alcuni testi.

campo solenoidale, se non nel vuoto — per cui è necessario conoscere la *relazione costitutiva* che lo collega con $\mathbf{B}$.

A questo fine, se si assume che $\mathbf{M}$ ed $\mathbf{H}$ siano paralleli e proporzionali tra loro, si definisce un coefficiente χ , adimensionale, detto *suscettibilità magnetica*, dato da $\mathbf{M} = \chi \mathbf{H}$. Tramite la relazione $\mathbf{B} = \mu_0(\mathbf{H} + \mathbf{M}) = \mu_0(1 + \chi)\mathbf{H} = \mu \mathbf{H}$ si definisce infine un nuovo coefficiente μ detto *permeabilità magnetica* che gioca un ruolo simile alla costante dielettrica ϵ in elettrostatica: indica cioè di quanto viene alterato il campo magnetico in presenza di un materiale. Infatti

$$\oint \mathbf{H} \cdot d\mathbf{l} = \oint \frac{\mathbf{B}}{\mu} \cdot d\mathbf{l} = I_{libera} \Rightarrow \oint \mathbf{B} \cdot d\mathbf{l} = \mu I_{libera}$$

e confrontandola con la 2.7 si vede che l'effetto del materiale sul campo generato da una corrente data è di sostituire μ_0 con μ . (Ancora, è istruttivo confrontare con la legge di Gauss espressa come $\oint \mathbf{E} \cdot d\mathbf{S} = q_{libera}/\epsilon$.)

Si noti che queste relazioni sono del tutto generali — si tratta più che altro di definizioni — purché valgano due ipotesi: che $\mathbf{M}$ ed $\mathbf{H}$ siano paralleli, e che la relazione tra essi sia lineare. La prima può essere eliminata permettendo all'occorrenza che χ sia un tensore; la seconda, come si vedrà tra poco, è violata in casi importanti ed il discorso si fa più delicato.

A differenza del caso elettrostatico non è detto che la *permeabilità relativa* $\mu_r = \mu/\mu_0$ sia sempre maggiore o sempre minore di 1. Infatti, come già detto, esistono sia materiali *diamagnetici*, con $\chi < 0 \Rightarrow \mu_r < 1$, che *paramagnetici*, con $\chi > 0 \Rightarrow \mu_r > 1$ (i materiali *ferromagnetici* meritano un discorso a parte). I valori tipici di χ , per paramagneti e diamagneti, sono dell'ordine di $10^{-4} \div 10^{-5}$, per cui μ_r è nella maggioranza dei casi praticamente uguale ad 1. Si confronti questo con ϵ_r che raggiunge facilmente valori elevati; per l'acqua, per esempio, $\epsilon_r = 78$.

2.8.2 Il meccanismo della magnetizzazione

Una spiegazione microscopica del fenomeno della magnetizzazione non può essere data in termini semplici. Bisogna tener conto che il magnetismo nella materia è un fenomeno che richiede *necessariamente* una trattazione quantistica; si può dimostrare facilmente che in un sistema all'equilibrio termico governato dalla meccanica classica non è possibile avere magnetizzazione di alcun tipo. Molte trattazioni che si trovano sui libri di testo sono ingannevoli, perché sebbene formalmente corrette non dimostrano ciò che sembrano voler dimostrare.

In termini molto generici, si può dire che la magnetizzazione nella materia è dovuta a due tipi di fenomeno:

- Magnetismo orbitale. Gli orbitali e le energie degli elettroni nel materiale vengono leggermente alterati dalla presenza del campo magnetico e di conseguenza gli elettroni, essendo particelle cariche in movimento, a loro volta generano un campo magnetico.
- Magnetismo di spin. Gli elettroni (anche i nuclei, ma qui importano poco) possiedono un momento magnetico intrinseco. Sottoposti ad un campo magnetico questi spin tendono ad allinearsi lungo la direzione di esso generando di conseguenza un campo netto aggiuntivo.

Nei materiali gli elettroni si raccolgono in quelle che sono chiamate *bande di energia*. Se tutti gli elettroni appartengono a bande totalmente riempite (*shell chiuse*, nel linguaggio della fisica atomica), il che accade negli isolanti e nei semiconduttori, si ha diamagnetismo; se alcuni elettroni appartengono a bande parzialmente riempite (*shell aperte*), come nei metalli, si ha spesso paramagnetismo. Spesso, non sempre, perché un contributo diamagnetico lo si ha comunque (per via delle shell chiuse) e tutto dipende da quale dei due è più forte.

Il caso dei ferromagneti è un po' diverso. In essi il magnetismo è essenzialmente di spin ed esistono interazioni - elettrostatiche, non magnetiche - fra gli elettroni che tendono a mantenere allineati gli spin stessi. Evidentemente se si applica un campo magnetico esterno, più spin elettronici si allineeranno lungo di esso generando un campo indotto più forte quanto più forte è il campo esterno inducente. In virtù delle interazioni elettroniche accennate in precedenza, questo allineamento ed il relativo campo magnetico indotto rimane, almeno in parte, anche quando il campo esterno viene rimosso (*isteresi magnetica*), dando origine ai magneti permanenti.

Perché gli spin non sono sempre necessariamente allineati - in altre parole, perché un pezzo di ferro dolce non è sempre una calamita? Le ragioni sono fondamentalmente due.

Da una parte, c'è l'agitazione termica. Il moto caotico degli atomi "squassa" il sistema di spin, un po' come l'agitare una cassetta piena di biglie le mette in disordine. Man mano che la temperatura sale, la tendenza degli spin ad allinearsi viene soffocata e al di sopra di una temperatura critica T_c , detta *temperatura di Curie*, scompare del tutto ed il materiale è completamente privo di magnetizzazione.

Dall'altra parte, l'effetto del campo magnetico indotto è stato finora del tutto trascurato. Ogni elettrone contribuisce al campo magnetico a seconda della direzione del proprio spin; se gli spin sono antiparalleli gli effetti si cancellano, ma se sono paralleli si sommano; è proprio questo il campo magnetico indotto. Il contributo è piccolo, e finché gli spin allineati

sono pochi si può trascurare. Però l'energia magnetica (quella che si deve spendere per generare un campo magnetico H) è proporzionale ad H^2 . Quando il numero di spin allineati supera una certa soglia, la "spesa magnetica" risulta dominante ed il sistema trova sfavorevole, dal punto di vista energetico, allineare gli spin. Il risultato netto è che nel materiale si formano zone, dette *domini magnetici*, in cui gli spin sono allineati in un verso, ma ai bordi (*muri di Bloch*) si girano e si passa ad un altro dominio in cui sono allineati nel verso opposto. Questo permette di avere tanti spin orientati nel verso giusto per massimizzare il guadagno dovuto allo scambio, ma al tempo stesso di minimizzare il campo magnetico macroscopico, perché ogni dominio genera un campo che tende ad essere cancellato dal campo generato dal dominio adiacente.

Se si applica un campo magnetico esterno i domini paralleli al campo esterno saranno favoriti energeticamente rispetto agli altri, dunque si espanderanno a discapito dei secondi; il materiale si magnetizza in direzione parallela al campo esterno. Se poi tale campo è veramente intenso, tutti gli spin puntano nella sua direzione e da quel punto in poi la magnetizzazione non cresce più; è la *saturazione*.

Quando viene tolto il campo esterno, la situazione dovrebbe tornare in equilibrio; però di fatto la riconversione non è facile, ed il magnete rimane intrappolato in una situazione in cui i domini che puntano in una direzione rimangono in maggioranza. Questo fenomeno, detto *isteresi magnetica*, è alla base dell'esistenza dei magneti permanenti: le calamite.

Capitolo 3

Elettrodinamica

3.1 L'induzione elettromagnetica

Finora sono stati considerati solo campi statici. Si è visto che in queste condizioni i campi elettrici e magnetici sono due entità indipendenti tra loro e non vi è alcuna legge, eccettuata la forza di Lorentz, che li vede comparire insieme: essi possono essere determinati in modo separato, il campo elettrico da un lato e quello magnetico dall'altro.

Quando si ammette che i campi possano variare nel tempo le cose cambiano radicalmente. I due campi elettrico e magnetico risultano essere in stretto collegamento, come si poteva sospettare dal fatto che le sorgenti di $\mathbf{E}$ sono le cariche e quelle di $\mathbf{B}$ le correnti, ovvero cariche in movimento, e che entrambi i campi agiscono sulle cariche elettriche. Da qui in poi quindi non si potrà più a rigore parlare di campi elettrici e magnetici separatamente, ma solo di un'unica entità nota come *campo elettromagnetico*.

3.1.1 Campo variabile, circuito fisso

Ai tempi in cui il fisico inglese Michael Faraday era attivo era già noto che una corrente genera quello che noi ora chiamiamo campo magnetico¹. Egli si chiese se, così come una corrente genera un campo magnetico, un campo magnetico genera una corrente. Dopo molti sforzi trovò che la risposta sperimentale è negativa: anche campi molto intensi non producono correnti misurabili. Trovò anche, però, che un campo magnetico *variabile nel tempo* genera una corrente elettrica. Ricordiamo che per fare circolare una corrente c'è bisogno di un campo elettrico: sappiamo infatti

¹ Il concetto ed il nome stesso di campo lo dobbiamo a lui.

che un campo magnetico può solo variare la velocità delle particelle e non aumentarne (o diminuirne) il modulo, quindi non può metterle in moto. Dalle osservazioni di Faraday si evince quindi che un campo $\mathbf{B}$ variabile col tempo genera una *forza elettromotrice*, o *tensione*, e precisamente risulta che tale tensione è pari alla variazione col tempo, cambiata di segno, del flusso di $\mathbf{B}$ concatenato col circuito:

$$V = \oint \mathbf{E} \cdot d\mathbf{l} = -\frac{d}{dt} \int \mathbf{B} \cdot d\mathbf{S} = -\frac{d}{dt} \Phi_B \quad (3.1)$$

dove il flusso di $\mathbf{B}$ è appunto calcolato su di una superficie avente come bordo (o “che si appoggia a”) il percorso sul quale si calcola l'integrale a primo membro. Questa legge è nota come *legge di Faraday(-Henry)*.

È lecito chiedersi: quale delle infinite superfici aventi tale percorso come bordo bisogna scegliere? Una qualunque; come si è visto nella sezione 2.3 il risultato non dipende da tale scelta.

3.1.2 Campo fisso, circuito in moto

Consideriamo adesso il caso, in apparenza del tutto diverso, di una sottile sbarra metallica di lunghezza l che si muove con velocità $\mathbf{v}$, misurata in un determinato sistema di riferimento S , che supponiamo inerziale, e immersa in un campo magnetico $\mathbf{B}$ costante e uniforme. Per semplicità — ma nulla nel ragionamento dipende crucialmente da questo — poniamo che $\mathbf{v}$ sia costante e che i vettori $\mathbf{v}$, $\mathbf{B}$ e la direzione della sbarra $\hat{\mathbf{l}}$ siano tutti perpendicolari tra loro e formino una terna destrorsa ($\hat{\mathbf{l}} = \hat{\mathbf{v}} \times \hat{\mathbf{B}}$). Le cariche mobili nella sbarra potranno, in pratica, spostarsi solo nella direzione parallela ad essa. Per la legge di Lorentz 2.1, ognuna di esse sentirà una forza pari a $\mathbf{F} = q\mathbf{v} \times \mathbf{B} = qvB\hat{\mathbf{l}}$.

Mettiamoci ora nel sistema di riferimento S' solidale con la sbarra. Qui $\mathbf{v} = 0$ e pertanto non si ha forza magnetica. Per il principio di relatività però la forza misurata in S' , che è anch'esso un sistema inerziale, deve essere la stessa che in S ; se ne conclude che in S' esiste una forza $\mathbf{F} = qvB\hat{\mathbf{l}}$, sentita solo dalle cariche, che si esercita anche sulle cariche ferme. Deve perciò, per definizione, trattarsi di una forza elettrica, derivante da un campo elettrico $\mathbf{E}' = \mathbf{F}/q = vB\hat{\mathbf{l}} = \mathbf{v} \times \mathbf{B}$. Quindi *il campo magnetico in S è visto come campo elettrico in S'* .

Immaginiamo che la sbarra formi il quarto lato di un circuito, reale od immaginario, di forma rettangolare. Il lato opposto alla sbarra ha evidentemente lunghezza l ; quelli perpendicolari avranno lunghezza $a + vt$, con a un valore arbitrario che scegliamo positivo (e siccome v è il modulo di un

vettore, quindi anch'esso positivo, stiamo considerando un circuito che si ingrandisce). Allora la forza elettromotrice nel circuito vale, visto che $\mathbf{E}'$ è diverso da zero solo sulla sbarra ed è ivi uniforme, $\oint \mathbf{E}' \cdot d\mathbf{l} = -E'l = -Blv$, dove il segno meno deriva dal fatto che il verso di percorrenza positivo del circuito è antiorario se visto da un osservatore orientato come $\mathbf{B}$.

Il flusso magnetico all'interno di questo circuito risulta funzione del tempo: $\Phi_B = BS = Bl(a + vt)$, e la sua derivata è $d\Phi_B/dt = Blv$. Pertanto anche in questo caso, oltre che in quello del circuito fisso, abbiamo

$$\oint \mathbf{E}' \cdot d\mathbf{l} = -\frac{d\Phi_B}{dt}.$$

Il ragionamento appena svolto può essere generalizzato a tutti i possibili casi di variazione del flusso, che però si riducono in essenza ai due visti finora (il ragionamento lo si può trovare nei testi indicati in Bibliografia). La legge di Faraday dunque vale sempre: *si ha forza elettromotrice ogniqualvolta varia il flusso magnetico attraverso un circuito, qualunque sia la causa di questa variazione, e tale forza elettromotrice è sempre data dalla 3.1.*

L'equazione 3.1 naturalmente ammette una forma differenziale

$$\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t} \quad (3.2)$$

che ci dice, tra le altre cose, che la legge di Faraday non è legata all'esistenza di un circuito, ma è *una relazione tra i campi valida in generale*.

Segue inoltre da questa legge che al di fuori del caso statico $\mathbf{E}$ non è più irrotazionale o, che è lo stesso, che il suo integrale lungo una superficie chiusa non è più necessariamente nullo. A sua volta questo implica che in generale non è possibile ricavare $\mathbf{E}$ come gradiente di una funzione scalare. La questione dei potenziali deve quindi essere riesaminata da capo, come vedremo nel seguito.

3.1.3 La conservazione dell'energia

Nel caso appena visto della sbarra che si muove, la 2.11 ci dice che la sbarra stessa è soggetta ad una forza. Infatti, ponendo per fissare le idee $\mathbf{B} = B\hat{\mathbf{y}}$, $\mathbf{v} = v\hat{\mathbf{x}}$ e $\hat{\mathbf{l}} = \hat{\mathbf{z}}$, si ha

$$\mathbf{F} = I\mathbf{a} \times \mathbf{B} = I(-a\hat{\mathbf{y}} \times B\hat{\mathbf{z}}) = -IaB\hat{\mathbf{x}},$$

ovvero una forza diretta in modo tale da *opporsi* al movimento della sbarra, quindi alla variazione del flusso magnetico. Perché la sbarra si possa

muovere di velocità costante è quindi necessario applicare una forza eguale ed opposta, o in altri termini spendere una potenza (energia per unità di tempo) pari a $W = |\mathbf{F} \cdot \mathbf{v}| = I a B v$.

Dove va a finire questa energia? Immaginiamo per esempio che la sbarra sia connessa ad un circuito di resistenza R ; allora la corrente vale $I = V/R = a B v/R$ e la potenza da fornire è $W = I a B v = a^2 B^2 v^2 / R$. Sulla resistenza peraltro si dissipa una potenza pari a $P = R I^2 = R a^2 B^2 v^2 / R^2 = W$. Il bilancio energetico dunque è in pareggio: l'energia fornita per far muovere la sbarra viene tutta dissipata sulla resistenza, e viceversa.

Si noti che per questo è essenziale che la legge di Faraday abbia il segno negativo davanti alla derivata del flusso (un risultato detto *legge di Lenz*). Se così non fosse, la sbarra riceverebbe una forza non opposta, ma concorde al movimento, e che tenderebbe a farla accelerare invece che rallentare. Di conseguenza la sbarra si autoaccelererebbe: non appena messa in moto aumenterebbe spontaneamente la propria velocità senza più bisogno di una forza esterna, producendo energia dal nulla.

Si noti che questo esempio ci mostra che *per variare il flusso magnetico bisogna spendere energia*. Questa conclusione è valida non solo per il caso qui esposto del circuito in moto, ma in generale.

3.1.4 L'alternatore

Consideriamo ora una spira di area S e costituita da N avvolgimenti che ruota in un campo magnetico uniforme $\mathbf{B}$ con velocità angolare costante ω attorno ad un asse perpendicolare al campo. Il flusso magnetico sarà dato, scegliendo opportunamente l'origine dei tempi, da $\Phi_B = \mathbf{B} \cdot \mathbf{S} = B S \cos \omega t$. Su ogni avvolgimento della spira viene dunque generata una tensione pari a $V = B S \omega \sin \omega t$, e la tensione totale è dunque

$$V = N B S \omega \sin \omega t.$$

Questo è il modo più semplice di generare una tensione, dunque una corrente elettrica. Dispositivi di questo tipo sono i prototipi dei generatori di tensione (corrente), e come si vede la tensione (corrente) generata è *alternata*, vale a dire cambia periodicamente verso, da cui il nome di *alternatore* dato al dispositivo. Per generare corrente continua $I = \text{costante}$ con questa tecnica è necessario ricorrere a qualche accorgimento.

Per generare la tensione è necessario spendere energia. Infatti si ottiene una corrente, e questo richiede una potenza istantanea $P(t) = V(t)I(t) = I(t)NBS\omega \sin \omega t$. Da dove viene questa energia? Da quanto detto nel paragrafo precedente, deve derivare dall'energia necessaria

per mantenere in moto la spira; sappiamo infatti dall'esempio della sbarra che per muovere un conduttore in un campo magnetico c'è bisogno di applicare una forza. L'energia non potrebbe venire dal campo magnetico: basti pensare che questo potrebbe essere generato (anche se nei generatori pratici di solito non è così) da dei magneti permanenti, che proprio in quanto permanenti non necessitano di energia per mantenere il loro campo. Deve quindi venire da una forza esterna.

Calcoliamo dunque l'energia trasferita alla spira nell'unità di tempo. Derivando la 2.17 si ottiene

$$\frac{dU}{dt} = -\frac{d}{dt} \mathbf{m} \cdot \mathbf{B} = I(t) N S B \omega \sin \omega t$$

che è proprio l'energia richiesta per mantenere in circolo la corrente. Il generatore genera energia elettrica, ma non energia *tout-court* (altrimenti avremmo il moto perpetuo): si limita a convertire l'energia fornita per far girare la spira — che può derivare da una turbina a vapore alimentata a carbone, petrolio o energia nucleare, o da turbine ad acqua come nelle centrali idroelettriche, o da pale di mulini a vento, o dal motore di un'automobile, o magari dalla forza muscolare — in energia elettrica che può poi essere facilmente distribuita.

3.2 Induttanza

3.2.1 Autoinduzione

In tutti gli esempi visti finora, il flusso del campo magnetico è risultato proporzionale alla corrente I . Pertanto in casi simili, piuttosto frequenti, quando B varia col tempo la tensione V è proporzionale a dI/dt :

$$V = L \frac{dI}{dt}. \quad (3.3)$$

Il coefficiente L è chiamato *coefficiente di autoinduzione* e si misura in Henry (unità già incontrata in precedenza nella definizione della permeabilità magnetica) $H = m^2 \text{ kg } C^{-2}$. Dispositivi con $L \neq 0$ sono chiamate induttanze; la tipica induttanza è il solenoide, per il quale nel caso ideale il coefficiente di autoinduzione per unità di lunghezza vale $L/l = \mu_0 n^2 S$ (*Esercizio*: dimostrarlo).

Le induttanze accumulano energia, come i condensatori, ma a differenza di questi ultimi, in cui l'energia è dovuta all'attrazione elettrostatica tra le cariche sulle armature, nelle induttanze l'energia è dovuta all'opposizione

(dovuta all'induzione elettromagnetica) che esercita il flusso al cambiamento della propria intensità, e che in questi casi è equivalente a dire che vi è un' opposizione della corrente al cambiamento della propria intensità. Il valore dell'energia accumulata lo si ottiene facilmente considerando che, se $v = v(t)$ e $i = i(t)$ sono rispettivamente la tensione e la corrente istantanee nell'induttanza, l'energia necessaria per instaurare una corrente I partendo dallo stato di assenza di corrente è

$$U = \int_0^I P(t) dt = \int_0^I v(t)i(t) dt = L \int_0^I i(t) \frac{di(t)}{dt} dt = \frac{1}{2}LI^2. \quad (3.4)$$

Si noti l'affinità della 3.4 con l'equazione 1.26 che esprime l'energia accumulata in un condensatore $U = q^2/2C$, mentre l' analogia con la legge di Joule $P = RI^2$ è ingannevole in quanto la resistenza non accumula energia ma la dissipa, mentre condensatori ed induttanze (ideali) accumulano energia e non la dissipano.

In effetti le induttanze

$$V = L dI/dt,$$

i condensatori

$$V = q/C = \int I dt/C,$$

e le resistenze

$$V = RI$$

hanno in comune la caratteristica di avere una relazione *lineare* tra corrente e tensione. Si tratta dei dispositivi lineari più importanti e come tali sono alle fondamenta dell'elettrotecnica. Non si dimentichi però che nell'elettronica hanno amplissimo e fondamentale uso dispositivi *non lineari* come diodi e transistor. Un diodo a semiconduttore, per esempio, ha una relazione tra tensione e corrente data da $I = I_0 (e^{kV} - 1)$, dove I_0 e k sono funzioni della temperatura caratteristiche dello specifico dispositivo.

3.2.2 Mutua induzione. Il trasformatore

Vi è un quarto importante tipo di dispositivo lineare, detto *mutua induttanza*. Nella discussione dell' induttanza non vi è nulla che imponga che il circuito che genera il flusso magnetico debba essere lo stesso nel quale

si misura la tensione; se il flusso generato da un dispositivo ne attraversa un secondo, avremo una tensione V_2 nel secondo proporzionale alla variazione di corrente I_1 nel primo:

$$V_2 = M_{21} \frac{dI_1}{dt}. \quad (3.5)$$

ed il coefficiente M_{21} è detto *coefficiente di mutua induzione*. Dispositivi del genere accoppiano circuiti o parti di circuiti che in corrente continua sarebbero separati. L'unità di misura della mutua induttanza ovviamente è ancora l'Henry.

Un esempio semplice ed importante di mutua induttanza è quello di due solenoidi 1 e 2, formati da N_1 e N_2 spire rispettivamente, avvolti attorno allo stesso traferro, di sezione S , piegato a forma di anello. Questa sistemazione permette di convogliare praticamente tutte le linee di campo magnetico generate da un solenoide all'interno dell'altro. Supporremo che si possano utilizzare le formule per solenoidi infiniti. Nel solenoide 1, detto *primario*, di lunghezza l_1 , scorre una corrente I_1 che genera un campo magnetico $B = \mu N_1 I_1 / l_1 = \mu n_1 I_1$ e di conseguenza un flusso $\Phi_1 = \mu n_1 S I_1$. Questo flusso è, per l'ipotesi, concatenato con ogni spira del solenoide 2 — il *secondario* — per cui il flusso totale nel secondario è $\Phi_2 = \mu n_1 N_2 S I_1$. La mutua induttanza perciò è $M_{21} = \mu n_1 N_2 S$.

Siccome ai capi del primario si ha una tensione pari a $V_1 = L_1 dI_1 / dt = \mu n_1^2 S l_1 dI_1 / dt$, tra le due tensioni V_1 e V_2 vale la relazione

$$V_2 / V_1 = M_{21} / L_1 = \mu n_1 N_2 S / \mu n_1^2 S l_1 = N_2 / N_1. \quad (3.6)$$

Questo dispositivo permette dunque di trasformare una tensione alternata V_1 in un'altra V_2 più alta o più bassa a seconda del numero di avvolgimenti nel primario e nel secondario. È perciò chiamato *trasformatore*. La possibilità di trasformare la tensione a piacimento è probabilmente il maggior vantaggio offerto dalle correnti alternate rispetto a quelle continue.

3.3 Circuiti oscillanti

Il circuito basilare dell'elettrotecnica è il circuito ad una maglia in cui si suppone che le varie caratteristiche del circuito:

- Tensione applicata V
- Resistenza R

- Capacità C
- Induttanza L

siano *concentrate* ciascuna in un elemento circuitale. Perciò si suppone che un'induttanza abbia resistenza trascurabile, che una resistenza abbia capacità trascurabile e così via. È facile scrivere l'equazione differenziale che descrive il comportamento di un circuito del genere, chiamato per ovvie ragioni *circuito RLC a costanti concentrate*. Basta imporre che la somma algebrica delle variazioni di tensione sia nulla:

$$V = L \frac{dI}{dt} + RI + q/C \quad (3.7)$$

da cui derivando

$$\frac{dV}{dt} = L \frac{d^2 I}{dt^2} + R \frac{dI}{dt} + I/C. \quad (3.8)$$

Naturalmente si può anche scrivere l'equazione per la carica accumulata (che, per le ipotesi fatte, è tutta sul condensatore)

$$V = L \frac{d^2 q}{dt^2} + R \frac{dq}{dt} + q/C. \quad (3.9)$$

Si noti l'analogia formale della 3.8 e della 3.9 con l'equazione del moto di un oscillatore meccanico forzato e con attrito: quest'ultima si ottiene da quella per il circuito *RLC* identificando le varie grandezze seguendo la tabella:

Oscillatore meccanico	3.8			3.9		
Coordinata	x	$\leftrightarrow$	I	x	$\leftrightarrow$	q
Velocità	v	$\leftrightarrow$	dI/dt	v	$\leftrightarrow$	I
Forzante	F_{ext}	$\leftrightarrow$	dV/dt	F_{ext}	$\leftrightarrow$	V
Massa	m	$\leftrightarrow$	L	m	$\leftrightarrow$	L
Coefficiente di attrito	λ	$\leftrightarrow$	R	λ	$\leftrightarrow$	R
Costante elastica	k	$\leftrightarrow$	$1/C$	k	$\leftrightarrow$	$1/C$

Ne segue che l'energia accumulata sull' induttanza corrisponde all'energia cinetica dell'oscillatore e quella sul condensatore all'energia potenziale elastica.

3.3.1 Oscillazioni libere e risonanza

La soluzione della 3.8 si ottiene con i metodi elementari di soluzione delle equazioni differenziali lineari a coefficienti costanti. Per prima cosa si risolve l'equazione omogenea associata, che corrisponde ad un sistema senza tensione applicata e che quindi esegue oscillazioni libere (nel parallelo con l'oscillatore meccanico si parlerebbe di un oscillatore armonico smorzato non forzato). L'equazione caratteristica è

$$L\omega^2 + R\omega + 1/C = 0$$

ed ha soluzioni

$$\omega_{\pm} = -\frac{R}{2L} \pm \sqrt{\left(\frac{R}{2L}\right)^2 - \frac{1}{LC}}.$$

Le soluzioni dell'equazione 3.8, per radici distinte, sono della forma $I(t) = Ae^{\omega_+ t} + Be^{\omega_- t}$ che a seconda del discriminante $\Delta = \left(\frac{R}{2L}\right)^2 - \frac{1}{LC}$ si dividono in due tipi:

- se $\Delta \geq 0$, le $\omega_{\pm}$ sono radici reali negative ed $I(t)$ ha un decadimento esponenziale, tanto più rapido quanto maggiore è $|\omega_{\pm}|$; per $\Delta = 0 \Rightarrow \omega_{\pm} = -\frac{R}{2L}$ si parla di *smorzamento critico* (in questo secondo caso peraltro la soluzione non è una semplice somma di esponenziali);
- se $\Delta < 0$, $I(t)$ ha un andamento oscillante smorzato. Lo smorzamento dipende dalla parte reale di $\omega_{\pm}$, $-\frac{R}{2L}$: tanto più essa è grande, tanto più rapido è lo smorzamento, mentre la frequenza delle oscillazioni è proporzionale a $\sqrt{-\Delta}$.

Nel caso limite di assenza di smorzamento, $R = 0$, si ha sempre $\Delta < 0$ ed I oscilla con frequenza $\omega_0 = 1/\sqrt{LC}$, detta *frequenza propria*, *auto-frequenza* o *frequenza di risonanza* del circuito. La corrente ha un andamento semplicemente sinusoidale con frequenza ω_0 , e così pure la carica sul condensatore, ma sfasata di $\pi/2$ rispetto alla corrente. Dalle espressioni 3.4 e 1.26 si vede che l'energia si rimpalla tra l'induttanza e la capacità, esattamente come in un oscillatore armonico si rimpalla tra l'energia cinetica e quella potenziale.

3.3.2 Oscillazioni forzate. Impedenza

Passiamo adesso al caso in cui vi sia una tensione applicata a guidare il circuito. Un modo generale di risolvere questa equazione è esposto nell'Appendice B. Qui usiamo tecniche più elementari e consideriamo il caso

in cui tale forza sia periodica e precisamente tale che $V = V_0 \cos \omega t$ e pertanto $dV/dt = -\omega V_0 \sin \omega t$. La soluzione dell'equazione 3.8 è dunque data dall'integrale generale dell'omogenea associata sommato ad un integrale particolare dell'equazione completa che soddisfi alle condizioni iniziali; se ci limitiamo al caso a regime, in cui $t \rightarrow \infty$ e pertanto $t \gg L/R$, il primo integrale è trascurabile e rimaniamo solo con l'integrale particolare. Quest'ultimo, come è noto, è in questo caso della forma $I(t) = I_0 \cos(\omega t - \phi)$ (corrente alternata).

Introducendo questa soluzione di prova nella 3.8 otteniamo le seguenti relazioni:

$$I_0 = V_0/Z; \tan \phi = X/R; Z = \sqrt{R^2 + X^2}; X = \omega L - \frac{1}{\omega C} \quad (3.10)$$

dove X è detta *reattanza* e Z è detta *impedenza*. L'angolo ϕ è detto *angolo di fase* ed indica, com'è evidente, lo sfasamento tra tensione e corrente istantanee. Si vede che l'induttanza tende a dare un $\phi > 0$, ovvero a far sì che la corrente "resti indietro" rispetto alla tensione esterna, mentre la capacità ha l'effetto opposto. La resistenza invece da questo punto di vista ha un effetto neutro. L'impedenza è la generalizzazione della resistenza al caso delle correnti alternate e la 3.10 è da parte sua la generalizzazione della legge di Ohm, ma si badi bene che si tratta di una relazione tra le *ampiezze* di tensione e corrente e non tra i valori istantanei.

Data la presenza di due grandezze rilevanti, l'ampiezza e lo sfasamento, è in effetti impossibile scrivere una relazione del tipo di Ohm tra numeri reali nel caso di correnti alternate. È però possibile farlo introducendo una relazione tra grandezze *complesse*, che di per sé corrispondono a coppie di numeri reali. Definiamo allora una *tensione complessa* $\hat{V} = V_0 e^{i\omega t}$ ed una *corrente complessa* $\hat{I} = I_0 e^{i(\omega t + \phi)}$; il significato di questa posizione sta nel fatto che le tensioni e le correnti misurabili, che ovviamente sono grandezze reali, possono essere ottenute prendendo, alla fine dei conti, la parte reale (o immaginaria) del risultato. La linearità delle equazioni ci garantisce che il procedimento è corretto. Inserendo queste grandezze nell'equazione del circuito 3.8 si ottiene

$$L\hat{I}'' + R\hat{I}' + \hat{I}/C = \hat{V}' \quad (3.11)$$

ovvero, eseguendo le derivate e raggruppando,

$$\hat{V} = \hat{I} \left[R + i \overbrace{\left(\omega L - \frac{1}{\omega C} \right)}^{\hat{X}} \right] = \hat{I} \hat{Z} \quad (3.12)$$

equazione che ha effettivamente la forma della legge di Ohm e vale per i valori istantanei di tensione e corrente. Le grandezze Z e X sono dunque il modulo delle grandezze complesse $\hat{Z}$ e $\hat{X}$.

Si noti che tutte le grandezze nella 3.10 sono funzioni della frequenza ω . In particolare Z è minima $\Rightarrow I_0$ è massima per $\omega = 1/\sqrt{LC} = \omega_0$, ovvero proprio per la frequenza di risonanza definita in precedenza. Tracciando un grafico di $I_0 = I_0(\omega)$ si vede che (*Esercizio*) l'ampiezza della corrente ha un picco attorno ad ω_0 tanto più alto quanto più R è piccola, arrivando a V_0/R alla risonanza esatta. Ne consegue che un circuito RLC con bassa resistenza può essere usato come *sintonizzatore*: le frequenze della forzante molto prossime alla risonanza produrranno una corrente notevole, mentre le altre saranno pressoché ignorate.

3.4 La legge di Ampère-Maxwell

Riprendiamo in considerazione la legge di Ampère 2.6. È ancora valida nel caso in cui la corrente non sia stazionaria ma vari col tempo? Per capirlo, calcoliamo la divergenza di entrambi i membri. Il primo membro dà zero, mentre il secondo diventa $\mu_0 \nabla \cdot \mathbf{j}$ che, per l'equazione di continuità 1.24, coincide con $-\mu_0 \partial \rho / \partial t$; prese insieme queste due osservazioni implicano che la densità di carica deve essere indipendente dal tempo. Se ne conclude che *l'equazione di Ampère vale solo nei casi stazionari*, in cui non si hanno mutamenti della densità di carica spaziale. Questo naturalmente non significa che le cariche non si muovano, ma soltanto che ogni spostamento di una carica deve essere allo stesso momento compensato esattamente dal movimento delle altre. In altre parole, si possono avere correnti, ma non devono cambiare nel tempo, dunque anche i campi magnetici devono essere statici.

Come è possibile generalizzare la 2.6 in modo che valga anche nel caso dinamico? Il ragionamento esposto sopra ce ne fornisce la chiave. Bisogna aggiungere al secondo membro un termine la cui divergenza cancelli esattamente il termine $-\mu_0 \partial \rho / \partial t$. Ma un termine del genere è facile da trovare: basta semplicemente ricordarsi della legge di Gauss per il campo elettrico, $\nabla \cdot \mathbf{E} = \rho / \epsilon_0$ e derivarla rispetto al tempo. Il termine da

aggiungere al secondo membro dell'equazione 2.6 è quindi $\mu_0\epsilon_0\partial\mathbf{E}/\partial t$ e si ottiene infine

$$\nabla \times \mathbf{B} = \mu_0\mathbf{j} + \epsilon_0\mu_0\frac{\partial\mathbf{E}}{\partial t} \quad (3.13)$$

che, come è agevole verificare, rispetta l'equazione di continuità della carica in qualunque circostanza e si riduce alla legge di Ampère per situazioni stazionarie. Il termine $c = 1/\sqrt{\epsilon_0\mu_0}$ ha le dimensioni di una velocità e vale $3 \cdot 10^8$ m/s; come si vedrà ha un'importanza fondamentale in ciò che segue.

L'equazione così modificata prende il nome di *legge di Ampère-Maxwell* e la sua forma integrale è

$$\oint \mathbf{B} \cdot d\mathbf{l} = \mu_0 I + \epsilon_0\mu_0 \frac{d}{dt} \int \mathbf{E} \cdot d\mathbf{S}. \quad (3.14)$$

È istruttivo verificare (*Esercizio*) che con la legge di Ampère-Maxwell in questa forma la scelta della superficie su cui calcolare l'integrale a secondo membro è ininfluente, purché si scelga la stessa superficie per il calcolo di I . Il ragionamento fatto nella sezione 3.1 non è infatti direttamente trasponibile a questo caso in quanto la divergenza del campo elettrico non è identicamente nulla.

3.4.1 La corrente di spostamento

Il termine $\epsilon_0 \frac{d}{dt} \int \mathbf{E} \cdot d\mathbf{S} = \epsilon_0 \frac{d\Phi_E}{dt}$ ha le dimensioni di una corrente e viene chiamato *corrente di spostamento*, I_{spost} . Questo nome ha ragioni storiche, ma in realtà non rappresenta una corrente in quanto esiste anche nel vuoto, dove non ci sono cariche né correnti. Riscrivendo la 3.13 sotto la forma

$$\nabla \times \mathbf{B} = \mu_0 (\mathbf{j} + \mathbf{j}_{spost})$$

si vede però che includendo $\mathbf{j}_{spost}$ nella densità di corrente totale otteniamo che *tutti i circuiti sono chiusi*. Infatti la divergenza del primo membro è nulla, quindi lo deve essere anche quella del secondo e ne segue che la corrente $\mathbf{j} + \mathbf{j}_{spost}$ ha divergenza nulla: per il significato geometrico della divergenza, significa che la densità di corrente non “nasce” né “muore” in nessun punto dello spazio.

Verifichiamo questa affermazione in un caso particolare: un circuito costituito da un generatore di tensione e da un condensatore nel vuoto a facce piane e parallele, di superficie S e distanti d , pertanto di capacità

$C = \epsilon_0 S/d$. Resistenze ed induttanze si considerano trascurabili. Siccome nel condensatore $V = q/C$ e $E = V/d =$ uniforme, possiamo calcolare

$$I_{spost} = \epsilon_0 \frac{d}{dt} \int \mathbf{E} \cdot d\mathbf{S} = \epsilon_0 S \frac{dE}{dt} = \frac{\epsilon_0 S}{d} \frac{dV}{dt} = \frac{d(CV)}{dt} = \frac{dq}{dt} = I.$$

Pertanto laddove non scorre la corrente vera “scorre” una corrente di spostamento esattamente dello stesso valore, CVD.

Infine si noti che con l'introduzione della corrente di spostamento l'equazione di Ampère-Maxwell diventa analoga alla legge di Faraday: nel vuoto ($\mathbf{j} = 0$) infatti le due equazioni sono perfettamente simmetriche in $\mathbf{E}$ e $\mathbf{B}$, se si eccettua il segno ed il termine $\epsilon_0 \mu_0 = 1/c^2$ che deriva esclusivamente dalla scelta delle unità di misura per i campi. La differenza principale però sta nel fatto che tipicamente il termine I_{spost} è piccolo rispetto ad I , quando quest'ultimo sia presente (nel caso appena visto del condensatore I_{spost} non è trascurabile perché I è assente), da cui si spiega la validità ampia della legge di Ampère 2.7 anche nei casi di campi variabili, tranne che nei casi in cui le variazioni siano effettivamente molto rapide. In questi ultimi casi la corrente di spostamento gioca un ruolo veramente fondamentale, come vedremo.

3.5 Le equazioni di Maxwell

Ci si può chiedere se nel caso dinamico sia necessario alterare le equazioni alla divergenza come abbiamo fatto con le equazioni al rotore. La risposta è no: il sistema di equazioni che abbiamo raggiunto, dette *equazioni di Maxwell*, rappresenta la forma definitiva delle equazioni dell'elettromagnetismo². Come tale si tratta di uno dei massimi risultati ottenuti dalla Fisica e vale la pena di riportare tutte le equazioni insieme in forma sia differenziale che integrale:

$$\begin{array}{ll} \nabla \cdot \mathbf{E} = \rho/\epsilon_0 & \oint \mathbf{E} \cdot d\mathbf{S} = q/\epsilon_0 \\ \nabla \cdot \mathbf{B} = 0 & \oint \mathbf{B} \cdot d\mathbf{S} = 0 \\ \nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t} & \oint \mathbf{E} \cdot d\mathbf{l} = -\frac{d}{dt} \int \mathbf{B} \cdot d\mathbf{S} \\ \nabla \times \mathbf{B} = \mu_0 \mathbf{j} + \epsilon_0 \mu_0 \frac{\partial \mathbf{E}}{\partial t} & \oint \mathbf{B} \cdot d\mathbf{l} = \mu_0 I + \epsilon_0 \mu_0 \frac{d}{dt} \int \mathbf{E} \cdot d\mathbf{S} \end{array} \quad (3.15)$$

La forma qui sopra vale per le equazioni nel vuoto o, comunque, ogni volta che si considerino esplicitamente tutte le cariche e le correnti pre-

²Almeno nel mondo macroscopico. La trattazione deve essere modificata, come si accennerà in seguito, quando entra in campo la Meccanica Quantistica.

senti. Laddove siano presenti mezzi materiali, per molti usi risulta più conveniente utilizzare una forma che utilizzi i vettori $\mathbf{D}$ ed $\mathbf{H}$:

$$\begin{aligned}
\nabla \cdot \mathbf{D} &= \rho_{libera} & \oint \mathbf{D} \cdot d\mathbf{S} &= q_{libera} \\
\nabla \cdot \mathbf{B} &= 0 & \oint \mathbf{B} \cdot d\mathbf{S} &= 0 \\
\nabla \times \mathbf{E} &= -\frac{\partial \mathbf{B}}{\partial t} & \oint \mathbf{E} \cdot d\mathbf{l} &= -\frac{d}{dt} \int \mathbf{B} \cdot d\mathbf{S} \\
\nabla \times \mathbf{H} &= \mathbf{j}_{libera} + \frac{\partial \mathbf{D}}{\partial t} & \oint \mathbf{H} \cdot d\mathbf{l} &= I_{libera} + \frac{d}{dt} \int \mathbf{D} \cdot d\mathbf{S}
\end{aligned} \tag{3.16}$$

dove il pedice “libera” indica che non si considerano cariche di polarizzazione né correnti di magnetizzazione.

Questa seconda forma, nonostante l'apparenza, è in realtà concettualmente più complicata della 3.15, in quanto bisogna anche conoscere le *relazioni costitutive* tra $\mathbf{D}$ ed $\mathbf{E}$ e tra $\mathbf{B}$ ed $\mathbf{H}$, e queste relazioni, che variano da materiale a materiale, sono spesso difficili da conoscere, se non sperimentalmente. Normalmente si assume una relazione lineare, $\mathbf{D} = \epsilon \mathbf{E}$ e $\mathbf{B} = \mu \mathbf{H}$, ma non di rado questa ipotesi non è valida. Per esempio la relazione tra $\mathbf{B}$ ed $\mathbf{H}$ nei ferromagneti è, come discusso nel paragrafo 2.8, fortemente non lineare. Si danno poi casi in cui i vettori $\mathbf{D}$ ed $\mathbf{E}$, $\mathbf{B}$ ed $\mathbf{H}$ non sono nemmeno paralleli fra loro per cui ϵ e μ non sono scalari ma tensori, eccetera. Nel seguito faremo uso solo delle equazioni 3.15.

A queste equazioni va aggiunta l'espressione 2.2 per la forza di Lorentz. Non c'è invece bisogno di stabilire esplicitamente la conservazione della carica, perché come si è già visto le equazioni di Maxwell implicano l'equazione di continuità (*Esercizio*: verificarlo per le equazioni 3.16).

3.5.1 L'equazione delle onde

Dalle 3.15 è possibile ottenere un insieme di equazioni più semplici, almeno nel vuoto ($\rho = 0, \mathbf{j} = 0$). Prendendo il rotore della terza, la derivata rispetto al tempo della quarta e combinandole si ottiene

$$\nabla \times \nabla \times \mathbf{E} + \epsilon_0 \mu_0 \frac{d^2 \mathbf{E}}{dt^2} = 0$$

e ricordando la D.13 e la prima delle 3.15 con $\rho = 0$

$$\square^2 \mathbf{E} = \nabla^2 \mathbf{E} - \frac{1}{c^2} \frac{\partial^2 \mathbf{E}}{\partial t^2} = 0 \tag{3.17}$$

dove si è definito l'*operatore d'alembertiano*

$$\square^2 = \nabla^2 - \frac{1}{c^2} \frac{\partial^2}{\partial t^2}. \tag{3.18}$$

Analogamente si trova (*Esercizio*), per il campo magnetico,

$$\square^2 \mathbf{B} = \nabla^2 \mathbf{B} - \frac{1}{c^2} \frac{\partial^2 \mathbf{B}}{\partial t^2} = 0. \quad (3.19)$$

Le equazioni della forma 3.17 e 3.19 sono note come *equazioni delle onde* (si noti come queste equazioni non sarebbero state ottenibili se non si fosse introdotta la corrente di spostamento), e verranno studiate nel prossimo capitolo. Qui basti ricordare che queste equazioni impongono condizioni necessarie, ma non sufficienti, sul campo elettromagnetico. Infatti sono sei equazioni (due equazioni vettoriali per vettori a tre componenti) mentre le equazioni di Maxwell in forma differenziale sono otto equazioni (due vettoriali più due scalari).

Non è difficile ottenere un'equazione corrispondente alla 3.17 valida all'interno dei metalli. Siccome in questo caso le cariche sono libere, è legittimo supporre $\rho = 0$, almeno per campi che non variano troppo velocemente, in quanto la carica è libera di spostarsi e di annullare rapidamente qualunque squilibrio si formi. Per la stessa ragione non è invece ammissibile trascurare la corrente; questa produce nella 3.17 un termine addizionale $-\mu_0 d\mathbf{j}/dt$, che si esprime in termini del campo elettrico tramite la legge di Ohm in forma vettoriale $\mathbf{j} = \sigma \mathbf{E}$. Si ottiene infine (*Esercizio*)

$$\square^2 \mathbf{E} - \mu_0 \sigma \frac{\partial \mathbf{E}}{\partial t} = 0. \quad (3.20)$$

3.6 Potenziali elettrodinamici

Ora che conosciamo le equazioni definitive del campo elettromagnetico, possiamo tornare sul problema della derivazione dei campi da dei potenziali. L'equazione di Faraday 3.1, 3.2, come si è visto, impedisce di considerare conservativo il campo elettrico nel caso generale, quindi di poterlo derivare da un potenziale scalare. D'altro canto la solenoidalità di $\mathbf{B}$ non è in discussione e la definizione del potenziale vettore $\mathbf{A}$ è ancora valida. Possiamo definire un potenziale anche per il campo elettrico non statico? Certamente. Infatti se riscriviamo l'equazione 3.2 in termini del potenziale vettore

$$\nabla \times \mathbf{E} = -\partial/\partial t (\nabla \times \mathbf{A}) \Rightarrow \nabla \times \left(\mathbf{E} + \frac{\partial \mathbf{A}}{\partial t} \right) = 0$$

vediamo che il termine tra parentesi nella seconda equazione è irrotazionale e può pertanto essere considerato come il gradiente di una funzione

scalare, $\mathbf{E} + \partial \mathbf{A} / \partial t = -\nabla \phi$, da cui si ottiene l'espressione generale per il campo elettrico

$$\mathbf{E} = -\nabla \phi - \frac{\partial \mathbf{A}}{\partial t}. \quad (3.21)$$

Il potenziale scalare ϕ che compare nella 3.21 si riduce evidentemente al potenziale elettrostatico nel caso in cui i campi non varino nel tempo. Non bisogna però dimenticare che ϕ ha un significato più ampio e può essere una funzione esplicita del tempo.

Rimane vero il fatto che i potenziali non sono definiti univocamente. Il fatto che ϕ ed $\mathbf{A}$ siano connessi nella 3.21 significa anzi che abbiamo più libertà di scelta; il potenziale scalare non è determinato semplicemente a meno di una costante additiva come nel caso statico. Infatti, possiamo senz'altro aggiungere un termine a ϕ purché ne sottraiamo uno opportuno ad $\mathbf{A}$ in modo tale che la 3.21 rimanga immutata. Questo si ottiene, come è facile verificare (*Esercizio*) introducendo una funzione arbitraria Λ delle coordinate e del tempo in modo tale che

$$\begin{aligned} \phi' &= \phi + \frac{\partial \Lambda}{\partial t} \\ \mathbf{A}' &= \mathbf{A} - \nabla \Lambda. \end{aligned} \quad (3.22)$$

La trasformazione dei potenziali definita dalla 3.22 si chiama *trasformazione di gauge*³ e l'invarianza dei campi che ne segue è detta *invarianza di gauge*. Di solito tale invarianza viene sfruttata per imporre una o più condizioni supplementari direttamente sui potenziali, senza esplicitare la funzione Λ . Una scelta comune, già vista in precedenza quando è stato introdotto il potenziale vettore, consiste nell'imporre

$$\nabla \cdot \mathbf{A} = 0 \quad (3.23)$$

scelta che prende il nome di *gauge di Coulomb* o *gauge trasverso* (le ragioni per questi nomi saranno chiarite più avanti). Un'altra scelta utile è quella del *gauge di Lorentz*

$$\nabla \cdot \mathbf{A} + \frac{1}{c^2} \frac{\partial \phi}{\partial t} = 0. \quad (3.24)$$

Si verifica (*Esercizio*) che in questo gauge la funzione Λ necessaria per passare da $(\mathbf{A}, \phi)$ a $(\mathbf{A}', \phi')$ soddisfa all'equazione delle onde.

³Il termine "gauge" in inglese significa "misura" o "apparecchio di misura". Si pronuncia più o meno "gheig", con la "g" finale dolce (come la parola inglese "engage" privata della prima sillaba).

3.6.1 L'equazione d'onda per i potenziali

Per definire i potenziali sono state usate la terza e la quarta equazione di Maxwell, ovvero quelle che non contengono le sorgenti dei campi. Questo è logico: la definizione dei potenziali deve consistere in una pura relazione tra essi ed i campi. Le altre equazioni, come si è già visto nel caso statico, forniscono la relazione tra i potenziali e le sorgenti. Per trovare questa relazione introduciamo le definizioni di $\mathbf{A}$ e ϕ nella prima e quarta equazione di Maxwell

$$\begin{aligned}\rho/\epsilon_0 &= \nabla \cdot \mathbf{E} = -\nabla^2 \phi - \frac{\partial}{\partial t} \nabla \cdot \mathbf{A} \\ \nabla \times (\nabla \times \mathbf{A}) &= \nabla \times \mathbf{B} = \mu_0 \mathbf{j} + \frac{1}{c^2} \frac{\partial \mathbf{E}}{\partial t} = \mu_0 \mathbf{j} - \frac{1}{c^2} \nabla \frac{\partial \phi}{\partial t} - \frac{1}{c^2} \frac{\partial^2 \mathbf{A}}{\partial t^2}.\end{aligned}$$

Utilizzando la D.13 e riordinando

$$\begin{aligned}\nabla^2 \phi + \frac{\partial}{\partial t} \nabla \cdot \mathbf{A} &= -\rho/\epsilon_0 \\ \nabla^2 \mathbf{A} - \nabla \left(\nabla \cdot \mathbf{A} + \frac{1}{c^2} \frac{\partial \phi}{\partial t} \right) - \frac{1}{c^2} \frac{\partial^2 \mathbf{A}}{\partial t^2} &= -\mu_0 \mathbf{j}.\end{aligned}$$

Non si tratta di equazioni particolarmente attraenti. Possiamo però renderle più simmetriche utilizzando il gauge di Lorentz 3.24 per eliminare $\mathbf{A}$ dalla prima equazione e ϕ dalla seconda:

$$\begin{aligned}\square^2 \phi &= -\rho/\epsilon_0 \\ \square^2 \mathbf{A} &= -\mu_0 \mathbf{j}\end{aligned}\tag{3.25}$$

e queste invece sono quanto di meglio ci si poteva aspettare: ogni potenziale ha la sua equazione d'onda, ciascuna contenente la sorgente "giusta": cariche per ϕ , correnti per $\mathbf{A}$. Inoltre sono evidentemente una generalizzazione delle equazioni 1.16 e 2.19. Si noti come invece nel gauge di Coulomb 3.23 l'equazione per ϕ sia formalmente identica a quella di Poisson 1.16, da cui il nome di questo gauge.

3.6.2 I potenziali ritardati

Le 3.25 sono equazioni d'onda, come le 3.17 e 3.19 per i campi, ma questa volta inhomogenee. Trovarne la soluzione formale è possibile in modo completamente generale, ma qui ci limitiamo a fornirla senza dimostrazione (una giustificazione si trova nell'Appendice B):

$$\begin{aligned}
\phi(\mathbf{r}, t) &= \frac{1}{4\pi\epsilon_0} \int \frac{\rho(\mathbf{r}', t')}{|\mathbf{r} - \mathbf{r}'|} d\Omega' \\
\mathbf{A}(\mathbf{r}, t) &= \frac{\mu_0}{4\pi} \int \frac{\mathbf{j}(\mathbf{r}', t')}{|\mathbf{r} - \mathbf{r}'|} d\Omega'
\end{aligned} \tag{3.26}$$

identica a quella statica salvo che per la sostituzione di t con $t' = t \pm |\mathbf{r} - \mathbf{r}'|/c$. Possiamo quindi notare che se le sorgenti si estendono in una regione finita dello spazio e se ci poniamo a grande distanza da esse i relativi potenziali andranno a zero come $1/r$.

Le equazioni 3.26 dicono inoltre che i potenziali sono determinati dalla distribuzione di cariche e correnti *non com'è* all'istante in cui sono misurati ma da come essa *era* ad un istante anteriore di $|\mathbf{r} - \mathbf{r}'|/c$ se si sceglie il segno $-$, o da come *sarà* ad un istante posteriore di altrettanto se si sceglie il segno $+$. Nel primo caso si parla di *potenziali ritardati*, nel secondo di *potenziali anticipati*. È però difficile immaginare come possa il potenziale “sapere in anticipo” come si distribuirà la carica o la corrente che lo genera, per cui questa seconda soluzione viene sempre scartata ed i potenziali che si considerano sono solo quelli ritardati, $t' = t - |\mathbf{r} - \mathbf{r}'|/c$.

Si vede quindi come le interazioni elettromagnetiche facciano sentire i loro effetti non istantaneamente, ma con un ritardo proporzionale alla distanza, con coefficiente di proporzionalità $1/c$; questo significa che *le interazioni elettromagnetiche, nel vuoto, si propagano a velocità c* . Nel caso di campi statici questo è irrilevante, visto che non vi sono dipendenze dal tempo; il campo “è sempre stato così” e non si propaga. Nel caso dinamico invece vedremo che le conseguenze sono di enorme importanza.

3.7 Lagrangiana ed Hamiltoniana di una particella in campo elettromagnetico

3.7.1 Formulazione lagrangiana

Siccome la forza di Lorentz dipende esplicitamente dalla velocità, ed i potenziali elettrodinamici non sono semplici potenziali scalari, non è ovvio che sia possibile porre l'equazione del moto di una particella (o più particelle) sotto forma Lagrangiana. Perché questo sia possibile bisogna trovare una funzione $\mathcal{L} = \mathcal{L}(\mathbf{r}, \mathbf{v}, t)$ tale che l'equazione di Eulero-Lagrange⁴

⁴In questo paragrafo usiamo la comune convenzione secondo cui la derivata rispetto ad un vettore rappresenta il gradiente rispetto alle componenti di quel vettore.

$$\frac{d}{dt} \frac{\partial \mathcal{L}}{\partial \mathbf{v}} = \frac{\partial \mathcal{L}}{\partial \mathbf{r}} \quad (3.27)$$

coincida con l'equazione di Newton

$$m \frac{d\mathbf{v}}{dt} = q (\mathbf{E} + \mathbf{v} \times \mathbf{B}) = q (-\partial\phi/\partial\mathbf{r} - \partial\mathbf{A}/\partial t + \mathbf{v} \times \nabla \times \mathbf{A}). \quad (3.28)$$

Qui bisogna tenere ben presente la differenza tra derivata parziale e derivata totale di una funzione: la prima riguarda solo la dipendenza esplicita dalla specifica variabile rispetto alla quale si sta derivando, la seconda tiene anche conto dell'eventuale dipendenza implicita tramite le altre variabili. Così, la derivata di $\mathbf{A}$ rispetto al tempo nella 3.28 è una derivata parziale, perché connette il *campo* $\mathbf{A}$, pensato come funzione di quattro variabili indipendenti (tre spaziali ed una temporale), con il *campo* $\mathbf{E}$. Ma da qui in poi abbiamo bisogno di calcolare la variazione col tempo del valore di $\mathbf{A}$ *calcolato nel punto in cui si trova la particella*, quindi dobbiamo considerare non solo che $\mathbf{A}$ può variare col tempo, ma che la particella può esplorare in tempi diversi regioni di spazio in cui $\mathbf{A}$ assume valori diversi.

In altri termini, dobbiamo considerare $\mathbf{A} = \mathbf{A}(\mathbf{r}(t), t)$ come funzione di una sola variabile, il tempo. Pertanto c'è differenza tra derivata parziale e totale rispetto al tempo di $\mathbf{A}$, e la relazione tra di esse è data, per esempio per la componente x , da

$$\frac{dA_x}{dt} = \frac{\partial A_x}{\partial t} + \left(v_x \frac{\partial A_x}{\partial x} + v_y \frac{\partial A_x}{\partial y} + v_z \frac{\partial A_x}{\partial z} \right).$$

Torniamo ora all'equazione 3.28 e scriviamo esplicitamente la componente x dell'ultimo termine

$$\mathbf{v} \times \nabla \times \mathbf{A}|_x = v_y \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) - v_z \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right)$$

ed aggiungiamo e sottraiamo $v_x \partial A_x / \partial x$ per avere un'espressione più simmetrica:

$$\begin{aligned} \mathbf{v} \times \nabla \times \mathbf{A}|_x &= v_x \frac{\partial A_x}{\partial x} + v_y \frac{\partial A_y}{\partial x} + v_z \frac{\partial A_z}{\partial x} - v_x \frac{\partial A_x}{\partial x} - v_y \frac{\partial A_x}{\partial y} - v_z \frac{\partial A_x}{\partial z} \\ &= \frac{\partial}{\partial x} (\mathbf{A} \cdot \mathbf{v}) + \frac{\partial A_x}{\partial t} - \frac{dA_x}{dt}. \end{aligned}$$

Inserendo quest'espressione, componente per componente, nell'equazione 3.28 otteniamo

$$m \frac{d\mathbf{v}}{dt} = q \left[\frac{\partial}{\partial \mathbf{r}} (-\phi + \mathbf{A} \cdot \mathbf{v}) - \frac{d\mathbf{A}}{dt} \right] \Rightarrow \frac{d}{dt} (m\mathbf{v} + q\mathbf{A}) = \frac{\partial}{\partial \mathbf{r}} (-q\phi + q\mathbf{A} \cdot \mathbf{v})$$

e si vede che per soddisfare la 3.27 è sufficiente trovare $\mathcal{L}$ tale che

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial \mathbf{v}} &= m\mathbf{v} + q\mathbf{A} \\ \frac{\partial \mathcal{L}}{\partial \mathbf{r}} &= \frac{\partial}{\partial \mathbf{r}} (-q\phi + q\mathbf{A} \cdot \mathbf{v}) \end{aligned}$$

Ricordando che né $\mathbf{A}$ né ϕ dipendono da $\mathbf{v}$ si ha come possibile Lagrangiana:

$$\mathcal{L} = \frac{1}{2}mv^2 - q\phi + q\mathbf{A} \cdot \mathbf{v}. \quad (3.29)$$

La forma della Lagrangiana 3.29 è piacevolmente semplice e simmetrica, ma non è l'unica possibile. Sappiamo infatti che una trasformazione di gauge 3.22 lascia invariati i campi, e di conseguenza anche le equazioni del moto. Che succede alla Lagrangiana sotto una tale trasformazione? Succede che diventa un'altra

$$\begin{aligned} \mathcal{L}' &= \frac{1}{2}mv^2 - q\phi' + q\mathbf{A}' \cdot \mathbf{v} = \frac{1}{2}mv^2 - q\phi - q\frac{\partial \Lambda}{\partial t} + q\mathbf{A} \cdot \mathbf{v} - q\mathbf{v} \cdot \nabla \Lambda \\ &= \mathcal{L} - q \left(\frac{\partial \Lambda}{\partial t} + \mathbf{v} \cdot \nabla \Lambda \right) = \mathcal{L} + \frac{d(-q\Lambda)}{dt} \end{aligned}$$

che però differisce dalla prima solo per la derivata totale di una funzione delle coordinate e del tempo, e fornisce di conseguenza le stesse equazioni del moto. Questo è perfettamente coerente col fatto che la scelta del gauge non può influenzare nulla di osservabile.

3.7.2 Formulazione hamiltoniana

Ottenere l'Hamiltoniana ora è semplice. Seguendo le normali prescrizioni, il momento generalizzato diventa

$$\mathbf{p} = \frac{\partial \mathcal{L}}{\partial \mathbf{v}} = m\mathbf{v} + q\mathbf{A} \quad (3.30)$$

e non coincide più con la familiare espressione $m\mathbf{v}$. L'hamiltoniana allora diventa

$$\mathcal{H} = \mathbf{p} \cdot \mathbf{v} - \mathcal{L} = mv^2 + q\mathbf{A} \cdot \mathbf{v} - \frac{1}{2}mv^2 + q\phi - q\mathbf{A} \cdot \mathbf{v} = \frac{1}{2}mv^2 + q\phi$$

da cui si vede che l'introduzione di un potenziale vettore non cambia la forma dell'hamiltoniana, se espressa in termini di v . Nel caso statico $\mathcal{H}$ coincide con l'energia ed $\mathbf{A}$ determina solo il campo magnetico, per cui vediamo confermato quanto detto nella sezione 2.2: un campo magnetico statico non cambia l'energia delle particelle.

Ovviamente questo non significa che le equazioni del moto siano le stesse che in assenza di campo, visto che tali equazioni si devono derivare dall'hamiltoniana espressa in funzione delle variabili "giuste" $\mathbf{p}$ e $\mathbf{r}$. Visto che $\mathbf{v} = (\mathbf{p} - q\mathbf{A})/m$, $\mathcal{H}$ diventa

$$\mathcal{H} = \frac{(\mathbf{p} - q\mathbf{A})^2}{2m} + q\phi. \quad (3.31)$$

Si vede quindi che l'unico effetto della presenza di un potenziale vettore è di ridefinire il momento generalizzato $\mathbf{p}$.

3.8 Leggi di conservazione per il campo elettromagnetico

3.8.1 Energia

Un sistema di cariche elettriche interagenti possiede, evidentemente, una certa energia potenziale. Si è visto nel paragrafo 1.5.1 che se prendiamo un condensatore nel vuoto a facce piane e parallele la densità di energia accumulata $u = U/\Omega$ è proporzionale al quadrato del campo elettrico.

Analogamente, in un solenoide (*Esercizio*)

$$U = \frac{1}{2}LI^2 = \Omega \frac{B^2}{2\mu_0} \quad (3.32)$$

ed anche in questo caso la densità di energia (che, a rigore, non è un'energia potenziale nel senso della meccanica in quanto non esiste un potenziale scalare magnetico) è proporzionale al quadrato del campo, in questo caso magnetico.

Quanto sono generali questi risultati? Per scoprirlo, calcoliamo la potenza $W = \sum_i \mathbf{F}_i \cdot \mathbf{v}_i$ trasferita dalle cariche, o alle cariche, tramite il campo elettromagnetico:

$$W = \sum_i q_i (\mathbf{E}_i + \mathbf{v}_i \times \mathbf{B}_i) \cdot \mathbf{v}_i = \sum_i q_i \mathbf{E}_i \cdot \mathbf{v}_i = \int \rho \mathbf{E} \cdot \mathbf{v} d\Omega = \int \mathbf{E} \cdot \mathbf{j} d\Omega.$$

Esprimiamo l'integrando $w = \mathbf{E} \cdot \mathbf{j}$, che è una densità di potenza trasferita da/alle cariche, in funzione dei soli campi, sfruttando le equazioni di Maxwell 3.15:

$$\begin{aligned} w &= \mathbf{E} \cdot \left(\frac{1}{\mu_0} \nabla \times \mathbf{B} - \epsilon_0 \frac{\partial \mathbf{E}}{\partial t} \right) = \frac{1}{\mu_0} \mathbf{B} \cdot \nabla \times \mathbf{E} - \frac{1}{\mu_0} \nabla \cdot \mathbf{E} \times \mathbf{B} - \frac{\epsilon_0}{2} \frac{\partial E^2}{\partial t} \\ &= -\frac{\mathbf{B}}{\mu_0} \frac{\partial \mathbf{B}}{\partial t} - \frac{\epsilon_0}{2} \frac{\partial E^2}{\partial t} - \frac{1}{\mu_0} \nabla \cdot \mathbf{E} \times \mathbf{B} = -\frac{1}{2} \frac{\partial}{\partial t} \left(\frac{B^2}{\mu_0} + \epsilon_0 E^2 \right) - \nabla \cdot \mathbf{N} \end{aligned}$$

dove è stata usata l'eq. D.12 ed è stato introdotto il *vettore di Poynting*

$$\mathbf{N} = \frac{1}{\mu_0} \mathbf{E} \times \mathbf{B}. \quad (3.33)$$

Definendo $u = B^2/2\mu_0 + \epsilon_0 E^2/2$ e $U = \int u d\Omega$ otteniamo

$$w + \frac{\partial u}{\partial t} + \nabla \cdot \mathbf{N} = 0 \quad (3.34)$$

da cui integrando sul volume e sfruttando il teorema di Gauss otteniamo la forma integrale

$$W + \frac{dU}{dt} + \oint \mathbf{N} \cdot d\mathbf{S} = 0. \quad (3.35)$$

Queste equazioni, se le priviamo del primo termine, hanno un aspetto familiare. Infatti sono formalmente identiche alle due forme dell'equazione di continuità 1.23 e 1.24, se interpretiamo il flusso del vettore di Poynting come flusso di energia attraverso una superficie⁵. Ovvero, possiamo interpretare le 3.34 e 3.35 (con $W = w = 0$) come la legge che stabilisce la conservazione dell'energia per un sistema di campi elettromagnetici che non siano in interazione con particelle cariche.

Se ora reintroduciamo il termine W (o w), vediamo che l'energia del campo elettromagnetico non si conserva più, e che la mancata conservazione è interamente dovuta all'energia che viene trasferita dalle o alle cariche. Se ne conclude quindi che *il campo elettromagnetico possiede una*

⁵Una superficie *chiusa*. L'interpretazione non vale per superfici aperte, e tutta la deduzione riportata infatti fa riferimento a superfici chiuse.

sua energia e che in presenza di particelle cariche, questa energia si conserva se, e solo se, viene tenuto conto dell'energia meccanica posseduta dalle cariche.

Lo stesso vale per le particelle: la legge di conservazione dell'energia per un sistema di particelle cariche è valida se e solo se si tiene conto anche dell'energia del campo elettromagnetico.

Il termine di flusso può spesso essere annullato scegliendo una superficie tale che il vettore di Poynting si annulli, per esempio scegliendo una superficie che va all'infinito se i campi vanno a zero più rapidamente di $1/r$. Questo peraltro, per campi statici, è sempre il caso in assenza di cariche o correnti all'infinito.

Dalla discussione precedente emerge una visione alternativa della trasmissione di energia nei circuiti. Prendiamo ad esempio un semplicissimo circuito in corrente continua, formato da un generatore di tensione V e da una resistenza R . Quest'ultima la schematizziamo come un cilindretto uniforme di raggio r e lunghezza l . La corrente che scorre è $I = V/R$ e la potenza trasferita dal generatore alla resistenza, e da questa dissipata in calore, è $W = RI^2$.

Siccome i campi sono statici, $U = \text{costante}$ ed il secondo termine della 3.35 è nullo. Il campo elettrico nella resistenza vale in modulo $E = V/l = RI/l$ ed è longitudinale, mentre il campo magnetico sulla superficie laterale vale, per la 2.8, $B = \mu_0 I/(2\pi r)$ ed è ad essa tangente ed è perpendicolare alla lunghezza del cilindro. Sulle basi invece $\mathbf{B}$ è parallelo ad esse.

Il vettore di Poynting perciò è parallelo alle basi del cilindro, ed ivi il suo flusso è nullo. Sulla superficie laterale invece $\mathbf{N}$ è ad essa perpendicolare e punta verso l'interno. Il modulo è $N = EB/\mu_0 = RI^2/(2\pi r l)$ ed il flusso $\oint \mathbf{N} \cdot d\mathbf{S} = 2\pi r l RI^2/(2\pi r l) = RI^2$. Ne segue che l'energia che il generatore di tensione trasmette alla resistenza la si può considerare trasmessa tramite il campo elettromagnetico, attraverso la superficie della resistenza (e non tramite il filo).

3.8.2 Quantità di moto

Ragionamenti simili si possono svolgere per la quantità di moto. Vediamo che la densità di quantità di moto $\mathcal{P}$ è esprimibile in termini della densità di forza di Lorentz $\mathbf{f}$, e tramite essa e le equazioni di Maxwell in termini del campo elettromagnetico:

$$\frac{\partial \mathcal{P}}{\partial t} = \mathbf{f} = \rho \mathbf{E} + \mathbf{j} \times \mathbf{B} = \epsilon_0 \mathbf{E} \nabla \cdot \mathbf{E} + \left(\frac{1}{\mu_0} \nabla \times \mathbf{B} - \epsilon_0 \frac{\partial \mathbf{E}}{\partial t} \right) \times \mathbf{B}$$

$$= \epsilon_0 \mathbf{E} \nabla \cdot \mathbf{E} + \underbrace{\frac{1}{\mu_0} \mathbf{B} \nabla \cdot \mathbf{B}}_{=0} + \frac{1}{\mu_0} (\nabla \times \mathbf{B}) \times \mathbf{B} + \epsilon_0 \mathbf{E} \times \frac{\partial \mathbf{B}}{\partial t} - \epsilon_0 \mu_0 \frac{\partial \mathbf{N}}{\partial t}$$

dove si è sfruttato il fatto che $\mu_0 \frac{\partial \mathbf{N}}{\partial t} = \frac{\partial}{\partial t} (\mathbf{E} \times \mathbf{B}) = \frac{\partial \mathbf{E}}{\partial t} \times \mathbf{B} + \mathbf{E} \times \frac{\partial \mathbf{B}}{\partial t}$. Mettendo tutto in forma più simmetrica, utilizzando la D.14 e ricordando la definizione di c ,

$$\begin{aligned} \frac{\partial \mathcal{P}}{\partial t} &= \epsilon_0 \mathbf{E} \nabla \cdot \mathbf{E} + \frac{1}{\mu_0} \mathbf{B} \nabla \cdot \mathbf{B} + \frac{1}{\mu_0} (\nabla \times \mathbf{B}) \times \mathbf{B} + \epsilon_0 (\nabla \times \mathbf{E}) \times \mathbf{E} - \frac{1}{c^2} \frac{\partial \mathbf{N}}{\partial t} \\ &= \epsilon_0 \mathbf{E} \nabla \cdot \mathbf{E} + \epsilon_0 \left[(\mathbf{E} \cdot \nabla) \mathbf{E} - \frac{1}{2} \nabla E^2 \right] \\ &\quad + \frac{1}{\mu_0} \mathbf{B} \nabla \cdot \mathbf{B} + \frac{1}{\mu_0} \left[(\mathbf{B} \cdot \nabla) \mathbf{B} - \frac{1}{2} \nabla B^2 \right] - \frac{1}{c^2} \frac{\partial \mathbf{N}}{\partial t}. \end{aligned}$$

Si tratta di un'espressione abbastanza complicata, per cui calcoliamola per una sola componente $x_\alpha \in \{x, y, z\}$, limitandoci alla parte relativa al campo elettrico (per quello magnetico il calcolo sarà identico). Abbiamo:

$$\begin{aligned} &\left\{ \mathbf{E} \nabla \cdot \mathbf{E} + \left[(\mathbf{E} \cdot \nabla) \mathbf{E} - \frac{1}{2} \nabla E^2 \right] \right\}_\alpha = \\ &= \sum_\beta \left\{ E_\alpha \frac{\partial E_\beta}{\partial x_\beta} + \left[E_\beta \frac{\partial E_\alpha}{\partial x_\beta} - \frac{1}{2} \frac{\partial E_\beta^2}{\partial x_\alpha} \right] \right\} = \sum_\beta \frac{\partial}{\partial x_\beta} \left(E_\alpha E_\beta - \frac{1}{2} E^2 \delta_{\alpha\beta} \right). \end{aligned}$$

L'ultimo termine può essere visto come la divergenza di un tensore di ordine 2 che ha la forma, reinserendo il fattore ϵ_0 ed il termine magnetico,

$$\mathcal{T}_{\alpha\beta} = \epsilon_0 E_\alpha E_\beta + \frac{1}{\mu_0} B_\alpha B_\beta - \frac{1}{2} \left(\epsilon_0 E^2 + \frac{1}{\mu_0} B^2 \right) \delta_{\alpha\beta} \quad (3.36)$$

ed è noto come *tensore degli sforzi di Maxwell*. L'equazione del bilancio della quantità di moto assume finalmente la forma differenziale

$$\frac{\partial}{\partial t} \left(\mathcal{P} + \frac{1}{c^2} \mathbf{N} \right) = \nabla \cdot \mathcal{T} \quad (3.37)$$

e quella integrale

$$\frac{d}{dt} \left(\mathbf{p} + \frac{1}{c^2} \int \mathbf{N} d\Omega \right) = \oint \mathcal{T} \cdot d\mathbf{S}. \quad (3.38)$$

Anche qui come nell'espressione dell'energia il termine di flusso può spesso essere annullato scegliendo una superficie tale che il tensore degli

sforzi sia nullo: per esempio scegliendo una superficie che va all'infinito, se i campi vanno a zero più rapidamente di $1/r$.

Si possono ripetere qui gli stessi commenti fatti per le equazioni 3.35 e 3.34: *il campo elettromagnetico possiede una sua quantità di moto, e in presenza di particelle cariche tale quantità di moto si conserva se, e solo se, viene tenuto conto della quantità di moto meccanica posseduta dalle particelle.*

Da un altro punto di vista, vediamo che *la legge di conservazione della quantità di moto per un sistema di particelle cariche è valida se e solo se si tiene conto anche della quantità di moto del campo elettromagnetico.* In particolare, se c'è un certo ammontare di quantità di moto "in volo" tra due cariche interagenti, *non possiamo aspettarci che valga la terza legge di Newton* (di azione e reazione), che deve quindi essere soppiantata dalla legge più generale 3.37 o 3.38.

Si noti che la densità di quantità di moto del campo elettromagnetico è data dal vettore di Poynting che svolge il ruolo di "flusso di energia" nell'equazione di conservazione dell'energia; il flusso della quantità di moto è dato dal flusso del tensore degli sforzi. Vediamo così da questo ruolo duale di $\mathbf{N}$ che nel campo elettromagnetico energia e quantità di moto sono strettamente collegate, esattamente come per una particella materiale libera, per la quale $U = p^2/2m$. Questa connessione si vede anche dal fatto che la traccia (somma dei termini diagonali) di $\mathcal{T}$ non è altro che l'energia del campo elettromagnetico.

Il campo elettromagnetico possiede dunque proprietà dinamiche indipendentemente dalla presenza o meno di materia ponderabile, e questo conduce a vederlo come dotato di una sua esistenza al di là del ruolo di puro intermediario tra le particelle per cui era stato introdotto.

Capitolo 4

Onde

4.1 Propagazione per onde

Si è visto che dalle equazioni di Maxwell è possibile dedurre delle equazioni per il campo elettrico, quello magnetico e per i potenziali, equazioni che, nel vuoto, hanno tutte la stessa forma:

$$\square^2 f(\mathbf{r}, t) = \nabla^2 f(\mathbf{r}, t) - \frac{1}{c^2} \frac{\partial^2 f(\mathbf{r}, t)}{\partial t^2} = 0 \quad (4.1)$$

dove $f(\mathbf{r}, t)$ nel caso elettromagnetico è una funzione vettoriale, ma in altri casi può essere, *mutatis mutandis*, una funzione scalare. Infatti tale equazione descrive molti altri fenomeni fisici: la propagazione del suono in un gas, delle onde di superficie in un fluido, delle vibrazioni di una corda e nei solidi, eccetera. La grandezza f di volta in volta rappresenta una diversa grandezza fisica: la pressione locale del gas, il livello del fluido, lo spostamento della corda o degli atomi. L'equazione 4.1 è dunque una delle equazioni più importanti e versatili della fisica. Perciò verrà qui trattata, per quanto possibile, nella sua generalità senza fare necessariamente sempre riferimento al caso elettromagnetico.

4.1.1 Onde piane e onde sferiche

Per studiare il tipo di soluzioni che tale equazione può avere, partiamo dal caso scalare e unidimensionale. Con quest'ultima espressione qui s'intende che consideriamo una funzione $f = f(x, t)$ che, ad un dato tempo, dipenda solo da x , per cui abbia lo stesso valore indipendentemente da y e z . I punti dello spazio in cui f ha un determinato valore rappresentano dei fasci di piani paralleli, da cui il nome di *onde piane*.

L'equazione 4.1 diventa, per le onde piane,

$$\left(\frac{\partial^2}{\partial x^2} - \frac{1}{c^2} \frac{\partial^2}{\partial t^2} \right) f(x, t) = 0 \quad (4.2)$$

e l'operatore a primo membro somiglia vagamente al prodotto notevole algebrico $a^2 - b^2 = (a + b)(a - b)$. Sfruttiamo questa somiglianza cercando di fattorizzarlo, ed a tal fine poniamo $s = x - ct$ e $\sigma = x + ct$, ottenendo, in funzione di queste nuove variabili,

$$\begin{aligned} \frac{\partial}{\partial x} &= \frac{\partial s}{\partial x} \frac{\partial}{\partial s} + \frac{\partial \sigma}{\partial x} \frac{\partial}{\partial \sigma} = \frac{\partial}{\partial s} + \frac{\partial}{\partial \sigma}; & \frac{1}{c} \frac{\partial}{\partial t} &= \frac{\partial}{\partial s} - \frac{\partial}{\partial \sigma} \\ \frac{\partial^2}{\partial x^2} &= \left(\frac{\partial}{\partial s} + \frac{\partial}{\partial \sigma} \right)^2 = \frac{\partial^2}{\partial s^2} + \frac{\partial^2}{\partial \sigma^2} + 2 \frac{\partial^2}{\partial s \partial \sigma}; & \frac{1}{c^2} \frac{\partial^2}{\partial t^2} &= \frac{\partial^2}{\partial s^2} + \frac{\partial^2}{\partial \sigma^2} - 2 \frac{\partial^2}{\partial s \partial \sigma} \end{aligned}$$

per cui, posto $f(x, t) = f[x(s, \sigma), t(s, \sigma)] = g(s, \sigma)$ l'equazione diventa

$$\frac{\partial^2 g(s, \sigma)}{\partial s \partial \sigma} = 0$$

che ha soluzioni della forma

$$g(s, \sigma) = h_1(s) + h_2(\sigma) \Rightarrow f(x, t) = f_1(x - ct) + f_2(x + ct).$$

La soluzione generale dell'equazione 4.2 è perciò composta da funzioni in sé arbitrarie, ma che dipendono da una combinazione particolare delle variabili indipendenti. In special modo si vede che, supponendo per esempio $f_2 = 0$, il profilo della funzione $f(x - ct)$ è identico a quello della condizione iniziale $f(x, 0) = \psi(x)$ ma spostato verso destra di una distanza pari a ct . In altri termini, la funzione f_1 rappresenta una funzione che si sposta rigidamente in direzione dell'asse x , nel verso delle x crescenti, a velocità c . Similmente, la funzione f_2 rappresenta una funzione che si sposta rigidamente nel senso delle x decrescenti a velocità c . Si noti che questo dipende solo dal segno *relativo* di x e t ; una funzione $f(ct - x)$ rappresenta ugualmente un'onda in moto nel verso positivo dell'asse x .

Caratteristica delle onde è quindi di essere una perturbazione di un campo esteso, di solito un campo continuo, che si sposta più o meno rigidamente nello spazio. In senso stretto, comunque, si chiamerà onda ogni perturbazione di un campo che obbedisce all'equazione 4.1. In senso lato, viene solitamente chiamata onda qualunque perturbazione in un campo di grandezze.

Un altro tipo importante di onde è quello delle onde sferiche, ovvero onde che si diramano simmetricamente da un punto, poniamo il punto $r = 0$.

In questo caso, f è una funzione solo del modulo di r , e nel laplaciano dell'equazione 4.1 espresso in termine delle coordinate polari sferiche sopravvive solo il termine radiale (cfr. D.5):

$$\nabla^2 f(r) = \frac{1}{r^2} \frac{\partial}{\partial r} \left[r^2 \frac{\partial f(r)}{\partial r} \right] = \frac{1}{r} \frac{\partial^2}{\partial r^2} [r f(r)]. \quad (4.3)$$

Ne segue che la soluzione generale dell'equazione 4.1 nel caso di onde sferiche si ottiene da quella della 4.2 mediante la sostituzione $f \rightarrow r f$, ovvero

$$f(x, t) = \frac{1}{r} f_1(x - ct) + \frac{1}{r} f_2(x + ct).$$

Si tratta quindi di onde la cui ampiezza decade come $1/r$. È evidente che a grandi distanze dall'origine un'onda sferica è approssimabile *localmente* con un'onda piana.

4.1.2 Onde monocromatiche

Considereremo in particolare le onde piane (co)sinusoidali:

$$f(x, t) = f_0 e^{i(kx - \omega t)} \quad (4.4)$$

le cui parti reali e immaginarie hanno evidentemente la forma richiesta: basta riscrivere il termine tra parentesi, detto *fase*, come $k(x - \omega t/k)$ e porre $\omega = ck$. La grandezza f_0 si chiama *ampiezza* dell'onda e può anche essere, in questo formalismo, una grandezza complessa. I punti dello spazio in cui la fase ha un valore fissato si chiamano *fronti d'onda* e nel caso di onde piane sono evidentemente dei piani perpendicolari alla direzione di propagazione.

La periodicità spaziale — ovvero la distanza tra due punti con lo stesso valore di f a t fissato — è data dalla *lunghezza d'onda* $\lambda = 2\pi/k$; k è detto *numero d'onda*. La periodicità temporale — ovvero l'intervallo di tempo che intercorre tra il passaggio di due punti corrispondenti di f in un punto fissato — è data dal *periodo* $T = 2\pi/\omega = 1/\nu$; le grandezze ω e ν si misurano entrambe in s^{-1} (l'unità di misura corrispondente è detta Hertz) e sono chiamate rispettivamente *pulsazione* e *frequenza* dell'onda (a volte si usa il termine frequenza anche per ω). È pure immediato (*Esercizio*) vedere che $\lambda\nu = c$.

L'ipotesi di onda sinusoidale non è limitativa perché, come spiegato nell'Appendice B, ogni funzione dotata di sufficienti regolarità può essere scritta come combinazione lineare, detta anche *sovrapposizione*, di funzioni

del tipo della 4.4; data la linearità dell'equazione delle onde 4.1 le conclusioni tratte dallo studio di una di queste funzioni periodiche potranno essere facilmente estese ad una funzione qualsiasi. Qualora però l'onda sia esattamente del tipo 4.4 e non una sovrapposizione di esse, e si abbia quindi un ben determinato valore di ω , l'onda è detta *monocromatica*; la ragione di questa denominazione sarà vista più avanti.

Le considerazioni precedenti si estendono facilmente al caso di onde piane in tre dimensioni ed al caso di campi vettoriali. È sufficiente scrivere, invece dell'equazione 4.4, l'equazione

$$\mathbf{f}(\mathbf{r}, t) = \mathbf{f}_0 e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)} \quad (4.5)$$

dove sia $\mathbf{f}$ che $\mathbf{f}_0$ che $\mathbf{k}$ sono ora vettori. In particolare $\mathbf{k}$ prende il nome di *vettore d'onda* ed è evidentemente perpendicolare ai piani a fase costante. Siccome, come si è visto, l'onda si propaga perpendicolarmente a questi piani, ne segue che $\mathbf{k}$ indica la direzione di propagazione dell'onda. Il modulo di $\mathbf{k}$ evidentemente vale $2\pi/\lambda$.

4.2 Velocità di fase e velocità di gruppo

La definizione di velocità data nel paragrafo precedente si riferisce alla velocità di propagazione di un determinato punto di un'onda; per esempio, in un'onda sinusoidale, la velocità di spostamento di una cresta o di un ventre. Tuttavia, quando la forma dell'onda subisce deformazioni importanti durante la propagazione, questa velocità diventa non solo di dubbia utilità, ma anche difficile da definire.

Deformazioni nella forma d'onda sono possibili quando il coefficiente c nell'equazione 4.1 non è costante. Di solito conviene descrivere questa occorrenza definendo una velocità diversa per ogni onda sinusoidale, vale a dire ponendo $c = c(\omega)$. In questo modo le componenti a determinate frequenze viaggeranno più velocemente di quelle a frequenze differenti e tenderanno a "sorpassarle"; la forma d'onda totale, somma di tutte le forme d'onda delle componenti, di conseguenza non si conserverà (si pensi al gruppo formato dai concorrenti ad una maratona, inizialmente compatto ed in seguito sgranato per via delle differenze di velocità dei corridori).

Vediamo questo fenomeno in un semplice caso specifico, quello di un'onda scalare piana composta da due onde sinusoidali di frequenze e lunghezze d'onda poco diverse ($k' = k + \Delta k$, $\Delta k \ll k$; $\omega' = \omega + \Delta\omega$, $\Delta\omega \ll \omega$) e con la stessa ampiezza:

$$\begin{aligned}
f &= f_0 \sin(kx - \omega t) + f_0 \sin(k'x - \omega' t) \\
&= 2f_0 \cos \frac{1}{2}[(k' - k)x - (\omega' - \omega)t] \sin \frac{1}{2}[(k' + k)x - (\omega' + \omega)t] \\
&\simeq 2f_0 \cos \frac{1}{2}(\Delta kx - \Delta \omega t) \sin(kx - \omega t).
\end{aligned}$$

L'onda complessiva quindi risulta essere un'onda con approssimativamente la stessa frequenza (*frequenza portante*) e lunghezza d'onda delle onde componenti, ma *modulata* da un termine che a sua volta si propaga con velocità $\Delta\omega/\Delta k \rightarrow d\omega/dk$, che non necessariamente coincide con c . In effetti la *velocità di gruppo*

$$v_g = \frac{d\omega}{dk} = c + k \frac{dc}{dk} \quad (4.6)$$

coincide con la *velocità di fase* $c = \omega/k$ solo se quest'ultima non dipende da k , quindi da ω . Qualora questa condizione non si verifichi si dice di essere in presenza di un *mezzo dispersivo*.

Come si vede, velocità di fase e di gruppo in linea generale sono grandezze diverse. Dal punto di vista pratico la velocità di fase è quella con meno importanza. Essa si riferisce infatti ad onde esattamente monocromatiche, le quali a rigore non esistono: esse infatti avrebbero durata infinita. Di solito si è interessati a *treni d'onde* di durata limitata, e la velocità con cui si propagano effetti fisici ed informazione è quella di gruppo. Non è quindi vietato, e di fatto si verifica in diversi casi, che la velocità di fase delle onde elettromagnetiche sia superiore a quella della onde elettromagnetiche nel vuoto.

4.2.1 Onde di compressione nei solidi

Un caso importante di onde dispersive si ha nel caso dei solidi. Schematizziamo il solido come una sequenza infinita di atomi tutti di massa m , costretti a muoversi in una sola dimensione — diciamo l'asse x — connessi da molle di costante elastica G solo con i primi vicini (cioè l' n -mo solo con l' $n-1$ -mo e l' $n+1$ -mo), e con posizioni di riposo $x_n = an$, $n \in \mathbb{Z}$. La posizione dell' n -mo atomo, all'istante t , sia $x_n(t) = x_n + u_n(t) = an + u_n(t)$, ed ipotizzeremo se necessario piccole oscillazioni: $u_n \ll a \forall n$. L'equazione del moto per l'atomo n -mo è data dunque da

$$m \frac{d^2 u}{dt^2} = F_n = G(u_{n+1} - u_n) - G(u_n - u_{n-1}) \quad (4.7)$$

visto che le forze agenti sono proporzionali alle differenze degli spostamenti degli atomi dalle loro posizioni di equilibrio. Cerchiamo soluzioni del tipo “onda viaggiante”:

$$u_n(t) = \varepsilon e^{i(kan - \omega t)} \quad (4.8)$$

dove compare l'indice intero n invece della coordinata reale x in quanto stiamo trattando un sistema discreto e non continuo. Introducendo la 4.8 nella 4.7 otteniamo la relazione

$$\begin{aligned} - m\omega^2 \varepsilon e^{i(kan - \omega t)} &= G\varepsilon [e^{i(ka(n+1) - \omega t)} - 2e^{i(kan - \omega t)} + e^{i(ka(n-1) - \omega t)}] \\ &= G\varepsilon e^{i(kan - \omega t)} [e^{ika} + e^{-ika} - 2] \end{aligned}$$

da cui si ottiene la *relazione di dispersione*

$$\omega^2 = \frac{2G}{m} [1 - \cos(ka)] \Rightarrow \omega = 2\sqrt{\frac{G}{m}} \sin(ka/2). \quad (4.9)$$

Questa relazione non è lineare, e pertanto la velocità di gruppo delle onde di compressione varia con la lunghezza d'onda:

$$v_g(k) = \frac{d\omega}{dk} = a\sqrt{\frac{G}{m}} \cos(ka/2).$$

Per lunghezze d'onda estremamente corte, paragonabili alla spaziatura degli atomi nel reticolo, le onde si muovono a velocità via via più basse, fino a raggiungere velocità nulla quando $\lambda = 2a$ (*Esercizio: verificarlo*). Per onde di lunghezza macroscopica invece la *velocità del suono* $v = \lim_{\lambda \rightarrow \infty} (\omega/k) = a\sqrt{\frac{G}{m}}$ è indipendente dalla frequenza e pertanto velocità di fase e di gruppo coincidono.

4.3 Effetto Doppler

Uno degli effetti più vistosi e noti della meccanica delle onde — vistoso e noto, però, solo nel mondo moderno dove siamo abituati a velocità dell'ordine delle decine di chilometri all'ora — consiste nella variazione apparente della frequenza di un'onda a seconda dello stato di moto della sorgente o dell'osservatore.

Consideriamo una sorgente S di onde, per esempio sonore, ed un osservatore O , distanti l al tempo $t = 0$. Sia c la velocità dell'onda e ν la

sua frequenza, ed immaginiamo che S ed O siano in moto con velocità rispettivamente v_s e v_o lungo la stessa direzione. Supporremo per fissare le idee che le velocità siano concordi, ma nulla cambia se sono discordi, salvo qualche segno. Possiamo invece scartare il caso $v_o \geq c$ in quanto in questo caso i fronti emessi da S non raggiungono mai O .

Un fronte d'onda emesso da S a $t = 0$ raggiungerà O non al tempo $t = l/c$ ma ad un tempo successivo t_1 , in quanto l'osservatore durante il tempo di volo dell'onda si è allontanato. Si deve avere

$$ct_1 = l + v_o t_1 \Rightarrow t_1 = \frac{l}{c - v_o}$$

Al tempo successivo $\tau = 1/\nu$ viene emesso un secondo fronte d'onda. In questo momento S si è spostato di una distanza $v_s \tau$ rispetto al punto in cui si trovava a $t = 0$, ed O di una distanza $v_o \tau$. Siccome O continua a spostarsi durante il viaggio del secondo fronte d'onda, questo lo raggiunge quando O si è spostato di un'ulteriore distanza $c(t_2 - \tau)$, ed il tempo t_2 in cui questo avviene è di conseguenza dato dall'equazione

$$v_s \tau + c(t_2 - \tau) = l + v_o t_2 \Rightarrow t_2 = \frac{l + (c - v_s)\tau}{c - v_o}$$

per cui O sente i fronti d'onda consecutivi arrivare ad intervalli

$$\tau' = t_2 - t_1 = \frac{c - v_s}{c - v_o} \tau. \quad (4.10)$$

Si noti che per $v_s = c$ i fronti d'onda arrivano con intervallo nullo. Questa stranezza è dovuta al fatto che la sorgente, muovendosi alla stessa velocità dell'onda, emette un fronte senza che il precedente abbia potuto "staccarsi". I fronti pertanto si accavallano, e questo corrisponde alla formazione di un'onda d'urto, che nel caso delle onde sonore genera il cosiddetto *bang supersonico*.

L'intervallo di tempo τ' è proprio il periodo dell'onda percepito da O , per cui O percepisce una frequenza *diversa* da quella emessa, e precisamente

$$\nu' = \frac{c - v_o}{c - v_s} \nu. \quad (4.11)$$

La frequenza percepita è perciò più alta di quella emessa se sorgente ed osservatore si avvicinano, mentre è più bassa quando si allontanano, e questo è bene in accordo con l'esperienza quotidiana. La distanza invece non ha alcuna importanza (influirà ovviamente sull'intensità dell'onda), né importano le accelerazioni.

Per velocità v_s e v_o entrambe molto più piccole di c possiamo applicare lo sviluppo binomiale

$$(1+x)^\alpha = 1 + \alpha x + \frac{\alpha(\alpha-1)}{2!}x^2 + \dots \quad (4.12)$$

all'ordine più basso e pertanto si ottiene

$$\nu' \simeq [1 - (v_o - v_s)/c]\nu \quad (4.13)$$

e questa formula si può facilmente generalizzare al caso in cui S ed O non si muovano sulla stessa retta; se θ è l'angolo fra le direzioni dei moti si ha (*Esercizio*)

$$\nu' \simeq [1 - (v_o - v_s) \cos \theta / c]\nu. \quad (4.14)$$

Si noti che è solo in questa approssimazione che la formula è simmetrica nelle velocità di osservatore e sorgente; in generale lo spostamento di frequenza dipende non solo dalla velocità relativa ma anche dalla “velocità assoluta” di S ed O , ovvero la velocità che O ed S hanno rispetto al mezzo in cui la velocità dell'onda è c .

Se immaginiamo di avere una sirena che emette un La centrale (circa 440 Hz) e che si muova a 20 m/s=72 km/h rispetto ad un osservatore fermo rispetto all'aria, considerando che la velocità del suono è di circa 330 m/s, la frequenza percepita sarà di circa 25 Hz più alta o più bassa. Tale differenza corrisponde a circa un semitono ed è quindi perfettamente percepibile. Velocità relative inferiori ai 5 m/s = 18 km/h danno d'altra parte differenze di meno di un ottavo di tono, difficilmente distinguibili da parte di un orecchio non allenato. Non sorprende quindi che la teoria dell'effetto Doppler, pur semplice, risalga solo al 1842 e la sua verifica sperimentale a due anni dopo; in precedenza non si aveva avuto modo di prendere in considerazione il problema.

4.4 Le onde elettromagnetiche

4.4.1 Onde elettromagnetiche piane

Vediamo cosa si può dire di specifico sulle onde piane elettromagnetiche nel vuoto. Non perdiamo in generalità se supponiamo che l'onda si propaghi lungo l'asse $\hat{x} = \hat{k}$ e che sia monocromatica. I campi pertanto possono essere scritti come

$$\begin{aligned}\mathbf{E} &= \mathbf{E}_0 e^{i(kx - \omega t)} \\ \mathbf{B} &= \mathbf{B}_0 e^{i(kx - \omega t)}.\end{aligned}\quad (4.15)$$

Le equazioni alla divergenza nel vuoto $\nabla \cdot \mathbf{E} = \nabla \cdot \mathbf{B} = 0$ implicano immediatamente (*Esercizio*) che $\mathbf{E}$ e $\mathbf{B}$ sono perpendicolari alla direzione di propagazione $\mathbf{k}$ ($E_x = B_x = 0$)¹, a meno di un campo uniforme.

Peraltro scrivendo le equazioni al rotore nel caso di spazio vuoto, ovvero $\rho = 0$, $\mathbf{j} = 0$, otteniamo

$$\begin{aligned}\nabla \times \mathbf{E} &= -\frac{\partial \mathbf{B}}{\partial t} \\ \nabla \times \mathbf{B} &= \epsilon_0 \mu_0 \frac{\partial \mathbf{E}}{\partial t} = \frac{1}{c^2} \frac{\partial \mathbf{E}}{\partial t}.\end{aligned}$$

Siccome stiamo considerando onde piane propagantisi lungo x , non si ha dipendenza da y né da z in $\mathbf{E}$ e $\mathbf{B}$, per cui i rotori di $\mathbf{E}$ e $\mathbf{B}$ non hanno componente x . Ne segue immediatamente che B_x e E_x sono indipendenti anche dal tempo. Ma un campo indipendente da spazio e tempo non è un'onda viaggiante, per cui lo scarteremo. Ne segue l'importante conclusione che *le onde elettromagnetiche piane nel vuoto sono trasversali*, ovvero che i campi $\mathbf{E}$ e $\mathbf{B}$ oscillano perpendicolarmente alla direzione di propagazione.

Nel caso di onde monocromatiche del tipo descritto dall'equazione 4.15, la prima delle equazioni scritte sopra diventa (analogamente si procede per la seconda), per componenti,

$$\begin{aligned}(0, -ikE_{0z}e^{i(kx - \omega t)}, ikE_{0y}e^{i(kx - \omega t)}) &= \left(0, -\frac{\partial E_z}{\partial x}, \frac{\partial E_y}{\partial x}\right) = \nabla \times \mathbf{E} \\ &= -\frac{\partial \mathbf{B}}{\partial t} = (0, i\omega B_{0y}e^{i(kx - \omega t)}, i\omega B_{0z}e^{i(kx - \omega t)})\end{aligned}$$

ovvero, in forma vettoriale, e ricordando che $\mathbf{k} = k\hat{\mathbf{x}}$

$$i\mathbf{k} \times \mathbf{E} = i\omega \mathbf{B} \longrightarrow \mathbf{B} = \frac{\mathbf{k}}{\omega} \times \mathbf{E} = \frac{\hat{\mathbf{k}}}{c} \times \mathbf{E} \quad (4.16)$$

che significa che *in un'onda piana $\mathbf{k}$, $\mathbf{E}$, $\mathbf{B}$ formano una terna destrorsa* ed in particolare che *$\mathbf{k}$, $\mathbf{E}$, $\mathbf{B}$ sono ortogonali tra loro*.

¹Questa connessione tra divergenza nulla e trasversalità giustifica il nome di *gauge trasverso* dato in precedenza al gauge in cui $\nabla \cdot \mathbf{A} = 0$.

Si noti che $B = E/c$ e che perciò B in un certo senso è “piccolo”; infatti se calcoliamo la forza di Lorentz magnetica F_B su di una carica q su di cui incide un’onda elettromagnetica piana otteniamo

$$F_B = q |\mathbf{v} \times \mathbf{B}| = q |\mathbf{v} \times \frac{\hat{\mathbf{k}}}{c} \times \mathbf{E}| \leq q \frac{v}{c} E = \frac{v}{c} F_E. \quad (4.17)$$

Ovvero la forza magnetica è più debole della forza elettrica F_E generata dalla stessa onda sulla stessa particella di un fattore v/c , fattore che solitamente è $\ll 1$.

4.4.2 Polarizzazione delle onde elettromagnetiche

Il fatto che le onde elettromagnetiche piane siano onde vettoriali trasversali implica che esse sono caratterizzate non da una sola direzione — quella di propagazione $\hat{\mathbf{k}}$ — ma anche da una seconda, quella del campo elettrico² $\mathbf{E}$. Questa direzione è detta *direzione di polarizzazione*. Il piano formato da questa direzione e da $\mathbf{k}$ è detto *piano di polarizzazione* ed in generale è una funzione del tempo.

Prendiamo al solito un’onda monocromatica con $\hat{\mathbf{x}} = \hat{\mathbf{k}}$, e scriviamo le componenti del campo elettrico direttamente in forma reale come

$$\begin{aligned} E_x &= 0 \\ E_y &= \mathcal{I}m [E_{0y} e^{i\phi} e^{(kx - \omega t)}] = E_{0y} \sin(kx - \omega t + \phi) \\ E_z &= \mathcal{I}m [E_{0z} e^{i\theta} e^{(kx - \omega t)}] = E_{0z} \sin(kx - \omega t + \theta) \end{aligned} \quad (4.18)$$

dove, per far sì che le componenti di $\mathbf{E}_0$ (E_{0y} ed E_{0z}) siano reali, si sono esplicitate le loro fasi θ e ϕ , che possono essere diverse. Il modulo del vettore $\mathbf{E}$ è dato da

$$E^2 = E_{0y}^2 \sin^2(kx - \omega t + \phi) + E_{0z}^2 \sin^2(kx - \omega t + \theta)$$

e ovviamente in genere non è costante. La sua direzione — ovvero la direzione di polarizzazione — è individuata da un versore con al più due componenti non nulle

$$\hat{\mathbf{n}} = (0, E_y/E, E_z/E). \quad (4.19)$$

In generale questo versore non è fisso, a meno che $\phi = \theta + m\pi$. In questo caso la direzione di polarizzazione è data da $\hat{\mathbf{n}} = \mathbf{E}_0/E_0$ e si parla

²Si tratta di una convenzione; la direzione del campo magnetico andrebbe altrettanto bene.

di *onda polarizzata linearmente*. Tutte le onde elettromagnetiche, in virtù della linearità delle equazioni dell'elettromagnetismo, sono dunque scrivibili come combinazione lineare di onde polarizzate “orizzontalmente” o “verticalmente”, od anche “diagonalmente”:

$$\hat{\mathbf{n}} = \frac{E_{01}}{E_0} \hat{\mathbf{n}}_1 + \frac{E_{02}}{E_0} \hat{\mathbf{n}}_2$$

con $\mathbf{n}_1$ e $\mathbf{n}_2$ versori arbitrari ma fissati.

Se invece le componenti del vettore $\mathbf{E}_0$ hanno fasi diverse, le equazioni 4.18 sono le equazioni parametriche di un'ellisse; in altre parole, l'estremo del vettore $\mathbf{E}$ percorre un'ellisse, anzi un cerchio se $E_{0y} = E_{0z}$. In questi casi si parla di polarizzazione *ellittica* o *circolare*. Se l'estremo del vettore $\mathbf{E}$ percorre la traiettoria in senso antiorario, visto da un osservatore che vede l'onda procedere verso di sé, allora si parla di *polarizzazione destrorsa*. È evidente che il vettore di propagazione $\mathbf{k}$ è dato in questo caso dalla solita regola della mano destra.

Naturalmente si parla anche di polarizzazione *sinistrorsa* nel caso opposto. Anche in questo caso si dimostra (con gli strumenti dell'Appendice B) che tutte le onde sono scrivibili come combinazione lineare di onde polarizzate ellitticamente sia destrorse che sinistrorse.

Come si vedrà, la luce è un'onda elettromagnetica, dunque può essere polarizzata. Di solito la luce che vediamo è una mescolanza casuale di onde di polarizzazione varia; si parla quindi di luce non polarizzata. In considerazione di questo non ci si sorprende che il nostro occhio, a differenza di certi insetti come le api che ne hanno bisogno per l'orientamento in volo, non sia in grado di distinguere le direzioni di polarizzazione. Ci sono però diversi casi, cui si accennerà anche in seguito, in cui la polarizzazione produce fenomeni visibili. Per esempio, la luce riflessa da una superficie è composta in misura preponderante da onde polarizzate parallelamente a detta superficie; è quindi possibile eliminare selettivamente la componente riflessa, vale a dire la più fastidiosa per la vista, usando lenti costituite da un materiale che assorba selettivamente luce di una particolare polarizzazione (*polaroid*) lasciando passare l'altra. Un secondo fenomeno è la cosiddetta *attività ottica*, esibita da alcune sostanze, che consiste nella rotazione del piano di polarizzazione della luce quando li attraversa. Se la rotazione avviene in senso orario si parla di sostanza *levogira*, in caso contrario di sostanza *destrogira*.

4.5 Energia e quantità di moto di un'onda

Finora si è parlato di “onde viaggianti”. È legittimo chiedersi *che cosa viaggia* in un'onda, soprattutto se ci si ricorda che anche in onde aventi un evidente sostrato materiale come le onde del mare, la materia stessa in effetti non viaggia ma si limita ad oscillare più o meno verticalmente, eccettuato che nelle immediate vicinanze della costa. Chi peraltro abbia tentato di nuotare col mare grosso si rende conto che le onde del mare hanno la capacità di spostare anche masse considerevoli. Viene quindi naturale pensare che ciò che è trasportato da un'onda non è altro che *energia e quantità di moto*. Verifichiamo questa ipotesi nel caso delle onde elettromagnetiche piane nel vuoto.

L'equazione per la densità energia del campo elettromagnetico in questo caso ci dà, ricordando che per un'onda piana $cB = E$,

$$u = \epsilon_0 E^2 / 2 + B^2 / 2\mu_0 = \frac{\epsilon_0}{2} E^2 + \frac{1}{2\mu_0 c^2} E^2 = \epsilon_0 E^2. \quad (4.20)$$

L'onda elettromagnetica pertanto trasporta effettivamente energia, suddivisa in parti uguali tra i contributi elettrico e magnetico. In molti casi il campo oscilla ad una frequenza molto più alta delle frequenze di interesse nella materia, per cui si è interessati al valor medio di tale energia per unità di volume. Tale valor medio è, per onde monocromatiche, $\langle u \rangle = \epsilon_0 E_0^2 / 2$. Da qui in poi si sottintenderà che u e U si riferiscono a valori medi.

Il flusso di energia deve essere dato dal flusso del vettore di Poynting $\mathbf{N} = \mathbf{E} \times \mathbf{B} / \mu_0 = \mathbf{E} \times \frac{\hat{\mathbf{k}}}{c} \times \mathbf{E} / \mu_0 = \hat{\mathbf{k}} E^2 / \mu_0 c = \sqrt{\frac{\epsilon_0}{\mu_0}} E^2 \hat{\mathbf{k}}$, e questo ci dice che *in un'onda elettromagnetica l'energia fluisce nella direzione di propagazione dell'onda*, come si poteva intuire. Per un'onda piana monocromatica,

$$\mathbf{N} = \frac{1}{2} \sqrt{\frac{\epsilon_0}{\mu_0}} E_{0x}^2 \hat{\mathbf{k}} = uc \hat{\mathbf{k}}.$$

Il flusso di energia per unità di tempo W lo possiamo calcolare prendendo un parallelepipedo avente un lato di lunghezza L lungo x e sezione trasversale S . Chiamando U l'energia totale ivi contenuta e considerando che il flusso di $\mathbf{N}$ lungo le facce laterali e posteriore è nullo,

$$W = \int_S \mathbf{N} \cdot d\mathbf{S} = ucS = Uc/L.$$

La grandezza “flusso di energia per unità di tempo e di superficie” è detta *intensità* dell'onda

$$I = uc = c \frac{\epsilon_0}{2} E_0^2 \quad (4.21)$$

e la seconda eguaglianza vale per onde monocromatiche. Questo flusso di energia, come si vede, *dipende dal quadrato dell'ampiezza dell'onda*.

Uguualmente si mostra che un'onda elettromagnetica trasporta quantità di moto; prendendo il parallelepipedo usato in precedenza,

$$\mathbf{p} = \frac{1}{c^2} \int \mathbf{N} d\Omega = \frac{U}{c^2} c \hat{\mathbf{k}} = \frac{U}{c} \hat{\mathbf{k}} \quad (4.22)$$

e si vede che $\mathbf{p}$ è diretta, come previsto, lungo la direzione di propagazione dell'onda.

Tutti i ragionamenti qui esposti sono facilmente estendibili ad altri casi, in particolare alle onde sferiche. Anche in questo caso l'energia e la quantità di moto si propagano in direzione perpendicolare ai fronti d'onda. Siccome l'ampiezza di un'onda sferica decresce come $1/r$ ne segue che il flusso di energia per unità di superficie decresce come $1/r^2$. Questo è logico perchè la superficie dei fronti d'onda a sua volta cresce come r^2 ed il flusso di energia attraverso ogni superficie sferica, in assenza di dissipazione, deve essere costante.

4.5.1 La pressione di radiazione

Se un'onda elettromagnetica trasporta quantità di moto, e quest'onda “rimbalza” — ovvero viene riflessa — su di un qualche oggetto (uno “specchio”), la legge di conservazione della quantità di moto ci dice che l'oggetto in questione deve ricevere un impulso $J = \Delta p = \int F dt$ pari alla variazione della quantità di moto dell'onda stessa. Prendiamo in considerazione il caso in cui un'onda elettromagnetica piana viene totalmente riflessa da uno specchio su cui incide perpendicolarmente. Consideriamo un parallelepipedo con area di base S ed altezza $L = c\tau$, con τ un intervallo di tempo arbitrario; esso conterrà una quantità di moto elettromagnetica pari a

$$p = u\Omega/c = uSL/c = \epsilon_0 E_0^2 S\tau/2.$$

L'impulso trasmesso allo specchio per unità di tempo e di superficie, pertanto, nel caso di riflessione totale vale

$$J = \epsilon_0 E_0^2 = 2I/c = 2u$$

e nel caso di assorbimento totale

$$J = \epsilon_0 E_0^2 = 2I/c = u. \quad (4.23)$$

Questa è dunque la forza per unità di superficie, vale a dire una pressione nota come *pressione di radiazione*.

Si tratta di una pressione estremamente debole. Per esempio, l'intensità della radiazione solare sulla Terra vale $I_S = 1.37 \cdot 10^3 \text{ W m}^{-2}$, quindi la pressione di radiazione solare è di soli $9.1 \cdot 10^{-6}$ Pascal, ovvero oltre 10 miliardi di volte più piccola della pressione atmosferica. Tuttavia questa debolissima pressione è sufficiente a far sì che le code delle comete siano sempre rivolte in direzione opposta al Sole³.

4.6 La luce come onda elettromagnetica

Come si è visto nella sezione 3.5.1 i campi elettromagnetici (e, come si è visto altrove, ovviamente anche i potenziali) obbediscono all'equazione delle onde, e le onde relative si muovono, nel vuoto, con velocità $c = 1/\sqrt{\epsilon_0 \mu_0} = 3 \cdot 10^8 \text{ m/s}$. Questa è proprio la velocità della luce nel vuoto; è difficile quindi sottrarsi alla conclusione, sostenuta anche da una vasta serie di altre considerazioni (per esempio, che sia le onde elettromagnetiche che la luce sono onde trasversali), che *la luce è un'onda elettromagnetica*, esattamente come le onde radio, i raggi X ed altre forme di radiazione che vedremo tra poco.

Non si tratta di un risultato da poco: le equazioni di Maxwell sono state ottenute considerando sistemi di fili, magneti, condensatori eccetera, tutte cose che apparentemente con la luce hanno poco a che fare. La connessione tra questi fenomeni dunque è un qualcosa di inaspettato ed è una delle grandi predizioni della teoria elettromagnetica classica.

4.6.1 Quantizzazione del campo elettromagnetico: il fotone

Si è visto che i campi elettromagnetici trasportano energia e quantità di moto: due caratteristiche “meccaniche” che fanno sospettare che sia possibile una qualche trattazione del campo sotto forma particellare. Tuttavia

³Per la precisione, le code cometarie, composte sia da polveri che da gas, sono affette sia dalla pressione di radiazione che dal cosiddetto *vento solare*, ovvero il flusso di particelle che emana dalla nostra stella.

nulla di ciò che è stato visto finora sembra sostenere questa tesi: il campo elettromagnetico è un campo continuo e tutta la fenomenologia fin qui discussa può essere spiegata senza fare ricorso ad altri concetti.

Una serie di risultati sperimentali ottenuti tra la fine del XIX e l'inizio del XX secolo costrinsero però ad un ripensamento di tutto il trattamento dell'elettromagnetismo prima, e di tutta la fisica poi. Questo processo ebbe sbocco nello sviluppo della *meccanica quantistica* e non è qui possibile illustrarlo. Vediamo solo alcune caratteristiche fondamentali.

Consideriamo il caso in cui della radiazione elettromagnetica incide su di una superficie metallica. Siccome l'onda trasporta energia ci si attende che, al momento in cui abbastanza energia sia stata depositata sul metallo, degli elettroni comincino ad avere energia sufficiente per essere emessi dal metallo. Questo si verifica ed è noto come *effetto fotoelettrico*.

Dalla teoria di Maxwell ci si aspetta che sia il numero di elettroni emessi sia la loro energia dipenda dall'energia totale depositata, dunque dall'intensità dell'onda incidente, e che ci sia un intervallo di tempo, che cresce col diminuire dell'intensità stessa, prima che un elettrone abbia accumulato abbastanza energia per potere uscire.

Non è questo in realtà quello che si osserva. Il numero di elettroni emessi è sì proporzionale all'intensità dell'onda, ma l'energia dipende dalla sua *frequenza*. Per onde incidenti monocromatiche in effetti l'energia massima degli elettroni emessi è la stessa per qualunque intensità. Inoltre al di sotto di una frequenza minima (*frequenza di soglia*) non si ha alcuna emissione, anche per intensità elevate di radiazione. Quando però si ha emissione, si osservano elettroni emessi ad intervalli di tempo anche molto brevi dopo l'inizio dell'illuminazione; solo il tempo *medio* di inizio dell'emissione aumenta.

Nel suo lavoro del 1905 che gli fruttò il Nobel, Albert Einstein diede una spiegazione semplice e rivoluzionaria di queste anomalie. Einstein fece l'ipotesi che l'onda elettromagnetica fosse costituita in realtà da uno sciame di "particelle" chiamate *fotoni*, ognuna dotata di quantità di moto $\mathbf{p}$ e di energia data dalla 4.22

$$U = pc$$

(si confronti con la relazione $U = p^2/2m$ valida per particelle materiali libere), secondo quanto prescritto dall'elettrodinamica maxwelliana per le onde piane monocromatiche. A differenza di Maxwell però Einstein assegnò ad ogni onda di frequenza ν un'energia

$$U = h\nu \tag{4.24}$$

dove n è un intero e $h = 6.6256 \cdot 10^{-34}$ J s è la costante che Planck aveva introdotto alcuni anni prima per spiegare altri comportamenti della radiazione⁴. Posto $\hbar = h/2\pi$ si ha di conseguenza $p = \frac{U}{c} = \frac{h}{\lambda} = \hbar k$. L'intensità dell'onda fornisce il numero di fotoni presenti nell'onda stessa, che quindi è vista come formata da una massa di “granellini di luce”. In altri termini, l'energia elettromagnetica è *quantizzata*. Si noti però che non esiste un'energia minima per la radiazione, in quanto è sempre possibile generare onde di frequenza arbitrariamente bassa, dunque fotoni di energia arbitrariamente bassa.

Per esempio, un fascio luminoso giallo ($\lambda = 6 \cdot 10^{-7}$ m, dunque $\nu = 5 \cdot 10^{14}$ s⁻¹) dell'intensità di 100 W e della durata di un secondo risulta composto da $n = U/h\nu = 3 \cdot 10^{20}$ fotoni. È chiaro che con numeri così alti di fotoni la granularità della radiazione è normalmente impercettibile, esattamente come è normalmente impercettibile la granularità della materia⁵. Solo per frequenze relativamente elevate della radiazione il numero di fotoni è sufficientemente piccolo perchè la loro natura discreta sia avvertibile.

Tramite l'ipotesi del fotone tutte le anomalie summenzionate trovano una spiegazione. Un elettrone è colpito da un singolo fotone che gli trasmette la sua energia. Se w è il *lavoro di estrazione* dal metallo, ovvero l'energia necessaria per farlo uscire con velocità nulla, allora l'energia cinetica degli elettroni emessi sarà, per la conservazione dell'energia,

$$U = h\nu - w \quad (4.25)$$

che mostra un'evidente soglia per $\nu = w/h$: fotoni con frequenza inferiore a questa, anche se numerosi, non possono estrarre alcun elettrone. Viceversa, un fotone con ν superiore alla soglia può estrarre un elettrone, anche se l'intensità è bassissima: ne segue che il periodo d'attesa prima che cominci l'emissione può essere molto breve anche per intensità basse.

4.6.2 Lo spettro elettromagnetico

Le onde elettromagnetiche possono essere classificate equivalentemente secondo la lunghezza d'onda, la frequenza o l'energia del singolo fotone. A seconda di questi parametri le onde hanno effetti diversi nell'interazione

⁴Si trattava del problema dell'*emissione di corpo nero*, per la cui trattazione si rimanda ai testi introduttivi di Meccanica Quantistica.

⁵La granularità della materia è dovuta alle molecole. Si noti come la cifra ottenuta non sia troppo dissimile dal numero di Avogadro $N_A = 6 \cdot 10^{23}$. Una mole di fotoni — N_A fotoni — è talvolta chiamata un *Einstein*.

con la materia, sia che si tratti di generazione che di assorbimento. L'insieme di tutte le frequenze (o lunghezze d'onda o energie...) è noto come *spettro*, ed è convenzionalmente diviso in un certo numero di *bande*. Le suddivisioni principali sono come segue:

Raggi X e γ . Prendono questo nome le radiazioni di lunghezza d'onda inferiore ai 10^{-9} m, ovvero di frequenza superiore a $10^{17} \div 10^{18}$ Hz. L'energia dei fotoni in questione supera il keV. La differenza tra i due tipi consiste essenzialmente nel modo di generazione: i raggi X solitamente tramite eccitazione degli elettroni più legati degli atomi, i γ attraverso reazioni nucleari, e per conseguenza questi ultimi sono più energetici dei primi. Un esempio semplice è la generazione di due raggi γ di 511 MeV per mezzo dell'annichilazione di un elettrone e di un positrone (antielettrone).

È noto l'uso dei raggi X per analisi mediche ed ingegneristiche, dovuto al fatto che tessuti e materiali diversi assorbono diversamente tali radiazioni. Le parti più dense risultano quindi più scure — o meglio chiare, visto che ci si ferma al negativo — nelle *radiografie*.

Ultravioletti. Prendono questo nome le radiazioni di lunghezza d'onda compresa tra i $4 \cdot 10^{-7}$ ed i 10^{-9} m, ovvero di frequenza compresa tra i 10^{15} ed i $3.8 \cdot 10^{17}$ Hz circa. L'energia dei fotoni in questione è compresa tra i 3.2 ed i 1000 eV. Si tratta di radiazioni meno penetranti delle precedenti, in larga parte assorbite dall'atmosfera, ma ancora abbastanza energetiche da poter facilmente causare reazioni chimiche e rompere legami chimici. Da questo deriva la loro pericolosità per la salute ed il loro potere sterilizzante.

Visibile. Comprende le radiazioni di lunghezza d'onda compresa tra i $7.6 \cdot 10^{-7}$ ed i $4 \cdot 10^{-7}$ m, ovvero di frequenza compresa tra i 10^{15} ed i $4 \cdot 10^{14}$ Hz. L'energia dei fotoni in questione è compresa tra gli 1.6 ed i 3.2 eV. Il colore degli oggetti dipende dalla lunghezza d'onda: quelle più corte danno la sensazione del violetto, poi via via che si scende si ottengono il blu, il verde, il giallo, l'arancione ed il rosso. Questo fatto giustifica l'aggettivo *monocromatica* (in greco "di un solo colore") applicato alle onde di una frequenza ben determinata, anche se non appartengono al visibile. Si faccia attenzione però che la percezione del colore è una questione ben più complessa di quanto qui accennato.

Infrarossi. Radiazioni di lunghezza d'onda compresa tra i 10^{-3} ed i $7.6 \cdot 10^{-7}$ m, ovvero di frequenza compresa tra i $3 \cdot 10^{11}$ ed i $4 \cdot 10^{14}$

Hz. L'energia dei fotoni in questione è compresa tra gli 0.001 e gli 1.6 eV. Queste radiazioni sono particolarmente efficienti nell'interagire con i gradi di libertà rotazionali e vibrazionali delle molecole, che hanno frequenze proprie — si ricordi la discussione dei circuiti oscillanti e dei corrispondenti oscillatori meccanici! — proprio di questi ordini di grandezza. Inoltre i corpi a temperature non troppo elevate, fino a qualche migliaio di gradi, emettono principalmente radiazione elettromagnetica in questa banda; si noti che a temperatura ambiente $k_B T \simeq 0.025$ eV. Ne segue che gli infrarossi sono i principali responsabili, anche se non gli unici, del trasferimento del calore per *irraggiamento*.

Microonde. Radiazioni di lunghezza d'onda compresa tra i 0.3 ed i 10^{-3} m, ovvero di frequenza compresa tra i 10^9 ed i $3 \cdot 10^{11}$ Hz. L'energia dei fotoni in questione è compresa tra i il meV ed il centesimo di meV. Le frequenze di lavoro dei forni a microonde (2.5 GHz), dei radar (circa 10 GHz) e dei cellulari odierni (intorno ai 2 GHz) cadono in questa banda. È una banda importante anche per l'astronomia, in quanto alcune radiazioni importanti emesse dall'idrogeno cadono proprio qui.

Onde radio. Le onde di maggiore lunghezza d'onda, superiore agli 0.3 m, e di minore frequenza, inferiore al GHz. All'estremo superiore si trovano le frequenze televisive (100 – 1000 MHz), mentre le varie bande radio vanno dai 500 kHz ai 200 MHz. L'atmosfera è generalmente trasparente ad esse, ma ad alte quote si ha una zona, detta *ionosfera*, dove la presenza di atomi ionizzati ed elettroni liberi crea una zona conduttrice in grado di rifletterli. Questo permette le trasmissioni radio anche al di là dell'orizzonte.

4.7 Cenni di ottica geometrica

4.7.1 Le onde elettromagnetiche nei mezzi materiali

L'affermazione secondo cui le onde elettromagnetiche si muovono con velocità c è valida solo nel vuoto. È però elementare vedere che in un mezzo materiale trasparente con costante dielettrica $\epsilon = \epsilon_r \epsilon_0$ e permeabilità magnetica $\mu = \mu_r \mu_0$ un'onda elettromagnetica si propaga con velocità

$$c' = \frac{1}{\sqrt{\epsilon\mu}} = \frac{c}{\sqrt{\epsilon_r\mu_r}} = \frac{c}{n}$$

avendo introdotto la grandezza *indice di rifrazione (assoluto)*

$$n = \sqrt{\epsilon_r \mu_r} \simeq \sqrt{\epsilon_r}. \quad (4.26)$$

La seconda eguaglianza deriva dal fatto che nella grande maggioranza dei mezzi trasparenti $\mu_r \simeq 1$. È necessario tenere presente che, sebbene finora si sia considerata ϵ_r come una costante, essa è in realtà una funzione della frequenza, e pertanto *i mezzi materiali sono mezzi dispersivi*. La dipendenza non è affatto trascurabile; per esempio, per l'acqua $\epsilon_r = 78 \Rightarrow n = 8.8$ a basse frequenze, ma alle frequenze del visibile $\epsilon_r = 1.8 \Rightarrow n = 1.33$.

Il fatto che la velocità delle onde elettromagnetiche vari da un mezzo all'altro ha conseguenze importanti. La principale è che all'interfaccia tra due materiali (uno dei quali può essere un gas, o il vuoto) le onde non si propagano più in maniera rettilinea.

4.7.2 Le leggi dell'ottica geometrica

Per discutere degli effetti all'interfaccia conviene introdurre il concetto di *raggio*. Un raggio è una linea a cui il vettore $\mathbf{k}$ di propagazione dell'onda è sempre tangente, ed orientata in verso concorde ad esso. In pratica, tra $\mathbf{k}$ ed il raggio intercorre la stessa relazione che intercorre tra una linea di campo elettrico o magnetico ed il vettore campo elettrico o magnetico. È immediato vedere che i raggi sono ad ogni punto ortogonali ai fronti d'onda. Per un'onda piana i raggi sono fasci di rette parallele; per una sferica sono fasci di rette passanti per un punto. I punti di due fronti d'onda connessi da un raggio vengono chiamati *punti corrispondenti* ed è chiaro che *l'intervallo di tempo che separa punti corrispondenti di due superficie d'onda è lo stesso per tutte le coppie di punti corrispondenti (Teorema di Malus)*.

Si dimostra che, purché la lunghezza d'onda sia sufficientemente grande (la ragione di questa limitazione la si vedrà tra poco), all'interfaccia tra due mezzi trasparenti di indici di rifrazione n_1 e n_2 , un raggio luminoso (*raggio incidente*) che forma un angolo θ_i con la normale alla superficie si sdoppia: si genera un *raggio riflesso*, formante un angolo θ_R con la detta normale, ed un *raggio rifratto*, che continua nel secondo mezzo formando un angolo θ_r con la normale (orientata verso l'interno) alla superficie.

Le relazioni tra questi raggi, nelle ipotesi fatte, sono riassunte in tre leggi:

1. *I raggi incidente, riflesso e rifratto giacciono tutti sullo stesso piano, che comprende la normale alla superficie.*

2. *L'angolo di incidenza e l'angolo di riflessione sono uguali.*

3. *(Legge di Snell) Tra l'angolo d'incidenza e l'angolo di rifrazione vale la relazione*

$$n_1 \sin \theta_i = n_2 \sin \theta_r \quad (4.27)$$

dove n_i è l'indice di rifrazione del mezzo i -mo.

Il rapporto $n_{12} = n_1/n_2$ tra gli indici di rifrazione è anche indicato come *indice di rifrazione relativo del mezzo 1 rispetto al mezzo 2* ed è chiaramente il reciproco del rapporto delle velocità della luce nei due mezzi. Spesso gli indici di rifrazione vengono dati relativi all'aria, per ovvi motivi. La differenza comunque è piccola, visto che per l'aria $n = 1.0003$.

La giustificazione delle tre leggi può essere data sulla base del principio di Malus. La prima e la seconda in realtà seguono da semplici considerazioni di simmetria. Per dimostrare la prima, immaginiamo per assurdo che un raggio riflesso (per quello rifratto il ragionamento è identico) non giaccia sul piano definito dal raggio incidente e dalla normale, ma che sia p.es. "spostato a destra". Ribaltiamo ora il sistema rispetto al piano "raggio incidente/normale": ora il raggio riflesso dovrebbe essere ugualmente ribaltato ed essere, quindi, "spostato a sinistra". Ma in virtù del ribaltamento il sistema non è mutato: quindi il raggio riflesso dovrebbe essere comunque "spostato a destra". L'unico modo in cui può essere entrambe le cose è che non sia spostato affatto.

Per dimostrare la seconda, si può procedere in maniera analoga, ma sfruttando la simmetria per inversione temporale. Di nuovo per assurdo, supponiamo che l'angolo di riflessione sia p.es. maggiore di quello di incidenza (lo stesso ragionamento varrà per l'ipotesi che sia minore). Questo vuol dire che il raggio "si abbassa" durante la riflessione, cioè che è più radente alla superficie. Ora immaginiamo che il raggio torni sui suoi passi: allora il raggio, retrocedendo, dovrebbe "alzarsi" durante la riflessione. Ma non c'è alcuna ragione di distinguere un caso dall'altro: le leggi di propagazione ondosa sono indifferenti al segno del tempo, ovvero se il tempo scorre in un senso o nell'altro⁶. Quindi il raggio che "retrocede" dovrebbe "abbassarsi" nella riflessione, e l'unico modo in cui può farlo mentre al contempo "si alza" è che non faccia nessuna delle due cose.

Per dimostrare la terza, consideriamo due raggi R_1 e R_2 di un'onda piana incidente, con R_1 che incide prima di R_2 , e prendiamo due punti A (su R_1) e B (su R_2) appartenenti allo stesso fronte d'onda; A giace

⁶Questa è una conseguenza della presenza di una derivata *seconda* rispetto al tempo nell'equazione 4.1.

all'interfaccia tra i due mezzi e B nel primo mezzo. Prendiamo anche due punti C e D , sempre su R_1 e R_2 rispettivamente, ma su di un altro fronte d'onda, e con D all'interfaccia. C si trova all'interno del secondo mezzo.

Il triangolo ACD è rettangolo in C , e l'angolo ADC è uguale a θ_r (il segmento CD è un fronte d'onda, dunque perpendicolare a R_2 nel mezzo 2), perciò $\overline{AD} \sin \theta_r = \overline{AC}$. Il triangolo ABD è rettangolo in B , e l'angolo ADB è uguale a θ_i , perciò $\overline{AD} \sin \theta_i = \overline{BD}$. Usiamo il principio di Malus per affermare che i tempi per andare da A a C e da B a D sono gli stessi. Ne segue che $\overline{AC}/v_2 = \overline{BD}/v_1$, e da questo $v_2 \sin \theta_i = v_1 \sin \theta_r$, ovvero la legge di Snell.

Si noti come la legge di Snell, vista come equazione per l'angolo del raggio rifratto, non ammetta sempre soluzioni. Se l'angolo d'incidenza è abbastanza grande, vale a dire se l'incidenza è radente, in modo tale che $n_{12} \sin \theta_i > 1$ il raggio rifratto non può esistere. Questa condizione è detta di *riflessione totale* ed è la ragione per cui i subacquei in immersione non hanno una visuale completa di quanto accade al di sopra del pelo dell'acqua. Nel caso dell'interfaccia acqua-aria la riflessione totale si ha per angoli superiori a circa 49° .

Poste queste leggi, e noti gli indici di rifrazione dei vari mezzi, il calcolo delle proprietà ottiche di un sistema di lenti, specchi, o altro si riduce ad un calcolo puramente geometrico. Da qui il nome di *ottica geometrica* data a questa branca della fisica.

Fin qui si è sottinteso che i valori n_i fossero costanti. In realtà, come si è già accennato, la velocità di un'onda elettromagnetica in un mezzo dipende dalla frequenza e pertanto la rifrazione di un raggio formato da più frequenze — più colori, nel visibile — produce un fascio di raggi ad angolazioni diverse, ogni angolo con un colore diverso. Questo è il familiare fenomeno di scomposizione della luce da parte dei prismi.

Un altro caso in cui si ha più di un raggio rifratto è quello in cui si ha dipendenza dell'indice di rifrazione dalla polarizzazione dell'onda incidente (fenomeno della *doppia rifrazione*). Questo fenomeno si osserva in determinati cristalli a struttura anisotropa, che vengono detti *birifrangenti*. Se, come normalmente si dà, la luce è una miscelanza di polarizzazioni diverse, allora si hanno due raggi rifratti, detti *ordinario* e *straordinario*, corrispondenti a due direzioni di polarizzazione distinte.

Infine si noti come le tre leggi si riferiscano a superfici *liscie*. Se le superfici sono "ruvide", ovvero formate da un gran numero di superfici più piccole orientate in varie direzioni, bisognerebbe applicare le tre leggi a ciascuna delle piccole superfici in questione. Il risultato è una riflessione ed una rifrazione *diffusa*, il che spiega perchè non tutte le superfici sono degli specchi e come in particolare si possano rendere tali levigandole a

sufficienza. È questa la ragione per la limitazione dell'ottica geometrica a lunghezze d'onda "non troppo piccole" menzionata all'inizio del paragrafo.

Ma che cosa significa "liscia" in questo caso? Significa che le irregolarità della superficie devono essere molto più piccole della lunghezza d'onda della luce incidente; in questo modo l'onda non è sensibile alle irregolarità della superficie.

4.7.3 Il principio di Fermat

Le leggi dell'ottica geometrica possono essere riassunte in un solo, semplice principio variazionale, detto di Fermat:

La luce segue sempre una traiettoria che rende estrema il tempo di percorrenza.

Per esempio, per andare da un punto A ad un punto B, distanti d ed entrambi immersi in un mezzo omogeneo, la luce deve rendere estrema il tempo $t = x/c$, dove x è la distanza percorsa e c in questo caso è costante. Quindi bisogna rendere estrema la distanza percorsa, e questo si ottiene ovviamente tramite un percorso rettilineo.

Una deduzione più complicata si ha nel caso della rifrazione. Diamo per dimostrata la prima legge dell'ottica geometrica e lavoriamo direttamente nel piano definito dai raggi incidente e rifratto, che chiamiamo piano xy . Poniamo $A = (0, h)$, immerso nel mezzo omogeneo 1, con indice di rifrazione pari a n_1 , che occupa il semispazio $y > 0$, e $B = (p, q)$ nel mezzo 2. La luce si propagerà in maniera rettilinea all'interno di ognuno dei due mezzi, ed i raggi si piegheranno solo al momento di attraversare l'interfaccia nel punto $(x, 0)$. Il tempo totale dunque sarà dato da

$$t = t_1 + t_2 = x_1/(c/n_1) + x_2/(c/n_2) = \frac{1}{c} \left[n_1 \sqrt{x^2 + h^2} + n_2 \sqrt{(x-p)^2 + q^2} \right]$$

e la condizione di estremo $\delta t = 0$ diventa

$$0 = n_1 \frac{x}{\sqrt{x^2 + h^2}} + n_2 \frac{x-p}{\sqrt{(x-p)^2 + q^2}} = n_1 \sin \theta_i - n_2 \sin \theta_r$$

(il segno negativo deriva dal fatto che θ_r è positivo se $p > x$) che è proprio la legge di Snell. Analogamente (*Esercizio*) si procede nel caso della riflessione.

Si faccia attenzione a non considerare il principio di Fermat come un principio di *minimo*: il percorso seguito dalla luce può ugualmente bene essere un massimo od un punto di sella. Si consideri per esempio un *diottro sferico*: un cilindro trasparente sormontato da una semisfera, pure

trasparente, entrambi di raggio R e con indice di rifrazione n , nel vuoto. Poniamo l'asse del cilindro coincidente con l'asse delle x positive, col centro della semisfera nell'origine, e chiediamoci quale sia il tragitto seguito dalla luce per andare da un punto esterno al diottero $A = (-a - R, 0)$ ed il punto ad esso interno $B = (a - R, 0)$. Evidentemente si tratterà di un percorso rettilineo passante per l'apice della semisfera. Calcoliamo ora il tempo impiegato dalla luce per andare da A a B attraverso un percorso vicino a quello reale ma passante per un punto $(-R \cos \phi, R \sin \phi)$ sulla semisfera. Si ottiene

$$\begin{aligned} ct &= ct_1 + ct_2 \\ &= \sqrt{(a + R - R \cos \phi)^2 + (R \sin \phi)^2} + n \sqrt{(a - R + R \cos \phi)^2 + (R \sin \phi)^2} \\ &= \sqrt{a^2 + 2R(R + a)(1 - \cos \phi)} + n \sqrt{a^2 + 2R(R - a)(1 - \cos \phi)} \end{aligned}$$

ed effettuando la variazione rispetto al parametro, piccolo, ϕ

$$\begin{aligned} 0 &= c \frac{dt}{d\phi} = \frac{R(R + a) \sin \phi}{\sqrt{a^2 + 2R(R + a)(1 - \cos \phi)}} + n \frac{R(R - a) \sin \phi}{\sqrt{a^2 + 2R(R - a)(1 - \cos \phi)}} \\ &\simeq \frac{R}{a} [R + a + n(R - a)] \phi = R \left[\frac{R}{a} (n + 1) - (n - 1) \right] \phi. \end{aligned}$$

Questa equazione è soddisfatta, come previsto, per $\phi = 0$, e ci dice, per esempio prendendo la derivata seconda, che tale valore è un punto di minimo se $R/a > (n - 1)/(n + 1)$, massimo se $R/a < (n - 1)/(n + 1)$, e flesso se $R/a = (n - 1)/(n + 1)$.

4.8 Interferenza e diffrazione

Il fatto che l'equazione delle onde sia lineare ci dice che la somma di due onde è ancora un'onda. Tuttavia le equazioni 4.21 e 4.20 ci dicono anche che l'energia trasferita dipende dal quadrato dell'ampiezza. Questi fattori, lineari e non lineari, si combinano per dare alcuni dei fenomeni più caratteristici connessi alla propagazione per onde.

4.8.1 Interferenza

Il primo e più fondamentale fenomeno consiste nel fatto che la somma di due onde deve dare in taluni punti dello spazio un valore nullo. Prendiamo infatti in considerazione cosa succede se due sorgenti di onde scalari

monocromatiche $f_i = f_{0i} e^{i(kr_i - \omega t)}$, $i = 1, 2$, $f_{0i} \in \mathbb{R}$, si trovano a distanza a tra di loro e misuriamo l'onda risultante ad una distanza d da entrambe. Il valore trovato sarà

$$f(t) = f_{01} e^{i(kr_1 - \omega t)} + f_{02} e^{i(kr_2 - \omega t)}$$

se r_1 e r_2 sono le rispettive distanze delle sorgenti. Definendo lo *sfasamento*

$$\delta = k(r_1 - r_2) \quad (4.28)$$

si vede che l'onda risultante vale

$$\begin{aligned} f(t) &= (f_{01} e^{i\delta} + f_{02}) e^{i(kr_2 - \omega t)} \\ &= \sqrt{(f_{01} \cos \delta + f_{02})^2 + f_{01}^2 \sin^2 \delta} e^{i(kr_2 - \omega t + \phi)} \\ &= \sqrt{f_{01}^2 + f_{02}^2 + 2f_{01}f_{02} \cos \delta} e^{i(kr_2 - \omega t + \phi)} \end{aligned} \quad (4.29)$$

dove ϕ è un numero reale il cui valore qui non interessa. Interessa invece notare che l'onda risultante ha un'ampiezza dipendente dallo sfasamento δ , quindi dalla posizione in cui l'onda è rilevata. Si vede che tale ampiezza è massima per

$$\cos \delta = 1 \Rightarrow \delta = 2n\pi \quad (4.30)$$

e minima per

$$\cos \delta = -1 \Rightarrow \delta = (2n + 1)\pi. \quad (4.31)$$

La somma di due onde dunque può dare come risultato, in certi punti, un campo eventualmente anche nullo (quando $f_{01} = f_{02}$). Questo è un fenomeno familiare negli auditorium, che devono essere progettati accuratamente in modo che non vi siano punti in cui l'interferenza crei problemi all'ascolto; ma si può osservare anche ascoltando la radio, notando come in casa certi punti siano favorevoli alla ricezione ed altri invece meno.

Si parla di *interferenza costruttiva* quando è soddisfatta la condizione 4.30 e l'ampiezza totale del campo è maggiore della somma dei campi delle singole onde, *distruttiva* quando invece è soddisfatta la condizione 4.31. Il luogo geometrico di questi punti (*superfici nodali*) è, in questo caso, una famiglia di iperbolioidi $r_1 - r_2 = (n + 1)\pi/k$ per l'interferenza distruttiva ed analogamente per quella costruttiva.

Si noti come gli iperboloidi appena definiti siano superfici fisse nel tempo. Questa è una diretta conseguenza del fatto che abbiamo scelto sorgenti monocromatiche con la stessa frequenza, o come si suol dire *coerenti*. Per sorgenti non coerenti non si può dare un risultato così ben definito in quanto le superfici nodali si accavallano in modo più o meno caotico e non sono generalmente osservabili.

Un modo di ottenere due sorgenti sincrone consiste nell'illuminare (prendiamo per fissare le idee il caso di un'onda luminosa) uno schermo opaco in cui sono state ricavate due fenditure, che supporremo molto sottili, distanti a con una radiazione monocromatica. In questo modo ogni fenditura risulta essere una sorgente monocromatica in fase con l'altra. Poniamo ora un secondo schermo, integro, dalla parte "oscura" del primo ad una distanza D , e calcoliamo l'ampiezza dell'onda su di esso, prendendo come parametro la distanza x tra il punto sullo schermo ed il punto direttamente di fronte alla coppia di fenditure. Se $x \ll D$ possiamo porre $r_1 \simeq r_2$ e $f_{01} \simeq f_{02}$, per cui dalla 4.29 otteniamo

$$f_0 = f_{01} \sqrt{2(1 + \cos \delta)} = 2f_{01} \cos(\delta/2).$$

L'angolo θ sotto il quale è visto x guardando dalle fenditure è tale per cui $\sin \theta \simeq \tan \theta = x/D$. Pertanto $r_1 - r_2 = a \sin \theta \simeq ax/D$ e perciò

$$\delta = \frac{2\pi}{\lambda}(r_1 - r_2) = \frac{2\pi ax}{\lambda D}. \quad (4.32)$$

Si vede dunque che sul secondo schermo si osserveranno regioni, dette *frange d'interferenza*, in cui l'illuminazione è maggiore ed altre in cui è pressoché nulla. Per la precisione, l'intensità dell'onda sullo schermo è data da

$$I = I_0 \cos^2 \left(\frac{\pi ax}{\lambda D} \right) \quad (4.33)$$

dove I_0 è una costante opportuna. Misurando la separazione delle frange si può ottenere la lunghezza d'onda della radiazione.

4.8.2 Diffrazione

Un secondo fenomeno prettamente ondulatorio è dato dalla caratteristica delle onde di "girare gli angoli". Per capirlo, esaminiamo un caso molto simile al precedente, considerando però una sola fenditura, pensata questa volta come avente una larghezza b non trascurabile. Per semplicità non faremo la trattazione completa. Vogliamo solo dimostrare che nelle direzioni

tali che $b \sin \theta = n\lambda$, con n intero non nullo, si ha interferenza distruttiva; un argomento simile mostra che per $2b \sin \theta = n\lambda$ si ha interferenza costruttiva.

Ricordiamo a questo proposito la condizione di interferenza distruttiva 4.31, che in termini di lunghezza d'onda diventa

$$r_1 - r_2 = (n + \frac{1}{2})\lambda$$

per cui in questo caso, se prendiamo coppie di punti nella fessura distanti di $b/2$ — per esempio un'estremità ed il punto di mezzo — si ha

$$\frac{b}{2} \sin \theta = r_1 - r_2 = \frac{m}{2} \lambda.$$

Così per m dispari questi raggi interferiscono distruttivamente. Per m pari ripetiamo il ragionamento ma con coppie di punti separate di $b/4$ e otteniamo

$$\frac{b}{4} \sin \theta = \frac{m}{2} \frac{\lambda}{2}.$$

Il procedimento può essere esteso a tutti gli interi. Non è però applicabile, come si vede facilmente, per $\theta = 0$, e a quest'angolo l'interferenza è costruttiva.

Si osservi come questo effetto, detto *diffrazione*, dipenda in maniera cruciale dalla lunghezza d'onda della luce, così come ne dipende l'interferenza. Per $\lambda \rightarrow 0$ tali effetti scompaiono, e la luce procede semplicemente in linea retta formando un'immagine fedele della fenditura sul secondo schermo. Abbiamo così che la condizione per cui l'ottica geometrica sia valida è che la lunghezza d'onda sia molto più piccola delle dimensioni degli oggetti, fenditure eccetera che la luce trova sul suo cammino (ma non troppo piccola, altrimenti abbiamo la diffusione discussa in precedenza). Questa condizione è ben soddisfatta dalla luce visibile in circostanze ordinarie, ma molto meno da onde radio e da onde sonore le cui lunghezze d'onda sono dell'ordine dei metri. Non esiste perciò un' "acustica geometrica".

Capitolo 5

La teoria della Relatività

5.1 Contrasto tra Meccanica Classica ed elettromagnetismo

Nella discussione fin qui vista dell'elettrodinamica, è stata data poca attenzione alla questione dei sistemi di riferimento; di fatto il solo punto in cui essa è stata menzionata è stata nella discussione della legge di Faraday. Non si tratta di un'omissione casuale, perché l'unione di elettromagnetismo e di meccanica classica porta ad alcune profonde difficoltà.

Ricordiamo per prima cosa i due risultati fondamentali della teoria classica dei moti relativi:

R) *Principio di relatività: tutte le leggi fisiche hanno la stessa forma in ogni sistema di riferimento inerziale.*

Questo evidentemente non implica che ogni fenomeno fisico sia descritto nello stesso modo in ogni riferimento inerziale: le coordinate di una particella per esempio saranno diverse in sistemi diversi. La legge di gravitazione universale però conserva la forma $\mathbf{F} = -Gm_1m_2\hat{\mathbf{r}}/r^2$ in ogni riferimento inerziale, visto che dipende solo dalla coordinata relativa $\mathbf{r}$ dei corpi interagenti, coordinata che non varia con la scelta del riferimento.

G) *Trasformazioni di Galileo: la posizione $\mathbf{r}$ di una particella misurata in un sistema di riferimento inerziale O e la posizione $\mathbf{r}'$ della stessa particella in un altro sistema inerziale O' , in moto rispetto al primo con velocità costante $\mathbf{V}$ sono collegate dalla relazione*

$$\mathbf{r}' = \mathbf{r} - \mathbf{V}t \quad (5.1)$$

e di conseguenza le velocità dalla relazione

$$\mathbf{v}' = \mathbf{v} - \mathbf{V}. \quad (5.2)$$

Questo secondo principio è coerente col primo. Infatti la legge fondamentale della meccanica classica $\mathbf{F} = m\mathbf{a}$ dipende solo dall'accelerazione, che evidentemente è invariante¹ per trasformazioni di Galileo, e dalla forza, invariante per il principio di Relatività. Di conseguenza, è sufficiente supporre che le masse siano invarianti per non incorrere in contraddizioni.

Con l'elettrodinamica però ci confrontiamo con un fatto nuovo ed incompatibile coi precedenti:

M) *Le equazioni di Maxwell non sono invarianti per trasformazioni di Galileo.*

Il modo più semplice di vederlo è di considerare che da esse si ottiene un'equazione d'onda che prevede onde la cui velocità è, nel vuoto, sempre $c = 1/\sqrt{\epsilon_0\mu_0}$. Se, per il principio R, ammettiamo che le equazioni di Maxwell valgano in tutti i riferimenti inerziali, dobbiamo concludere che

M₁) *La velocità della luce è la stessa in tutti i riferimenti inerziali.*

E questa conclusione è chiaramente in contraddizione con G.

5.1.1 Tentativi di misurare la velocità rispetto all'etere

Il contrasto tra i principi R, G ed M è insanabile, e ci costringe a concludere che uno dei tre deve essere modificato o rigettato.

Una prima possibilità che si presenta viene suggerita dalla discussione dell'effetto Doppler per il suono. Anche in quel caso ci siamo imbattuti in una formula che dipendeva non dalle velocità relative, ma dalla velocità rispetto ad un sistema di riferimento privilegiato, ovvero quello in cui il mezzo in cui si propaga il suono, l'aria, è a riposo. La velocità del suono effettivamente dipende dal riferimento; se anche le onde elettromagnetiche si propagassero in un mezzo, l'*etere*, di cui rappresentano le perturbazioni, sarebbe naturale supporre che le leggi che le regolano siano valide solo nel sistema di riferimento in cui l'etere è in quiete, e la velocità della luce

¹Qui ed in seguito chiamiamo invariante una grandezza il cui valore non cambia nel passaggio da un sistema inerziale ad un altro. Si rifletta sulla differenza tra "invariante" e "conservato".

sarebbe diversa nei vari sistemi di riferimento inerziali secondo la 5.2. In altre parole, per l'elettrodinamica non varrebbe il principio di relatività.

Siccome l'ipotesi dell'etere era già stata avanzata molto tempo prima per altre ragioni, questa ipotesi fu la prima ad essere presa in considerazione. Michelson e Morley si proposero quindi di misurare la velocità della Terra rispetto a questa ipotetica sostanza. Il problema principale consiste nel fatto che la velocità della Terra è molto piccola rispetto a c : possiamo aspettarci che essa sia dell'ordine della velocità orbitale attorno al sole, $v_t \simeq 30$ km/s. (Se anche per un caso fortuito la velocità rispetto all'etere fosse pressoché nulla in un certo momento dell'anno, basterebbe aspettare sei mesi per avere una velocità pari circa a $2v_t$.) Purtroppo questo significa che $v_t/c \sim 10^{-4}$. Non solo, ma si può dimostrare² che non è possibile rivelare effetti del primo ordine in v/c , per cui la precisione degli esperimenti deve essere dell'ordine di $(v_t/c)^2 \sim 10^{-8}$.

Si tratta di una precisione estremamente elevata, ma Michelson trovò un modo per riuscire ad ottenerla, mediante un dispositivo noto come *interferometro di Michelson*. Il dispositivo consiste in una sorgente luminosa S , il più possibile coerente, il cui fascio luminoso viene indirizzato lungo un braccio rigido sul quale è posto a 45° uno specchio semiargentato. Con questo termine intendiamo uno specchio che lascia passare circa metà della radiazione incidente e riflette il resto. Il fascio luminoso incidendo sullo specchio dunque viene diviso in due: una parte prosegue lungo il braccio 1 per una lunghezza l_1 , l'altra percorre un secondo braccio, ad angolo retto rispetto al primo, di lunghezza l_2 . I due fasci poi vengono riflessi alle estremità di detti bracci e fatti ricombinare su di uno schermo, formando una figura di interferenza. Tale figura è dovuta allo sfasamento dei due raggi luminosi che effettuano percorsi diversi in tempi diversi.

Possiamo calcolare facilmente lo sfasamento tra i due raggi nell'ipotesi che la velocità v rispetto all'etere sia parallela al braccio 1. L'analisi è analoga, ma algebricamente più complicata, nel caso generale. Il raggio 1 per andare e tornare dallo specchio semiriflettente impiega un tempo $t_1 = l_1/(c+v) + l_1/(c-v) = 2cl_1/(c^2 - v^2)$. Il raggio 2 compie un tragitto un po' più complicato: bisogna infatti tener conto che l'apparato durante il viaggio della luce si muove rispetto all'etere e che pertanto il raggio, nel riferimento dell'etere, compie un tragitto triangolare. Se il tempo di andata e ritorno è t_2 , il tragitto percorso vale $s = 2\sqrt{v^2 \left(\frac{t_2}{2}\right)^2 + l_2^2}$ ed il tempo è quindi dato dall'equazione $t_2 = s/c = \frac{2}{c}\sqrt{v^2 \left(\frac{t_2}{2}\right)^2 + l_2^2}$. Risolvendola troviamo $t_2 = 2l_2/\sqrt{c^2 - v^2}$.

²P. es. Ref.[11], nota 7 al capitolo 1.

Lo sfasamento Δ è dato dalla differenza $t_2 - t_1$ tra i tempi di arrivo moltiplicata per la frequenza della luce $\nu = c/\lambda$:

$$\begin{aligned}\Delta &= c(t_2 - t_1)/\lambda = \frac{2c}{\lambda} \left(l_2/\sqrt{c^2 - v^2} - cl_1/(c^2 - v^2) \right) \\ &= \frac{2}{\lambda\sqrt{1 - \beta^2}} \left(l_2 - \frac{l_1}{\sqrt{1 - \beta^2}} \right)\end{aligned}$$

dove si è definito $\beta = v/c$.

Questa espressione non è adatta per una misura sperimentale, in quanto il termine tra parentesi richiederebbe una misura accuratissima di l_1 e l_2 . Infatti, all'ordine più basso in β^3 , tale termine diventa $(l_2 - l_1) - l_1\beta^2/2$, ed il termine da misurare è dell'ordine di $l_1 \times 10^{-8}$. La differenza $l_2 - l_1$ dovrebbe essere misurata quindi con un errore inferiore a questo, il che non è praticamente possibile.

Michelson trovò un modo per aggirare la difficoltà. Ruotando tutto il dispositivo di 90° si ottiene un diverso sfasamento Δ' , e per ottenerlo basta scambiare i ruoli di 1 e 2:

$$\Delta' = \frac{2}{\lambda\sqrt{1 - \beta^2}} \left(\frac{l_2}{\sqrt{1 - \beta^2}} - l_1 \right).$$

Dopo la rotazione si ha dunque uno spostamento di frange pari a

$$n = \Delta' - \Delta = \frac{2}{\lambda\sqrt{1 - \beta^2}} \left[\frac{l_1 + l_2}{\sqrt{1 - \beta^2}} - (l_1 + l_2) \right] \simeq \frac{l_1 + l_2}{\lambda} \beta^2 \quad (5.3)$$

dove l'eguaglianza approssimata si riferisce al termine di ordine più basso per $\beta \rightarrow 0$. Si vede dunque che il termine $l_2 - l_1$ scompare, e dunque un errore relativo ε anche relativamente grande — per esempio 10^{-4} — nelle misure delle lunghezze dei bracci influisce solo all'ordine ε sul risultato finale. L'importante è che le lunghezze l_1 e l_2 non varino durante la rotazione

Con un apparato avente $l_1 \simeq l_2 \simeq 10$ m, in luce gialla ($\lambda \simeq 5 \cdot 10^{-7}$ m) e stimando $\beta \sim 10^{-4}$ si può prevedere uno spostamento di circa 0.4 frange; l'apparato permetteva di apprezzare lo spostamento di 0.01 frange. Ripetizioni moderne permettono risoluzioni molto più elevate.

³Il modo più semplice di vederlo è di utilizzare lo sviluppo binomiale 4.12. Lo si tenga sempre presente: in questo capitolo verrà usato spesso.

5.1.2 Reazioni ai risultati sperimentali

Il trascinamento dell'etere

Il risultato è noto: non venne rilevato alcuno spostamento di frange. Questo implica che la velocità della Terra rispetto all'ipotetico etere è nulla, contrariamente ad ogni ragionevole aspettativa. Una via d'uscita consiste nel notare che ciò che si misura in realtà è la velocità della Terra rispetto all'etere che la circonda. Se l'etere fosse trascinato in qualche modo dai corpi materiali, allora la Terra sarebbe *localmente* in quiete ed il risultato dell'esperimento sarebbe ovvio.

Esistono però osservazioni ed esperimenti che contraddicono questa ipotesi. Qui ne verrà menzionato solo uno, l'osservazione dell'*aberrazione stellare*. Si tratta di un fenomeno osservato già nel XVIII secolo e che consiste in un moto apparente annuale, grossomodo circolare, delle stelle. Questo moto è facilmente spiegabile supponendo che la velocità c della luce che ci arriva dalle stelle si sommi — vettorialmente — con la velocità terrestre v_t ; in questo modo si ha che l'angolo della direzione di arrivo della luce varia come

$$\tan \theta = v_t/c. \quad (5.4)$$

Se invece l'etere fosse trascinato dalla Terra la sua velocità osservata sarebbe indipendente dal moto orbitale e l'aberrazione non si osserverebbe.

L'ipotesi della contrazione

Un altro modo di uscire dalla contraddizione è di postulare che la lunghezza l del braccio dell'interferometro che di volta in volta si trova diretto lungo la direzione del moto rispetto all'etere sia ridotta (*contrazione di Lorentz*) di un fattore $\gamma = 1/\sqrt{1 - \beta^2} > 1$, ovvero

$$l' = \gamma l. \quad (5.5)$$

È immediato rendersi conto che questo escamotage rende conto del risultato nullo di Michelson e Morley *se i due bracci sono uguali*. Quando però si ripete l'esperimento con interferometri a bracci disuguali si ottiene ugualmente un risultato nullo. La contrazione di Lorentz perciò non riesce nell'intento di salvare l'ipotesi dell'etere.

Possibili modifiche dell'elettrodinamica

Se non sembra possibile rinunciare al postulato di relatività R, forse è possibile rinunciare alla validità delle equazioni di Maxwell M. Le possibilità di farlo senza entrare in stridente contrasto con quanto già noto, ed in modo da spiegare il risultato dell'esperimento di Michelson-Morley, ed altri, non sono molte. Il tentativo più serio consiste nell'ipotizzare che la velocità della luce dipenda dal moto della sorgente (teorie *emissive* o *balistiche*), includendo nel concetto di sorgente eventualmente gli specchi.

Ipotizzando che la velocità della luce sia sempre c rispetto alla sorgente originale, si poté spiegare il risultato dei vari esperimenti fin qui visti. Tuttavia non è possibile spiegare il risultato di altri esperimenti in cui la velocità della sorgente viene fatta variare in vari modi, ed in cui di conseguenza si avrebbero diversi valori della velocità della luce con effetti facilmente rilevabili. Anche questi effetti, tuttavia, non sono mai stati trovati.

5.2 I postulati della teoria della Relatività Ristretta

L'unico modo di salvare l'ipotesi dell'etere risultò essere quella di renderlo inosservabile. Lorentz riuscì infine a formulare una teoria che, postulando una serie di proprietà della materia e dell'etere, rendeva quasi impossibile⁴ realizzare esperimenti di elettrodinamica che ne rivelassero l'esistenza. Si disse che la funzione dell'etere si era ridotta all'essere il soggetto del verbo "ondulare".

Ma può una teoria fisica essere basata su di una necessità grammaticale? In effetti, pare strano che dopo essersi dati tanta pena per spiegare le proprietà elastiche della materia in termini di proprietà elettromagnetiche, si debba ritenere necessario spiegare le proprietà elettromagnetiche in termini di proprietà elastiche. . . In assenza di ragioni empiriche cogenti che portino necessariamente ad ipotizzarlo, il concetto di etere, come qualunque altro concetto, diventa nel migliore dei casi superfluo se non sospetto, e comunque vittima del rasoio di Occam *Entia non sunt multiplicanda praeter necessitatem*.

⁴Ma non del tutto. Esistono almeno due esperimenti in grado di discriminare la teoria di Lorentz dalla Relatività Ristretta: il primo è la rilevazione dell'effetto Doppler trasverso 5.27; il secondo è il cosiddetto esperimento di Laub, riguardante il moto di un dielettrico magnetico. Si veda [11], nota 27 al capitolo 7. Naturalmente ne esistono altri che possono falsificare la teoria di Lorentz; per esempio esperimenti che rivelino un comportamento dell'elettrone diverso da quello postulato.

La strada seguita da Einstein nel creare la teoria della Relatività Ristretta (o Speciale)⁵ fu quella di rinunciare del tutto ad un concetto ormai diventato superfluo come quello di etere, e di accettare sia R che M, rifiutando G. Eleviamo quindi a livello di postulato, alla stessa stregua dei principi della dinamica, l'affermazione secondo la quale le equazioni di Maxwell sono valide in tutti i riferimenti inerziali, ovvero che la velocità della luce è la stessa in tutti i riferimenti inerziali (M_1).

Questa proprietà si esprime anche dicendo che la velocità della luce è un *invariante*. La rinuncia alle trasformazioni di Galileo di per sé non sarebbe particolarmente sconvolgente; non sarebbe stata la prima volta che concetti fisici apparentemente assodati venivano rivisti. Ciò che rende sorprendente e dura da accettare la posizione di Einstein è il fatto che le trasformazioni 5.1, 5.2 sembrano essere — sono — una conseguenza immediata delle nostre concezioni di spazio e di tempo. In altre parole, *la teoria della Relatività richiede una revisione delle nostre concezioni dello spazio e del tempo*. Pertanto è necessario procedere ad un'accurata disamina di tali nozioni ed in particolare delle procedure tramite le quali coordinate e tempi sono misurate, cercando di dare meno cose possibili per scontate.

5.2.1 Velocità della luce e simultaneità

Esaminiamo il processo di misura di una lunghezza. Come si sa, si dice che un oggetto ha lunghezza l se, posto di fianco ad esso un regolo graduato, tra un'estremità dell'oggetto e l'altra giacciono l suddivisioni di lunghezza pari ad una lunghezza campione.

Si noti come questo procedimento fornisca un risultato univoco solo se il regolo ed il corpo misurato sono in quiete l'uno rispetto all'altro; o, detto altrimenti, se sono entrambi in quiete nello stesso riferimento (inerziale). Se i due oggetti, regolo e corpo, sono in moto relativo, è fondamentale aggiungere la specificazione che i confronti alle estremità vanno eseguiti *simultaneamente*. Se infatti misurassimo, per esempio, la lunghezza di un treno in moto a 100 km/h facendone coincidere l'estremità posteriore con una tacca del regolo e l'estremità anteriore con un'altra tacca *ma un'ora dopo* troveremmo che il treno è lungo 100 km!

⁵Il nome non è molto appropriato, come Einstein stesso ebbe ad osservare; un nome più adatto sarebbe *teoria dell'invarianza*. I termini "ristretta" o "speciale" si usano per distinguerla dalla Relatività Generale, che si applica anche a sistemi di riferimento non inerziali.

Si vede dunque che le misure di lunghezza dipendono in qualche modo dalle misure di tempo, e più specificamente dalle misure di simultaneità. (Chiamiamo *evento* ogni cosa che accade in un ben determinato punto dello spazio ad un tempo ben determinato. In pratica, un evento è un punto nello spazio-tempo.) Fintanto che tutti gli osservatori concordano nel definire due eventi “simultanei” o “non simultanei” questo non crea problemi. Tuttavia questa ipotesi, per quanto naturale, è tutt’altro che ovvia ed in effetti, quando analizzata criticamente, discutibile.

Immaginiamo che due eventi — per esempio la caduta di due fulmini — avvengano simultaneamente a distanza L in un riferimento O — per esempio un binario ferroviario. Un osservatore solidale con O e situato a metà strada tra i due punti in cui i fulmini sono caduti vedrà i lampi arrivare simultaneamente. Dato che la velocità della luce — per il postulato M_1 — è una costante universale, ed invariante l’osservatore ha tutti i diritti di concludere che i due eventi sono simultanei.

Consideriamo invece la stessa coppia di eventi osservati da un riferimento O' in moto con velocità v rispetto ad O lungo il binario. I due lampi saranno allora osservati in istanti *diversi*: la luce proveniente dal fulmine verso il quale O' si sta dirigendo arriverà *prima* dell’altro, dovendo percorrere un intervallo di lunghezza $L/2 - vt$ a velocità c (di nuovo, per il postulato M_1). L’altro lampo invece deve percorrere un intervallo di lunghezza $L/2 + vt$ alla stessa velocità.

Ne consegue che i due osservatori *non concordano* sulla simultaneità dei due eventi. Ovviamente O' è in grado di dire che O vede simultanei i due eventi, ma per il principio di relatività R nessuno dei due può dire che la propria misura sia “più corretta” dell’altra. Ne consegue che *la simultaneità tra eventi è relativa al riferimento* e che pertanto anche *le misure di lunghezza sono relative al riferimento*.

5.2.2 Misure di lunghezza trasversali

Consideriamo come variano le misure di un oggetto effettuate in un sistema di riferimento in cui esso è in moto. Prendiamo due sbarre identiche A e B e poniamo alle loro estremità degli oggetti — lamette, gessi, o altro — tali che possano lasciare un segno permanente sull’altra sbarra. Allontaniamo le sbarre e poi facciamole avvicinare in modo tale che 1) le sbarre siano parallele, 2) la velocità relativa sia perpendicolare ad entrambe le sbarre e 3) il punto di mezzo delle sbarre coincida quando queste si incrociano.

Nel sistema di riferimento solidale con A è solo B che si muove. Se ci fosse contrazione delle lunghezze trasversali allora la sbarra B sarebbe vista più corta di A e lascerebbe su quest'ultima una coppia di segni, mentre su B non ne rimarrebbero.

Nel sistema di riferimento solidale con B è solo A che si muove. Se ci fosse contrazione delle lunghezze trasversali allora la sbarra A sarebbe vista più corta di B e lascerebbe su quest'ultima una coppia di segni, mentre su A non ne rimarrebbero.

Ma per il principio di relatività A le due descrizioni devono essere equivalenti; e siccome in realtà esse portano ad effetti fisici misurabili (e permanenti) diversi, ne consegue che l'ipotesi di contrazione trasversale non può essere valida. Analogo ragionamento vale per l'ipotesi di allungamento trasversale. Quindi le misure delle dimensioni perpendicolari al moto di un oggetto sono immutate rispetto alle misure delle stesse dimensioni quando l'oggetto è a riposo.

5.2.3 Misure di tempi

Consideriamo ora il caso di un raggio luminoso proiettato da una sorgente S solidale con un sistema di riferimento O' in moto con velocità $\mathbf{v}$ rispetto ad un sistema O , ortogonalmente a $\mathbf{v}$, e fatto riflettere su di uno specchio, pure in quiete rispetto ad O' , posto a distanza d da S fino a tornarvi.

In O' il raggio percorre un semplice percorso di andata-ritorno di lunghezza $2d$; per il postulato M_1 il tempo trascorso dall'emissione all'assorbimento vale $\Delta t' = 2d/c$. In questo riferimento i due fenomeni di emissione e di assorbimento avvengono nello stesso posto; un riferimento di questo tipo è detto *riferimento proprio*. In O invece il tragitto ha la forma dei due lati di un triangolo (si ricordi la discussione dell'esperimento di Michelson-Morley), ognuno dei quali ha lunghezza $\delta = \sqrt{(v \Delta t'/2)^2 + d^2}$. Si noti che abbiamo usato il risultato dell'esperimento mentale precedente per affermare che d è lo stesso. Abbiamo perciò

$$\Delta t = \frac{2d}{c} \sqrt{1 - \left(\frac{v}{c}\right)^2} = \gamma \Delta t'. \quad (5.6)$$

Evidentemente si ha che $\Delta t' < \Delta t$. Ovvero si ottiene che *l'intervallo di tempo tra due eventi, misurato da un osservatore in moto, è più lungo dell'intervallo di tempo proprio*, un fenomeno noto come *dilatazione dei tempi*.

5.2.4 Misure di lunghezza longitudinali

Siamo ora pronti per affrontare il problema della misura di lunghezze in direzione parallela al moto relativo di due sistemi di riferimento. Come per il caso delle misure di simultaneità, il modo migliore per definire le lunghezze di oggetti in moto è di sfruttare l'invarianza della velocità della luce. Possiamo dunque misurare la lunghezza di un oggetto, fermo od in moto rispetto al nostro sistema di riferimento poco importa, ponendo uno specchio ad una delle sue estremità e misurando il tempo che la luce impiega per andare dall'altra estremità dell'oggetto allo specchio e ritorno e moltiplicando per la velocità della luce. Per l'osservatore O' solidale con l'oggetto si ha

$$\Delta t' = \frac{2l'}{c}.$$

La lunghezza l' è chiamata, ovviamente, *lunghezza propria*. Per un osservatore O rispetto al quale l'oggetto si muove a velocità v l'analisi è un po' più complicata. Se x è la posizione dello specchio al momento $t_1 = x/c$ della riflessione,

$$x = l + vt_1 = l + \frac{v}{c}x \Rightarrow x = \frac{l}{1 - \beta}.$$

Riguardo al tragitto di ritorno, se y è la distanza della sorgente al momento t_2 del riassorbimento dalla posizione dello specchio al momento della riflessione,

$$y = l - v(t_2 - t_1) = l - \frac{v}{c}y \Rightarrow y = \frac{l}{1 + \beta}.$$

Il tempo totale impiegato è dunque

$$t = (x + y)/c = 2l/c(1 - \beta^2) \Rightarrow l' = \gamma l > l. \quad (5.7)$$

Ne viene dunque confermato e quantificato quanto detto in precedenza: la lunghezza misurata di un oggetto in moto risulta diversa, e precisamente più corta, della lunghezza propria (quella, ricordiamo, misurata dello stesso oggetto effettuata in un sistema di riferimento in cui l'oggetto è in quiete). Questo fatto è noto, come abbiamo già visto, come *contrazione di Lorentz*.

5.2.5 Sincronizzazione degli orologi

Una procedura che dipende in maniera molto evidente dalla definizione di simultaneità è quella di sincronizzazione degli orologi. Due orologi infatti si dicono sincronizzati se forniscono la stessa indicazione nello stesso momento. Nello stesso spirito dei punti precedenti sincronizziamo due orologi distanti l' ponendo una sorgente di luce in un punto equidistante tra loro ed emettendo un lampo; i due orologi verranno fatti segnare lo stesso tempo quando ricevono il lampo.

Seguendo il ragionamento del paragrafo 5.2.1, sappiamo che questa procedura fornisce una sincronizzazione diversa per orologi in quiete in un sistema inerziale O' e per orologi solidali con un riferimento O rispetto a cui O' è in moto. Infatti i lampi viaggiano per spazi più grandi o più piccoli a seconda che la velocità v di O' sia discorde o concorde a quella del lampo di luce. Ci attendiamo dunque un "errore" di sincronizzazione δ degli orologi in O' se visti da O .

Sia $t' = l'/2c$ il tempo che il lampo di luce impiega per raggiungere gli orologi in O' . Mettiamoci ora nel riferimento O : osserveremo orologi posti ad una distanza $l = l'/\gamma$. Con ragionamenti analoghi a quelli del paragrafo precedente, si trova che il tempo che il lampo di luce impiega per raggiungere gli orologi, osservato da O , sarà diverso: per uno sarà $t_1 = l'/2c\gamma - \beta t_1$, per l'altro $t_2 = l'/2c\gamma + \beta t_2$. L'osservatore in O vede quindi un "errore di sincronizzazione" pari a

$$\Delta t = t_2 - t_1 = \frac{l'}{2c\gamma} \overbrace{\left(\frac{1}{1-\beta} - \frac{1}{1+\beta} \right)}^{2\beta\gamma^2} = \frac{l'\beta\gamma}{c}. \quad (5.8)$$

Applicando all'inverso la dilatazione dei tempi

$$\Delta t' = \Delta t/\gamma = \frac{l'v}{c^2}. \quad (5.9)$$

Secondo O gli orologi a riposo in O' e che, secondo quest'ultimo, sono sincronizzati, sono *sfasati* di un fattore $l'v/c^2$, e precisamente l'orologio che si avvicina alla sorgente luminosa *va avanti* rispetto all'altro.

5.3 Le trasformazioni di Lorentz

I risultati degli esperimenti mentali visti in precedenza possono essere sintetizzati nelle trasformazioni di coordinate che prendono il posto delle 5.1,

le cosiddette *trasformazioni di Lorentz*⁶:

$$\begin{aligned}x' &= \frac{x - vt}{\sqrt{1 - (v/c)^2}} \\y' &= y \\z' &= z \\t' &= \frac{t - (v/c^2)x}{\sqrt{1 - (v/c)^2}}\end{aligned}\tag{5.10}$$

Le trasformazioni inverse si ottengono invertendo il segno di v e scambiando le variabili primite con quelle non primite:

$$\begin{aligned}x &= \frac{x' + vt'}{\sqrt{1 - (v/c)^2}} \\y &= y' \\z &= z' \\t &= \frac{t' + (v/c^2)x'}{\sqrt{1 - (v/c)^2}}\end{aligned}\tag{5.11}$$

La prima equazione esprime il fenomeno della contrazione delle lunghezze. Le due successive seguono dall'invarianza della misura delle lunghezze perpendicolari al moto. La quarta rappresenta la dilatazione dei tempi e l'asincronia degli orologi. Si noti che *non è possibile avere* $v \geq c$, perché in questo caso le trasformazioni diventerebbero prive di senso. Ne segue che *la velocità della luce è una velocità limite*: non esiste nessun sistema di riferimento che, visto da un altro, si muove a velocità c o superiore e, di conseguenza, la velocità di nessun corpo materiale è in grado di eguagliare o superare c secondo nessun osservatore.

Con un po' d'algebra si dimostra facilmente (*Esercizio*) che

$$c^2t^2 - x^2 - y^2 - z^2 = c^2t'^2 - x'^2 - y'^2 - z'^2\tag{5.12}$$

Ovvero *un fronte d'onda sferico che si espande a velocità c in un sistema di riferimento si espande sfericamente alla stessa velocità in tutti i sistemi di riferimento*. Si verifica così che la nostra trattazione è coerente con il postulato M_1 , come deve essere.

⁶Come spesso accade, il nome non dice tutta la verità. Le trasformazioni in oggetto ebbero una gestazione lunga ed il nome di Lorentz vi rimase attaccato perché, come accennato in precedenza, egli fu costretto a scriverle e ad elevarle a postulato della sua teoria. Si noti come nella Relatività siano invece *dedotte* dal semplice postulato M_1 .

5.3.1 Le trasformazioni della velocità

Non è difficile ora ricavare le trasformazioni che devono sostituire le 5.2. Siccome la velocità $\mathbf{u}$ in un dato riferimento è semplicemente la derivata dello spazio rispetto al tempo, le trasformazioni delle velocità quando sono viste da un riferimento ad un altro in moto rispetto al primo con velocità $\mathbf{v} = v_x \hat{\mathbf{x}}$ si ottengono, in perfetta analogia con la meccanica classica, esprimendo le derivate di x , y e z rispetto a t in funzione delle derivate di x' , y' e z' rispetto a t' . Notiamo preliminarmente che

$$\frac{dt'}{dt} = \gamma \left(1 - u_x \frac{v}{c^2} \right) \quad (5.13)$$

e si vede subito che per conservare l'ordinamento temporale degli eventi tra due sistemi di riferimento qualsiasi, e di conseguenza per conservare il carattere invariante della causalità, deve essere sempre $\frac{dt'}{dt} > 0$, quindi $vu_x < c^2 \Rightarrow u_x < c$, coerentemente con quanto detto al paragrafo precedente,

Prendendo allora le derivate nel sistema di riferimento primato ed esprimendole in quello non primato, otteniamo

$$\begin{aligned} u'_x &= \frac{dx'}{dt'} = \frac{dx'}{dt} \frac{dt}{dt'} = \frac{u_x - v}{1 - u_x v/c^2} \\ u'_y &= \frac{u_y \sqrt{1 - (v/c)^2}}{1 - u_x(v/c^2)} \\ u'_z &= \frac{u_z \sqrt{1 - (v/c)^2}}{1 - u_x(v/c^2)} \end{aligned} \quad (5.14)$$

e le relative inverse

$$\begin{aligned} u_x &= \frac{u'_x + v}{1 + u'_x v/c^2} \\ u_y &= \frac{u'_y \sqrt{1 - (v/c)^2}}{1 + u'_x v/c^2} \\ u_z &= \frac{u'_z \sqrt{1 - (v/c)^2}}{1 + u'_x v/c^2}. \end{aligned} \quad (5.15)$$

A differenza delle trasformazioni delle coordinate, le trasformazioni delle velocità non lasciano invariate le componenti trasversali a causa della dilatazione dei tempi che altera le misure in ogni direzione.

Da queste equazioni si vede, di nuovo, che nessuna velocità minore di c in un determinato riferimento viene misurata come maggiore di c in un qualunque altro riferimento, e che (*Esercizio*) la velocità della luce è la stessa in ogni riferimento. La trattazione fin qui fatta è dunque del tutto coerente.

5.4 Spazio di Minkowski e quadrivettori

Le trasformazioni di Lorentz sono esprimibili convenientemente in forma matriciale se si definisce un vettore a quattro dimensioni in uno spazio tetradimensionale:

$$x^1 = x; x^2 = y; x^3 = z; x^0 = ct \quad (5.16)$$

o, compattamente,

$$x^i = (ct, \mathbf{r}) \quad (5.17)$$

dove si mettono in evidenza la componente zeresima, la cosiddetta componente *temporale*, e le altre, cosiddette *spaziali*⁷. Un vettore del genere è detto *quadrivettore* e lo spazio in cui è definito viene detto *spazio di Minkowski*. Le trasformazioni di Lorentz si esprimono come trasformazioni lineari in questo spazio, ovvero matrici:

$$\mathbf{L} = \begin{vmatrix} \gamma & -\gamma\beta & 0 & 0 \\ -\gamma\beta & \gamma & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{vmatrix} \quad (5.18)$$

È immediato verificare (*Esercizio*) che $\mathbf{L}$ è unitaria e che la trasformazione inversa è data dall'inversa di $\mathbf{L}$. Si verifica anche che trasformazioni di Lorentz consecutive sono date dal prodotto delle matrici corrispondenti, per cui *le trasformazioni di Lorentz costituiscono un gruppo (abeliano)*, detto *gruppo di Lorentz*, il cui elemento neutro è la trasformazione con $v = 0$.

Il vantaggio di operare con i quadrivettori — e con la loro generalizzazione, i quadritensori — è lo stesso che si ha operando con i normali vettori

⁷Un'altra convenzione molto seguita consiste nel porre l'indice 4 alla componente temporale. In questo caso nella scrittura 5.17 si inverte l'ordine della parte spaziale e temporale.

tridimensionali: un'equazione espressa tramite di essi, se valida in un riferimento qualunque, è immediatamente valida in *ogni* riferimento. Una relazione tra componenti, invece, cambia da un riferimento all'altro perchè in generale le componenti cambiano. Nel caso della Relatività si parla di *covarianza* dell'equazione, ed in particolare se l'equazione è espressa tramite operazioni covarianti su quadrivettori si parla di *covarianza a vista*.

5.4.1 Gli intervalli quadridimensionali

Segue allo stesso modo della 5.12 che

$$c^2 dt^2 - dx^2 - dy^2 - dz^2 = c^2 dt'^2 - dx'^2 - dy'^2 - dz'^2 \quad (5.19)$$

Dunque nello spazio di Minkowski gli intervalli invarianti non sono più i semplici intervalli spaziali $dr^2 = dx^2 + dy^2 + dz^2$ e temporali dt presi separatamente, ma la loro combinazione

$$ds^2 = c^2 dt^2 - dr^2 = c^2 dt^2 - dx^2 - dy^2 - dz^2. \quad (5.20)$$

La metrica dello spazio di Minkowski perciò è diagonale a coefficienti costanti⁸:

$$g_{\mu\nu} = \begin{cases} 1, & \mu = \nu = 0 \\ -1, & \mu = \nu \neq 0 \\ 0, & \mu \neq \nu \end{cases} \quad (5.21)$$

ma *non è euclidea*: infatti non è nemmeno definita positiva. Viene quindi chiamata *pseudoeuclidea*⁹. Un quadrivettore può non avere modulo positivo, visto che il prodotto quadriscolare tra due quadrivettori $a^i = (a^0, \mathbf{a})$ e $b^i = (b^0, \mathbf{b})$ è, con questa metrica,

$$\sum_{ij} a^i g_{ij} b_j = a^0 b_0 - a^1 b_1 - a^2 b_2 - a^3 b_3 = a^0 b_0 - \mathbf{a} \cdot \mathbf{b} \quad (5.22)$$

e la norma di a^i vale $a^{02} - a^2$, una grandezza che può benissimo essere negativa. Il prodotto scalare nello spazio di Minkowski è dunque degenere. L'esempio più semplice di vettore di modulo eventualmente nullo o

⁸Anche qui si trova spesso usata una convenzione diversa, che consiste nell'usare per la $g_{\mu\nu}$ segni opposti a quelli qui usati.

⁹La metrica può essere formalmente resa euclidea ponendo $x^0 = ict$ invece che $x^0 = ct$, e modificando di conseguenza il formalismo. L'uso di coordinate immaginarie però tende ad oscurare il fatto importante che le caratteristiche dello spazio di Minkowski sono sensibilmente diverse da quello euclideo, ed è perciò ormai scomparso.

negativo è dato proprio dal quadriintervallo ds^2 che può essere positivo, negativo o nullo (e ds può essere reale positivo, nullo o immaginario puro. Naturalmente sia dt^2 che dr^2 rimangono reali; *devono* esserlo, in quanto si tratta di quantità misurabili, mentre il ds non lo è.)

I quadriintervalli quindi si dividono con naturalezza in tre classi.

- $ds^2 > 0$: intervalli *di tipo tempo*. Per ogni coppia di eventi separati da uno di questi intervalli esiste un sistema di riferimento in cui gli eventi accadono nello stesso punto ($dr = 0$). In questo riferimento si ha che $ds/c = dt$, da cui otteniamo che ds/c è l'intervallo di tempo proprio $d\tau$. Ne discende che $d\tau$ è un invariante. Si vede anche subito che $d\tau$ è sempre più piccolo del corrispondente intervallo dt misurato in un altro riferimento.

Non esiste invece alcun riferimento in cui i due eventi accadono simultaneamente ($dt = 0$). Infatti, data l'invarianza del quadriintervallo ds^2 , esso dovrebbe essere positivo anche in questo ipotetico riferimento, ma allora si avrebbe $dr^2 < 0$, che è impossibile.

Se due eventi sono separati da un intervallo di tipo tempo allora è in linea di principio possibile che i due eventi siano connessi dalla traiettoria di un corpo materiale; infatti in quel riferimento la velocità di un corpo che "va da un evento all'altro" è dr/dt , un numero reale.

- $ds^2 = 0$: intervalli *di tipo luce*. Non esiste alcun sistema di riferimento in cui entrambi gli eventi, se distinti, avvengono nello stesso punto, o simultaneamente. È possibile andare da un evento all'altro, ma solo muovendosi alla velocità della luce.
- $ds^2 < 0$: intervalli *di tipo spazio*. Per ogni coppia di eventi separati da uno di questi intervalli esiste un sistema di riferimento in cui gli eventi accadono simultaneamente ($dt = 0$). Non esiste invece alcun riferimento in cui i due eventi accadono nello stesso luogo ($dr = 0$).

Se due eventi sono separati da un intervallo di tipo spazio allora è impossibile che i due eventi siano connessi dalla traiettoria di un corpo materiale o dalla luce; infatti in quel riferimento la velocità di un corpo che "va da un evento all'altro" non è un numero reale. Due eventi connessi da un intervallo di tipo spazio non possono influenzarsi reciprocamente; non esiste un'influenza causale di uno sull'altro.

Data l'invarianza del ds , la classificazione suesposta è *invariante*, dunque *assoluta*: ogni osservatore inerziale concorda con ogni altro sul tipo di intervallo che connette due eventi.

Dato un evento E, l'insieme di tutti i punti ad esso connessi mediante intervalli di tipo tempo è detto *cono di luce*. Questo cono quadridimensionale divide, per ogni evento, lo spazio di Minkowski in parti (oltre al cono stesso):

1. il *passato*, ovvero l'insieme degli eventi che possono avere influenzato causalmente E (non rileva se lo abbiano effettivamente fatto o meno); sono tutti gli eventi contenuti nella “parte posteriore” del cono;
2. il *futuro*, ovvero l'insieme degli eventi che possono essere influenzati causalmente da E (non rileva se lo abbia effettivamente fatto o meno); sono tutti gli eventi contenuti nella “parte anteriore” del cono;
3. *l'altrove*, ovvero l'insieme degli eventi che non possono avere influenzato causalmente né possono esserne stati influenzati; sono tutti gli eventi “esterni” al cono.

È fondamentale ricordare che, da quanto detto, *il cono di luce di uno specifico evento è invariante*. Questo significa che tutti gli osservatori inerziali concordano sul fatto che due eventi possano o meno influenzarsi causalmente o meno, e se sì su quale dei due possa essere la causa dell'altro. Non sorgono dunque problemi nel definire la causalità, problemi se sorgerebbero se fosse possibile affermare, magari da riferimenti diversi, sia “A causa B” che “B causa A”.

5.4.2 Le trasformazioni generali tra quadrivettori

È uso comune, per semplificare le notazioni, adottare la *convenzione di Einstein* per le sommatorie, che consiste nel sottintendere che in presenza di indici ripetuti si somma su di essi:

$$a^i b_i = \sum_i a^i b_i$$

$$a^{ijk} b_{jkl} = \sum_{jk} a^{ijk} b_{jkl}.$$

Questa operazione è anche detta *contrazione* e produce, da una coppia di quadritensori di ordine N , un quadritensore di ordine inferiore ($N - m$ se la contrazione è su m indici). Un quadritensore di ordine 0 è detto, logicamente, *quadriscalare*, ed è un invariante.

Un quadrivettore, passando da un riferimento ad un altro, si trasforma secondo le trasformazioni di Lorentz¹⁰

$$a'_i = (L^i_j) a_j. \quad (5.23)$$

Evidentemente la quadricordinata $x^i = (ct, \mathbf{r})$ è un quadrivettore, così come il *quadrigradiente*

$$\square_i = \left(\frac{1}{c} \frac{\partial}{\partial t}, -\nabla \right). \quad (5.24)$$

Il d'alembertiano $\square^2 = \nabla^2 - \frac{1}{c^2} \frac{\partial^2}{\partial t^2}$, in quanto prodotto quadriscale del quadrivettore quadrigradiente (a meno del segno) con se stesso, è invece un operatore invariante.

Esempio: effetto Doppler

La fase di un'onda è invariante, dunque

$$\mathbf{k} \cdot \mathbf{r} - \omega t = \mathbf{k}' \cdot \mathbf{r}' - \omega' t'$$

e siccome $x^i = (ct, \mathbf{r})$ è un quadrivettore ed il prodotto con $k_x, k_y, k_z, \omega/c$ è un invariante — pertanto è un quadriscale — ne segue che

$$k_i = (\omega/c, \mathbf{k}) \quad (5.25)$$

è un quadrivettore. Applicando le trasformazioni di Lorentz a k_i , ovvero moltiplicando il quadrivettore per $\mathbf{L}$ si ottiene, supponendo come al solito un moto relativo rispetto all'asse x ,

$$\begin{aligned} \omega' &= \gamma(\omega - vk_x) \\ k'_x &= \gamma(k_x - \omega\beta/c) \\ k'_y &= k_y \\ k'_z &= k_z \end{aligned} \quad (5.26)$$

Per un'onda elettromagnetica nel vuoto $\omega = ck$ e siccome $k_x = k \cos \theta$, dove θ è l'angolo tra la direzione di propagazione dell'onda e la velocità relativa di sorgente ed osservatore, e $\omega = 2\pi\nu$, si ha

¹⁰Solitamente viene fatta una distinzione tra quadrivettori *covarianti* e *controvarianti*. Nel caso della Relatività Ristretta, però, la differenza è di fatto inessenziale ed il presente trattamento è tale da rendere tale distinzione non necessaria. Si noti anche che purtroppo il termine "covariante" è entrato nell'uso con due significati distinti, ed è probabilmente troppo tardi per rimediare. Si faccia attenzione a non confondersi.

$$\nu' = \gamma(1 - \beta \cos \theta)\nu \simeq (1 - \beta \cos \theta)\nu \quad (5.27)$$

L'eguaglianza approssimata vale per $v \ll c$. Va notato che l'effetto Doppler relativistico si ha anche nel caso in cui $\theta = \pi/2$. Classicamente non si avrebbe alcuna variazione di frequenza; relativisticamente si ha $\nu' = \gamma\nu$. Questo è detto *effetto Doppler trasverso*, ed è causato dalla dilatazione dei tempi. Infatti i due osservatori “sorgente” e “ricevitore” se in moto l'uno rispetto all'altro vedono i rispettivi orologi andare più piano e pertanto vedono la luce di una frequenza diversa.

Anche la formula relativistica per l'aberrazione della luce si ottiene dalle 5.26.

5.5 Dinamica relativistica

5.5.1 La lagrangiana di una particella

È possibile dedurre tutta la dinamica relativistica, così come quella classica, a partire dal principio di Hamilton

$$\delta \int_A^B \mathcal{L} d\tau = 0 \quad (5.28)$$

dove A e B sono in questo caso eventi, cioè punti nello spaziotempo, e non semplicemente punti nello spazio, mentre il tempo τ è il tempo proprio della particella (in modo da avere un'espressione invariante del principio variazionale). La variazione quindi è effettuata ad eventi estremi fissati.

Per ottenere l'espressione per $\mathcal{L}$ utilizziamo un procedimento generalizzato dal caso non relativistico¹¹. La Lagrangiana è, per ipotesi, funzione solo delle coordinate, delle velocità e del tempo. Nel caso classico τ coincide con il tempo universale t ; inoltre per una particella isolata, dall'ipotesi di omogeneità ed isotropia dello spazio e del tempo segue che essa dipende solo dal modulo della velocità, ovvero, che è lo stesso, dal suo quadrato:

$$\mathcal{L} = \mathcal{L}(v^2).$$

Se valgono le trasformazioni di Galileo la lagrangiana $\mathcal{L}'$ della stessa particella vista in un altro sistema di riferimento deve dare le stesse equazioni del moto, quindi deve differire dalla prima per una derivata totale di

¹¹L.D. Landau e E.M. Lifšits, *Corso di Fisica Teorica voll. 1 e 2*, Editori Riuniti.

una funzione delle coordinate e del tempo. Prendendo una velocità relativa dei due sistemi ε infinitesima, si deve dunque avere, all'ordine lineare

$$\mathcal{L}'(v^2) + \frac{d}{dt}f(\mathbf{r}, \mathbf{v}, t) = \mathcal{L}[(v + \varepsilon)^2] = \mathcal{L}(v^2 + 2v\varepsilon + \varepsilon^2) \simeq \mathcal{L}(v^2) + \frac{d}{dv^2}\mathcal{L}(v^2)2\varepsilon$$

con f opportuna. Perciò $\frac{d}{dv^2}\mathcal{L}(v^2)$ deve essere una costante, ovvero $\mathcal{L}$ deve essere lineare in v^2 , visto che nessuna funzione solo di v^2 può essere una derivata totale del tempo di una funzione di coordinate, tempo e velocità. Ne segue il risultato noto $\mathcal{L} = \frac{1}{2}mv^2$ (si dimostra facilmente che $m > 0$).

Nel caso relativistico lo stesso ragionamento ci porta ad ipotizzare una lagrangiana proporzionale al quadrato del quadrivettore *quadrivelocità* definito nella maniera ovvia¹²:

$$u^i = \frac{dx^i}{d\tau} = \frac{d}{d\tau}(ct, \mathbf{r}) = (\gamma c, \gamma \mathbf{u}) \quad (5.29)$$

(la derivazione rispetto ad un invariante di un quadrivettore evidentemente produce un quadrivettore), quadrato che vale $\gamma^2 c^2 - \gamma^2 v^2 = c^2$. Perciò il principio variazionale diventa

$$\delta \int_A^B \mathcal{L} d\tau = \delta \int_A^B \alpha c^2 d\tau = \delta \int_A^B \alpha c^2 / \gamma dt = 0. \quad (5.30)$$

Il coefficiente costante α , che per definizione è un invariante, rimane determinato prendendo il limite classico $\beta \rightarrow 0$ della lagrangiana vista in uno specifico riferimento, ovvero dell'integrando dell'ultimo integrale, $\mathcal{L}/\gamma$:

$$(\alpha c^2 / \gamma)_{\beta \rightarrow 0} \rightarrow \alpha c^2 (1 - \beta^2 / 2) = \frac{1}{2}mv^2$$

a meno di costanti additive, per cui $\alpha = -m$ e $\mathcal{L} = -mc^2 \sqrt{1 - v^2/c^2}$.

5.5.2 Il quadrimomento. Massa ed energia

Sappiamo che il momento generalizzato $p = \partial \mathcal{L} / \partial q$ coniugato ad una coordinata ciclica q si conserva. Nel caso relativistico di una particella libera abbiamo dunque che si conserva il *quadrimomento*

$$p^i = \partial \mathcal{L} / \partial u^i = m u^i = (m\gamma c, m\gamma \mathbf{u}) = (m\gamma c, \mathbf{p}) \quad (5.31)$$

dove si è definita la *quantità di moto relativistica*, componente spaziale del quadrimomento

¹²Anche qui si incontra talvolta una convenzione diversa: $u^i = dx^i/ds = (\gamma, \mathbf{u}/c)$.

$$\mathbf{p} = m\gamma\mathbf{u}. \quad (5.32)$$

La conservazione del quadrimomento implica quindi la conservazione del momento tridimensionale $\mathbf{p}$, versione relativistica del momento classico $m\mathbf{u}$, cui si riduce per piccole velocità¹³.

Una seconda quantità conservata è la componente p^0 del quadrimomento. È logico sospettare che sia strettamente collegata con l'altra importante grandezza meccanica conservata (per sistemi isolati), cioè l'energia. Per verificarlo, calcoliamone la derivata temporale in uno specifico riferimento:

$$\begin{aligned} c \frac{dp^0}{dt} &= \frac{d}{dt} \gamma m c^2 = m \frac{d}{dt} [\gamma(c^2 - u^2) + \gamma u^2] = m \frac{d}{dt} \left[\frac{c^2}{\gamma} + \gamma u^2 \right] \\ &= -m\gamma\mathbf{u} \cdot \frac{d\mathbf{u}}{dt} + m \frac{d}{dt} \gamma u^2 = \mathbf{u} \cdot \frac{d}{dt} m\gamma\mathbf{u} = \mathbf{u} \cdot \frac{d\mathbf{p}}{dt} = \mathbf{u} \cdot \mathbf{F}. \end{aligned}$$

Vediamo perciò che questa derivata, moltiplicata per c , non è altro che la potenza trasferita da o alla particella da una forza esterna. Perciò concludiamo che

$$E = cp^0 = \gamma m c^2 \quad (5.33)$$

è l'energia del corpo in uno specifico riferimento. Per corpi fermi questa espressione prende la notissima forma

$$E = mc^2. \quad (5.34)$$

Otteniamo quindi il risultato che *la massa di un corpo è una misura del suo contenuto di energia a riposo*, dove l'espressione "a riposo" è relativa al centro di massa. Perciò *ogni mutamento dell'energia di un corpo si riflette in un mutamento della sua massa a riposo*. Questo mutamento è però sensibile solo nel caso in cui si abbiano enormi variazioni di energia, come nel caso delle reazioni nucleari, vista la presenza del termine $c^2 = 9 \cdot 10^{16} \text{ m}^2/\text{s}^2$.

Per esempio, se mettiamo un litro di acqua a bollire sul fuoco, e trascuriamo l'energia e la massa trasferite all'aria per evaporazione ed altri modi,

¹³Spesso si trova definita una *massa relativistica* $m_r = \gamma m_0$, con m_0 (*massa a riposo*) coincidente con la massa, invariante, m qui definita. Per quanto comunissimo, tale concetto è fonte di equivoci ed oscura il reale significato fisico delle equazioni. Inoltre tale grandezza è virtualmente assente nei testi più avanzati e di ricerca. Non ne faremo dunque uso.

la sua massa varierà di $\Delta m = Q/c^2$, dove Q è l'energia trasferita come calore. Nel caso in esame, $Q = C_p \Delta T$, e preso il calore specifico $C_p = 1 \text{ cal/K} = 4.2 \text{ J/K}$ e $\Delta T = 80 \text{ K}$ abbiamo $\frac{\Delta m}{m} = 3.7 \cdot 10^{-15}$, abbondantemente al di sotto del limite di misurabilità. Nel caso della reazione nucleare di fusione di quattro protoni in un nucleo di He l'emissione di energia è pari a 26.73 Mev e la massa cala di $\frac{\Delta m}{m} = 0.007$. Questo è il principale meccanismo di generazione di energia nel Sole, che di conseguenza perde una massa pari a circa 4 milioni di tonnellate *al secondo*. (*Esercizio*: ottenere questo risultato, sapendo che l'intensità della radiazione solare sulla Terra, detta *costante solare*, vale $1.37 \cdot 10^3 \text{ W/m}^2$ e che la distanza media Terra-Sole è di $1.51 \cdot 10^{11} \text{ m}$).

L'energia cinetica relativistica E_k è definita ovviamente come quella frazione di energia che dipende dalla velocità. Essa si riduce, come deve, all'energia cinetica classica per moti lenti:

$$E_k = E - mc^2 = \gamma mc^2 - mc^2 = (\gamma - 1)mc^2 \rightarrow \frac{1}{2}mv^2. \quad (5.35)$$

Con gli stessi ragionamenti del caso classico si trova che *il quadrimento di un sistema isolato* (e non solo di una particella) è *costante*. È dunque costante anche il suo modulo, cosa peraltro ovvia visto che

$$p^i p_i = m^2 \gamma^2 c^2 - m^2 \gamma^2 u^2 = m^2 c^2. \quad (5.36)$$

Moltiplicando per c^2 questa equazione ed inserendovi l'espressione di E 5.33 e di $\mathbf{p} = \gamma m \mathbf{u}$ otteniamo la relazione più generale tra E e p

$$E^2 = p^2 c^2 + m^2 c^4. \quad (5.37)$$

Si noti che mentre la 5.33 vale solo per corpi con massa non nulla, la 5.37 non soffre di questa limitazione. Infatti, ponendovi $m = 0$ otteniamo una relazione perfettamente sensata, $E = pc$, che è anche la relazione 4.22 tra energia e quantità di moto per un fotone. La relatività permette perciò di interpretare il concetto di fotone come quello di una *particella a massa nulla*.

Si noti come in meccanica classica tale concetto sarebbe assurdo; massa nulla significherebbe energia e quantità di moto nulle. In relatività questo non è necessario, purchè si consideri che *una particella a massa nulla deve sempre muoversi a velocità c* . Vale evidentemente anche il viceversa: *le uniche particelle che si possono muovere a velocità c sono quelle di massa nulla*, in particolare i fotoni. Questo giustifica il termine di "tipo luce" assegnato ai quadriintervalli con $ds = 0$: solo la luce (ed altre particelle a massa nulla) possono percorrerli.

L'equazione 5.33 ha una seconda importante conseguenza. Siccome $\gamma_{\beta \rightarrow 1} \rightarrow \infty$, anche l'energia di un corpo tende all'infinito all'avvicinarsi della sua velocità a c . Fanno eccezione le particelle a massa nulla, che però sappiamo costrette a muoversi sempre esattamente a velocità c . Otteniamo così un'interpretazione dinamica dell'insuperabilità della velocità della luce: *nessun corpo può essere accelerato a velocità superiore a c perché questo richiederebbe di fornire un'energia infinita*.

Notiamo infine un'altra sorprendente conseguenza della 5.37. Siccome la relazione tra E , p ed m non è lineare, e siccome sia E che p sono grandezze additive, abbiamo che *la massa non è additiva*: un sistema di particelle può avere una massa che non è la somma delle masse delle particelle costituenti.

Per esempio, prendiamo un sistema di due fotoni in moto lungo la stessa direzione ma in versi opposti. Mettiamoci nel riferimento in cui entrambi i fotoni hanno la stessa frequenza, dunque la stessa energia E e quantità di moto uguali ed opposte. L'energia totale del sistema è $2E$ e la quantità di moto totale è nulla. Pertanto $4E^2 = 0 + m^2 c^4 \Rightarrow m = 2E/c^2 \neq 0$. (*Esercizio*: si dimostri che se i due fotoni si muovessero nello stesso verso avremmo il risultato, più intuitivo, $m = 0$).

L'effetto Compton

Un esempio storicamente importante dell'applicazione delle leggi di conservazione relativistiche, perché contribuì all'accettazione definitiva del concetto di fotone, è lo studio della diffusione di un fotone da parte di un elettrone libero. Si noti che quanto segue è altrettanto applicabile all'emissione di un fotone da parte di un elettrone libero; è sufficiente immaginare il processo invertendo il segno del tempo.

Mettiamoci nel riferimento in cui l'elettrone è in riposo prima dell'urto e scegliamo l'asse x lungo la direzione di arrivo del fotone, ed il piano xy coincidente con quello definito dalle direzioni dell'elettrone e del fotone dopo l'urto. Se ν è la frequenza del fotone ed m la massa dell'elettrone, le equazioni di conservazione della quantità di moto lungo x e y e dell'energia danno

$$\begin{aligned}\frac{h\nu}{c} &= \frac{h\nu'}{c} \cos \theta + p \cos \phi \\ 0 &= \frac{h\nu'}{c} \sin \theta - p \sin \phi \\ h\nu + mc^2 &= h\nu' + \sqrt{p^2 c^2 + m^2 c^4}\end{aligned}$$

dove si è inteso che dopo l'urto il fotone possiede una frequenza ν' e si dirige ad un angolo θ rispetto all'asse x , e l'elettrone possiede una quantità di moto (tridimensionale) di modulo p e si muove ad un angolo ϕ rispetto all'asse x .

Abbiamo a disposizione tre equazioni per le quattro incognite ν' , θ , ϕ e p , potremo perciò esprimere una qualunque delle incognite in funzione di una qualunque delle altre. Solitamente si sceglie di esprimere ν' in funzione di θ . Per farlo, eliminiamo ϕ dalle prime due

$$h\nu = h\nu' \cos \theta + \sqrt{p^2 c^2 - h^2 \nu'^2 \sin^2 \theta} \Rightarrow p^2 c^2 = h^2 (\nu^2 + \nu'^2 - 2\nu\nu' \cos \theta)$$

e poi eliminiamo p dalla terza equazione e semplifichiamo:

$$\begin{aligned} [h(\nu - \nu') + mc^2]^2 &= h^2 (\nu^2 + \nu'^2 - 2\nu\nu' \cos \theta) + m^2 c^4 \\ \Downarrow \\ h\nu\nu' - mc^2(\nu - \nu') &= h\nu\nu' \cos \theta \end{aligned}$$

Il risultato finale è dunque

$$\nu' = \frac{\nu}{1 + \frac{h\nu}{mc^2}(1 - \cos \theta)} \quad (5.38)$$

Si noti che non è possibile avere assorbimento completo, o una emissione, di un fotone da parte di un elettrone libero: ciò corrisponderebbe a $\nu' = 0$, che non è una soluzione delle equazioni scritte¹⁴. Per poter avere tale assorbimento è essenziale che l'elettrone sia in interazione con un altro corpo — per esempio un nucleo atomico — con cui possa essere scambiata energia e quantità di moto.

La relazione di Compton diventa più semplice espressa in termini delle lunghezze d'onda λ e λ' del fotone prima e dopo l'urto:

$$\lambda' = \lambda + \lambda_c(1 - \cos \theta) \quad (5.39)$$

La grandezza $\lambda_c = h/mc = 2.42 \cdot 10^{-12} \text{ m}$ è detta *lunghezza d'onda Compton dell'elettrone*. Le variazioni di lunghezza d'onda del fotone sono quindi dell'ordine di λ_c , e per quanto piccole sono agevolmente misurabili per fotoni X.

¹⁴Si noti (*esercizio*) che questo sarebbe possibile mediante un trattamento non relativistico del problema, ma darebbe un risultato insensato: la velocità dell'elettrone risulterebbe superiore a c .

5.6 Elettrodinamica covariante

5.6.1 Il campo magnetico come effetto relativistico

Torniamo ora all'argomento fondamentale del presente corso: l'elettrodinamica. Per prima cosa consideriamo il problema delle proprietà di trasformazione della carica. Se il valore di una carica dipendesse dal sistema di riferimento, dovremmo osservare uno sbilanciamento di cariche in ogni materiale, anche se il numero di elettroni coincide col numero di protoni. Infatti negli atomi gli elettroni si muovono a velocità molto alte¹⁵, una frazione apprezzabile di c , ed i nuclei molto più basse. Siccome, come si è visto nel capitolo 1, basta un minutissimo sbilanciamento di carica per osservare effetti vistosi, possiamo concludere che *il valore della carica è invariante*. Esperimenti più sofisticati confortano questa conclusione.

Stabilito questo, riprendiamo in considerazione uno dei casi elementari e fondamentali: il filo rettilineo infinito percorso da corrente, visto però da un punto di vista relativistico.

Consideriamo un filo giacente lungo l'asse z percorso da una corrente che, nel sistema O a riposo rispetto agli ioni che formano il filo, ha intensità I ed è formata, come di consueto, da elettroni. Supponiamo che il filo in O sia neutro e sia λ_0^+ la densità di carica per unità di lunghezza dovuta agli ioni in questo sistema (è la densità di carica propria degli ioni). La densità di carica elettronica, sempre in O , vale λ^- ed è uguale ed opposta a λ_0^+ per l'ipotesi di neutralità. Ne segue che $I = \lambda^- v = -\lambda_0^+ v$ dove $\mathbf{v}$ è la velocità degli elettroni in O .

La densità di carica λ^- però *non* è la densità di carica propria λ_0^- degli elettroni; data la contrazione di Lorentz,

$$\lambda^- = \gamma \lambda_0^- = \frac{\lambda_0^-}{\sqrt{1 - \left(\frac{v}{c}\right)^2}}.$$

Questa relazione la si può comprendere pensando che l'osservatore in O vede gli elettroni più vicini di un fattore γ di quanto non li veda un osservatore solidale con gli elettroni stessi. La densità propria di carica degli ioni è dunque più alta della densità propria di carica degli elettroni.

Questo risultato è in apparente contraddizione col principio d'invarianza della carica, ma la contraddizione è semplicemente un artefatto del nostro modello di filo infinito. Se prendessimo seriamente questo modello dovremmo concludere che le cariche positive e negative $\int_{-\infty}^{\infty} \lambda^+ dz$ e

¹⁵Questo ragionamento va rivisto alla luce della Meccanica Quantistica, ma rimane comunque corretto.

$\int_{-\infty}^{\infty} \lambda^- dz$ appartenenti al filo sono entrambe infinite, con tutte le difficoltà che conseguono da un raffronto tra infiniti. Se invece interpretiamo più fisicamente il modello come un filo molto lungo ma finito collegato alle estremità a due serbatoi di cariche positive e negative, possiamo semplicemente attribuire lo sbilanciamento alla differenze nel numero di cariche cedute al (od assorbite dal) filo dai serbatoi, viste nei due sistemi di riferimento. Questa differenza si spiega con la relatività della simultaneità: per poter dire che “i serbatoi hanno ceduto una quantità di carica netta q al filo” dobbiamo essere in grado di dire *quando* l’hanno ceduta. Se due osservatori sono in disaccordo sulla simultaneità, essi forniranno numeri diversi per q e di conseguenza può capitare che uno di essi osservi che il filo sia carico ed un altro no.

Prendiamo ora una carica positiva di prova q che viaggia a distanza r dal filo con velocità $\mathbf{u}$ rispetto ad O . Per cominciare prenderemo $\mathbf{u}$ parallela al filo e concorde a $\mathbf{v}$. Nel riferimento O' solidale alla carica le densità di carica di ioni ed elettroni sono rispettivamente¹⁶

$$\lambda^{+'} = \frac{\lambda_0^+}{\sqrt{1 - \left(\frac{u}{c}\right)^2}}$$

$$\lambda^{-'} = \frac{\lambda_0^-}{\sqrt{1 - \left(\frac{u-v}{c}\right)^2}} = -\lambda_0^+ \frac{\sqrt{1 - \left(\frac{v}{c}\right)^2}}{\sqrt{1 - \left(\frac{u-v}{c}\right)^2}}$$

e la densità di carica totale del filo non è più nulla ma vale, sviluppando all’ordine più basso in v/c e u/c ,

$$\lambda = \lambda^{+'} + \lambda^{-'} \simeq \lambda_0^+ \frac{uv}{c^2} = -I \frac{u}{c^2}. \quad (5.40)$$

Nel sistema O' esiste perciò un campo elettrico di modulo dato dalla 1.6, e la carica di prova q è dunque *attirata* dal filo da una forza di modulo

$$F = qE = qI \frac{u}{2\pi c^2 \epsilon_0 r} = qu \frac{\mu_0 I}{2\pi r}. \quad (5.41)$$

Ma questa è proprio la forza che agisce su q per opera del campo magnetico dato dalla 2.8. La stessa coincidenza si ottiene ripetendo il ragionamento per altre direzioni di $\mathbf{u}$: per esempio, se la carica si muove

¹⁶Per semplicità vengono usate le trasformazioni classiche delle velocità. L’errore commesso è di ordine β^2 all’interno di un termine che è già di per sé moltiplicato per β ; siccome terremo solo il termine di ordine più basso in β , non si commette nessun errore nel risultato finale.

perpendicolarmente al filo le contrazioni di Lorentz sono nulle ed il filo continua ad essere visto neutro, dunque non c'è nessuna forza su q , ed anche questo è in accordo con la 2.8.

La conclusione che se ne può trarre è che *il campo magnetico è un effetto relativistico*, dovuto al diverso modo di vedere le distribuzioni di cariche e correnti da parte dei diversi osservatori inerziali. Non per nulla il rapporto tra le forze elettriche e magnetiche dipende dal rapporto v/c , come si vede per esempio dalla 4.17. Se tali forze sono importanti nel nostro mondo “a bassa velocità” è solo perchè la forza elettrostatica è talmente intensa che i minutissimi sbilanciamenti di densità di carica che ne derivano sono sufficienti a dare effetti avvertibili.

5.6.2 Trasformazioni di cariche e correnti

Poniamo le considerazioni semiquantitative del paragrafo precedente nel formalismo relativistico. Una semplice estensione del ragionamento fatto poco sopra porta a concludere che cariche e correnti siano legate tra loro in una trasformazione di Lorentz. È infatti possibile definire un quadri-vettore che sia estensione del vettore densità di corrente semplicemente moltiplicando la densità di carica propria ρ_0 — che, come ogni grandezza propria, è un invariante — per la quadri-velocità:

$$j_i = \rho_0 u^i = (\gamma \rho_0 c, \gamma \rho_0 \mathbf{u}). \quad (5.42)$$

È facile vedere che un corpo che in un riferimento è neutro ma percorso da corrente, è visto da un altro osservatore come carico, e viceversa un corpo carico non percorso da corrente viene visto da altri osservatori come percorso da corrente. (*Esercizio:* verificare che questo è in accordo con il caso del filo del paragrafo precedente.)

L'equazione di continuità diventa semplicemente la quadridivergenza di j_i , e come tale è covariante a vista:

$$\frac{\partial j_i}{\partial x^i} = \left(\frac{1}{c} \frac{\partial}{\partial t}, -\nabla \right) \cdot (\gamma \rho_0 c, \gamma \rho_0 \mathbf{u}) = \gamma \left(\frac{\partial \rho_0}{\partial t} + \nabla \cdot \mathbf{j} \right) = 0. \quad (5.43)$$

Si vede che la carica totale è un invariante:

$$q' = \int_{\Omega'} j'_0 d\Omega' = \int_{\Omega} \gamma j_0 d\Omega / \gamma = \int_{\Omega} j_0 d\Omega = q \quad (5.44)$$

come da ipotesi.

5.6.3 Il quadripotenziale

Il D'Alembertiano è un operatore invariante per cui le equazioni per i potenziali 3.25 diventano

$$\square^2 A_i = -j_i/\epsilon_0 \quad (5.45)$$

avendo definito

$$A_i = (\phi, c\mathbf{A}). \quad (5.46)$$

Nella 5.45 il primo membro contiene un operatore invariante applicato ad A_i , mentre il secondo membro è un quadrivettore. Ne consegue che *i potenziali elettrodinamici effettivamente formano un quadrivettore definito dalla 5.46.*

La condizione di Lorentz 3.24 diventa quindi

$$\frac{\partial A_i}{\partial x^i} = \left(\frac{1}{c} \frac{\partial}{\partial t}, -\nabla\right) \cdot (\phi, c\mathbf{A}) = \frac{1}{c} \frac{\partial \phi}{\partial t} + c\nabla \cdot \mathbf{A} = 0. \quad (5.47)$$

e quindi è covariante a vista, essendo la quadridivergenza di un quadrivettore.

5.6.4 Trasformazioni dei campi

I campi si ottengono dai potenziali 5.46 tramite le solite equazioni. Per brevità omettiamo il calcolo e forniamo solo il risultato:

$$\begin{aligned} E'_{\parallel} &= E_{\parallel} \\ \mathbf{E}'_{\perp} &= \gamma[\mathbf{E}_{\perp} + (\mathbf{v} \times \mathbf{B}_{\perp})] \\ B'_{\parallel} &= B_{\parallel} \\ \mathbf{B}'_{\perp} &= \gamma[\mathbf{B}_{\perp} - \frac{1}{c^2}(\mathbf{v} \times \mathbf{E}_{\perp})]. \end{aligned} \quad (5.48)$$

dove i pedici $\parallel$ e $\perp$ indicano rispettivamente le componenti parallele e perpendicolari al moto tra i sistemi di riferimento primati e non. Nel limite classico, $\beta \rightarrow 0$, $\gamma \rightarrow 1$, e le 5.48 diventano

$$\begin{aligned} \mathbf{E}' &= \mathbf{E} + \mathbf{v} \times \mathbf{B} \\ \mathbf{B}' &= \mathbf{B} \end{aligned} \quad (5.49)$$

relazioni che conoscevamo già. La Relatività ristretta fornisce dunque una descrizione del tutto unificata dell'induzione elettromagnetica.

Possiamo inoltre facilmente dare una descrizione dei campi di una particella in moto uniforme con velocità $\mathbf{v}$. Per semplicità considereremo trascurabile il ritardo di propagazione delle interazioni (risulta che ciò non influenza il risultato). Se prendiamo un punto lungo la direzione del moto della particella, nel sistema di riferimento O' solidale con la particella il campo elettrico $\mathbf{E}$ è semplicemente il campo coulombiano ed è parallelo a $\mathbf{v}$, per cui nel riferimento O , ricordando la contrazione di Lorentz

$$E_{\parallel} = E'_{\parallel} = \frac{q}{4\pi\epsilon_0 r'^2} = \frac{q}{4\pi\epsilon_0 \gamma^2 r^2}. \quad (5.50)$$

Il campo visto da O lungo la direzione di marcia è dunque *più debole* di quello coulombiano di un fattore γ^2 . Analogamente si procede per il calcolo del campo in un punto situato al di fuori della linea di propagazione nel momento in cui la carica “ci passa accanto” (vale a dire, in cui la posizione della carica sulla sua traiettoria coincide con la proiezione ortogonale del punto campo sulla traiettoria stessa; ricordiamo che stiamo trascurando il ritardo):

$$E_{\perp} = \gamma E'_{\perp} = \gamma \frac{q}{4\pi\epsilon_0 r'^2} = \gamma \frac{q}{4\pi\epsilon_0 r^2}. \quad (5.51)$$

Il campo ortogonale — nel senso descritto prima — è dunque *più grande* di quello coulombiano per un fattore γ . Se ne deduce che le linee di campo, viste dall'osservatore O , risultano addensate perpendicolarmente alla direzione del moto della particella e, al limite $v \rightarrow c$, il campo elettrico è trasversale.

Il campo magnetico è evidentemente trasversale e vale, quando la particella “ci passa accanto”,

$$B_{\perp} = \frac{\gamma v}{c^2} E'_{\perp} = \frac{\beta}{c} E_{\perp}. \quad (5.52)$$

Il campo di una particella in moto veloce somiglia dunque sempre più a quello di un'onda piana — $\mathbf{E}$ e $\mathbf{B}$ perpendicolari tra loro ed alla direzione di propagazione, e moduli nel rapporto $E/B = c$ — all'approssimarsi della sua velocità a c .

Si noti che $B_{\perp}$ non si annulla nel limite classico $\beta \rightarrow 0$ ma vale

$$B_{\perp} \simeq \frac{qv}{4\pi\epsilon_0 c^2 r^2} = \left| \frac{\mu_0 q \mathbf{v} \times \mathbf{r}}{4\pi r^3} \right| \quad (5.53)$$

che coincide con la 2.8 se si interpreta la carica in moto come un “elemento di corrente” $I dl$ (si ricorderà come questa interpretazione fosse scorretta partendo dalla magnetostatica, in quanto la legge 2.8 è stata dedotta per correnti chiuse; qui vediamo come in realtà il risultato in sé sia corretto).

I campi dunque sono relativi all’osservatore. Esistono però alcune grandezze che sono invarianti. La prima è semplicemente il prodotto scalare tra i campi elettrico e magnetico:

$$\mathbf{E}' \cdot \mathbf{B}' = E'_{\parallel} B'_{\parallel} + \mathbf{E}'_{\perp} \cdot \mathbf{B}'_{\perp} = E_{\parallel} B_{\parallel} + \gamma^2 [\mathbf{E}_{\perp} + (\mathbf{v} \times \mathbf{B}_{\perp})] \cdot [\mathbf{B}_{\perp} - \frac{1}{c^2} (\mathbf{v} \times \mathbf{E}_{\perp})]. \quad (5.54)$$

Ricordando che il prodotto vettore è ortogonale ad entrambi i fattori e notando che da questo fatto si ha che i vettori $\mathbf{a} \times \mathbf{b}$ e $\mathbf{a} \times \mathbf{c}$ formano tra di loro lo stesso angolo che i vettori $\mathbf{b}$ e $\mathbf{c}$, si ottiene

$$\mathbf{E}' \cdot \mathbf{B}' = E_{\parallel} B_{\parallel} + \gamma^2 [\mathbf{E}_{\perp} \cdot \mathbf{B}_{\perp} - \frac{v^2}{c^2} \mathbf{B}_{\perp} \cdot \mathbf{E}_{\perp}] = E_{\parallel} B_{\parallel} + \mathbf{E}_{\perp} \cdot \mathbf{B}_{\perp} = \mathbf{E} \cdot \mathbf{B} \quad (5.55)$$

come affermato. Analogamente si vede che (*Esercizio*)

$$E^2 - c^2 B^2 = E'^2 - c^2 B'^2. \quad (5.56)$$

Se ne deduce quindi che la proprietà di un’onda elettromagnetica di essere piana è invariante.

Capitolo 6

Termodinamica

6.1 I sistemi termodinamici

Un sistema termodinamico è un sistema di N particelle, moltissime, tipicamente dell'ordine del *numero di Avogadro* $N_A = 6.022 \cdot 10^{23}$. Essendo ovviamente impossibile trattare nel dettaglio il moto di tutte le particelle, bisogna trovare un modo di comprimere le informazioni in modo che la descrizione del sistema sia ridotta a un insieme di gradi di libertà macroscopici — le *coordinate termodinamiche* — necessario e sufficiente per descrivere il comportamento di esso. La Termodinamica tratta di come si può effettuare questa riduzione, a volte a soli tre numeri. A questo fine, si considerano solo le interazioni interne: di un sasso per esempio non consideriamo l'energia potenziale gravitazionale né l'energia cinetica del centro di massa, che sono di competenza della meccanica. Quando in termodinamica si parla di energia si intende quindi l'energia interna.

Uno stato descritto in questo modo si definisce *stato macroscopico* o *macrostato*, e da quanto detto si capisce che ogni macrostato corrisponde a un numero enorme di *microstati* — tipicamente dell'ordine di e^N — ovvero degli stati determinati da posizioni e quantità di moto di tutte le particelle componenti. Lo scopo della Termodinamica è di stabilire relazioni e leggi a cui obbediscano i sistemi. In particolare è importante trovare relazioni tra le grandezze termodinamiche di un sistema. Queste relazioni sono dette *equazioni di stato* e in particolare una relazione che collega esplicitamente la pressione P con le altre coordinate viene spesso chiamata, per antonomasia, l'equazione di stato.

Le coordinate in questione sono state originariamente definite in maniera empirica e macroscopica, e solo in seguito interpretate alla luce delle leggi microscopiche della meccanica tramite la *Meccanica Statistica*. Qui

le tratteremo dapprima in maniera intuitiva e informale per poi rendere i concetti più precisi e rigorosi con lo svilupparsi della trattazione, restando però sempre a livello macroscopico, salvo per alcuni accenni finali.

6.1.1 Grandezze estensive ed intensive

La dimensione di un sistema si può misurare in vari modi. Un modo è di indicarne la massa totale m , oppure il numero di particelle N . Si possono combinare questi dati definendo il concetto di *mole*, ovvero il numero di particelle diviso per il numero di Avogadro: $n = N/N_A$. Questo infatti, nel caso di un sistema omogeneo composto di molecole, è connesso alla massa perché il numero di moli è pari a tanti grammi quanto è il peso μ_a della molecola espresso in unità di massa atomica¹. Il valore di N_A fu fissato come sopra esattamente per soddisfare questa relazione: $N_A = \mu_a/m$ dove m è la massa molecolare in grammi. Per l'idrogeno p.es. questo rapporto vale $1/1.6726 \cdot 10^{-24} = 5.9787 \cdot 10^{23}$, la differenza col valore di N_A essendo data dal fatto che un nucleo di ^{12}C non ha una massa esattamente 12 volte quella di un protone. In prima approssimazione, comunque, si può considerare μ_a pari al numero di nucleoni nella molecola, per esempio porre $\mu_{\text{H}_2\text{O}} = 18$. La massa m (in grammi) di un corpo omogeneo quindi è connessa al numero di moli dalla relazione $m = n\mu_a$.

Un altro modo di misurare la grandezza di un sistema consiste nel darne il volume Ω . Quest'altro modo in realtà è connesso al precedente, in quanto per sistemi definiti e omogenei la massa risulta proporzionale al volume tramite la densità $\rho = m/\Omega$, il numero tramite la *densità numero* $\rho_N = N/\Omega$, e le moli tramite la *densità molare* $\rho_n = n/\Omega$. Esiste quindi un insieme di grandezze proporzionali alla massa m : N , Ω , ma anche l'energia interna U , la carica elettrica e altre grandezze che vedremo. Queste grandezze si chiamano *estensive*.

Esistono anche grandezze che sono indipendenti dalla dimensione del sistema. Se si mescolano due bicchieri di acqua alla stessa temperatura, la temperatura non raddoppia, ma rimane invariata: T pertanto è detta *intensiva*. È intensiva anche la pressione P , il potenziale elettrostatico, la densità ρ , la concentrazione di una soluzione e in generale ogni rapporto tra grandezze estensive. Ovviamente una grandezza intensiva moltiplicata per una intensiva dà una grandezza estensiva.

¹Una unità di massa atomica corrisponde all'incirca alla massa di un atomo di idrogeno, o di un nucleone. Per l'esattezza corrisponde a 1/12 della massa di un nucleo di ^{12}C .

Si potrebbe pensare che esistano grandezze la cui dipendenza dalle dimensioni del sistema non sia né $O(N)$, né $O(N^0)$, ma per esempio $O(N^{2/3})$, cioè proporzionali alla superficie del sistema ($N^{2/3} \sim \Omega^{2/3} \sim L^2$ se $L \sim \Omega^{1/3}$ ne è una dimensione lineare). Tuttavia, a meno di studiare esplicitamente i fenomeni di superficie, quindi se si studiano i fenomeni relativi al grosso del sistema, bisogna ricordare che stiamo trattando $N \sim N_A \gg 1$, e anche, di conseguenza, Ω molto maggiore delle dimensioni volumiche delle molecole componenti. Stiamo quindi considerando il *limite termodinamico*: $N \rightarrow \infty, \Omega \rightarrow \infty$, con $\rho_N = N/\Omega$ tenuto costante e pari alla densità numero del sistema. In queste circostanze, anche $N \gg N^{2/3}$, e di conseguenza tutte le grandezze che crescono meno rapidamente di N sono asintoticamente trascurabili.

6.1.2 Equilibrio, trasformazioni, cicli. Funzioni di stato

La termodinamica, almeno al livello trattato qui, riguarda principalmente sistemi in equilibrio. In Termodinamica, questo termine indica sistemi le cui coordinate termodinamiche 1) sono ben definite: temperatura, pressione ecc. sono le stesse in ogni punto, e 2) in assenza di perturbazioni esterne rimangono invariate. Questo in particolare implica l'assenza di movimento macroscopico all'interno. Nella Termodinamica, quindi, a dispetto del nome, raramente si menziona il tempo, con qualche notevole e importantissima eccezione che vedremo. Per studiare il comportamento dei sistemi al cambiare dei parametri si fa normalmente — ma non sempre! — l'ipotesi che il sistema passi attraverso una successione di stati di equilibrio, o meglio di quasi equilibrio, in modo che i valori delle coordinate termodinamiche siano sempre ben definiti. Si suppone, in altre parole, che il passaggio da uno stato all'altro sia infinitamente piccolo o, se si preferisce, estremamente lento. Si parla in questi casi di *trasformazioni quasi-statiche*, o anche *trasformazioni reversibili*, visto che passare da uno stato A a uno stato B infinitamente vicino è altrettanto possibile che il passaggio da B ad A.

Una trasformazione può essere tale per cui lo stato iniziale coincida con lo stato finale. In questo caso si parla di *trasformazione ciclica* o *ciclo*. In un ciclo, per definizione, il sistema ritorna alla situazione originaria e pertanto le grandezze che sono univocamente determinate dallo stato del sistema, ad esempio l'energia, tornano ad avere lo stesso valore. Queste funzioni sono dette *funzioni di stato*. Oltre all'energia già menzionata sono funzioni di stato la pressione, la temperatura, il volume, eccetera. Non tutte le grandezze però, come vedremo, sono funzioni di stato.

Un modo estremamente utile di rappresentare stati e trasformazioni consiste nel tracciare grafici bidimensionali aventi due coordinate termodinamiche sugli assi, spesso P e Ω o P e T . Gli stati sono rappresentati da punti in questo grafico e le trasformazioni da linee continue che connettono i punti iniziali e finali. Un ciclo viene ovviamente rappresentato da una curva chiusa. In questi grafici una trasformazione a pressione costante, detta *isobara* risulta rappresentata da una retta perpendicolare all'asse P . Altre trasformazioni notevoli sono quelle che mantengono costante T , dette *isoterme*, e quelle che mantengono costante Ω , dette *isocore*.

6.1.3 Lavoro

Quando un sistema è in interazione con l'esterno può scambiare energia in vari modi. Uno di questi consiste nel modificare uno dei parametri macroscopici, meccanicamente definibili, del sistema: nel caso più semplice, quello di un fluido chimicamente definito, è il volume. Sappiamo che l'energia viene trasferita tramite lo svolgimento di lavoro, $dL = \mathbf{F} \cdot d\mathbf{l}$. Applicando questa definizione a ogni pezzo infinitesimo dS della superficie del sistema, e ricordando che la pressione è per definizione la forza ortogonale per unità di superficie, si ottiene $L = \int_S P d\mathbf{l} dS = \int_\Omega P d\Omega$. Si vede quindi che il lavoro è compiuto *dal* sistema quando il sistema si espande, e in questo caso è positivo, e *sul* sistema quando si contrae, e in questo caso è negativo. In una trasformazione isocora invece non si compie lavoro. In una isobara il lavoro è banalmente $P\Delta\Omega$.

Si vede subito che il lavoro dipende dalla trasformazione, pertanto non può essere una funzione di stato, e in effetti non è una proprietà del sistema, ma del modo in cui l'energia è trasferita tra sistemi. Ne consegue anche che il lavoro compiuto in un ciclo non è nullo (lo è solo se la trasformazione di "andata" è identica a quella di "ritorno"), e in effetti è facile vedere che è dato, nella rappresentazione grafica, dall'area inclusa dal ciclo: positivo se il ciclo è percorso in senso orario e negativo se percorso in senso antiorario. Infatti nel primo caso la parte che compie lavoro positivo sta al di sopra di quella che compie lavoro negativo, e nel secondo è il contrario.

6.1.4 Principio Zero e temperatura

Una delle relazioni più importanti tra sistemi è quella dell' *equilibrio termico* fra essi. Lo si definisce come la condizione in cui due sistemi, messi a contatto tra loro ma senza eseguire lavoro l'uno sull'altro, non scambia-

no energia. Data questa definizione, si può porre un principio, verificato sperimentalmente, detto **Principio Zero della Termodinamica**²:

Se due sistemi sono in equilibrio termico con un terzo, allora sono in equilibrio termico tra loro.

La relazione “essere in equilibrio termico con” è evidentemente riflessiva, simmetrica e anche, dato il Principio Zero, transitiva. Si tratta quindi di una *relazione di equivalenza* e come tale divide i sistemi termodinamici in equilibrio in classi di equivalenza. Tutti i sistemi appartenenti alla stessa classe si dicono avere la stessa temperatura T : la temperatura di un corpo quindi è per definizione il numero che etichetta la specifica classe a cui il corpo appartiene. Ovviamente la temperatura di un corpo in generale varia col tempo e di conseguenza il corpo cambia classe di continuo, ma ad ogni istante, se in equilibrio, appartiene ad una e una sola. Queste classi le ordiniamo in modo tale che qualunque elemento appartenente a una classe con T maggiore, messo a contatto con un qualsiasi elemento appartenente a una classe con T minore, causi un flusso di calore dal primo al secondo. Che questa definizione sia ben posta lo si vedrà al paragrafo 6.3.1.

La temperatura T si misura mettendo in contatto il corpo di cui si vuole misurarla a contatto con un corpo campione e misurando la variazione di una opportuna grandezza fisica, funzione di T , di quest'ultimo, confrontando questa con una curva di calibrazione. Il caso più familiare è quello dei termometri a metallo liquido — una volta mercurio, ora varie leghe — consistenti in un sottile tubo di vetro riempito del metallo; la misura consiste nel misurare la dilatazione termica della colonnina e nel confrontarla con la sua dilatazione tra il caso in cui il tubo sia immerso in una soluzione di ghiaccio in equilibrio con acqua liquida e il caso dell'ebollizione. La dilatazione fra i due casi estremi è posta convenzionalmente pari a 100. La temperatura misurata in questo modo, che poi viene estesa anche all'esterno dell'intervallo 0–100, è la familiare *temperatura Celsius* e si misura in *gradi Celsius*, simbolo $^{\circ}C$.

Le leggi della Termodinamica risultano però molto più semplici se si utilizza la *scala Kelvin*, simbolo K , ottenuta dalla Celsius aggiungendo 273.16 a questa. Le differenze tra temperature Celsius sono ovviamente identiche alle differenze tra Kelvin, ma i rapporti altrettanto ovviamente no. Lo zero della scala Kelvin è noto come *zero assoluto* ed ha un'importanza fondamentale, essendo una temperatura limite alla quale non è possibile

²La numerazione un po' strana deriva dal fatto che l'importanza di questo principio fu riconosciuta solo dopo che gli altri principi erano già stati enunciati e numerati.

arrivare e sotto la quale non è possibile andare, come vedremo. Le temperature Kelvin sono pertanto sempre strettamente positive³. Da qui in poi, salvo esplicita indicazione, le temperature saranno intese in Kelvin.

6.2 Il Primo Principio

6.2.1 La conservazione dell'energia e il calore

Nella Meccanica si mostra che sistemi soggetti a forze conservative conservano una quantità nota come *energia meccanica* definita come la somma delle energie cinetiche più le energie potenziali. Anche quando si considerano casi in cui si ha attrito o viscosità si trova che il moto è in ultima analisi riconducibile a forze conservative. Infine nella parte di elettrodinamica si è visto che questo concetto si può estendere anche ai casi di interazioni elettromagnetiche. Lo si eleva dunque al rango di Principio — in effetti, forse il principio più importante della Fisica — detto *Principio di conservazione dell'energia* asserendo che

L'energia totale di un sistema isolato si conserva.

Se invece abbiamo un sistema non isolato ma in contatto con l'esterno — o, come si usa dire, l'universo⁴ — la conservazione dell'energia deve tenere conto della possibilità di scambi energetici con l'esterno. Alcuni di questi scambi avranno luogo tramite la variazione di parametri meccanici macroscopici, per esempio il volume, altri no. Il primo tipo lo abbiamo già chiamato *lavoro* e indicato con L , il secondo lo chiamiamo *calore* e lo si denota generalmente con la lettera Q . Possiamo anche dire che il calore non è altro che l'energia trasferita tra sistemi attraverso meccanismi diversi dalla mutazione di parametri meccanici macroscopici, il che in pratica significa che si trasferisce a causa di differenze di temperatura. Il calore quindi è una grandezza estensiva con le dimensioni di un'energia e si misura in Joule, anche se persiste ampiamente la vecchia unità di misura detta *caloria*, che vale circa 4.18J. Per convenzione si pone $Q > 0$ se l'energia entra nel sistema e $Q < 0$ se ne esce.

³Esistono eccezioni a questa regola, ma riguardano sistemi abbastanza esotici di cui non ci occuperemo.

⁴In Termodinamica per universo non si intende la stessa cosa che nel comune vocabolario, ovvero stelle, galassie eccetera, ma più modestamente un sistema e tutto l'ambiente che lo circonda, nella misura in cui questo si possa considerare isolato da tutto il resto.

La conservazione dell'energia in questi casi pertanto si esprime tramite il **Primo Principio della Termodinamica**:

$$\Delta U = Q - L \quad (6.1)$$

che in forma differenziale diventa

$$dU = dQ - dL. \quad (6.2)$$

Si noti che il primo membro dell'equazione 6.1 è una funzione di stato, ma che i termini al secondo membro non lo sono, visto che come abbiamo visto il lavoro dipende dalla trasformazione, e di conseguenza ne dipende anche il calore. Si vede che il calore, non essendo una funzione di stato **non** è e non può essere una proprietà di un sistema. Espressioni del tipo "il calore posseduto dal sistema" perciò non hanno senso; il calore è energia in transito tra sistemi, non appartenente a uno specifico, e non ha senso nemmeno l'espressione ΔQ perché non è la variazione di qualcosa. Si deve anche intendere, di conseguenza, che il differenziale dQ non è il differenziale di una funzione, quindi non è un differenziale esatto.

Volendo la 6.1 può essere vista come definizione del calore. Il lavoro invece è sempre definibile indipendentemente, e nel caso frequente in cui lo si ottenga dalla variazione del volume del sistema, la 6.2 prende la forma

$$dU = dQ - Pd\Omega. \quad (6.3)$$

Una trasformazione in cui non vi sia trasferimento di calore è detta *adiabatica*.

6.2.2 Calori specifici e latenti

Aggiungere o togliere calore da un sistema ha effetti fisici importanti. Tra questi il più comune consiste nel farne variare la temperatura. Il rapporto tra calore fornito e variazione di temperatura è noto come *capacità termica* $C = Q/\Delta T$ ed è una grandezza estensiva. Il valore di C dipende da come si trasferisce il calore, e varia con i valori delle coordinate termodinamiche. Risulta quindi conveniente definire C in maniera differenziale e specificare se, per esempio, è stata tenuta costante una delle coordinate. Si trovano perciò definite capacità termiche a *volume costante* $\partial Q/\partial T|_{\Omega}$ e a *pressione costante* $\partial Q/\partial T|_P$.

Se si divide C per una misura della massa si ottiene una grandezza indipendente da essa, e dipendente solo dalle coordinate termodinamiche

(spesso questa dipendenza è piccola e si può trascurare) e soprattutto dalla sostanza considerata, e dalla sua fase termodinamica. Queste grandezze sono dette *calori specifici*, e a seconda dei casi possiamo avere calori specifici *molari a volume costante* $C_\Omega = \frac{1}{n} \partial Q / \partial T|_\Omega$, *molari a pressione costante* $C_P = \frac{1}{n} \partial Q / \partial T|_P$, *di massa a volume costante* $C_\Omega = \frac{1}{m} \partial Q / \partial T|_\Omega$, eccetera. Qui useremo quasi solo calori specifici molari, visto che quelli di massa si ottengono immediatamente notando che $\frac{1}{n} = \frac{1}{m} \frac{m}{n} = \frac{\mu_a}{m}$.

Possiamo vedere subito che in generale $C_P > C_\Omega$, almeno nel caso in cui il lavoro sia dato solo da variazioni di volume. Infatti usando la 6.3 abbiamo $\partial Q / \partial T|_P = (dU + Pd\Omega) / dT|_P = \partial U / \partial T|_\Omega + P \partial \Omega / \partial T|_P > \partial U / \partial T|_\Omega$, se consideriamo che il volume di un sistema cresce sempre con l'aumentare della temperatura. Questa disuguaglianza si capisce facilmente se consideriamo che l'incremento di volume richiede energia, e che pertanto mantenere il volume costante significa avere più energia disponibile per fare aumentare la temperatura: ma avere più aumento di temperatura a parità di calore fornito significa avere un calore specifico più basso. Si vede subito che la differenza tra i calori specifici a pressione e volume costante è tanto più importante quanto maggiore il *coefficiente di espansione termica* $\alpha = 1/\Omega \partial \Omega / \partial T|_P$, che è molto maggiore nei gas che nei liquidi o nei solidi. Di fatto la differenza è significativa principalmente nei primi, per cui è comune parlare di calore specifico di liquidi o solidi senza specificare se a pressione o a temperatura costante. Per l'acqua diciamo perciò che, alla temperatura di $15^\circ C$, $C_\Omega \simeq C_P \simeq 4.18 \cdot 10^3 JK^{-1} Kg^{-1}$, ovvero una caloria al grammo.

Un secondo modo in cui la fornitura o sottrazione di calore influisce su di un corpo è di farne variare lo *stato di aggregazione*, o *fase termodinamica*. Riscaldare dell'acqua che si trova a $100^\circ C$ e pressione atmosferica porta all'ebollizione e quindi al passaggio da liquido a vapore. Questo passaggio avviene a P e T costanti, quindi la fornitura di calore non può essere descritta tramite il calore specifico: in un certo senso il calore si "nasconde" (in realtà viene usato per separare le molecole d'acqua tra loro, evento che richiede energia), e questo viene descritto tramite il concetto di *calore latente* Λ , anche qui meglio definito per unità di massa o per mole: $\lambda = \Lambda / m$ o $\lambda = \Lambda / n$. È evidente che il calore latente di una trasformazione è uguale e opposto a quello della trasformazione inversa — in questo esempio, la condensazione del vapore — altrimenti potremmo guadagnare energia indefinitamente eseguendo opportunamente più volte una trasformazione e la sua inversa.

6.2.3 Leggi e trasformazioni dei gas perfetti

Definiamo *gas perfetto* un gas in cui le interazioni tra molecole sono trascurabili, o equivalentemente uno in cui l'energia potenziale delle molecole è trascurabile rispetto a quella cinetica. Questo tipo di gas ha la particolarità di essere descritto da equazioni di stato generali e semplici. Detto a parole, la prima equazione dice che a temperatura costante pressione e volume sono inversamente proporzionali, mentre pressione e temperatura sono direttamente proporzionali se il volume costante e similmente per volume e temperatura. La seconda invece dice che l'energia dipende solo dalla temperatura, non da pressione o volume, e ne è anzi direttamente proporzionale. In formule:

$$P\Omega = nRT \quad (6.4)$$

$$\Delta U = nC_\Omega \Delta T \quad (6.5)$$

valide per tutti i gas perfetti eccetto che per i diversi valori che assume la costante C_Ω , che è facile vedere essere effettivamente il calore specifico a volume costante; in particolare si osserva che per i gas monoatomici $C_\Omega = \frac{3}{2}R$ e per i biatomici $C_\Omega = \frac{5}{2}R$. La costante R è detta *costante dei gas* e vale circa $8.3 \text{ J K}^{-1} \text{ mol}^{-1}$. Da queste, notando che U dipende solo da T e quindi $\frac{1}{n} \frac{\partial U}{\partial T} \Big|_P = \frac{1}{n} \frac{\partial U}{\partial T} \Big|_\Omega = C_\Omega$ si ricava anche una relazione esplicita tra i calori specifici:

$$C_P = \frac{1}{n} \frac{\partial Q}{\partial T} \Big|_P = \frac{1}{n} \frac{\partial U}{\partial T} \Big|_P + \frac{1}{n} P \frac{\partial \Omega}{\partial T} \Big|_P = C_\Omega + \frac{1}{n} \frac{\partial P\Omega}{\partial T} \Big|_P = C_\Omega + R$$

che soddisfa la relazione generale, già vista, $C_P > C_\Omega$. Strada facendo abbiamo anche calcolato il coefficiente di espansione termica $\alpha = 1/T$. Segue che per i gas monoatomici $C_\Omega = \frac{5}{2}R$ e per i biatomici $C_\Omega = \frac{7}{2}R$.

Osserviamo anche come lo zero delle temperature (Kelvin) sia effettivamente la temperatura minima, visto che in quel caso il volume del gas si ridurrebbe a zero.

In un diagramma $P - \Omega$ le trasformazioni isoterme sono delle iperboli. Il calcolo del lavoro tra uno stato iniziale i e uno finale f è semplice:

$$L = \int_i^f P d\Omega = \int_i^f nRT/\Omega d\Omega = nRT \int_i^f d\Omega/\Omega = nRT \log(\Omega_f/\Omega_i).$$

Per un gas perfetto è semplice anche l'equazione della curva rappresentante una trasformazione adiabatica. Prendendo la 6.3 con $dQ = 0$ e ricordando le 6.4 e 6.5 si ottiene

$$nC_{\Omega}dT = -Pd\Omega = -nRTd\Omega/\Omega$$

e separando le variabili e integrando, indi applicando la 6.4 si trovano le relazioni fra loro equivalenti

$$\begin{aligned} T\Omega^{\gamma-1} &= \text{costante} \\ P\Omega^{\gamma} &= \text{costante} \\ P^{(1-\gamma)/\gamma}T &= \text{costante} \end{aligned} \quad (6.6)$$

dove si è posto $\gamma = C_P/C_{\Omega} = 1 + R/C_{\Omega}$ e le costanti a secondo membro sono diverse fra loro. Dalla seconda di queste relazioni si vede che in un diagramma $P - \Omega$ le adiabatiche dei gas perfetti sono simili alle isoterme, ma con pendenza più accentuata visto che $\gamma > 1$.

Una applicazione di queste equazioni sta nella spiegazione del perchè l'aria ad alta quota è più fredda che a livello del mare. Infatti, data la scarsa capacità dell'aria nel trasmettere calore, è possibile trattare i moti verticali di aria atmosferica, in prima approssimazione, come adiabatici. Consideriamo quindi uno straterello di spessore dh in una colonna d'aria che supponiamo secca, così da trascurare l'energia rilasciata dalla condensazione del vapore acqueo. La differenza di pressioni tra le basi sarà data dal peso dell'aria, per unità di superficie, nello straterello, ovvero $dP = -mgdh/\Omega = -mgPdh/(nRT)$. Differenziando la terza delle 6.6 si ha $dT/T = \frac{\gamma-1}{\gamma}dP/P$, da cui finalmente otteniamo la derivata

$$dT/dh = -\frac{\gamma-1}{\gamma}gm/nR = -\frac{\gamma-1}{\gamma}g\mu/R$$

che è negativa, ovvero l'aria si raffredda salendo. Per l'aria, il peso molecolare medio è $m/n = \mu = 28,9$ e i gas componenti, ossigeno e azoto, sono principalmente biatomici, per cui $\gamma = 7/5$. Quindi dT/dh è circa 9.7 K/km. Il valore reale, a causa delle approssimazioni fatte, è inferiore: circa 7 K/km.

6.2.4 Gas reali

A differenza dei gas perfetti, nei sistemi reali le particelle interagiscono, e le interazioni sono spesso intense. Se ci limitiamo al caso dei gas, che sono sostanze diluite (*Esercizio*: stimare la distanza media tra molecole

d'aria a condizioni ambiente, partendo dalla 6.4, e sapendo che la pressione atmosferica ordinaria vale circa $1.013 \cdot 10^5 \text{ N/m}^2$ — 1 N/m^2 è l'unità di misura della pressione nel Sistema Internazionale e si chiama Pascal, simbolo Pa) in cui pertanto le interazioni non sono particolarmente forti — ovvero in cui le energie potenziali sono meno importanti di quelle cinetiche — possiamo semplificare la trattazione di sistemi di particelle interagenti partendo dall'equazione dei gas perfetti e modificandola considerando le principali caratteristiche qualitative delle interazioni. Queste caratteristiche, ignorandone i dettagli, sono:

- a breve raggio, dell'ordine di un Angstrom (10^{-10} m), le molecole si respingono;
- a lungo raggio, le molecole si attraggono con una forza che, per molecole neutre, varia⁵ in modo inversamente proporzionale a r^{-6} .

La prima considerazione porta a ipotizzare che il volume totale a disposizione delle molecole sia in realtà minore del volume del sistema, e che quindi Ω debba essere ridotto di una quantità proporzionale al numero di molecole. La seconda, visto che il volume per particella è proporzionale al cubo della distanza media tra particelle $r^3 \sim \Omega/N$, porta a ipotizzare un contributo alla pressione inversamente proporzionale al quadrato del volume molare. Questa seconda si aggiungerà a P in modo che a parità di T e Ω la pressione del gas non perfetto sia minore di quella di un corrispondente gas perfetto. Introducendo queste correzioni nell'equazione di stato dei gas perfetti otteniamo l'*equazione di stato di van der Waals*:

$$(P + an^2/\Omega^2)(\Omega - nb) = nRT \quad (6.7)$$

che descrive abbastanza bene il comportamento dei gas reali, anche in prossimità della condensazione. A differenza dei gas perfetti i gas di van der Waals hanno isoterme che, nel grafico $P - \Omega$, non sono iperboli. L'equazione $P = P(\Omega) = \text{costante}$ è un'equazione di terzo grado e le curve relative hanno, a seconda dei parametri a e b , o due punti di stazionarietà (un massimo e un minimo), oppure nessuno (come i gas perfetti), oppure hanno un flesso, che corrisponde, eseguendo i calcoli, a una *temperatura critica* $T_c = \frac{8}{27} \frac{a}{Rb}$ e *pressione critica* $P_c = \frac{a}{27b^2}$.

⁵Questa dipendenza ha a che fare con l'interazione tra dipoli elettrici. Una spiegazione completa richiede la Meccanica Quantistica.

6.2.5 Termodinamica della radiazione

Si è visto che la radiazione elettromagnetica esercita una pressione, data, se il fascio di radiazione ha una direzione ben definita, dalla 4.23. Si vede inoltre che i corpi non solo assorbono radiazione, ma se caldi (ovvero se a temperatura superiore allo zero assoluto) ne emettono anche. Possiamo quindi trattare la radiazione come un sistema termodinamico, cui possiamo assegnare, se in equilibrio, una temperatura pari a quella di un corpo in equilibrio termico con essa. Lo scambio di energia dipende dal tipo di corpo in questione; si può dimostrare però che un corpo che assorbe tutta la radiazione che vi giunge sopra indipendentemente dalla frequenza, un *corpo nero*, emette anche energia in modo universale, indipendentemente dalla sua composizione, e in funzione della sola temperatura. Useremo perciò il corpo nero per lo studio delle proprietà termodinamiche della radiazione.

Se abbiamo un sistema all'equilibrio la radiazione si distribuisce uniformemente sulle tre coordinate cartesiane, equivalenti, quindi la 4.23 deve essere divisa per 3:

$$P\Omega = \frac{U}{3} \quad (6.8)$$

da confrontarsi con quella dei gas perfetti monoatomici

$$P\Omega = nRT = n\frac{2}{3}C_{\Omega}T = \frac{2}{3}U. \quad (6.9)$$

Per l'irradianza I (potenza emessa al metro quadro) di un corpo nero vale la legge di Štefan-Boltzmann

$$I = \sigma T^4 \quad (6.10)$$

dove σ è la costante di Štefan che vale $5,672 \cdot 10^{-8} \text{ J K}^{-4}/\text{sm}^2$, e che possiamo scrivere in termini di densità di energia

$$u = U/\Omega = \frac{\sigma}{c}T^4, \quad (6.11)$$

quindi

$$P\Omega = \frac{\sigma\Omega}{3c}T^4 \quad (6.12)$$

è l'altra equazione di stato della radiazione⁶.

⁶Quantisticamente, come vedremo più avanti, la radiazione elettromagnetica può essere vista come un gas di fotoni.

Partendo da questo e dalla *costante solare* (intensità della radiazione solare sulla Terra su di una superficie di 1 metro quadro perpendicolare ai raggi solari) $W = 1.37 \text{ kJ/sm}^2$ e sapendo che l' *albedo* a (frazione della radiazione riflessa) della Terra è circa 0.3, si può stimare la temperatura media del nostro pianeta. L'intensità media J della radiazione solare è data da W moltiplicata l'area del disco terrestre vista dal Sole divisa la superficie totale della Terra: $J = W \frac{\pi R^2}{4\pi R^2} = W/4$. Questa va ridotta di a ed eguagliata all'irradianza della Terra, approssimata a un corpo nero, ottenendo $T = [W(1 - a)/4\sigma]^{1/4} = 255 \text{ K}$. La differenza col valore reale di 288 K è dovuta all'effetto serra.

6.2.6 Motori termici, frigoriferi, pompe di calore

Un motore termico è un sistema che percorre un numero indefinito di volte un ciclo termodinamico, scambiando calore con l'esterno e producendo lavoro. Bisogna tener presente che l'"esterno" va inteso in senso generale e può essere anche qualcosa che materialmente sta entro il motore stesso, come per esempio la miscela di carburante e aria che brucia nei motori a combustione interna (benzina, Diesel, eccetera). L'importante è che ci sia una sorgente di calore, un pozzo di calore dove scaricarne, e un sistema meccanico che permetta di usare il lavoro generato. Ovviamente il lavoro generato da un ciclo non può mai essere superiore al calore immesso, e in generale sarà anzi minore. Siccome in un ciclo $\Delta U = 0$, la somma algebrica del calore immesso $Q_1 > 0$, di quello emesso $Q_2 < 0$ e del lavoro compiuto $L > 0$ cambiato di segno deve essere 0. Si definisce *rendimento* η il rapporto tra il lavoro ottenuto e il calore immesso, e si cerca ovviamente di renderlo massimo, che significa il più possibile vicino a 1, visto che per definizione

$$\eta = L/Q_1 = (Q_1 + Q_2)/Q_1 = 1 + Q_2/Q_1 \in [0, 1] \quad (6.13)$$

(si ricordi che $Q_2 < 0$). Il ciclo più semplice che si può usare è quello che usa due sorgenti a temperature diverse, una T_1 da cui ricevere il calore Q_1 e una $T_2 < T_1$ in cui scaricare il calore Q_2 . Siccome le sorgenti sono due e il calore è scambiato solo con quelle, il ciclo deve essere composto da due isoterme a T_1 e T_2 e due adiabatiche che le colleghino. Un ciclo così fatto si chiama *ciclo di Carnot* e ha una particolare importanza perché si dimostra che tutti i cicli possono essere visti come una successione, eventualmente infinita, di cicli di Carnot.

Siccome il rendimento dipende dai calori scambiati e non dal fluido utilizzato, possiamo analizzare il ciclo di Carnot assumendo che il sistema

sia composto da un gas perfetto. Chiamando AB l'isoterma superiore, CD quella inferiore e BC e DA le adiabatice, otteniamo $L = L_{AB} + L_{BC} + L_{CD} + L_{DA}$, e visto che in un' isoterma di un gas perfetto l'energia non cambia si ha pure $L_{AB} = Q_1$ e $L_{CD} = Q_2$, mentre per le adiabatice $L_{BC} = -L_{DA}$, visto che collegano punti con differenze di temperatura uguali e opposte — si ricordi la 6.5. Pertanto

$$\begin{aligned}\eta &= (L_{AB} + L_{BC} + L_{CD} + L_{DA}) / L_{AB} = 1 + L_{CD} / L_{AB} = \\ &= 1 + T_2 \log (\Omega_D / \Omega_C) / T_1 \log (\Omega_B / \Omega_A) .\end{aligned}$$

Sappiamo inoltre che i punti B e C giacciono sulla stessa adiabatice, così come i punti D e A . Applicando la prima delle 6.6 otteniamo

$$\begin{aligned}T_1 \Omega_B^{\gamma-1} &= T_2 \Omega_C^{\gamma-1} \\ T_1 \Omega_A^{\gamma-1} &= T_2 \Omega_D^{\gamma-1}\end{aligned}$$

da cui, dividendo, $\Omega_B / \Omega_A = \Omega_C / \Omega_D$. I logaritmi di conseguenza sono uguali e opposti e il rendimento di un ciclo di Carnot reversibile vale

$$\eta = 1 - T_2 / T_1 . \quad (6.14)$$

Esistono pertanto due modi, non esclusivi, di massimizzare il rendimento. Il primo consiste nel massimizzare la temperatura T_1 alla quale il calore viene fornito, la seconda, meno intuitiva, nel minimizzare la temperatura T_2 del pozzo in cui viene scaricato il calore di scarto. Tuttavia questa seconda strategia trova un limite nel fatto che il pozzo è di fatto l'ambiente in cui si trova il motore, e non è perciò facilmente manipolabile.

Ogni ciclo può essere percorso nei due sensi, nel qual caso i segni di lavoro e calore si invertono. Pertanto un motore termico fatto "girare al contrario" assorbe calore dalla sorgente a temperatura inferiore e lo scarica alla temperatura superiore, utilizzando, invece di produrre, energia che arriva dall'esterno. In questi casi quindi abbiamo un *frigorifero* o una *pompa di calore*. Il concetto di rendimento non è più adeguato, visto che ciò che interessa ora è di massimizzare il trasferimento di calore a parità di energia immessa. Si definisce pertanto un *coefficiente di prestazione* come

$$\omega = Q_2 / |L| \quad (6.15)$$

nel caso di un frigorifero, dove l'interesse sta nel sottrarre calore al sistema freddo, e

$$\text{COP} = Q_1/L \quad (6.16)$$

nel caso di una pompa di calore, dove l'interesse sta nel riscaldare il sistema più caldo. Si vede subito che per un ciclo di Carnot si ha $\omega = Q_2/|Q_1 + Q_2| = 1/|Q_1/Q_2 + 1| = 1/|1 - T_1/T_2| = T_2/(T_1 - T_2)$ e $\text{COP} = T_1/(T_1 - T_2)$. Per una stessa macchina di Carnot tra le stesse temperature, $\text{COP} = \omega + 1$. Questi valori possono essere anche sensibilmente maggiori di 1: per esempio per un frigorifero con temperatura interna 277 K ed esterna 300 K si ha $\omega = 12$, per una pompa che opera tra 275 K e 296 K $\text{COP} = 14$. Si tratta ovviamente di valori ideali, quelli delle macchine reali sono in pratica più bassi, spesso la metà.

6.3 Il Secondo Principio

6.3.1 Formulazioni di Clausius e Kelvin-Planck

Fino a qui abbiamo visto sistemi che eseguono trasformazioni reversibili. In questi casi il senso del tempo non ha importanza: ogni trasformazione che procede in un senso può procedere ugualmente bene nel senso opposto. La nostra esperienza però ci dice che non è sempre così: esistono *processi irreversibili*, che avvengono, riconoscibilmente, solo in un senso e mai nell'altro. Questi processi rompono la simmetria passato-futuro e introducono quindi una *freccia del tempo*.

Il processo più semplice di questo tipo consiste nel passaggio di calore tra due corpi semplicemente a contatto, che avviene sempre dal corpo a temperatura maggiore verso il corpo a temperatura minore. Non si osserva mai il fenomeno opposto. Allora eleviamo questo fatto empirico a principio enunciando così il **Secondo Principio della Termodinamica** nella forma di Clausius:

Non esiste nessuna trasformazione il cui unico risultato sia di trasferire calore da un corpo a temperatura minore a uno a temperatura maggiore.

Si deve fare attenzione, qui come altrove, alla formulazione esatta. Se si omette la clausola "unico risultato" si trovano infatti subito dei controesempi, il più vistoso essendo quello dei frigoriferi. In questo caso il principio non è applicabile in quanto la trasformazione — il ciclo in effetti — richiede l'introduzione di energia dall'esterno, e pertanto il trasferimento di calore non è l'unico risultato.

L'esistenza di questo principio permette di ordinare le temperature, cosa che finora avevamo dato per scontata senza particolare analisi. Possiamo adesso dire che una temperatura T_1 è maggiore di una temperatura T_2 se, mettendo un corpo 1 con temperatura T_1 a contatto con un corpo 2 a temperatura T_2 , in assenza qualunque altro effetto, il calore fluisce sempre dal corpo 1 al corpo 2.

Lo stesso principio può essere enunciato diversamente. La formulazione di Kelvin-Planck afferma che

Non esiste nessun motore termico che trasformi tutto il calore ricevuto in lavoro.

In altri termini, un motore termico deve necessariamente espellere calore in un pozzo termico. Se opera tra due temperature diverse, quindi, ciò che abbiamo chiamato Q_2 non può mai essere nullo. Altrimenti detto, il rendimento di un motore termico non può mai essere 1: $L < Q_1$.

Sorge evidentemente la domanda: le due formulazioni viste del Secondo Principio sono equivalenti? Ovvero, enunciano davvero lo stesso principio? Per dimostrarlo, dimostriamo che se è falso Clausius è falso Kelvin-Planck, e che se è falso Kelvin-Planck è falso Clausius.

Per dimostrare la prima implicazione, prendiamo un motore di Carnot che sfrutta una sorgente di calore a T_1 e un pozzo a T_2 , ovviamente con $T_1 > T_2$. Se esistesse un ciclo che viola Clausius, potremmo utilizzarlo per prendere il calore di scarto Q_2 , che si trova a T_2 e reimmetterlo interamente nella sorgente alla temperatura più alta T_1 , senza nessun'altra conseguenza. Ma allora il motore più questo ciclo formerebbero un motore che produrrebbe lavoro senza avere calore di scarto, in contraddizione con Kelvin-Planck.

Per dimostrare la seconda prendiamo un motore che, per l'ipotesi della violazione di Kelvin-Planck, converte integralmente calore, ricevuto da una sorgente a T_1 , in lavoro L . A questo punto possiamo prendere l'energia L e immetterla in un sistema a temperatura T_0 arbitraria: infatti non c'è nulla che impedisca la trasformazione di lavoro in calore (pensiamo ad attriti, resistenze elettriche eccetera). Ma se prendiamo $T_0 > T_1$ abbiamo un motore che come unico risultato ha il trasferimento di calore da un sistema a temperatura più bassa T_1 a uno a temperatura più alta T_0 , in contraddizione con Clausius. Le due formulazioni sono pertanto equivalenti.

6.3.2 Entropia

Torniamo per un momento alle trasformazioni reversibili e consideriamo un ciclo di Carnot. Sappiamo che il suo rendimento è $\eta = 1 + Q_2/Q_1 = 1 - T_2/T_1$. Questo implica $Q_2/T_2 + Q_1/T_1 = 0$. Scrivendo ora un generico ciclo come somma, eventualmente infinita, di cicli di Carnot, otteniamo il risultato generale $\sum_i Q_i/T_i = 0$. In generale, per cicli reversibili arbitrari,

$$\oint dQ/T = 0. \quad (6.17)$$

Sappiamo che ogni volta che una forma differenziale, integrata su di un percorso chiuso, dà zero, esiste una funzione primitiva della forma stessa. In altre parole, la 6.17 dice che la temperatura è il fattore integrante del calore⁷. Questa funzione si chiama *entropia* e normalmente si indica con S , ed è definita a meno di una costante, che possiamo prendere pari al valore di S in un punto di riferimento O scelto arbitrariamente (si noti la similitudine formale con l'energia potenziale della meccanica). Pertanto il valore si può calcolare tramite l'integrale, preso su di una qualunque trasformazione, purchè reversibile, da O a X :

$$S(X) = \int_O^X dQ/T + S(O). \quad (6.18)$$

L'entropia è evidentemente una grandezza estensiva, visto che lo è il calore ma non la temperatura, ed ha dimensioni fisiche J/K, e come appena visto è una funzione di stato.

Risultano immediatamente alcuni casi particolari: lungo un'adiabatica reversibile $\Delta S = 0$; lungo un'isoterma reversibile $\Delta S = Q/T$; nel caso di una transizione di fase, che avviene a T costante, $\Delta S = \Lambda/T$; lungo un'isobara reversibile $\Delta S = \int nC_P dT/T$, che nel caso di calore specifico costante diventa $\Delta S = nC_P \log(T_f/T_i)$.

Anche l'espressione generale per la variazione dell'entropia di un gas perfetto si ottiene facilmente:

$$\begin{aligned} \Delta S_{AB} &= \int_A^B dQ/T = \int_A^B (dU + Pd\Omega)/T \\ &= \int_A^B (nC_\Omega dT/T + nRd\Omega/\Omega) \\ &= nC_\Omega \log(T_B/T_A) + nR \log(\Omega_B/\Omega_A). \end{aligned} \quad (6.19)$$

⁷Nella formulazione assiomatica della Termodinamica di Carathéodory questa è la definizione stessa, bisogna dire un po' astratta, della temperatura.

Questa relazione si può anche ottenere integrando prima lungo un' isocora e poi lungo un' isoterma, o viceversa prima lungo un' isoterma e poi lungo un' isocora. Il risultato, trattandosi di una funzione di stato, ovviamente non cambia.

6.3.3 Il Terzo Principio

Si è visto che l'entropia si definisce a meno di una costante arbitraria. Questo solitamente non crea problemi, visto che nella maggioranza dei casi si è interessati alle differenze di entropia e non ai valori assoluti. Tuttavia esistono anche casi in cui fissare uno zero all'entropia è necessario, o almeno utile. È possibile fissarlo in modo non arbitrario? Il **Terzo Principio della Termodinamica** risponde a questo quesito. In una delle sue formulazioni infatti recita:

L'entropia di un sistema allo zero assoluto è indipendente da ogni altro parametro termodinamico.

Questo implica che si possano prendere gli stati a $T = 0$ come stati di riferimento per il calcolo dell'entropia e di conseguenza porre l'entropia uguale a 0.

Come il Secondo, anche il Terzo Principio ha varie formulazioni equivalenti. La più celebre è la seguente:

È impossibile portare un sistema allo zero assoluto in un numero finito di trasformazioni.

Non riportiamo qui la dimostrazione dell'equivalenza con l'altra formulazione.

6.3.4 Trasformazioni irreversibili e aumento dell'entropia

Fino a qui si sono viste quasi esclusivamente trasformazioni reversibili, altrimenti dette quasi-statiche. Ma cosa succede quando le trasformazioni, come praticamente sempre nel mondo reale, sono irreversibili? Per analizzare il problema, consideriamo un ciclo composto da quattro trasformazioni:

AB Adiabatica irreversibile

BC Adiabatica reversibile

CD Isoterma reversibile a temperatura T

DA Adiabatica reversibile.

Il calore Q quindi è scambiato solo in CD . Trattandosi di un ciclo, $0 = \Delta U = Q - L \rightarrow L = Q$. Inoltre Q ed L devono essere negativi, altrimenti avremmo un ciclo con trasformazione totale di calore in lavoro, che violerebbe il Secondo Principio. Pertanto è negativo anche $\Delta S_{CD} = Q/T$. Deve essere nulla anche la variazione totale di entropia nel ciclo: $0 = \Delta S = \Delta S_{AB} + Q/T \leq \Delta S_{AB}$. Ne segue che la variazione di entropia in un'adiabatica irreversibile è positiva, ovvero superiore alla variazione — nulla — che si avrebbe se fosse reversibile. Concludiamo quindi che

La variazione di entropia in un sistema isolato non è mai negativa, ed è nulla solo nel caso di trasformazioni reversibili.

Questa è una terza formulazione del Secondo Principio. Anche qui bisogna badare a formularlo bene: la condizione di essere un sistema isolato non può assolutamente essere omessa, altrimenti si avrebbe un semplice controesempio nel congelamento di un bicchiere d'acqua posto nel congelatore di un frigorifero, che vede l'acqua diminuire la propria entropia. Bisogna, in altri termini, considerare il sistema più l'universo (come definito in Termodinamica: i dintorni del sistema che possono essere considerati isolati da tutto il resto) e solo questo avrà variazione positiva o nulla di entropia. I vari sottosistemi possono avere variazioni di entropia di ambo i segni. È in particolare assurda l'affermazione che a volte si sente fare che i sistemi viventi contraddicono il Secondo Principio: un sistema vivente non è affatto isolato, visto che si nutre, respira, elimina calore eccetera. Un essere che non fa queste cose è un essere morto.

Le considerazioni appena fatte chiariscono quale sia il significato fisico dell'entropia. Se prendiamo la formulazione di Kelvin-Planck del secondo principio, vediamo che da un flusso di calore si può estrarre lavoro ma non con efficienza 1: una parte del calore va sempre scartata e trasferita a un pozzo più freddo. Una volta immagazzinata a questa nuova, più bassa, temperatura l'energia è meno disponibile a generare lavoro, in quanto necessita di un altro pozzo a temperatura ancora minore per far funzionare un altro motore termico. Questo nuovo motore scaricherà calore a un pozzo ancora più freddo e così via. L'energia quindi si degrada: a ogni

passaggio diventa sempre meno utile a generare lavoro. L'entropia quantifica questo processo: ci dice quanta energia iniziale si sia degradata e quanto. Siccome il processo è irreversibile la direzione di aumento dell'entropia (in un sistema isolato) ci dice in che direzione scorre il tempo: è la cosiddetta freccia del tempo termodinamica.

6.3.5 Fasi e passaggi di stato

Immaginiamo di seguire una isoterma di van der Waals a partire da volumi molto grandi, a fissato numero di particelle. Se $T > T_c$ queste isoterme somigliano a quelle dei gas perfetti, e il sistema è in una fase fluida, senza avere caratteristiche particolari di un gas o di un liquido. Al di sotto di T_c si osserva che a una certa densità il fluido comincia a condensare. Si forma cioè una nuova *fase termodinamica*, o meglio una coppia di fasi, con caratteristiche distinte: la fase liquida e quella di vapore. Siamo pertanto a una *transizione di fase*, che come si è già accennato avviene a P e T costanti, con l'energia liberata — il calore latente Λ — che "appare dal nulla", venendo in realtà dall'energia interna del fluido, le cui molecole entrano in uno stato legato e pertanto hanno energia potenziale inferiore che nello stato gassoso. In questa condizione, se non si fornisce né sottrae ulteriore calore, liquidi e vapore sono in equilibrio. La transizione da una fase all'altra è determinata dal bilanciamento tra minimizzazione dell'energia e la massimizzazione dell'entropia.

Siccome P rimane costante, la curva reale del fluido di van der Waals non è quella descritta dall'equazione, ma è una retta orizzontale. All'ulteriore calo di volume questa rimane una retta fino a quando tutto il gas è condensato, e a questo punto riprende la forma già vista. Si può dimostrare che le due curve, van der Waals pura e quella reale, hanno la stessa area sottesa: in questo modo si può calcolare la lunghezza del tratto a P costante. La regione coperta dall'insieme di queste rette orizzontali è detta *regione di coesistenza liquido-gas* e termina al punto critico già descritto. Se invece riportiamo tutto su di un grafico $P - \Omega$, questa regione diventa una curva di coesistenza che inizia al punto critico e termina al punto in cui incontra una curva analoga ma che separa lo stato solido da quelli fluidi. Questo punto è detto *punto triplo* perché in quel punto coesistono non solo il liquido e il gas, ma anche il solido. Nel caso dell'acqua, come è noto, il punto triplo si trova a $T = 100^\circ \text{C}$ e $P = 1$ atmosfera.

A differenza della curva di coesistenza liquido-gas, quella che separa il solido dalle altre fasi termina solo su uno degli assi coordinati, mentre dall'altra parte si estende all'infinito. Questo significa che mentre è possibile

passare dallo stato liquido a quello gassoso tramite una trasformazione che gira attorno al punto critico e senza pertanto incontrare una transizione di fase, non è possibile fare una cosa simile partendo, o arrivando, da o allo stato solido. La ragione sta nel fatto che mentre liquido e gas sono entrambi disordinati, e il passaggio da un tipo di disordine all'altro può essere effettuato con continuità, il solido è ordinato (cristallino)⁸ e una tale opzione non può esistere. Va inoltre tenuto presente che il solido — e in alcuni casi anche il liquido — non di rado ha più di una fase, ed esistono di conseguenza molte transizioni. Il ghiaccio per esempio ha la bellezza di 22 fasi note. Un caso ben noto è quello del carbonio, che ammette due fasi solide: la grafite a temperature e pressioni basse, e il diamante a temperature e pressioni alte. Il diamante perciò è, a rigore, metastabile e a differenza del noto motto non è per sempre. Anche in questo caso, però, i tempi di transizione sono astronomicamente lunghi.

Le transizioni di fase che avvengono lungo le curve di coesistenza hanno, come già detto, un calore latente non nullo, mentre viceversa risulta che altre grandezze termodinamiche, come i calori specifici, hanno una discontinuità finita. Queste transizioni sono dette *transizioni del primo ordine*. Diverso è il comportamento al termine della curva di coesistenza liquido-gas: lì la transizione ha calore latente nulla e i calori specifici hanno un asintoto verticale. Si parla di *transizioni del secondo ordine*. Le reincontreremo vedendo i ferromagneti.

6.4 Cenni all' interpretazione statistica

6.4.1 I gas perfetti

Consideriamo una scatola cubica di lato L e volume $\Omega = L^3$, in cui si muovono particelle puntiformi di massa m la cui interazione reciproca è trascurabile ma che si trovino comunque in una condizione di equilibrio termico. Per iniziare consideriamo una singola particella che si muove a velocità $\mathbf{v} = (v_x, v_y, v_z)$. Il suo moto è rettilineo uniforme fino a quando non urta contro una delle pareti, che ipotizziamo liscia e sulla quale supponiamo rimbalzi elasticamente. Considerando il moto lungo l'asse x , a ogni rimbalzo la componente x della quantità di moto cambia di $\Delta p_x = 2mv_x$.

⁸Prescindiamo qui da sistemi come i vetri, che macroscopicamente si comportano come solidi ma hanno una struttura disordinata. La ragione sta nel fatto che i vetri, a rigore, non sono in uno stato stabile ma *metastabile*, e su tempi lunghissimi fluiscono come liquidi. Questi tempi però sono generalmente enormi, molte volte l'età stimata dell'universo.

Questa variazione di p trasferisce un impulso $I = \Delta p_x$ alla parete, e si ripete dopo che la particella ha rimbalzato sulla parete opposta, quindi dopo un tempo $\tau = 2L/|v_x|$. Ora notiamo che ci sono $N \sim N_A \gg 1$ particelle nella scatola, e ognuna urta a suo tempo contro tutte le pareti. Questi urti sono estremamente frequenti, quindi da un punto di vista macroscopico possiamo prendere la media sul tempo e sulle particelle. In particolare, ricordando che $I = \int_{\tau} F dt$, significa che le particelle esercitano mediamente una forza

$$\langle F \rangle = \langle NI/\tau \rangle = N \langle \Delta p_x |v_x| \rangle / 2L = Nm \langle v_x^2 \rangle / L.$$

Notiamo ora che le tre direzioni cartesiane sono equivalenti e che pertanto $\langle v^2 \rangle = \langle v_x^2 + v_y^2 + v_z^2 \rangle = \langle v_x^2 \rangle + \langle v_y^2 \rangle + \langle v_z^2 \rangle = 3 \langle v_x^2 \rangle$ e l'energia media - per ipotesi interamente cinetica - di una particella pertanto vale $\epsilon = \frac{1}{2}m \langle v^2 \rangle = \frac{3}{2}m \langle v_x^2 \rangle$. Si può quindi scrivere

$$\langle F \rangle = \frac{2N\epsilon}{3L}.$$

Dividendo la forza media per l'area L^2 della parete infine otteniamo la pressione esercitata dalle particelle di gas:

$$P = \frac{2N\epsilon}{3L^3} = \frac{2N\epsilon}{3\Omega} = \frac{2U}{3\Omega}.$$

che coincide con la 6.9, e porta a identificare la temperatura come proporzionale all'energia cinetica per particella del gas:

$$T = \frac{2N\epsilon}{3nR} = \frac{2\epsilon}{3k_B}$$

dove si è introdotta la *costante di Boltzmann* $k_B = R/N_A = 1.38 \cdot 10^{-23} \text{ J/K}$.

6.4.2 Equipartizione dell'energia

La relazione vista sopra è un caso particolare di un teorema della Meccanica Statistica classica, il *Teorema di Equipartizione dell'Energia*, che afferma:

Ad ogni grado di libertà quadratico dell' Hamiltoniana (energia), corrisponde, all'equilibrio termodinamico, un'energia pari a $\frac{1}{2}k_B T$.

Di conseguenza il calore specifico molare a volume costante $C_\Omega = \left. \frac{1}{n} \frac{\partial U}{\partial T} \right|_\Omega$, per un sistema con M gradi di libertà quadratici, vale $C_\Omega = Mk_B/2n =$

$MR/2N$. Il calore specifico molare di un sistema di oscillatori armonici vale quindi $3R$. Questo risultato, applicato ai solidi, che a temperature non troppo alte possono essere trattati proprio come un insieme di oscillatori armonici (si veda il par.4.2.1) si chiama *Legge di Dulong e Petit*.

Un gas perfetto monoatomico, avendo tre gradi di libertà per molecola, ha invece $C_\Omega = \frac{3}{2}R$, come già visto. Se biatomico, supponendo la molecola rigida, ha due gradi addizionali di libertà connessi alla rotazione (il terzo è bloccato dal vincolo di rigidità), e otteniamo, ancora, $C_\Omega = \frac{5}{2}R$.

Va sottolineato che il teorema vale specificamente per sistemi che possono essere ben descritti dalla Meccanica Classica. Non vale per la meccanica relativistica e soprattutto non vale per i sistemi descritti dalla Meccanica Quantistica, dove la connessione tra energia cinetica e temperatura **non** è affatto così diretta. Lo si vede dal fatto, che i calori specifici costanti sono incompatibili col Terzo Principio della Termodinamica, che invece come vedremo è conseguenza della Meccanica Quantistica. Abbiamo visto, infatti, che la variazione di entropia di un sistema con calore specifico costante varia come $\log T$, e di conseguenza a $T = 0$ non è 0 ma anzi è infinito. I calori specifici pertanto devono tendere a 0 per $T \rightarrow 0$. E infatti gli esperimenti mostrano che i calori specifici dei solidi all'abbassarsi della temperatura tendono a zero: proporzionalmente a T per i metalli, a T^3 per gli isolanti, come si vedrà.

6.4.3 Entropia e disordine

Si è visto che in un sistema isolato l'entropia non può decrescere. Dalla 6.19 vediamo che per un gas perfetto S cresce con il volume e con la temperatura. Se quindi consideriamo una scatola adiabatica e rigida contenente un gas perfetto, questo gas avrà temperatura costante, vista la 6.5 e il fatto che il sistema è isolato. L'entropia pertanto è proporzionale al logaritmo del volume. Se il gas fosse vincolato in qualche modo a restare in un settore del contenitore e fatto equilibrare lì la sua entropia sarebbe minore che nel caso di gas occupante tutto il contenitore. Rimuovendo il vincolo il gas, dato il moto caotico delle molecole, finisce per espandersi in modo evidentemente irreversibile raggiungendo la sua entropia finale, maggiore di prima: ecco un esempio di trasformazione di non equilibrio che fa aumentare l'entropia.

Come possiamo caratterizzare questo dal punto di vista della meccanica delle particelle componenti? Notiamo che, come detto in 6.1, la Termodinamica comprime l'informazione da quella dettagliatissima dei microstati in pochi valori che caratterizzano un macrostato. Viceversa, a un macro-

stato corrispondono moltissimi microstati, dati dall'insieme delle posizioni e velocità istantanee di tutte le particelle. Questo portò Boltzmann a porre una relazione tra il numero di microstati di un sistema isolato Γ_α compatibili con un macrostato α e l'entropia di α :

$$S_\alpha = k_B \log \Gamma_\alpha, \quad (6.20)$$

un'equazione così famosa da essere incisa sulla lapide della tomba di Boltzmann stesso. Si noti che se i microstati sono equiprobabili — *Principio di equiprobabilità a priori*, che si pone per sistemi isolati — Γ_α è una probabilità, non normalizzata, che il sistema si trovi in uno dei microstati compatibili con α , e quindi che il sistema si trovi in α . Si può vedere che l'entropia così definita è estensiva. È inoltre immediato vedere che se ad α corrisponde un solo microstato allora $S_\alpha = 0$, il che ci permette di giustificare il Terzo principio come l'osservazione che a $T = 0$ il solo stato possibile è quello a energia minima, che risulta essere unico.

Se ora torniamo al gas perfetto visto appena sopra, vediamo anche che il processo che porta all'espansione porta anche ad avere molti più stati compatibili col macrostato, visto che un volume maggiore è compatibile con molte più traiettorie che un volume minore. Pertanto l'espansione porta a un Γ maggiore e, come deve essere, a un'entropia maggiore.

Viene spontaneo quindi associare uno stato a entropia maggiore con uno stato "più disordinato", visto che siamo soliti associare l'ordine con oggetti raccolti in posti ben precisi e non dispersi per ogni dove... E tuttavia ci sono da tenere presenti alcune considerazioni⁹:

- Il concetto di ordine (e la sua negazione, il disordine) è apparentemente familiare a tutti. Per esempio, una successione di lettere è generalmente riconosciuta come ordinata se corrisponde ad una parola elencata sul dizionario. Lettere disordinate fanno pensare ad una scimmia che preme a caso i tasti di un sistema di scrittura. Tuttavia, un testo apparentemente disordinato o casuale secondo il linguaggio ordinario, potrebbe essere il risultato di una codifica di un messaggio perfettamente "ordinato". Gli esempi potrebbero continuare ma il punto chiave è che ordine o disordine sono concetti vaghi finché non si dice secondo quale criterio. I criteri però possono essere diversi ed anche contrastanti.
- Occorre perciò definire bene il tipo di ordine a cui siamo interessati dal punto di vista della termodinamica statistica. Di fatto quello che è

⁹Ringrazio il professor Giorgio Pastore per avermi permesso di riportare quanto segue. Per una discussione più ampia si veda il libro di Callen citato in bibliografia.

possibile fare in meccanica statistica è di definire il grado di disordine nella distribuzione del sistema tra i possibili microstati.

- Avendo così ristretto il senso della parola ordine (disordine), possiamo, dal punto di vista tecnico, introdurre una misura del disordine di questa distribuzione mediante lo stesso formalismo sviluppato da Shannon per la teoria dell'informazione.
- Come è stato dimostrato da Khinchin, poche richieste ragionevoli su questa misura del disordine in termini della distribuzione di frequenze tra microstati, f_j , permettono di arrivare ad una formula per il "disordine":

$$\text{"disordine"} = -k \sum_j f_j \log f_j$$

e dove k è una costante positiva. Questo "disordine" è massimo per microstati equiprobabili. In tal caso vale:

$$\text{"disordine"} = k \log \Gamma$$

con Γ = numero di microstati.

Quindi per un sistema isolato, l'entropia corrisponde alla misura quantitativa di Shannon del massimo possibile disordine nella distribuzione del sistema sui possibili microstati. Pertanto:

- a) entropia = disordine sì, ma solo nel senso tecnico di misura del numero dei microstati accessibili al sistema isolato.
- b) entropia = disordine no, se pensiamo all'ordine "macroscopico" del sistema.

Capitolo 7

Elementi di Meccanica Quantistica

7.1 Origine

Alla fine del XIX secolo la Fisica Classica, nonostante i suoi enormi successi, doveva affrontare anche alcune serie difficoltà. Alcune di queste, si è visto in precedenza, collegate all'elettromagnetismo e alle alte velocità portarono allo sviluppo della Teoria della Relatività. Altre invece riguardavano il mondo microscopico. Elenchiamo le principali.

Corpo nero

Sulla base di considerazioni termodinamiche Planck riuscì a stabilire che la forma della densità di energia della radiazione termica in equilibrio con un corpo nero a temperatura T , nell'intervallo $[\nu, \nu + \Delta\nu]$ dovesse essere della forma

$$u \sim \frac{\nu^3 \Delta\nu}{e^{h\nu/k_B T} - 1} \quad (7.1)$$

dove k_B è la costante di Boltzmann e h è una nuova costante, detta *costante di Planck*, che vale $6.63 \cdot 10^{-34}$ Js. Il tentativo di dedurla dalla elettrodinamica più la Meccanica Statistica classica si arenò, e infine si vide che in effetti un calcolo classico portava a un risultato non solo errato, ma assurdo: la densità di energia divergeva per $\nu \rightarrow \infty$ - la cosiddetta "catastrofe ultravioletta" - e l'energia totale della radiazione risultava infinita. Planck, alla fine dell'anno 1900, riuscì a ottenere la 7.1 usando l'ipotesi

che l'energia elettromagnetica venisse scambiata tra corpo nero e radiazione sotto forma di *quanti* di energia pari a $\epsilon = h\nu$. L'energia irradiata a ogni frequenza, in altre parole, era un multiplo di un quanto minimo, fissato per quella particolare frequenza. Questa ipotesi però non poté essere giustificata teoricamente in alcuna maniera.

Effetto fotoelettrico

Proiettando della luce ultravioletta su di un metallo si osserva l'emissione di elettroni, il cosiddetto *effetto fotoelettrico*. Classicamente ci si aspetta che il numero e l'energia degli elettroni emessi siano proporzionali al quadrato dell'ampiezza del campo elettrico E dell'onda elettromagnetica, diminuita dell'energia potenziale necessaria per emettere l'elettrone dal metallo, il cosiddetto *lavoro di estrazione* φ . L'emissione dovrebbe essere sempre presente, anche se molto debole, in presenza di onde anche molto deboli, ma dovrebbe cominciare, in questo caso, dopo un certo tempo in cui l'energia si accumula nel sistema di elettroni. Gli esperimenti mostrano un quadro piuttosto diverso: solo il numero degli elettroni emessi dipende dall'ampiezza di E , mentre la loro energia dipende dalla frequenza dell'onda. Inoltre non si ha emissione, nemmeno a grandi intensità dell'onda, sotto una certa soglia di frequenza, ma quando la si supera si osservano eventi di emissione anche dopo brevissimo tempo. In un articolo del 1905 Einstein riprese la proposta dei pacchetti di energia, detti *fotoni*, affermando che gli eventi di emissione corrispondevano all'assorbimento da parte di un singolo elettrone di un singolo fotone: se l'energia $h\nu > \varphi$, l'emissione ha luogo e l'elettrone viene emesso con energia $\epsilon = h\nu - \varphi$. Il valore soglia per la frequenza è quindi $\nu_0 = \varphi/h$. Superata la soglia, maggiore è il numero di quanti, maggiore il numero di elettroni emessi. Nel 1921 Einstein ricevette il premio Nobel "Per i contributi alla Fisica teorica, in particolare per la spiegazione dell'effetto fotoelettrico".

Calori specifici

Si è visto che l'esistenza di calori specifici costanti non è compatibile col Terzo Principio della Termodinamica. Sfortunatamente la meccanica Statistica Classica porta alla legge di Dulong e Petit, che afferma che per i solidi a bassa temperatura accade esattamente questo. Sempre Einstein, nel 1906, mostrò che una quantizzazione delle energie di oscillazione corrispondenti alle frequenze di oscillazione atomiche ω secondo la solita legge $\epsilon = \hbar\omega = h\nu$, con $\hbar = h/2\pi = 1.05 \cdot 10^{-34}$ Js, portava a calori specifici

che vanno a zero per $T \rightarrow 0$. Poco dopo Debye, raffinando il modello, mostrò che per gli isolanti $C_\Omega \sim T^3$.

Modello atomico

Le esperienze di Rutherford, negli primi anni del secolo, mostrarono che l'atomo era composto di un nucleo pesante, carico positivamente, circondato dagli elettroni. Questo generò una serie di problemi: cariche accelerate irradiano in continuazione perdendo energia, ma questo nell'atomo non succede, e se succedesse l'atomo collasserebbe in una frazione di secondo¹. Invece si osserva emissione e assorbimento solo a frequenze (*righe spettrali*) ben definite, date dalla *formula di Rydberg*

$$\frac{1}{\lambda_{mn}} = R \left(\frac{1}{m^2} - \frac{1}{n^2} \right) \quad (7.2)$$

dove R è una costante pari a circa $1.10 \cdot 10^{-7} \text{ m}^{-1}$. Nel 1913 Bohr propose che esistesse una *regola di quantizzazione* per le orbite circolari, poi generalizzata da Sommerfeld anche per le orbite ellittiche, che imponesse come uniche orbite possibili quelle aventi un momento angolare multiplo intero di $\hbar$: $l = n\hbar$, $n \geq 1$. Gli elettroni avrebbero in questo modello avrebbero avuto energie quantizzate $E_n = -e^2/8\pi\epsilon_0 a_0 n^2$, con $a_0 = 4\pi\epsilon_0 \hbar^2 / m e^2 = 0.529 \text{ \AA}$ la *lunghezza di Bohr*, e avrebbero potuto abbandonare la loro orbita caratterizzata da n solo per andare, tramite un *salto quantico*, a un'altra caratterizzata da m , assorbendo o cedendo la differenza di energia tramite assorbimento o emissione di un fotone di energia corrispondente $\epsilon = E_m - E_n = h\nu = hc/\lambda$. In questo modo vengono spiegate sia la stabilità degli atomi, sia i loro spettri, sia la formula 7.2.

7.1.1 Fotoni e onde di materia

L'esistenza di queste quantizzazioni, in fenomeni tanto diversi, era una sfida per la Fisica Classica, che non poteva renderne conto. Eppure l'evidenza sperimentale sempre crescente non lasciava dubbi, e richiedeva un mutamento radicale del quadro teorico. Il punto fondamentale era: come connettere la Fisica nota, basata sul continuo, con questi fenomeni che richiedono invece una discretizzazione?

¹Oggi noi sappiamo che il sistema sarebbe anche fortemente caotico e incompatibile, anche per questo, con la stabilità osservata degli atomi.

Nel 1924 De Broglie propose di rovesciare il percorso fatto per i fotoni. Questi ultimi potevano essere visti come onde elettromagnetiche che si comportavano come particelle, prive di massa in quanto moventesi a c ma con energia $\epsilon = h\nu$ e quantità di moto $p = \epsilon/c$. Perchè non supporre che le particelle potessero comportarsi come onde? In questo modo nei casi in cui il moto risulta limitato, come nelle orbite atomiche, le condizioni al contorno porterebbero a delle limitazioni alle lunghezze d'onda ammissibili: dovrebbero infatti essere tali per cui la lunghezza dell'orbita (di raggio r) sia un multiplo intero di tale lunghezza d'onda, $n\lambda = 2\pi r$. Questo implica una quantizzazione delle orbite, $r_n = n\lambda/2\pi = nv/2\pi\nu$ dove v è la velocità dell'onda. Svolgendo i calcoli, risultò che questo portava alle condizioni di Bohr.

7.2 L'equazione di Schrödinger

Nel 1926 Schrödinger propose, sulla base di considerazioni euristiche, un'equazione per le onde di materia:

$$\hat{H}\psi = E\psi \quad (7.3)$$

dove $\hat{H}$ è un operatore differenziale chiamato *hamiltoniano*, scrivibile come somma di un termine cinetico $\hat{K} = -\frac{\hbar^2}{2m}\nabla^2$ e uno potenziale $\hat{V}$. Si tratta, come accennato, di un'equazione agli autovalori, e l'autovalore E è l'energia del sistema. Nel caso dell'atomo di idrogeno l'equazione diventa

$$-\frac{\hbar^2}{2m}\nabla^2\psi - \frac{e^2}{4\pi\epsilon_0 r}\psi = E\psi. \quad (7.4)$$

Qui m a rigore sarebbe la massa ridotta del sistema nucleo-elettrone, ma identificarla con la massa dell'elettrone dà un'ottima approssimazione. Cercando le soluzioni a simmetria sferica, rimangono solo le parti dipendenti da r e l'equazione si semplifica in

$$-\frac{\hbar^2}{2mr^2}\frac{\partial}{\partial r}\left(r^2\frac{\partial\psi}{\partial r}\right) - \frac{e^2}{4\pi\epsilon_0 r}\psi = E\psi. \quad (7.5)$$

Le soluzioni possono essere cercate, per esempio, tramite uno sviluppo in serie, ma qui seguiamo una strada più semplice ipotizzando soluzioni del tipo $\psi(r) = P_n(r)e^{-r/\alpha_n}$, dove P_n è un opportuno polinomio di grado n . Inserendo questa soluzione nella 7.5 vediamo che queste sono effettivamente le soluzioni, purché si prenda $\alpha_n = n4\pi\epsilon_0\hbar^2/mc^2 = na_0$ e $P_n = c_n L_n^1$ con L_n^1 i c.d. polinomi di Laguerre associati e c_n opportune

costanti che assicurano la normalizzazione $\int_{\infty} |\psi|^2 d\mathbf{r} = 1$. La cosa che più interessa qui è che gli autovalori sono esattamente quelli previsti da Bohr: $E_n = -e^2/8\pi\epsilon_0 a_0 n^2$. L'equazione di Schrödinger quindi descrive esattamente l'atomo di idrogeno e ha avuto da allora in poi innumerevoli conferme in moltissimi campi della Fisica e della Chimica.

7.3 Stati e osservabili

L'equazione 7.3, base della Meccanica Quantistica, è come abbiamo visto un'equazione agli autovalori. Queste equazioni si definiscono su spazi vettoriali, e in questo caso su di uno spazio di funzioni ψ su $\mathbb{C}$, che possiamo prendere con un prodotto scalare $(\psi, \phi) = \int_{\infty} \psi^* \phi d\Omega$. Ma che cosa rappresentano queste funzioni? Scartata l'idea, per considerazioni che via via vedremo, che rappresentino la densità di massa o di carica della particella, anche perchè risulta che in generale ψ è complessa, si è visto che il significato sta nel fornire la probabilità di trovare la particella in uno specifico posto. Per l'esattezza, *la probabilità di trovare la particella in un volume Ω attorno al punto $\mathbf{R}$ è data da $P(\mathbf{R}) = \int_{\Omega} |\psi(\mathbf{r})|^2 d\mathbf{r}$* . La grandezza $\mathcal{P}(\mathbf{r}) = |\psi(\mathbf{r})|^2$ è una *densità di probabilità*, e $\psi(\mathbf{r})$ un'*ampiezza di densità di probabilità*. Diciamo quindi che la particella *non ha una posizione determinata*, e pertanto *non percorre una traiettoria*. Le funzioni ψ in un atomo perciò non sono orbite, ma oggetti ancora più astratti - funzioni a valori complessi - detti *orbitali*². Possiamo descrivere questa situazione dicendo che una particella non ha una posizione precisa in nessun tempo, ma è in qualche modo potenzialmente in ogni luogo ove $\psi \neq 0$, e una misura di posizione la "costringe a scegliere" un punto dove stare. Questo viene detto *Principio di Sovrapposizione*.

Data questa interpretazione, risulta immediato che debba essere soddisfatta la *condizione di normalizzazione*

$$||\psi|| = (\psi, \psi) = \int_{\infty} |\psi(\mathbf{r})|^2 d\mathbf{r} = \int_{\infty} \mathcal{P}(\mathbf{r}) d\mathbf{r} = 1 \quad (7.6)$$

in quanto è certo che la particella sia da qualche parte nell'Universo.

Data questa interpretazione, se A è una grandezza fisica misurabile — un' *osservabile* — che assume il valore $A(\mathbf{r})$ quando la particella si trova in $\mathbf{r}$, il valor medio di A quando la particella occupa un orbitale deve essere

²In alcuni testi didattici purtroppo si trovano definizioni di orbitale del tipo "la regione dello spazio in cui la probabilità di trovare la particella è il 90%". Si tratta di una definizione completamente errata: questo al più è un criterio per visualizzare gli orbitali, che sono, va ribadito, funzioni a valori complessi e **non** regioni dello spazio.

$$\langle A \rangle = \int \mathcal{P}(\mathbf{r}) A(\mathbf{r}) d\mathbf{r} = \int |\psi(\mathbf{r})|^2 A(\mathbf{r}) d\mathbf{r} = \int \psi^*(\mathbf{r}) A(\mathbf{r}) \psi(\mathbf{r}) d\mathbf{r} = (\psi, A\psi). \quad (7.7)$$

7.4 Formulazione generale

Quanto detto nel paragrafo precedente si può generalizzare rendendolo più astratto. A fini di concisione e chiarezza, la trattazione non sarà particolarmente rigorosa. In particolare, saranno spesso trascurate questioni di dominio di definizione. Quanto detto qui sotto di fatto è corretto per spazi di Hilbert a dimensione finita e richiederebbe alcune modifiche tecniche per essere esteso a spazi di dimensione infinita, ma per semplicità questa estensione non viene esplicitata.

Poniamo quindi questi postulati:

1. *Ad ogni sistema fisico viene associato un opportuno spazio di Hilbert $\mathcal{H}$ su $\mathbf{C}$.* Quale spazio di Hilbert si debba utilizzare va deciso di volta in volta; questo è il corrispettivo quantistico della procedura che si segue in Meccanica Classica di identificazione dei gradi di libertà del sistema. In Meccanica Quantistica però lo spazio di Hilbert è solitamente infinito, spesso lo spazio delle funzioni a quadrato sommabile (L^2) con l'integrale di Lebesgue, mentre quello classico è $\mathbf{R}^N$.
2. *Lo stato di un sistema è completamente descritto da un sottospazio unidimensionale di $\mathcal{H}$.* Potremmo dire da un vettore, senonché un vettore può essere sempre normalizzato ponendo la sua lunghezza pari ad 1 (vale a dire dividendo per la sua norma: è per questo che si richiedono funzioni a quadrato sommabile). La normalizzazione risulta necessaria solo in specifiche circostanze, per cui spesso si opera con vettori non normalizzati. I vettori di $\mathcal{H}$ li indicheremo, usando la cosiddetta notazione di Dirac, tramite una parentesi angolare: $|\psi\rangle$.
3. *Ad ogni osservabile di un sistema fisico viene associato un operatore lineare autoaggiunto definito su $\mathcal{H}$ a valori in $\mathcal{H}$.* Stante il teorema di Riesz, ogni operatore lineare (limitato) può essere espresso in modo unico tramite un prodotto scalare con un opportuno vettore, ed esiste di conseguenza un isomorfismo tra lo spazio dei vettori su $\mathcal{H}$ e quello degli operatori su $\mathcal{H}$ (detto *spazio duale*). Se indichiamo il

prodotto scalare tra vettori $|u\rangle$ e $|v\rangle$ con la notazione $\langle u|v\rangle$ (in L^2 per definizione $\langle u|v\rangle = \int u^*(x)v(x)dx$), possiamo identificare il "vettore" $\langle u|$ con l'operatore ad esso associato.

Un operatore $\hat{A}$ si dice *autoaggiunto* o *hermitiano* se coincide col suo aggiunto $\hat{A}^\dagger$, definito dalla relazione³ $\langle \hat{A}^\dagger u|v\rangle = \langle u|\hat{A}v\rangle$. Un operatore autoaggiunto quindi si può trasferire da un lato all'altro di un prodotto scalare: $\langle \hat{A}u|v\rangle = \langle u|\hat{A}v\rangle$.

Se $\mathcal{H}$ è a dimensione finita n , i vettori sono n -ple di numeri complessi, il prodotto scalare la somma dei prodotti delle componenti dei vettori (coniugando quelle del primo), gli operatori sono matrici $n \times n$, e gli operatori autoaggiunti sono le matrici hermitiane (ovvero quelle per cui $a_{ij} = a_{ji}^*$).

4. *I possibili risultati della misura di un'osservabile sono gli autovalori dell'operatore associato ad essa.* Ovvero, se A è un'osservabile e $\hat{A}$ l'operatore associato, i valori a che si possono trovare misurando A sono quelli che soddisfano la

$$\hat{A}|\phi_a\rangle = a|\phi_a\rangle. \quad (7.8)$$

Il vettore $|\phi_a\rangle$ si chiama *autostato associato all'autovalore a* .

Da questo postulato si vede la necessità della condizione di autoaggiunzione, in quanto si dimostra che *gli autovalori di un operatore autoaggiunto sono reali*, e i risultati delle misure devono ovviamente essere reali. Inoltre ricordiamo che *gli autovettori di un operatore autoaggiunto costituiscono una base ortogonale di $\mathcal{H}$* .

Va notato che la relazione tra autostato e autovalore non è biunivoca. Ad ogni autostato corrisponde un autovalore, ma a un autovalore possono corrispondere più stati. In questo caso è facile vedere che tutte le combinazioni lineari di autostati corrispondenti a un determinato autovalore sono anche autostati corrispondenti allo stesso autovalore. Nel caso si parla di *stati degeneri*.

5. *A ogni osservabile è associata una statistica di misure, vale a dire che per ogni autovalore a associato a un'autovettore $|\phi_a\rangle$ e ogni funzione $|\psi\rangle$ la probabilità di ottenere a eseguendo una misura di A su $|\psi\rangle$ è $|\langle \phi_a|\psi\rangle|^2$, ovvero il modulo della proiezione del vettore di*

³Questo uso della notazione di Dirac è un po' eterodosso, ma confido che sia comprensibile.

stato sul corrispondente autovettore di $\hat{A}$, e la proiezione stessa è un' *ampiezza di probabilità*. Ne segue che *Il valor medio delle misure di un' osservabile A su di uno stato ψ è dato da $\langle \psi | \hat{A} | \psi \rangle$* ⁴. Infatti, sviluppando $|\psi\rangle = \sum_a |\phi_a\rangle \langle \phi_a | \psi \rangle$ e $\hat{A}|\psi\rangle = \sum_a \hat{A}|\phi_a\rangle \langle \phi_a | \psi \rangle = \sum_a a|\phi_a\rangle \langle \phi_a | \psi \rangle$:

$$\begin{aligned} \langle \psi | \hat{A} | \psi \rangle &= \sum_{a,a'} a \langle \phi_{a'} | \langle \psi | \phi_{a'} \rangle \langle \phi_a | \psi \rangle | \phi_a \rangle = \sum_{a,a'} a \langle \psi | \phi_{a'} \rangle \langle \phi_a | \psi \rangle \overbrace{\langle \phi_{a'} | \phi_a \rangle}^{\delta_{a,a'}} \\ &= \sum_a a \langle \psi | \phi_a \rangle \langle \phi_a | \psi \rangle = \sum_a a |\langle \psi | \phi_a \rangle|^2 = \sum_a a P_a = \langle a \rangle. \end{aligned}$$

La notazione $\langle A \rangle = \langle \psi | \hat{A} | \psi \rangle$ è quella comunemente usata, e mette in risalto il fatto che se $\hat{A}$ è autoaggiunto, può essere indifferentemente applicato al vettore a destra come a quello a sinistra.

6. *Se si effettua una misura di A al tempo t ottenendo un valore a e la si ripete dopo un tempo $t' \rightarrow t + 0^+$, il risultato sarà con certezza di nuovo a . Visti gli altri postulati, questo significa che al momento in cui si effettua una misura di A con risultato l'autovalore a , il vettore di stato $|\psi\rangle$ diventa immediatamente il corrispondente autovettore $|\phi_a\rangle$. Questo evento viene chiamato *collasso della funzione d'onda*.*

Dal fatto che gli autostati di un operatore associato a un' osservabile qualunque A costituiscono una base di $\mathcal{H}$ segue che ogni stato $|\psi\rangle$ può essere sviluppato in una combinazione lineare di autostati di un' arbitraria A . Questa è una formulazione più generale del Principio di Sovrapposizione visto prima. Il suo significato sta quindi nel fatto che ogni stato contiene, in linea di principio, altri stati e ha, in linea di principio, la possibilità di dare un qualunque risultato per la misura di un' osservabile.

Questi non sono tutti i postulati richiesti dalla nostra formulazione; gli altri verranno enunciati in seguito.

7.5 Operatori con analogo classico

Quali siano gli operatori corrispondenti alle grandezze misurabili è questione empirica. Questa ricerca viene resa più facile se le osservabili hanno un analogo classico, ovvero se continuano a esistere anche nel limite

⁴Si tenga conto che la somma è sugli stati e non sugli autovalori: un autovalore degenero va contato tante volte quanti sono gli autostati indipendenti ad esso associati.

formale in cui $\hbar \rightarrow 0$ e al tempo stesso gli autovalori tendono all'infinito. Vedremo in seguito osservabili senza analogo classico.

In primo luogo, abbiamo visto che l'operatore "energia potenziale" $\hat{V}$ è semplicemente l'operatore moltiplicativo per la funzione $V(\mathbf{r})$. Concludiamo che l'operatore "posizione x " $\hat{x}$ non è altro che la moltiplicazione per la coordinata. L'equazione agli autovalori diventa

$$\hat{x}|x_\xi\rangle = \xi|x_\xi\rangle \quad (7.9)$$

con ξ l'autovalore, ovvero la posizione possibile della particella. Non è un caso fortunato perché l'autofunzione corrispondente a ξ non è un vettore di $\mathcal{H}$ in senso stretto. Questa è una ragione per cui la formulazione qui presentata non è in realtà sufficiente, ma per semplicità ignoreremo questa difficoltà, limitandoci a osservare che la generalizzazione a questi casi esiste. Siccome non esistono limitazioni su ξ , che può appartenere a un segmento dell'asse reale (in effetti, a tutto $\mathbb{R}$) diciamo che l'operatore posizione ha *spettro continuo*.

Applicando i postulati espressi sopra, vediamo che il modulo quadro del prodotto scalare $\langle x|\psi\rangle$ definisce la probabilità di trovare una particella situata in uno stato $|\psi\rangle$, o meglio l'ampiezza di densità di probabilità, e si identifica quindi con la funzione $\psi(x)$ usata da Schrödinger: $\psi(x) = \langle x|\psi\rangle$. Si parla di $|\psi\rangle$ in *rappresentazione x* . Da questo vediamo che gli autostati di $\hat{x}$, in questa rappresentazione, sono istanze della δ di Dirac, $\langle x|x_\xi\rangle = \delta(x - \xi)$, descritta in appendice C.

Ricordiamo poi che in Meccanica Classica $K = p^2/2m$, e abbiamo visto che quantisticamente vale $\hat{K} = -\hbar^2\nabla^2/2m$. La corrispondente relazione tra energia cinetica e quantità di moto quindi è $-\hbar^2\nabla^2 = \hat{p}^2$, ovvero

$$\hat{\mathbf{p}} = -i\hbar\nabla \quad (7.10)$$

e la corrispondente equazione agli autovalori è, per una componente,

$$\hat{\mathbf{p}}_x\psi = -i\hbar\partial\psi/\partial x = q\psi \rightarrow \psi = Ae^{iqx/\hbar} = Ae^{ikx} \quad (7.11)$$

dove $k = q/\hbar$ è il numero d'onda, esattamente come per le onde che abbiamo già visto, e A la costante di normalizzazione. Si tratta di una funzione periodica: se imponiamo condizioni periodiche alle estremità di un segmento di lunghezza L allora $k = 2\pi j/L$, $j \in \mathbb{N}$, e siccome $|\psi|^2 = A^2$, il valore di A è $= 1/\sqrt{L}$. In tre dimensioni quindi gli autostati sono

$$\psi_{\mathbf{k}} = \frac{1}{\sqrt{\Omega}}e^{i\mathbf{k}\cdot\mathbf{r}} \quad (7.12)$$

e per $\Omega \rightarrow \infty$ corrispondono ad autovalori $\mathbf{k}$ continui.

Queste autofunzioni lo sono anche dell'operatore energia cinetica, con in questo caso $E_k = \hbar^2 k^2 / 2m$. Sono quindi le autofunzioni che descrivono una particella libera, o quantomeno la parte dipendente dalle coordinate dell'autofunzione.

7.6 Commutatori

Siccome tra operatori sono definiti una somma, un prodotto e la moltiplicazione per uno scalare (come detto, sono uno spazio duale dello spazio di Hilbert su cui sono definiti), gli operatori definiscono un'algebra. Mentre l'addizione struttura l'insieme degli operatori come gruppo abeliano, la moltiplicazione genera sì un gruppo, ma non abeliano visto che in generale la moltiplicazione tra operatori non è commutativa: $\hat{A}\hat{B} \neq \hat{B}\hat{A}$. Si pensi, per esempio alla moltiplicazione tra matrici. Dati due operatori qualunque $\hat{A}, \hat{B}$ chiamiamo quindi *commutatore* l'operatore

$$[\hat{A}, \hat{B}] = \hat{A}\hat{B} - \hat{B}\hat{A}. \quad (7.13)$$

Si definisce anche l' *anticommutatore* $\{\hat{A}, \hat{B}\} = \hat{A}\hat{B} + \hat{B}\hat{A}$.

Per una coordinata e la corrispondente componente della quantità di moto,

$$[\hat{x}, \hat{p}_x] = -i\hbar x \partial/\partial x + i\hbar \partial(x\cdot)/\partial x = i\hbar \quad (7.14)$$

mentre, ovviamente, $[\hat{x}, \hat{p}_y] = 0$ e così via per le altre componenti. È immediato vedere che il commutatore tra un operatore $\hat{A}$ e un operatore funzione solo di $\hat{A}$, $\hat{f}(\hat{A})$ è nullo. Per esempio, il commutatore della quantità di moto con l'energia cinetica è nullo.

Il significato della commutazione o meno di due operatori è molto più profondo di quanto potrebbe sembrare, ed è collegato al fatto che l'applicazione di un operatore significa effettuare una misura dell'osservabile corrispondente. La mancata commutazione di due operatori significa che l'ordine in cui si prendono le misure non è, in generale, influente. Infatti effettuare una misura dell'osservabile associato a $\hat{B}$, ottenendo l'autovalore b , significa far collassare lo stato iniziale $|\psi\rangle$ su di un autostato $|\phi_b\rangle$ di $\hat{B}$. Prendere di seguito una misura dell'osservabile $\hat{A}$ significa trovare un autostato $|\chi_a\rangle$ di $\hat{A}$ corrispondente all'autovalore a , con probabilità $p(a) = \langle \chi_a | \phi_b \rangle$. Se $\hat{A}$ e $\hat{B}$ non hanno una base comune di autostati, esisteranno diversi valori di a con probabilità $p(a) \neq 0$; se ce l'hanno, ci sarà un solo valore di a . Se ora si inverte l'ordine di $\hat{A}$ e $\hat{B}$ avremo viceversa

che esisteranno diversi valori di b con probabilità $p(b) \neq 0$; se ce l'hanno, ci sarà un solo valore di b .

Si conclude pertanto che *due operatori commutano se e solo se ammettono una base di autovettori comune*.

7.6.1 Il Principio di Indeterminazione

Consideriamo due osservabili A e B che per semplicità prenderemo a media nulla (i risultati non cambiano rispetto al caso generale). Chiamando σ_A^2 il quadrato della deviazione standard delle misure di A in uno stato qualunque $|\psi\rangle$, $\sigma_A^2 = \langle A^2 \rangle - \langle A \rangle^2 = \langle A^2 \rangle$, il prodotto delle deviazioni è

$$\sigma_A^2 \sigma_B^2 = \langle \hat{A}^2 \rangle \langle \hat{B}^2 \rangle = \langle \psi | \hat{A}^2 | \psi \rangle \langle \psi | \hat{B}^2 | \psi \rangle = \|\hat{A}|\psi\rangle\|^2 \|\hat{B}|\psi\rangle\|^2 \geq |\langle \psi | \hat{A}\hat{B} | \psi \rangle|^2$$

dove si è usata la disuguaglianza di Schwarz. Scrivendo il prodotto $\hat{A}\hat{B}$ come la semisomma del commutatore con l'anticommutatore, e notando che (*Esercizio*) il valor medio del commutatore di due operatori autoaggiunti è immaginario puro mentre quello del loro anticommutatore è reale, si ottiene

$$\sigma_A^2 \sigma_B^2 \geq |\langle \hat{A}\hat{B} \rangle|^2 = \left| \left\langle \frac{1}{2} \left([\hat{A}, \hat{B}] + \{\hat{A}, \hat{B}\} \right) \right\rangle \right|^2 \geq \frac{1}{4} \left| \langle [\hat{A}, \hat{B}] \rangle \right|^2$$

ovvero

$$\sigma_A \sigma_B \geq \frac{1}{2} \left| \langle [\hat{A}, \hat{B}] \rangle \right|. \quad (7.15)$$

Questa relazione è nota come *Principio di Indeterminazione*, o *Principio di Heisenberg*, anche se come si vede si tratta in realtà di un teorema. Questa relazione proibisce che misurando contemporaneamente su di uno stesso stato due osservabili i cui operatori non commutano (o, come anche si dice, *non compatibili*), si possano ottenere risultati infinitamente precisi⁵ per entrambi. Non vieta, però, di ottenerne uno infinitamente preciso, purché l'altro sia completamente indeterminato. Non vieta nemmeno di misurare contemporaneamente con infinita precisione grandezze rappresentate da operatori che commutano.

Possiamo dare un significato a questa relazione notando che la prima misura porta lo stato originale $|\psi\rangle$ a collassare su di un autostato $|\phi\rangle$ di $\hat{B}$,

⁵Si tenga presente la distinzione tra *accuratezza*, ovvero la vicinanza di una misura al valore vero, e *precisione*, ovvero il numero di cifre significative con cui si ottiene un risultato. La prima pertiene agli errori sistematici, la seconda a quelli casuali.

il quale però, se due operatori non commutano, non è un autostato di $\hat{A}$, e il risultato della seconda misura conseguentemente non sarà determinato. Questo è invece possibile se gli operatori commutano e hanno quindi una base di autovettori comune. Si noti inoltre che questa limitazione è intrinseca e non dipende, come a volte si sente dire, da caratteristiche degli apparati di misura, che possono introdurre incertezze a loro peculiari ma non possono mai violare il limite di Heisenberg.

Il caso più noto è la relazione di indeterminazione tra una coordinata e la corrispondente componente della quantità di moto:

$$\sigma_x \sigma_{p_x} \geq \frac{\hbar}{2}. \quad (7.16)$$

Si trova anche una relazione di indeterminazione simile tra il valore dell'energia in uno stato e il tempo medio τ di permanenza della particella in quello stato:

$$\sigma_E \tau \geq \frac{\hbar}{2}. \quad (7.17)$$

Si tratta però di una somiglianza in parte ingannevole, in quanto il significato è, come si è visto, leggermente diverso, e la deduzione non è quella vista sopra, in quanto il tempo in sé non è un'osservabile e di conseguenza non ha un operatore associato.

7.7 L'oscillatore armonico

Autovalori e autovettori dell'oscillatore armonico si possono ottenere risolvendo l'equazione di Schrödinger corrispondente. Tuttavia esiste un metodo più semplice, elegante e istruttivo che non richiede di risolvere equazioni differenziali ma solo di conoscere le regole di commutazione di opportuni operatori.

L'Hamiltoniana dell'oscillatore armonico unidimensionale è, in corrispondenza con quello classico,

$$\hat{H} = \hat{p}^2/2m + \frac{1}{2}m\omega^2 x^2 = -\hbar^2/2m d^2/dx^2 + \frac{1}{2}m\omega^2 x^2. \quad (7.18)$$

Definiamo

$$\hat{a} = \sqrt{m\omega/2\hbar} x + \sqrt{\hbar/2m\omega} d/dx \quad (7.19)$$

$$\hat{a}^\dagger = \sqrt{m\omega/2\hbar} x - \sqrt{\hbar/2m\omega} d/dx \quad (7.20)$$

che sono l'uno l'aggiunto dell'altro, e il cui commutatore è (*Esercizio*)

$$[\hat{a}, \hat{a}^\dagger] = 1. \quad (7.21)$$

L'Hamiltoniana si scrive come

$$\hat{H} = \frac{\hbar\omega}{2} (\hat{a}\hat{a}^\dagger + \hat{a}^\dagger\hat{a}) = \hbar\omega \left(\hat{a}^\dagger\hat{a} + \frac{1}{2} \right) \quad (7.22)$$

e i suoi commutatori con $\hat{a}$ e $\hat{a}^\dagger$ sono $[\hat{H}, \hat{a}] = -\hbar\omega\hat{a}$ e $[\hat{H}, \hat{a}^\dagger] = \hbar\omega\hat{a}^\dagger$.

Chiamiamo $|m\rangle$ un arbitrario autostato di $\hat{H}$ con energia E_m . Applicando in sequenza $\hat{a}$ e $\hat{H}$ ad $|m\rangle$ si ottiene $\hat{H}\hat{a}|m\rangle = \hat{a}\hat{H}|m\rangle + [\hat{H}, \hat{a}]|m\rangle = E_m|m\rangle - \hbar\omega\hat{a}|m\rangle = (E_m - \hbar\omega)\hat{a}|m\rangle$, ovvero $\hat{a}|m\rangle$ è un altro autostato dell'energia, ma con autovalore inferiore al precedente di $\hbar\omega$, oppure il vettore nullo. Similmente si ottiene un altro autostato, con energia accresciuta di $\hbar\omega$, da $\hat{H}\hat{a}^\dagger|m\rangle = (E_m + \hbar\omega)\hat{a}^\dagger|m\rangle$, se $\hat{a}^\dagger|m\rangle$ non è il vettore nullo. Gli operatori $\hat{a}^\dagger$ e $\hat{a}$ vengono quindi chiamati rispettivamente di *innalzamento* e *abbassamento*, o *creazione* e *distruzione*. Tutti gli autovalori dell'energia sono dunque equispaziati di $\hbar\omega$.

L'applicazione di $\hat{a}$ non può essere ripetuta all'infinito, visto che sia l'operatore cinetico che quello potenziale sono quadrati di operatori autoaggiunti e pertanto gli autovalori dell'energia non possono essere negativi⁶. L'unico modo in cui l'applicazione ripetuta di $\hat{a}$ si può interrompere è quando il vettore di stato $\hat{a}^{\dagger n}|m\rangle$ si annulla. Per stabilire il valore dello stato fondamentale, che indichiamo con $|0\rangle$, consideriamo che non esistendone, per definizione, uno a energia più bassa, deve essere $\hat{a}|0\rangle =$ vettore nullo, quindi $E_0 = \langle 0|\hat{H}|0\rangle = \hbar\omega\langle 0|(\hat{a}^\dagger\hat{a} + \frac{1}{2})|0\rangle = \hbar\omega(\langle 0|\hat{a}^\dagger\hat{a}|0\rangle + \langle 0|\frac{1}{2}|0\rangle) = \frac{1}{2}\hbar\omega$. Pertanto, etichettando gli stati con il numero quantico n , e applicando n volte l'operatore $\hat{a}^\dagger$ allo stato fondamentale $n = 0$, si ottiene

$$E_n = \left(n + \frac{1}{2} \right) \hbar\omega. \quad (7.23)$$

Lo stato fondamentale quindi non ha energia nulla ma ha un' *energia di punto zero* $\frac{1}{2}\hbar\omega$, che è una diretta conseguenza del principio di Heisenberg: avere un'energia nulla significherebbe avere un'energia potenziale nulla, quindi localizzare esattamente l'oscillatore in $x = 0$. L'autostato fondamentale lo otteniamo dalla condizione $\hat{a}\psi = \sqrt{m\omega/2\hbar} x\psi + \sqrt{\hbar/2m\omega} d\psi/dx = 0$, che dà

⁶Se un'osservabile B è il quadrato di A , $\hat{A}$ e $\hat{B}$ hanno gli stessi autovettori. Chiamando α l'autovalore di $\hat{A}$ su di un autostato qualunque $|\psi\rangle$ e β l'autovalore di $\hat{B}$ su $|\psi\rangle$, $\beta = \langle \psi|\hat{B}|\psi\rangle = \langle \psi|\hat{A}^2|\psi\rangle = \langle \hat{A}^\dagger\psi|\hat{A}\psi\rangle = \langle \hat{A}\psi|\hat{A}\psi\rangle = \alpha^2 \geq 0$.

$$\psi_0(x) = A_0 e^{-\frac{m\omega x^2}{2\hbar}}. \quad (7.24)$$

Gli stati eccitati si ottengono applicando ripetutamente $\hat{a}^\dagger = \sqrt{m\omega/2\hbar} x - \sqrt{\hbar/2m\omega} d/dx$ alla 7.24, e si ottengono funzioni del tipo

$$\psi_n(x) = A_n H_n(x) e^{-\frac{m\omega x^2}{2\hbar}}, \quad (7.25)$$

con H_n il polinomio di Hermite di grado n . Si vede subito che:

1. Le autofunzioni sono pari o dispari a seconda della parità di n ;
2. il numero di nodi cresce con l'energia, l'unico stato senza nodi essendo quello fondamentale;
3. l'energia cinetica e quella potenziale coincidono (*Esercizio*), e non sono mai nulle in nessuno stato.

Questi fatti sono conseguenze di considerazioni generali. La prima deriva dal fatto che l'Hamiltoniana è pari. La seconda può essere dimostrata per ogni stato a una particella a partire dalla condizione di ortogonalità. La terza è conseguenza di un teorema più generale detto *teorema del viriale (quantistico)*.

7.8 Il momento angolare

L'operatore corrispondente all'osservabile momento angolare si ottiene tramite gli operatori già visti, e in coordinate cartesiane diventa

$$\hat{\mathbf{L}} = \hat{\mathbf{r}} \times \hat{\mathbf{p}} = -i\hbar \begin{vmatrix} \hat{x} & \hat{y} & \hat{z} \\ x & y & z \\ \partial/\partial x & \partial/\partial y & \partial/\partial z \end{vmatrix}. \quad (7.26)$$

Se invece consideriamo la rotazione lungo uno specifico asse, per esempio l'asse z , di un angolo θ , dalla 7.26 si trova

$$\hat{L}_z = -i\hbar \partial/\partial\theta \quad (7.27)$$

che ricorda da vicino, e non per caso, la 7.10. Gli autostati e autovalori di $\hat{L}_z$ si ottengono risolvendo l'equazione agli autovalori corrispondente, $\hat{L}_z\psi = -i\hbar \partial\psi/\partial\theta = J\psi$:

$$\psi_m(\theta) = \frac{1}{\sqrt{2\pi}} e^{iJ\theta/\hbar} = \frac{1}{\sqrt{2\pi}} e^{im\theta} \quad (7.28)$$

e ricordando che, trattandosi di una rotazione in spazio reale, $\psi(\theta + 2\pi) = \psi(\theta)$, si deve avere $m \in \mathbb{Z}$. Spesso per semplicità si indica direttamente l'intero m come autovalore invece di $J = m\hbar$.

L'osservabile "modulo quadro" di L si ottiene come

$$\hat{L}^2 = \hat{L}_x^2 + \hat{L}_y^2 + \hat{L}_z^2 \propto \hat{\Lambda} \quad (7.29)$$

dove Λ è definito come la parte angolare dell'operatore laplaciano (cfr. D.5). Gli autostati di $\hat{L}^2$ si ottengono quindi dall'equazione agli autovalori $\hat{\Lambda}\psi = \lambda\psi$, e sono le cosiddette *armoniche sferiche*.

Così facendo si trovano anche gli autovalori, che però si possono ottenere anche, in analogia con l'oscillatore armonico, facendo uso solo delle relazioni di commutazione. Tali relazioni si trovano facilmente tramite calcolo diretto:

$$\begin{aligned} [\hat{L}_x, \hat{L}_y] &= i\hbar\hat{L}_z, [\hat{L}_y, \hat{L}_z] = i\hbar\hat{L}_x, [\hat{L}_z, \hat{L}_x] = i\hbar\hat{L}_y \\ [\hat{L}_\alpha, \hat{L}^2] &= 0, \alpha \in \{x, y, z\}. \end{aligned} \quad (7.30)$$

Salta all'occhio che le componenti di $\hat{\mathbf{L}}$ commutano tutte con $\hat{L}^2$, ma non commutano tra loro. In termini di operatori, questo significa che è sempre possibile trovare una base di autostati comune a $\hat{L}^2$ e a una qualunque delle sue componenti, ma che questa base comune non è comune con le altre componenti. Stante il principio di Heisenberg, quindi, è possibile misurare simultaneamente con precisione il modulo del momento angolare e una delle sue componenti, per esempio $\hat{L}_z$, ma le altre rimarranno indeterminate.

Stabilito questo, per trovare gli autovalori di $\hat{L}^2$ e $\hat{L}_z$ si procede in maniera simile a come già visto in 7.7 ovvero, a grandi linee:

- si definiscono operatori di innalzamento e abbassamento $L_\pm = L_x \pm iL_y$;
- si mostra che applicando $L_\pm$ a un autostato di $\hat{L}_z$ si aumenta o diminuisce l'autovalore di $\hbar$;
- ripetute applicazioni di $L_\pm$, unite al fatto che $\lambda \geq 0$, portano a porre condizioni su J e su λ .

Risulta che gli autovalori di $\hat{L}^2$ sono

$$\lambda = \hbar^2 l(l+1) \quad (7.31)$$

con $l \in \mathbb{N}_0$. Anche qui spesso per semplicità si indica direttamente l'intero l come autovalore. Vigge inoltre la limitazione $-l \leq m \leq l$, che è abbastanza intuitiva: una componente di un vettore non può essere maggiore del modulo. Ad ogni l fissato corrispondono quindi $2l + 1$ possibili valori di m .

7.9 Il modello dell'atomo

7.9.1 Atomi idrogenoidi

Avendo ora a disposizione il formalismo del momento angolare possiamo completare il trattamento dell'atomo di idrogeno. L'equazione di Schrödinger completa diventa, ricordando D.5 e moltiplicando per r^2

$$-\frac{\hbar^2}{2m} \left\{ \frac{\partial}{\partial r} \left(r^2 \frac{\partial \psi}{\partial r} \right) + \left[\frac{1}{\sin^2 \phi} \frac{\partial^2 \psi}{\partial \theta^2} + \frac{1}{\sin \phi} \frac{\partial}{\partial \phi} \left(\sin \phi \frac{\partial \psi}{\partial \phi} \right) \right] \right\} = \left(Er^2 + \frac{Ze^2 r}{4\pi\epsilon_0} \right) \psi \quad (7.32)$$

dove $\psi = \psi(r, \theta, \phi)$ e si è introdotta una carica nucleare Ze per generalità (atomi idrogenoidi). L'equazione è evidentemente a variabili separabili e si può risolvere col metodo standard di fattorizzare la funzione d'onda $\psi(r, \theta, \phi) = R(r)\Phi(\theta, \phi)$, inserire questa forma nella 7.32 e dividere per ψ . Si ottiene

$$\begin{aligned} -\frac{\hbar^2}{2m} \frac{\partial}{\partial r} \left(r^2 \frac{\partial R}{\partial r} \right) / R - \left(Er^2 + \frac{Ze^2 r}{4\pi\epsilon_0} \right) &= -\lambda/2m \\ -\frac{\hbar^2}{2m} \left[\frac{1}{\sin^2 \phi} \frac{\partial^2 \Phi}{\partial \theta^2} + \frac{1}{\sin \phi} \frac{\partial}{\partial \phi} \left(\sin \phi \frac{\partial \Phi}{\partial \phi} \right) \right] / \Phi &= \lambda/2m \end{aligned}$$

Si è divisa l'equazione in due in quanto somma di due termini, uno funzione solo di r e l'altro solo degli angoli. Questi termini devono essere uguali per ogni valore di r da una parte e di θ e ϕ dall'altra, che sono però variabili indipendenti, e l'unico modo in cui possono esserlo è che siano separatamente uguali a una costante che abbiamo scritto come $\pm\lambda/2m$. Il valore λ si vede immediatamente essere proprio l'autovalore del quadrato del momento angolare. La parte radiale dell'equazione è la stessa che per le soluzioni a simmetria sferica già viste, ma con un "termine centrifugo", analogo a quello classico, pari a $\lambda/2mr^2$.

Sappiamo già che $m \leq l$, e si può dimostrare che vale anche la disuguaglianza $n > l$. Gli stati con $l = 0$ sono chiamati *stati s*, quelli con $l = 1$ *stati p*, e a seguire *d*, *f*, *g*, *h*, ... Come abbiamo visto gli stati *s*

corrispondono a stati a simmetria sferica, $\psi_n^s = R_n(r)$ e sono convenzionalmente indicati come stati $ns = 1s, 2s, \dots$. Quelli p , oltre ad avere una dipendenza angolare, hanno una parte radiale diversa, conseguentemente alla presenza del termine centrifugo, e possono essere scritti, per esempio, come $\psi_2^{p_x} = xR_2(r)$, $\psi_2^{p_y} = yR_2(r)$, $\psi_2^{p_z} = zR_2(r)$, e sono indicati come $2p_x, 2p_y, 2p_z$. Gli stati d contengono termini quadratici nelle coordinate e via seguendo.

Gli autostati idrogenoidi hanno la particolarità⁷ di essere degeneri in m e l . Ovvero, gli autovalori dipendono solo da n e sono quelli che abbiamo visto prima, $E_n = Ze^2/8\pi\epsilon_0 a_0 n^2$. Abbiamo quindi che la degenerazione k degli stati aumenta con n , e in particolare vale n^2 :

n	$l(s, p, d, f)$	m	k	E_n
1	0	0	1	E_1
2	0, 1	0; -1, 0, 1	4	$E_1/4$
3	0, 1, 2	0; -1, 0, 1; -2, -1, 0, 1, 2	9	$E_1/9$
4	0, 1, 2, 3	0; -1, 0, 1; -2, -1, 0, 1, 2; -3, -2, -1, 0, 1, 2, 3	16	$E_1/16$

7.9.2 Atomi a molti elettroni

La trattazione degli atomi più pesanti dell'idrogeno richiede l'inserimento di termini di interazione tra gli elettroni e non può essere svolta analiticamente. Tuttavia trattando gli elettroni come indipendenti ma sottoposti ciascuno a un *campo medio* definito come il potenziale elettrico dato da una distribuzione di carica ρ uguale a $-e$ moltiplicato per la somma dei moduli quadri delle funzioni d'onda elettroniche, si trovano risultati qualitativamente simili a quelli idrogenoidi. Si può quindi procedere a una costruzione degli atomi con $Z > 1$ immaginando di aggiungere via via funzioni d'onda elettroniche fino ad averne Z .

Il livello fondamentale si dovrebbe ottenere mettendo tutti gli elettroni nello stato a energia più bassa, che per quanto detto fin qui è lo stato con $n = 1$. Questo genera una seria difficoltà: gli atomi così generati non assomigliano affatto a quelli reali. In particolare sarebbero semplicemente oggetti sferici qualitativamente uguali tra loro. Empiricamente in realtà si osserva che l'occupazione di ogni orbitale può essere di uno o due elettroni, o nessuno. Abbiamo quindi un nuovo postulato detto *Principio di esclusione di Pauli*:

⁷Questa particolarità è condivisa con l'oscillatore armonico tridimensionale, e nessun altro sistema, ed è connessa al fatto che in Meccanica Classica campo newtoniano-coulombiano e armonico sono gli unici ad ammettere esclusivamente orbite chiuse (se finite).

- *Uno stato non può essere occupato da più di due elettroni simultaneamente.*

A questo punto possiamo procedere con la costruzione degli atomi a più elettroni. Mettendo due elettroni nello stato a $n = 1$, otteniamo l'elio. L'elettrone successivo, stante Pauli, va in uno stato $n = 2$. Calcoli più accurati che tengono meglio conto dell'interazione elettrone-elettrone mostrano che l'energia dello stato $2s$ è più bassa di quella degli stati $2p$, pertanto il $2s$ si riempie prima. Abbiamo così gli atomi di litio e berillio, la cui struttura elettronica, mettendo i numeri di occupazione come apici, si indicano con $1s^2 2s$ e $1s^2 2s^2$. Aggiungendo elettrone a elettrone otteniamo via via il boro $1s^2 2s^2 2p$, il carbonio $1s^2 2s^2 2p^2$, l'azoto $1s^2 2s^2 2p^3$, l'ossigeno $1s^2 2s^2 2p^4$, il fluoro $1s^2 2s^2 2p^5$ e il neon $1s^2 2s^2 2p^6$. Avendo ora riempito gli stati fino a $n = 2, l = 1$ diciamo di avere una *shell chiusa*. Da una shell chiusa possiamo continuare con il riempimento della shell successiva e via procedendo.

L'importanza delle shell chiuse (completamente riempite) o non chiuse sta nel fatto che si può vedere che un sistema di due (o più) atomi in cui le shell siano o tutte chiuse o tutte vuote ha energia inferiore, ed è quindi più stabile, di un sistema in cui le shell siano parzialmente riempite. Così accostando un atomo di litio e uno di fluoro il sistema sarà energeticamente favorito se un elettrone lascia il litio per trasferirsi nel fluoro. Abbiamo quindi una molecola di LiF (fluoruro di litio), composta da due ioni Li^+ e F^- . Questo corrisponde al concetto chimico di *valenza*, positiva o negativa. La valenza, in altri termini, è quantomeccanicamente — in prima approssimazione, nel mondo reale le cose sono come sempre un po' più complicate, specie all'aumentare del numero di elettroni — il numero di elettroni che devono essere ricevuti o ceduti per avere shell completamente piene o vuote. Infatti il berillio ha valenza 2, il boro 3, il carbonio 4, l'azoto 3, l'ossigeno 2, eccetera.

Siccome la valenza determina le proprietà chimiche di un elemento, e le shell si riempiono nello stesso modo una volta che le shell sottostanti sono piene, abbiamo che *le proprietà chimiche degli elementi sono periodiche*. Elementi con una shell chiusa più un elettrone, come litio, sodio, potassio, rubidio e cesio, hanno proprietà chimiche simili e si dicono appartenere allo stesso *gruppo*, indicato da un numero romano (i metalli alcalini, appena menzionati, sono il gruppo I, carbonio, silicio, germanio sono nel gruppo IV eccetera). Lo stesso quelli con shell chiusa più due elettroni eccetera. Man mano che una shell va riempiendosi quindi le proprietà chimiche dei vari atomi si ripetono: l'alternarsi di proprietà nel *secondo*

periodo Li-Be-B-C-N-O-F-Ne è la stessa, o quantomeno simile, che nel terzo Na-Mg-Al-Si-P-S-Cl-Ar eccetera.

Gli elementi quindi si possono ordinare utilmente in una tavola ordinata per *periodi* orizzontalmente e *gruppi* verticalmente: è la celebre *tavola periodica degli elementi* di Mendeleev.

7.10 Dinamica quantistica

7.10.1 Il propagatore

Fin qui abbiamo visto quella che potrebbe essere chiamata la cinematica quantistica, in quanto non è stato ancora trattato il problema di determinare l'evoluzione temporale dei vettori di stato. Poniamo quindi un postulato apposito:

- *In assenza di misure, l'evoluzione temporale di uno stato è deterministica e unitaria.*

La condizione di unitarietà deriva dall'interpretazione probabilistica e in particolare dal fatto che la norma del vettore deve rimanere 1 (la particella deve essere trovata da qualche parte!). Il determinismo sull'evoluzione del vettore di stato - che non implica determinismo in assoluto, visto che il processo di misura fa collassare il vettore d'onda, e questo processo, dipendendo dal risultato della misura, è stocastico - dice che il vettore a un tempo t è funzione del vettore a tempo 0:

$$|t\rangle = \hat{U}_t|0\rangle. \quad (7.33)$$

Gli operatori $\hat{U}_t$, come detto, devono essere unitari e formano un gruppo a un parametro isomorfo a $(\mathbf{R}, +)$, visto che chiaramente $\hat{U}_{t+t'} = \hat{U}_t\hat{U}_{t'}$. La proprietà di gruppo ci dice anche che $\hat{I} = \hat{U}_0 = \hat{U}_{t+(-t)} = \hat{U}_t\hat{U}_{-t}$, per cui l'evoluzione effettuata all'indietro nel tempo riporta lo stato esattamente alla condizione iniziale. Questo ci dice che l'evoluzione temporale è esattamente invertibile, ovvero che l'evoluzione temporale di un sistema quantistico è invariante per inversione temporale, in esatta corrispondenza con l'evoluzione di un sistema classico. Questa proprietà si perde però in presenza di una misura, in quanto il collasso di una funzione d'onda non permette di ricostruire la funzione d'onda ante collasso da quella post collasso.

Per gruppi del tipo suddetto, ipotizzando come al solito una regolarità matematica sufficiente, vale il *teorema di Stone*: esiste un operatore autoaggiunto $\hat{H}$ tale che

$$\hat{U}_t = e^{-i\hat{H}t/\hbar}. \quad (7.34)$$

Pertanto la dinamica quantistica è descritta tramite un *propagatore* dato da

$$|t\rangle = \hat{U}_t|0\rangle = e^{-i\hat{H}t/\hbar}|0\rangle. \quad (7.35)$$

7.10.2 L'equazione di Schrödinger dipendente dal tempo

Dalla notazione usata si capisce che intendiamo identificare $\hat{H}$ con l'hamiltoniano. Verifichiamolo derivando la 7.35 rispetto al tempo, ottenendo

$$\begin{aligned} \frac{d|t\rangle}{dt} &= \frac{d}{dt}e^{-i\hat{H}t/\hbar}|0\rangle = -i\hat{H}/\hbar e^{-i\hat{H}t/\hbar}|0\rangle = -i\hat{H}|t\rangle/\hbar \\ i\hbar \frac{d|t\rangle}{dt} &= \hat{H}|t\rangle \end{aligned} \quad (7.36)$$

dove si è ipotizzato che $\hat{H}$ non dipenda dal tempo. In rappresentazione x , in una dimensione per semplicità di notazione

$$\hat{H}\psi(x, t) = i\hbar \frac{\partial \psi(x, t)}{\partial t}. \quad (7.37)$$

Avendo ipotizzato un'indipendenza di $\hat{H}$ dal tempo, possiamo separare le variabili nella 7.37, usando il metodo già visto di fattorizzare $\psi(x, t) = \phi(x)\chi(t)$ e ottenendo

$$\chi(t)\hat{H}\phi(x) = i\hbar\phi(x)\frac{\partial \chi(x, t)}{\partial t} \Rightarrow \hat{H}\phi(x)/\phi(x) = i\hbar\frac{\partial \chi(t)}{\partial t}/\chi(t)$$

visto che per ipotesi $\hat{H}$ non contiene il tempo. I due membri dell'equazione devono essere uguali per ogni valore di x e di t , che sono però variabili indipendenti, e l'unico modo in cui possono esserlo è che siano separatamente uguali a una costante E :

$$\begin{aligned} \hat{H}\phi(x) &= E\phi(x) \\ i\hbar \frac{\partial \chi(t)}{\partial t} &= E\chi(t). \end{aligned} \quad (7.38)$$

La prima delle 7.38 è l'equazione di Schrödinger già vista, che chiameremo quindi *equazione di Schrödinger indipendente dal tempo*, mentre la 7.37 è detta *equazione di Schrödinger dipendente dal tempo*. La seconda delle 7.38 ci dice che nel caso di Hamiltoniana indipendente dal tempo la funzione d'onda totale è data dal prodotto dell' autofunzione dell'hamiltoniana $f_E(x)$ con un fattore di fase:

$$\psi(x, t) = f_E(x)e^{-iEt/\hbar} = f_E(x)e^{-i\omega t}. \quad (7.39)$$

Pertanto la funzione che descrive una particella a energia definita che si muove liberamente è un'onda piana monocromatica $e^{i(\mathbf{k}\cdot\mathbf{r}-\omega t)}$, con $\omega = \hbar k^2/2m$. Si noti che questa *relazione di dispersione* è qualitativamente diversa - quadratica - da quella lineare delle onde elettromagnetiche per le quali $\omega = ck$.

7.10.3 Leggi di conservazione

È chiaro che i valori medi in generale evolveranno nel tempo. Derivando abbiamo, nel caso in cui A non dipenda esplicitamente dal tempo,

$$\begin{aligned} \frac{d}{dt} \langle \hat{A} \rangle_t &= \frac{d}{dt} \left(\langle t | \hat{A} | t \rangle \right) = \left(\frac{d}{dt} \langle t | \right) \hat{A} | t \rangle + \langle t | \hat{A} \left(\frac{d}{dt} | t \rangle \right) = \\ &= -\frac{i}{\hbar} \langle \hat{H} t | \hat{A} | t \rangle + \frac{i}{\hbar} \langle t | \hat{H} \hat{A} | t \rangle = \frac{i}{\hbar} \langle t | [\hat{H}, \hat{A}] | t \rangle = \frac{i}{\hbar} \langle [\hat{H}, \hat{A}] \rangle_t. \end{aligned}$$

In particolare ne segue che *un'osservabile è conservata se l'operatore corrispondente commuta con l'hamiltoniano*. Per esempio la quantità di moto si conserva se $\hat{H}$ non dipende dalle coordinate e il momento angolare se $\hat{H}$ non dipende dagli angoli (come visto nel caso atomico).

Il calcolo appena fatto suggerisce una considerazione. Potremmo infatti sostituire il ket $|t\rangle$ con $U_t|0\rangle$ e il bra $\langle t|$ con $\langle 0|U_t^\dagger = \langle 0|U_t^{-1}$, ottenendo $\langle t | \hat{A} | t \rangle = \langle 0 | \hat{U}_t^\dagger \hat{A} \hat{U}_t | 0 \rangle$: in questo modo abbiamo una descrizione delle osservabili tali per cui è il vettore di stato del sistema a essere fisso e sono le osservabili che cambiano col tempo secondo la legge $\hat{A}(t) = \hat{U}_t^\dagger \hat{A} \hat{U}_t$. Questo modo di descrivere la dinamica prende il nome di *rappresentazione di Heisenberg*, essendo di fatto il modo usato originariamente da Heisenberg, mentre quella illustrata in precedenza, in cui gli stati cambiano nel tempo ma gli operatori generalmente no è nota come *rappresentazione di Schrödinger*.

7.11 Spin e sistemi a molte particelle

7.11.1 Lo spin

La discussione fatta in precedenza del momento angolare sembra completa e conclusa. Non è così. Sperimentalmente si vede che a ogni particella è associata una grandezza, detta *spin*, generalmente indicata con σ o S , che si comporta per molti versi come il momento angolare, ma ne differisce per vari aspetti. Cominciamo con le somiglianze.

In primo luogo, lo spin si somma con il momento angolare. Anche nelle circostanze in cui $\mathbf{L}$ dovrebbe conservarsi, questa conservazione in generale non si ha se non si tiene conto del termine addizionale di spin. In altri termini, la grandezza conservata è il vettore $\mathbf{J} = \mathbf{L} + \mathbf{S}$, dove $\mathbf{L}$ e $\mathbf{S}$ sono il momento angolare e lo spin totale (ovvero la somma dei momenti angolari e la somma degli spin di tutte le particelle)

In secondo luogo, le relazioni di commutazione tra gli operatori associati allo spin sono le stesse che per il momento angolare:

$$\begin{aligned} [\hat{S}_x, \hat{S}_y] &= i\hbar\hat{S}_z, [\hat{S}_y, \hat{S}_z] = i\hbar\hat{S}_x, [\hat{S}_z, \hat{S}_x] = i\hbar\hat{S}_y \\ [\hat{S}_\alpha, \hat{S}^2] &= 0, \alpha \in \{x, y, z\}. \end{aligned} \quad (7.40)$$

dove ovviamente

$$\hat{S}^2 = \hat{S}_x^2 + \hat{S}_y^2 + \hat{S}_z^2. \quad (7.41)$$

Avendo le stesse regole di commutazione del momento angolare, lo spin ha anche le stesse relazioni per gli autovalori. In particolare, gli autovalori Σ di $\hat{S}^2$ hanno la stessa forma:

$$\Sigma = \hbar^2 s(s+1) \quad (7.42)$$

e quelli della componente, per esempio, z sono $\hbar s, \hbar(s-1), \dots, -\hbar s$.

In terzo luogo, particelle cariche (ma anche neutre, purché composte da particelle cariche, come il neutrone) dotate di spin esibiscono un momento magnetico, esattamente come i corpi classici composti da particelle cariche in movimento, ovvero che sono percorsi da corrente elettrica.

Queste somiglianze sono sufficienti per considerare lo spin una sorta di *momento angolare intrinseco*. E tuttavia ci sono alcune differenze che rendono difficile la visualizzazione dello spin come una "rotazione della particella su sé stessa" tout-court. A parte il fatto che non è affatto chiaro

che senso possa avere la rotazione di una particella puntiforme su sé stessa, ci sono almeno due ostacoli sostanziali:

- Il modulo s dello spin, per una data particella, è fissato. Non esiste modo di aumentarlo o diminuirlo.
- I valori di s non sono necessariamente interi, ma possono essere semiinteri, e così di conseguenza gli autovalori delle componenti. Siccome il fatto che gli autovalori di $\hat{L}_z$ sono interi deriva direttamente dall'essere conseguenza di rotazioni nello spazio fisico, tridimensionale, ne consegue che lo spin non può essere interpretato come una rotazione nello spazio fisico.

La conclusione è che lo spin è una proprietà intrinseca senza un analogo classico. Lo vediamo anche dal fatto che prendendo il limite formale, già menzionato, per $\hbar \rightarrow 0$ e per autovalori che simultaneamente tendono all'infinito lo spin scompare in quanto, banalmente, gli autovalori non possono andare all'infinito.

I vettori d'onda di un sistema quindi appartengono al prodotto cartesiano dello spazio di Hilbert, per esempio, basato sulle coordinate con lo spazio di Hilbert degli spin, che ha dimensione pari a $2s + 1$. Nel caso importante $s = 1/2$ la dimensione è 2 e gli operatori sono perciò (proporzionali a) matrici 2×2 , dette *matrici di Pauli*:

$$\hat{\sigma}_x = \begin{vmatrix} 0 & 1 \\ 1 & 0 \end{vmatrix}, \quad \hat{\sigma}_y = \begin{vmatrix} 0 & -i \\ i & 0 \end{vmatrix}, \quad \hat{\sigma}_z = \begin{vmatrix} 1 & 0 \\ 0 & -1 \end{vmatrix} \quad (7.43)$$

con

$$\hat{S}_\alpha = \frac{1}{2}\hbar\hat{\sigma}_\alpha, \quad S_\alpha^2 = \frac{1}{4}\hbar^2\hat{\mathbf{I}} \rightarrow \Sigma = \frac{3}{4}\hbar^2\hat{\mathbf{I}}. \quad (7.44)$$

e si è presa come direzione per la misura della componente di $\hat{\mathbf{S}}$ il solito asse z . Gli autovalori di $\hat{S}_z$ sono quindi $\pm\hbar/2$, comunemente indicati come "su" e "giù".

7.11.2 Particelle identiche

Fino a qui la trattazione ha riguardato sistemi a singola particella. Il formalismo non cambia se ne abbiamo più di una, e i mutamenti riguardano la realizzazione specifica. Se abbiamo N particelle, lo spazio di Hilbert da usare è il prodotto cartesiano di quelli delle singole particelle

$\mathcal{H} = \mathcal{H}_1 \otimes \mathcal{H}_2 \dots \otimes \mathcal{H}_N$, che in rappresentazione $\mathbf{r}$ significa (tralasciamo nell'immediato seguito la parte di spin)

$$\psi = \psi(\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N).$$

Il significato di questa ψ è quello ovvio: il modulo quadro della ψ indica la densità di probabilità di trovare la particella 1 in un intorno di $\mathbf{r}_1$ e la particella 2 in un intorno di $\mathbf{r}_2$ e la particella 3 in un intorno di $\mathbf{r}_3$ eccetera. Se si vuole la probabilità incondizionata di trovare la particella 1 in un intorno di $\mathbf{r}_1$ bisogna integrare su tutte le altre variabili: $P(\mathbf{r}_1) = \int \psi(\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N) d\mathbf{r}_2, \dots, d\mathbf{r}_N$. Se le particelle sono indipendenti, vale a dire se non esiste interazione tra di loro, il loro moto produrrà delle probabilità indipendenti, quindi la probabilità totale sarà il prodotto delle probabilità di trovare una specifica particella, e di conseguenza la funzione d'onda sarà il prodotto di N funzioni d'onda, ognuna dipendente da una sola coordinata:

$$\psi(\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N) = \psi(\mathbf{r}_1)\psi(\mathbf{r}_2)\dots\psi(\mathbf{r}_N).$$

Quanto detto presuppone però che si possa dire quale particella sia stata trovata in $\mathbf{r}_1$, quale in $\mathbf{r}_2$ eccetera, vale a dire che si possa distinguere ogni particella da ogni altra. Questo è il caso per oggetti macroscopici, che possono essere etichettati in maniera così leggera da essere ininfluente sul loro comportamento, o per particelle che possono essere seguite con assoluta precisione. Per particelle quantistiche, che non seguono traiettorie definite, questo non è possibile. Bisogna perciò ammettere che le funzioni d'onda di un sistema di più particelle debba avere un modulo quadro che non cambia per una qualunque permutazione delle coordinate. Per esempio, se abbiamo due particelle identiche non interagenti in due stati ϕ_α e ϕ_β , la funzione d'onda totale $\psi(\mathbf{r}_1, \mathbf{r}_2)$ deve essere tale che

$$|\psi(\mathbf{r}_1, \mathbf{r}_2)|^2 = |\psi(\mathbf{r}_2, \mathbf{r}_1)|^2 \Rightarrow |\phi_\alpha(\mathbf{r}_1)\phi_\beta(\mathbf{r}_2)|^2 = |\phi_\alpha(\mathbf{r}_2)\phi_\beta(\mathbf{r}_1)|^2 \quad (7.45)$$

che implica, se vogliamo conservare l'indipendenza delle due particelle, una combinazione o simmetrica o antisimmetrica delle funzioni d'onda:

$$\psi(\mathbf{r}_1, \mathbf{r}_2) = \phi_\alpha(\mathbf{r}_1)\phi_\beta(\mathbf{r}_2) \pm \phi_\alpha(\mathbf{r}_2)\phi_\beta(\mathbf{r}_1). \quad (7.46)$$

Questa conclusione a rigore vale solo per coppie di particelle, mentre per sistemi con più di due particelle sono possibili altre combinazioni oltre alla completa simmetrizzazione o antisimmetrizzazione. Queste combinazioni però non sono mai state osservate, e si conclude che la

funzione d'onda di uno stato deve essere o completamente simmetrica o completamente antisimmetrica per scambio di particelle.

Quanto qui esposto rimane valido anche reintroducendo le variabili di spin. L'unico cambiamento sta nel fatto che la simmetria o asimmetria deve coinvolgere anche queste variabili:

$$\psi(\mathbf{r}_1, \mathbf{r}_2, \sigma_1, \sigma_2) = \phi_\alpha(\mathbf{r}_1, \sigma_1)\phi_\beta(\mathbf{r}_2, \sigma_2) \pm \phi_\alpha(\mathbf{r}_2, \sigma_2)\phi_\beta(\mathbf{r}_1, \sigma_1). \quad (7.47)$$

Pertanto il requisito di antisimmetrizzazione può essere soddisfatto anche avendo parti spaziali identiche ma parti di spin antisimmetriche.

7.11.3 Statistiche quantistiche

I tipi di particelle le cui funzioni d'onda sono antisimmetriche vengono dette *fermioni*, quelli che hanno funzioni d'onda simmetriche *bosoni*. Dalla 7.47 si vede immediatamente che una combinazione antisimmetrica si annulla identicamente se le due funzioni componenti sono identiche, mentre questo non accade nel caso di combinazioni simmetriche. Pertanto *due fermioni non possono avere simultaneamente la stessa funzione d'onda totale*. Possono avere la stessa funzione d'onda spaziale purché quelle di spin siano diverse. Si dice anche che i fermioni seguono la *statistica di Fermi-Dirac*, mentre i bosoni quella di *Bose-Einstein*.

Quali particelle seguono una statistica e quali l'altra? Detto che la trattazione fatta finora si applica sia a particelle elementari sia a corpi composti, purché si possano considerare identici, risulta dall'osservazione⁸ che (*Postulato di connessione spin-statistica*)

- *Le particelle con spin intero sono bosoni, quelle con spin semiintero sono fermioni.*

Per particelle composte vale la somma degli spin: una particella composta da un numero pari di fermioni quindi è un bosone.

Se consideriamo elettroni, quindi con spin 1/2, e fattorizziamo la funzione d'onda di due di essi come $\psi(\mathbf{r}, \sigma) = \phi(\mathbf{r})\chi(\sigma)$, possiamo dire che due elettroni in un atomo possono stare nello stesso orbitale purché abbiano uno spin "su" e l'altro spin "giù". Non possiamo però averne più di due nello stesso stato $\phi(\mathbf{r})$: otteniamo dunque il principio di esclusione di Pauli.

⁸Questo postulato è in realtà deducibile dalla Meccanica Quantistica Relativistica, che qui non viene trattata.

È istruttivo applicare la 7.47 al caso di onde piane, $\phi_\alpha(x) = \frac{1}{\sqrt{2}}e^{i\alpha x}$, con lo stesso spin (che quindi ignoriamo in quanto fattore comune):

$$\psi(x_1, x_2) = e^{i(\alpha x_1 + \beta x_2)} \pm e^{i(\alpha x_2 + \beta x_1)}$$

Posto il centro di massa $\eta = (x_1 + x_2)/2$ e la distanza relativa $\xi = x_1 - x_2$, svolgendo i calcoli si trova, nel caso dei fermioni

$$\psi(x_1, x_2) = ie^{i(\alpha + \beta)\eta} \sin[(\alpha - \beta)\xi/2]$$

e dei bosoni

$$\psi(x_1, x_2) = e^{i(\alpha + \beta)\eta} \cos[(\alpha - \beta)\xi/2].$$

In entrambi i casi, il primo fattore indica che il centro di massa si muove come una particella libera con quantità di moto uguale alla somma delle quantità di moto delle particelle singole, come c'era da aspettarsi. Il secondo fattore, per i fermioni, indica in primo luogo che non si può avere $\alpha = \beta$, come da principio di esclusione. In secondo luogo indica che i fermioni, *anche non interagenti*, tendono a restare lontani ($P(\xi = 0) = 0$), mentre viceversa i bosoni tendono ad avvicinarsi.

7.11.4 Fermioni

Siccome per i fermioni è impossibile occupare stati già occupati da altri fermioni identici, la costruzione dello stato fondamentale procede come già visto per l'atomo. Si calcolano gli autostati del sistema e si riempiono di particelle uno a uno partendo da quello a energia più bassa e via via risalendo.

Consideriamo per semplicità il caso più semplice e importante, quello di fermioni a spin $s = 1/2$. Questa categoria include le particelle che costituiscono la materia ordinaria: elettroni, protoni, neutroni, oltre a svariate particelle esotiche. Per fissare le idee parliamo di elettroni e consideriamoli non interagenti. Questo è in realtà un modello ragionevolmente adeguato alla descrizione degli elettroni nei metalli, o più correttamente di quegli elettroni, detti *di conduzione*, meno legati ai singoli ioni e che pertanto costituiscono normalmente la corrente elettrica. Un modello analogo vale per le stelle di neutroni.

Immaginiamo dunque elettroni liberi di muoversi in una scatola di lunghezza L in ogni direzione cartesiana. Come visto in precedenza, sono descritti dalle onde piane 7.12. Se abbiamo N elettroni, lo stato fondamentale è costituito dagli $N/2$ stati più bassi in energia, interamente occupati.

In una scatola lunga L e di volume $\Omega = L^3$ le componenti del vettore $\mathbf{k}$ sono separate da $\Delta k = 2\pi/L$, e possiamo vedere i vettori $\mathbf{k}$ come un reticolo di punti in un apposito spazio tridimensionale, avente come coordinate k_x, k_y, k_z , detto *spazio reciproco*; la densità di questo reticolo in questo spazio è $1/\Delta k^3 = L^3/8\pi^3 = \Omega/8\pi^3$. Gli stati occupati saranno quelli inferiori a una energia $\epsilon(k_F) = \hbar^2 k_F^2/2m$, dove m è la massa dell'elettrone e k_F è detto vettore d'onda di Fermi, e il loro numero è pari alla densità suddetta moltiplicata per il volume di una sfera di raggio k_F . Questo numero va eguagliato alla metà degli elettroni disponibili:

$$N/2 = \frac{4}{3}\pi k_F^3/\Delta k^3 = \frac{4}{3}k_F^3\Omega/8\pi^3 \Rightarrow k_F = (3\pi^2 N/\Omega)^{1/3}.$$

L'energia del livello occupato più alto, che nel limite termodinamico ($L, N \rightarrow \infty, N/\Omega = \text{costante}$) coincide con quello del livello non occupato più basso, è detta *energia di Fermi*:

$$\epsilon_F = \frac{\hbar^2}{2m} \left(\frac{3\pi^2 N}{\Omega} \right)^{2/3}. \quad (7.48)$$

Si vede quindi che anche in questo caso, come in quello dell'oscillatore armonico, che nemmeno nello stato fondamentale l'energia — in questo caso, interamente cinetica — è nulla. Nel caso dei metalli l'energia di Fermi è dell'ordine di alcuni eV, ovvero qualche centinaio di volte più alta dell'energia di agitazione termica, e siamo quindi giustificati a considerare gli elettroni come se fossero a temperatura praticamente zero.

A temperatura finita alcuni elettroni ricevono abbastanza energia da occupare stati a energia più alta. Si può far vedere, con le tecniche della Meccanica Statistica, che il numero medio di elettroni che occupa uno stato a energia ϵ è dato dalla *distribuzione di Fermi*

$$n(\epsilon) = \frac{1}{e^{(\epsilon-\mu)/k_B T} + 1} \quad (7.49)$$

dove μ è il *potenziale chimico* degli elettroni, che a $T \ll T_f = \epsilon_F/k_B$ tende a ϵ_F . Per basse temperature, e i metalli da questo punto di vista sono sempre a bassa temperatura, il numero di elettroni che si muovono da stati sotto ϵ_F a stati sopra è proporzionale a T (basta studiare la funzione per $T/T_f \ll 1$ per vederlo), e ognuno di questi elettroni guadagna un'energia proporzionale a T . Ne segue che la variazione di energia interna rispetto allo stato fondamentale è proporzionale a T^2 e di conseguenza il calore specifico a volume costante è proporzionale a T , risolvendo così il contrasto col Terzo Principio della Termodinamica.

7.11.5 Bosoni

I bosoni ricadono generalmente in una di due categorie: particelle mediatrici di interazioni, come il fotone; oppure sistemi composti da un numero pari di fermioni. Per tutti vale la *distribuzione di Bose*

$$n(\epsilon) = \frac{1}{e^{(\epsilon-\mu)/k_B T} - 1} \quad (7.50)$$

che come si vede differisce solo per un segno da quella di Fermi. Una prima conseguenza di questo mutamento apparentemente piccolo sta nei calori specifici, che non sono lineari ma cubici.

Conseguenze anche più importanti si hanno per $T \rightarrow 0$, dove vediamo che se $\epsilon > \mu$ il numero di occupazione va giustamente a 0, ma per $\epsilon < \mu$ diventerebbe negativo. Pertanto deve essere $\mu < 0$. Non solo, ma μ diventa 0 sotto una temperatura critica, dipendente dal sistema, non nulla, e a questa temperatura una frazione macroscopica del gas occupa lo stato fondamentale, fenomeno che chiamato *condensazione di Bose-Einstein*.

I condensati di Bose-Einstein hanno caratteristiche molto particolari, che portano a livello macroscopico caratteristiche prettamente quantistiche. L'elio ha una temperatura critica di circa 2 K e sotto questo valore diventa un *superfluido* e manifesta proprietà peculiari, per esempio

- conducibilità termica (praticamente) infinita
- viscosità trascurabile
- se il contenitore contenente elio superfluido è immerso in altro elio superfluido, il fluido all'interno scavalca le pareti del contenitore (effetto fontana).

In alcuni solidi conduttori, non tutti, possono darsi condizioni, troppo complesse per essere qui illustrate, che conducono al formarsi di coppie di elettroni (*coppie di Cooper*), che quindi si comportano come bosoni che a basse temperature formano un superfluido. Questo superfluido viene chiamato *superconduttore* in quanto trasporta corrente senza perdere energia. Ha anche varie altre proprietà, come quella di espellere ogni campo magnetico in cui fosse eventualmente immerso inizialmente (*effetto Meissner*).

Appendice A

I sistemi di unità di misura dell'elettromagnetismo

I sistemi di unità di misura, come è noto, sono frutto di una scelta arbitraria, e questo fatto permette di utilizzarne uno qualunque, a seconda dei gusti, a patto di essere coerenti nel suo utilizzo. Questo fatto ha storicamente portato alla coesistenza di diversi sistemi di unità, coesistenza che in parte permane tuttora.

Consideriamo qui i due sistemi principali, gli unici ormai ad essere utilizzati: il Sistema Internazionale (SI) ed il cgs. Finché si considerano solo grandezze meccaniche, in cui le grandezze fondamentali sono solo tre: Lunghezza, Tempo e Massa, la differenza tra i due sistemi è ridotta. Infatti in essi vengono scelte come unità fondamentali che differiscono tra loro al più solo di potenze di 10: metro/centimetro, kilogrammo/grammo. Il risultato è che i due sistemi definiscono grandezze che differiscono tra loro solamente di potenze di 10, e tutte le equazioni, definizioni, ecc. assumono la stessa forma.

A.1 Il Sistema Internazionale

Le cose cambiano radicalmente quando si passa allo studio dell' elettromagnetismo. In questo caso infatti il SI definisce un'unità di misura addizionale, l' Ampère (A), definito come "quella corrente che, scorrendo in due fili infiniti paralleli posti a distanza di 1 m, genera una forza tra di essi pari a $2 \cdot 10^{-7}$ N/m". Invertendo il ragionamento che porta alla 2.12, si può dire che 1 A è la corrente che, presa la costante μ_0 esattamente pari al valore di $4\pi \cdot 10^{-7}$ m kg/C², rende valide congiuntamente le espressioni

2.1 della forza di Lorentz e 2.7 della legge di Ampère. Così facendo, si è anche definita l'unità di misura del campo magnetico, il Tesla.

Il Coulomb (C), unità di misura della carica, è nel SI un'unità derivata che vale 1 A s. La costante k nella legge di Coulomb $\mathbf{F} = kq_1q_2\hat{\mathbf{r}}/r^2$ va dunque misurata, visto che non c'è più nulla di definibile arbitrariamente nelle altre grandezze che compaiono nell'equazione. Il SI pone, come è noto, $k = 1/4\pi\epsilon_0$. Rimane così fissata anche l'unità di misura del campo elettrico $\mathbf{E} = \mathbf{F}/q$, il Volt/m = kg m s⁻² C⁻¹.

A.2 Il sistema di Gauss

Nell sistema cgs, almeno nella sua versione di gran lunga più comune che prende il nome di Gauss, si evita di definire un'ulteriore unità fondamentale, e si procede invece a definire l'unità di carica, detta *unità elettrostatica* (ues, o all'inglese esu), ponendo $k = 1$ nella legge di Coulomb. L'equazione di Gauss in forma integrale diventa $\oint \mathbf{E} \cdot d\mathbf{S} = 4\pi q$.

Le dimensioni fisiche dell'ues sono g^{1/2} cm^{3/2} s⁻¹; 1 C corrisponde a 3 10⁹ ues. Il fattore 3 che compare in questa ed altre conversioni ha a che fare, come vedremo, con la velocità della luce $c = 2.997925 \cdot 10^8 \approx 3 \cdot 10^8$ m/s; per avere una conversione più precisa si sostituisca dunque al 3 il valore 2.997925.

Per quanto la comparsa di dimensioni con esponente frazionario possa sembrare strana, non c'è nulla di sbagliato, bizzarro, o viceversa molto profondo nella loro presenza. Piuttosto è molto importante notare che non solo le stesse grandezze fisiche sono espresse da numeri diversi nel SI e nel sistema di Gauss, ma sono in generale *misurate con unità diverse*, e *possono avere forme diverse*. Per esempio, la conducibilità elettrica si misura in s⁻¹ — o, se si preferisce, in Hertz — nel sistema di Gauss, e l'equazione 2.12 diventa $F/L = \frac{2I_1I_2}{c^2s}$.

Nel sistema di Gauss l'unità di corrente, lo statAmpère, è derivata ed ha un valore ben definito (1 A corrisponde a 3 10⁹ stA). Nella legge di Ampère viene introdotta una costante, la costante c , avente le dimensioni di una velocità, e l'equazione assume la forma $\oint \mathbf{B} \cdot d\mathbf{l} = \frac{4\pi}{c} I$. In questo modo $\mathbf{B}$ e $\mathbf{E}$ sono misurati nelle stesse unità, lo statV/cm = 3 10⁴ V/m, che nel caso del campo magnetico prende il nome di Gauss, equivalente a 10⁻⁴ T. Anche $\mathbf{D}$ ed $\mathbf{H}$ hanno queste stesse unità di misura.

Se ora si sceglie la seguente espressione per la forza di Lorentz

$$\mathbf{F} = q \left(\mathbf{E} + \frac{\mathbf{v}}{c} \times \mathbf{B} \right)$$

il valore di c risulta non arbitrario, ma determinato (e determinabile) empiricamente (e, fra le altre cose, il valore misurato deve far sì che valga l'equazione 2.12). Come è logico, il sistema di Gauss contiene lo stesso numero — uno — di costanti empiriche del SI (che contiene ϵ_0).

Riassumendo, le equazioni di Maxwell in unità di Gauss appaiono come segue:

$$\begin{array}{ll} \nabla \cdot \mathbf{E} = 4\pi\rho & \oint \mathbf{E} \cdot d\mathbf{S} = 4\pi q \\ \nabla \cdot \mathbf{B} = 0 & \oint \mathbf{B} \cdot d\mathbf{S} = 0 \\ \nabla \times \mathbf{E} = -\frac{1}{c} \frac{\partial \mathbf{B}}{\partial t} & \oint \mathbf{E} \cdot d\mathbf{l} = -\frac{1}{c} \frac{d}{dt} \int \mathbf{B} \cdot d\mathbf{S} \\ \nabla \times \mathbf{B} = \frac{4\pi}{c} \mathbf{j} + \frac{1}{c} \frac{\partial \mathbf{E}}{\partial t} & \oint \mathbf{B} \cdot d\mathbf{l} = \frac{4\pi}{c} I + \frac{1}{c} \frac{d}{dt} \int \mathbf{E} \cdot d\mathbf{S} \end{array}$$

Deducendo (*Esercizio*) l'equazione delle onde da questa forma delle equazioni di Maxwell si vede subito che la costante c risulta essere proprio la velocità della luce.

Infine, conoscendo le equazioni di Maxwell in entrambe le forme, si ricava una ricetta pratica per la conversione di una qualunque equazione dell'elettromagnetismo nel vuoto (non solo le equazioni di Maxwell e di Lorentz, ma tutte, visto che ne discendono) da SI a Gauss e viceversa:

SI		Gauss		Gauss		SI
ϵ_0	$\rightarrow$	$1/4\pi$		E	$\rightarrow$	$\sqrt{4\pi\epsilon_0}E$
μ_0	$\rightarrow$	$4\pi/c^2$		B	$\rightarrow$	$\sqrt{4\pi/\mu_0}B$
B	$\rightarrow$	B/c		q	$\rightarrow$	$q/\sqrt{4\pi\epsilon_0}$
				j	$\rightarrow$	$j/\sqrt{4\pi\epsilon_0}$

Appendice B

L'analisi di Fourier

B.1 La serie di Fourier

Consideriamo l'insieme $\mathcal{F}$ di tutte le funzioni a valori complessi, limitate e con un numero finito di punti di discontinuità, definite sull'intervallo reale $[-T/2, T/2]$ o, equivalentemente, periodiche con periodo T . Come è noto questo insieme può essere strutturato come spazio di Hilbert $\mathcal{H}$ su $\mathbf{C}$ con il prodotto scalare

$$(f, g) = \int_{-T/2}^{T/2} f^*(t)g(t) dt. \quad (\text{B.1})$$

Di conseguenza è sempre possibile sviluppare una qualunque funzione $u(t) \in \mathcal{F}$ su di una base di funzioni ortogonali (che possiamo se necessario sempre supporre ortonormali) $\{\mathbf{v}_l\}$:

$$\mathbf{u} = \sum_l (\mathbf{v}_l, \mathbf{u}) \mathbf{v}_l = \sum_l \int_{-T/2}^{T/2} v_l^*(\tau) u(\tau) d\tau \mathbf{v}_l. \quad (\text{B.2})$$

Una tale base è fornita dall'insieme di funzioni esponenziali

$$v_l(t) = \frac{1}{\sqrt{T}} e^{i2\pi lt/T}, l \in \mathbf{Z} \quad (\text{B.3})$$

che in effetti appartengono ad $\mathcal{F}$ e costituiscono un insieme ortonormale:

$$(\mathbf{v}_l, \mathbf{v}_{l'}) = \frac{1}{T} \int_{-T/2}^{T/2} e^{i2\pi(l'-l)t/T} dt = \delta_{ll'}. \quad (\text{B.4})$$

Evidentemente lo stesso vale per le funzioni seno e coseno; se l'insieme di funzioni di base viene definito come

$$\begin{aligned}
v_0(t) &= \sqrt{\frac{1}{T}} \\
v_l(t) &= \sqrt{\frac{2}{T}} \cos(2\pi lt/T), \quad l \in \mathbf{N} \\
w_l(t) &= \sqrt{\frac{2}{T}} \sin(2\pi lt/T), \quad l \in \mathbf{N}
\end{aligned} \tag{B.5}$$

anch'esso risulta ortonormale.

Si noti che tutte le funzioni risultano periodiche di periodo T o, se si preferisce, aventi frequenze $\omega_l = 2\pi l/T$ che sono tutte multiple intere di una frequenza $\omega = 2\pi/T$ detta *fondamentale*. Le altre frequenze sono dette *armoniche*.

Si dimostra che sia l'insieme B.3 che l'insieme B.5 formano un insieme completo, dunque ciascuno è una base dello spazio di Hilbert $\mathcal{H}$.

Ogni funzione $u \in \mathcal{F}$ può dunque essere sviluppata in una serie della forma

$$u(t) = \sum_{l=-\infty}^{\infty} c_l e^{i\omega_l t} = \frac{1}{T} \sum_{l=-\infty}^{\infty} \int_{-T/2}^{T/2} e^{-i2\pi l\tau/T} u(\tau) d\tau e^{i2\pi lt/T} \tag{B.6}$$

oppure

$$\begin{aligned}
u(t) &= a_0 + \sum_{l=1}^{\infty} [a_l \cos(\omega_l t) + b_l \sin(\omega_l t)] = \frac{1}{T} \int_{-T/2}^{T/2} u(\tau) d\tau \\
&+ \frac{2}{T} \sum_{l=1}^{\infty} \left[\int_{-T/2}^{T/2} \cos(2\pi l\tau/T) u(\tau) d\tau \cos(2\pi lt/T) \right. \\
&\left. + \int_{-T/2}^{T/2} \sin(2\pi l\tau/T) u(\tau) d\tau \sin(2\pi lt/T) \right].
\end{aligned} \tag{B.7}$$

Il coefficiente a_0 o c_0 denota il valor medio della funzione u .

Se la funzione u è a valori reali i coefficienti a_l e b_l sono reali, ma i c_l non lo sono necessariamente [lo sono, per esempio, se (*Esercizio*) la funzione u è pari]. In generale vale solo la relazione $c_{-l}^* = c_l$. Lo sviluppo converge ad $u(t)$ in tutti i punti di continuità ed al valor medio dei limiti destro e sinistro (che, per le ipotesi fatte, esistono sempre) nei punti di discontinuità. Evidentemente, per funzioni pari si avrà $b_l = 0$ e per funzioni dispari $a_l = 0 \forall l$.

Ognuno di questi sviluppi è *unico* ed è noto come *sviluppo in serie di Fourier*. Tale sviluppo definisce dunque una relazione biunivoca tra l'insieme $\mathcal{F}$ ed un opportuno sottoinsieme dell'insieme delle successioni.

B.2 La trasformata di Fourier

Possiamo considerare il caso di una funzione non periodica definita su tutto l'asse reale prendendo il limite $T \rightarrow \infty$ delle formule precedenti. In questo limite le frequenze ω_l tendono ad una distribuzione continua, infatti la loro distanza è $\Delta\omega = \omega_{l+1} - \omega_l = 2\pi/T \rightarrow 0$. Passiamo quindi al limite del continuo identificando $1/T$ con $\Delta\omega/2\pi \rightarrow d\omega/2\pi$. Prendendo la base di esponenziali complessi B.3 si dimostra che lo sviluppo di Fourier diventa

$$u(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{u}(\omega) e^{i\omega t} d\omega = \frac{1}{2\pi} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-i\omega\tau} u(\tau) d\tau e^{i\omega t} d\omega \quad (\text{B.8})$$

dove si è introdotta la *trasformata di Fourier*

$$\hat{u}(\omega) = \int_{-\infty}^{\infty} e^{-i\omega\tau} u(\tau) d\tau. \quad (\text{B.9})$$

La funzione originale $u(t)$ viene anche detta *antitrasformata* di $\hat{u}(\omega)$ e, come si vede, l'antitrasformazione di Fourier è del tutto analoga alla trasformazione, cambiando solo il segno all'esponente, la variabile di integrazione ed il coefficiente (anche quest'ultimo avrebbe potuto essere scelto simmetrico). È frequente l'uso di omettere il "cappelletto" sulla trasformata e di distinguerla dall' antitrasformata solo tramite il simbolo della variabile indipendente.

La trasformata di Fourier può essere vista, oltre che come uno sviluppo su di una base più che numerabile di funzioni, come una funzione avente come dominio un sottoinsieme dell'insieme delle funzioni da $\mathbf{R}$ in $\mathbf{C}$ e come codominio un sottoinsieme dell'insieme delle funzioni da $\mathbf{C}$ in $\mathbf{R}$. Si tratta di una funzione *lineare*: se $f(t) = \alpha g(t) + \beta h(t)$, allora $\hat{f}(\omega) = \alpha \hat{g}(\omega) + \beta \hat{h}(\omega)$.

Si dimostra che se si chiama $\mathcal{G}$ lo spazio di Hilbert costruito, nella maniera usuale, sull'insieme delle funzioni infinitamente derivabili $u(t) : [-\infty, \infty] \rightarrow \mathbf{C}$ tali che per ognuna di esse esiste un insieme $\{C_{pq}(u)\}$ tale che $|t^p u^{(q)}| < C_{pq}(u)$ (funzioni "rapidamente decrescenti"), allora la *trasformata di Fourier* è un isomorfismo da $\mathcal{G}$ in $\mathcal{G}$.

Valgono le proprietà, di facile verifica (*Esercizio*)

$$\begin{aligned}
\widehat{\frac{d^n u}{dt^n}} &= (i\omega)^n \widehat{u}(\omega) \\
\widehat{u}(\omega) &= \widehat{u}^*(-\omega) \\
f(t) &= \int_{-\infty}^{\infty} g(\tau)h(t-\tau) d\tau = (g * h)(t) \Rightarrow \widehat{f}(\omega) = \widehat{g}(\omega)\widehat{h}(\omega).
\end{aligned} \tag{B.10}$$

Quest'ultima è nota come *teorema della convoluzione*.

L'estensione della trasformata di Fourier a funzioni di più variabili è immediata. Se $\mathbf{r}$ è un vettore a tre dimensioni, allora

$$\widehat{u}(\mathbf{k}) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-i\mathbf{k}\cdot\mathbf{r}} u(\mathbf{r}) dx dy dz \tag{B.11}$$

$$u(\mathbf{r}) = \frac{1}{(2\pi)^3} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{i\mathbf{k}\cdot\mathbf{r}} \widehat{u}(\mathbf{k}) dk_x dk_y dk_z. \tag{B.12}$$

B.3 Applicazioni

La linearità della trasformata di Fourier è alla base delle sue principali applicazioni. In virtù di questa proprietà è possibile risolvere un'equazione differenziale (o integrale) lineare trasformando secondo Fourier l'equazione stessa per poi lavorare separatamente su ogni componente ω . Al termine si ricombinano linearmente le soluzioni trovate, funzioni di ω , tramite l'antitrasformata, ottenendo così la soluzione generale dell'equazione.

B.3.1 Il circuito RLC

Un esempio semplice è dato dalla soluzione dell'equazione 3.8 del circuito *RLC*. Possiamo scrivere l'equazione come

$$\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-i\omega\tau} \frac{dV}{d\tau} d\tau e^{i\omega t} d\omega = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-i\omega\tau} \left(L \frac{d^2 I}{dt^2} + R \frac{dI}{dt} + \frac{I}{C} \right) d\tau e^{i\omega t} d\omega$$

ovvero, sfruttando le proprietà della trasformata,

$$\int_{-\infty}^{\infty} i\omega \widehat{V}(\omega) e^{i\omega t} d\omega = \int_{-\infty}^{\infty} \left(-\omega^2 L + i\omega R + \frac{1}{C} \right) \widehat{I}(\omega) e^{i\omega t} d\omega$$

e notando che gli esponenziali $e^{i\omega t}$ sono linearmente indipendenti possiamo eguagliare gli integrandi:

$$i\omega\hat{V}(\omega) = \left(-\omega^2 L + i\omega R + \frac{1}{C}\right) \hat{I}(\omega) \Rightarrow \hat{I}(\omega) = \frac{\hat{V}(\omega)}{R + i[\omega L - 1/(\omega C)]}. \quad (\text{B.13})$$

Di solito i passaggi intermedi vengono sottintesi e si scrive direttamente la B.13. Come si vede, ci si è ricondotti da un'equazione differenziale ad una algebrica. Questa equazione dà il valore dell'ampiezza di ogni componente di Fourier della corrente $\hat{I}(\omega)$ in funzione dell'ampiezza della componente di Fourier della tensione applicata $\hat{V}(\omega)$. Altrimenti detto, fornisce la risposta in corrente del circuito per una tensione con una sola frequenza ω e di ampiezza $\hat{V}(\omega)$; per la linearità dell'equazione la risposta in corrente per una tensione variabile in maniera arbitraria si ottiene semplicemente sommando tutte le risposte a tutte le frequenze, vale a dire integrando la B.13:

$$I(t) = \int_{-\infty}^{\infty} \hat{I}(\omega) e^{i\omega t} d\omega = \int_{-\infty}^{\infty} \frac{\hat{V}(\omega) e^{i\omega t}}{R + i[\omega L - 1/(\omega C)]} d\omega.$$

L'integrale, data la forma esplicita di $V(t) \Leftrightarrow V(\omega)$, si calcola o esplicitamente o facendo ricorso ad opportune tavole. Nel caso semplice $V(t) = V_0 \cos(\omega' t) = V_0(e^{i\omega' t} + e^{-i\omega' t})/2$ l'integrale si riduce alla somma di due termini, quelli con $\omega = \pm\omega'$, per cui, ricordando le definizioni di ϕ , X e Z :

$$\begin{aligned} I(t) &= \frac{V_0}{2} \left\{ \frac{e^{i\omega' t}}{R + i[\omega' L - 1/(\omega' C)]} + \frac{e^{-i\omega' t}}{R - i[\omega' L - 1/(\omega' C)]} \right\} \\ &= \frac{V_0}{2} \left\{ \frac{R - iX}{R^2 + X^2} e^{i\omega' t} + \frac{R + iX}{R^2 + X^2} e^{-i\omega' t} \right\} = \frac{V_0}{Z} \left\{ \frac{R}{Z} \cos \omega' t + \frac{X}{Z} \sin \omega' t \right\} \\ &= \frac{V_0}{Z} \{ \cos \omega' t \cos \phi + \sin \omega' t \sin \phi \} = \frac{V_0}{Z} \cos(\omega' t - \phi) \end{aligned}$$

coincidente con la 3.10. Naturalmente, avendo considerato fin dal principio funzioni periodiche troviamo soltanto le soluzioni a regime e non i transitori.

B.3.2 I potenziali ritardati

Un esempio più complicato, ma importante, è quello delle equazioni per i potenziali 3.25. Prendendo per semplicità quella per il potenziale scalare,

e trasformando con Fourier rispetto alla variabile temporale, l'equazione diventa

$$\nabla^2 \phi(\mathbf{r}, \omega) + \frac{\omega^2}{c^2} \phi(\mathbf{r}, \omega) = -\rho(\mathbf{r}, \omega)/\epsilon_0. \quad (\text{B.14})$$

Immaginiamo ora di avere una carica puntiforme di valore q nell'origine. La densità di carica è allora tale che $\rho(\mathbf{r})$ sia nulla ovunque tranne che in un intorno molto piccolo (infinitesimo) attorno a $\mathbf{r} = 0$ e che $\int \rho d\Omega = q$, ovvero è descritta da una δ (si veda la discussione nell'appendice C).

La soluzione $G(\mathbf{r}, \omega)$ dell'equazione B.14 in questo caso particolare, a simmetria sferica, dev'essere anch'essa a simmetria sferica e si trova agevolmente in spazio di Fourier:

$$G(k, \omega) = \frac{q}{\epsilon_0 (k^2 - \omega^2/c^2)} \quad (\text{B.15})$$

Ma il caso dipendente dal tempo differisce da quello statico per la presenza di un fattore $\omega \neq 0$; vediamo dunque che la differenza fra i due casi si riduce alla sostituzione di k^2 con $k^2 - \omega^2/c^2 = (k - \omega/c)(k + \omega/c)$. Dato che, nelle trasformate, k corrisponde ad r ed ω a t , possiamo perciò aspettarci che la trasformata inversa della B.15 dia un risultato che differisce dal caso statico per la sostituzione di r con $r \pm ct$. Si verifica (*Esercizio*) che è proprio così. D'altro canto $G(r, \omega = 0)$ deve coincidere con la soluzione dell'eq. 1.16, per cui si deve avere (*Esercizio*: verificarlo)

$$G(r, 0) = \frac{q}{4\pi\epsilon_0 r} \Leftarrow \frac{q}{\epsilon_0 k^2}$$

e di conseguenza

$$G(r, t) = \frac{q}{4\pi\epsilon_0 (r \pm ct)} \Leftarrow \frac{q}{\epsilon_0 (k^2 - \omega^2/c^2)}.$$

Nel caso generale, dunque, se sommiamo tutti i contributi alla $\rho(\mathbf{r}, t)$ delle cariche puntiformi, in completa analogia con il paragrafo 3.6.2, otteniamo un risultato che è proprio quello espresso dalla 3.26.

Appendice C

La “funzione” δ di Dirac

In molti casi risulta necessario, o almeno comodo, introdurre una funzione che rappresenti la distribuzione di densità di una qualche proprietà (massa, carica, probabilità, ...) di una particella puntiforme localizzata in un punto x_0 . Tale funzione dovrebbe avere la proprietà di essere nulla per $x \neq x_0$ ma di integrare ad un valore finito σ su tutto lo spazio, anzi in ogni intervallo che includa x_0 . In una dimensione, questo significherebbe $\int_{-\infty}^{\infty} \rho(x) dx = \sigma$ o, posto $\rho(x) = \sigma \delta(x)$,

$$\int_{-\infty}^{\infty} \delta(x - x_0) dx = 1 \quad (\text{C.1})$$

La primitiva di questa funzione si potrebbe ottenere integrando da $-\infty$ a t ed è, a meno di costanti, la funzione segno:

$$\int_{-\infty}^t \delta(\tau - t_0) d\tau = \frac{1}{2} [\text{segno}(t - t_0) + 1] \quad (\text{C.2})$$

e viceversa la δ sarebbe la derivata della funzione segno (che però, essendo discontinua, non ha derivata ovunque). Un uso di questa funzione in meccanica classica, dove l'argomento come suggerito dalla notazione è il tempo, sta nella rappresentazione delle cosiddette forze impulsive, ovvero forze che agiscono solo per un tempo infinitesimo ma hanno un impulso $I = \Delta p = \int F dt$ finito. In questo caso $F = I \delta$ e ha origine da potenziali discontinui, per esempio durante l'urto di sfere dure.

Dalla C.1 si vede che dovrebbe essere

$$\int_{-\infty}^{\infty} \delta(x - x_0) f(x) dx = f(x_0) \quad (\text{C.3})$$

per ogni $f(x)$ integrabile. Infatti potremmo suddividere l'asse x in tre parti, un intervallino Δ attorno a x_0 e le due semirette rimanenti, per cui

$$\begin{aligned}\int_{-\infty}^{\infty} \delta(x - x_0) f(x) dx &= \int_{\Delta} \delta(x - x_0) f(x) dx = \lim_{\Delta \rightarrow 0} \int_{\Delta} \delta(x - x_0) f(x) dx \\ &= \lim_{\Delta \rightarrow 0} f(x_0) \int_{\Delta} \delta(x - x_0) dx = f(x_0) \int_{-\infty}^{\infty} \delta(x - x_0) dx = f(x_0).\end{aligned}$$

È evidente che nessuna funzione ordinaria può soddisfare le condizioni richieste, visto che una funzione identicamente 0 ovunque tranne che su un insieme di misura nulla integra a 0. Possiamo però rendere rigorosi i calcoli fatti su questa base vedendoli come l'applicazione di un particolare tipo di funzionale lineare, detto *distribuzione*, definito su di uno spazio di funzioni e a valori in $\mathbf{R}$ o in $\mathbf{C}$. In questo caso il funzionale associa ad ogni $f(x)$ il suo valore in $f(x_0)$. Le distribuzioni sono oggetto di una branca dell'Analisi Funzionale detta appunto Teoria delle Distribuzioni, e si rimanda ai testi specifici per un approfondimento. In generale, si possono rendere rigorosi tutti i ragionamenti che coinvolgono una δ vedendola come una distribuzione nel senso accennato: cioè come un modo per estrarre un particolare valore da una funzione. In questo modo si guadagna in rigore perdendo però in immediatezza e intuizione fisica.

Un modo più elementare ma meno rigoroso di vedere la "funzione" δ è di considerarla come il limite di un insieme di funzioni. Per esempio possiamo definire una famiglia di funzioni a scalino:

$$\delta_{\alpha}(x) = \begin{cases} 0, & x < x_0 - \alpha/2 \\ 1/\alpha, & x_0 - \alpha/2 < x < x_0 + \alpha/2 \\ 0, & x > x_0 + \alpha/2 \end{cases}$$

(potremmo scegliere anche p.es. delle funzioni gaussiane $\delta_{\alpha}(x) = \sqrt{\frac{\pi}{\alpha}} e^{-x^2/2\alpha^2}$), ognuna delle quali integra a 1, i cui membri diventano sempre più stretti e alti nel limite di $\alpha \rightarrow 0$. Il limite $\alpha \rightarrow 0$ di questa famiglia lo chiamiamo $\delta(x)$. Pertanto, visto che

$$\lim_{\alpha \rightarrow 0} \int_{-\infty}^{\infty} \delta_{\alpha}(x - x_0) f(x) dx = \lim_{\alpha \rightarrow 0} \frac{1}{\alpha} \int_{x_0 - \alpha/2}^{x_0 + \alpha/2} f(x) dx = f(x_0),$$

si ottiene la C.3 tramite una disinvolta (e a rigore ingiustificata) inversione di limiti: quello su α e l'integrale.

Quanto detto si generalizza facilmente a una qualunque dimensionalità. Dal punto di vista fisico, le dimensioni della δ sono, vista la C.1, il reciproco

del suo argomento: in una dimensione il reciproco di una lunghezza, in tre dimensioni il reciproco di un volume eccetera.

Infine, è immediato vedere che la trasformata di Fourier di una δ è una costante; viceversa, la trasformata di una costante è una δ . Questo fornisce un'altra possibile rappresentazione della δ :

$$\delta(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{ikx} dk.$$

Si vede immediatamente che questo corrisponde alla definizione intuitiva di δ : per $x = 0$ l'integrale diverge, mentre per ogni altro valore di x si possono raggruppare i valori di k in coppie $(k, k' = k + \pi/x)$. Per queste coppie $e^{ikx} + e^{ik'x} = 0$ e la somma integrale pertanto è nulla.

Appendice D

Relazioni vettoriali

D.1 Definizioni e relazioni algebriche

$$\mathbf{u} \cdot \mathbf{v} = u_x v_x + u_y v_y + u_z v_z \quad (\text{D.1})$$

$$\mathbf{u} \times \mathbf{v} = \begin{vmatrix} \hat{\mathbf{x}} & \hat{\mathbf{y}} & \hat{\mathbf{z}} \\ u_x & u_y & u_z \\ v_x & v_y & v_z \end{vmatrix} \quad (\text{D.2})$$

$$\mathbf{u} \cdot (\mathbf{v} \times \mathbf{w}) = \mathbf{w} \cdot (\mathbf{u} \times \mathbf{v}) = \mathbf{v} \cdot (\mathbf{w} \times \mathbf{u}) \quad (\text{D.3})$$

$$\hat{\mathbf{x}} \times \hat{\mathbf{y}} = \hat{\mathbf{z}}$$

$$\hat{\mathbf{y}} \times \hat{\mathbf{z}} = \hat{\mathbf{x}}$$

$$\hat{\mathbf{z}} \times \hat{\mathbf{x}} = \hat{\mathbf{y}}$$

$$\nabla f = \hat{\mathbf{x}} \frac{\partial f}{\partial x} + \hat{\mathbf{y}} \frac{\partial f}{\partial y} + \hat{\mathbf{z}} \frac{\partial f}{\partial z} \quad (\text{D.4})$$

$$\nabla^2 f(x, y, z) = \nabla \cdot \nabla f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} + \frac{\partial^2 f}{\partial z^2} \quad (\text{D.5})$$

$$\nabla^2 f(r, \theta, \phi) = R(r; f) + \frac{1}{r^2} \Lambda(\theta, \phi; f) \quad (\text{D.6})$$

$$\text{con } R(r; f) = \frac{1}{r^2} \frac{\partial}{\partial r} \left[r^2 \frac{\partial f}{\partial r} \right] = \frac{1}{r} \frac{\partial^2}{\partial r^2} [r f]$$

$$\text{e } \Lambda(\theta, \phi; f) = \frac{1}{\sin^2 \phi} \frac{\partial^2 f}{\partial \theta^2} + \frac{1}{\sin \phi} \frac{\partial}{\partial \phi} \left[\sin \phi \frac{\partial f}{\partial \phi} \right]$$

$$\nabla^2 f(r, \phi, z) = R(r; f) + \frac{1}{r^2} \frac{\partial^2 f}{\partial \phi^2} + \frac{\partial^2 f}{\partial z^2} \quad (\text{D.7})$$

$$\nabla \cdot \mathbf{v} = \frac{\partial v_x}{\partial x} + \frac{\partial v_y}{\partial y} + \frac{\partial v_z}{\partial z} \quad (\text{D.8})$$

$$\nabla^2 \mathbf{v} = \nabla^2 v_x \hat{\mathbf{x}} + \nabla^2 v_y \hat{\mathbf{y}} + \nabla^2 v_z \hat{\mathbf{z}} \quad (\text{D.9})$$

$$\nabla \times \mathbf{v} = \begin{vmatrix} \hat{\mathbf{x}} & \hat{\mathbf{y}} & \hat{\mathbf{z}} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ v_x & v_y & v_z \end{vmatrix} \quad (\text{D.10})$$

D.2 Relazioni differenziali

$$(\mathbf{u} \cdot \nabla) \mathbf{v} = u_x \frac{\partial \mathbf{v}}{\partial x} + u_y \frac{\partial \mathbf{v}}{\partial y} + u_z \frac{\partial \mathbf{v}}{\partial z} \quad (\text{D.11})$$

$$\nabla \cdot \mathbf{u} \times \mathbf{v} = \mathbf{v} \cdot \nabla \times \mathbf{u} - \mathbf{u} \cdot \nabla \times \mathbf{v} \quad (\text{D.12})$$

$$\nabla \times \nabla \times \mathbf{v} = \nabla (\nabla \cdot \mathbf{v}) - \nabla^2 \mathbf{v} \quad (\text{D.13})$$

$$(\nabla \times \mathbf{u}) \times \mathbf{u} = \mathbf{u} \cdot \nabla \mathbf{u} - \frac{1}{2} \nabla u^2 \quad (\text{D.14})$$

$$\nabla \cdot \nabla \times \mathbf{v} = 0 \quad (\text{D.15})$$

$$\nabla \times \nabla f = 0 \quad (\text{D.16})$$

D.3 Relazioni integrali

$$\oint_{\Sigma} \mathbf{v} \cdot d\mathbf{S} = \int_{\Omega} \nabla \cdot \mathbf{v} d\Omega \quad (\text{D.17})$$

$$\int_{\Sigma} \nabla \times \mathbf{v} \cdot d\mathbf{S} = \oint_{\Gamma} \mathbf{v} \cdot d\mathbf{l} \quad (\text{D.18})$$

dove nella D.17 Ω è il volume incluso da Σ e nella D.18 Σ è la superficie avente come bordo la linea Γ .

Bibliografia

- [1] J. Walker, *Fondamenti di Fisica*, CEA, ottava edizione, 2023. Versione rivista del classico testo di D. Halliday e R. Resnick, abbastanza elementare (non usa le forme differenziali per le equazioni dell'elettromagnetismo), ricco di esempi ed esercizi. Sostituibile con qualunque testo di Fisica Generale per i primi due anni di Fisica o Ingegneria.
- [2] R.P. Feynman, R.B. Leighton e M. Sands, *La fisica di Feynman*, Zanichelli. Un libro eccellente, ma non un libro di testo: da leggere *dopo* avere studiato gli altri, per capirli più a fondo.
- [3] H. Goldstein, *Meccanica Classica*, Zanichelli. Rimane la presentazione più chiara a me nota della meccanica analitica. Include una trattazione degli elementi della meccanica analitica relativistica.
- [4] D.J. Griffiths, *Introduction to electrodynamics*, Prentice Hall. È il testo che per la parte di elettromagnetismo più si avvicina a questo corso, sia per tipo di argomenti trattati che per livello di approfondimento. Scritto in stile molto chiaro e discorsivo.
- [5] G.D. Jackson, *Elettrodinamica classica*, Zanichelli. Il testo standard per l'elettrodinamica, utilizzabile anche per la Relatività, molto completo ed aggiornato. Strana la scelta di utilizzare, nell'ultima edizione, il SI per i primi capitoli e quello di Gauss per gli ultimi.
- [6] P. Panofski e M. Phillips, *Elettricità e magnetismo*, CEA. Un altro testo classico, ora fuori commercio in Italia ma disponibile a prezzo contenuto in inglese (Dover). Ottima la parte di relatività.
- [7] Bo Thidé, *Electromagnetic Field Theory*, reperibile gratuitamente in rete su vari siti. Abbastanza conciso, contiene però tutti gli argomenti e le equazioni più importanti.
- [8] E. Purcell, *La Fisica di Berkeley - vol. 2*, Zanichelli. Un testo meno avanzato dei precedenti ma molto buono. Usa un approccio diverso

da quello qui seguito (la relatività è data per nota). Sfortunatamente scritto nel sistema di Gauss.

- [9] R. Resnick, *Introduzione alla Relatività Ristretta*, CEA. Introduzione chiara ed estesa, a livello non troppo avanzato. Attenzione: fa largo uso della “massa relativistica”, concetto che ritengo poco utile se non deleterio.
- [10] E.F. Taylor e J.A. Wheeler, *Fisica dello spazio-tempo*, Zanichelli (mi è nota direttamente però solo l'edizione inglese). È un testo anticonvenzionale — non sono molti i libri di Fisica che includono i Klingon nell'indice analitico — focalizzato sugli aspetti fisici e con il minimo di matematica necessaria. Si tenga presente che nel testo si parla di “sistema in caduta libera” intendendo “sistema inerziale”.
- [11] A.I. Miller, *Albert Einstein's Special Theory of Relativity*, Addison-Wesley. Ampia illustrazione della genesi della Teoria della Relatività, nel contesto scientifico (sia teorico che sperimentale) dell'epoca.
- [12] Callen, *Thermodynamics and an introduction to thermostatistics*,. Una presentazione assiomatica, ma senza eccedere in rigore, della termodinamica. L'approccio adottato, abbastanza inusuale, rende più chiari svariati punti problematici della termodinamica.
- [13] L.D. Landau e E.M. Lifšits, *Corso di Fisica Teorica; vol. 2: Teoria dei Campi, vol. 3: Meccanica Quantistica — teoria non relativistica*, Editori Riuniti. Valgono le stesse avvertenze che per il Feynman. Più di consultazione che di studio.
- [14] C. Kittel, *Introduzione alla Fisica dello Stato Solido*, Bollati Boringhieri. In un paio di capitoli discute il magnetismo nella materia sforzandosi di non usare troppo la Meccanica Quantistica.
- [15] F.W. Byron e R.W. Fuller, *Mathematics of classical and quantum physics*, Dover 1992. Il titolo dice tutto. Copre dagli spazi vettoriali alla teoria delle funzioni analitiche e oltre, con un occhio particolare per le applicazioni fisiche.