

Università di Pisa
Corso di Laurea Magistrale in Fisica



Appunti di Struttura della Materia 2

Giovanni Moruzzi

Anno accademico 2016/2017

Versione del 18 maggio 2017

Premessa

Tutti i Corsi di Studi in Fisica delle università italiane hanno a statuto un corso chiamato *Struttura della Materia*, a volte suddiviso in due corsi semestrali. Ma una rapida ricerca su internet mostra che da un'università all'altra il programma del corso varia moltissimo, tanto che spesso programmi svolti in università diverse hanno sovrapposizione piccolissima o addirittura nulla. Estendendo poi la ricerca anche a università straniere si scopre invece che, fuori dall'Italia, e in particolare nei paesi anglosassoni, praticamente non esistono corsi con nome equivalente.

Questo perché con “struttura della materia”, o “fisica della materia”, in Italia si intende un'area vastissima della fisica, che va dalla fisica a scala atomica in su, fermandosi, senza comprenderla, alla scala astronomica. Sono quindi escluse, per esempio, la fisica nucleare e la fisica delle particelle elementari a un estremo, e l'astrofisica all'altro. Ma quasi tutto il resto è compreso nel concetto di struttura della materia: la fisica atomica, la fisica molecolare, la fisica dei plasmi, la fisica delle varie forme di materia condensata . . . Non è ovviamente possibile trattare tutti questi argomenti in un singolo corso. Mentre c'è sostanziale accordo sul fatto che, se non tutti, almeno gran parte di questi argomenti debba essere affrontata prima di conseguire un serio titolo di studio in fisica, resta una certa arbitrarietà su come ripartirli tra vari possibili corsi. La ripartizione viene effettuata ovunque tra vari corsi specializzati, con confini più o meno ben definiti tra l'uno e l'altro. A questi, in Italia, si aggiunge un corso “generico” chiamato “Struttura della Materia”. Questo spiega sia le differenze tra i programmi di università italiane diverse, dovute a diverse ripartizioni degli argomenti, sia la legittimità della scelta dei paesi anglosassoni di ripartire tutto solo tra i corsi specializzati. Come è legittima la scelta di avere anche un corso generico. Questo spiega anche perché, purtroppo, il programma di Struttura della materia comprende argomenti non sempre molto legati tra loro, e perché c'è qualche sovrapposizione con altri corsi.

Questi appunti coprono interamente il programma da me svolto nel corso di *Struttura della Materia 2* del Corso di Laurea Magistrale in Fisica dell'Università di Pisa (*Struttura della Materia 1* appartiene invece al Corso di Laurea Triennale). Nella scelta degli argomenti trattati ho seguito in buona parte la tradizione che si è consolidata negli ultimi decenni presso il nostro Corso di Studi, con alcune omissioni e aggiunte determinate sia da suggerimenti di colleghi dell'area “strutturistica” del nostro Dipartimento che da considerazioni mie personali.

Sicuramente alcune parti di questi appunti sono migliorabili, e inevitabilmente sono presenti qua e là errori di battitura. Sarò quindi grato agli studenti che suggeriranno correzioni o miglioramenti. Per forza di cose la trattazione dei vari argomenti non è sempre esauriente, e gli studenti sono invitati a consultare anche le fonti citate in bibliografia, in particolare per gli argomenti di loro maggiore interesse.

In fine, per chi fosse interessato, questi appunti sono stati scritti in $\text{\LaTeX}$ su computer che girano sotto Linux, e le figure sono state realizzate con il programma Xfig. Per certe figure, i files Xfig

sono stati generati da programmi in C/C++ compilati da g++. Tutto il software usato è di pubblico dominio.

Pisa, 12 febbraio 2015

Giovanni Moruzzi

Indice

Premessa	i
1 Fisica e informazione	1
1.1 Informazione disponibile e informazione mancante	1
1.2 Informazione mancante per un numero finito di scelte equiprobabili	2
1.3 Scelta tra possibilità con probabilità note	4
1.4 Informazione mancante per una scelta nel continuo	5
2 Meccanica statistica	9
2.1 Statistica classica	9
2.1.1 Stato di un sistema classico e grandezze fisiche	9
2.1.2 Evoluzione temporale	10
2.1.3 Trattamento statistico	10
2.1.4 Rappresentazioni di Schrödinger e di Heisenberg	11
2.2 Statistica quantistica	13
2.2.1 Stato di un sistema quantistico e valori di aspettazione delle grandezze fisiche	13
2.2.2 Evoluzione temporale	14
2.2.3 Trattamento statistico e matrice densità	14
3 Scelta della probabilità degli stati	17
3.1 Primo postulato: riproduzione dell'informazione disponibile	17
3.2 Secondo postulato: massimo condizionato dell'informazione mancante	18
3.3 Moltiplicatori di Lagrange	19
3.4 Informazione disponibile in statistica classica	20
3.5 Particelle identiche in fisica classica	22
3.6 Informazione disponibile in statistica quantistica	24
3.7 Particelle identiche in meccanica quantistica	25
3.8 Funzione di partizione	25
4 Informazione mancante e entropia	27
4.1 Costanti del moto	27
4.2 Gas perfetto e entropia	28
4.3 Teorema di equipartizione dell'energia	31
5 Insieme macrocanonico	33
5.1 Introduzione	33

5.2	Caso classico	33
5.2.1	Numero non noto, o variabile, di particelle identiche	33
5.2.2	Entropia nell'insieme macrocanonico	34
5.2.3	Gas perfetto monoatomico classico nell'insieme macrocanonico	36
5.2.4	Ampiezza delle fluttuazioni in statistica classica	38
5.3	Caso quantistico	39
5.3.1	Gas perfetto quantistico, impostazione del problema	39
5.3.2	Potenziale chimico	42
5.3.3	Numeri di occupazione	42
5.3.4	Fluttuazioni dei numeri di occupazione	43
6	Random walk	47
6.1	Introduzione	47
6.2	Impostazione matematica in una dimensione	47
6.3	Valor medio della posizione	49
6.4	Dispersione	50
6.5	Distribuzione di probabilità per grandi N	51
6.6	Limite continuo	53
6.7	Muro riflettente e muro assorbente	54
6.8	Momenti di una variabile aleatoria	56
7	Processi stocastici	59
7.1	Introduzione	59
7.2	Densità spettrale di una funzione stocastica	62
7.3	Autocorrelazione di una funzione stocastica	63
7.4	L'equazione di Fokker-Planck	66
7.4.1	Equazione di drift	68
7.4.2	Equazione di diffusione	69
7.5	Moto browniano	70
7.5.1	Distribuzione delle velocità	71
7.5.2	Diffusione delle particelle	74
7.6	Shot noise	75
7.7	Rumore $1/f$	77
7.8	Fotoconteggi	78
7.9	Teorema di Nyquist	80
8	Fenomeni di interferenza	83
8.1	Formalismo dei campi complessi	83
8.2	Beam-splitter	85
8.3	Interferometro di Mach-Zehnder	86
8.4	Interferometro di Michelson	89
8.5	Spettroscopia a trasformata di Fourier	91
8.6	Forma di riga	95
8.7	Esperimento di Michelson e Morley	97
8.8	Interferometro di Young e coerenza spaziale	98

8.9	Volume di coerenza	100
8.10	Interferometro stellare di Michelson	102
8.11	Interferometro stellare di Hanbury-Brown e Twiss	103
8.11.1	Spiegazione classica	103
8.11.2	Spiegazione quantistica	105
9	Principi di funzionamento del laser	109
9.1	Coefficienti di Einstein e formula di Planck	109
9.2	Propagazione di un fascio di radiazione in un mezzo materiale	112
9.3	Equazioni di bilancio per l'effetto laser	114
9.3.1	Inversione di popolazione	114
9.3.2	Equazioni di bilancio per i fotoni e fattore di qualità	115
9.3.3	Equazioni di bilancio per le densità di popolazione dei livelli	117
9.3.4	Altri possibili schemi di pompaggio	119
9.4	Onde parassiali	121
9.4.1	Equazione di Helmholtz e approssimazione parassiale	121
9.4.2	Fasci gaussiani	122
9.5	Stabilità della cavità risonante	126
9.6	Instaurazione del regime oscillatorio	128
9.7	Mode-locking	130
9.8	Laser a semiconduttore	133
9.9	Laser a buca di potenziale e laser a cascata quantica	136
10	Interazione radiazione-materia	141
10.1	Perturbazioni dipendenti dal tempo	141
10.1.1	Introduzione	141
10.1.2	Probabilità di transizione	141
10.1.3	Perturbazioni periodiche	144
10.1.4	Transizioni in uno spettro continuo	147
10.1.5	Regola d'oro di Fermi	148
10.2	Perturbazioni periodiche al secondo ordine	150
10.3	Transizioni di dipolo elettrico	153
10.4	Assorbimento, emissione stimolata e spontanea	155
10.5	Interazione con campo forte	158
11	Matrice densità	165
11.1	Matrice densità, popolazioni e coerenze	165
11.2	Modello di Feynman-Vernon-Hellwarth	168
11.3	Rilassamento e equazioni di Bloch ottiche	171
11.4	Equazioni di rate e forme di riga	175
11.5	Rate equations e bilancio dettagliato	176
11.6	Accoppiamento atomo-bagno termico	177
11.7	Allargamento omogeneo	179
11.8	Allargamento disomogeneo: effetto Doppler	180
11.9	Regole di selezione: considerazioni di simmetria	181

11.10	Regole di selezione: transizioni di dipolo elettrico	183
12	Superconduttività	187
12.1	Un richiamo di magnetismo classico	187
12.2	Lavoro di magnetizzazione	189
12.2.1	Lavoro “meccanico” sul campione	189
12.2.2	Lavoro fatto dal generatore	190
12.3	Un richiamo di termodinamica	192
12.4	Aspetti sperimentali della superconduttività	193
12.5	Campo magnetico critico	195
12.6	Diamagnetismo perfetto	196
12.7	Corrente nei superconduttori, trattazione classica	200
12.8	Corrente nei superconduttori, trattazione quantistica	203
12.8.1	Equazione di continuità per la funzione d’onda	203
12.8.2	Densità di corrente in un superconduttore	204
12.8.3	Quantizzazione del flusso attraverso un anello superconduttore	207
12.9	Coppie di Cooper	208
13	Raffreddamento magnetico	211
13.1	Introduzione	211
13.2	Energia	211
13.3	Entropia e raffreddamento magnetico	213
14	Risonanza magnetica	217
14.1	Descrizione classica	217
14.2	Descrizione quantistica	218
14.3	Campo magnetico rotante	221
14.4	Campo oscillante	223
14.5	Rilassamento e equazioni di Bloch	224
14.6	Equazioni di Bloch a campi oscillanti deboli	225
14.7	Misura del tempo di rilassamento trasversale T_2	226
14.8	Spettroscopia NMR	229
14.9	Spettroscopia EPR	230
Appendici		
A	Derivazione alternativa dell’equazione di Fokker-Planck	231
B	Teorema centrale del limite	235
C	Progressione geometrica	239
D	Emissione spontanea e seconda quantizzazione	241
E	Carica in campo elettromagnetico	245
Bibliografia		247

Capitolo 1

Fisica e informazione

1.1 Informazione disponibile e informazione mancante

Secondo la meccanica, sia classica che quantistica, se conosciamo lo stato in cui si trova un sistema fisico costituito da N_p particelle a un certo istante t_0 , e ne conosciamo l'hamiltoniana, siamo in grado di calcolarne la successiva evoluzione temporale. O almeno lo siamo in linea di principio. Nel caso di un sistema che possa essere trattato secondo le leggi della meccanica classica dovremo risolvere, per esempio, le equazioni di Hamilton (2.3). In meccanica quantistica dovremo invece risolvere l'equazione di Schrödinger. Notiamo che nell'ambito di questo corso, per quel che riguarda la meccanica quantistica, non ci interesserà il problema di se, e come, la funzione d'onda collassi quando viene effettuata una misura.

Ricordiamo che in meccanica classica lo stato di un sistema è completamente determinato dalla conoscenza delle coordinate generalizzate indipendenti $\{q_1^{(i)}, \dots, q_{f_i}^{(i)}\}$ e dei corrispondenti momenti coniugati $\{p_1^{(i)}, \dots, p_{f_i}^{(i)}\}$, qui le parentesi graffe denotano gli insiemi. L'indice i , con $1 \leq i \leq N_p$, corre su tutte le particelle del sistema ed f_i è il numero dei gradi di libertà della i -esima particella. L'unione delle $q_j^{(i)}$ e delle $p_j^{(i)}$ costituisce l'insieme delle *coordinate canoniche*, che identifica univocamente il punto rappresentativo dello stato del nostro sistema nello spazio delle fasi Γ . Se le particelle hanno tutte lo stesso numero di gradi di libertà f , cioè se $f_i = f$ per ogni i , lo spazio delle fasi avrà $2N_p f$ dimensioni. Nel caso quantistico lo stato del sistema all'istante t_0 è completamente determinato dalla conoscenza della sua funzione d'onda

$$\psi(q_1^{(1)}, \dots, q_{f_1}^{(1)}, q_1^{(2)}, \dots, q_{f_2}^{(2)}, \dots, q_1^{(N_p)}, \dots, q_{f_{N_p}}^{(N_p)}, t_0),$$

cui corrisponde un vettore $|\psi\rangle$ nello spazio di Hilbert degli stati del sistema.

Rimangono però due problemi assolutamente non trascurabili: 1) conoscere lo stato del sistema e la sua hamiltoniana, 2) risolvere le equazioni. Se il sistema è macroscopico N_p è dell'ordine del numero di Avogadro ($N_A = 6.022 \times 10^{23}$), e questo esclude ogni possibilità pratica di determinare il punto esatto dello spazio delle fasi che lo rappresenta nel caso classico, o la sua funzione d'onda nel caso quantistico. Per non parlare della scrittura dell'hamiltoniana e della soluzione delle equazioni. Dobbiamo così lavorare con un'informazione incompleta, insufficiente per un trattamento "dinamico" del problema. Quello che possiamo fare è trattare il sistema, e la sua evoluzione temporale, utilizzando invece che parametri microscopici, come le coordinate e i momenti delle singole particelle, parametri macroscopici come, per esempio, pressione, volume, temperatura, posizione del baricentro, impulso

totale ... Abbiamo un *difetto di informazione*, o un' *informazione mancante*, corrispondente alla quasi totalità della quantità di informazione necessaria per una completa trattazione dinamica del problema. Per impostare matematicamente il problema da affrontare dobbiamo cominciare quantificando il concetto di informazione mancante.

1.2 Informazione mancante per un numero finito di scelte equiprobabili

Per semplicità cominciamo considerando un sistema che possa trovarsi in un numero discreto e finito W di stati. Un esempio quantistico può essere la componente S_z dello spin $\mathbf{S}$ di una particella, in questo caso avremmo $W = 2S + 1$. Facciamo l'ipotesi di non sapere in quale stato si trovi il sistema, ma di sapere solo che ha la stessa probabilità $1/W$ di trovarsi in uno qualunque degli stati possibili. Quanta informazione ci manca per conoscere lo stato del sistema? Per analogia, la stessa informazione che ci mancherebbe se volessimo trovare una pallina nascosta in una di W scatole chiuse, esternamente tutte uguali. Per semplificare i calcoli e i ragionamenti iniziali, ci conviene partire dal caso particolare in cui W è una potenza di 2, cioè $W = 2^n$.

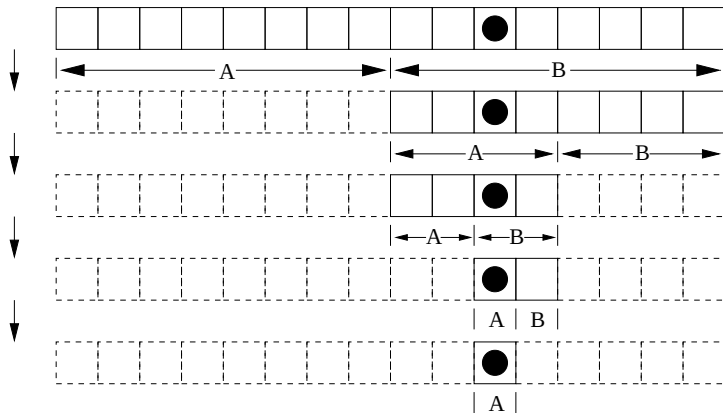


Figura 1.1 Informazione mancante per la localizzazione di una pallina nascosta in una di $W = 16 = 2^4$ scatole.

numero i ?” ci permetterebbe di arrivare alla soluzione in un colpo solo nel caso più fortunato. Ma ci porterebbe a fare $W - 1$ domande nel caso meno fortunato ($W - 1$ perché, dopo esserci sentiti rispondere $W - 1$ volte *no*, non abbiamo bisogno di fare un'ulteriore domanda: la pallina può solo trovarsi nell'unica scatola che resta). Il numero medio di domande $\bar{N}_D$ cui ci porta questa strategia vale quindi

$$\bar{N}_D = \frac{1}{W-1} \sum_{k=1}^{W-1} k = \frac{W}{2}.$$

Si può fare di meglio. La migliore strategia è cominciare dividendo le scatole in due insiemi di $W/2$ scatole l'uno, A e B , e poi chiedere, come prima domanda, “la pallina si trova nell'insieme A ?”. Sia che la risposta sia *sì*, o che sia *no*, abbiamo dimezzato il numero delle scatole in cui la pallina può trovarsi. Cioè abbiamo diminuito l'informazione mancante. Iterando il procedimento arriviamo a trovare la pallina con $n = \log_2 W$ domande. Notate che, se il numero di scatole raddoppia, il numero di domande da fare aumenta solo di una. Data la corrispondenza tra una risposta del tipo *sì/no* e un bit

Definiamo quantitativamente l'informazione mancante I come il minimo numero di domande, a cui è permesso rispondere solo con un *sì* o con un *no*, che dobbiamo fare per arrivare a individuare in quale scatola si trovi la pallina. Per il momento escludiamo quindi la domanda diretta “in quale scatola si trova la pallina?”, che però considereremo tra poco. Dobbiamo scegliere una strategia per organizzare le nostre domande in modo da minimizzarne il numero. Per esempio, una domanda del tipo “la pallina si trova nella scatola numero i ?”

di informazione (o uno dei due stati di un circuito elettronico bistabile tipo *flip-flop*), l'informazione mancante I del problema, misurata in bit, è

$$I = n = \log_2 W. \quad (1.1)$$

Torniamo adesso al problema di quanta informazione ci viene data rispondendo alla domanda diretta “*in quale scatola si trova la pallina?*”. Per comodità etichettiamo le scatole con un numero compreso tra 0 e $W - 1$. Se la risposta ci viene data comunicandoci l'*etichetta* della scatola contenente la pallina in notazione binaria dovranno essere usati $n = \log_2 W$ bit, cioè 4 bit nel caso delle 16 scatole di Fig. 1.1. La combinazione 0000 corrisponderà alla scatola numero 0, 0001 alla numero 1, 0010 alla numero 2, fino alla combinazione 1111 per la scatola numero 15. Quindi $\log_2 W$ bit è anche la quantità di informazione che ci viene data rispondendo alla domanda diretta “*in quale scatola si trova la pallina?*”, in accordo con la (1.1). Questa definizione dell'informazione mancante ha le proprietà seguenti:

1. E' una funzione crescente di W : maggiore è il numero dei casi possibili, maggiore è l'informazione mancante.
2. Consideriamo il caso di due palline, nascoste la prima in una di W_1 scatole identiche, e la seconda in una di altre W_2 scatole, sempre identiche, con W_1 e W_2 entrambi potenze di 2. Le due palline possono essere nascoste in $W_1 W_2$ modi equiprobabili, e l'informazione mancante è $I = \log_2 W_1 W_2$. Questa informazione ci può essere fornita in due mandate: una prima I_1 per localizzare la prima pallina, poi I_2 per localizzare, indipendentemente, la seconda. Deve quindi essere $I = I_1 + I_2$, in perfetto accordo con la relazione

$$\log_2 W_1 W_2 = \log_2 W_1 + \log_2 W_2. \quad (1.2)$$

Nel caso generale il numero di scatole W sarà un numero naturale qualunque, non necessariamente una potenza di 2, quindi $\log_2 W$ non sarà necessariamente intero. Però, se ripetiamo più volte l'esperimento, ovviamente troveremo sempre la pallina dopo un numero intero N_D di domande. Ma nel corso della sequenza delle domande/risposte relative al singolo esperimento capiterà talvolta di trovarci di fronte ad un numero W_r dispari di scatole residue. Per la domanda successiva saremo costretti a dividere W_r in una parte contenente $(W_r + 1)/2$ scatole e una $(W_r - 1)/2$. A seconda della ripartizione, N_D potrà variare da esperimento a esperimento, valendo a volte $N_D = \lfloor \log_2 W \rfloor$, a volte $N_D = \lceil \log_2 W \rceil$. Qui con $\lfloor x \rfloor$ abbiamo indicato il massimo intero minore di x (*parte intera*, o funzione *floor*), e con $\lceil x \rceil$ il minimo intero maggiore di x (*parte intera superiore*, o funzione *ceiling*). Nel caso di scelta tra un numero qualunque W di casi equiprobabili, possiamo estendere il concetto di informazione mancante I definendola come il *numero medio* di domande da fare per individuare la scatola che contiene la pallina: $I = \overline{N_D}$. Così I non sarà più necessariamente un numero intero di bit. Dato che $\lfloor \log_2 W \rfloor \leq \overline{N_D} \leq \lceil \log_2 W \rceil$, viene istintivo estendere la (1.1) considerandola valida per W numero naturale qualunque. Questa però non è ancora una dimostrazione.

La dimostrazione si può fare mediando su un numero molto grande N_E di esperimenti equivalenti e indipendenti l'uno dall'altro. In ogni esperimento la pallina è nascosta in una di W scatole identiche. Il numero complessivo di risultati possibili, per l'insieme degli esperimenti, è W^{N_E} . Se W , e quindi W^{N_E} , fosse una potenza di 2, il numero complessivo di domande necessarie per trovare tutte le N_E

palline sarebbe $N_D^{(\text{tot})} = \log_2(W^{N_E}) = N_E \log_2 W$. Essendo invece W un numero naturale qualunque, per quanto visto sopra possiamo dire solo che

$$\lfloor N_E \log_2 W \rfloor \leq N_D^{(\text{tot})} \leq \lceil N_E \log_2 W \rceil. \quad (1.3)$$

D'altra parte valgono le disuguaglianze

$$N_E \log_2 W - 1 \leq \lfloor N_E \log_2 W \rfloor \quad \text{e} \quad \lceil N_E \log_2 W \rceil \leq N_E \log_2 W + 1, \quad (1.4)$$

che, combinate con la (1.3), ci danno le disuguaglianze

$$N_E \log_2 W - 1 \leq N_D^{(\text{tot})} \leq N_E \log_2 W + 1. \quad (1.5)$$

Sopra abbiamo definito l'informazione mancante per il singolo esperimento come

$$I = \overline{N_D} = \lim_{N_E \rightarrow \infty} \frac{N_D^{(\text{tot})}}{N_E}, \quad (1.6)$$

quindi, dividendo la (1.5) per N_E e facendo il limite $N_E \rightarrow \infty$, otteniamo finalmente

$$I = \log_2 W. \quad (1.7)$$

1.3 Scelta tra possibilità con probabilità note

Supponiamo adesso di avere ancora una pallina nascosta in una di W scatole, ma di avere un po' di informazione in più. Ogni scatola ha una probabilità diversa P_i di contenere la pallina, e noi conosciamo queste probabilità. Ovviamente dovrà essere

$$\sum_{i=1}^W P_i = 1 \quad \text{e} \quad P_i \geq 0 \quad \forall i. \quad (1.8)$$

Prima di procedere ci conviene riscrivere l'equazione (1.7) nella forma

$$I = \log_2 W = r_2 \ln W, \quad (1.9)$$

dove $\ln$ indica il logaritmo naturale, e abbiamo posto $r_2 = \log_2 e = 1/\ln 2$.

Come nel paragrafo precedente, pensiamo di fare un grande numero N_E di esperimenti identici. La legge dei grandi numeri ci dice che, se N_E tende all'infinito, il numero degli esperimenti in

Figura 1.2 Ogni colonna corrisponde a uno di N_E esperimenti, in ognuno dei quali una pallina è nascosta in una di W scatole, quindi ogni colonna contiene una pallina. La riga i corrisponde alla i -esima scatola, che ha probabilità P_i di contenere la pallina. Al limite di grandi N_E la riga i contiene $P_i N_E$ palline.

cui troveremo la pallina nella scatola i -esima tende a $N_E P_i$. Ma la legge dei grandi numeri non ci dice in che ordine si presentano i risultati degli N_E esperimenti. Quali sono gli $N_E P_i$ esperimenti in cui la pallina si trova nella scatola numero i ? I risultati degli esperimenti si possono presentare in qualunque ordine con uguale probabilità. Il numero di modi in cui gli esperimenti possono essere ordinati

vale $N_E! / \prod_{i=1}^W (N_E P_i)!$, poiché, per ogni i , tutti gli esperimenti in cui la pallina si trova nella scatola i sono equivalenti. Una volta specificato l'ordine dei risultati degli esperimenti, non ci manca più informazione. L'informazione mancante complessiva $I^{(\text{tot})}$, relativa a tutti gli N_E esperimenti, è così

$$I^{(\text{tot})} = r_2 \ln \frac{N_E!}{\prod_{i=1}^W (N_E P_i)!} = r_2 \left[\ln N_E! - \sum_{i=1}^W \ln(N_E P_i)! \right] \quad (1.10)$$

Volendo applicare la legge dei grandi numeri, $I^{(\text{tot})}$ ci interessa solo al limite $N_E \rightarrow \infty$, quindi possiamo usare l'approssimazione di Stirling

$$\ln N! \simeq N \ln N - N \quad (1.11)$$

per ottenere

$$I^{(\text{tot})} \simeq r_2 \left[N_E \ln N_E - N_E - \sum_{i=1}^W N_E P_i (\ln N_E + \ln P_i) + \sum_{i=1}^W N_E P_i \right] = -N_E r_2 \sum_{i=1}^W P_i \ln P_i. \quad (1.12)$$

L'informazione mancante complessiva è così proporzionale al numero N_E degli esperimenti, ed essendo gli esperimenti tutti equivalenti, l'informazione mancante per ogni esperimento è

$$I = \lim_{N_E \rightarrow \infty} \frac{1}{N_E} I^{(\text{tot})} = -r_2 \sum_{i=1}^W P_i \ln P_i. \quad (1.13)$$

Se tutte le scatole hanno la stessa probabilità di contenere la pallina abbiamo $P_i = 1/W$ per ogni i , e l'equazione (1.13) coincide con la (1.9).

Passiamo adesso a un numero di scatole W che tende all'infinito. Se tutte le scatole continuano ad avere la stessa probabilità di contenere la pallina, l'informazione mancante sarà

$$I = \lim_{W \rightarrow \infty} r_2 \ln W = +\infty. \quad (1.14)$$

Cioè, l'informazione che ci manca per scegliere tra infinite possibilità equiprobabili è infinita, come ci potevamo aspettare. Ma se le probabilità di trovare la pallina è diversa da scatola a scatola, l'equazione (1.13) ci dà per l'informazione mancante

$$I = -r_2 \sum_{i=1}^{\infty} P_i \ln P_i, \quad (1.15)$$

che, a seconda dei valori delle delle probabilità P_i , può essere tranquillamente finita.

1.4 Informazione mancante per una scelta nel continuo

In questo caso il nostro compito non è più equivalente ad individuare una pallina nascosta in una scatola, ma ad individuare la posizione x di un punto nell'intervallo $a \leq x \leq b$. Supponiamo che esista, e di conoscere, una funzione continua di distribuzione della probabilità $P(x)$ tale che la probabilità che il punto si trovi tra x_1 e x_2 , con $a \leq x_1 \leq x_2 \leq b$ sia data da

$$W(x_1, x_2) = \int_{x_1}^{x_2} P(x) dx, \quad (1.16)$$

con

$$P(x) \geq 0 \quad \forall x \in [a, b], \quad (1.17)$$

e

$$\int_a^b P(x) dx = 1. \quad (1.18)$$

Per calcolare l'informazione mancante in questo caso possiamo partire dai risultati precedenti per il caso discreto, e fare un passaggio al limite. Dividiamo il segmento $[a, b]$ in W celle uguali, e calcoliamo quanta informazione ci manca per sapere in quale cella si trova il punto cercato. La cella i -esima avrà estremi

$$a_i(W) = a + (i-1) \frac{b-a}{W} \quad \text{e} \quad b_i(W) = a + i \frac{b-a}{W}, \quad \text{con } a_{i+1} = b_i \text{ se } 1 \leq i \leq W-1, \quad (1.19)$$

e la probabilità che questa cella contenga la posizione x cercata sarà

$$\mathcal{P}_i(W) = \int_{a_i(W)}^{b_i(W)} P(x) dx. \quad (1.20)$$

L'informazione che ci manca per determinare la cella in cui si trova il nostro punto vale così

$$\begin{aligned} I(W) &= -r_2 \sum_{i=1}^W \mathcal{P}_i(W) \ln \mathcal{P}_i(W) \\ &= -r_2 \sum_{i=1}^W \left[\int_{a_i(W)}^{b_i(W)} P(x) dx \right] \\ &\quad \cdot \ln \left[\int_{a_i(W)}^{b_i(W)} P(x) dx \right]. \end{aligned} \quad (1.21)$$

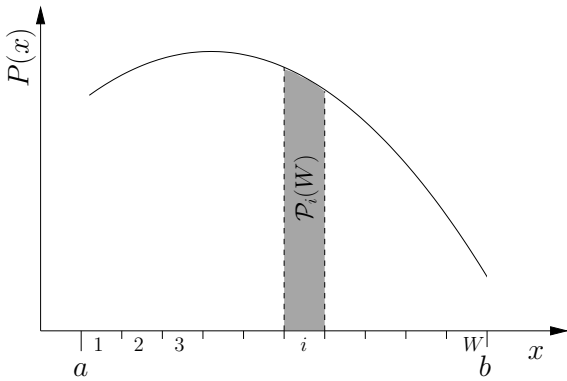


Figura 1.3 Scelta nel continuo. Cerchiamo un punto di coordinata x , con $a \leq x \leq b$. La curva $P(x)$ rappresenta la densità di probabilità. Dividiamo l'intervallo $[a, b]$ in W intervallini uguali. L'area $\mathcal{P}_i(W)$ corrisponde alla probabilità di trovare x nello i -esimo intervallino.

Per passare al caso continuo dobbiamo far tendere W all'infinito e, corrispondentemente, l'ampiezza delle singole celle a zero. La probabilità di trovare x nella cella i -esima diventa

$$\mathcal{P}_i(W) = \int_{a_i(W)}^{b_i(W)} P(x) dx = \bar{P}_i \frac{b-a}{W}, \quad (1.22)$$

dove $\bar{P}_i$ è un valore che $P(x)$ assume in un punto all'interno dell' i -esima cella. Sostituendo la (1.22) nella (1.21) otteniamo

$$I(W) = -r_2 \sum_{i=1}^W \bar{P}_i \frac{b-a}{W} \ln \left(\bar{P}_i \frac{b-a}{W} \right) = -r_2 \sum_{i=1}^W \frac{b-a}{W} \bar{P}_i \ln \bar{P}_i - r_2 \ln \frac{b-a}{W}. \quad (1.23)$$

L'ultimo termine a destra si ottiene ricordando che

$$\sum_{i=1}^W \frac{b-a}{W} \bar{P}_i = 1.$$

Facendo separatamente i limiti per $W \rightarrow \infty$ dei due addendi dell'ultimo membro della (1.23) otteniamo

$$\lim_{W \rightarrow \infty} -r_2 \sum_{i=1}^W \frac{b-a}{W} \bar{P}_i \ln \bar{P}_i = -r_2 \int_a^b P(x) \ln P(x) dx \quad \text{e} \quad \lim_{W \rightarrow \infty} -r_2 \ln \frac{b-a}{W} = +\infty. \quad (1.24)$$

Così, l'informazione mancante per individuare un punto in un intervallo continuo è infinita. Ma il termine che diverge nel secondo limite della (1.24) è del tutto indipendente dalla particolare distribuzione di probabilità $P(x)$. In queste condizioni la differenza di informazione mancante tra due distribuzioni di probabilità diverse $P_1(x)$ e $P_2(x)$ conserva un significato e vale

$$\Delta I = I_1 - I_2 = -r_2 \int_a^b P_1(x) \ln P_1(x) dx + r_2 \int_a^b P_2(x) \ln P_2(x) dx$$

Possiamo quindi *rinormalizzare* l'informazione mancante per la scelta nel continuo sottraendone la costante infinita

$$\lim_{W \rightarrow \infty} -r_2 \ln \frac{b-a}{W}$$

e definire ¹

$$I = -r_2 \int_a^b P(x) \ln P(x) dx. \quad (1.25)$$

Vedremo che questa rinormalizzazione ci permetterà di trattare in maniera perfettamente coerente la fisica statistica classica. Ci sono però due prezzi da pagare

1. L'informazione mancante non è nulla, come dovrebbe, se la risposta è nota. Se sappiamo che la risposta è x_0 , $P(x)$ tende a $\delta(x - x_0)$, che sostituita nella (1.25) porta a $I \rightarrow -\infty$.
2. L'integrale al secondo membro della (1.25) contiene il logaritmo di $P(x)$, che non è un numero puro, ma ha dimensioni $[x^{-1}]$. A parte l'aspetto matematicamente "poco elegante", questo ha una conseguenza più seria. La risposta al nostro problema di scelta sul continuo avrebbe potuto essere scritta, invece che in termini della variabile x , in termini della variabile $\eta = \eta(x)$, con $\eta(x)$ continua e monotona nell'intervallo tra $\eta(a)$ e $\eta(b)$. La distribuzione di probabilità da usare sarebbe stata $P(\eta) = P[x(\eta)](dx/d\eta)$. Ma in questo caso, a seconda della funzione $\eta = \eta(x)$, a celle uguali in η non corrispondono necessariamente celle uguali in x , e a celle equiprobabili in η non corrispondono necessariamente celle equiprobabili in x , come si può facilmente controllare. Vedremo in seguito che, per il caso continuo, dovremo scegliere una variabile per la quale celle di estensione uguale possano essere considerate *equiprobabili a priori*, cioè equiprobabili quando non abbiamo alcuna informazione sul sistema.

L'informazione mancante definita in questo capitolo è legata sia all'*entropia di Shannon*, introdotta da Claude E. Shannon nel 1948 nell'ambito della teoria delle comunicazioni, che alla *entropia di von Neumann*, introdotta da John (nato János) von Neumann nel 1955 per estendere alla meccanica quantistica l'entropia classica di Gibbs. Nel paragrafo 4.2 discuteremo il legame tra entropia e informazione mancante.

¹ **Nota matematica:** La (1.25) ci dice che, nel caso della scelta nel continuo, matematicamente l'informazione mancante non è una funzione, ma un *funzionale*. Questo perché una funzione associa un valore numerico a un altro valore numerico come, per esempio, la funzione $\sin x$ associa un valore reale nell'intervallo $[-1, +1]$ a ogni x reale. Un funzionale, invece, associa un valore numerico non a un altro valore numerico, ma a una funzione. Così, nel nostro caso, la (1.25) associa un valore numerico dell'informazione mancante non a un altro valore numerico, ma a una funzione di distribuzione della probabilità $P(x)$.

Capitolo 2

Meccanica statistica

2.1 Statistica classica

2.1.1 Stato di un sistema classico e grandezze fisiche

Come abbiamo detto nel capitolo precedente, in meccanica classica lo stato di un sistema è completamente specificato dalla conoscenza delle sue coordinate canoniche, cioè dall'insieme delle coordinate generalizzate q_i e dei corrispondenti momenti coniugati p_i , con $1 \leq i \leq N_f$, dove N_f è il numero totale di gradi di libertà del sistema. L'esempio più semplice è costituito da un sistema di N punti materiali, per il quale le coordinate generalizzate $\{q_i\}$ sono

$$\begin{aligned}\{q_i\} &= q_1, q_2, q_3, q_4, q_5, q_6, \dots, q_{3j-2}, q_{3j-1}, q_{3j}, \dots, q_{3N-2}, q_{3N-1}, q_{3N} \\ &= x^{(1)}, y^{(1)}, z^{(1)}, x^{(2)}, y^{(2)}, z^{(2)}, \dots, x^{(j)}, y^{(j)}, z^{(j)}, \dots, x^{(N)}, y^{(N)}, z^{(N)},\end{aligned}\quad (2.1)$$

e i corrispondenti momenti coniugati $\{p_i\}$ sono

$$\begin{aligned}\{p_i\} &= p_1, p_2, p_3, p_4, p_5, p_6, \dots, p_{3j-2}, p_{3j-1}, p_{3j}, \dots, p_{3N-2}, p_{3N-1}, p_{3N} = \\ &= m^{(1)}\dot{x}^{(1)}, m^{(1)}\dot{y}^{(1)}, m^{(1)}\dot{z}^{(1)}, m^{(2)}\dot{x}^{(2)}, m^{(2)}\dot{y}^{(2)}, m^{(2)}\dot{z}^{(2)}, \dots, \\ &\quad m^{(j)}\dot{x}^{(j)}, m^{(j)}\dot{y}^{(j)}, m^{(j)}\dot{z}^{(j)}, \dots, m^{(N)}\dot{x}^{(N)}, m^{(N)}\dot{y}^{(N)}, m^{(N)}\dot{z}^{(N)},\end{aligned}\quad (2.2)$$

dove l'indice (j) corre su tutte le particelle del nostro sistema, e i gradi di libertà sono $N_f = 3N$. L'insieme delle coordinate e dei momenti $(q, p) = (q_1, \dots, q_{N_f}, p_1, \dots, p_{N_f})$ determina un vettore posizione in uno spazio delle fasi a $2N_f$ dimensioni. In altre parole, ad ogni possibile stato del sistema corrisponde un punto (o vettore posizione) nello spazio delle fasi.

Ricordiamo che le coordinate e i momenti obbediscono alle equazioni del moto canoniche

$$\dot{q}_i = \frac{\partial H}{\partial p_i}, \quad \dot{p}_i = -\frac{\partial H}{\partial q_i}, \quad (2.3)$$

(equazioni di Hamilton) dove $H(q, p, t)$ è l'hamiltoniana del sistema.

Ogni grandezza fisica relativa al nostro sistema è rappresentata da una funzione delle coordinate generalizzate e dei loro momenti coniugati del tipo $G(q, p)$. Per esempio, la quantità di moto totale sarà scritta $\mathbf{P} = \sum_{j=1}^N \mathbf{p}_j$, e la sua componente P_x sarà $P_x = \sum_{j=1}^N p_{3j-2}$, dove l'indice j corre su tutte le particelle.

2.1.2 Evoluzione temporale

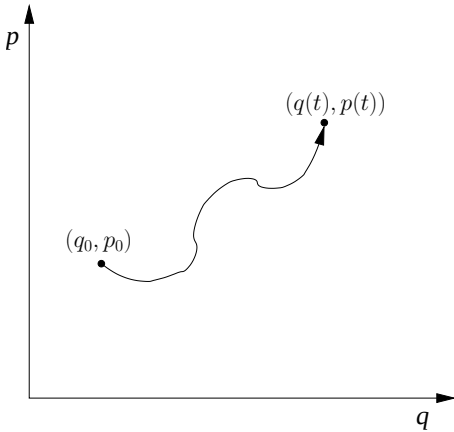


Figura 2.1 Traiettoria nello spazio delle fasi.

dove il termine $\{G, H\}$ indica la parentesi di Poisson della funzione che rappresenta la grandezza, $G(q, p)$, con l'hamiltoniana del sistema. In altre parole, il punto dello spazio delle fasi che rappresenta il nostro sistema descrive una certa traiettoria al passare del tempo, come in figura 2.1, e il valore della nostra grandezza G è dato, istante per istante, dal valore che l'espressione analitica $G(q, p)$ assume nel punto $[q_i(q_0, p_0, t), p_i(q_0, p_0, t)]$.

Notiamo che nel corso dell'evoluzione temporale rappresentata dall'equazione (2.5) manteniamo immutata la dipendenza analitica della funzione $G(q, p)$ dalle coordinate dello spazio delle fasi, e che la funzione $G(q, p)$ non ha alcuna dipendenza esplicita dal tempo t . Chi varia nel tempo sono le q e le p .

2.1.3 Trattamento statistico

Abbiamo già detto che la trattazione dinamica rigorosa di un sistema fisico composto da un numero di particelle dell'ordine del numero di Avogadro non è possibile. Questo implicherebbe conoscere esattamente il punto (q, p) dello spazio delle fasi corrispondente allo stato del sistema, e determinarne l'evoluzione temporale. Per farlo dovremmo utilizzare una quantità di informazione che non è assolutamente pensabile né ottenere, né, comunque, maneggiare. Trattando problemi macroscopici siamo quindi costretti a lavorare in presenza di *informazione mancante*, e dobbiamo sviluppare un formalismo che sfrutti al meglio l'informazione disponibile.

Per far questo ci accontenteremo di determinare non il punto esatto (q, p) che rappresenta lo stato del sistema, ma solo una funzione di distribuzione di probabilità su tutto lo spazio delle fasi. Il nostro obiettivo sarà quindi trovare una funzione $f(q, p)$ tale che, dato un qualunque volume τ dello spazio delle fasi, la quantità

$$W(\tau) = \int_{\tau} f(q, p) dq dp = \int_{\tau} f(q_1, \dots, q_{N_f}, p_1, \dots, p_{N_f}) dq_1 dq_2 \dots dq_{N_f} dp_1 dp_2 \dots dp_{N_f} \quad (2.6)$$

sia la probabilità che il punto rappresentativo dello stato del sistema sia all'interno di τ . Trattandosi di densità di probabilità devono valere le condizioni

$$\int_{\text{tutto lo spazio delle fasi}} f(q, p) dq dp = 1, \quad f(q, p) \geq 0 \quad \forall (q, p). \quad (2.7)$$

Data una grandezza fisica $G(q, p)$, l'informazione contenuta nella funzione di distribuzione di probabilità $f(q, p)$, inferiore all'informazione contenuta nella conoscenza del punto esatto (q, p) , non ci permetterà più di calcolarne il “valore esatto”, ma solo un *valore di aspettazione* $\langle G \rangle$, dato dalla media dei valori possibili pesata dalle loro probabilità, cioè

$$\langle G \rangle = \int G(q, p) f(q, p) dq dp, \quad (2.8)$$

con l'integrale esteso a tutto lo spazio delle fasi.

Le equazioni del moto canoniche (2.3) determinano poi l'evoluzione temporale del valore d'aspettazione $\langle G \rangle$ della (2.8). Questa evoluzione temporale può essere trattata secondo due rappresentazioni diverse ma equivalenti, descritte nel prossimo paragrafo.

2.1.4 Rappresentazioni di Schrödinger e di Heisenberg

Le rappresentazioni di Schrödinger e di Heisenberg sono nate per descrivere l'evoluzione temporale di un sistema quantistico. Nella rappresentazione di Schrödinger gli operatori corrispondenti alle osservabili, rappresentati da operatori differenziali, sono considerati costanti, mentre lo stato del sistema, rappresentato da una funzione d'onda, evolve nel tempo. Nella rappresentazione di Heisenberg, invece, lo stato del sistema è rappresentato da un vettore costante nel tempo in uno spazio di Hilbert, e la dipendenza dal tempo è incorporata negli operatori corrispondenti alle osservabili, in forma di matrici. I due modelli differiscono solo per un cambio di base rispetto alla dipendenza temporale, e le previsioni delle due rappresentazioni per l'evoluzione temporale dei valori di aspettazione delle variabili coincidono. Vedremo più sotto che cosa questo comporti. Queste rappresentazioni possono essere estese alla fisica classica.

Rappresentazione di Schrödinger

Cominciamo dalla rappresentazione di Schrödinger. Supponiamo di conoscere, a $t = 0$, la densità di probabilità $f(q, p)$ di localizzare il punto rappresentativo del nostro sistema nello spazio delle fasi. Questa densità di probabilità è considerata dipendente dal tempo, avremo cioè una $f(q, p, t)$. La probabilità che all'istante $t = 0$ il punto rappresentativo del nostro sistema sia in un certo volumetto $\tau(0)$ attorno a un certo punto (q_0, p_0) sarà, al primo ordine, $W = f(q_0, p_0, 0) \tau(0)$. Se aspettiamo un intervallo di tempo t ogni punto inizialmente in $\tau(0)$ avrà seguito una sua traiettoria determinata dalle equazioni canoniche di Hamilton (2.3), e si troverà in un nuovo volumetto $\tau(t)$, di forma, in generale, diversa da quella di $\tau(0)$, come schematizzato in Fig. 2.2. In particolare, il punto (q_0, p_0) si sarà spostato in un punto $[q(q_0, p_0, t), p(q_0, p_0, t)]$ contenuto in $\tau(t)$. Poiché ogni punto di $\tau(0)$ si è spostato in $\tau(t)$, la probabilità di trovare il punto rappresentativo del sistema in $\tau(t)$ all'istante t sarà

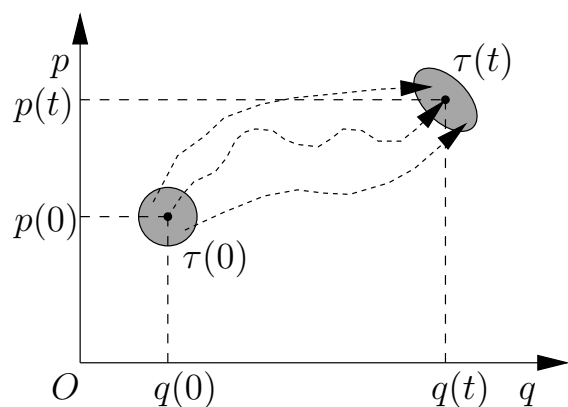


Figura 2.2 Trasformazione canonica nello spazio delle fasi.

uguale alla probabilità che inizialmente il punto fosse in $\tau(0)$. Avremo così

$$W = f(q_0, p_0, 0) \tau(0) = f(q(q_0, p_0, t), p(q_0, p_0, t)) \tau(t). \quad (2.9)$$

Ma le trasformazioni canoniche conservano gli elementi di volume dello spazio delle fasi, e avremo quindi $\tau(t) = \tau(0)$, da cui segue immediatamente che

$$f(q(q_0, p_0, t), p(q_0, p_0, t)) = f(q_0, p_0, 0) \quad \forall t \quad (2.10)$$

(teorema di Liouville). Questo implica che sia

$$0 = \frac{df}{dt} = \frac{\partial f}{\partial t} + \sum_{i=1}^{N_f} \left(\frac{\partial f}{\partial q_i} \frac{\partial q_i}{\partial t} + \frac{\partial f}{\partial p_i} \frac{\partial p_i}{\partial t} \right) = \frac{\partial f}{\partial t} + \sum_{i=1}^{N_f} \left(\frac{\partial f}{\partial q_i} \frac{\partial H}{\partial p_i} - \frac{\partial f}{\partial p_i} \frac{\partial H}{\partial q_i} \right), \quad (2.11)$$

da cui otteniamo per l'evoluzione temporale della funzione di distribuzione

$$\frac{\partial f}{\partial t} = - \sum_{i=1}^{N_f} \left(\frac{\partial f}{\partial q_i} \frac{\partial H}{\partial p_i} - \frac{\partial f}{\partial p_i} \frac{\partial H}{\partial q_i} \right) = -\{f, H\}. \quad (2.12)$$

E' importante notare che la parentesi di Poisson della (2.12) ha segno opposto a quello che appare nella (2.5). Dall'evoluzione temporale di $f(q, p, t)$ possiamo ricavare l'evoluzione temporale del valore di aspettazione di qualunque grandezza fisica $G(q, p)$ attraverso la relazione

$$\langle G(t) \rangle = \int G(q, p) f(q, p, t) dq dp. \quad (2.13)$$

Rappresentazione di Heisenberg

Nella rappresentazione di Heisenberg, invece, si considera costante nel tempo la distribuzione di probabilità $f(q, p)$ sullo spazio delle fasi, e si trasferisce la dipendenza temporale alla grandezza fisica G , che scriveremo quindi $G(q, p, t)$, e che evolverà secondo la (2.5). La quantità $dW = f(q, p) dq dp$ è così la probabilità che il nostro punto rappresentativo si trovi nell'intorno infinitesimo $dq dp$ di (q, p) sia all'istante iniziale che a ogni istante successivo. Questo corrisponde a dare probabilità dW al valore $G(q, p, t)$ della grandezza all'istante t . Per ottenere l'evoluzione temporale del valore di aspettazione effettueremo quindi la media pesata

$$\langle G(t) \rangle = \int G(q, p, t) dW = \int G(q, p, t) f(q, p) dq dp. \quad (2.14)$$

Ovviamente i valori per $\langle G(t) \rangle$ dati dalle (2.13) e (2.14) coincidono.

Relazione tra le due rappresentazioni

Per chiarire le idee può essere utile considerare un esempio elementare di elettromagnetismo classico. Supponiamo di avere una carica elettrica q ferma nel laboratorio, e di misurare il campo elettrico che questa genera in un punto P , che si allontana da essa con velocità costante v (supponiamo $v \ll c$). Con un'opportuna scelta dell'origine dei tempi calcoliamo il risultato della misura sia nel sistema di riferimento S solidale con la carica, in cui i vettori posizione di q e P sono, rispettivamente, $\mathbf{r}_q \equiv (0, 0, 0)$ e $\mathbf{r}_P \equiv (vt, 0, 0)$, che nel sistema di riferimento S' solidale con P , dove i vettori posizione sono $\mathbf{r}'_q \equiv (-vt, 0, 0)$ e $\mathbf{r}'_P \equiv (0, 0, 0)$.

1. Nel sistema di riferimento S (rappresentazione di Schrödinger) la posizione del punto in cui misuriamo il campo, $\mathbf{r}_p \equiv (vt, 0, 0)$, dipende esplicitamente dal tempo, in analogia con una funzione di distribuzione $f(q, p, t)$. Invece il campo elettrico in un punto qualunque fisso in S , di coordinate $\mathbf{r} \equiv (x, y, z)$, vale $\mathbf{E}(\mathbf{r}) = \frac{1}{4\pi\epsilon_0} \frac{q}{r^2} \hat{\mathbf{r}}$, senza alcuna dipendenza esplicita dal tempo.

Il campo elettrico misurato in $\mathbf{r}_p$ vale quindi $\mathbf{E}(\mathbf{r}_p) = \frac{1}{4\pi\epsilon_0} \frac{q}{(vt)^2} \hat{\mathbf{x}}$.

2. Nel sistema S' (rappresentazione di Heisenberg) la posizione del punto in cui misuriamo il campo è $\mathbf{r}'_p \equiv (0, 0, 0)$, indipendente dal tempo in analogia con il vettore che rappresenta lo stato di un sistema nella rappresentazione di Heisenberg. A variare nel tempo è la posizione della carica $\mathbf{r}'_q \equiv (-vt, 0, 0)$, e di conseguenza il campo elettrico nel generico punto $\mathbf{r}' \equiv (x', y', z')$, fisso in S' , che vale $\mathbf{E}'(\mathbf{r}') = \frac{1}{4\pi\epsilon_0} \frac{q}{v^2 t^2 + y'^2 + z'^2} \hat{\mathbf{r}}'$. Per il campo in $\mathbf{r}'_p$ abbiamo così $\mathbf{E}'(\mathbf{r}'_p) = \frac{1}{4\pi\epsilon_0} \frac{q}{(vt)^2} \hat{\mathbf{x}}'$. Poiché i versori $\hat{\mathbf{x}}$ e $\hat{\mathbf{x}}'$ coincidono, coincidono i valori della misura calcolati in S e in S' .

2.2 Statistica quantistica

2.2.1 Stato di un sistema quantistico e valori di aspettazione delle grandezze fisiche

In meccanica quantistica lo stato di un sistema è completamente specificato da un vettore $|\psi\rangle$ in uno spazio di Hilbert, normalizzato in modo che

$$\langle\psi|\psi\rangle = 1. \quad (2.15)$$

Alle grandezze fisiche osservabili corrispondono operatori $\hat{G}$ che operano su questo spazio. Il valore di aspettazione di una grandezza $\hat{G}$ quando il sistema si trova nello stato $|\psi\rangle$ è

$$\langle\hat{G}\rangle = \langle\psi|\hat{G}|\psi\rangle. \quad (2.16)$$

La grandezza assume un valore definito G_i (autovalore di $\hat{G}$) solo se il sistema si trova in un autostato $|i\rangle$ di $\hat{G}$

$$\hat{G}|i\rangle = G_i|i\rangle. \quad (2.17)$$

Se un sistema è confinato in un volume finito, gli autostati di ogni grandezza fisica $\hat{G}$ costituiscono insiemi normalmente infiniti, ma discreti. Autostati corrispondenti ad autovalori diversi della stessa grandezza fisica sono ortogonali tra loro. In presenza di autovalori degeneri, è comunque possibile scegliere gli autostati di $\hat{G}$ in modo che formino un insieme ortonormale completo. Questo significa che $\langle i|j\rangle = \delta_{ij}$ e che l'operatore identità $\hat{E}$ può essere scritto

$$\hat{E} = \sum_i |i\rangle\langle i|. \quad (2.18)$$

2.2.2 Evoluzione temporale

L'evoluzione temporale dello stato $|\psi\rangle$ di un sistema quantistico è data dall'equazione di Schrödinger dipendente dal tempo

$$\frac{d}{dt}|\psi\rangle = -\frac{i}{\hbar}\hat{H}|\psi\rangle, \quad (2.19)$$

dove $\hat{H}$ è l'hamiltoniana (operatore hamiltoniano) del sistema. L'equazione (2.19) può essere integrata formalmente per dare

$$|\psi(t)\rangle = e^{-\frac{i}{\hbar}\hat{H}t}|\psi(0)\rangle. \quad (2.20)$$

Nella (2.20) per esponenziale di un operatore $\hat{A}$ si intende l'espressione

$$e^{\hat{A}} = \hat{E} + \hat{A} + \frac{1}{2}\hat{A}^2 + \frac{1}{6}\hat{A}^3 + \dots + \frac{1}{n!}\hat{A}^n + \dots \quad (2.21)$$

dove $\hat{E} = \hat{A}^0$ è, come al solito, l'operatore identità.

Partendo dalla (2.16) l'evoluzione temporale del valore di aspettazione di una grandezza fisica $\hat{G}$ può essere scritta

$$\langle\hat{G}(t)\rangle = \langle\psi(t)|\hat{G}|\psi(t)\rangle. \quad (2.22)$$

Questa è la rappresentazione di Schrödinger, in cui l'operatore $\hat{G}$ viene considerato costante mentre evolve nel tempo lo stato del sistema fisico $|\psi(t)\rangle$.

D'altra parte, inserendo la (2.20) nella (2.22) otteniamo

$$\langle\hat{G}(t)\rangle = \langle\psi(0)|e^{+\frac{i}{\hbar}\hat{H}t}\hat{G}e^{-\frac{i}{\hbar}\hat{H}t}|\psi(0)\rangle = \langle\psi(0)|\hat{G}(t)|\psi(0)\rangle. \quad (2.23)$$

Questa è la rappresentazione di Heisenberg, in cui i vettori dello spazio di Hilbert sono considerati costanti nel tempo, mentre gli operatori rappresentanti le osservabili fisiche evolvono secondo la legge

$$\hat{G}(t) = e^{+\frac{i}{\hbar}\hat{H}t}\hat{G}(0)e^{-\frac{i}{\hbar}\hat{H}t}. \quad (2.24)$$

Calcolando la derivata temporale della (2.24) otteniamo

$$\frac{d}{dt}\hat{G}(t) = \frac{i}{\hbar}[\hat{H}, \hat{G}(t)]. \quad (2.25)$$

Sia che usiamo la rappresentazione di Schrödinger, sia che usiamo quella di Heisenberg, l'evoluzione temporale del valore di aspettazione di $\hat{G}$ rimane ovviamente la stessa.

2.2.3 Trattamento statistico e matrice densità

Quando abbiamo a che fare con un sistema quantistico composto da un grande numero di particelle diventa impossibile determinare esattamente in quale stato $|\psi\rangle$ il sistema si trovi. Anche qui ci troviamo quindi ad operare in assenza di un'informazione completa. Ci accontenteremo di conoscere, o meglio, di cercare di determinare, le probabilità W_i che il sistema si trovi in ognuno degli stati $|i\rangle$ di una certa base ortonormale completa $\{|i\rangle\}$ dello spazio di Hilbert. Dovrà essere $\sum_i W_i = 1$ e $W_i \geq 0 \forall i$. Tenendo presente che se fossimo sicuri che il sistema si trova nello stato $|i\rangle$ il valore d'aspettazione

dell'osservabile $\hat{G}$ sarebbe $\langle i|\hat{G}|i\rangle$, la stima migliore che possiamo fare per il risultato di una misura è data dalla media pesata

$$\langle \hat{G} \rangle = \sum_i W_i \langle i|\hat{G}|i\rangle. \quad (2.26)$$

A questo punto definiamo un operatore $\hat{\rho}$, detto *matrice densità* e che tratteremo in maggior dettaglio nel capitolo 11, diagonale sulla base $\{|i\rangle\}$ e definito dalla relazione

$$\langle i|\hat{\rho}|j\rangle = W_i \delta_{ij}. \quad (2.27)$$

In termini di $\hat{\rho}$ l'equazione (2.26) per il valore di aspettazione si riscrive

$$\langle \hat{G} \rangle = \sum_i \langle i|\hat{\rho}|i\rangle \langle i|\hat{G}|i\rangle = \underbrace{\sum_{ij} \langle i|\hat{\rho}|j\rangle \langle j|\hat{G}|i\rangle}_{\text{perché } \langle i|\hat{\rho}|j\rangle = W_i \delta_{ij}} = \sum_i \langle i|\hat{\rho} \hat{G}|i\rangle = \text{Tr}(\hat{\rho} \hat{G}) \quad (2.28)$$

E' da notare che, se le W_i non sono tutte uguali tra loro, $\hat{\rho}$ non rimane necessariamente diagonale per un cambiamento di base dello spazio di Hilbert. Però un cambiamento di coordinate conserva il valore della traccia degli operatori, quindi la relazione $\langle \hat{G} \rangle = \text{Tr}(\hat{\rho} \hat{G})$ rimane valida indipendentemente dalla base scelta.

Per quel che riguarda l'evoluzione temporale del valore di aspettazione, la rappresentazione di Schrödinger considera l'operatore $\hat{G}$ costante nel tempo, mentre fa evolvere la matrice densità $\hat{\rho}$. Nella base in cui è diagonale, la matrice densità può essere scritta

$$\hat{\rho}(t) = \sum_i |i(t)\rangle W_i \langle i(t)|, \quad \text{da cui} \quad \hat{\rho}(t) = \sum_i e^{-\frac{i}{\hbar} \hat{H} t} |i\rangle W_i \langle i| e^{+\frac{i}{\hbar} \hat{H} t} = e^{-\frac{i}{\hbar} \hat{H} t} \hat{\rho}(0) e^{+\frac{i}{\hbar} \hat{H} t}. \quad (2.29)$$

La rappresentazione di Heisenberg, invece, considera la matrice densità costante nel tempo, facendo evolvere l'operatore $\hat{G}(t)$ secondo la relazione (2.24). E' importante notare che, rispetto alla (2.24), la (2.29) fa evolvere la matrice densità "all'indietro nel tempo".

Capitolo 3

Scelta della probabilità degli stati

3.1 Primo postulato: riproduzione dell'informazione disponibile

Abbiamo visto che, disponendo di una quantità di informazione molto inferiore a quella necessaria per una trattazione dinamica del problema, dobbiamo ricorrere a una *trattazione probabilistica*. Per farlo, dobbiamo introdurre due principi, o postulati, su come costruire una funzione di distribuzione di probabilità sullo spazio delle fasi nel caso classico, o una matrice densità nel caso quantistico, partendo dall'informazione a noi disponibile. Questi postulati, come tutti i postulati che si incontrano in fisica, non sono ricavabili da principi preesistenti, ma saranno ritenuti giustificati a posteriori se porteranno a risultati in accordo con l'esperienza.

Il primo postulato ha bisogno di una premessa, o di un “postulato zero”, su cosa fare in assenza totale di informazione. *In assenza completa di informazione disponibile in statistica classica assegniamo uguale probabilità a priori a uguali volumi dello spazio delle fasi*. Questa scelta, nel caso di un sistema fisico costituito da particelle non identiche, privilegia una trattazione hamiltoniana rispetto, per esempio, a una trattazione lagrangiana. Infatti, in una trattazione hamiltoniana le variabili sono le q_i e le p_i , in una trattazione lagrangiana le variabili sono le q_i e le $\dot{q}_i$. Se le particelle del sistema sono tutte identiche, a uguali volumi nello spazio qp corrispondono uguali volumi nello spazio $q\dot{q}$, essendo $p_i = m\dot{q}_i \forall i$, ma se il sistema contiene particelle di massa diversa questo non è più vero. In assenza di informazione disponibile sul sistema, la funzione di distribuzione $f(q, p)$ deve quindi avere la proprietà che sia

$$\int_{\tau_1} f(q, p) dq dp = \int_{\tau_2} f(q, p) dq dp, \quad (3.1)$$

per ogni coppia di domini τ_1 e τ_2 dello spazio delle fasi che abbiano uguale volume.

Sempre in assenza di informazione disponibile, ma *in statistica quantistica, assegniamo uguale probabilità a priori a tutti gli elementi di ogni base ortonormale completa dello spazio di Hilbert degli stati del sistema*. Se gli stati di questa base sono N , e li etichettiamo $|i\rangle$, con $1 \leq i \leq N$, avremo $W_i = 1/N$ per ogni i . Questo corrisponde a una matrice densità diagonale

$$\hat{\rho} = \frac{1}{N} \hat{E}, \quad (3.2)$$

dove $\hat{E}$ è la matrice identità, indipendentemente dalla base scelta.

Fatta questa premessa, resta da decidere che cosa richiedere a $f(q, p)$ e $\hat{\rho}$ quando invece disponiamo di qualche informazione sul sistema. Disporre di informazione sul sistema vuol dire che, a un

certo istante, ci è noto il valore di qualche grandezza fisica che abbiamo misurato. La richiesta ovvia, che conclude il primo postulato, è che, *se esistono grandezze fisiche che abbiamo misurato, i loro valori di aspettazione previsti dalla nostra trattazione probabilistica devono coincidere con i valori noti dalle misure*. Quindi, se nel caso classico abbiamo n grandezze fisiche $G_i(q, p)$ di ognuna delle quali conosciamo il valore g_i , la $f(q, p)$ dovrà essere tale che

$$\langle G_i \rangle = \int G_i(q, p) f(q, p) dq dp = g_i \quad \forall i. \quad (3.3)$$

La corrispondente informazione mancante rinormalizzata sarà

$$I = -r_2 \int f(q, p) \ln f(q, p) dq dp. \quad (3.4)$$

Analogamente, se nel caso quantistico abbiamo n osservabili $\hat{G}_i$ di ognuna delle quali conosciamo con certezza il valore g_i , la matrice densità $\hat{\rho}$ dovrà essere tale che

$$\langle G_i \rangle = \text{Tr}(\hat{G}_i \hat{\rho}) = g_i \quad \forall i. \quad (3.5)$$

L'informazione mancante nel caso quantistico si ottiene partendo da una base in cui $\hat{\rho}$ è diagonale, con $\rho_{ii} = W_i$

$$I = -r_2 \sum_i W_i \ln W_i = -r_2 \sum_i \rho_{ii} \ln \rho_{ii} = -r_2 \text{Tr}(\hat{\rho} \ln \hat{\rho}), \quad (3.6)$$

dove il logaritmo di un operatore è definito come l'operatore il cui esponenziale (2.21) è uguale all'operatore di partenza (funzione inversa). E' da notare che, in matematica, non tutte le matrici hanno un logaritmo, e quelle che ce l'hanno possono averne più di uno.

3.2 Secondo postulato: massimo condizionato dell'informazione mancante

Il nostro primo postulato ci lascia ancora molta libertà, e non è sufficiente a determinare univocamente la $f(q, p)$ o la $\hat{\rho}$. Ci serve un ulteriore criterio di scelta. Consideriamo, nel caso classico, due diverse funzioni di distribuzione $f_1(q, p)$ e $f_2(q, p)$ che riproducano ambedue, come valori di aspettazione, gli n valori noti sperimentalmente g_i di n grandezze fisiche $G_i(q, p)$, abbiamo cioè

$$\langle G_i \rangle = \int G_i(q, p) f_1(q, p) dq dp = \int G_i(q, p) f_2(q, p) dq dp = g_i \quad \forall i. \quad (3.7)$$

In base a quale criterio preferiremo f_1 rispetto a f_2 , o vice versa? Diamo un'occhiata alle informazioni mancanti relative alle due distribuzioni:

$$I_1 = -r_2 \int f_1(q, p) \ln f_1(q, p) dq dp, \quad I_2 = -r_2 \int f_2(q, p) \ln f_2(q, p) dq dp, \quad (3.8)$$

e supponiamo che sia $I_1 > I_2$. Tutte e due le funzioni di distribuzione riproducono l'informazione nota, perché i valori di aspettazione delle grandezze che abbiamo misurato coincidono con i valori sperimentali, ma f_2 , avendo un'informazione mancante minore, per definizione contiene una quantità

di informazione maggiore. Questo eccesso di informazione non ha niente a che fare con l'informazione disponibile sperimentalmente. Quindi, preferiremo f_1 , che ci riproduce tutti i dati sperimentali "inventandosi meno cose".

Arriviamo così al nostro secondo postulato: *in statistica classica sceglieremo la funzione di distribuzione di probabilità sullo spazio delle fasi, $f(q, p)$, che massimizza l'informazione mancante, rispettando la condizione di riprodurre tutta l'informazione sperimentale a noi nota. In statistica quantistica sceglieremo la matrice densità $\hat{\rho}$ che massimizza l'informazione mancante, sempre rispettando la condizione di riprodurre tutta l'informazione sperimentale a noi nota.*

Dovremo quindi risolvere problemi di massimo condizionato, e, per far questo, ricorreremo al metodo dei moltiplicatori di Lagrange, descritto nel paragrafo seguente.

3.3 Moltiplicatori di Lagrange

Per semplicità, supponiamo di dover cercare il massimo di una funzione continua e derivabile di due variabili, $f(x, y)$, soggetto alla condizione $g(x, y) = k$. Un'equazione del tipo $\varphi(x, y) = \rho$, con ρ costante, determina una linea sul piano xy . Così la condizione $g(x, y) = k$ è rappresentata, per esempio, dalla linea continua della Fig. 3.1. Su questa linea dobbiamo cercare gli estremi condizionati di $f(x, y)$. D'altra parte, al variare di h , le equazioni $f(x, y) = h$ determinano la famiglia delle curve tratteggiate disegnate sulla stessa figura. Alcune di queste curve non incontrano mai la curva $g(x, y) = k$, come la $f(x, y) = h_1$, alcune la intersecano, come le $f(x, y) = h_3$, $f(x, y) = h_4$ e $f(x, y) = h_5$. Altre, infine, le sono tangenti, come la $f(x, y) = h_2$ nel punto P_4 . Dalla figura vediamo che la curva $g(x, y) = k$ in ogni suo punto o interseca una linea del tipo $f(x, y) = h$, come in P_1, P_2, P_3, P_5, P_6 e P_7 , o è tangente a una linea di questo tipo, come in P_4 . Gli estremi condizionati di $f(x, y)$ possono trovarsi solo in punti del tipo P_4 .

Infatti vediamo che nei punti P_3 e P_5 la funzione $f(x, y)$ assume lo stesso valore h_3 , mentre nel punto intermedio P_4 ha un valore diverso h_2 . Ci troveremo quindi in presenza di un massimo condizionato se $h_2 > h_3$, o di un minimo condizionato se $h_2 < h_3$. Il problema di trovare estremi condizionati di $f(x, y)$ soggetti al vincolo $g(x, y) = k$ si risolve quindi individuando i punti in cui la curva $g(x, y) = k$ è tangente a una curva del tipo $f(x, y) = h$. Per individuare questi punti notiamo che il gradiente di $f(x, y)$, denotato $\nabla f(x, y)$, è un vettore sempre perpendicolare alle curve $f(x, y) = h$, mentre il gradiente $\nabla g(x, y)$ è sempre perpendicolare alla curva $g(x, y) = k$. Quindi i vettori $\nabla f(x, y)$ e $\nabla g(x, y)$ sono paralleli nei punti in cui la curva $g(x, y) = k$ è tangente a una delle curve $f(x, y) = h$. In questi punti avremo $\nabla f(x, y) = \lambda \nabla g(x, y)$, con λ valore scalare opportuno.

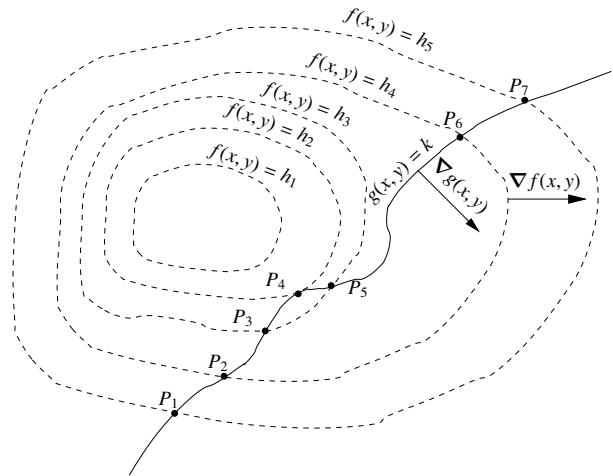


Figura 3.1 Ricerca di estremo condizionato.

Introduciamo una nuova funzione delle tre variabili x, y e λ , detta funzione di Lagrange, $F(x, y, \lambda) = f(x, y) - \lambda [g(x, y) - k]$, e cerchiamo i punti dello spazio $xy\lambda$ cui $F(x, y, \lambda)$ ha estremi incondizionati.

In questi punti le sue derivate parziali rispetto a x , y e λ devono essere tutte e tre nulle:

$$\begin{aligned} \frac{\partial F(x, y, \lambda)}{\partial x} &= \frac{\partial f(x, y)}{\partial x} - \lambda \frac{\partial g(x, y)}{\partial x} = 0, & \text{da cui} & \quad \frac{\partial f(x, y)}{\partial x} = \lambda \frac{\partial g(x, y)}{\partial x} \\ \frac{\partial F(x, y, \lambda)}{\partial y} &= \frac{\partial f(x, y)}{\partial y} - \lambda \frac{\partial g(x, y)}{\partial y} = 0, & \text{da cui} & \quad \frac{\partial f(x, y)}{\partial y} = \lambda \frac{\partial g(x, y)}{\partial y} \\ \frac{\partial F(x, y, \lambda)}{\partial \lambda} &= -[g(x, y) - k] = 0, & \text{da cui} & \quad g(x, y) = k, \end{aligned} \quad (3.9)$$

ci troviamo quindi in punti in cui i gradienti di $f(x, y)$ e $g(x, y)$ sono paralleli, ed è soddisfatta la condizione $g(x, y) = k$.

La ricerca degli estremi condizionati della funzione di due variabili $f(x, y)$ con il vincolo $g(x, y) = k$ si riduce quindi alla ricerca degli estremi incondizionati della funzione di tre variabili $F(x, y, \lambda)$. Al parametro λ si dà il nome di *moltiplicatore di Lagrange*. Ma attenzione: i punti così trovati sono *punti stazionari* di F , non necessariamente estremi: potrebbero essere *punti di sella*. Quindi, in generale, bisognerà controllare che corrispondano effettivamente a estremi condizionati di f .

Consideriamo adesso la ricerca degli estremi di una funzione di n variabili $f(x_1, x_2, \dots, x_n)$ sottoposti a m vincoli (con, ovviamente, $m < n$) $g_1(x_1, x_2, \dots, x_n) = k_1$, $g_2(x_1, x_2, \dots, x_n) = k_2 \dots g_m(x_1, x_2, \dots, x_n) = k_m$. Si può dimostrare che il problema può essere risolto cercando gli estremi incondizionati della nuova funzione di $n + m$ variabili

$$F(x_1, x_2, \dots, x_n, \lambda_1, \lambda_2, \dots, \lambda_m) = f(x_1, x_2, \dots, x_n) - \sum_{i=1}^m \lambda_i [g_i(x_1, x_2, \dots, x_n) - k_i],$$

dove i λ_i sono moltiplicatori di Lagrange. Per il seguito, ci sarà utile notare che gli estremi condizionati di $f(x_1 \dots x_n)$ si trovano nei punti in cui è nulla la variazione al primo ordine di $F(x_1 \dots x_n, \lambda_1 \dots \lambda_m)$

$$\delta F = \sum_{i=1}^n \frac{\partial F}{\partial x_i} \delta x_i + \sum_{j=1}^m \frac{\partial F}{\partial \lambda_j} \delta \lambda_j = 0, \quad (3.10)$$

dove le variazioni δx_i e i $\delta \lambda_j$ sono piccole, ma, per il resto, completamente arbitrarie e indipendenti l'una dall'altra.

3.4 Informazione disponibile in statistica classica

Nel caso della statistica classica, l'informazione mancante è data dalla (3.4). Quindi, come abbiamo visto nella nota 1 di pagina 7, quello che dobbiamo massimizzare non è una funzione, ma un *funzionale*. Ci troviamo così nel campo della matematica chiamato *calcolo delle variazioni*. Dobbiamo trovare la funzione di distribuzione $f(q, p)$ che massimizzi il funzionale $I[f(q, p)]$ corrispondente all'informazione mancante

$$I[f(q, p)] = -r_2 \int f(q, p) \ln f(q, p) dq dp$$

con le ipotesi che la $f(p, q)$ sia una funzione continua e derivabile, e che rispetti i vincoli

$$\int f(q, p) dq dp = 1 \quad (\text{condizione di normalizzazione di } f), \text{ e} \quad (3.11)$$

$$\int f(q, p) G_i(q, p) dq dp = g_i \quad \text{con } i = 1, \dots, n \quad (3.12)$$

dove $G_1 \dots G_n$ sono le grandezze fisiche di cui conosciamo il valore.

Se la nostra funzione di distribuzione $f(q, p)$ corrisponde a un estremo vincolato del funzionale I , deve essere nulla, al primo ordine, la variazione δI corrispondente a piccole variazioni $\delta f(q, p)$, variazioni soggette alla condizione che $f(q, p) + \delta f(q, p)$ continui a rispettare i vincoli (3.11) e (3.12). Analogamente a quanto visto nella 3.10, il problema può essere risolto introducendo $n + 1$ moltiplicatori di Lagrange scritti, per comodità, nella forma $-r_2(\Omega + 1)$ e $r_2 \lambda_i$ per $i = 1, \dots, n$, e imponendo che, per variazioni piccole, ma per il resto del tutto arbitrarie (cioè non sottoposte a vincoli) $\delta f(q, p)$ della funzione di distribuzione, sia nulla al primo ordine la variazione del funzionale I' definito da

$$\begin{aligned} I' &= I + r_2(\Omega + 1) \left[\int f(q, p) dq dp - 1 \right] - r_2 \sum_{i=1}^n \lambda_i \left[\int f(q, p) G_i(q, p) dq dp - g_i \right] \\ &= -r_2 \int f(q, p) \ln f(q, p) dq dp + r_2(\Omega + 1) \left[\int f(q, p) dq dp - 1 \right] \\ &\quad - r_2 \sum_{i=1}^n \lambda_i \left[\int f(q, p) G_i(q, p) dq dp - g_i \right] \\ &= -r_2 \int f(q, p) \left[\ln f(q, p) - \Omega - 1 + \sum_{i=1}^n \lambda_i G_i(q, p) \right] dq dp - r_2 \left(\Omega + 1 - \sum_{i=1}^n \lambda_i g_i \right), \end{aligned} \quad (3.13)$$

I valori dei moltiplicatori di Lagrange verranno poi determinati in modo da imporre le condizioni sulla normalizzazione di f e sui valori di aspettazione delle grandezze note. Per la variazione di I' abbiamo, tenendo presente che la variazione dell'ultimo termine tra parentesi tonda della (3.13) è nulla perché indipendente da $f(q, p)$

$$\begin{aligned} \delta I' &= -r_2 \int \delta f(q, p) \left[\ln f(q, p) - \Omega - 1 + \sum_{i=1}^n \lambda_i G_i(q, p) \right] dq dp - r_2 \int f(q, p) \frac{\delta f(q, p)}{f(q, p)} dq dp \\ &= -r_2 \int \delta f(q, p) \left[\ln f(q, p) - \Omega + \sum_{i=1}^n \lambda_i G_i(q, p) \right] dq dp, \end{aligned} \quad (3.14)$$

dove $\delta f(q, p)$ è una variazione infinitesima arbitraria di $f(q, p)$. L'arbitrarietà di δf impone che, perché $\delta I'$ sia nulla, debba essere nulla la parentesi quadra all'ultimo membro della (3.14), e otteniamo così

$$\ln f(q, p) = \Omega - \sum_{i=1}^n \lambda_i G_i(q, p), \quad \text{ovvero} \quad f(q, p) = e^{\Omega - \sum_{i=1}^n \lambda_i G_i(q, p)}. \quad (3.15)$$

Questa soluzione soddisfa la condizione $f(q, p) \geq 0$ per ogni punto dello spazio delle fasi. Adesso determiniamo i moltiplicatori di Lagrange. La condizione di normalizzazione impone

$$1 = \int f(q, p) dq dp = \int e^{\Omega - \sum_{i=1}^n \lambda_i G_i(q, p)} dq dp = e^{\Omega} \int e^{-\sum_{i=1}^n \lambda_i G_i(q, p)} dq dp$$

da cui otteniamo

$$e^{\Omega} = \left[\int e^{-\sum_{i=1}^n \lambda_i G_i(q, p)} dq dp \right]^{-1}, \quad \text{e} \quad \Omega = -\ln \int e^{-\sum_{i=1}^n \lambda_i G_i(q, p)} dq dp. \quad (3.16)$$

Dagli altri vincoli otteniamo

$$g_i = \int f(q, p) G_i(q, p) dq dp = \int e^{\Omega - \sum_{j=1}^n \lambda_j G_j(q, p)} G_i(q, p) dq dp = \frac{\int e^{-\sum_{j=1}^n \lambda_j G_j(q, p)} G_i(q, p) dq dp}{\int e^{-\sum_{j=1}^n \lambda_j G_j(q, p)} dq dp}, \quad (3.17)$$

che, guardando la (3.16), può essere riscritta nella forma

$$\frac{\partial \Omega}{\partial \lambda_i} = g_i, \quad i = 1, \dots, n. \quad (3.18)$$

Le equazioni (3.16) e (3.18) ci permettono di ottenere Ω e gli altri moltiplicatori di Lagrange λ_i in funzione dei valori noti delle grandezze fisiche g_i e delle espressioni analitiche delle grandezze $G_i(q, p)$.

Per quel che riguarda l'informazione mancante abbiamo

$$\begin{aligned} I &= -r_2 \int f(q, p) \ln f(q, p) dq dp = -r_2 \int e^{\Omega - \sum_{i=1}^n \lambda_i G_i(q, p)} \left[\Omega - \sum_{i=1}^n \lambda_i G_i(q, p) \right] dq dp \\ &= -r_2 \Omega + r_2 \int \sum_{i=1}^n \lambda_i G_i(q, p) e^{\Omega - \sum_{i=1}^n \lambda_i G_i(q, p)} dq dp = -r_2 \Omega + r_2 \sum_{i=1}^n \lambda_i g_i. \end{aligned} \quad (3.19)$$

Nel caso importante in cui l'unica grandezza fisica di cui conosciamo il valore è l'energia, chiamando β il moltiplicatore di Lagrange relativo all'hamiltoniana $H(q, p)$, e U il valore dell'energia stesso, abbiamo

$$f(q, p) = e^{\Omega - \beta H(q, p)}, \quad \Omega = -\ln \int e^{-\beta H(q, p)} dq dp, \quad I = -r_2 \Omega + r_2 \beta U \quad (3.20)$$

3.5 Particelle identiche in fisica classica

Finora abbiamo trattato le particelle come *distinguibili*. Questo non vuole solo dire attribuire a ogni singola particella i la sua posizione $\mathbf{q}_i$ e il suo momento coniugato $\mathbf{p}_i$. Questo vuole anche dire considerare come distinti due punti dello spazio delle fasi che si ottengono scambiando tra loro simultaneamente le posizioni e i momenti di una o più coppie di particelle.

Ma in natura particelle dello stesso tipo (cioè, per esempio, due atomi di ^4He , ma non un atomo di ^4He e uno di ^3He) sono *indistinguibili*. Non esiste un esperimento che permetta di distinguere un atomo di ^4He dall'altro. Questo implica che punti diversi dello spazio delle fasi che si trasformano l'uno nell'altro per una qualunque permutazione di particelle identiche non corrispondono a stati diversi, ma a un unico stato del sistema. In altre parole, nel caso di particelle identiche, un unico stato del sistema non è rappresentato da uno, ma da più punti dello spazio delle fasi. Il fatto di non sapere se è la particella A a trovarsi nella posizione $\mathbf{q}_1$ con momento $\mathbf{p}_1$, e la particella B a trovarsi nella posizione $\mathbf{q}_2$ con momento $\mathbf{p}_2$, piuttosto che viceversa, non corrisponde a mancanza di informazione relativa a una scelta tra due stati diversi, ma a una scelta tra due punti dello spazio delle fasi corrispondenti allo stesso stato del sistema, quindi non è vera informazione mancante.

Consideriamo un sistema costituito da N particelle classiche identiche. Esistono $N!$ permutazioni possibili delle N particelle, quindi ogni singolo stato del sistema è rappresentato da $N!$ punti distinti dello spazio delle fasi. L'informazione mancante relativa a una scelta tra $N!$ casi equiprobabili è

$$I_{N!} = r_2 \ln N!, \quad (3.21)$$

e questa quantità che dovrà essere sottratta dalla (3.19), ottenendo

$$I_{\text{indist}} = -r_2 \Omega + r_2 \sum_i^n \lambda_i g_i - r_2 \ln N!, \quad (3.22)$$

o, se l'unica grandezza nota è l'energia

$$I_{\text{indist}} = -r_2 \Omega + r_2 \beta U - r_2 \ln N! \quad (3.23)$$

In realtà, così facendo sovraestimiamo la correzione a I ogni volta che ci sono due o più particelle che hanno gli stessi $\mathbf{q}$ e $\mathbf{p}$. Questo, però, succede solo su di un insieme di misura zero.

Il termine $-r_2 \ln N!$ della (3.23) non dipende dalle variabili canoniche q e p , e può essere incorporato in Ω , vedi (3.16)

$$\begin{aligned} \Omega_{\text{indist}} &= \Omega_{\text{dist}} + \ln N! = -\ln \int e^{-\sum_{i=1}^n \lambda_i G_i(q,p)} dq dp + \ln N! \\ &= -\ln \left[\frac{1}{N!} \int e^{-\sum_{i=1}^n \lambda_i G_i(q,p)} dq dp \right] \end{aligned} \quad (3.24)$$

in modo che l'informazione mancante continui a essere scritta

$$I_{\text{indist}} = -r_2 \Omega_{\text{indist}} + r_2 \sum_{i=1}^n \lambda_i g_i, \quad \text{o, se è nota solo l'energia} \quad I_{\text{indist}} = -r_2 \Omega_{\text{indist}} + r_2 \beta U. \quad (3.25)$$

Per la funzione di distribuzione $f(q, p)$ avremo

$$f_{\text{indist}}(q, p) = e^{\Omega_{\text{indist}} - \beta H(q,p)} = e^{\Omega_{\text{dist}} + \ln N! - \beta H(q,p)} = N! e^{\Omega_{\text{dist}} - \beta H(q,p)} = N! f_{\text{dist}}(q, p). \quad (3.26)$$

Avremo quindi

$$\int f_{\text{indist}}(q, p) dq dp = N! \int f_{\text{dist}}(q, p) dq dp = N!, \quad (3.27)$$

e il valore di aspettazione di una grandezza $G(q, p)$ sarà

$$\langle G \rangle = \frac{1}{N!} \int G(q, p) f_{\text{indist}}(q, p) dq dp, \quad (3.28)$$

dove il fattore $1/N!$ tiene conto che a ogni stato del sistema corrispondono non uno, ma $N!$ punti distinti dello spazio delle fasi. E' da notare che, fin che il numero N delle particelle è costante, il tener conto dell'indistinguibilità delle particelle non altera i risultati ottenuti precedentemente. La correzione all'informazione mancante moltiplica semplicemente la funzione di distribuzione per $N!$, mentre, quando vengono calcolate le medie sullo spazio delle fasi viene introdotto un fattore $1/N!$ per tener conto degli $N!$ punti che corrispondono a un singolo stato fisico. L'importanza dell'indistinguibilità delle particelle apparirà quando introdurremo l'*insieme macrocanonico* per trattare sistemi con numero di particelle variabile.

3.6 Informazione disponibile in statistica quantistica

Anche qui dobbiamo massimizzare l'informazione mancante I , la cui espressione questa volta è data dalla (3.6), con i vincoli $\text{Tr}(\hat{\rho}) = 1$ (condizione di normalizzazione di $\hat{\rho}$) e $\text{Tr}(\hat{\rho} \hat{G}_i) = g_i$ con $i = 1, \dots, n$, dove $\hat{G}_1 \dots \hat{G}_n$ sono gli operatori relativi alle osservabili di cui conosciamo il valore. Introduciamo ancora gli $n+1$ moltiplicatori di Lagrange nella forma $-r_2(\Omega+1)$ e $r_2 \lambda_i$ per $i = 1, \dots, n$, e cerchiamo la matrice densità che dia il massimo non vincolato della funzione

$$\begin{aligned}
 I' &= I + r_2(\Omega+1) [\text{Tr}\hat{\rho} - 1] - r_2 \sum_{i=1}^n \lambda_i [\text{Tr}(\hat{\rho} \hat{G}_i) - g_i] \\
 &= -r_2 \text{Tr}(\hat{\rho} \ln \hat{\rho}) + r_2(\Omega+1) \text{Tr}\hat{\rho} - r_2 \sum_{i=1}^n \lambda_i \text{Tr}(\hat{\rho} \hat{G}_i) - r_2 \left(\Omega+1 - \sum_{i=1}^n \lambda_i g_i \right) \\
 &= -r_2 \text{Tr} \left[\hat{\rho} \left(\ln \hat{\rho} - \Omega - 1 + \sum_{i=1}^n \lambda_i \hat{G}_i \right) \right] - r_2 \left(\Omega+1 - \sum_{i=1}^n \lambda_i g_i \right). \tag{3.29}
 \end{aligned}$$

Si può dimostrare che la stazionarietà di I' richiede che gli operatori $\hat{\rho}$ e $\sum_{i=1}^n \lambda_i \hat{G}_i$ commutino, per cui potremo lavorare su una base $\{|\alpha\rangle\}$ su cui sia $\hat{\rho}$ che $\sum_{i=1}^n \lambda_i \hat{G}_i$ (ma non necessariamente i singoli $\hat{G}_i$) sono diagonali. Per “stenografia”, definiamo l'operatore $\hat{\mathcal{G}}$

$$\hat{\mathcal{G}} = \sum_{i=1}^n \lambda_i \hat{G}_i = \sum_{\alpha} |\alpha\rangle \mathcal{G}_{\alpha\alpha} \langle\alpha|,$$

e introduciamolo nella (3.29)

$$\begin{aligned}
 I' &= -r_2 \sum_{\alpha} \langle\alpha| \left[\hat{\rho} (\ln \hat{\rho} - \Omega - 1 + \hat{\mathcal{G}}) \right] |\alpha\rangle - r_2 \left(\Omega+1 - \sum_{i=1}^n \lambda_i g_i \right) \\
 &= -r_2 \sum_{\alpha} [\rho_{\alpha\alpha} (\ln \rho_{\alpha\alpha} - \Omega - 1 + \mathcal{G}_{\alpha\alpha})] - r_2 \left(\Omega+1 - \sum_{i=1}^n \lambda_i g_i \right),
 \end{aligned}$$

che deve essere stazionario rispetto a piccole variazioni arbitrarie $\delta\rho_{\alpha\alpha}$. La variazione dell'ultimo termine tra parentesi tonda della (3.29) è nulla perché non vi compare la $\hat{\rho}$. Avremo quindi

$$\begin{aligned}
 0 = \delta I' &= -r_2 \sum_{\alpha} \left[\delta\rho_{\alpha\alpha} (\ln \rho_{\alpha\alpha} - \Omega - 1 + \mathcal{G}_{\alpha\alpha}) + \rho_{\alpha\alpha} \left(\frac{\delta\rho_{\alpha\alpha}}{\rho_{\alpha\alpha}} \right) \right] \\
 &= -r_2 \sum_{\alpha} [\delta\rho_{\alpha\alpha} (\ln \rho_{\alpha\alpha} - \Omega + \mathcal{G}_{\alpha\alpha})],
 \end{aligned}$$

che ci porta a

$$\ln \rho_{\alpha\alpha} - \Omega + \mathcal{G}_{\alpha\alpha} = 0, \quad \text{che implica} \quad \rho_{\alpha\alpha} = e^{\Omega - \mathcal{G}_{\alpha\alpha}},$$

ovvero

$$\hat{\rho} = \sum_{\alpha} |\alpha\rangle e^{\Omega - \mathcal{G}_{\alpha\alpha}} \langle\alpha| = e^{\Omega - \sum_i \lambda_i \hat{G}_i}. \tag{3.30}$$

Notiamo che tutti gli autovalori $\rho_{\alpha\alpha}$ di $\hat{\rho}$ sono positivi, e che l'ultimo membro della (3.30) è invariante per cambio di base. La condizione di normalizzazione impone che

$$e^{\Omega} = \frac{1}{\text{Tr}(e^{-\sum_i \lambda_i \hat{G}_i})}, \quad \text{ovvero} \quad \Omega = -\ln \text{Tr}(e^{-\sum_i \lambda_i \hat{G}_i}), \quad (3.31)$$

che è l'analogo quantistico della (3.16). Anche le (3.18) hanno l'analogo quantistico:

$$\frac{\partial \Omega}{\partial \lambda_i} = \frac{\partial}{\partial \lambda_i} \left(-\ln \sum_{\alpha} e^{-\mathcal{G}_{\alpha\alpha}} \right) = \frac{\sum_{\alpha} \frac{\partial \mathcal{G}_{\alpha\alpha}}{\partial \lambda_i} e^{-\mathcal{G}_{\alpha\alpha}}}{\sum_{\alpha} e^{-\mathcal{G}_{\alpha\alpha}}}.$$

D'altra parte, per definizione

$$\mathcal{G}_{\alpha\alpha} = \langle \alpha | \sum_{i=1}^n \lambda_i \hat{G}_i | \alpha \rangle, \quad \text{da cui} \quad \frac{\partial \mathcal{G}_{\alpha\alpha}}{\partial \lambda_i} = \langle \alpha | \hat{G}_i | \alpha \rangle.$$

Sostituendo nell'equazione precedente otteniamo

$$\frac{\partial \Omega}{\partial \lambda_i} = \frac{\sum_{\alpha} e^{-\mathcal{G}_{\alpha\alpha}} \langle \alpha | \hat{G}_i | \alpha \rangle}{\sum_{\alpha} e^{-\mathcal{G}_{\alpha\alpha}}} = e^{-\Omega} \text{Tr}(e^{-\sum_i \lambda_i \hat{G}_i} \hat{G}_i) = \text{Tr}(\hat{\rho} \hat{G}_i) = g_i. \quad (3.32)$$

3.7 Particelle identiche in meccanica quantistica

Nel caso quantistico l'indistinguibilità delle particelle implica che si incontrano solo funzioni d'onda completamente simmetriche (statistica di Bose-Einstein) o completamente antisimmetriche (statistica di Fermi-Dirac) per scambio di particelle. Quindi anche i vettori della base dello spazio di Hilbert che usiamo per costruire la matrice densità dovranno essere o completamente simmetrici, o completamente antisimmetrici, per scambio di particelle.

3.8 Funzione di partizione

In molti testi di fisica statistica si preferisce usare la funzione di partizione Z anziché il moltiplicatore di Lagrange Ω relativo alla condizione di normalizzazione. Le due quantità sono legate dalle relazioni

$$Z = e^{-\Omega}, \quad \Omega = -\ln Z. \quad (3.33)$$

In statistica classica abbiamo così

$$Z = \int e^{-\sum_{i=1}^n \lambda_i G_i(q,p)} dq dp, \quad (3.34)$$

mentre in statistica quantistica abbiamo

$$Z = \text{Tr}(e^{-\sum_i \lambda_i \hat{G}_i}). \quad (3.35)$$

La (3.15) può quindi essere riscritta in termini di Z come

$$f(q, p) = \frac{1}{Z} e^{-\sum_{i=1}^n \lambda_i G_i(q, p)}, \quad (3.36)$$

mentre la (3.30) può essere riscritta

$$\hat{\rho} = \frac{1}{Z} e^{-\sum_{i=1}^n \lambda_i \hat{G}_i(q, p)}. \quad (3.37)$$

Queste ultime due equazioni rendono forse più immediata l'idea che Z è legata alla normalizzazione della probabilità. Sia in statistica classica che in statistica quantistica abbiamo, per il valor medio di una grandezza $\langle G_i \rangle$

$$\langle G_i \rangle = -\frac{\partial}{\partial \lambda_i} \ln Z. \quad (3.38)$$

Capitolo 4

Informazione mancante e entropia

4.1 Costanti del moto

Se le uniche grandezze fisiche di cui conosciamo il valore sono costanti del moto, ci aspettiamo che la funzione di distribuzione $f(q, p)$ nel caso classico, o la nostra densità $\hat{\rho}$ nel caso quantistico, siano costanti nel tempo. Le più importanti costanti del moto che conosciamo sono l'energia, la quantità di moto $\mathbf{P}$ e il momento angolare $\mathbf{L}$. Per quanto visto nel Capitolo 3 avremo quindi, nel caso classico, una funzione di distribuzione del tipo

$$f(q, p) = e^{\Omega - \beta H(q, p) - \boldsymbol{\sigma} \cdot \mathbf{P}(q, p) - \boldsymbol{\lambda} \cdot \mathbf{L}(q, p)}, \quad (4.1)$$

dove β è il moltiplicatore di Lagrange relativo all'energia, mentre $\boldsymbol{\sigma}$ e $\boldsymbol{\lambda}$ sono i moltiplicatori di Lagrange vettoriali relativi alla quantità di moto e al momento angolare. Nel caso quantistico avremo una matrice densità

$$\hat{\rho} = e^{\Omega - \beta \hat{H} - \boldsymbol{\sigma} \cdot \hat{\mathbf{P}} - \boldsymbol{\lambda} \cdot \hat{\mathbf{L}}}, \quad (4.2)$$

dove $\hat{\mathbf{P}}$ è l'operatore quantità di moto totale del sistema e $\hat{\mathbf{L}}$ è l'operatore momento angolare totale. Se ci mettiamo in un sistema di riferimento solidale con il baricentro del nostro sistema, la quantità di moto complessiva sarà nulla. Se supponiamo poi che in questo sistema di riferimento il momento angolare totale del sistema sia nullo, abbiamo

$$f(q, p) = e^{\Omega - \beta H(q, p)} \quad \text{nel caso classico, e} \quad \hat{\rho} = e^{\Omega - \beta \hat{H}} \quad \text{nel caso quantistico.} \quad (4.3)$$

Non comparendo né $\boldsymbol{\sigma}$ né $\boldsymbol{\lambda}$ nelle espressioni, le derivate parziali di Ω rispetto ad una qualunque loro componente saranno infatti tutte nulle, e quindi, in base alle (3.18) e (3.32), saranno nulli i valori di aspettazione delle grandezze fisiche corrispondenti.

Nel caso in cui l'unica grandezza fisica su cui abbiamo informazione sia l'energia, rappresentato dalle (4.3), il moltiplicatore di Lagrange Ω nel caso classico si scrive

$$\Omega = -\ln \int e^{-\beta H(q, p)} dq dp \quad (+ \ln N!), \quad (4.4)$$

dove l'addendo tra parentesi compare solo se le particelle sono indistinguibili. A questa espressione per Ω corrisponde una funzione di partizione

$$Z = e^{-\Omega} = \left(\frac{1}{N!} \right) \int e^{-\beta H(q, p)} dq dp, \quad (4.5)$$

dove, ancora una volta, il fattore tra parentesi compare solo per particelle indistinguibili. Nel caso quantistico avremo invece

$$\Omega = -\ln \text{Tr}(e^{-\beta \hat{H}}) \quad \text{e} \quad Z = \text{Tr}(e^{-\beta \hat{H}}). \quad (4.6)$$

4.2 Gas perfetto e entropia

Per un paragone con la termodinamica classica consideriamo il sistema macroscopico più semplice che conosciamo, cioè il gas perfetto monoatomico. Questo è schematizzabile come un insieme di N punti materiali non interagenti tra loro, ognuno di massa m . Se escludiamo anche la presenza di forze esterne, ma teniamo conto del fatto che il gas è confinato in un contenitore di volume V , nel caso classico l'hamiltoniana si scrive

$$H(p, q) = \sum_{i=1}^N \frac{\mathbf{p}_i^2}{2m} + U^{\text{pot}}(q), \quad (4.7)$$

dove l'indice i corre su tutte le particelle del sistema, e $U^{\text{pot}}(q)$ è l'energia potenziale. Mentre la sommatoria dipende solo dalle quantità di moto, l'energia potenziale dipende solo dalle coordinate spaziali e, nel caso di un gas confinato in una “scatola” di volume V , ha la forma

$$U^{\text{pot}}(q) = \begin{cases} 0 & \text{se tutte le particelle sono all'interno del volume permesso;} \\ +\infty & \text{se almeno una particella è al di fuori del volume permesso.} \end{cases} \quad (4.8)$$

La funzione di distribuzione diventa così

$$f(q, p) = \exp\left(\Omega - \beta \sum_{i=1}^N \frac{\mathbf{p}_i^2}{2m} - \beta U^{\text{pot}}(q)\right). \quad (4.9)$$

Considerando per il momento le particelle come “uguali ma distinguibili” (considereremo il caso delle particelle indistinguibili più avanti) il moltiplicatore di Lagrange Ω vale

$$\begin{aligned} \Omega(\beta) &= -\ln \int d^3 q_1 \dots d^3 q_N e^{-\beta U^{\text{pot}}(q)} \int d^3 p_1 \dots d^3 p_N \exp\left(-\beta \sum_i \frac{\mathbf{p}_i^2}{2m}\right) \\ &= -\ln \left\{ V^N [\varphi(\beta)]^N \right\} = -N \ln V - N \ln \varphi(\beta), \end{aligned} \quad (4.10)$$

dove V è il volume occupato dal gas, e abbiamo posto

$$\varphi(\beta) = \int d^3 p e^{-\beta \mathbf{p}^2/2m} = \left(\frac{2\pi m}{\beta}\right)^{3/2}, \quad \text{ricordando che} \quad \int_{-\infty}^{+\infty} e^{-ax^2} dx = \sqrt{\frac{\pi}{a}}. \quad (4.11)$$

Analogamente, nel caso quantistico abbiamo

$$\Omega(\beta) = -\ln \left[\sum_{\alpha} \langle \alpha | \exp\left(-\beta \sum_i \frac{\hat{\mathbf{p}}_i^2}{2m} - \beta U^{\text{pot}}(q)\right) | \alpha \rangle \right] \quad (4.12)$$

dove gli $|\alpha\rangle$ costituiscono un insieme ortonormale e completo di autostati simultanei di tutti gli operatori $\hat{\mathbf{p}}_i^2$. Per sistemi di grandi dimensioni la traccia può essere sostituita da un integrale su uno “spazio delle fasi quantizzato” con celle di volume h^{3N} , quindi abbiamo

$$\begin{aligned}\Omega(\beta) &= -\ln \left[\frac{V^N}{h^{3N}} \int \exp\left(-\beta \sum_i \frac{\hat{\mathbf{p}}_i^2}{2m}\right) d^3 p_1 \dots d^3 p_N \right] \\ &= -\ln \left\{ \frac{V^N}{h^{3N}} [\varphi(\beta)]^N \right\} = -N \ln V - N \ln \varphi(\beta) + \ln h^{3N},\end{aligned}\quad (4.13)$$

che differisce dalla (4.10) solo per la presenza del termine additivo $\ln h^{3N}$, differenza dovuta alla rinormalizzazione (1.25). L'informazione mancante per un gas perfetto diventa così

$$I = -r_2 \Omega + r_2 \beta U = r_2 N \ln V + r_2 N \ln \varphi(\beta) - 3Nr_2 \ln h + r_2 \beta U, \quad (4.14)$$

dove U è il valore di aspettazione dell'energia (da non confondere con l'energia potenziale U^{pot} della (4.7) e seguenti), e il termine $-3Nr_2 \ln h$ è presente solo nel caso quantistico. Per il valore di aspettazione dell'energia U abbiamo dalle (3.18), (4.10) e (4.11)

$$U = \frac{\partial \Omega}{\partial \beta} = N \frac{3}{2\beta}, \quad (4.15)$$

che, confrontato con il valore noto dalla termodinamica classica $U = \frac{3}{2}NkT$, ci porta a porre

$$\beta = \frac{1}{kT}, \quad (4.16)$$

dove k è la costante di Boltzmann. Sostituendo la (4.15) nella (4.14) otteniamo per l'informazione mancante

$$I = r_2 N \ln V + r_2 N \ln \varphi(\beta) + r_2 \frac{3}{2}N - 3Nr_2 \ln h. \quad (4.17)$$

Dalla trattazione statistica svolta fin qui possiamo ottenere l'equazione di stato del gas perfetto classico. Cominciamo notando che la (4.9) può essere fattorizzata in un prodotto di funzioni di distribuzione di particella singola, e che ognuna di queste funzioni può, a sua volta, essere fattorizzata in una parte spaziale e una parte di impulso

$$f(q, p) = \prod_{i=1}^N s(\mathbf{q}_i) g(\mathbf{p}_i), \quad \text{con} \quad s(\mathbf{q}_i) = \frac{1}{V} e^{-\beta U_i^{\text{pot}}(\mathbf{q}_i)}, \quad g(\mathbf{p}_i) = \left(\frac{\beta}{2\pi m} \right)^{3/2} e^{-\beta \mathbf{p}_i^2 / 2m}, \quad (4.18)$$

dove $U_i^{\text{pot}}(\mathbf{q}_i) = 0$ se la particella è dentro al volume permesso, $U_i^{\text{pot}}(\mathbf{q}_i) = +\infty$ se la particella ne è fuori. Consideriamo adesso una porzione ΔS della superficie della scatola che contiene il gas, e scegliamo l'asse x delle nostre coordinate perpendicolare a ΔS . Calcoliamo la probabilità che una particella con componente x dell'impulso nell'intervallo $(p_x, p_x + dp_x)$ urti ΔS nell'intervallo di tempo $(t, t + \Delta t)$. Essendo $v_x = p_x/m$, all'istante t la particella dovrà trovarsi in un elemento di volume pari a $\Delta V = \Delta S (p_x/m) \Delta t$, come mostrato in Fig. 4.1. La probabilità che questo avvenga è $\Delta W_S = \Delta V/V = \Delta S p_x \Delta t / (mV)$. Dobbiamo moltiplicare per la probabilità dW_p che la componente x

dell'impulso sia nell'intervallo $(p_x, p_x + dp_x)$, indipendentemente dai valori di p_y e di p_z . Dall'ultima delle (4.18) abbiamo

$$\begin{aligned} dW_p &= \left(\frac{\beta}{2\pi m}\right)^{3/2} e^{-\beta p_x^2/2m} dp_x \int_{-\infty}^{+\infty} e^{-\beta p_y^2/2m} dp_y \int_{-\infty}^{+\infty} e^{-\beta p_z^2/2m} dp_z \\ &= \sqrt{\frac{\beta}{2\pi m}} e^{-\beta p_x^2/2m} dp_x. \end{aligned} \quad (4.19)$$

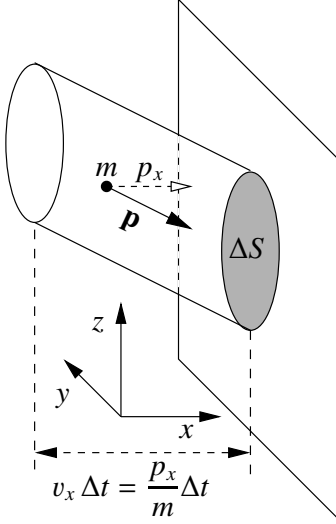


Figura 4.1 Volume in cui deve trovarsi all'istante t una particella di impulso $\mathbf{p}$ per urtare la superficie ΔS nell'intervallo $(t, t + \Delta t)$.

Il numero di particelle con impulso nell'intervallo $(p_x, p_x + dp_x)$ che urtano contro ΔS durante l'intervallo $(t, t + \Delta t)$ si ottiene moltiplicando il numero totale di particelle N per $\Delta W_S dW_p$

$$dN = N \Delta W_S dW_p = N \frac{\Delta S p_x \Delta t}{m V} \sqrt{\frac{\beta}{2\pi m}} e^{-\beta p_x^2/2m} dp_x. \quad (4.20)$$

Per urti perfettamente elastici, ogni particella subisce una variazione di impulso $\Delta \mathbf{p} = -2 p_x \hat{\mathbf{x}}$. Per urti non perfettamente elastici, ma in presenza di equilibrio termico, $-2 p_x \hat{\mathbf{x}}$ sarà comunque la variazione media dell'impulso. Per il terzo principio della dinamica, questo corrisponde a una forza media compressiva esercitata sulla parete

$$\begin{aligned} dF_x &= dN \frac{2p_x}{\Delta t} = 2p_x N \frac{\Delta S p_x}{m V} \sqrt{\frac{\beta}{2\pi m}} e^{-\beta p_x^2/2m} dp_x \\ &= \frac{2N\Delta S}{m V} \sqrt{\frac{\beta}{2\pi m}} p_x^2 e^{-\beta p_x^2/2m} dp_x. \end{aligned} \quad (4.21)$$

Integrando in dp_x tra 0 e $+\infty$, e ricordando che $p_x^2 e^{-\beta p_x^2/2m} = -2m \frac{\partial}{\partial \beta} e^{-\beta p_x^2/2m}$, otteniamo

$$F_x = \frac{2N\Delta S}{m V} \sqrt{\frac{\beta}{2\pi m}} \int_0^{+\infty} p_x^2 e^{-\beta p_x^2/2m} dp_x = \frac{2N\Delta S}{m V} \sqrt{\frac{\beta}{2\pi m}} \frac{m \sqrt{2\pi m}}{2\beta \sqrt{\beta}} = \frac{N\Delta S}{V} \frac{1}{\beta}. \quad (4.22)$$

Ricordando che la pressione è $p = F_x/\Delta S$, e inserendo la (4.16), abbiamo finalmente

$$pV = NkT. \quad (4.23)$$

Adesso passiamo alla relazione tra la nostra informazione mancante I e l'entropia S definita dalla termodinamica classica. Cominciamo considerando una trasformazione reversibile (molto lenta) isoterma (T , e quindi β , costanti) tra un volume iniziale V_i e un volume finale V_f . Secondo la termodinamica classica la variazione di entropia vale

$$\Delta S_{\text{isoterma}} = \int \frac{\delta Q}{T} = \frac{1}{T} \int_{V_i}^{V_f} p dV = \frac{1}{T} \int_{V_i}^{V_f} \frac{NkT}{V} dV = Nk \ln \left(\frac{V_f}{V_i} \right), \quad (4.24)$$

mentre dalla (4.17) abbiamo per la variazione di informazione mancante

$$\Delta I_{\text{isoterma}} = Nr_2 \ln \left(\frac{V_f}{V_i} \right). \quad (4.25)$$

Analogamente, la variazione di entropia per una trasformazione isocora reversibile da una temperatura iniziale T_i ad una temperatura finale T_f è, secondo la termodinamica classica,

$$\Delta S_{\text{isocora}} = \frac{3}{2}Nk \int_{T_i}^{T_f} \frac{dT}{T} = \frac{3}{2}Nk \ln \left(\frac{T_f}{T_i} \right), \quad (4.26)$$

per la stessa trasformazione la variazione di informazione mancante è

$$\Delta I_{\text{isocora}} = r_2 N \ln \left[\frac{\varphi(\beta_f)}{\varphi(\beta_i)} \right] = r_2 \frac{3}{2}N \ln \left(\frac{\beta_i}{\beta_f} \right) = r_2 \frac{3}{2}N \ln \left(\frac{T_f}{T_i} \right). \quad (4.27)$$

Questo ci porta a pensare che esista il legame $S = kI/r_2 = (k \ln 2)I$ tra entropia e informazione mancante. Poiché $k \ln 2$ è costante, vediamo che l'entropia introdotta in termodinamica non è altro che l'informazione mancante misurata in altre unità di misura, anziché in bit: un bit di informazione mancante corrisponde, in termini di entropia, a $(k \ln 2)$ J/K. Partendo dalla (4.17), l'entropia di un gas perfetto può essere scritta

$$S = -k\Omega + k\beta U = kN \ln V + kN \ln \varphi(\beta) + k\frac{3}{2}N - 3Nk \ln h \quad (4.28)$$

dove, ancora, l'ultimo addendo $-3Nk \ln h$ è presente solo nel caso quantistico.

Per un gas perfetto classico monoatomico di particelle identiche avremo

$$S_{\text{indist}} = -k\Omega_{\text{indist}} + k\beta U = kN \ln V + kN \ln \varphi(\beta) + k\frac{3}{2}N - 3Nk \ln h - k \ln N! \quad (4.29)$$

4.3 Teorema di equipartizione dell'energia

Consideriamo un sistema macroscopico di N particelle per il quale una certa coordinata canonica ξ compare nell'hamiltoniana in forma quadratica. La ξ può essere indifferentemente una coordinata generalizzata (una q) o un momento coniugato (una p), e non importa se nell'hamiltoniana anche altre coordinate canoniche (eventualmente tutte) compaiono in forma quadratica. Quello che importa è che l'hamiltoniana del sistema $H(q, p)$ possa essere scomposta nella somma

$$H(q, p) = H'(q', p') + A(\xi - \xi_0)^2, \quad (4.30)$$

dove A e ξ_0 sono costanti opportune (con ξ_0 che può anche essere nulla), mentre $\{q', p'\}$ è l'insieme di tutte le coordinate canoniche del nostro spazio delle fasi esclusa ξ . Combinando le (4.30) e (4.4) otteniamo per il moltiplicatore di Lagrange Ω

$$\begin{aligned} \Omega &= -\ln \int e^{-\beta H(q, p)} dq dp (+ \ln N!) = -\ln \int e^{-\beta [H'(q', p') + A(\xi - \xi_0)^2]} dq' dp' d\xi (+ \ln N!) \\ &= -\ln \int e^{-\beta H'(q', p')} dq' dp' (+ \ln N!) - \ln \int_{-\infty}^{+\infty} e^{-\beta A(\xi - \xi_0)^2} d\xi \\ &= \Omega' - \ln \int_{-\infty}^{+\infty} e^{-\beta A \eta^2} d\eta = \Omega' - \ln \sqrt{\frac{\pi}{A\beta}} = \Omega' - \frac{1}{2} \ln \frac{\pi}{A\beta}, \end{aligned} \quad (4.31)$$

dove Ω' è il moltiplicatore di Lagrange che avremmo in assenza di ξ , abbiamo introdotto la variabile $\eta = \xi - \xi_0$, e abbiamo usato l'ultima delle (4.11) per il calcolo dell'integrale. Per il valore di aspettazione dell'energia abbiamo così

$$U = \frac{\partial \Omega}{\partial \beta} = \frac{\partial \Omega'}{\partial \beta} - \frac{1}{2} \frac{\partial}{\partial \beta} \ln \frac{\pi}{A\beta} = U' + \frac{1}{2\beta} = U' + \frac{1}{2} kT, \quad (4.32)$$

dove U' è il valore che avremmo se la coordinata ξ non comparisse nell'hamiltoniana. Abbiamo ottenuto il *teorema di equipartizione dell'energia*: in statistica classica ogni coordinata canonica (che sia coordinata spaziale o momento coniugato) che compare in forma quadratica nell'espressione dell'hamiltoniana porta un contributo $\frac{1}{2} kT$ all'energia di aspettazione totale del sistema. Per esempio, l'hamiltoniana di un gas perfetto monoatomico è data dalla (4.7), in cui ogni componente cartesiana della quantità di moto di ogni particella compare in forma quadratica. La (4.15) è quindi una conseguenza del teorema di equipartizione dell'energia.

Come ulteriore esempio consideriamo un solido cristallino schematizzabile come un sistema di N atomi ognuno di massa m e oscillante attorno alla sua posizione di equilibrio nel reticolo. Ogni atomo del cristallo può così essere considerato, in buona approssimazione, come un oscillatore armonico tridimensionale. Per lo i -esimo oscillatore l'energia potenziale è $V_i = C_x(x_i - x_i^0)^2 + C_y(y_i - y_i^0)^2 + C_z(z_i - z_i^0)^2$, dove x_i , y_i e z_i sono le coordinate dell'atomo, mentre x_i^0 , y_i^0 e z_i^0 sono le coordinate della sua posizione di equilibrio. Abbiamo ipotizzato tre costanti di Hooke diverse C_x , C_y e C_z . L'energia cinetica sarà invece $K_i = (p_{ix}^2 + p_{iy}^2 + p_{iz}^2)/2m$. Abbiamo così $3N$ coordinate spaziali e $3N$ momenti coniugati che compaiono tutti in forma quadratica nell'hamiltoniana, portando a un valore di aspettazione dell'energia

$$U = 6N \frac{1}{2} kT = 3NkT, \quad (4.33)$$

in accordo con la legge scoperta empiricamente nel 1819 da Pierre R. Dulong e Alexis T. Petit, che attribuisce il valore $3R = 3N_A k$ al calore molare dei solidi. In realtà, Dulong e Petit non conoscevano il valore, e nemmeno l'esistenza, della costante R . Misurando il calore specifico di alcune sostanze avevano però trovato che questo era più piccolo per sostanze di maggiore "peso atomico". Avevano poi notato che se i calori specifici veniva moltiplicato per i corrispondenti "pesi atomici" si ottenevano valori in buona approssimazione uguali tra loro. A temperatura ambiente questa legge è rispettata con buona precisione dai 13 elementi di cui Dulong e Petit avevano originariamente misurato il calore specifico (S, Fe, Co, Ni, Cu, Zn, Ag, Sn, Te, Pt, Au, Pb, Bi), oltre che da molti altri solidi di struttura cristallina relativamente semplice. La legge di Dulong e Petit non è invece valida quando diventano importanti fenomeni quantistici. Questo accade, per esempio, già a temperatura ambiente nel caso di di atomi leggeri fortemente legati tra loro, e per tutti i solidi a temperature sufficientemente basse.

E' importante notare che il teorema di equipartizione attribuisce un'energia pari a $\frac{1}{2} kT$ non a ogni grado di libertà, come talvolta viene erroneamente enunciato, ma a ogni coordinata canonica, sia coordinata spaziale generalizzata che momento coniugato, che compaia in forma quadratica nell'espressione dell'hamiltoniana.

Capitolo 5

Insieme macrocanonico

5.1 Introduzione

Finora abbiamo sempre considerato noto e costante il numero N delle particelle che costituiscono il sistema sotto esame. Non è però sempre così. A parte il problema pratico di determinare dei numeri dell'ordine del numero di Avogadro, N può non essere noto a priori o variare nel tempo. Per esempio in presenza di reazioni chimiche, che fanno “sparire” certe molecole facendone “apparire” altre, o nel caso di un “gas” di fotoni, che possono essere emessi e assorbiti.

5.2 Caso classico

5.2.1 Numero non noto, o variabile, di particelle identiche

Supponiamo di avere un sistema di particelle identiche di cui non conosciamo il numero esatto N , ma solo una sua stima, o valore medio, $\bar{N}$. A priori non escludiamo nessuna possibilità $0 \leq N < \infty$. Non avendo a disposizione uno spazio delle fasi con numero non definito delle particelle, dobbiamo lavorare simultaneamente su tutti gli spazi delle fasi con $0 \leq N < \infty$ (insieme *macrocanonico*, o *gran canonico*), ovviamente tutti relativi allo stesso tipo di particelle. Per ogni spazio delle fasi con $N > 0$ avremo una funzione di distribuzione della probabilità

$$f_N(\{\mathbf{q}_i\}, \{\mathbf{p}_i\}), \quad \text{con} \quad 1 \leq i \leq N, \quad (5.1)$$

mentre per $N = 0$ avremo una f_0 costante. Imponiamo le condizioni di normalizzazione

$$\frac{1}{N!} \int f_N(q, p) \, dq \, dp = 1 \quad \forall N, \quad (5.2)$$

e, per ogni N , chiamiamo W_N la probabilità che il sistema abbia N particelle. Dovrà essere

$$W_N \geq 0 \quad \forall N, \quad \text{e} \quad \sum_{N=0}^{\infty} W_N = 1. \quad (5.3)$$

Avremo così bisogno di un insieme infinito $\{f_N\}$ di funzioni di distribuzione sui singoli spazi delle fasi, accompagnato dall'insieme delle probabilità $\{W_N\}$ di avere proprio N particelle. Ogni

grandezza fisica G , dipendendo dalle coordinate e dai momenti delle particelle del sistema, non è più rappresentata da una singola funzione ma da un insieme di funzioni del tipo

$$G_{\text{mc}} \equiv \{G_0, G_1(\mathbf{q}_1, \mathbf{p}_1), G_2(\mathbf{q}_1, \mathbf{p}_1, \mathbf{q}_2, \mathbf{p}_2), \dots, G_N(\mathbf{q}_1, \mathbf{p}_1, \mathbf{q}_2, \mathbf{p}_2, \dots, \mathbf{q}_N, \mathbf{p}_N), \dots\}, \quad (5.4)$$

per esempio, l'energia cinetica traslazionale T sarà scritta

$$T_{\text{mc}} \equiv \left\{0, \frac{\mathbf{p}_1^2}{2m}, \frac{\mathbf{p}_1^2}{2m} + \frac{\mathbf{p}_2^2}{2m}, \dots, \sum_{j=1}^N \frac{\mathbf{p}_j^2}{2m} \cdot \dots\right\}. \quad (5.5)$$

Il valore di aspettazione di G sarà

$$\langle G \rangle = \sum_{N=0}^{\infty} \frac{W_N}{N!} \int f_N(q, p) G_N(q, p) dq dp. \quad (5.6)$$

E' da notare che nella (5.6) ogni f_N compare sempre moltiplicata per la W_N corrispondente, cosa che ci permetterà una semplificazione di notazione nel paragrafo 5.2.2.

Infine introduciamo una nuova grandezza fisica, il *numero delle particelle* $\mathcal{N}$, definita da

$$\mathcal{N} = \{0, 1, 2, \dots, N, \dots\}, \quad \text{con valor medio} \quad \langle \mathcal{N} \rangle = \sum_{N=0}^{\infty} W_N N = \bar{N}, \quad (5.7)$$

pari al valore di aspettazione del numero delle particelle del sistema.

5.2.2 Entropia nell'insieme macrocanonico

Se le particelle fossero esattamente N , l'entropia sarebbe, come abbiamo visto

$$S_N = -k \frac{1}{N!} \int f_N \ln f_N dq dp. \quad (5.8)$$

Se mediamo la (5.8) sui possibili numeri di particelle N , ognuno con probabilità W_N , otteniamo

$$\langle S \rangle = -k \sum_N \frac{W_N}{N!} \int f_N \ln f_N dq dp, \quad (5.9)$$

a questo valore va ancora aggiunta l'entropia corrispondente al difetto di informazione relativo alla scelta tra i possibili valori di N , cioè

$$S(N) = -k \sum_N W_N \ln W_N. \quad (5.10)$$

Complessivamente, nell'insieme macrocanonico l'entropia vale quindi

$$S = -k \sum_N W_N \left[\ln W_N + \frac{1}{N!} \int f_N \ln f_N dq dp \right], \quad (5.11)$$

e questa espressione può essere riscritta, ricordando che $\frac{1}{N!} \int f_N dq dp = 1$ per ogni N ,

$$\begin{aligned}
 S &= -k \sum_N W_N \left[\frac{1}{N!} \int (\ln W_N) f_N dq dp + \frac{1}{N!} \int f_N \ln f_N dq dp \right] \\
 &= -k \sum_N \frac{W_N}{N!} \int f_N \ln(W_N f_N) dq dp \\
 &= -k \sum_N \frac{1}{N!} \int W_N f_N \ln(W_N f_N) dq dp.
 \end{aligned} \tag{5.12}$$

Cerchiamo adesso il massimo vincolato di S , considerando noti il valor medio del numero delle particelle $\bar{N}$ e il valor medio dell'energia U . Come abbiamo già notato sotto alla (5.6), ogni f_N compare sempre moltiplicata per la W_N corrispondente, così che, per “stenografia”, ci conviene definire le nuove funzioni

$$g_N(q, p) = W_N f_N(q, p), \quad \text{con} \quad \frac{1}{N!} \int g_N(q, p) dq dp = W_N. \tag{5.13}$$

Sostituendo le $g_N(q, p)$ nella (5.12) otteniamo

$$S = -k \sum_N \frac{1}{N!} \int g_N \ln g_N dq dp. \tag{5.14}$$

Dobbiamo cercare l'insieme $\{g_N\}$ che corrisponde al massimo vincolato di S con le condizioni

$$1 = \sum_N W_N = \sum_N \frac{1}{N!} \int g_N dq dp \tag{5.15}$$

$$\bar{N} = \sum_N W_N N = \sum_N \frac{N}{N!} \int g_N dq dp \tag{5.16}$$

$$U = \sum_N \frac{1}{N!} \int H_N g_N dq dp. \tag{5.17}$$

Usando il solito metodo dei moltiplicatori di Lagrange (con un unico moltiplicatore relativo alla normalizzazione!), cerchiamo il massimo non vincolato dell'espressione

$$\begin{aligned}
 S' &= -k \sum_N \frac{1}{N!} \int g_N \ln g_N dq dp + k(\Omega + 1) \left(\sum_N \frac{1}{N!} \int g_N dq dp - 1 \right) \\
 &\quad - k\alpha \left(\sum_N \frac{N}{N!} \int g_N dq dp - \bar{N} \right) - k\beta \left(\sum_N \frac{1}{N!} \int H_N g_N dq dp - U \right)
 \end{aligned} \tag{5.18}$$

dove $-k(\Omega + 1)$ è il moltiplicatore di Lagrange relativo alla normalizzazione, $k\alpha$ quello relativo al numero delle particelle, e $k\beta$ quello relativo all'energia. Deve essere nulla la variazione al primo

ordine

$$\begin{aligned}
 0 &= \delta S' = -k \sum_N \frac{1}{N!} \left[\int \delta g_N \ln g_N dq dp + \int \delta g_N dq dp - (\Omega + 1) \int \delta g_N dq dp \right. \\
 &\quad \left. + \alpha N \int \delta g_N dq dp + \beta \int H_N \delta g_N dq dp \right] \\
 &= -k \sum_N \frac{1}{N!} \int \delta g_N (\ln g_N - \Omega + \alpha N + \beta H_N) dq dp.
 \end{aligned} \tag{5.19}$$

Data l'arbitrarietà delle variazioni δg_N , la (5.19) implica che, per ogni N , si abbia

$$\ln g_N - \Omega + \alpha N + \beta H_N = 0, \quad \text{ovvero} \quad g_N = e^{\Omega - \alpha N - \beta H_N}, \tag{5.20}$$

con Ω tale che

$$1 = \sum_N \frac{1}{N!} \int g_N dq dp = e^\Omega \sum_N \frac{e^{-\alpha N}}{N!} \int e^{-\beta H_N} dq dp.$$

Otteniamo così per Ω

$$\Omega = -\ln \left(\sum_N \frac{e^{-\alpha N}}{N!} \int e^{-\beta H_N} dq dp \right) = -\ln \left(\sum_N e^{-\alpha N - \Omega_N} \right), \tag{5.21}$$

mentre la funzione di partizione Z diventa

$$Z = e^{-\Omega} = \sum_N \frac{e^{-\alpha N}}{N!} \int e^{-\beta H_N} dq dp = \sum_N e^{-\alpha N} Z_N. \tag{5.22}$$

Nelle (5.21) e (5.22) le grandezze Ω_N e Z_N corrispondono, rispettivamente, al moltiplicatore di Lagrange relativo alla condizione di normalizzazione (4.4) e alla funzione di partizione (4.5) per un gas con un numero fisso N di particelle identiche. Avremo poi

$$\frac{\partial \Omega}{\partial \alpha} = \bar{N}, \quad \text{e la solita} \quad \frac{\partial \Omega}{\partial \beta} = U. \tag{5.23}$$

5.2.3 Gas perfetto monoatomico classico nell'insieme macrocanonico

Le espressioni diventano particolarmente semplici in assenza di interazione tra le particelle. Se consideriamo un gas perfetto monoatomico classico, l'hamiltoniana sarà

$$H = \{H_N\} = \left\{ \sum_{i=1}^N \frac{\mathbf{p}_i^2}{2m} + U_N^{\text{pot}}(\mathbf{q}_1 \dots \mathbf{q}_N) \right\}, \tag{5.24}$$

con le singole $U_N^{\text{pot}}(\mathbf{q}_1 \dots \mathbf{q}_N)$, corrispondenti ognuna all'energia potenziale di un gas perfetto di N particelle, analoghe alla (4.8). Avremo

$$\Omega = -\ln \left[\sum_{N=0}^{\infty} \frac{e^{-\alpha N}}{N!} \left(\int e^{-\beta p^2/2m} d^3 p \cdot V \right)^N \right] = -\ln \left[\sum_{N=0}^{\infty} \frac{e^{-\alpha N}}{N!} V^N \left(\frac{2\pi m}{\beta} \right)^{\frac{3}{2}N} \right], \tag{5.25}$$

dove V è il volume occupato dal gas.

Dalla relazione $\sum_{N=0}^{\infty} \frac{x^N}{N!} = e^x$ segue che $\Omega = -\ln \left\{ \exp \left[e^{-\alpha} V \left(\frac{2\pi m}{\beta} \right)^{\frac{3}{2}} \right] \right\}$, ovvero

$$\Omega = -V \left(\frac{2\pi m}{\beta} \right)^{\frac{3}{2}} e^{-\alpha}. \quad (5.26)$$

Dalla prima delle (5.23) segue

$$\bar{N} = \frac{\partial \Omega}{\partial \alpha} = V \left(\frac{2\pi m}{\beta} \right)^{\frac{3}{2}} e^{-\alpha} = -\Omega, \quad (5.27)$$

e da qui

$$e^{-\alpha} = \frac{\bar{N}}{V} \left(\frac{\beta}{2\pi m} \right)^{\frac{3}{2}}, \quad \alpha = -\ln \frac{\bar{N}}{V} + \ln \left(\frac{2\pi m}{\beta} \right)^{\frac{3}{2}}. \quad (5.28)$$

Le componenti della funzione di distribuzione sono

$$g_N = e^{-\bar{N}} \left[\frac{\bar{N}}{V} \left(\frac{\beta}{2\pi m} \right)^{\frac{3}{2}} \right]^N e^{-\beta \sum \frac{p_i^2}{2m}} e^{-\beta U_N^{\text{pot}}}. \quad (5.29)$$

La (5.29) ci permette di calcolare le probabilità W_N di avere esattamente N particelle per un gas perfetto con numero medio di particelle $\bar{N}$

$$\begin{aligned} W_N &= \frac{1}{N!} \int g_N \, dq \, dp = \frac{1}{N!} e^{-\bar{N}} \left[\frac{\bar{N}}{V} \left(\frac{\beta}{2\pi m} \right)^{\frac{3}{2}} \right]^N \int dq \, e^{-\beta U_N^{\text{pot}}} \, dp \, e^{-\beta \sum \frac{p_i^2}{2m}} \\ &= e^{-\bar{N}} \frac{(\bar{N})^N}{N!}, \end{aligned} \quad (5.30)$$

che è la distribuzione di Poisson con valore medio $\bar{N}$. Da qui possiamo ottenere

$$f_N = \frac{g_N}{W_N} = N! \left[\frac{1}{V} \left(\frac{\beta}{2\pi m} \right)^{\frac{3}{2}} \right]^N e^{-\beta \sum \frac{p_i^2}{2m}} e^{-\beta U_N^{\text{pot}}}, \quad (5.31)$$

che è la normale funzione di distribuzione di un gas perfetto monoatomico con un numero fisso N di particelle indistinguibili.

Avremo inoltre per il valor medio dell'energia del gas perfetto monoatomico U

$$U = \frac{\partial \Omega}{\partial \beta} = \frac{3}{2} V \left(\frac{2\pi m}{\beta} \right)^{\frac{3}{2}} \frac{1}{\beta} e^{-\alpha} = \frac{3}{2} V \left(\frac{2\pi m}{\beta} \right)^{\frac{3}{2}} \frac{1}{\beta} \frac{\bar{N}}{V} \left(\frac{\beta}{2\pi m} \right)^{\frac{3}{2}} = \frac{3}{2} \frac{\bar{N}}{\beta} = \frac{3}{2} \bar{N} kT. \quad (5.32)$$

Per l'entropia otteniamo infine, utilizzando le (5.27) e (5.28),

$$S = -k\Omega + k\beta U + k\alpha \bar{N} = k\bar{N} \left[\ln \left(\frac{2\pi m}{\beta} \right)^{\frac{3}{2}} + \frac{5}{2} - \ln \frac{\bar{N}}{V} \right]. \quad (5.33)$$

5.2.4 Ampiezza delle fluttuazioni in statistica classica

Consideriamo un sistema di N particelle, non importa se distinguibili o indistinguibili, nel caso in cui l'unica grandezza nota a priori è l'energia. Dalla (3.24) abbiamo

$$\Omega = -\ln \int e^{-\beta H(q,p)} dq dp \quad (+ \ln N!), \quad (5.34)$$

dove il termine tra parentesi compare solo nel caso di particelle indistinguibili. Per il valor medio dell'energia abbiamo in ogni caso (N non dipende da β)

$$\langle U \rangle = \frac{\partial \Omega}{\partial \beta} = \frac{\int H e^{-\beta H} dq dp}{\int e^{-\beta H} dq dp}. \quad (5.35)$$

Vogliamo calcolare quanto l'energia del sistema fluttui attorno al suo valore medio $\langle U \rangle$. Usando, come al solito, le parentesi triangolari per denotare i valori medi, definiamo la *fluttuazione quadratica media* dell'energia come

$$\langle (\Delta U)^2 \rangle = \langle (U - \langle U \rangle)^2 \rangle = \langle U^2 - 2U\langle U \rangle + \langle U \rangle^2 \rangle = \langle U^2 \rangle - \langle U \rangle^2. \quad (5.36)$$

D'altra parte, se calcoliamo la derivata seconda di Ω rispetto a β vediamo che

$$\begin{aligned} \frac{\partial^2 \Omega}{\partial \beta^2} &= \frac{\partial}{\partial \beta} \frac{\int H e^{-\beta H} dq dp}{\int e^{-\beta H} dq dp} = \frac{-\int H^2 e^{-\beta H} dq dp \int e^{-\beta H} dq dp + \int H e^{-\beta H} dq dp \int H e^{-\beta H} dq dp}{\left[\int e^{-\beta H} dq dp \right]^2} \\ &= -\frac{\int H^2 e^{-\beta H} dq dp}{\int e^{-\beta H} dq dp} + \left[\frac{\int H e^{-\beta H} dq dp}{\int e^{-\beta H} dq dp} \right]^2 = -\langle U^2 \rangle + \langle U \rangle^2 = -\langle (\Delta U)^2 \rangle. \end{aligned} \quad (5.37)$$

Possiamo quindi scrivere

$$\langle (\Delta U)^2 \rangle = -\frac{\partial^2 \Omega}{\partial \beta^2} = -\frac{\partial \langle U \rangle}{\partial \beta} = -\frac{\partial \langle U \rangle}{\partial T} \frac{1}{\frac{\partial \beta}{\partial T}} = kT^2 \frac{\partial \langle U \rangle}{\partial T} = kT^2 C_V, \quad (5.38)$$

dove C_V è il calore specifico a volume costante del sistema. La fluttuazione quadratica media relativa è definita come

$$\frac{\langle (\Delta U)^2 \rangle}{\langle U \rangle^2} = \frac{kT^2 C_V}{\langle U \rangle^2}. \quad (5.39)$$

Nel caso di un gas perfetto monoatomico dalla (4.10) abbiamo

$$\begin{aligned} \Omega &= -N \ln V - N \ln \left(\frac{2\pi m}{\beta} \right)^{3/2}, \\ \langle U \rangle &= \frac{\partial \Omega}{\partial \beta} = \frac{3}{2} N \frac{1}{\beta} = \frac{3}{2} N kT, \\ \langle (\Delta U)^2 \rangle &= -\frac{\partial^2 \Omega}{\partial \beta^2} = \frac{3}{2} N \frac{1}{\beta^2} = \frac{3}{2} N (kT)^2, \end{aligned} \quad (5.40)$$

da cui otteniamo per la fluttuazione quadratica media relativa dell'energia

$$\frac{\langle(\Delta U)^2\rangle}{\langle U\rangle^2} = \frac{3}{2}N(kT)^2 \left(\frac{2}{3NkT}\right)^2 = \frac{2}{3N},$$

la cui radice quadrata vale

$$\sqrt{\frac{\langle(\Delta U)^2\rangle}{\langle U\rangle^2}} = \sqrt{\frac{2}{3N}}. \quad (5.41)$$

Quindi le fluttuazioni relative sono inversamente proporzionali a $\sqrt{N}$.

5.3 Caso quantistico

5.3.1 Gas perfetto quantistico, impostazione del problema

L'indistinguibilità delle particelle ha come conseguenza che lo stato di un sistema costituito da particelle identiche deve restare lo stesso per ogni scambio di due particelle tra loro. In meccanica quantistica questo implica che, se scambiamo tra loro due particelle, la funzione d'onda ψ può solo cambiare per un fattore di fase $e^{i\varphi}$. D'altra parte, se scambiamo due volte le stesse particelle, la ψ deve restare immutata, quindi $e^{i2\varphi} = 1$, e può essere solo $e^{i\varphi} = \pm 1$. Dopo lo scambio di due particelle tra loro ci sono quindi due possibilità: i) la ψ rimane invariata (funzione d'onda simmetrica), oppure ii) la ψ cambia segno (funzione d'onda antisimmetrica). Queste due possibilità corrispondono alle due diverse statistiche della meccanica quantistica:

- i. sistemi di particelle rappresentati da funzioni d'onda simmetriche per scambio di particelle obbediscono alla statistica di Bose-Einstein;
- ii. sistemi di particelle rappresentati da funzioni d'onda antisimmetriche obbediscono alla statistica di Fermi-Dirac.

In ambedue i casi è conveniente per motivi di calcolo, come vedremo, effettuare i conti statistici nell'insieme macrocanonico, cioè fissare un valore di aspettazione per il numero $\bar{N}$ delle particelle, ma non vincolare N a un valore fisso.

Lo spazio di Hilbert relativo all'insieme macrocanonico sarà la somma diretta di tutti gli spazi di Hilbert relativi a un sistema di N particelle identiche con $N = 0, 1, 2, \dots, \infty$. Cerchiamo adesso di costruirci una base per questo spazio di Hilbert.

Cominciamo considerando un'unica particella confinata nel volume che dovrà poi contenere l'intero sistema. Questa particella avrà uno spettro discreto di livelli energetici E_r , e una base possibile per lo spazio di Hilbert dei suoi stati possibili sarà data dall'insieme ortonormale

$$\{|r\rangle\} \equiv \{|E_r \alpha_r\rangle\}, \quad (5.42)$$

dove gli α_r sono gli altri numeri quantici eventualmente necessari per identificare lo stato, e più E_r possono avere lo stesso valore in caso di degenerazione.

Se supponiamo che il volume contenga più di una particella, ma che le particelle non interagiscano tra loro, una base completa e stazionaria per il nostro spazio di Hilbert è data dall'insieme dei vettori del tipo

$$|\{n_r\}\rangle = |n_0 E_0 \alpha_0\rangle |n_1 E_1 \alpha_1\rangle |n_2 E_2 \alpha_2\rangle \dots |n_r E_r \alpha_r\rangle \dots, \quad (5.43)$$

detti *stati di Fock* (da Влади́мир Алекса́ндрович Фо́к), dove n_r è il numero di particelle che si trovano nello stato $|r\rangle$. Questo è il motivo per cui i calcoli sono più facili nell'insieme macrocanonico: infatti qui sono possibili tutti gli insiemi $\{n_r\}$, senza il limite posto dalla condizione $\sum n_r = N$. Gli $|\{n_r\}\rangle$ sono autostati sia dell'hamiltoniana $\hat{H}$, con autovalori

$$E_{\{n_r\}} = \sum_r n_r E_r \quad (5.44)$$

che dell'operatore *numero delle particelle* $\hat{N}$, con autovalori

$$N_{\{n_r\}} = \sum_r n_r. \quad (5.45)$$

Per i bosoni lo stato è simmetrico per lo scambio di una qualunque coppia di particelle, questo non ha conseguenze sui numeri di occupazione degli stati di particella singola e tutti gli $|\{n_r\}\rangle$ sono permessi. Per i fermioni lo stato deve cambiare segno per ogni scambio di due particelle tra loro, e questo impone che ogni n_r possa assumere solo i valori 0 oppure 1.

Se introduciamo l'interazione tra le particelle, gli $|\{n_r\}\rangle$ restano una base ortonormale completa per lo spazio di Hilbert del nostro sistema, ma i singoli elementi della base, pur restando autostati di $\hat{N}$, non sono più autostati di $\hat{H}$.

La matrice densità sarà

$$\hat{\rho} = e^{\Omega - \beta \hat{H} - \alpha \hat{N}}, \quad \text{con} \quad \Omega(\beta, \alpha) = -\ln \text{Tr} \left(e^{-\beta \hat{H} - \alpha \hat{N}} \right). \quad (5.46)$$

Consideriamo adesso un gas perfetto quantistico confinato in un contenitore di volume V . Gli stati di particella singola $|r\rangle$ saranno autostati, oltre che dell'hamiltoniana di particella singola con autovalore E_r , dell'operatore *quantità di moto al quadrato* $\hat{\mathbf{p}}^2$ (ma non dell'operatore *quantità di moto* $\hat{\mathbf{p}}$) con autovalore p_r^2 , e avremo

$$E_r = \frac{p_r^2}{2m}. \quad (5.47)$$

A seconda della geometria del contenitore, i livelli energetici di particella singola potranno essere degeneri. In assenza di interazioni tra le particelle, gli stati del gas $|\{n_r\}\rangle$ saranno autostati simultanei dell'hamiltoniana $\hat{H}$ con autovalori

$$E_{\{n_r\}} = \sum_r n_r \frac{p_r^2}{2m} = \sum_r n_r E_r, \quad (5.48)$$

e dell'operatore numero delle particelle $\hat{N}$ con autovalori $\sum_r n_r$, dove, in tutte le somme, l'indice r corre su tutti gli stati di particella singola. Avremo così

$$\hat{H}|\{n_r\}\rangle = E_{\{n_r\}}|\{n_r\}\rangle = \left(\sum_r n_r \frac{p_r^2}{2m} \right) |\{n_r\}\rangle = \left(\sum_r n_r E_r \right) |\{n_r\}\rangle. \quad (5.49)$$

Avremo poi

$$\begin{aligned}
 \Omega(\beta, \alpha) &= -\ln \sum_{\{n_r\}} \langle \{n_r\} | e^{-\beta \hat{H} - \alpha \hat{N}} | \{n_r\} \rangle = -\ln \sum_{\{n_r\}} e^{-\beta \sum_r n_r E_r - \alpha \sum_r n_r} \\
 &= -\ln \sum_{\{n_r\}} e^{-\sum_r n_r (\beta E_r + \alpha)} = -\ln \sum_{\{n_r\}} \prod_r e^{-n_r (\beta E_r + \alpha)} \\
 &= -\ln \left[\sum_{n_0} \sum_{n_1} \cdots \sum_{n_r} \cdots e^{-n_0 (\beta E_0 + \alpha)} e^{-n_1 (\beta E_1 + \alpha)} \cdots e^{-n_r (\beta E_r + \alpha)} \cdots \right] \\
 &= -\ln \left[\sum_{n_0} e^{-n_0 (\beta E_0 + \alpha)} \sum_{n_1} e^{-n_1 (\beta E_1 + \alpha)} \cdots \sum_{n_r} e^{-n_r (\beta E_r + \alpha)} \cdots \right] \\
 &= -\ln \prod_r \sum_n e^{-n (\beta E_r + \alpha)}, \tag{5.50}
 \end{aligned}$$

dove il prodotto su r corre su tutti gli stati di particella singola, e la somma su n corre su tutti i numeri di occupazione permessi per uno stato di particella singola. Possiamo poi riscrivere

$$\Omega(\beta, \alpha) = - \sum_r \ln \sum_n e^{-n (\beta E_r + \alpha)}. \tag{5.51}$$

In base alla relazione $\sum_{n=0}^{\infty} q^n = \frac{1}{1-q}$, valida per $q < 1$, abbiamo per i bosoni

$$\sum_{n=0}^{\infty} e^{-n (\beta E_r + \alpha)} = \frac{1}{1 - e^{-\beta E_r - \alpha}}, \tag{5.52}$$

mentre per i fermioni abbiamo

$$\sum_{n=0}^1 e^{-n (\beta E_r + \alpha)} = 1 + e^{-\beta E_r - \alpha}. \tag{5.53}$$

Possiamo compattare l'espressione di Ω per bosoni e fermioni nella forma

$$\Omega(\beta, \alpha) = - \sum_r \ln \left(1 - \varepsilon e^{-\beta E_r - \alpha} \right)^{-\varepsilon}, \tag{5.54}$$

con $\varepsilon = +1$ per i bosoni e $\varepsilon = -1$ per i fermioni.

Per la matrice densità abbiamo

$$\begin{aligned}
 \hat{\rho} &= e^{\Omega - \beta \hat{H} - \alpha \hat{N}} = \left[\prod_r \left(1 - \varepsilon e^{-\beta E_r - \alpha} \right)^{\varepsilon} \right] e^{-\beta \hat{H} - \alpha \hat{N}} \\
 &= \left[\prod_r \left(1 - \varepsilon e^{-\beta E_r - \alpha} \right)^{\varepsilon} \right] \sum_{\{n_r\}} e^{-\sum_r (\beta E_r + \alpha) n_r} | \{n_r\} \rangle \langle \{n_r\} | \tag{5.55}
 \end{aligned}$$

5.3.2 Potenziale chimico

Spesso, invece del moltiplicatore di Lagrange relativo al numero delle particelle α si preferisce usare il *potenziale chimico* μ , che ha le dimensioni di un'energia ed è definito dalla relazione

$$\mu = -\frac{\alpha}{\beta}, \quad (5.56)$$

così che il termine $-\beta E_r - \alpha$ viene riscritto $-\beta(E_r - \mu)$. In questo modo, per esempio, la (5.55) si riscrive

$$\hat{\rho} = \left[\prod_r (1 - \varepsilon e^{-\beta(E_r - \mu)})^\varepsilon \right] \sum_{\{n_r\}} e^{-\sum_r \beta(E_r - \mu)n_r} |\{n_r\}\rangle \langle \{n_r\}| \quad (5.57)$$

Il concetto di potenziale chimico è stato introdotto nel 1876 da Josiah W. Gibbs. Se un sistema termodinamico è costituito da più *sostanze* (per esempio contiene atomi o molecole di tipo diverso), la sostanza i -esima sarà caratterizzata dal suo potenziale chimico μ_i . Il potenziale chimico è una funzione termodinamica di stato intensiva, legata alla capacità della sostanza di reagire con altre sostanze (reazioni chimiche), effettuare trasformazioni di fase, distribuirsi nello spazio (diffusione). Secondo l'equazione fondamentale di Gibbs i potenziali chimici μ_i sono legati all'energia interna del sistema U , alla sua entropia S , alla pressione p e alle *quantità di sostanza* (numeri di molecole, o di atomi ...) n_i dalla relazione

$$U = TS - pV + \sum_i \mu_i n_i. \quad (5.58)$$

Il altre parole, il potenziale chimico della sostanza i contenuta nel sistema termodinamico è pari alla variazione dell'energia interna del sistema ΔU che si avrebbe se venisse aggiunta una piccola quantità di sostanza Δn_i , a S e V costanti, divisa per la quantità di sostanza aggiunta:

$$\mu_i = \left. \frac{\partial U}{\partial n_i} \right|_{S, V, n_{j \neq i} \text{ costanti}}. \quad (5.59)$$

5.3.3 Numeri di occupazione

Dalla (5.57) otteniamo che la probabilità che un gas perfetto quantistico si trovi nello stato $|\{n_r\}\rangle = |n_0, n_1, \dots, n_r, \dots\rangle$ è

$$\begin{aligned} \langle n_0, n_1, \dots, n_r, \dots | \hat{\rho} | n_0, n_1, \dots, n_r, \dots \rangle &= \left[\prod_r (1 - \varepsilon e^{-\beta(E_r - \mu)})^\varepsilon \right] e^{-\sum_r \beta(E_r - \mu)n_r} \\ &= \left[\prod_r (1 - \varepsilon e^{-\beta(E_r - \mu)})^\varepsilon \right] \prod_r e^{-\beta(E_r - \mu)n_r} \\ &= \prod_r \frac{e^{-\beta(E_r - \mu)n_r}}{[1 - \varepsilon e^{-\beta(E_r - \mu)}]^{-\varepsilon}} \\ &= \prod_r W_r(n_r), \end{aligned} \quad (5.60)$$

dove il fattore $W_r(n_r)$ può essere interpretato come la probabilità che lo stato di particella singola $|r\rangle$ sia occupato da n_r particelle. Scritto in forma esplicita abbiamo

$$W_r(n) = \frac{e^{-n\beta(E_r - \mu)}}{[1 - \varepsilon e^{-\beta(E_r - \mu)}]^{-\varepsilon}} \quad (5.61)$$

dove abbiamo tralasciato l'indice r che compariva in n_r . Questa fattorizzazione della matrice densità (fattorizzazione sugli stati di particella singola, non sulle particelle!) può essere usata per calcolare i valori medi di tutte le grandezze di particella singola, che sono le uniche permesse per un gas perfetto. Per esempio, possiamo calcolare $\langle n_r \rangle$, il numero medio delle particelle che si trovano nello stato $|r\rangle$

$$\langle n_r \rangle = \sum_n W_r(n) n = \sum_n \frac{n e^{-n\beta(E_r-\mu)}}{[1 - \varepsilon e^{-\beta(E_r-\mu)}]^{-\varepsilon}} = [1 - \varepsilon e^{-\beta(E_r-\mu)}]^\varepsilon \sum_n n e^{-n\beta(E_r-\mu)}. \quad (5.62)$$

La somma su n vale per i fermioni

$$\sum_{n=0}^1 n e^{-n\beta(E_r-\mu)} = e^{-\beta(E_r-\mu)}, \quad (5.63)$$

mentre per i bosoni abbiamo

$$\sum_{n=0}^{\infty} n e^{-n\beta(E_r-\mu)} = \frac{e^{-\beta(E_r-\mu)}}{[1 - e^{-\beta(E_r-\mu)}]^2} \quad \text{in base alla relazione} \quad \sum_{n=0}^{\infty} n q^n = \frac{q}{(1-q)^2} \quad (q < 1) \quad (5.64)$$

Inserendo queste somme nella (5.62) abbiamo per i fermioni

$$\langle n_r \rangle^{\text{fermioni}} = [1 + e^{-\beta(E_r-\mu)}]^{-1} e^{-\beta(E_r-\mu)} = \frac{e^{-\beta(E_r-\mu)}}{1 + e^{-\beta(E_r-\mu)}} \quad (5.65)$$

e per i bosoni

$$\langle n_r \rangle^{\text{bosoni}} = [1 - e^{-\beta(E_r-\mu)}]^{-1} \frac{e^{-\beta(E_r-\mu)}}{[1 - e^{-\beta(E_r-\mu)}]^2} = \frac{e^{-\beta(E_r-\mu)}}{1 - e^{-\beta(E_r-\mu)}} \quad (5.66)$$

che possono essere compattate nell'espressione

$$\langle n_r \rangle = \frac{e^{-\beta(E_r-\mu)}}{1 - \varepsilon e^{-\beta(E_r-\mu)}}, \quad (5.67)$$

moltiplicando numeratore e denominatore per $e^{\beta(E_r-\mu)}$ otteniamo infine

$$\langle n_r \rangle = \frac{1}{e^{\beta(E_r-\mu)} - \varepsilon}. \quad (5.68)$$

5.3.4 Fluttuazioni dei numeri di occupazione

Per le fluttuazioni dei numeri di occupazione degli stati di particella singola dei gas perfetti quantistici abbiamo

$$\langle (\Delta n_r)^2 \rangle = \langle n_r^2 \rangle - \langle n_r \rangle^2 = \sum_n W_r(n) n^2 - \left[\frac{1}{e^{\beta(E_r-\mu)} - \varepsilon} \right]^2. \quad (5.69)$$

Ricordando la (5.61), per i fermioni il termine $\langle n_r^2 \rangle$ vale

$$\langle n_r^2 \rangle = \sum_{n=0}^1 W_r(n) n^2 = \sum_{n=0}^1 \frac{e^{-n\beta(E_r-\mu)}}{1 + e^{-\beta(E_r-\mu)}} n^2 = \frac{e^{-\beta(E_r-\mu)}}{1 + e^{-\beta(E_r-\mu)}} = \frac{1}{e^{\beta(E_r-\mu)} + 1}, \quad (5.70)$$

quindi, sempre per i fermioni

$$\langle (\Delta n_r)^2 \rangle = \langle n_r^2 \rangle - \langle n_r \rangle^2 = \frac{1}{e^{\beta(E_r - \mu)} + 1} - \left[\frac{1}{e^{\beta(E_r - \mu)} + 1} \right]^2 = \frac{e^{\beta(E_r - \mu)}}{[e^{\beta(E_r - \mu)} + 1]^2}. \quad (5.71)$$

Per i bosoni invece abbiamo

$$\begin{aligned} \langle n_r^2 \rangle &= \sum_{n=0}^{\infty} W_r(n) n^2 = \sum_{n=0}^{\infty} [1 - e^{-\beta(E_r - \mu)}] e^{-n\beta(E_r - \mu)} n^2 \\ &= [1 - e^{-\beta(E_r - \mu)}] \sum_{n=0}^{\infty} e^{-n\beta(E_r - \mu)} n^2 = \frac{e^{-\beta(E_r - \mu)} + e^{-2\beta(E_r - \mu)}}{[1 - e^{-\beta(E_r - \mu)}]^2}, \end{aligned} \quad (5.72)$$

dove l'ultimo passaggio è stato ottenuto utilizzando la relazione

$$\sum_{n=0}^{\infty} n^2 e^{-nx} = \frac{\partial^2}{\partial x^2} \sum_{n=0}^{\infty} e^{-nx} = \frac{\partial^2}{\partial x^2} \left(\frac{1}{1 - e^{-x}} \right) = \frac{e^{-x} + e^{-2x}}{(1 - e^{-x})^3}.$$

Moltiplicando numeratore e denominatore dell'ultima delle (5.72) per $e^{2\beta(E_r - \mu)}$ abbiamo

$$\langle n_r^2 \rangle = \frac{e^{\beta(E_r - \mu)} + 1}{[e^{\beta(E_r - \mu)} - 1]^2},$$

inserendo nella (5.69)

$$\langle (\Delta n_r)^2 \rangle = \langle n_r^2 \rangle - \langle n_r \rangle^2 = \frac{e^{\beta(E_r - \mu)} + 1}{[e^{\beta(E_r - \mu)} - 1]^2} - \left[\frac{1}{e^{\beta(E_r - \mu)} - 1} \right]^2 = \frac{e^{\beta(E_r - \mu)}}{[e^{\beta(E_r - \mu)} - 1]^2}. \quad (5.73)$$

Le espressioni per i bosoni e per i fermioni possono essere compattate nella forma

$$\langle (\Delta n_r)^2 \rangle = \frac{e^{\beta(E_r - \mu)}}{[e^{\beta(E_r - \mu)} - \varepsilon]^2} \quad (5.74)$$

Il limite classico si ha per $\alpha \rightarrow \infty$, quindi, essendo $\mu = -\alpha/\beta$, per $\mu \rightarrow -\infty$

$$\lim_{\mu \rightarrow -\infty} \langle (\Delta n_r)^2 \rangle = e^{-\beta(E_r - \mu)} = \lim_{\mu \rightarrow -\infty} \langle n_r \rangle, \quad (5.75)$$

per cui, al limite classico

$$\Delta n_r \simeq \sqrt{\langle n_r \rangle}, \quad \text{e, per le fluttuazioni relative} \quad \frac{\langle (\Delta n_r)^2 \rangle}{\langle n_r \rangle^2} = \frac{1}{\langle n_r \rangle}, \quad \frac{\sqrt{\langle (\Delta n_r)^2 \rangle}}{\langle n_r \rangle} = \frac{1}{\sqrt{\langle n_r \rangle}}. \quad (5.76)$$

Per la statistica di Fermi-Dirac abbiamo

$$\begin{aligned} \langle (\Delta n_r)^2 \rangle &= \frac{e^{\beta(E_r - \mu)}}{[e^{\beta(E_r - \mu)} + 1]^2} = \frac{[e^{\beta(E_r - \mu)} + 1] - 1}{[e^{\beta(E_r - \mu)} + 1]^2} = \frac{1}{e^{\beta(E_r - \mu)} + 1} - \left[\frac{1}{e^{\beta(E_r - \mu)} + 1} \right]^2 \\ &= \langle n_r \rangle - \langle n_r \rangle^2, \quad \text{da cui} \quad \frac{\langle (\Delta n_r)^2 \rangle}{\langle n_r \rangle^2} = \frac{1}{\langle n_r \rangle} - 1, \end{aligned} \quad (5.77)$$

mentre per la statistica di Bose-Einstein

$$\begin{aligned}
 \langle (\Delta n_r)^2 \rangle &= \frac{e^{\beta(E_r - \mu)}}{[e^{\beta(E_r - \mu)} - 1]^2} = \frac{[e^{\beta(E_r - \mu)} - 1] + 1}{[e^{\beta(E_r - \mu)} - 1]^2} = \frac{1}{e^{\beta(E_r - \mu)} - 1} + \left[\frac{1}{e^{\beta(E_r - \mu)} - 1} \right]^2 \\
 &= \langle n_r \rangle + \langle n_r \rangle^2, \quad \text{da cui} \quad \frac{\langle (\Delta n_r)^2 \rangle}{\langle n_r \rangle^2} = \frac{1}{\langle n_r \rangle} + 1 \quad .
 \end{aligned} \tag{5.78}$$

Cioè, per i fermioni la dispersione relativa è minore che per il limite classico, mentre per i bosoni è maggiore. In particolare, per i fermioni la dispersione relativa tende 0 per $\langle n_r \rangle \rightarrow 1$, mentre per i bosoni la dispersione non tende a 0 per $\langle n_r \rangle \rightarrow \infty$.

Capitolo 6

Random walk

6.1 Introduzione

Il problema del *random walk* (esiste, ma è poco usata, anche la traduzione italiana *passeggiata aleatoria*), è un'introduzione semplice, ma importante, a vari aspetti della teoria della probabilità. Una formulazione tradizionale del problema è la seguente: un ubriaco parte da sotto un lampione situato a metà di una strada. L'ubriaco fa passi ognuno di lunghezza l . Tuttavia l'uomo è così ubriaco che la direzione di ogni passo, verso l'estremità della strada dei numeri civici bassi (diciamo verso sinistra per chi guarda dall'altro marciapiede) o quella dei numeri civici alti (diciamo verso destra), è del tutto indipendente dalla direzione del passo precedente. Tutto quello che possiamo dire è che ogni passo ha probabilità p di essere verso destra, e probabilità $q = 1 - p$ di essere verso sinistra. Scegliamo un sistema di riferimento con l'asse x lungo la strada, e l'origine $x = 0$ al lampione da cui il nostro ubriaco parte. Dopo ogni passo l'uomo si troverà a una posizione del tipo $x = ml$, con m numero intero. Il nostro problema è: dopo che l'uomo ha effettuato N passi, qual è la probabilità $P_N(m)$ che si trovi alla posizione $x = ml$? Questo problema unidimensionale può essere facilmente generalizzato a due (un ubriaco che parte da un lampione al centro di una piazza?) o più dimensioni.

Ovviamente il problema fisico non riguarderà la passeggiata di un ubriaco, ma la somma di N vettori uguali in modulo ma di direzione diversa, come, solo per fare un esempio, la somma di N spin $\frac{1}{2}$ ognuno dei quali può essere "su" o "giù".

6.2 Impostazione matematica in una dimensione

In termini più fisici, poniamoci il problema di dove possa trovarsi una particella vincolata a muoversi lungo l'asse x che, partendo da $x = 0$, abbia effettuato N spostamenti successivi, ognuno uguale a l in modulo. Chiaramente la posizione della particella sarà del tipo

$$x = ml, \quad \text{con } m \text{ intero tale che } -N \leq m \leq N. \quad (6.1)$$

Vogliamo calcolare la probabilità $P_N(m)$ per ogni posizione possibile dopo N passi. Se chiamiamo n_1 il numero di passi verso destra, ed $n_2 = N - n_1$ il numero di passi verso sinistra, lo spostamento complessivo, misurato in "unità di passo" l , sarà

$$m = n_1 - n_2 = n_1 - (N - n_1) = 2n_1 - N. \quad (6.2)$$

Questo dimostra che, se N è pari, anche la posizione m deve essere pari. O, in altre parole, N ed m devono avere la stessa parità. Siamo partiti dall'ipotesi che passi successivi siano statisticamente indipendenti, con probabilità p per i passi verso destra e q per i passi verso sinistra. Quindi, la probabilità di una *singola* sequenza di N passi, di cui n_1 sono verso destra ed n_2 verso sinistra, sarà $p^{n_1} q^{n_2}$. Ma non esiste un'unica sequenza costituita da n_1 passi verso destra ed n_2 verso sinistra, ne esistono $\frac{N!}{n_1! n_2!}$. Quindi la probabilità $W_N(n_1)$ di aver fatto complessivamente n_1 passi verso destra ed $n_2 = N - n_1$ verso sinistra in un ordine qualunque è

$$W_N(n_1) = \frac{N!}{n_1! n_2!} p^{n_1} q^{n_2}. \quad (6.3)$$

Notare che

$$\frac{N!}{n_1! n_2!} = \frac{N!}{n_1! (N - n_1)!} = \binom{N}{n_1} \quad (6.4)$$

è il *coefficiente binomiale* che compare nell'espressione dello sviluppo della potenza ennesima del binomio $(a + b)$

$$(a + b)^N = \sum_{n_1=0}^N \binom{N}{n_1} a^{N-n_1} b^{n_1}. \quad (6.5)$$

Ovviamente, noti N ed n_1 , sono automaticamente noti anche n_2 ed m , con

$$n_2 = N - n_1 \quad \text{e} \quad m = 2n_1 - N, \quad \text{da cui} \quad n_1 = \frac{1}{2}(N + m) \quad \text{e} \quad n_2 = \frac{1}{2}(N - m). \quad (6.6)$$

Ricordiamo che, avendo N ed m la stessa parità, sia $N + m$ che $N - m$ sono numeri pari. E' quindi possibile scrivere la probabilità di trovarsi nella posizione m dopo N passi nella forma

$$P_N(m) = W_N\left(\frac{N + m}{2}\right) = \frac{N!}{\left(\frac{N + m}{2}\right)! \left(\frac{N - m}{2}\right)!} p^{(N+m)/2} q^{(N-m)/2} \quad (6.7)$$

Nel caso particolare in cui la probabilità di fare un passo verso destra sia uguale alla probabilità di fare un passo verso sinistra, cioè se $p = q = 1/2$, abbiamo

$$P_N^{(p=q)}(m) = \frac{N!}{\left(\frac{N + m}{2}\right)! \left(\frac{N - m}{2}\right)!} \left(\frac{1}{2}\right)^N. \quad (6.8)$$

Si noti che la (6.8) è il rapporto tra il numero di sequenze che in N passi portano dalla posizione 0 alla posizione m , che vale

$$\frac{N!}{\left(\frac{N + m}{2}\right)! \left(\frac{N - m}{2}\right)!}, \quad (6.9)$$

e il numero totale di possibili sequenze di N passi, che vale 2^N .

6.3 Valor medio della posizione

Dalla (6.3) abbiamo visto che la probabilità di fare n_1 passi verso destra, e $n_2 = N - n_1$ verso sinistra, su di un totale di N passi è

$$W(n_1) = \frac{N!}{n_1! (N - n_1)!} p^{n_1} q^{N-n_1} \quad (6.10)$$

(d'ora in poi trascureremo il pedice N da W_N quando questo è chiaro dal contesto). Per prima cosa controlliamo che sia soddisfatta la condizione di normalizzazione

$$\sum_{n_1=0}^N W(n_1) = 1. \quad (6.11)$$

Abbiamo

$$\sum_{n_1=0}^N \frac{N!}{n_1! (N - n_1)!} p^{n_1} q^{N-n_1} = (p + q)^N = 1 \quad (6.12)$$

dato che $p + q = 1$. Il valore di aspettazione $\bar{n}_1$ del numero di passi verso destra è, per definizione,

$$\bar{n}_1 = \sum_{n_1=0}^N W(n_1) n_1 = \sum_{n_1=0}^N \frac{N!}{n_1! (N - n_1)!} p^{n_1} q^{N-n_1} n_1. \quad (6.13)$$

Osservando che

$$n_1 p^{n_1} = p \frac{\partial}{\partial p} p^{n_1} \quad (6.14)$$

il valore di aspettazione di $\bar{n}_1$ può essere riscritto

$$\bar{n}_1 = \sum_{n_1=0}^N \frac{N!}{n_1! (N - n_1)!} \left[p \frac{\partial}{\partial p} (p^{n_1}) \right] q^{N-n_1} = p \frac{\partial}{\partial p} \left[\sum_{n_1=0}^N \frac{N!}{n_1! (N - n_1)!} p^{n_1} q^{N-n_1} \right]. \quad (6.15)$$

Questo, ricordando il teorema binomiale, può essere riscritto

$$\bar{n}_1 = p \frac{\partial}{\partial p} (p + q)^N = p N (p + q)^{N-1}. \quad (6.16)$$

Questo risultato è valido per valori arbitrari di p e q . In particolare è valido nel caso che ci interessa, cioè per $q \equiv 1 - p$, quindi abbiamo

$$\bar{n}_1 = Np, \quad (6.17)$$

risultato che potevamo aspettarci, dato che, se p è la probabilità di fare un passo verso destra, il numero medio di passi verso destra su un totale di N passi sarà semplicemente Np . Per quel che riguarda lo spostamento medio $\bar{m}$ abbiamo

$$\bar{m} = \bar{n}_1 - n_2 = N(p - q). \quad (6.18)$$

Nel caso particolare di $p = q = 1/2$ abbiamo, ovviamente, $\bar{m} = 0$.

6.4 Dispersione

Abbiamo per definizione

$$\overline{(\Delta n_1)^2} \equiv \overline{(n_1 - \bar{n}_1)^2} = \overline{n_1^2} - \bar{n}_1^2. \quad (6.19)$$

Conosciamo già il valore di $\bar{n}_1$, ci resta da calcolare il valore di $\overline{n_1^2}$, definito da

$$\overline{n_1^2} = \sum_{n_1=0}^N W(n_1) n_1^2 = \sum_{n_1=0}^N \frac{N!}{n_1! (N - n_1)!} p^{n_1} q^{N-n_1} n_1^2. \quad (6.20)$$

Usando un trucco analogo alla (6.14) possiamo scrivere

$$n_1^2 p^{n_1} = \left(p \frac{\partial}{\partial p} \right)^2 p^{n_1}, \quad (6.21)$$

e introducendo questa relazione nella (6.20) otteniamo

$$\begin{aligned} \overline{n_1^2} &= \sum_{n_1=0}^N \frac{N!}{n_1! (N - n_1)!} \left(p \frac{\partial}{\partial p} \right)^2 p^{n_1} q^{N-n_1} = \left(p \frac{\partial}{\partial p} \right)^2 \sum_{n_1=0}^N \frac{N!}{n_1! (N - n_1)!} p^{n_1} q^{N-n_1} \\ &= \left(p \frac{\partial}{\partial p} \right)^2 (p + q)^N = \left(p \frac{\partial}{\partial p} \right) [pN(p + q)^{N-1}] \\ &= p [N(p + q)^{N-1} + pN(N - 1)(p + q)^{N-2}]. \end{aligned} \quad (6.22)$$

Questa formula è valida per p e q generici. Il caso che ci interessa è quello di $p + q = 1$, per cui

$$\begin{aligned} \overline{n_1^2} &= p [N + pN(N - 1)] = pN(1 + pN - p) \\ &= (Np)^2 + Np(1 - p) = \bar{n}_1^2 + Np(1 - p), \end{aligned} \quad (6.23)$$

dove abbiamo inserito la (6.17). Notare che $1 - p = q$. Abbiamo così

$$\overline{(\Delta n_1)^2} = Np(1 - p) = Npq. \quad (6.24)$$

Per la dispersione quadratica relativa abbiamo

$$\frac{\overline{(\Delta n_1)^2}}{\bar{n}_1^2} = \frac{Np(1 - p)}{N^2 p^2} = \frac{1 - p}{p} \frac{1}{N} \quad \text{da cui} \quad \sqrt{\frac{\overline{(\Delta n_1)^2}}{\bar{n}_1^2}} = \sqrt{\frac{1 - p}{p}} \frac{1}{\sqrt{N}}. \quad (6.25)$$

Nel caso particolare di $p = q = 1/2$ abbiamo

$$\sqrt{\frac{\overline{(\Delta n_1)^2}}{\bar{n}_1^2}} = \frac{1}{\sqrt{N}}. \quad (6.26)$$

Per quel che riguarda la dispersione di m , cioè la dispersione della posizione, abbiamo

$$m = n_1 - n_2 = 2n_1 - N,$$

quindi

$$\Delta m = m - \bar{m} = (2n_1 - N) - (2\bar{n}_1 - N) = 2(n_1 - \bar{n}_1) = 2\Delta n_1,$$

da cui

$$(\Delta m)^2 = 4(\Delta n_1)^2.$$

Inserendo la (6.24)

$$\overline{(\Delta m)^2} = 4Npq, \quad \text{in particolare per } p = q = \frac{1}{2} \text{ abbiamo } \overline{(\Delta m)^2} = N. \quad (6.27)$$

6.5 Distribuzione di probabilità per grandi N

Al crescere di N la distribuzione di probabilità binomiale $W(n_1)$ presenta un massimo sempre più pronunciato corrispondente a un certo valore $n_1 = \nu_1$, e decresce molto rapidamente quando n_1 si allontana da ν_1 . Questo comportamento può essere sfruttato per trovare un'approssimazione di $W(n_1)$.

Al crescere di N , infatti, cresce anche il valore di ν_1 . E, per valori di ν_1 sufficientemente grandi e valori di n_1 non troppo lontani da ν_1 , vale la relazione

$$|W(n_1 + 1) - W(n_1)| \ll W(n_1). \quad (6.28)$$

In altre parole, la variazione relativa di $W(n_1)$ per l'argomento che passa da n_1 a $n_1 \pm 1$ è molto piccola. La funzione W può così essere considerata, con buona approssimazione, una funzione continua della variabile continua n_1 (ovviamente, solo i valori interi di n_1 hanno senso fisico). In queste condizioni il valore ν_1 a cui W presenta il massimo può essere determinato con buona approssimazione dalla condizione

$$\frac{dW}{dn_1} = 0, \quad \text{o, equivalentemente,} \quad \frac{d \ln W}{dn_1} = 0. \quad (6.29)$$

Il motivo di cercare il massimo di $\ln W$ anziché quello di W è che $\ln W$ è una funzione che varia molto più lentamente di W al variare di n_1 , quindi ci aspettiamo che l'espansione in serie di Taylor di $\ln W$ converga molto più rapidamente dell'espansione diretta di W . Per studiare il comportamento di $\ln W$ nei pressi del suo massimo $\ln W(\nu_1)$, poniamo $n_1 = \nu_1 + \eta$ ed espandiamo $\ln W$ in potenze di η , ottenendo

$$\ln W(n_1) = \ln W(\nu_1) + \sum_{k=1} \frac{1}{k!} \frac{d^k \ln W}{dn_1^k} \eta^k, \quad (6.30)$$

dove tutte le derivate sono calcolate per $n_1 = \nu_1$. Poiché stiamo espandendo attorno a un massimo, avremo per le derivate prima e seconda

$$\frac{d \ln W}{dn_1} = 0 \quad \text{e} \quad \frac{d^2 \ln W}{dn_1^2} = -2\alpha < 0, \quad (6.31)$$

con α valore positivo per il momento ignoto. Per η sufficientemente piccolo possiamo quindi approssimare

$$\ln W(n_1) = \ln W(\nu_1 + \eta) \simeq \ln W(\nu_1) - \alpha \eta^2, \quad \text{da cui} \quad W(n_1) \simeq A e^{-\alpha \eta^2}, \quad (6.32)$$

dove abbiamo posto $A = W(v_1)$ per brevità. Dalla (6.3) abbiamo

$$\ln W(n_1) = \ln N! - \ln n_1! - \ln(N - n_1)! + n_1 \ln p + (N - n_1) \ln q. \quad (6.33)$$

Ricordando che per un numero $n \gg 1$ possiamo considerare $\ln n!$ una funzione praticamente continua, abbiamo

$$\frac{d \ln n!}{dn} \simeq \frac{\ln(n+1)! - \ln n!}{1} = \ln \frac{(n+1)!}{n!} = \ln(n+1) \simeq \ln n, \quad (6.34)$$

dove abbiamo approssimato la derivata con il rapporto incrementale. Derivando la (6.33) rispetto a n_1 e sostituendo la (6.34) otteniamo

$$\frac{d \ln W}{dn_1} = -\ln n_1 + \ln(N - n_1) + \ln p - \ln q = \ln \left[\frac{(N - n_1)p}{n_1 q} \right]. \quad (6.35)$$

Uguagliando questa derivata a zero, troviamo il valore v_1 di n_1 per cui W presenta il massimo

$$\ln \left[\frac{(N - v_1)p}{v_1 q} \right] = 0, \quad \text{da cui} \quad (N - v_1)p = v_1 q, \quad (6.36)$$

e, ricordando che $p + q = 1$, abbiamo finalmente

$$v_1 = Np. \quad (6.37)$$

Derivando la (6.35) una seconda volta rispetto a n_1 abbiamo

$$\frac{d^2 \ln W}{dn_1^2} = -\frac{1}{n_1} - \frac{1}{N - n_1}. \quad (6.38)$$

Sostituendo il valore $n_1 = Np$ e ricordando che nella (6.32) abbiamo $\alpha = -\frac{1}{2} \frac{d^2 \ln W}{dn_1^2} \Big|_{v_1}$ otteniamo

$$\alpha = -\frac{1}{2} \frac{d^2 \ln W}{dn_1^2} \Big|_{v_1} = \frac{1}{2} \left(\frac{1}{Np} + \frac{1}{N - Np} \right) = \frac{1}{2N} \left[\frac{1 - p + p}{p(1 - p)} \right] = \frac{1}{2Npq}. \quad (6.39)$$

La costante A che compare nella (6.32) va determinata in modo che la distribuzione W sia normalizzata. La condizione di normalizzazione può essere approssimata

$$1 = \sum_{n_1=0}^N W(n_1) \simeq \int W(n_1) dn_1 \simeq \int_{-\infty}^{+\infty} W(v_1 + \eta) d\eta, \quad (6.40)$$

dove i limiti di integrazione in $d\eta$ sono stati estesi da $-\infty$ a $+\infty$ perché l'integrando è trascurabile per valori di $|\eta|$ sufficientemente lontani da 0. Abbiamo così

$$1 = A \int_{-\infty}^{+\infty} e^{-\alpha \eta^2} d\eta = A \sqrt{\frac{\pi}{\alpha}}, \quad \text{da cui} \quad A = \sqrt{\frac{\alpha}{\pi}} = \sqrt{\frac{1}{2\pi Npq}}. \quad (6.41)$$

L'approssimazione per W diventa così la gaussiana

$$W(n_1) \simeq \sqrt{\frac{1}{2\pi Npq}} e^{-\frac{(n_1 - Np)^2}{2Npq}}. \quad (6.42)$$

Nel caso del random walk simmetrico, $p = q = 1/2$, abbiamo

$$W(n_1) \simeq \sqrt{\frac{2}{\pi N}} e^{-\frac{2(n_1 - N/2)^2}{N}}, \quad (6.43)$$

che, tenendo conto che la posizione è data da $m = 2n_1 - N$, ci porta alla distribuzione di probabilità per la posizione

$$W(m) \simeq \sqrt{\frac{2}{\pi N}} e^{-\frac{m^2}{2N}}. \quad (6.44)$$

6.6 Limite continuo

Il caso di un random walk in cui al trascorrere del tempo la posizione vari con continuità, anziché a scatti, si ottiene come limite per N ed m molto grandi e passo l molto piccolo. Infatti per $m \rightarrow \infty$ e $l \rightarrow 0$ la posizione $x = ml$ tenderà a variare con continuità. Al limite del continuo non dovremo calcolare la probabilità che all'istante t la posizione sia esattamente x (probabilità nulla al limite), ma la probabilità che sia all'interno di un certo intervallo Δx , che prenderemo molto maggiore di l , ma ancora sufficientemente piccolo perché si abbia

$$W_N\left(m + \frac{\Delta x}{l}\right) \simeq W_N(m). \quad (6.45)$$

Per passare dal discreto al continuo dobbiamo passare da $W_N(m)$, probabilità che dopo N passi la posizione sia $x = ml$, a una densità di probabilità $P(x, t)$ tale che, dopo un intervallo di tempo t dall'inizio del moto, $P(x, t) dx$ sia la probabilità che la posizione sia nell'intervallo $(x, x + dx)$. In termini di $W_N(m)$ la probabilità che, dopo N passi, la posizione sia compresa tra x e $x + \Delta x$ può essere approssimata da

$$\begin{aligned} P(x, N) \Delta x &= \sum_{m' = m - [\Delta x/2l]}^{m + [\Delta x/2l]} W_N(m') \simeq W_N(m) \left[\frac{\Delta x}{2l} \right] \\ &\simeq W_N\left(\frac{x}{l}\right) \cdot \left(\frac{\Delta x}{2l}\right), \end{aligned} \quad (6.46)$$

dove $[\Delta x/2l]$ indica la parte intera di $\Delta x/2l$, approssimata all'ultimo passaggio con la quantità reale $\Delta x/2l$ in base all'ipotesi $\Delta x \gg l$. Il numero di posizioni possibili entro l'intervallo Δx sono $[\Delta x/2l]$ perché, per la (6.2), m deve avere la stessa parità di N , quindi la spaziatura tra una posizione possibile e la successiva vale $2l$. Supponiamo ancora che i passi vengano effettuati a frequenza costante, cioè che sia

$$N = [\eta t], \quad \text{da cui} \quad t \simeq \frac{N}{\eta}, \quad (6.47)$$

dove η è la frequenza, o “numero di passi al secondo”, e t il tempo trascorso dall'inizio della “passeggiata”. Per $N \rightarrow \infty$ dovrà essere anche $\eta \rightarrow \infty$. Al limite del continuo possiamo inserire la (6.47)

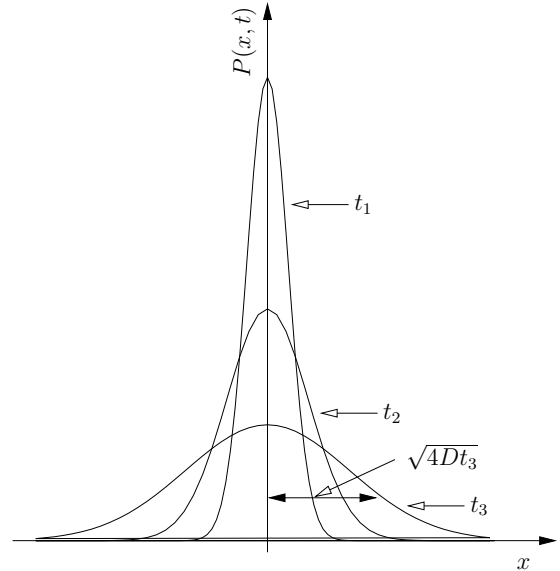


Figura 6.1 Evoluzione temporale della soluzione fondamentale dell'equazione di diffusione. Per $t = 0$ avremmo un δ di Dirac, poi la $P(x, t)$ si allarga proporzionalmente a $\sqrt{t}$. Qui abbiamo $t_2 = 4t_1$, $t_3 = 4t_2$.

nella (6.44) e questa nella (6.46). Sostituendo poi Δx con dx otteniamo finalmente

$$\begin{aligned} P(x, t) dx &= \sqrt{\frac{2}{\pi \eta t}} e^{-\frac{(x/l)^2}{2\eta t}} \left(\frac{dx}{2l} \right) \\ &= \frac{1}{\sqrt{2\pi \eta l^2 t}} e^{-x^2/(2\eta l^2 t)} dx, \end{aligned} \quad (6.48)$$

dove vediamo che i limiti $\eta \rightarrow \infty$ e $l \rightarrow 0$ devono avvenire in modo che il prodotto ηl^2 sia finito. Definendo il *coefficiente di diffusione traslazionale* come $D = \eta l^2/2$ otteniamo

$$P(x, t) = \frac{1}{\sqrt{2\pi D t}} e^{-x^2/(4Dt)}. \quad (6.49)$$

Questa espressione per $P(x, t)$ è la soluzione dell'equazione di diffusione

$$\frac{\partial P(x, t)}{\partial t} = D \frac{\partial^2 P(x, t)}{\partial x^2} \quad (6.50)$$

con la condizione iniziale $P(x, 0) = \delta(x)$, dove $\delta(x)$ è la delta di Dirac. La soluzione dell'equazione di diffusione con la condizione iniziale $P(x, 0) = \delta(x)$ si chiama *soluzione fondamentale*. L'evoluzione temporale di $P(x, t)$ è schematizzata in Fig. 6.1. Per la varianza abbiamo

$$\overline{x^2} = \frac{1}{\sqrt{2\pi D t}} \int_{-\infty}^{+\infty} x^2 e^{-x^2/(4Dt)} dx = -\frac{1}{\sqrt{2\pi D t}} \frac{\partial}{\partial \alpha} \int_{-\infty}^{+\infty} e^{-\alpha x^2} dx = -\frac{1}{\sqrt{2\pi D t}} \frac{\partial}{\partial \alpha} \sqrt{\frac{\pi}{\alpha}}, \quad (6.51)$$

dove abbiamo posto $\alpha = 1/(4Dt)$. Effettuando la derivata e risostituendo otteniamo

$$\sqrt{\overline{x^2}} = \sqrt{2Dt}. \quad (6.52)$$

6.7 Muro riflettente e muro assorbente

Supponiamo di avere un *muro riflettente* alla posizione kl , dove l è, come al solito, il modulo costante della lunghezza del passo. Senza perdita di generalità, supponiamo che sia $k > 0$. Supponiamo inoltre, per semplicità, di avere $p = q = 1/2$. Se la particella si trova in una qualunque posizione n , con $n < k$, al passo successivo può, con uguale probabilità, passare sia a $n - 1$ che a $n + 1$. Ma se si trova in $n = k$ subisce una riflessione elastica, e al passo successivo passa necessariamente alla posizione $k - 1$, abbiamo cioè $p = 0$ e $q = 1$ anziché $q = p = 1/2$. Il problema che ci poniamo è trovare la probabilità $P_N^{(r)}(m)$ che la particella, dopo N passi, si trovi nella posizione finale m in presenza del muro riflettente. Ancora una volta la probabilità sarà nulla se m ed N hanno parità diversa. Inoltre, anche se m ed N hanno la stessa parità, il vincolo del muro riflettente impone che la probabilità sia nulla se $m > k$. Ci aspettiamo invece che, per $m \leq k$ la probabilità sia maggiore di quella che avremmo in assenza del muro riflettente. Dimosteremo che, per $m < k$, la probabilità $P_N^{(r)}(m)$ è pari alla probabilità $P_N(m)$ di giungere in m con N passi in assenza del muro riflettente, data dalla (6.8), più la probabilità, sempre in assenza del muro riflettente, di raggiungere in N passi la posizione simmetrica a m rispetto al muro, cioè $(2k - m)$. La nostra tesi è quindi che sia

$$P_N^{(r)}(m) = P_N(m) + P_N(2k - m). \quad (6.53)$$

In presenza del muro la (6.8) deve essere modificata per tener conto del fatto che un percorso (sequenza di passi) che raggiunge m dopo n riflessioni deve essere contato 2^n volte. Questo perché ogni volta che si raggiunge k la probabilità di passare a $(k - 1)$ al passo successivo vale 1 anziché $1/2$. Consideriamo per esempio la sequenza di passi $OABXD$ della Fig. 6.2, che giunge in m dopo un'unica riflessione.

Nel punto A abbiamo $q = 1$ anziché $q = 1/2$, e questo percorso andrebbe contato due volte. D'altra parte, se riflettiamo la parte del percorso da A in poi rispetto a k otteniamo un percorso, permesso in assenza del muro, che porta al punto immagine D' in $(2k - m)$. Analogamente, per ogni percorso che porta in $(2k - m)$ con un solo attraversamento del muro, ce ne è uno, simmetrico rispetto a k dall'attraversamento in poi, che porta in m dopo un'unica riflessione. Così, invece di contare due volte ogni percorso che subisce una riflessione, possiamo aggiungergli un percorso univocamente determinato che porta in $(2k - m)$. Consideriamo

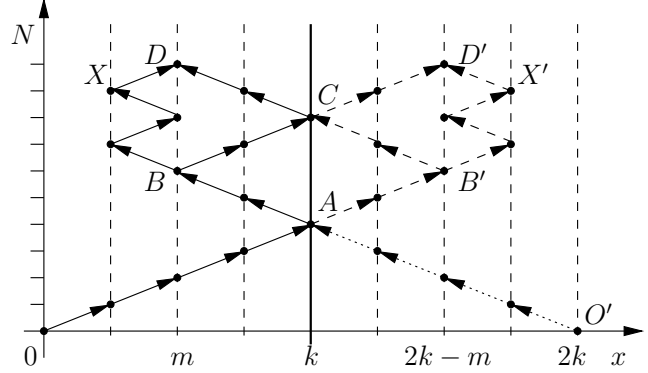


Figura 6.2 Muro riflettente. In questo grafico a ogni passo della particella il punto rappresentativo si sposta di un'unità verticale verso l'alto, e di un unità orizzontale verso destra o verso sinistra.

adesso la traiettoria $OABCD$, che porta in m dopo due riflessioni su k . Questa traiettoria andrebbe contata quattro volte. Ma ci sono le due traiettorie $OAB'CD'$ e $OABCD'$ che portano a $(2k - m)$, e la traiettoria $OAB'CD$ che conduce in m , che andrebbero tutte escluse a causa del muro. Continuando così si arriva a dimostrare la validità generale della (6.53). Analogamente, si può dimostrare che la probabilità di giungere con N passi in m in presenza di un muro riflettente è uguale alla probabilità di giungerci in assenza di muro, più la probabilità di giungerci, sempre in N passi e sempre in assenza del muro, partendo dall'immagine O' dell'origine O anziché dall'origine stessa.

Supponiamo adesso di avere ancora un muro, ma questa volta un *muro assorbente*, sempre nella posizione k , e che sia ancora $p = q = 1/2$. Quindi, se la particella si trova in una qualunque posizione n , con $n < k$, al passo successivo può passare, con uguale probabilità, sia a $n - 1$ che a $n + 1$. Ma se giunge in $n = k$, da quel momento in poi esce dal gioco. Per esempio, si può supporre sia che in k la particella resti attaccata alla parete, o che venga distrutta, in qualche modo “mangiata”. Ci poniamo due problemi:

1. l'analogo del problema del muro riflettente: calcolare la probabilità di giungere in m , con $m \leq k$, in N passi;
2. calcolare la probabilità che la particella sia ancora “libera” dopo N passi.

Calcoliamo prima la probabilità $P_N^{(a)}(m)$ di giungere in m in presenza del muro assorbente. Nel contare le possibili sequenze di passi questa volta dobbiamo escludere tutte quelle che contengono anche un singolo passaggio per k . Così possiamo prima contare tutte le sequenze di N passi che portano in m in assenza del muro, poi togliere un certo numero di “sequenze proibite”. Per quanto visto prima, ognuna di queste sequenze definisce univocamente un'altra sequenza che porta all'immagine $(2k - m)$ di m riflessa rispetto a k . Avremo quindi

$$P_N^{(a)}(m) = P_N(m) - P_N(2k - m). \quad (6.54)$$

Il secondo problema può essere riformulato partendo dalla probabilità $A_N(k)$ che la particella giunga per la prima volta in k , e vi sia assorbita, all' N -esimo passo. Da questa otteniamo la probabilità $L_N(k)$ che all' N -esimo passo la particella sia ancora libera come $L_N(k) = 1 - A_N(k)$. Se N ha parità diversa da k abbiamo $A_N(k) = 0$, perché per la (6.2) la particella non può essere in k all' N -esimo passo. In questo caso abbiamo $L_N(k) = 1 - A_{(N-1)}(k)$. Per N e k della stessa parità, chiamiamo $S_N(k)$ il numero di percorsi che, in presenza del muro assorbente, portano per la prima volta in k al passo N . In assenza di muro la (6.9), posto $m = k$, ci dà il numero totale di percorsi che portano da 0 a k in N passi. Tutti questi percorsi avevano portato la particella in $(k - 1)$ o in $(k + 1)$ al passo precedente, l' $(N - 1)$ -esimo. Ma, in presenza del muro, i percorsi che passano per $(k + 1)$ sono proibiti e vanno tolti dal conteggio. Inoltre sono proibiti tutti quei percorsi che arrivano in $(k - 1)$ dopo aver attraversato, o toccato, il muro una o più volte. Il numero di questi ultimi percorsi proibiti è uguale al numero di sequenze di $(N - 1)$ passi che portano da 0 a $(k + 1)$ in assenza di muro. Infatti, da ognuna di queste sequenze si ottiene, per riflessione sul muro dal primo passaggio per k in poi, una sequenza proibita che porta in $(k - 1)$. Così, per quanto appena detto, $S_N(k)$ è uguale al numero di sequenze di N passi che portano in k in assenza di muro, meno due volte il numero di sequenze di $(N - 1)$ passi che portano, sempre in assenza di muro, in $(k + 1)$. Usando la (6.9) abbiamo

$$\begin{aligned}
 S_N(k) &= \frac{N!}{\left(\frac{N-k}{2}\right)! \left(\frac{N+k}{2}\right)!} - 2 \frac{(N-1)!}{\left(\frac{N+k}{2}\right)! \left(\frac{N-k-2}{2}\right)!} \\
 &= \frac{N!}{\left(\frac{N-k}{2}\right)! \left(\frac{N+k}{2}\right)!} \left(1 - \frac{N-k}{N}\right) \\
 &= \frac{k}{N} \frac{N!}{\left(\frac{N-k}{2}\right)! \left(\frac{N+k}{2}\right)!}. \tag{6.55}
 \end{aligned}$$

Per ottenere la probabilità $A_N(k)$ che la particella sia assorbita dal muro all' N -esimo passo, ma non prima, dobbiamo dividere il numero di percorsi appena calcolato, $S_N(k)$, per il numero totale di sequenze di N passi possibili in assenza di muro, cioè 2^N , ottenendo

$$A_N(k) = \frac{S_N(k)}{2^N} = \frac{k}{N} \frac{N!}{\left(\frac{N-k}{2}\right)! \left(\frac{N+k}{2}\right)!} \left(\frac{1}{2}\right)^N = \frac{k}{N} P_N(k), \tag{6.56}$$

dove, all'ultimo passaggio, abbiamo sostituito la (6.8). Quindi la probabilità che, in presenza di muro assorbente in posizione k , la particella sia ancora libera dopo N passi vale

$$L_N(k) = 1 - A_N(k) = 1 - \frac{k}{N} P_N(k). \tag{6.57}$$

6.8 Momenti di una variabile aleatoria

In statistica, il *momento semplice* (o *momento attorno all'origine*) di ordine k di una variabile aleatoria x , denotato M_k , è definito come il valor medio, o valore di aspettazione, di x^k . Se x ha distribuzione

continua con densità di probabilità $p(x)$ abbiamo così, per definizione,

$$M_k = \int_{-\infty}^{+\infty} x^k p(x) dx. \quad (6.58)$$

Il *momento centrale* (o *momento rispetto alla media*) di ordine k della stessa variabile casuale x , denotato m_k , è definito invece come il valor medio della k -esima potenza dello scarto di x dal suo valor medio $\mu = M_1 = \int_{-\infty}^{+\infty} x p(x) dx$. In formula la definizione di m_k è

$$m_k = \int_{-\infty}^{+\infty} (x - \mu)^k p(x) dx. \quad (6.59)$$

I momenti centrali sono usati più spesso dei momenti semplici perché sono legati alle deviazioni rispetto al valor medio anziché alle deviazioni rispetto allo zero. Descrivono quindi la dispersione (o variabilità) e la forma della distribuzione, indipendentemente dalla sua localizzazione. E' da notare che, se una distribuzione è simmetrica (invariante per riflessione attorno al suo valor medio), tutti i suoi momenti centrali dispari sono nulli. Importanti caratteristiche dei momenti semplici e centrali sono:

1. M_0 e m_0 sono sempre uguali all'unità;
2. M_1 è la media aritmetica, indicata tradizionalmente con μ ;
3. m_1 è sempre nullo;
4. $m_2 = M_2 - M_1^2$ è la *varianza*, indicata tradizionalmente con σ^2 . Questa quantità fornisce una misura di quanto i valori assunti dalla variabile si scostino dal valor medio.

Per quanto sopra, in caso di distribuzione non simmetrica rispetto al proprio valor medio, il primo momento centrale dispari che può essere diverso da zero è m_3 . A partire da m_3 si costruisce l'*indice di asimmetria* (in inglese *skewness*) γ_1 definito da

$$\gamma_1 = \frac{m_3}{m_2^{3/2}} = \frac{m^3}{\sigma^3}. \quad (6.60)$$

Questo è il primo momento che *può* dare informazioni sull'asimmetria di una distribuzione. La divisione per $m_2^{3/2}$ rende γ_1 invariante per trasformazioni lineari della variabile del tipo $x' = ax + b$, con $a > 0$, che trasformano i momenti centrali come $m_k(ax + b) = a^k m_k(x)$.

Capitolo 7

Processi stocastici

7.1 Introduzione

Un processo casuale, o stocastico, è un processo in cui il valore di una certa variabile x dipende da un'altra variabile, che qui supporremo essere il tempo t (ma il concetto può essere facilmente generalizzato), secondo una legge $x(t)$ che comporta un certo grado di aleatorietà. Non esiste cioè una funzione $x(t)$ che descriva esattamente il variare di x al passare del tempo. Qui faremo però le ipotesi che i) la dipendenza di x da t , anche se imprevedibile, sia comunque continua, e che ii) sia possibile scrivere una *densità di probabilità al primo ordine* dipendente dal tempo

$$W_1(x, t) \quad (7.1)$$

tale che $W_1(x', t) dx$ sia la probabilità che il valore di x sia nell'intervallo $(x', x' + dx)$ all'istante t . Per trattare i casi in cui, ripetendo più volte lo stesso esperimento sotto le stesse condizioni *macroscopiche* iniziali, ma nell'impossibilità di controllare le condizioni *microscopiche*, si osservano evoluzioni temporali diverse, è stato introdotto il concetto di *ensemble* (francese per *insieme*), o *ensemble statistico* (in italiano si dice anche *insieme statistico* o *insieme rappresentativo*). L'ensemble è un'astrazione in cui si immagina di avere un grandissimo numero di copie del sistema fisico considerato, in ognuna delle quali la variabile x segue una delle evoluzioni temporali possibili. Il valor medio della variabile x all'istante t su un ensemble costituito da N copie del sistema, indicato da parentesi triangolari $\langle x(t) \rangle$, vale così

$$\langle x(t) \rangle = \frac{1}{N} \sum_{i=1}^N x_i(t), \quad (7.2)$$

dove l'indice i corre su tutte le copie del sistema, e $x_i(t)$ è il valore di x all'istante t nella i -esima copia. In termini della densità di probabilità al primo ordine (7.1) il valor medio, o valore di aspettazione, $\langle x(t) \rangle$ della (7.2) coincide con la media pesata

$$\langle x(t) \rangle = \int x W_1(x, t) dx. \quad (7.3)$$

Definita la densità di probabilità al primo ordine, definiamo poi la *densità di probabilità congiunta al secondo ordine*

$$W_2(x_1, t_1; x_2, t_2), \quad (7.4)$$

tale che $W_2(x_1, t_1; x_2, t_2) dx_1 dx_2$ sia la probabilità di trovare x nell'intervallo $(x_1, x_1 + dx_1)$ all'istante t_1 e nell'intervallo $(x_2, x_2 + dx_2)$ all'istante t_2 . Una grandezza che ci interesserà nel paragrafo 7.3 è l'*autocorrelazione* di x , cioè la correlazione tra i valori che la stessa variabile stocastica x assume a due diversi istanti t e $t + \tau$, definita dal valor medio sull'ensemble

$$\langle x(t) x(t + \tau) \rangle = \frac{1}{N} \sum_{i=1}^N x_i(t) x_i(t + \tau). \quad (7.5)$$

Mediante la $W_2(x_1, t_1; x_2, t_2)$ la funzione di autocorrelazione di x si riscrive

$$\langle x(t) x(t + \tau) \rangle = \int \int x_1 x_2 W_2(x_1, t + \tau; x_2, t) dx_1 dx_2. \quad (7.6)$$

In modo analogo alla (7.4) si definiscono le densità di probabilità congiunte di ogni ordine superiore n

$$W_n(x_1, t_1; x_2, t_2; \dots; x_n, t_n). \quad (7.7)$$

Qui e nel seguito, all'interno delle parentesi usiamo i punti e virgola per separare gli eventi, e la virgola per separare la coppia di valori x e t corrispondenti allo stesso evento. Inoltre facciamo sempre l'ipotesi che sia $t_1 > t_2 > \dots > t_{n-1} > t_n$. Attenzione: sfortunatamente esistono sia libri che adottano questa convenzione, sia libri che adottano la convenzione opposta ($t_n > \dots > t_1$). Un processo stocastico si dice *stazionario* se le densità di probabilità congiunte di qualunque ordine n sono invarianti per traslazioni dell'origine dei tempi, se cioè abbiamo

$$W_n(x_1, t_1 + \Delta t; x_2, t_2 + \Delta t; \dots; x_n, t_n + \Delta t) = W_n(x_1, t_1; x_2, t_2; \dots; x_n, t_n), \quad (7.8)$$

con Δt arbitrario. Quindi, in un processo stazionario, W_n dipende solo dalle differenze temporali $t_1 - t_n$, $t_2 - t_n$, $\dots$, $t_{n-1} - t_n$, e non da t_n . Per un processo stazionario possiamo così scrivere

$$W_n(x_1, t_1; x_2, t_2; \dots; x_n), \quad (7.9)$$

dove t_n è stato posto uguale a 0 e ignorato nella scrittura. Una condizione necessaria, ma non sufficiente, perché un processo stocastico sia stazionario, è che $W_1(x, t)$ sia indipendente dal tempo, e che quindi si possa scrivere $W_1(x, t) = W_1(x)$. Per un processo stazionario la funzione di correlazione temporale (7.6) dipende dall'intervallo τ tra i due eventi, ma non da t .

Oltre alle densità di probabilità congiunte di ordine n , introduciamo le *densità di probabilità condizionata* $P_{n|m}(x_1, t_1; x_2, t_2; \dots; x_n, t_n | y_1, \tau_1; y_2, \tau_2; \dots; y_m, \tau_m)$ definite in modo che

$$P_{n|m}(x_1, t_1; x_2, t_2; \dots; x_n, t_n | y_1, \tau_1; y_2, \tau_2; \dots; y_m, \tau_m) dx_1 dx_2 \dots dx_n \quad (7.10)$$

sia la probabilità che x si trovi nell'intervallo $(x_1, x_1 + dx_1)$ all'istante t_1 , nell'intervallo $(x_2, x_2 + dx_2)$ all'istante $t_2, \dots$, nell'intervallo $(x_n, x_n + dx_n)$ all'istante t_n , una volta noto che il valore di x era y_1 all'istante τ_1 , y_2 all'istante $\tau_2, \dots, y_m$ all'istante τ_m . Continuiamo a supporre che sia

$$t_1 > t_2 > \dots > t_n > \tau_1 > \tau_2 > \dots > \tau_m. \quad (7.11)$$

All'interno della parentesi continuiamo ad usare i punti e virgola per separare i singoli eventi, e usiamo la barra “|” per separare gli n eventi condizionati dagli m eventi condizionanti. La condizione

di stazionarietà continua a essere che $P_{n|m}(x_1, t_1; x_2, t_2; \dots; x_n, t_n | y_1, \tau_1; y_2, \tau_2; \dots; y_m, \tau_m)$ dipenda dal tempo solo attraverso le differenze $t_1 - \tau_m, t_2 - \tau_m, \dots, t_n - \tau_m, \tau_1 - \tau_m, \dots, \tau_{m-1} - \tau_m$. In condizioni stazionarie possiamo ancora una volta porre $\tau_m = 0$ e scrivere

$$P_{n|m}(x_1, t_1; x_2, t_2; \dots; x_n, t_n | y_1, \tau_1; y_2, \tau_2; \dots; y_m) . \quad (7.12)$$

Nel caso di processi stazionari abbiamo, per esempio,

$$\int_{-\infty}^{+\infty} P_{1|1}(x, t | y) dx = 1, \quad (7.13)$$

che rappresenta la certezza che, qualunque sia il valore iniziale y , dopo un qualunque intervallo di tempo t , la variabile abbia un valore reale qualunque. Poi abbiamo

$$W_2(x_1, t; x_2) = P_{1|1}(x_1, t | x_2) W_1(x_2), \quad \text{da cui otteniamo} \quad P_{1|1}(x_1, t | x_2) = \frac{W_2(x_1, t; x_2)}{W_1(x_2)}. \quad (7.14)$$

Valgono le relazioni

$$W_3(x_1, t_1; x_2, t_2; x_3) = P_{1|2}(x_1, t_1 | x_2, t_2; x_3) W_2(x_2, t_2; x_3) \quad (7.15)$$

e

$$\int_{-\infty}^{+\infty} W_3(x_1, t_1; x_2, t_2; x_3) dx_2 = W_2(x_1, t_1; x_3) \quad (7.16)$$

per ogni t_2 tale che $0 < t_2 < t_1$.

Un caso importante dei processi casuali è costituito dai processi markoviani (da Андрей Андреевич Марков), che sono definiti dalla relazione

$$P_{1|m}(x, t | y_1, \tau_1; y_2, \tau_2; \dots; y_m) = P_{1|1}(x, t - \tau_1 | y_1). \quad (7.17)$$

Questo significa che in un processo markoviano la probabilità di un evento condizionato è determinata solo dall'evento condizionante più recente, e non da tutta la storia precedente a quest'ultimo. In una *catena markoviana*, cioè, noto un evento, la conoscenza di tutta la storia anteriore non aggiunge informazione utile per determinare con maggiore precisione la probabilità degli eventi successivi. Un esempio di processo markoviano è quello di una variabile x che soddisfa l'equazione differenziale $\frac{dx}{dt} = f(x)$. Una volta noto il valore x_0 della variabile all'istante t_0 l'evoluzione temporale di x è determinata univocamente da una relazione del tipo $x = \varphi(x_0, t - t_0)$. La x soddisfa la relazione di Markov con la densità di probabilità condizionata

$$P_{1|1}(x, t | x_0, t_0) = \delta [x - \varphi(x_0, t - t_0)]. \quad (7.18)$$

Da qui segue che i processi deterministici sono casi particolari dei processi markoviani.

Nel caso di un processo markoviano la relazione (7.15) diventa

$$W_3(x_1, t_1; x_2, t_2; x_3) = P_{1|1}(x_1, t_1 | x_2, t_2) W_2(x_2, t_2; x_3) \quad (7.19)$$

e le relazioni (7.14) e (7.16) combinate ci danno

$$P_{1|1}(x_f, t | x_i) = \frac{W_2(x_f, t; x_i)}{W_1(x_i)} = \frac{\int_{-\infty}^{+\infty} W_3(x_f, t; x_2, t_2; x_i) dx_2}{W_1(x_i)} \quad (7.20)$$

Utilizzando adesso la (7.19) otteniamo per i processi markoviani

$$P_{1|1}(x_f, t|x_i) = \int_{-\infty}^{+\infty} \frac{W_2(x_2, t_2; x_i)}{W_1(x_i)} P_{1|1}(x_f, t|x_2, t_2) dx_2 = \int_{-\infty}^{+\infty} P_{1|1}(x_f, t|x_2, t_2) P_{1|1}(x_2, t_2|x_i) dx_2, \quad (7.21)$$

dove l'integrazione è su tutti possibili valori x_2 della variabile stocastica al tempo intermedio t_2 . Questa è l'equazione di Smoluchowski, o equazione di Chapman-Kolmogorov (Андрей Николаевич Колмогоров), valida per qualunque t_2 tale che $0 < t_2 < t$.

7.2 Densità spettrale di una funzione stocastica

Ricordiamo che, data una funzione *quadraticamente integrabile* del tempo $f(t)$, la sua trasformata di Fourier $\tilde{f}(\nu)$ è definita come

$$\tilde{f}(\nu) = \int_{-\infty}^{+\infty} f(t) e^{-2\pi i \nu t} dt, \quad \text{e vale la relazione inversa} \quad f(t) = \int_{-\infty}^{+\infty} \tilde{f}(\nu) e^{2\pi i \nu t} d\nu. \quad (7.22)$$

La densità spettrale, o *densità spettrale di energia*, $\Phi(\nu)$ di $f(t)$ è una funzione reale e positiva della frequenza definita come

$$\Phi(\nu) = \left| \int_{-\infty}^{+\infty} f(t) e^{-2\pi i \nu t} dt \right|^2 = |\tilde{f}(\nu)|^2. \quad (7.23)$$

Per le funzioni $f(t)$ e $\Phi(\nu)$ vale il *teorema di Parseval*

$$\int_{-\infty}^{+\infty} |f(t)|^2 dt = \int_{-\infty}^{+\infty} |\tilde{f}(\nu)|^2 d\nu = \int_{-\infty}^{+\infty} \Phi(\nu) d\nu. \quad (7.24)$$

Consideriamo adesso una funzione stocastica stazionaria reale $x(t)$. Se vogliamo farne la trasformata di Fourier, ed eventualmente applicare poi il teorema di Parseval (7.24), ci troviamo di fronte al problema che $x(t)$ non tende a zero per $t \rightarrow \pm\infty$, essendo, appunto, *stazionaria*. Quindi l'integrale $\int_{-\infty}^{+\infty} x^2(t) dt$ diverge, e la nostra $x(t)$ non è quadraticamente integrabile. Noi, però, normalmente, non saremo interessati a questo integrale, ma solo alla media temporale di x^2 , definita come

$$\overline{x^2} = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} x^2(t) dt. \quad (7.25)$$

Per trattare il problema, definiamo una funzione $x_T(t)$ che coincide con $x(t)$ per $-T < t < +T$, ed è nulla al di fuori di questo intervallo, $x_T(t)$ è quindi quadraticamente integrabile. Se chiamiamo $\tilde{x}_T(\nu)$ la sua trasformata di Fourier, il teorema di Parseval si scrive

$$\int_{-\infty}^{+\infty} x_T^2(t) dt = \int_{-T}^{+T} x^2(t) dt = \int_{-\infty}^{+\infty} |\tilde{x}_T(\nu)|^2 d\nu, \quad (7.26)$$

dove abbiamo posto

$$\tilde{x}_T(\nu) = \int_{-\infty}^{+\infty} x_T(t) e^{-2\pi i \nu t} dt = \int_{-T}^{+T} x(t) e^{-2\pi i \nu t} dt, \quad (7.27)$$

e il valore mediato sul tempo di $x_T^2(t)$ è

$$\overline{x_T^2} = \frac{1}{2T} \int_{-T}^{+T} x^2(t) dt = \frac{1}{2T} \int_{-\infty}^{+\infty} |\tilde{x}_T(\nu)|^2 d\nu = \int_{-\infty}^{+\infty} \frac{1}{2T} |\tilde{x}_T(\nu)|^2 d\nu. \quad (7.28)$$

Per il valor medio sul tempo della nostra $x^2(t)$ originale abbiamo così

$$\overline{x^2} = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} x^2(t) dt = \int_{-\infty}^{+\infty} \lim_{T \rightarrow \infty} \frac{1}{2T} |\tilde{x}_T(\nu)|^2 d\nu, \quad (7.29)$$

e se mediamo anche sull'ensemble otteniamo

$$\langle \overline{x^2} \rangle = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} \langle x^2(t) \rangle dt = \int_{-\infty}^{+\infty} \lim_{T \rightarrow \infty} \frac{1}{2T} \langle |\tilde{x}_T(\nu)|^2 \rangle d\nu. \quad (7.30)$$

Se definiamo la *densità spettrale di potenza* (anziché di *energia*) mediata sul tempo e sull'ensemble $G(\nu)$ come

$$G(\nu) = \lim_{T \rightarrow \infty} \frac{1}{2T} \langle |\tilde{x}_T(\nu)|^2 \rangle = \lim_{T \rightarrow \infty} \frac{1}{2T} \left\langle \left| \int_{-T}^{+T} x(t) e^{-2\pi i \nu t} dt \right|^2 \right\rangle \quad (7.31)$$

possiamo scrivere

$$\langle \overline{x^2} \rangle = \int_{-\infty}^{+\infty} G(\nu) d\nu. \quad (7.32)$$

Questa è l'estensione del teorema di Parseval ai processi stocastici stazionari. Per un processo stocastico stazionario le operazioni di media temporale e di media sull'ensemble commutano, e

$$\langle \overline{x^2} \rangle = \langle x^2 \rangle = \overline{x^2}. \quad (7.33)$$

Per esempio, nel caso del rumore Johnson, prendendo come x la differenza di potenziale V ai capi della resistenza R , avremo

$$G(\nu) = 4RkT. \quad (7.34)$$

7.3 Autocorrelazione di una funzione stocastica

Se abbiamo due funzioni del tempo reali e quadraticamente integrabili, $f(t)$ e $g(t)$, la loro funzione di correlazione $C_{fg}(\tau)$ è definita come

$$C_{fg}(\tau) = \int_{-\infty}^{+\infty} f(t) g(t + \tau) dt. \quad (7.35)$$

La trasformata di Fourier di $C_{fg}(\tau)$ ha una proprietà importante che ricaviamo qui sotto. Partiamo dalla definizione

$$\tilde{C}_{fg}(\nu) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} f(t) g(t + \tau) e^{-2\pi i \nu \tau} dt d\tau, \quad (7.36)$$

e introduciamo la nuova variabile $\vartheta = t + \tau$, così che $\tau = \vartheta - t$. La trasformata diventa

$$\begin{aligned} \tilde{C}_{fg}(\nu) &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} f(t) g(\vartheta) e^{-2\pi i \nu \vartheta} e^{2\pi i \nu t} dt d\vartheta = \int_{-\infty}^{+\infty} f(t) e^{2\pi i \nu t} dt \int_{-\infty}^{+\infty} g(\vartheta) e^{-2\pi i \nu \vartheta} d\vartheta \\ &= \tilde{f}(-\nu) \tilde{g}(\nu) = \tilde{f}^*(\nu) \tilde{g}(\nu), \end{aligned} \quad (7.37)$$

dove l'ultima riga tiene conto del fatto che, se $f(t)$ è reale, vale la proprietà

$$\tilde{f}(-\nu) = \tilde{f}^*(\nu).$$

Così, la trasformata di Fourier della funzione di correlazione di $f(t)$ e $g(t)$ è uguale al prodotto del complesso coniugato della trasformata di Fourier di $f(t)$ per la trasformata di Fourier di $g(t)$. Questo è il *teorema di correlazione*. La formula inversa è

$$C_{fg}(\tau) = \int_{-\infty}^{+\infty} \tilde{f}^*(\nu) \tilde{g}(\nu) e^{2\pi i \nu \tau} d\nu, \quad (7.38)$$

questa formula è di grande interesse pratico perché esistono algoritmi computazionali veloci per il calcolo delle trasformate (e antitrasformate) di Fourier, detti algoritmi di FFT (Fast Fourier Transform), che possono quindi essere usati anche per il calcolo delle funzioni di correlazione.

La funzione di autocorrelazione $C_{ff}(\tau)$ è la funzione di correlazione di $f(t)$ con sé stessa, definita da

$$C_{ff}(\tau) = \int_{-\infty}^{+\infty} f(t) f(t + \tau) dt. \quad (7.39)$$

La funzione di autocorrelazione dà la correlazione tra istanti separati da un intervallo di tempo τ di un processo descritto dalla funzione $f(t)$. In base alle (7.38) e (7.23) abbiamo

$$C_{ff}(\tau) = \int_{-\infty}^{+\infty} \tilde{f}^*(\nu) \tilde{f}(\nu) e^{2\pi i \nu \tau} d\nu = \int_{-\infty}^{+\infty} \Phi(\nu) e^{2\pi i \nu \tau} d\nu, \quad (7.40)$$

quindi la funzione di autocorrelazione di $f(t)$ è l'antitrasformata di Fourier della sua densità spettrale. Questo è il *teorema di Wiener-Khinchin* (Александр Яковлевич Хинчин).

Questo per una funzione quadraticamente integrabile. Ma se proviamo a definire la funzione di autocorrelazione per una funzione stocastica stazionaria $x(t)$ partendo dalla (7.39), cioè come

$$C_{xx}(\tau) = \int_{-\infty}^{+\infty} x(t) x(t + \tau) dt,$$

otteniamo una funzione di autocorrelazione che può divergere, e che, in particolare, divergerà per $\tau = 0$. Analogamente a quanto fatto nel paragrafo precedente 7.3, possiamo però definire una *funzione di autocorrelazione mediata nel tempo e mediata sull'ensemble*

$$C_{xx}(\tau) = \left\langle \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} x(t) x(t + \tau) dt \right\rangle. \quad (7.41)$$

Con procedimento analogo a quanto fatto per ottenere la (7.37) possiamo calcolare la trasformata di Fourier di $C_{xx}(\tau)$

$$\begin{aligned} \tilde{C}_{xx}(\nu) &= \int_{-\infty}^{+\infty} \left\langle \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} x(t) x(t + \tau) dt \right\rangle e^{-2\pi i \nu \tau} d\tau \\ &= \lim_{T' \rightarrow \infty} \int_{-T'}^{+T'} \left\langle \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} x(t) x(t + \tau) dt \right\rangle e^{-2\pi i \nu \tau} d\tau \\ &= \left\langle \lim_{T' \rightarrow \infty} \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T'}^{+T'} d\tau \int_{-T}^{+T} x(t) x(t + \tau) e^{-2\pi i \nu \tau} dt \right\rangle. \end{aligned} \quad (7.42)$$

A questo punto introduciamo il cambiamento di variabili $\vartheta = t + \tau$, così che $\tau = \vartheta - t$,

$$\begin{aligned}\tilde{C}_{xx}(\nu) &= \left\langle \lim_{T' \rightarrow \infty} \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T'+t}^{T'+t} d\vartheta \int_{-T}^{+T} x(t)x(\vartheta) dt e^{-2\pi i \nu \vartheta} e^{2\pi i \nu t} \right\rangle \\ &= \left\langle \lim_{T' \rightarrow \infty} \int_{-T'+t}^{T'+t} x(\vartheta) e^{-2\pi i \nu \vartheta} d\vartheta \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} x(t) e^{2\pi i \nu t} dt \right\rangle.\end{aligned}\quad (7.43)$$

Con l'ipotesi che la $x(t)$ sia una variabile aleatoria stazionaria possiamo traslare gli estremi di integrazione per $d\vartheta$, ottenendo

$$\tilde{C}_{xx}(\nu) = \left\langle \lim_{T' \rightarrow \infty} \int_{-T'}^{+T'} x(\vartheta) e^{-2\pi i \nu \vartheta} d\vartheta \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} x(t) e^{2\pi i \nu t} dt \right\rangle.$$

Poiché i due integrali sono indipendenti, possiamo sostituire T' con T negli estremi del primo integrale, e scrivere un limite unico

$$\begin{aligned}\tilde{C}_{xx}(\nu) &= \left\langle \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} x(\vartheta) e^{-2\pi i \nu \vartheta} d\vartheta \int_{-T}^{+T} x(t) e^{2\pi i \nu t} dt \right\rangle \\ &= \left\langle \lim_{T \rightarrow \infty} \frac{1}{2T} \tilde{x}_T(\nu) \tilde{x}_T(-\nu) \right\rangle = \lim_{T \rightarrow \infty} \frac{1}{2T} \langle \tilde{x}_T(\nu) \tilde{x}_T(-\nu) \rangle \\ &= \lim_{T \rightarrow \infty} \frac{1}{2T} \langle |\tilde{x}_T(\nu)|^2 \rangle = G(\nu),\end{aligned}\quad (7.44)$$

nel penultimo passaggio abbiamo tenuto conto del fatto che $x(t)$ è reale, e quindi $\tilde{x}(-\nu) = \tilde{x}^*(\nu)$, mentre nell'ultimo passaggio abbiamo usato la (7.31).

Abbiamo quindi

$$\tilde{C}_{xx}(\nu) = G(\nu), \quad \text{da cui} \quad C_{xx}(\tau) = \int_{-\infty}^{+\infty} G(\nu) e^{2\pi i \nu \tau} d\nu. \quad (7.45)$$

Se ricordiamo ancora che essendo $x(t)$ reale, abbiamo $\tilde{x}^*(\nu) = \tilde{x}(-\nu)$, vediamo che $\tilde{x}^*(\nu) \tilde{x}(\nu)$, e quindi $G(\nu)$, è una funzione pari, per cui

$$C_{xx}(\tau) = \int_{-\infty}^{+\infty} G(\nu) e^{2\pi i \nu \tau} d\nu = 2 \int_0^{+\infty} G(\nu) \cos(2\pi \nu \tau) d\nu, \quad (7.46)$$

mentre la relazione inversa è

$$G(\nu) = 2 \int_0^{+\infty} C_{xx}(\tau) \cos(2\pi \nu \tau) d\tau. \quad (7.47)$$

Questa è l'estensione del teorema di Wiener-Khinchin ai processi stocastici stazionari.

Come esempio facciamo l'ipotesi che la funzione di autocorrelazione di un sistema sia

$$C(\tau) = C_0 e^{-\tau/\tau_c}. \quad (7.48)$$

Da questa possiamo ricavare la densità spettrale $G(\nu)$ dalla relazione

$$G(\nu) = 2 \int_0^{+\infty} C_0 e^{-\tau/\tau_c} \cos(2\pi \nu \tau) d\tau = 2 C_0 \frac{1/\tau_c}{1/\tau_c^2 + 4\pi^2 \nu^2} = 2 C_0 \tau_c \frac{\Delta \nu^2}{\Delta \nu^2 + \nu^2}. \quad (7.49)$$

La rappresentazione grafica di $G(\nu)$ è quindi una lorentziana centrata sullo zero di semilarghezza $\Delta \nu = 1/(2\pi \tau_c)$, che corrisponde al *rumore bianco*.

7.4 L'equazione di Fokker-Planck

Consideriamo una variabile stocastica x che evolva sotto l'azione di processi aleatori, ma stazionari e markoviani. Sotto queste ipotesi esiste una densità di probabilità stazionaria $W(x)$ per i suoi valori possibili. Supponiamo però che all'istante t_0 la densità di probabilità $W(x, t_0)$ sia diversa dalla densità all'equilibrio $W(x)$, per esempio a causa di una perturbazione. Vogliamo studiare come, a partire da t_0 , $W(x, t)$ evolve nel tempo fino a raggiungere la forma di equilibrio $W(x)$. Per questo faremo uso della densità di probabilità condizionata $P_{1|1}(x, t|x_0)$ definita nel paragrafo 7.1. Come visto nel paragrafo 7.1, in presenza di processi evolutivi stazionari $P_{1|1}(x, t|x_0)$ dipenderà solo dall'intervallo di tempo trascorso t , e non dall'istante "condizionante" in cui si realizza il valore x_0 . Avendo sempre a che fare con un solo evento condizionato e un solo evento condizionante, in quanto segue trascureremo il pedice 1|1 e scriveremo $P(x, t|x_0)$. Continuando a supporre, come nel paragrafo 7.1, che la dipendenza di x dal tempo sia sì stocastica, ma continua, al decrescere di t il valore di x dovrà tendere a x_0 , per cui

$$\lim_{t \rightarrow 0} P(x, t|x_0) = \delta(x - x_0). \quad (7.50)$$

All'estremo opposto, per $t \rightarrow \infty$, la x perderà ogni memoria del suo valore iniziale x_0 , così che

$$\lim_{t \rightarrow \infty} P(x, t|x_0) = W(x), \quad (7.51)$$

dove $W(x)$ è la densità di probabilità all'equilibrio menzionata sopra. Quindi al limite $t \rightarrow 0$ prevale la continuità dalle condizioni iniziali, mentre al limite $t \rightarrow \infty$ prevale la stocasticità stazionaria. Quello che vogliamo studiare è come $P(x, t|x_0)$ evolve nel tempo dal limite (7.50) al limite (7.51), cioè il suo comportamento per valori di t intermedi, già lontani da zero ma ancora non sufficientemente grandi perché si sia persa completamente la memoria del valore iniziale x_0 . Cominciamo scrivendo la derivata di $P(x, t|x_0)$ come limite di rapporto incrementale

$$\frac{\partial P(x, t|x_0)}{\partial t} = \lim_{\Delta t \rightarrow 0} \frac{P(x, t + \Delta t|x_0) - P(x, t|x_0)}{\Delta t}. \quad (7.52)$$

Per l'equazione di Chapman-Kolmogorov, o Smoluchowski, (7.21), la $P(x, t + \Delta t|x_0)$ può essere scritta sotto forma di integrale su uno stato intermedio x' , e il rapporto incrementale diventa

$$\frac{\partial P(x, t|x_0)}{\partial t} = \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \left[\int_{-\infty}^{+\infty} P(x, \Delta t|x') P(x', t|x_0) dx' - P(x, t|x_0) \right], \quad (7.53)$$

che, introducendo la nuova variabile $\Delta x = x - x'$, diventa a sua volta

$$\frac{\partial P(x, t|x_0)}{\partial t} = \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \left[\int_{-\infty}^{+\infty} P(x, \Delta t|x - \Delta x) P(x - \Delta x, t|x_0) d\Delta x - P(x, t|x_0) \right]. \quad (7.54)$$

La condizione che x vari nel tempo con continuità impone che, per Δt piccoli, la densità di probabilità condizionata $P(x, \Delta t|x - \Delta x)$ tenda rapidamente a zero al crescere di Δx . All'integrale che compare nella (7.54) contribuiranno quindi solo i valori di Δx molto piccoli, e ci aspettiamo che il prodotto $P(x, \Delta t|x - \Delta x) P(x - \Delta x, t|x_0)$ possa essere sviluppato in serie di Taylor in Δx fermandoci a potenze basse. Cominciamo riscrivendo il prodotto nella forma

$$P(x, \Delta t|x - \Delta x) P(x - \Delta x, t|x_0) = P(x + \Delta x - \Delta x, \Delta t|x - \Delta x) P(x - \Delta x, t|x_0), \quad (7.55)$$

in modo che la variabile x compaia ovunque nella forma $x - \Delta x$, ed effettuiamo lo sviluppo in serie rispetto all'incremento $-\Delta x$ (una derivazione alternativa si trova nell'Appendice A)

$$P(x, \Delta t | x - \Delta x) P(x - \Delta x, t | x_0) = \sum_{n=0}^{\infty} \frac{(-\Delta x)^n}{n!} \frac{\partial^n}{\partial x^n} [P(x + \Delta x, \Delta t | x) P(x, t | x_0)], \quad (7.56)$$

dove abbiamo usato la solita convenzione per la derivata di ordine zero

$$\left. \frac{\partial^0 f(x)}{\partial x^0} \right|_{\Delta x=0} = f(x). \quad (7.57)$$

Con questa espansione l'integrale all'interno delle parentesi quadre della (7.54) diventa

$$\int_{-\infty}^{+\infty} P(x, \Delta t | x - \Delta x) P(x - \Delta x, t | x_0) d\Delta x = \sum_{n=0}^{\infty} \frac{(-1)^n}{n!} \frac{\partial^n}{\partial x^n} \left[P(x, t | x_0) \int_{-\infty}^{+\infty} P(x + \Delta x, \Delta t | x) \Delta x^n d\Delta x \right], \quad (7.58)$$

dove le operazioni di derivazione n -esima rispetto a x e la funzione $P(x, t | x_0)$, indipendenti da Δx , sono state portate fuori dall'integrale in Δx . L'addendo corrispondente a $n = 0$ vale

$$P(x, t | x_0) \int_{-\infty}^{+\infty} P(x + \Delta x, \Delta t | x) d\Delta x = P(x, t | x_0) \quad (7.59)$$

per la condizione di normalizzazione di P , e si cancella con l'ultimo termine all'interno delle parentesi quadre della (7.54), che può essere riscritta

$$\frac{\partial P(x, t | x_0)}{\partial t} = \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \sum_{n=1}^{\infty} \frac{(-1)^n}{n!} \frac{\partial^n}{\partial x^n} \left[P(x, t | x_0) \int_{-\infty}^{+\infty} P(x + \Delta x, \Delta t | x) \Delta x^n d\Delta x \right]. \quad (7.60)$$

A questo punto definiamo i *coefficienti di Kramers-Moyal* di ordine n della variabile x' , $D_n(x')$,

$$D_n(x') = \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \int_{-\infty}^{+\infty} dx (x - x')^n P(x, \Delta t | x') = \lim_{\Delta t \rightarrow 0} \frac{\langle (x - x')^n \rangle}{\Delta t}. \quad (7.61)$$

Il coefficiente D_1 è la velocità con cui il valor medio della variabile stocastica x si allontana dal valore iniziale x' (deriva, o *drift*), D_2 è la velocità con cui cresce la varianza di x , e così via. Introducendo i coefficienti (7.61) nella (7.60) abbiamo

$$\frac{\partial P(x, t | x_0)}{\partial t} = \sum_{n=1}^{\infty} \frac{(-1)^n}{n!} \frac{\partial^n}{\partial x^n} [D_n(x) P(x, t | x_0)], \quad (7.62)$$

che è l'equazione di Kramers-Moyal, che non considereremo nel caso generale. In molti casi di interesse, al tendere a zero di Δt la quantità $\langle (x - x')^n \rangle$ della (7.61) tende a zero più rapidamente di Δt se $n > 2$. In queste condizioni i termini che coinvolgono i coefficienti D_n con $n > 2$ possono essere trascurati, e ci rimane l'*equazione di Fokker-Planck*

$$\frac{\partial P(x, t | x_0)}{\partial t} = -\frac{\partial}{\partial x} [D_1 P(x, t | x_0)] + \frac{1}{2} \frac{\partial^2}{\partial x^2} [D_2 P(x, t | x_0)], \quad (7.63)$$

dove, come abbiamo visto commentando la (7.61), D_1 ed D_2 sono le velocità con cui variano, rispettivamente, il valor medio e la varianza di x .

In un caso semplice, che corrisponde a situazioni sperimentali di interesse, come, per esempio, il moto browniano che considereremo tra poco, i coefficienti D_1 e D_2 hanno la forma

$$D_1(x) = -\gamma x, \quad (7.64)$$

$$D_2(x) = 2D, \quad (7.65)$$

e l'equazione di Fokker-Planck diventa

$$\frac{\partial P(x, t|x_0)}{\partial t} = \gamma \frac{\partial}{\partial x} [xP(x, t|x_0)] + D \frac{\partial^2}{\partial x^2} P(x, t|x_0). \quad (7.66)$$

Il primo termine al secondo membro di questa equazione è detto *termine di drift* (o di *deriva*), e il secondo *termine di diffusione*. Infatti, se poniamo $D = 0$ la (7.66) si riduce a

$$\frac{\partial P(x, t|x_0)}{\partial t} = \gamma \frac{\partial}{\partial x} [xP(x, t|x_0)], \quad (7.67)$$

che è l'equazione di drift. Invece per $\gamma = 0$ la (7.66) diventa

$$\frac{\partial P(x, t|x_0)}{\partial t} = D \frac{\partial^2}{\partial x^2} P(x, t|x_0), \quad (7.68)$$

che è l'equazione di diffusione. Discuteremo queste due equazioni un po' più in dettaglio nei due paragrafi seguenti.

7.4.1 Equazione di drift

Per farci un'idea un po' più chiara del tipo di problema trattato, cerchiamo una soluzione per l'equazione di drift con la condizione iniziale $P(x, 0|x_0) = \delta(x - x_0)$, corrispondente alla certezza che la variabile stocastica abbia inizialmente il valore x_0 . Sviluppando la derivata rispetto a x la (7.67) diventa

$$\frac{\partial P(x, t|x_0)}{\partial t} = \gamma P(x, t|x_0) + \gamma x \frac{\partial P(x, t|x_0)}{\partial x}. \quad (7.69)$$

Controlliamo se esiste una soluzione nella forma $P(x, t|x_0) = \delta[x - \xi(t)]$, con $\xi(0) = x_0$, corrispondente a una densità di probabilità a delta di Dirac che si sposta con il tempo. Scriviamo la delta di Dirac ipotizzata come limite di una lorentziana centrata attorno a $\xi(t)$

$$\delta[x - \xi(t)] = \lim_{\varepsilon \rightarrow 0} \frac{\varepsilon}{\pi} \frac{1}{[x - \xi(t)]^2 + \varepsilon^2}. \quad (7.70)$$

Sostituiamo il secondo membro della (7.70) nella (7.69) prima di fare il limite per $\varepsilon \rightarrow 0$, limite che faremo alla fine di tutti i passaggi. Le derivate rispetto a t e rispetto a x sono rispettivamente

$$\begin{aligned} \frac{\partial P(x, t|x_0)}{\partial t} &= \frac{\varepsilon}{\pi} \frac{2[x - \xi(t)]}{\{[x - \xi(t)]^2 + \varepsilon^2\}^2} \frac{d\xi}{dt}, \quad \text{e} \\ \frac{\partial P(x, t|x_0)}{\partial x} &= -\frac{\varepsilon}{\pi} \frac{2[x - \xi(t)]}{\{[x - \xi(t)]^2 + \varepsilon^2\}^2}, \end{aligned} \quad (7.71)$$

che, sostituite nella (7.69), ci danno

$$\frac{\varepsilon}{\pi} \frac{2[x - \xi(t)]}{\{[x - \xi(t)]^2 + \varepsilon^2\}^2} \frac{d\xi}{dt} = \gamma \frac{\varepsilon}{\pi} \frac{1}{[x - \xi(t)]^2 + \varepsilon^2} - \gamma x \frac{\varepsilon}{\pi} \frac{2[x - \xi(t)]}{\{[x - \xi(t)]^2 + \varepsilon^2\}^2}. \quad (7.72)$$

Moltiplicando primo e secondo membro per $\frac{\pi\{[x - \xi(t)]^2 + \varepsilon^2\}^2}{2\varepsilon[x - \xi(t)]}$, operazione legittima finché $\varepsilon \neq 0$ e $x \neq \xi(t)$, ci rimane

$$\frac{d\xi}{dt} + \gamma x = \gamma \frac{[x - \xi(t)]^2 + \varepsilon^2}{2[x - \xi(t)]}. \quad (7.73)$$

Adesso possiamo effettuare il limite $\varepsilon \rightarrow 0$, ottenendo

$$\frac{d\xi}{dt} + \gamma x = \gamma \frac{x - \xi(t)}{2}, \quad (7.74)$$

ma a questo limite la nostra distribuzione di probabilità diventa $\delta[x - \xi(t)]$, quindi ci interessa solo il caso $x = \xi(t)$ e la nostra equazione si riduce a

$$\frac{d\xi}{dt} + \gamma \xi = 0, \quad \text{che ha soluzione} \quad \xi(t) = \xi(0) e^{-\gamma t} = x_0 e^{-\gamma t}. \quad (7.75)$$

La nostra distribuzione rimane quindi una delta di Dirac che si sposta asintoticamente (cioè deriva, o “drifta”) dal valore iniziale x_0 verso il valore di equilibrio $x = 0$.

7.4.2 Equazione di diffusione

Anche per l'equazione di diffusione (7.68) cerchiamo la soluzione con la condizione iniziale $P(x, 0|x_0) = \delta(x - x_0)$. Per sostituzione si può controllare che questa è

$$P(x, t|x_0) = \frac{1}{\sqrt{4\pi Dt}} e^{-(x-x_0)^2/4Dt}, \quad (7.76)$$

cioè una distribuzione gaussiana che tende a una δ di Dirac centrata su x_0 per $t \rightarrow 0$, in accordo con la condizione iniziale, mentre si allarga proporzionalmente a $\sqrt{Dt}$ al crescere del tempo (diffusione), sempre restando centrata attorno a x_0 . Se per semplicità spostiamo l'origine delle coordinate in modo che sia $x_0 = 0$, la (7.76) ci dice che il valor medio della variabile stocastica $\langle x \rangle$ è costantemente nullo, mentre per il valor medio di x^2 (varianza) abbiamo

$$\langle x^2 \rangle = \frac{1}{\sqrt{4\pi Dt}} \int_{-\infty}^{+\infty} x^2 e^{-x^2/4Dt} dx = 2Dt, \quad (7.77)$$

quindi la varianza $\langle x^2 \rangle$ aumenta proporzionalmente al tempo. La (7.76) può essere utilizzata per risolvere anche il problema con la condizione iniziale generale $P(x, 0|x_0) = g(x)$, con $g(x)$ densità di probabilità generica. Ponendo

$$\Phi(x, t) = \frac{1}{\sqrt{4\pi Dt}} e^{-(x-x_0)^2/4Dt}$$

e utilizzando il fatto che $g(x)$ non dipende dal tempo, possiamo infatti scrivere la soluzione per l'equazione di diffusione con la nuova condizione iniziale mediante la convoluzione

$$P(x, t|x_0) = \int_{-\infty}^{+\infty} \Phi(x - \xi, t) g(\xi) d\xi. \quad (7.78)$$

7.5 Moto browniano

Uno dei fenomeni stocastici più semplici è il *moto browniano*, cioè il moto disordinato di *particelle mesoscopiche* sospese in un fluido. Per particelle mesoscopiche intendiamo particelle con diametro dell'ordine del μm , quindi grandi rispetto alle molecole, ma ancora sufficientemente piccole da essere soggette a fluttuazioni non osservabili a livello macroscopico. Il fenomeno era già stato osservato dall'olandese Jan Ingenhousz nel 1785, ma venne riscoperto nel 1828 dal botanico scozzese Robert Brown studiando il moto del polline in una sospensione acquosa, e da allora porta il suo nome. Una trattazione matematica rigorosa si ebbe solo agli inizi del novecento, prima con la tesi di laurea (1900) di Louis Bachelier, considerato poi il fondatore della matematica finanziaria, e poco dopo con il lavoro di Albert Einstein del 1905.

Consideriamo una particella mesoscopica di massa m immersa in un fluido a temperatura T . Le molecole del fluido, ognuna di massa μ , per ipotesi molto minore di m , sono in agitazione termica e sottopongono la particella macroscopica a frequentissimi urti. Questa è così sottoposta ad una forza aleatoria $\mathbf{F}(t)$. Per semplicità consideriamo il moto in una dimensione, e scriviamo

$$m \frac{dv}{dt} = F(t). \quad (7.79)$$

Escludiamo che gli urti cambino gli stati interni della particella mesoscopica e/o delle molecole del fluido, e che quindi siano urti perfettamente elastici. Chiamando v_i e v_f le velocità della particella mesoscopica rispettivamente prima e dopo l'urto, e V la velocità della molecola del fluido prima dell'urto, la teoria degli urti elastici ci dice che

$$v_f = \frac{(m - \mu) v_i + 2\mu V}{m + \mu}. \quad (7.80)$$

Quindi, a causa dell'urto, la particella subisce una variazione della quantità di moto pari a

$$\Delta p = m(v_f - v_i) = \frac{2\mu m}{m + \mu} (V - v_i). \quad (7.81)$$

Se adesso facciamo una media sull'ensemble, la media $\langle V \rangle$ delle velocità delle molecole del fluido, in agitazione termica e in assenza di direzioni privilegiate, sarà nulla. Abbiamo così

$$\langle \Delta p \rangle = -\frac{2\mu m}{m + \mu} \langle v_i \rangle, \quad (7.82)$$

cioè la variazione media subita dalla quantità di moto della particella mesoscopica in un urto non è nulla, ma è proporzionale a meno la velocità media prima dell'urto. Infatti gli urti devono portare il sistema all'equilibrio termico, al quale, in assenza di direzioni privilegiate, la media vettoriale delle velocità delle particelle mesoscopiche (e delle molecole) deve essere nulla. Ci conviene così scomporre la forza aleatoria $F(t)$ della (7.79) nella somma di due termini

$$F(t) = -\gamma v + f(t), \quad (7.83)$$

di cui il primo corrisponde ad una forza non aleatoria di tipo viscoso a media non nulla, mentre il secondo corrisponde ad una forza aleatoria a media nulla. E' stato Einstein, nel 1905, a notare per primo che le stesse forze aleatorie responsabili del moto stocastico delle particelle nel moto browniano

erano anche all'origine della forza viscosa che compare su una particella che viene trascinata nel fluido. La (7.79) si trasforma così nell'equazione di Langevin

$$m \frac{dv}{dt} = -\gamma v + f(t). \quad (7.84)$$

Sulla forza aleatoria $f(t)$ risultante dagli urti contro le molecole del fluido, una volta sottratto il termine viscoso, possiamo fare le seguenti ipotesi:

1. *Isotropia*: gli urti, e quindi la forza, non hanno direzioni privilegiate così che

$$\langle f(t) \rangle = 0. \quad (7.85)$$

2. *Scorrelazione*: la forza è il risultato di collisioni casuali di durata τ_c molto corta, quindi fluttua continuamente e in ogni momento non è correlata con il suo valore ad un istante precedente, e possiamo scrivere

$$\langle f(t)f(t+\tau) \rangle = 2K\delta(\tau), \quad (7.86)$$

dove K è una costante e δ la delta di Dirac.

3. *Gaussianità*: la forza è il risultato di un numero molto alto di eventi tra di loro indipendenti (gli urti delle molecole del fluido contro la particella) e quindi, per il teorema centrale del limite (vedi Appendice B), avrà una distribuzione di probabilità gaussiana.

Se dividiamo primo e secondo membro della (7.84) per la massa della particella mesoscopica m , e introduciamo le quantità

$$A(t) = \frac{f(t)}{m}, \quad \alpha = \frac{\gamma}{m} \quad \text{e} \quad C = \frac{K}{m^2}, \quad (7.87)$$

(dove $A(t)$ e αv hanno le dimensioni di un'accelerazione) l'equazione di Langevin diventa

$$\frac{dv}{dt} = -\alpha v + A(t), \quad (7.88)$$

mentre le condizioni (7.85) e (7.86) si riscrivono

$$\langle A(t) \rangle = 0 \quad \text{e} \quad \langle A(t)A(t+\tau) \rangle = 2C\delta(\tau). \quad (7.89)$$

7.5.1 Distribuzione delle velocità

Il processo che stiamo descrivendo è markoviano. La (7.88) è un'equazione differenziale lineare a coefficienti costanti non omogenea, quindi risolviamo prima l'equazione omogenea associata, che ha soluzione generale $v = v_0 e^{-\alpha t}$, poi cerchiamo una soluzione particolare dell'equazione completa. Ipotizziamo l'esistenza di una soluzione del tipo $v(t) = w(t) e^{-\alpha t}$ e proviamo a sostituirla nella (7.88), ottenendo

$$\frac{dw}{dt} e^{-\alpha t} - \alpha w e^{-\alpha t} = -\alpha w e^{-\alpha t} + A(t), \quad (7.90)$$

da cui, semplificando e moltiplicando primo e secondo membro per $e^{\alpha t}$, otteniamo

$$\frac{dw}{dt} = A(t) e^{\alpha t}. \quad (7.91)$$

L'espressione ipotizzata $v(t) = w(t) e^{-\alpha t}$ è quindi effettivamente una soluzione particolare della (7.88) a condizione che sia

$$w(t) = \int_0^t A(t') e^{\alpha t'} dt'. \quad (7.92)$$

Sommando questa soluzione particolare alla soluzione generale dell'equazione omogenea associata otteniamo la soluzione generale dell'equazione completa (7.88)

$$v(t) = v_0 e^{-\alpha t} + e^{-\alpha t} \int_0^t A(t') e^{\alpha t'} dt', \quad (7.93)$$

dove v_0 è la velocità della particella a $t = 0$. Facendo la media delle velocità sull'ensemble abbiamo

$$\begin{aligned} \langle v(t) \rangle &= v_0 e^{-\alpha t} + e^{-\alpha t} \int_0^t \underbrace{\langle A(t') \rangle}_{= 0 \text{ per la (7.89)}} e^{\alpha t'} dt' = v_0 e^{-\alpha t}. \end{aligned} \quad (7.94)$$

Notiamo che la (7.94) è valida indipendentemente dall'istante iniziale. Per cui, se chiamiamo v la velocità della particella ad un certo istante t , chiamiamo u la velocità a un istante successivo $t + \Delta t$, e $\langle u \rangle$ la media sui possibili valori di u una volta nota v , abbiamo

$$\langle u \rangle = v e^{-\alpha \Delta t}. \quad (7.95)$$

Per scrivere l'equazione di Fokker-Planck per la distribuzione delle velocità ci servono i due coefficienti di Kramers-Moyal $D_1(v)$ e $D_2(v)$. Se continuiamo a chiamare v la velocità all'istante t e u la velocità a $t + \Delta t$, dalla (7.61) abbiamo per $D_1(v)$

$$D_1(v) = \lim_{\Delta t \rightarrow 0} \frac{\langle u - v \rangle}{\Delta t} = \lim_{\Delta t \rightarrow 0} \frac{\langle u \rangle - v}{\Delta t}, \quad (7.96)$$

dove la media è fatta sui possibili valori di u , dato il valore di v . Dovendo fare il limite $\Delta t \rightarrow 0$ possiamo approssimare la (7.95) ponendo $\langle u \rangle \simeq v(1 - \alpha \Delta t)$, da cui

$$D_1(v) = \lim_{\Delta t \rightarrow 0} \frac{v(1 - \alpha \Delta t) - v}{\Delta t} = -\alpha v, \quad (7.97)$$

e siamo così nel caso (7.64). Per $D_2(v)$ abbiamo invece, sempre dalla (7.61),

$$D_2(v) = \lim_{\Delta t \rightarrow 0} \frac{\langle (u - v)^2 \rangle}{\Delta t} = \lim_{\Delta t \rightarrow 0} \frac{\langle u^2 \rangle - 2v\langle u \rangle + v^2}{\Delta t}, \quad (7.98)$$

dove ancora una volta la media è fatta sui possibili valori di u , una volta nota v . Possiamo ancora scrivere $\langle u \rangle \simeq v(1 - \alpha \Delta t)$ mentre, per $\langle u^2 \rangle$, abbiamo dalla (7.93)

$$\begin{aligned} \langle u^2 \rangle &= v^2 e^{-2\alpha \Delta t} + 2e^{-2\alpha \Delta t} \int_0^{\Delta t} v \langle A(t') \rangle e^{\alpha t'} dt' + e^{-2\alpha \Delta t} \left\langle \int_0^{\Delta t} e^{\alpha t'} A(t') dt' \int_0^{\Delta t} e^{\alpha t''} A(t'') dt'' \right\rangle \\ &= v^2 e^{-2\alpha \Delta t} + 2e^{-2\alpha \Delta t} \int_0^{\Delta t} v \langle A(t') \rangle e^{\alpha t'} dt' \\ &\quad + e^{-2\alpha \Delta t} \int_0^{\Delta t} dt' \int_0^{\Delta t} dt'' e^{\alpha(t' + t'')} \langle A(t') A(t'') \rangle. \end{aligned} \quad (7.99)$$

Ricordando che, per le (7.89), abbiamo $\langle A(t') \rangle = 0$ indipendentemente da t' , e $\langle A(t') A(t'') \rangle = 2C\delta(t' - t'')$, dopo un'integrazione in dt'' la (7.99) si riduce a

$$\begin{aligned}\langle u^2 \rangle &= v^2 e^{-2\alpha\Delta t} + 2Ce^{-2\alpha\Delta t} \int_0^{\Delta t} e^{2\alpha t'} dt' = v^2 e^{-2\alpha\Delta t} + 2Ce^{-2\alpha\Delta t} \frac{e^{2\alpha\Delta t} - 1}{2\alpha} \\ &= v^2 e^{-2\alpha\Delta t} + \frac{C}{\alpha} (1 - e^{-2\alpha\Delta t}) \simeq v^2(1 - 2\alpha\Delta t) + 2C\Delta t.\end{aligned}\quad (7.100)$$

Inserendo questa espressione per $\langle u^2 \rangle$ nella (7.98) otteniamo

$$\begin{aligned}D_2(v) &= \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \left[v^2(1 - 2\alpha\Delta t) + 2C\Delta t - 2v^2(1 - \alpha\Delta t) + v^2 \right] \\ &= \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} 2C\Delta t = 2C,\end{aligned}\quad (7.101)$$

e siamo così nel caso (7.65). Sostituendo le espressioni trovate per $D_1(v)$ e $D_2(v)$ nell'equazione di Fokker-Planck (7.66) otteniamo l'equazione per il moto browniano

$$\frac{\partial P(v, t|v_0)}{\partial t} = -\alpha \frac{\partial}{\partial v} [v P(v, t|v_0)] + C \frac{\partial^2}{\partial v^2} P(v, t|v_0). \quad (7.102)$$

Se ipotizziamo una condizione iniziale per la distribuzione delle velocità a δ di Dirac, del tipo $\lim_{t \rightarrow 0} P(v, t|v_0) = \delta(v - v_0)$, si può controllare per sostituzione che la soluzione della (7.102) è

$$P(v, t|v_0) = \sqrt{\frac{\alpha}{2\pi C(1 - e^{-2\alpha t})}} \exp\left[-\frac{\alpha(v - v_0 e^{-\alpha t})^2}{2C(1 - e^{-2\alpha t})}\right]. \quad (7.103)$$

Questa espressione, al limite $t \rightarrow \infty$, ovvero per $t \gg 1/\alpha$, diventa

$$\lim_{t \rightarrow \infty} P(v, t|v_0) \simeq \sqrt{\frac{\alpha}{2\pi C}} e^{-\alpha v^2/(2C)}. \quad (7.104)$$

Questo limite è indipendente dalla velocità iniziale della particella v_0 , ed è ovviamente scomparsa anche la dipendenza dal tempo. Quindi se, dopo aver perturbato il sistema, aspettiamo un tempo lungo rispetto a $1/\alpha$, ritroviamo il sistema in una situazione stazionaria. Se lo stato stazionario corrisponde all'equilibrio termodinamico, la distribuzione delle velocità deve coincidere con la distribuzione di Maxwell-Boltzmann

$$W(v) = \sqrt{\frac{m}{2\pi kT}} e^{-mv^2/(2kT)}. \quad (7.105)$$

Imponendo l'uguaglianza tra la (7.104) e la (7.105) otteniamo

$$\frac{\alpha v^2}{2C} = \frac{mv^2}{2kT}, \quad \text{da cui} \quad C = \frac{kT\alpha}{m}, \quad \text{ovvero} \quad \alpha = \frac{m}{kT} C, \quad (7.106)$$

che è la relazione di Einstein tra il termine di fluttuazione C e il termine dissipativo α . Ricordando le (7.87) otteniamo le relazioni di Einstein per i coefficienti γ e D

$$\gamma = \frac{D}{kT}, \quad D = kT\gamma. \quad (7.107)$$

7.5.2 Diffusione delle particelle

Una volta studiata l'evoluzione temporale della funzione di distribuzione della velocità per il moto browniano, vediamo cosa si può dire sulla funzione di distribuzione della posizione delle particelle. Partiamo dall'equazione di Langevin (7.88), e scriviamo $\dot{x}$ anziché v

$$\frac{d\dot{x}}{dt} = -\alpha \dot{x} + A(t). \quad (7.108)$$

Trasformiamo la (7.108) moltiplicando per x primo e secondo membro

$$x \frac{d\dot{x}}{dt} = -\alpha x \dot{x} + xA(t), \quad (7.109)$$

poi, notando che

$$\frac{d}{dt}(x\dot{x}) = \dot{x}^2 + x \frac{d\dot{x}}{dt}, \quad \text{da cui} \quad x \frac{d\dot{x}}{dt} = \frac{d}{dt}(x\dot{x}) - \dot{x}^2, \quad (7.110)$$

possiamo scrivere

$$\frac{d}{dt}(x\dot{x}) - \dot{x}^2 = -\alpha x \dot{x} + xA(t). \quad (7.111)$$

Adesso effettuiamo la media sull'ensemble, o, che è lo stesso, supponiamo di avere nello stesso sistema un grande numero di particelle mesoscopiche, e mediamo su di esse:

$$\left\langle \frac{d}{dt}(x\dot{x}) \right\rangle - \langle \dot{x}^2 \rangle = -\alpha \langle x\dot{x} \rangle + \langle xA(t) \rangle. \quad (7.112)$$

Per prima cosa notiamo che, essendo $A(t)$ a media nulla indipendentemente dai valori di x e $\dot{x}$, avremo $\langle xA(t) \rangle = 0$. Poi, ricordando che stiamo considerando un caso unidimensionale, dal teorema di equipartizione dell'energia abbiamo che $\frac{1}{2}m\langle \dot{x}^2 \rangle = \frac{1}{2}kT$, indipendente dal tempo. Sostituendo tutto nella (7.112) otteniamo

$$\left\langle \frac{d}{dt}(x\dot{x}) \right\rangle - \frac{kT}{m} = -\alpha \langle x\dot{x} \rangle. \quad (7.113)$$

Poiché le operazioni di derivata rispetto al tempo e di media sull'ensemble commutano, abbiamo un'equazione differenziale per la quantità $\langle x\dot{x} \rangle$ nella forma

$$\frac{d}{dt}\langle x\dot{x} \rangle + \alpha \langle x\dot{x} \rangle = \frac{kT}{m}, \quad \text{con soluzione generale} \quad \langle x\dot{x} \rangle = Ke^{-\alpha t} + \frac{kT}{m\alpha}. \quad (7.114)$$

con K costante da determinare dalle condizioni iniziali. Facendo l'ipotesi che all'istante $t = 0$ tutte le particelle si trovino in $x = 0$, in modo che x misuri lo spostamento dalla posizione iniziale, otteniamo $K = -kT/(m\alpha)$. Ricordando poi che $dx^2/dt = 2x\dot{x}$ e integrando abbiamo

$$\langle x^2(t) \rangle = 2 \int_0^t \left(Ke^{-\alpha t'} + \frac{kT}{m\alpha} \right) dt' = \frac{2kT}{m\alpha} \left[t - \frac{1}{\alpha} (1 - e^{-\alpha t}) \right]. \quad (7.115)$$

Consideriamo due casi limite. Per $t \ll \alpha^{-1}$ possiamo porre $e^{-\alpha t} \simeq 1 - \alpha t + \frac{1}{2}\alpha^2 t^2$, da cui

$$\langle x^2(t) \rangle \simeq \frac{kT}{m} t^2, \quad (7.116)$$

in altre parole, per un breve intervallo di tempo iniziale la particella si comporta, in media, come una particella libera con velocità termica costante $v = \sqrt{kT/m}$. Consideriamo invece il caso opposto, ovvero $t \gg \alpha^{-1}$. Adesso l'esponenziale nell'ultimo membro della (7.115) può essere trascurato rispetto a 1, e abbiamo

$$\langle x^2(t) \rangle \simeq \frac{2kT}{m\alpha} t, \quad (7.117)$$

e la particella esegue un moto stocastico di diffusione, con $\langle x^2(t) \rangle \propto t$. Paragonando questo risultato alla (7.77) vediamo che abbiamo una diffusione con coefficiente

$$D = \frac{kT}{m\alpha}. \quad (7.118)$$

7.6 Shot noise

Una sorgente importante di rumore a bassa frequenza è dovuta alla quantizzazione della carica elettrica. A correnti sufficientemente basse è possibile identificare il contributo della carica dei singoli elettroni in un dispositivo elettronico. Questo rumore, dovuto quindi al passaggio discontinuo delle cariche, viene chiamato *shot noise*, o *rumore a colpi*, e fu identificato per la prima volta da Walter Schottky nelle valvole termoelettroniche. Va notato che il trattamento che segue sarà sviluppato sotto le ipotesi: i) che la corrente sia molto debole, e ii) che i singoli elettroni contribuiscano alla corrente in modo indipendente l'uno dall'altro. Questa indipendenza mutua si verifica, per esempio, per elettroni che attraversano una barriera, come nel caso di un diodo a giunzione, ma non nel caso di conduttori metallici, in cui esistono correlazioni a lungo range tra i portatori di carica.

Se $\langle n \rangle$ è il numero medio degli elettroni che raggiunge il rivelatore durante il tempo di misura t_m , l'intensità di corrente media vale

$$\langle I \rangle = \frac{\langle n \rangle e}{t_m}. \quad (7.119)$$

Le fluttuazioni della corrente I sono caratterizzate dalla deviazione quadratica media $\langle (\Delta I)^2 \rangle$, che è legata alla deviazione quadratica media del numero degli elettroni $\langle (\Delta n)^2 \rangle$.

Dividiamo il tempo della misura t_m in N intervallini uguali, ciascuno di durata t_m/N , sufficientemente brevi perché in ognuno di essi passi al massimo un singolo elettrone. Chiamiamo w la probabilità che nel singolo intervallino temporale passi un elettrone, e $(1 - w)$ la probabilità che non ci sia passaggio. La probabilità che durante il tempo totale t_m si verifichino n passaggi è la probabilità che in n intervallini ci sia un passaggio e in $(N - n)$ intervallini no, ed è data dall'espressione binomiale

$$P(n, t_m) = \binom{N}{n} w^n (1 - w)^{N-n} = \frac{N!}{n! (N - n)!} w^n (1 - w)^{N-n}. \quad (7.120)$$

Se $N \gg 1$ e $w \ll 1$, con la condizione che il prodotto $Nw = \langle n \rangle$ resti finito, la (7.120) può essere

approssimata dalla poissoniana. Infatti, sostituendo $w = \langle n \rangle / N$, abbiamo

$$\begin{aligned}
 \lim_{N \rightarrow \infty} P(n, t_m) &= \lim_{N \rightarrow \infty} \frac{N!}{n! (N-n)!} \left(\frac{\langle n \rangle}{N} \right)^n \left(1 - \frac{\langle n \rangle}{N} \right)^{N-n} \\
 &= \lim_{N \rightarrow \infty} \frac{N(N-1)\dots(N-n+1)}{n!} \frac{\langle n \rangle^n}{N^n} \left(1 - \frac{\langle n \rangle}{N} \right)^N \left(1 - \frac{\langle n \rangle}{N} \right)^{-n} \\
 &= \lim_{N \rightarrow \infty} \frac{N(N-1)\dots(N-n+1)}{N^n} \frac{\langle n \rangle^n}{n!} \left(1 - \frac{\langle n \rangle}{N} \right)^N \left(1 - \frac{\langle n \rangle}{N} \right)^{-n} \\
 &= 1 \cdot \frac{\langle n \rangle^n}{n!} e^{-\langle n \rangle} \cdot 1 = \frac{\langle n \rangle^n}{n!} e^{-\langle n \rangle}, \tag{7.121}
 \end{aligned}$$

avendo usato la relazione

$$\lim_{N \rightarrow \infty} \left(1 - \frac{x}{N} \right)^N = e^{-x}.$$

Per la deviazione quadratica media nella distribuzione di Poisson abbiamo

$$\langle (\Delta n)^2 \rangle = \langle n \rangle, \quad e \quad \frac{\langle (\Delta n)^2 \rangle}{\langle n \rangle^2} = \frac{1}{\langle n \rangle}, \tag{7.122}$$

da cui otteniamo per la corrente I

$$\frac{\langle (\Delta I)^2 \rangle}{\langle I \rangle^2} = \frac{\langle (\Delta n)^2 \rangle}{\langle n \rangle^2} = \frac{1}{\langle n \rangle} = \frac{e}{t_m \langle I \rangle} \tag{7.123}$$

e

$$\langle (\Delta I)^2 \rangle = \frac{e \langle I \rangle}{t_m}. \tag{7.124}$$

In questo caso la variabile stocastica non ha media nulla. Inoltre, non è vero che se aumento arbitrariamente $\langle I \rangle$ aumenta il rumore, perché se aumenta $\langle I \rangle$ gli elettroni non sono più completamente scorrelati, non sono più indipendenti. Come esempio di shot noise, consideriamo una corrente media $\langle I \rangle = 1$ mA, e un tempo di misura $t_m = 1$ ms, ricordando che la carica dell'elettrone è $|e| = 1.6 \times 10^{-19}$ C, otteniamo $\sqrt{\langle (\Delta I)^2 \rangle} = 4.0 \times 10^{-10}$ A.

Per il calcolo della distribuzione spettrale del rumore $G_I(\nu)$, utilizziamo la variabile stocastica $I(t) = e n(t)/t_m$, dove $n(t)$ è il numero di elettroni effettivamente passati nell'intervallo di tempo $(t, t + t_m)$ associato alla misura. Immaginiamo che il contributo di ogni singolo elettrone alla corrente totale sia rappresentato da una funzione $F(t)$ diversa da zero solo per il breve periodo di tempo $(-\tau_c, +\tau_c)$, con $\tau_c \ll t_m$, durante il quale si ha il passaggio dell'elettrone, e che si abbia

$$\int_{-\tau_c}^{+\tau_c} F(t) dt = e. \tag{7.125}$$

La dipendenza temporale “istantanea” della corrente durante l'intervallo di tempo di misura t_m , nel quale passano n elettroni, può così essere scritta

$$I(t) = \sum_{i=1}^n F(t - t_i), \tag{7.126}$$

dove t_i , “istante” del passaggio dello i -esimo elettrone, è una variabile stocastica vincolata solo dalla condizione $0 < t_i < t_m$. La media della corrente durante la misura vale

$$\langle \bar{I} \rangle = \left\langle \frac{1}{t_m} \int_0^{t_m} I(t) dt \right\rangle = \left\langle \frac{1}{t_m} \sum_{i=1}^n \int_0^{t_m} F(t - t_i) dt \right\rangle = \frac{\langle n \rangle e}{t_m}, \quad (7.127)$$

che è quello che ci aspettavamo. Per la funzione di autocorrelazione della corrente durante la misura abbiamo

$$C_H(\tau) = \overline{\langle I(t)I(t+\tau) \rangle} = \overline{\left\langle \sum_{i=1}^n F(t - t_i) \sum_{j=1}^n F(t + \tau - t_j) \right\rangle}. \quad (7.128)$$

Se scegliamo un valore qualunque di τ e fissiamo l'indice i della prima sommatoria, per ogni $j \neq i$ nella seconda sommatoria il prodotto $F(t - t_i)F(t + \tau - t_j)$ sarà diverso da zero solo se siamo nell'intervallo di tempo $t_i - \tau_c < t < t_i + \tau_c$ (passaggio dello i -esimo elettrone), e se, contemporaneamente, $t_i + \tau - \tau_c < t_j < t_i + \tau + \tau_c$, in modo che nello stesso intervallo di tempo passi anche il j -esimo elettrone. Ma, con le condizioni ipotizzate all'inizio, ogni elettrone è correlato solo con sé stesso, quindi la probabilità che $F(t - t_i)F(t + \tau - t_j) \neq 0$ per $j \neq i$ è molto piccola (nulla se τ_c e la corrente I sono sufficientemente piccoli perché sia praticamente nulla la probabilità del passaggio simultaneo di due elettroni in un intervallo di tempo di durata $2\tau_c$) e del tutto indipendente da τ . Quindi la media sugli addendi con $j \neq i$ fornisce solo un contributo costante C_0 , molto piccolo e, al limite, nullo, alla funzione di autocorrelazione. Aggiungendo anche il contributo degli addendi con $j = i$ otteniamo

$$C_H(\tau) = \overline{\left\langle \sum_{i=1}^n F(t - t_i)F(t + \tau - t_i) \right\rangle} + C_0 = \langle n \rangle \overline{F(t)F(t + \tau)} + C_0. \quad (7.129)$$

Il tempo di correlazione tipico è molto breve, dell'ordine di $\tau_c \simeq 10^{-12}$ s, e la $G(\nu)$, data dalla trasformata di Fourier di $C_H(\tau) - C_0$ ha quindi una larghezza $\Delta\nu_c \simeq 1/\tau_c \simeq 10^{12}$ Hz. Questo corrisponde ad un *rumore bianco*, cioè ad una distribuzione di frequenze praticamente piatta in tutta la zona di interesse.

7.7 Rumore 1/f

Nelle valvole termoelettroniche è stato scoperto un particolare tipo di rumore, detto rumore 1/f perché la sua distribuzione spettrale dipende dalla frequenza come $1/\nu^\gamma$, con γ molto vicino a 1, secondo la formula

$$G(\nu) = \alpha \frac{\langle I \rangle^2}{\nu^\gamma}, \quad (7.130)$$

dove il coefficiente di proporzionalità α è fortemente dipendente dalle caratteristiche del particolare componente elettronico. Questo rumore viene anche detto *rumore rosa*, per essere intermedio tra il rumore bianco, che, nella regione in cui è praticamente costante va, ovviamente, come $1/\nu^0$, e il *rumore rosso*, o rumore browniano, che va come $1/\nu^2$. Il rumore 1/f, oltre che in elementi elettronici, appare anche in una varietà di fenomeni naturali, come la velocità della corrente negli oceani, lo scorrere della sabbia in una clessidra, nel flusso annuale del Nilo misurato negli ultimi 2000 anni, e in altri esempi.

Non esiste una spiegazione convincente per tutto il rumore 1/f che appare in queste situazioni fisiche molto diverse tra loro.

7.8 Fotoconteggi

Questo è lo schema di un esperimento di conteggio di singoli fotoni: un fascio di luce incide su di un fotomoltiplicatore collegato, attraverso un'opportuna elettronica, ad un contatore che registra il numero degli impulsi generati dai singoli fotoni che sono stati rivelati. Un otturatore davanti al fotomoltiplicatore controlla l'intervallo di tempo t_m durante il quale il fascio di luce può raggiungere il fotocatodo. La misura viene ripetuta più volte, registrando ogni volta il numero n degli impulsi rivelati durante il

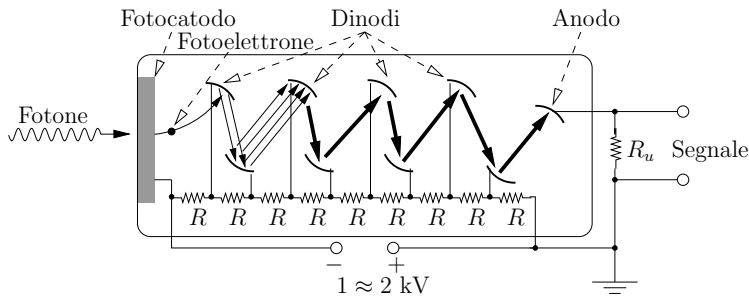


Figura 7.1 Fotomoltiplicatore. Quando un fotone incide sul *fotocatodo*, questo emette un *elettrone primario* (fotoelettrone) per effetto fotoelettrico. L'elettrone primario viene accelerato da una differenza di potenziale e, dopo aver acquisito una certa energia cinetica, colpisce il primo *dinodo*. L'urto provoca l'emissione di più *elettroni secondari*, ognuno di energia minore di quella dell'elettrone incidente, da parte del dinodo. Segue una successione di dinodi, ognuno ad un potenziale superiore a quello del dinodo precedente. Ogni elettrone secondario emesso da un dinodo viene accelerato verso il dinodo successivo, dove provoca l'emissione di più elettroni secondari, a loro volta accelerati verso un altro dinodo. Gli elettroni emessi dall'ultimo dinodo vengono raccolti dall'anodo, e si misura un impulso di corrente.

ve perché la probabilità di emissione di più di un elettrone sia trascurabile. La probabilità $W(t_m, n)$ si ottiene facendo la media dei conteggi ottenuti in un gran numero di misure, inizianti ognuna ad un diverso istante t , e ognuna di durata t_m . Consideriamo un determinato intervallo di conteggio $(t, t + t_m)$, chiamiamo $t + t'$ un istante all'interno dell'intervallo, e dt' un intervallino entro al quale possiamo considerare valida la (7.131): cioè, all'interno di dt' può essere contato al massimo un fotone. Definiamo

$$W_n(t, t' + dt') \quad (7.132)$$

la probabilità che n fotoni siano contati nell'intervallo $(t, t + t' + dt')$. Questi n conteggi possono essere realizzati in due modi diversi:

1. $n - 1$ conteggi nell'intervallo $(t, t + t')$ e 1 conteggio nell'intervallo $(t + t', t + t' + dt')$, con probabilità $W_{n-1}(t, t') w(t + t') dt'$;
2. n conteggi nell'intervallo $(t, t + t')$ e 0 conteggi nell'intervallo $(t + t', t + t' + dt')$, con probabilità $W_n(t, t')[1 - w(t + t') dt']$.

tempo t_m . Si risale così alla probabilità $W(t_m, n)$ di contare n eventi durante l'intervallo di tempo t_m . Se il fascio di radiazione è stazionario, la distribuzione statistica $W(t_m, n)$ contiene informazioni sulle proprietà statistiche della sorgente.

Se chiamiamo $w(t) dt$ la probabilità che il fotocatodo del nostro fotomoltiplicatore emetta un elettrone nell'intervallo di tempo $(t, t + dt)$, in prima approssimazione possiamo scrivere

$$w(t) dt = \alpha I(t) dt, \quad (7.131)$$

dove $I(t)$ è l'intensità del fascio all'istante t , la costante α rappresenta l'efficienza del tubo, e l'intervallo di tempo dt è supposto abbastanza lungo perché la teoria della probabilità di transizione sia valida, ma sufficientemente breve

La probabilità totale sarà la somma delle due probabilità

$$\begin{aligned} W_n(t, t' + dt') &= W_n(t, t')[1 - w(t + t') dt'] + W_{n-1}(t, t') w(t + t') dt' \\ &= W_n(t, t') + w(t + t') dt' [W_{n-1}(t, t') - W_n(t, t')] \end{aligned} \quad (7.133)$$

da cui otteniamo

$$W_n(t, t' + dt') - W_n(t, t') = w(t + t') dt' [W_{n-1}(t, t') - W_n(t, t')], \quad (7.134)$$

e, dividendo per dt' , abbiamo finalmente

$$\frac{dW_n(t, t')}{dt'} = w(t + t') [W_{n-1}(t, t') - W_n(t, t')] = \alpha I(t + t') [W_{n-1}(t, t') - W_n(t, t')]. \quad (7.135)$$

Quindi la probabilità di avere n conteggi nell'intervallo $(t, t + t')$ è legata alla probabilità di avere $n - 1$ conteggi nello stesso intervallo da un'equazione differenziale, e si può costruire una catena di equazioni differenziali che scende fino alla probabilità di avere zero conteggi. In questa catena di equazioni, quella per $n = 0$ è diversa dalle altre perché, ovviamente, la probabilità di avere un numero negativo di conteggi (in particolare $n = -1$) è nulla, quindi

$$\frac{dW_0(t, t')}{dt'} = -\alpha I(t + t') W_0(t, t'), \quad (7.136)$$

che può essere riscritta

$$\frac{dW_0(t, t')}{W_0(t, t')} = -\alpha I(t + t') dt', \quad \text{da cui, integrando,} \quad \ln \left[\frac{W_0(t, t_m)}{W_0(t, 0)} \right] = -\alpha \int_t^{t+t_m} I(t') dt'. \quad (7.137)$$

Imponendo la condizione iniziale che in un intervallo nullo si abbia la certezza di non avere conteggi

$$W_0(t, 0) = 1 \quad (7.138)$$

otteniamo

$$W_0(t, t_m) = \exp \left[-\alpha \int_t^{t+t_m} I(t') dt' \right]. \quad (7.139)$$

Se introduciamo l'intensità media della radiazione durante il tempo di conteggio

$$\bar{I}(t, t_m) = \frac{1}{t_m} \int_t^{t+t_m} I(t') dt' \quad (7.140)$$

possiamo scrivere

$$W_0(t, t_m) = e^{-\alpha t_m \bar{I}(t, t_m)}. \quad (7.141)$$

I successivi $W_n(t, t_m)$ possono essere determinati dalla catena di equazioni (7.135), cominciando con $n = 1$ e ogni volta imponendo la condizione che siano nulle le probabilità

$$W_n(t, 0) = 0 \quad \forall n \neq 0 \quad (7.142)$$

di avere più di zero conteggi in un tempo nullo. La soluzione generale risulta

$$W_n(t, t_m) = \frac{[\alpha \bar{I}(t, t_m) t_m]^n}{n!} e^{-\alpha \bar{I}(t, t_m) t_m}, \quad (7.143)$$

come si dimostra per induzione. Infatti la (7.143) verifica la relazione di ricorrenza (7.135):

$$\begin{aligned} \frac{dW_n(t, t')}{dt'} &= \frac{d}{dt'} \frac{[\alpha \bar{I}(t, t') t']^n}{n!} e^{-\alpha \bar{I}(t, t') t'} \\ &= \frac{n [\alpha \bar{I}(t, t') t']^{n-1}}{n!} \alpha \bar{I}(t, t') e^{-\alpha \bar{I}(t, t') t'} - \frac{[\alpha \bar{I}(t, t') t']^n}{n!} \alpha \bar{I}(t, t') e^{-\alpha \bar{I}(t, t') t'} \\ &= \alpha \bar{I}(t, t') [W_{n-1}(t, t') - W_n(t, t')], \end{aligned} \quad (7.144)$$

e, per $n = 0$, coincide con la (7.141).

La probabilità $W_n(t_m)$ si ottiene mediando su di un grande numero di misure

$$W_n(t_m) = \langle W_n(t, t_m) \rangle = \left\langle \frac{[\alpha \bar{I}(t, t_m) t_m]^n}{n!} e^{-\alpha \bar{I}(t, t_m) t_m} \right\rangle. \quad (7.145)$$

Nel caso che $\bar{I}(t, t_m)$ sia indipendente da t e da t_m possiamo porre

$$\bar{I} = \bar{I}(t, t_m) \quad (7.146)$$

e scrivere

$$W_n(t_m) = \frac{\langle n \rangle^n}{n!} e^{-\langle n \rangle}, \quad (7.147)$$

dove $\langle n \rangle$ è il numero medio dei conteggi nel tempo di misura t_m

$$\langle n \rangle = \alpha \bar{I} t_m. \quad (7.148)$$

In queste condizioni abbiamo così una distribuzione di Poisson. La distribuzione poissoniana, con deviazione quadratica media $\langle (\Delta n)^2 \rangle = \langle n \rangle$, può sembrare sorprendente se si pensa che questo risultato potrebbe essere stato ottenuto anche supponendo di avere una intensa radiazione incidente monocromatica, quindi sotto le condizioni richieste per l'applicazione della statistica di Bose-Einstein.

7.9 Teorema di Nyquist

Nel 1927 una serie di esperimenti condotti da John B. Johnson ha dimostrato che esiste un rumore elettrico su tutti i conduttori (rumore Johnson), e che questo rumore, dipendente dalla temperatura, è inevitabile e non dovuto a errori nella progettazione. La formula trovata sperimentalmente da Johnson per la potenza associata a questo rumore termico è semplice e non contiene alcuna dipendenza né dalle dimensioni né dalla resistenza del conduttore:

$$W = kT \Delta \nu,$$

dove $\Delta \nu$ è l'intervallo di frequenza entro il quale viene effettuata la misura. Nel 1928 Harry Nyquist ha dato la spiegazione teorica del fenomeno, che schematizziamo qui sotto.

Consideriamo un conduttore di lunghezza l e sezione A , come in Fig. 7.2. Supponiamo che il conduttore abbia n elettroni di conduzione per unità di volume, e chiamiamo R la sua resistenza. Per la differenza di potenziale V ai capi del conduttore abbiamo dalla legge di Ohm

$$V = IR = RAJ = RAen\bar{v}, \quad (7.149)$$

dove I è la corrente elettrica, J la densità di corrente, e il modulo della carica dell'elettrone e $\bar{v}$ la velocità di deriva degli elettroni, diretta lungo x . Il numero totale di elettroni di conduzione è nAl , quindi la velocità di deriva vale

$$\bar{v} = \frac{1}{nAl} \sum_i v_i, \quad (7.150)$$

dove l'indice i corre su tutti gli elettroni di conduzione e v_i è la componente x della velocità dell' i -esimo elettrone. La (7.149) può essere così riscritta

$$V = R \frac{e}{l} \sum_i v_i = \sum_i V_i, \quad \text{dove} \quad V_i = \frac{Re v_i}{l} \quad (7.151)$$

è il contributo dello i -esimo elettrone alla differenza di potenziale misurata tra i capi del conduttore. Poiché v_i è una variabile aleatoria stazionaria, anche V_i è una variabile aleatoria stazionaria. Introducendo la sua densità spettrale $G(\nu)$ data dalla (7.31) abbiamo per la (7.32)

$$\overline{V_i^2} = \int_{-\infty}^{+\infty} G(\nu) d\nu. \quad (7.152)$$

Se, con un filtro, escludiamo le componenti di frequenza al di fuori dell'intervallo $(\nu, \nu + \Delta\nu)$, il contributo medio del singolo elettrone all'intervallo di frequenza di osservazione vale

$$\overline{V_i^2} = G(\nu)\Delta\nu. \quad (7.153)$$

Supponiamo che la funzione di autocorrelazione di V_i sia del tipo

$$C(\tau) = \overline{V_i(t) V_i(t + \tau)} = \overline{V_i^2} e^{-\tau/\tau_c}, \quad (7.154)$$

dove τ_c è il tempo di rilassamento, o tempo di volo medio degli elettroni di conduzione. Dal teorema di Wiener-Khinchin (7.47) abbiamo

$$\begin{aligned} G(\nu) &= 4 \left(\frac{Re}{l} \right)^2 \overline{v^2} \int_0^\infty e^{-\tau/\tau_c} \cos 2\pi\nu\tau d\tau \\ &= 4 \left(\frac{Re}{l} \right)^2 \overline{v^2} \frac{\tau_c}{1 + (2\pi\nu\tau_c)^2}. \end{aligned} \quad (7.155)$$

Normalmente nei metalli a temperatura ambiente $\tau_c < 10^{-13}$ s, quindi dalla corrente continua fino alle microonde abbiamo che $2\pi\nu\tau_c \ll 1$, e questo termine può essere trascurato. Ricordando che

$$\frac{1}{2} m \overline{v^2} = \frac{1}{2} kT, \quad \text{da cui} \quad \overline{v^2} = \frac{kT}{m}, \quad (7.156)$$

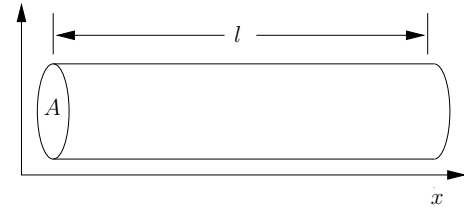


Figura 7.2 Resistenza a temperatura T .

per l'intervallo di frequenza $\Delta\nu$ possiamo scrivere

$$\overline{V^2} = N A l \overline{V_i^2} = N A l G(\nu) \Delta\nu = N A l 4 \left(\frac{kT}{m} \right) \left(\frac{Re}{l} \right)^2 \tau_c \Delta\nu, \quad (7.157)$$

ovvero

$$\overline{V^2} = 4kTR\Delta\nu, \quad \text{corrispondente ad una potenza media} \quad W = \frac{\overline{V^2}}{R} = 4kT\Delta\nu, \quad (7.158)$$

che è il teorema di Nyquist. Per arrivare a questo risultato abbiamo usato la relazione

$$R = \frac{l}{\sigma A} \quad (7.159)$$

e la relazione

$$\sigma = \frac{ne^2\tau_c}{m}, \quad (7.160)$$

che può essere ottenuta, per esempio, dal modello di Drude per la conduzione nei metalli.

Capitolo 8

Fenomeni di interferenza

8.1 Formalismo dei campi complessi

La trattazione matematica dei fenomeni di interferenza tra onde elettromagnetiche monocromatiche (o quasi monocromatiche) risulta più semplice se effettuata in termini dei *campi complessi* associati alle onde stesse. Dato il campo elettrico $\mathbf{E}(\mathbf{r}, t)$ di un'onda elettromagnetica osservato nel punto $\mathbf{r}$ dello spazio, il corrispondente campo complesso $\mathbf{A}(\mathbf{r}, t)$ è definito dalle relazioni

$$A_j(\mathbf{r}, t) = 2 \int_0^\infty \tilde{E}_j(\mathbf{r}, \nu) e^{+2\pi i \nu t} d\nu, \quad (8.1)$$

dove l'indice j si riferisce a una particolare componente del campo (per esempio $j = x, y, z$), e $\tilde{E}_j(\mathbf{r}, \nu)$ è la trasformata di Fourier della componente del campo elettrico dell'onda

$$\tilde{E}_j(\mathbf{r}, \nu) = \int_{-\infty}^{+\infty} E_j(\mathbf{r}, t) e^{-2\pi i \nu t} dt = \tilde{E}_j^*(\mathbf{r}, -\nu). \quad (8.2)$$

$A_j(\mathbf{r}, t)$ contiene così solo le frequenze positive di $E_j(\mathbf{r}, t)$, e si ha

$$E_j(\mathbf{r}, t) = \text{Re} [A_j(\mathbf{r}, t)], \quad (8.3)$$

La trasformata di Fourier di $A_j(\mathbf{r}, t)$ vale

$$\tilde{A}_j(\mathbf{r}, \nu) = 2\tilde{E}_j(\mathbf{r}, \nu) \quad \text{se } \nu \geq 0, \quad \text{e} \quad \tilde{A}_j(\mathbf{r}, \nu) = 0 \quad \text{se } \nu < 0. \quad (8.4)$$

D'ora in poi considereremo le singole componenti del campo separatamente, tralasciando, quando chiaro dal contesto, l'indice j . L'interferenza tra le onde ci interesserà in un punto determinato dello spazio, e, quando chiaro dal contesto, ometteremo l'argomento $\mathbf{r}$ nella nostra notazione per i campi. Come primo esempio consideriamo un'onda piana monocromatica polarizzata linearmente lungo l'asse x , che si propaghi lungo l'asse z . In un punto generico dello spazio abbiamo per il campo $E_x(t) = E(t) = E_0 \cos(2\pi\nu_0 t + \varphi)$. Qui la fase φ comprende anche la dipendenza dalla posizione, cioè $\varphi = \varphi_0 - 2\pi z/\lambda$. Ricordando la relazione

$$\int_{-\infty}^{+\infty} e^{-2\pi i \nu t} dt = \delta(\nu), \quad (8.5)$$

dove $\delta(\nu)$ è la delta di Dirac, la trasformata di Fourier della nostra onda vale

$$\begin{aligned}
 \tilde{E}(\nu) &= \int_{-\infty}^{+\infty} E_0 \cos(2\pi\nu_0 t + \varphi) e^{-2\pi i \nu t} dt \\
 &= \frac{E_0}{2} \left[\int_{-\infty}^{+\infty} e^{-2\pi i (\nu - \nu_0)t + i\varphi} dt + \int_{-\infty}^{+\infty} e^{-2\pi i (\nu + \nu_0)t - i\varphi} dt \right] \\
 &= \frac{E_0}{2} \left[e^{i\varphi} \delta(\nu - \nu_0) + e^{-i\varphi} \delta(\nu + \nu_0) \right], \tag{8.6}
 \end{aligned}$$

Inserendo la (8.6) nell'espressione (8.1) abbiamo per il campo complesso della nostra onda monocromatica

$$A(t) = E_0 \int_0^\infty \left[e^{i\varphi} \delta(\nu - \nu_0) - e^{-i\varphi} \delta(\nu + \nu_0) \right] e^{2\pi i \nu t} d\nu = E_0 e^{+2\pi i \nu_0 t + i\varphi} = A_0 e^{+2\pi i \nu_0 t + i\varphi}. \tag{8.7}$$

con $A_0 = E_0$. L'intensità istantanea di un'onda elettromagnetica classica, misurata in W/m^2 , è data dal vettore di Poynting $\mathbf{P}(t) = \frac{1}{\mu_0} \mathbf{E}(t) \times \mathbf{B}(t) = \varepsilon_0 c^2 \mathbf{E}(t) \times \mathbf{B}(t)$. Ricordando che, per esempio, per un'onda piana polarizzata lungo l'asse x , che si propaga nel vuoto lungo l'asse z , vale la relazione $B_y(t) = \frac{1}{c} E_x(t)$, possiamo scrivere il vettore di Poynting nella forma

$$\mathbf{P}(t) = \varepsilon_0 c E^2(t) \hat{\mathbf{z}}. \tag{8.8}$$

La 8.8 ci permette di scrivere l'intensità istantanea di un'onda monocromatica in funzione del suo campo complesso $A(t)$ come

$$I(t) = \overline{P(t)} = \frac{1}{2} \varepsilon_0 c \overline{A^*(t) A(t)} = \frac{1}{2} \varepsilon_0 c A^*(t) A(t) = \frac{1}{2} \varepsilon_0 c A_0^2, \tag{8.9}$$

dove la barra indica la media temporale su un periodo dell'onda, e nel penultimo passaggio si è tenuto conto del fatto che la dipendenza temporale scompare nel prodotto $A^*(t)A(t)$. Il fattore $1/2$ compare appunto perché, come vediamo dalla (8.7), $A^*(t)A(t) = \overline{A^*(t)A(t)} = E_0^2$, mentre da $E^2(t) = E_0^2 \cos^2(2\pi\nu_0 t + \varphi)$ segue $\overline{E^2(t)} = \frac{1}{2} E_0^2$. Per un'onda non monocromatica la media temporale andrà effettuata su un periodo lungo rispetto all'inverso della frequenza più bassa presente nell'onda. Se in un punto si sovrappongono due onde di ampiezza complessa rispettivamente $A_1(t)$ ed $A_2(t)$, il campo complessivo vale $A_{1+2}(t) = A_1(t) + A_2(t)$, mentre per l'intensità complessiva I_{1+2} abbiamo

$$I_{1+2} = \frac{1}{2} \varepsilon_0 c \overline{[A_1^*(t) + A_2^*(t)][A_1(t) + A_2(t)]} = I_1 + I_2 + \varepsilon_0 c \operatorname{Re} [A_1^*(t)A_2(t)], \tag{8.10}$$

dove la barra indica la media temporale sulla durata della misura, in generale molto lunga rispetto al periodo di oscillazione delle onde.

8.2 Beam-splitter

Un *beam-splitter* normalmente è costituito da un cubo formato da due prismi di vetro a base triangolare incollati mediante una resina (per esempio, balsamo del Canada) con indice di rifrazione diverso da quello del vetro, come in Fig. 8.1. Lo spessore dello strato di resina è tale che, per una certa lunghezza d'onda, metà della luce incidente viene trasmessa, e l'altra metà riflessa.

Consideriamo un generico beam-splitter su cui, sempre come in Fig. 8.1, incidano due fasci luminosi di ampiezza complessa rispettivamente $A_1(t)$ ed $A_2(t)$, ed escano i due fasci di ampiezza complessa $A_3(t)$ ed $A_4(t)$. I campi in uscita sono legati ai campi in ingresso dalla relazione lineare

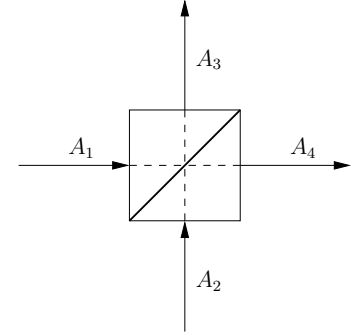


Figura 8.1 Beam-splitter ideale.

$$A_3 = R_{31}A_1 + T_{32}A_2 \quad \text{e} \quad A_4 = R_{42}A_2 + T_{41}A_1, \quad (8.11)$$

dove gli R_{ij} e T_{ij} sono, rispettivamente, i coefficienti di *riflessione* e *trasmissione* del beam splitter. In generale questi coefficienti saranno numeri complessi e dipenderanno dalla frequenza di radiazione. Le relazioni (8.11) possono esser scritte in forma matriciale

$$\begin{pmatrix} A_3 \\ A_4 \end{pmatrix} = \begin{pmatrix} R_{31} & T_{32} \\ T_{41} & R_{42} \end{pmatrix} \begin{pmatrix} A_1 \\ A_2 \end{pmatrix}, \quad (8.12)$$

dove la matrice 2×2 è nota come *matrice del beam-splitter*. Considerazioni di conservazione dell'energia tra ingresso e uscita portano ad importanti proprietà della matrice del beam-splitter. Se il beam-splitter non ha perdite dobbiamo avere

$$|A_1|^2 + |A_2|^2 = |A_3|^2 + |A_4|^2 \quad (8.13)$$

qualunque siano i valori di A_1 ed A_2 , quindi

$$|R_{31}|^2 + |T_{41}|^2 = |R_{42}|^2 + |T_{32}|^2 = 1, \quad \text{e} \quad R_{31}T_{32}^* + T_{41}R_{42}^* = 0. \quad (8.14)$$

I coefficienti di riflessione e trasmissione possono essere separati in ampiezza e fase secondo lo schema $R_{31} = |R_{31}|e^{i\varphi_{31}}$ e corrispondenti. Dalla seconda delle (8.14) segue che deve essere

$$\varphi_{31} + \varphi_{42} - \varphi_{32} - \varphi_{41} = \pm\pi \quad \text{e} \quad |R_{31}|/|T_{41}| = |R_{42}|/|T_{32}|. \quad (8.15)$$

Queste relazioni insieme alla prima delle (8.14) mostrano che i due coefficienti di riflessione devono avere uguale ampiezza, come pure i due coefficienti di trasmissione:

$$|R_{31}| = |R_{42}| \equiv |R| \quad \text{e} \quad |T_{41}| = |T_{32}| \equiv |T|. \quad (8.16)$$

Le equazioni (8.14) insieme alla (8.16) ci assicurano che la matrice del beam-splitter è unitaria, cioè che la sua inversa è uguale alla trasposta della sua complessa coniugata. Infatti

$$\begin{aligned}
 & \begin{pmatrix} R_{31} & T_{32} \\ T_{41} & R_{42} \end{pmatrix} \begin{pmatrix} R_{31}^* & T_{41}^* \\ T_{32}^* & R_{42}^* \end{pmatrix} = \begin{pmatrix} |R_{31}|^2 + |T_{32}|^2 & R_{31}T_{41}^* + T_{32}R_{42}^* \\ T_{41}R_{31}^* + R_{42}T_{32}^* & |T_{41}|^2 + |R_{42}|^2 \end{pmatrix} = \\
 & = \begin{pmatrix} |R|^2 + |T|^2 & RT[e^{i(\varphi_{31}-\varphi_{41})} + e^{i(\varphi_{32}-\varphi_{42})}] \\ RT[e^{i(\varphi_{41}-\varphi_{31})} + e^{i(\varphi_{42}-\varphi_{32})}] & |T|^2 + |R|^2 \end{pmatrix} = \\
 & = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad (8.17)
 \end{aligned}$$

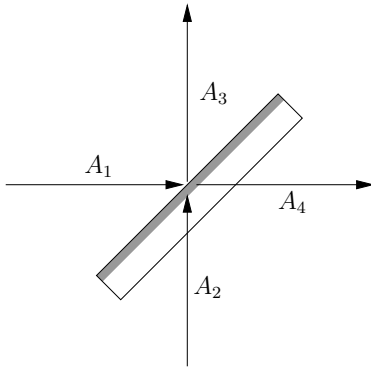


Figura 8.2 Specchio semiriflettente.

La struttura della matrice del beam-splitter può essere semplificata con ipotesi sui coefficienti di trasmissione e riflessione. Per esempio, i coefficienti possono essere presi tutti reali, con $\varphi_{31} = \varphi_{32} = \varphi_{41} = 0$, e $\varphi_{42} = \pi$, che implica

$$R_{31} = -R_{42} = |R| \quad \text{e} \quad T_{32} = T_{41} = |T|. \quad (8.18)$$

Un'alternativa, che useremo nel seguito, è prendere i coefficienti simmetrici, con $\varphi_{31} = \varphi_{42} = \varphi_R$ e $\varphi_{32} = \varphi_{41} = \varphi_T$, se

$$R_{31} = R_{42} = R = |R|e^{i\varphi_R} \quad \text{e} \quad T_{32} = T_{41} = T = |T|e^{i\varphi_T}. \quad (8.19)$$

In questo caso le relazioni (8.14) si riducono a

$$\begin{aligned}
 & |R|^2 + |T|^2 = 1 \quad \text{e} \quad RT^* + R^*T = 0, \\
 & \text{equivalente a} \quad \varphi_R - \varphi_T = \pm \frac{\pi}{2}. \quad (8.20)
 \end{aligned}$$

Un'ulteriore semplificazione si ha nel caso di un beam-splitter in cui di ogni fascio entrante il 50% venga riflesso ed il 50% venga trasmesso. In questo caso abbiamo

$$|R| = |T| = \frac{1}{\sqrt{2}}, \quad \text{con} \quad \varphi_R - \varphi_T = \frac{\pi}{2}. \quad (8.21)$$

Spesso il beam-splitter è sostituito da uno specchio semiargentato, che riflette metà della radiazione incidente e trasmette l'altra metà (Fig. 8.2).

8.3 Interferometro di Mach-Zehnder

L'interferometro di Mach-Zehnder, sviluppato da Ludwig Mach e Ludwig Zehnder nel 1891, è utile per illustrare alcuni principi generali dell'interferometria. Lo schema di principio è mostrato in Fig. 8.3. Il fascio di luce da analizzare, di ampiezza complessa A , incide sul primo beam-splitter BS_1 . Il fascio riflesso ed il fascio trasmesso, di ampiezza complesse rispettivamente A_1 ed A_2 , dopo riflessioni su specchi che generano due diversi cammini ottici, rispettivamente l_1 e l_2 , incidono agli ingressi 1 e 2 di del secondo beam-splitter BS_2 . La sovrapposizione della parte trasmessa di A_1 e della parte riflessa di A_2 , A_4 nella figura, viene osservata da un rivelatore R . Naturalmente, a rigor di termini, la configurazione di Fig. 8.3, con i cammini l_1 e l_2 che, con le riflessioni di 90° agli specchi formano

esattamente i lati un rettangolo, darebbe una differenza di cammino ottico $l_1 - l_2$ nulla, non proprio utile ai nostri scopi. Ma è possibile immaginare configurazioni diverse, in cui la differenza tra i due cammini ottici sia regolabile. Per esempio, nello schema di Fig. 8.4 è possibile variare il ritardo relativo tra i cammini l_1 e l_2 traslando il blocco degli specchi M_3 ed M_4 . Con la notazione del paragrafo 8.2 abbiamo per l'intensità complessa $A_4(t)$ del fascio che incide sul rivelatore

$$\begin{aligned} A_4(t) &= TR A(t_1) + RT A(t_2), \\ \text{con } t_1 &= t - \frac{l_1}{c} \\ \text{e } t_2 &= t - \frac{l_2}{c}. \end{aligned} \quad (8.22)$$

Qualunque sia il rivelatore usato, il tempo necessario per una misura sarà lungo rispetto al periodo della frequenza minima presente nella radiazione elettromagnetica. L'intensità "istantanea" misurata dal rivelatore sarà quindi proporzionale alla media dell'intensità sul tempo di misura. Chiamando $2T$ la durata della misura, che supponiamo duri da $t - T$ a $t + T$, il segnale "istantaneo" $S(t)$ può essere scritto

$$S(t) = \frac{1}{2T} \int_{t-T}^{t+T} I_4(t') dt' = \overline{I_4(t)}, \quad (8.23)$$

dove $\overline{I_4(t)}$ è il valor medio dell'intensità $I_4(t)$ durante la misura. Sempre indicando con una barra l'operazione di media temporale sulla misura abbiamo

$$\begin{aligned} S(t) = \overline{I_4(t)} &= \frac{1}{2} \varepsilon_0 c \overline{|A_4(t)|^2} = \frac{1}{2} \varepsilon_0 c |R|^2 |T|^2 \left\{ \overline{|A(t_1)|^2} + \overline{|A(t_2)|^2} + 2 \operatorname{Re} \left[\overline{A^*(t_1) A(t_2)} \right] \right\} \\ &= \frac{1}{2} \varepsilon_0 c \frac{1}{4} \left\{ \overline{|A(t_1)|^2} + \overline{|A(t_2)|^2} + 2 \operatorname{Re} \left[\overline{A^*(t_1) A(t_2)} \right] \right\}, \end{aligned} \quad (8.24)$$

dove nell'ultimo passaggio abbiamo fatto l'ipotesi di avere due beam-splitter ideali (senza perdite) con $|R| = |T| = 1/\sqrt{2}$. Il tipo di segnale osservato dipende dalla natura della luce incidente sull'interferometro, descritta dall'andamento temporale dell'ampiezza complessa $A(t)$. Se la luce è stazionaria, cioè, se le probabilità delle fluttuazioni non dipendono dal tempo, e se il tempo di misura $2T$ è lungo anche rispetto alla scala temporale delle fluttuazioni, i valori delle medie calcolate nella (8.24) non dipendono né dal valore di T , né dall'istante in cui la misura comincia. Delle tre medie temporali tra parentesi a graffia, le prime due corrispondono alle intensità che sarebbero prodotte da ognuno dei due

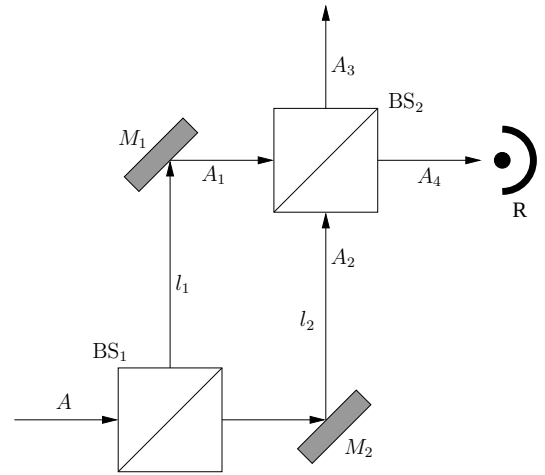


Figura 8.3 Schema di principio dell'interferometro di Mach-Zehnder.

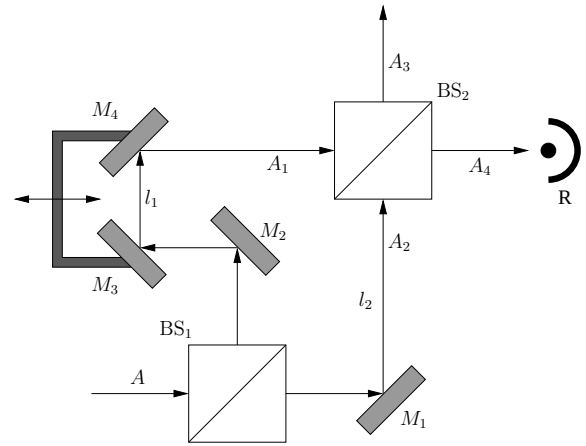


Figura 8.4 Variante dell'interferometro di Mach-Zehnder. Traslando il blocco degli specchi M_3 e M_4 si varia il ritardo relativo tra i percorsi l_1 e l_2 .

canali interni dell'interferometro (onda trasmessa e onda riflessa dal primo beam-splitter) in assenza dell'altro, e non generano frange di interferenza. Abbiamo infatti

$$\frac{1}{4} \varepsilon_0 c \overline{|A(t_1)|^2} = \frac{1}{4} \varepsilon_0 c \overline{|A(t_2)|^2} = \frac{1}{4} \varepsilon_0 c \frac{1}{2T} \int_{-T}^T |A(t')|^2 dt' = \frac{1}{2} S_i, \quad (8.25)$$

dove S_i è il segnale che leggeremmo se la radiazione in ingresso incidesse direttamente sul rivelatore, senza passare attraverso l'interferometro. Le frange sono originate dal terzo termine tra parentesi a graffa della (8.24). Ricordando le ultime delle (8.22), possiamo scrivere

$$\begin{aligned} \overline{A^*(t_1)A(t_2)} &= \frac{1}{2T} \int_{-T}^T A^* \left(t - \frac{l_1}{c} \right) A \left(t - \frac{l_2}{c} \right) dt \\ &\simeq \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^T A^*(t') A(t' + \tau) dt', \end{aligned} \quad (8.26)$$

dove abbiamo posto $t' = t - l_1/c$ e $\tau = (l_1 - l_2)/c$, e l'approssimazione è valida sotto l'ipotesi che T sia lungo rispetto alla scala temporale delle fluttuazioni. Questa è la *funzione di autocorrelazione temporale al primo ordine*, mediata sul tempo, del campo complesso $A(t)$, che abbiamo incontrato nella (7.41) del paragrafo 7.3, con la differenza che qui non compare ancora una media sull'ensemble. Un'importante proprietà di questa funzione in caso di statistica stazionaria è che essa dipende solo dal *tempo di ritardo* τ . Ne segue che

$$\overline{A^*(t)A(t + \tau)} = \overline{A^*(t - \tau)A(t)} = \left(\overline{A^*(t)A(t - \tau)} \right)^*, \quad (8.27)$$

e la parte reale della funzione di autocorrelazione ha lo stesso valore per τ positivo o negativo. Tenendo presente che la media temporale di $A(t)$ è nulla, sempre in caso di statistica stazionaria la parte reale della funzione di autocorrelazione avrà il suo valore massimo per $\tau = 0$. Possiamo normalizzare questa funzione di autocorrelazione ottenendo la *coerenza temporale al primo ordine*

$$g^{(1)}(\tau) = \frac{\overline{A^*(t)A(t + \tau)}}{\overline{A^*(t)A(t)}}, \quad \text{con} \quad g^{(1)}(-\tau) = g^{*(1)}(\tau) \quad \text{per la (8.27)}. \quad (8.28)$$

Abbiamo così $g^{(1)}(0) = 1$ e $\lim_{\tau \rightarrow \infty} g^{(1)}(\tau) = 0$. Notando che per la (8.25)

$$\frac{1}{4} \varepsilon_0 c \overline{A^*(t)A(t)} = \frac{1}{4} \varepsilon_0 c \overline{|A(t_1)|^2} = \frac{1}{4} \varepsilon_0 c \overline{|A(t_2)|^2} = \frac{1}{2} S_i, \quad (8.29)$$

possiamo scrivere per il segnale in uscita dell'interferometro di Mach-Zehnder

$$S_{\text{out}}(\tau) = \frac{1}{2} S_i \left\{ 1 + \text{Re}[g^{(1)}(\tau)] \right\}, \quad (8.30)$$

Le frange osservate in uscita dall'interferometro di Mach-Zehnder dipendono così dalla coerenza temporale al primo ordine della radiazione incidente $g^{(1)}(\tau)$.

8.4 Interferometro di Michelson

L'interferometro di Michelson, rappresentato nella fig. 8.5, è essenzialmente un interferometro di Mach-Zehnder “ripiegato su sé stesso”. In ingresso abbiamo il fascio di luce da analizzare che incide sul beam-splitter. I due fasci in uscita vengono rimandati riflessi su loro stessi da due specchi, in modo che lo stesso beam-splitter che separa la luce in ingresso serva anche per ricombinare la luce che viene inviata sul rivelatore. Nella figura abbiamo numerato i cammini della luce in accordo con la notazione di fig. 8.1. Traslando uno degli specchi (M_1 nella figura) si cambia la differenza di cammino ottico Δl tra i due fasci che giungono sul rivelatore. Per interpretare classicamente il funzionamento dell'interferometro supponiamo che la radiazione provenga da una lampada spettrale atomica. Da un punto di vista classico, nella lampada molti atomi emettono luce della stessa frequenza ω_0 . Per il momento ignoriamo l'effetto Doppler. Facciamo l'ipotesi che il singolo atomo emetta un treno di radiazione con fase costante finché non subisca una collisione. Facciamo le ulteriori ipotesi che la durata della collisione sia trascurabile e che immediatamente dopo la collisione l'atomo riprenda ad emettere con la stessa frequenza ω_0 , ma con fase che ha subito un brusco cambiamento casuale. L'andamento temporale del campo

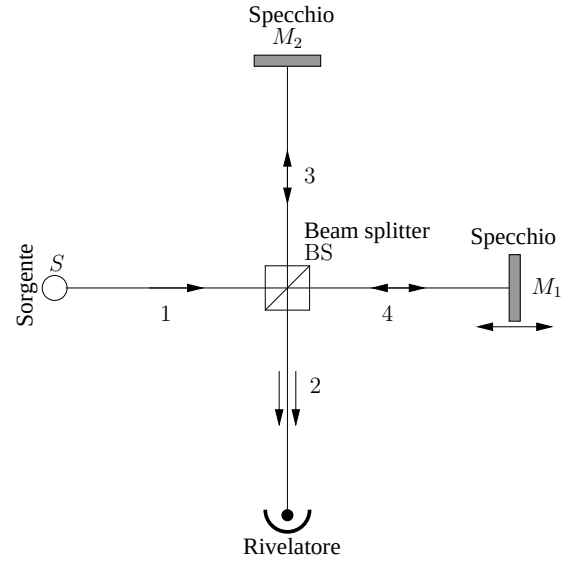


Figura 8.5 Interferometro di Michelson.

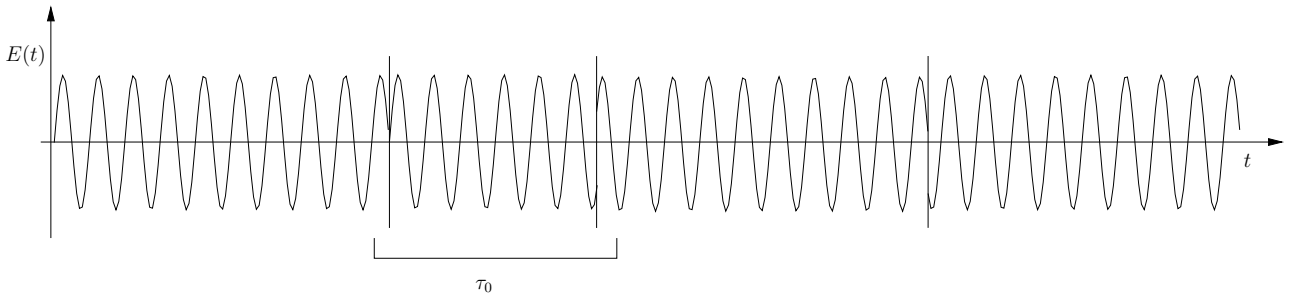


Figura 8.6 Campo elettrico del treno d'onda di frequenza ω_0 emesso da un singolo atomo. Le barre verticali corrispondono a collisioni, che cambiano la fase dell'onda. L'intervallo medio di tempo tra due collisioni è τ_0 . Il valore $\omega_0\tau_0$ usato in figura è esageratamente piccolo, per poter mostrare i cambiamenti casuali di fase.

elettrico $E(t)$ del treno d'onda emesso dal singolo atomo è rappresentato, schematicamente, in Fig. 8.6. L'ampiezza complessa dell'onda emessa dal j -esimo atomo può così essere scritta

$$A_j(t) = A_{j0} e^{-i[\omega_0 t + \varphi_j(t)]},$$

$$\text{con } E_j(t) = \text{Re}\{A_j(t)\}, \quad (8.31)$$

dove, per ogni atomo, $\varphi_j(t)$ ha un valore costante nell'intervallo di tempo tra due urti, e cambia in modo brusco e casuale ad ogni urto. Con opportuni aggiustamenti delle fasi $\varphi_j(t)$, possiamo sempre

prendere A_{j0} reale senza perdere di generalità. Per semplicità, faremo l'ipotesi che tutti i treni d'onda che giungono sull'interferometro, emessi dagli N atomi contenuti nella lampada, abbiano la stessa direzione di propagazione e la stessa intensità A_0 , ma fasi completamente scorrelate. In questo modo l'ampiezza complessa dell'onda incidente può essere scritta

$$A(t) = A_0 e^{-i\omega_0 t} \sum_{j=1}^N e^{-i\varphi_j(t)}, \quad (8.32)$$

dove l'indice j corre su tutti gli N atomi della lampada. La sua funzione di autocorrelazione temporale al primo ordine è

$$\overline{A^*(t)A(t+\tau)} = A_0^2 e^{-i\omega_0 \tau} \sum_{j,k=1}^N \overline{e^{-i[\varphi_j(t)-\varphi_k(t+\tau)]}}. \quad (8.33)$$

Dato che le fasi dei diversi atomi sono scorrelate, il contributo degli addendi con $j \neq k$ alla sommatoria a secondo membro è nullo, e l'espressione si riduce a

$$\overline{A^*(t)A(t+\tau)} = A_0^2 e^{-i\omega_0 \tau} \sum_{j=1}^N \overline{e^{-i[\varphi_j(t)-\varphi_j(t+\tau)]}} = N \overline{A_j^*(t)A_j(t+\tau)} \quad (8.34)$$

con, nell'ultimo passaggio, j qualunque, dato che tutti gli atomi sono equivalenti. Quindi la funzione di autocorrelazione globale per il fascio luminoso è uguale a N volte la funzione di autocorrelazione relativa alla radiazione emessa dal singolo atomo. La fase di ogni treno d'onda salta ad un valore casuale dopo che l'atomo ha subito una collisione, dopo di che dà un contributo nullo, in media, alla correlazione. Così la funzione di autocorrelazione di singolo atomo a secondo membro della (8.34) è proporzionale alla probabilità che l'atomo abbia un tempo di volo libero più lungo di τ . In base alla teoria cinetica dei gas, la probabilità $W(\tau) d\tau$ che un atomo abbia un tempo di volo compreso tra τ e $\tau + d\tau$ vale

$$W(\tau) d\tau = \frac{1}{\tau_f} e^{-\tau/\tau_f} d\tau, \quad \text{con} \quad \frac{1}{\tau_f} = \frac{4a^2 N}{V} \sqrt{\frac{\pi kT}{m}}, \quad (8.35)$$

dove τ_f è il tempo medio di volo libero, a la distanza tra i centri degli atomi durante una collisione, V il volume occupato dagli atomi e m la massa atomica. Avremo quindi

$$\begin{aligned} \overline{A_j^*(t)A_j(t+\tau)} &= A_0^2 e^{-i\omega_0 \tau} \overline{e^{-i[\varphi_j(t)-\varphi_j(t+\tau)]}} = A_0^2 e^{-i\omega_0 \tau} \int_{\tau}^{\infty} W(\tau') d\tau' \\ &= A_0^2 e^{-i\omega_0 \tau} \int_{\tau}^{\infty} \frac{1}{\tau_f} e^{-\tau'/\tau_f} d\tau' = A_0^2 e^{-i\omega_0 \tau - \tau/\tau_f}. \end{aligned} \quad (8.36)$$

La funzione di autocorrelazione (8.34) diventa in queste condizioni

$$\overline{A^*(t)A(t+\tau)} = N A_0^2 e^{-i\omega_0 \tau - \tau/\tau_f}, \quad (8.37)$$

ed il grado di coerenza al primo ordine della (8.28) diventa

$$g^{(1)}(\tau) = e^{-i\omega_0 \tau - \tau/\tau_f}. \quad (8.38)$$

Il fattore $e^{-\tau/\tau_0}$ determina lo smorzamento delle frange di interferenza al crescere del ritardo relativo τ tra i due fasci. Ad un risultato formalmente analogo si giunge anche sotto ipotesi diverse. Ricordiamo che il singolo atomo, emettendo radiazione, perde energia e che quindi, anche in assenza di collisioni, il treno d'onda emesso non può avere durata infinita. Come modello classico con ipotesi del tutto analoghe a quelle appena usate, abbiamo che il campo dell'onda emessa dal j -esimo atomo può essere scritto

$$A_j(t) = A_0 e^{-i(\omega_0 t - \varphi_j)} e^{-(t-t_j)/\tau_0}, \quad (8.39)$$

dove t_j è l'istante in cui il j -esimo atomo ha iniziato l'emissione, φ_j è la sua fase, e τ_0 è la vita media. Riprendendo la (8.34) abbiamo, tenendo conto della scorrelazione tra i vari t_j e φ_j ,

$$\overline{A^*(t)A(t+\tau)} = N \overline{A_j^*(t)A_j(t+\tau)} = N A_0^2 e^{-i\omega_0 \tau} e^{-\tau/\tau_0}, \quad (8.40)$$

con un grado di coerenza al primo ordine

$$g^{(1)}(\tau) = e^{-i\omega_0 \tau} e^{-\tau/\tau_0}. \quad (8.41)$$

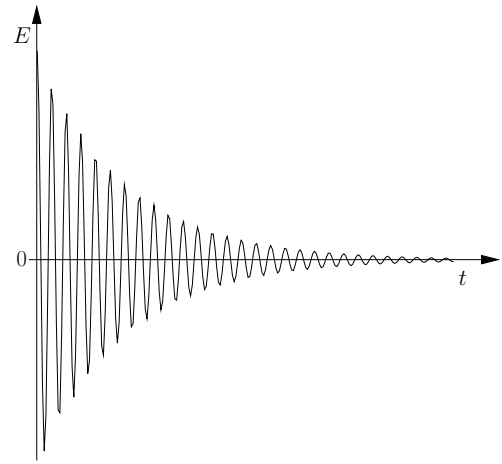


Figura 8.7 Decadimento esponenziale

8.5 Spettroscopia a trasformata di Fourier

Consideriamo adesso una radiazione non monocromatica, emessa da una sorgente, che incida su di uno spettroscopio di Michelson. Avendo a che fare con una radiazione non monocromatica, ci conviene ragionare in termini di campo elettrico dell'onda $E(t)$ anziché di ampiezza complessa $A(t)$. Riprendiamo quanto visto nel paragrafo 7.3 a proposito della trasformata di Fourier della funzione di autocorrelazione e della densità spettrale. Il campo elettrico $E(t)$ dell'onda che incide sull'interferometro è un segnale stocastico stazionario, e quindi ci aspettiamo che $\int_{-\infty}^{+\infty} |E(t)|^2 dt$ diverga, e che diverga la sua densità spettrale di energia $\Phi(\nu) = |\tilde{E}(\nu)|^2$. In altre parole, se aspettiamo un tempo abbastanza lungo, l'energia che incide sullo spettroscopio diventa grande quanto vogliamo. Ma riprendendo le equazioni (7.25) – (7.31), possiamo prima definire la funzione $E_T(t)$ tale che

$$E_T(t) = \begin{cases} E(t) & \text{se } -T < t < +T \\ 0 & \text{se } t < -T \text{ oppure } t > +T \end{cases} \quad (8.42)$$

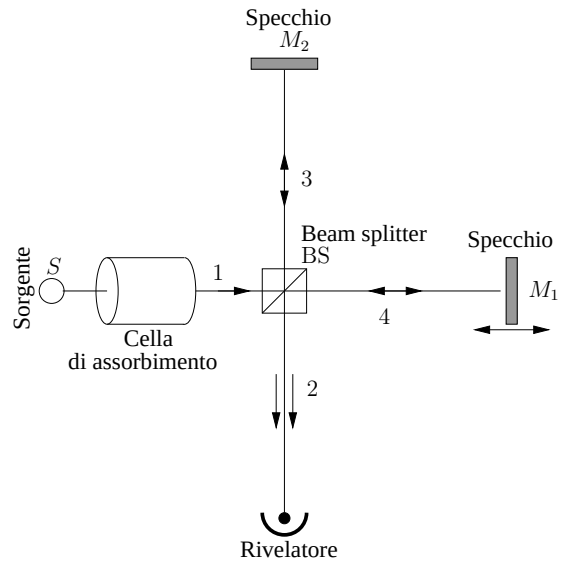


Figura 8.8 Schema di spettrometro a trasformata di Fourier.

di cui quindi esiste la trasformata di Fourier

$$\tilde{E}_T(\nu) = \int_{-\infty}^{+\infty} E_T(t) e^{-2\pi i \nu t} dt = \int_{-T}^T E(t) e^{-2\pi i \nu t} dt, \quad (8.43)$$

visto che l'intervallo temporale in cui la funzione è diversa da zero è limitato. Possiamo scrivere il teorema di Parseval per $E_T(t)$ nella forma

$$\int_{-\infty}^{+\infty} E_T^2(t) dt = \int_{-\infty}^{+\infty} |\tilde{E}_T(\nu)|^2 d\nu. \quad (8.44)$$

Se definiamo il valor medio di $E^2(t)$ come

$$\overline{E^2(t)} = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} E^2(t) dt \quad (8.45)$$

per il teorema di Parseval abbiamo

$$\overline{E^2(t)} = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-\infty}^{+\infty} |\tilde{E}_T(\nu)|^2 d\nu. \quad (8.46)$$

Moltiplicando entrambi i membri per $\varepsilon_0 c$ otteniamo

$$\varepsilon_0 c \overline{E^2(t)} = \varepsilon_0 c \int_{-\infty}^{+\infty} \lim_{T \rightarrow \infty} \frac{1}{2T} |\tilde{E}_T(\nu)|^2 d\nu, \quad \text{da cui} \quad S_i = \int_{-\infty}^{+\infty} G(\nu) d\nu, \quad (8.47)$$

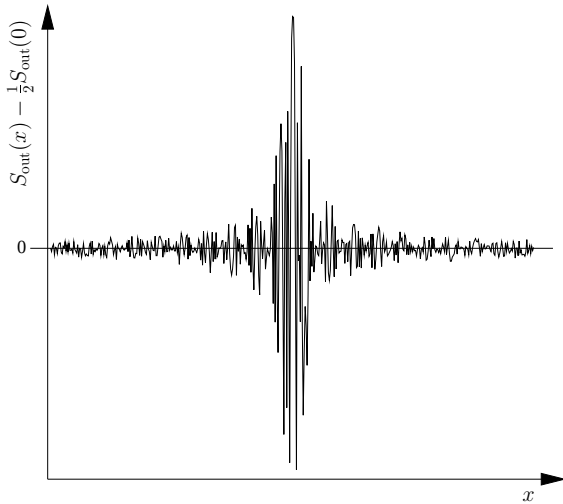


Figura 8.9 Esempio di interferogramma registrato da uno spettrometro a trasformata di Fourier.

dove S_i è l'intensità dell'onda in ingresso definita dalla (8.25), e $G(\nu) = \lim_{T \rightarrow \infty} \frac{1}{2T} \varepsilon_0 c |\tilde{E}_T(\nu)|^2$ è la densità spettrale di potenza della stessa onda in ingresso. Notare che qui non compare il fattore $1/2$ che appariva nella (8.9) perché stiamo usando il campo reale $E(t)$ anziché il campo complesso $A(t)$. Sempre come nel paragrafo 7.3 definiamo la *funzione di autocorrelazione mediata nel tempo* di $E(t)$

$$C_{EE}(\tau) = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} E(t) E(t + \tau) dt, \quad (8.48)$$

e, con passaggi del tutto analoghi a quelli del paragrafo 7.3 otteniamo per la trasformata di Fourier di $C_{EE}(\nu)$

$$\varepsilon_0 c \tilde{C}_{EE}(\nu) = G(\nu). \quad (8.49)$$

Essendo il campo elettrico $E(t)$ reale, abbiamo $\tilde{E}^*(\nu) = \tilde{E}(-\nu)$, per cui $|\tilde{E}(\nu)|^2 = \tilde{E}^*(\nu) \tilde{E}(\nu)$ è una funzione pari, e così anche $G(\nu)$. Se applichiamo l'antitrasformata di Fourier ai due membri

della (8.49) abbiamo

$$\begin{aligned}
 \varepsilon_0 c C_{EE}(\tau) &= \int_{-\infty}^{+\infty} G(\nu) e^{2\pi i \nu \tau} d\nu \\
 &= \int_0^{\infty} G(\nu) [e^{2\pi i \nu \tau} + e^{-2\pi i \nu \tau}] d\nu \\
 &= 2 \int_0^{\infty} G(\nu) \cos(2\pi \nu \tau) d\nu
 \end{aligned} \tag{8.50}$$

e la relazione inversa

$$G(\nu) = \varepsilon_0 c \, 2 \int_0^{\infty} C_{EE}(\tau) \cos(2\pi \nu \tau) d\tau. \tag{8.51}$$

Ricordando che l'interferometro di Michelson è un interferometro di Mach-Zehnder ripiegato su sé stesso, possiamo ripetere i calcoli del paragrafo 8.3 ed ottenere per il segnale in uscita

$$S_{\text{out}}(\tau) = \frac{1}{2} S_i [1 + g^{(1)}(\tau)], \tag{8.52}$$

dove adesso

$$g^{(1)}(\tau) = \frac{\overline{E(t)E(t+\tau)}}{\overline{E^2(t)}} = \frac{C_{EE}(\tau)}{\overline{E^2}} \tag{8.53}$$

e non c'è bisogno di prendere la parte reale a perché $E(t)$ è reale. Data la stazionarietà del segnale viene persa ogni dipendenza da t . La $g^{(1)}(\tau)$ è normalizzata, ed abbiamo

$$g^{(1)}(0) = 1, \quad \text{e} \quad \lim_{\tau \rightarrow \infty} g^{(1)}(\tau) = 0, \quad \text{da cui} \quad S_{\text{out}}(0) = S_i \quad \text{e} \quad \lim_{\tau \rightarrow \infty} S_{\text{out}}(\tau) = \frac{1}{2} S_i. \tag{8.54}$$

Per $\tau = 0$ il segnale in uscita è così uguale a tutto il segnale in entrata. Dalla figura Fig. 8.8 vediamo infatti che le due onde che si ricombinano sul rivelatore (percorso 2) hanno subito una prima riflessione poi una trasmissione da parte del beam-splitter (percorso 1-3-3-2), e l'altra prima una trasmissione poi una riflessione (percorso 1-4-4-2). Quindi i passaggi attraverso il beam-splitter, in base alla (8.21), non sfasano un'onda rispetto all'altra. L'unico sfasamento è dovuto alla differenza di cammino ottico, e, per $\tau = 0$, abbiamo interferenza costruttiva. Sempre dalla Fig. 8.8 vediamo che delle due onde riinviate verso la sorgente una ha subito due riflessioni sul beam-splitter (percorso 1-3-3-1), e l'altra due trasmissioni (percorso 1-4-4-1). Complessivamente, per la (8.21), le due onde vengono così sfasate l'una rispetto all'altra di π , e, per $\tau = 0$, abbiamo interferenza distruttiva. Quindi per $\tau = 0$ il segnale sul rivelatore è uguale a tutto il segnale in ingresso, mentre, per $\tau \rightarrow \infty$, con la perdita dell'autocorrelazione, metà dell'intensità in ingresso finisce sul rivelatore, e metà viene riinviata verso la sorgente.

La (8.52) può essere riscritta

$$S_{\text{out}}(\tau) = \frac{1}{2} S_{\text{out}}(0) \left[1 + \frac{C_{EE}(\tau)}{\overline{E^2}} \right], \tag{8.55}$$

da cui ricaviamo

$$\frac{1}{2} S_{\text{out}}(0) \frac{C_{EE}(\tau)}{\overline{E^2}} = S_{\text{out}}(\tau) - \frac{1}{2} S_{\text{out}}(0). \quad (8.56)$$

Ricordando che

$$S_{\text{out}}(0) = S_i = \varepsilon_0 c \overline{E^2} \quad \text{otteniamo}$$

$$C_{EE}(\tau) = \frac{2}{\varepsilon_0 c} \left[S_{\text{out}}(\tau) - \frac{1}{2} S_{\text{out}}(0) \right], \quad (8.57)$$

che inserita nella (8.51) ci dà

$$G(\nu) = 4 \int_0^{+\infty} \left[S_{\text{out}}(\tau) - \frac{1}{2} S_{\text{out}}(0) \right] \cos(2\pi\nu\tau) d\tau. \quad (8.58)$$

E' quindi possibile ricostruire la densità spettrale della radiazione facendo la trasformata di Fourier della differenza tra il segnale in uscita in funzione di τ e la metà del segnale che si ha per $\tau = 0$. Questo è il principio alla base della *spettroscopia a trasformata di Fourier*.

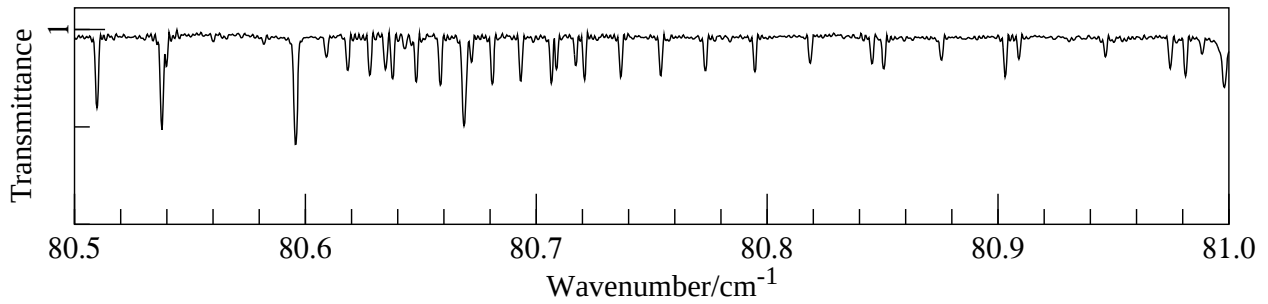


Figura 8.10 Porzione di spettro a trasformata di Fourier del metanolo.

Uno spettrometro a trasformata di Fourier è schematizzato in Fig. 8.8. E' costituito da un interferometro di Michelson con una corsa molto lunga dello specchio mobile M_1 (a seconda della precisione richiesta e della regione spettrale studiata, la corsa può raggiungere o superare i due metri, anche se corse molto più brevi sono comuni). Si sceglie una sorgente di radiazione che emetta uno spettro il più possibile continuo nella regione spettrale di interesse. La radiazione attraversa una *cella di assorbimento* contenente il campione da analizzare. Le righe di assorbimento modificheranno lo spettro emesso dalla sorgente. In realtà, spesso la cella di assorbimento viene posta all'uscita anziché all'ingresso dello spettrometro. Il segnale sul rivelatore viene misurato in funzione della differenza di cammino ottico x determinata dalla posizione dello specchio mobile. Viene quindi ricostruita, per campionamento, la funzione $S_{\text{out}}(x)$ anziché $S_{\text{out}}(\tau)$, con $\tau = x/c$. Conviene così riscrivere la 8.58 nella forma

$$G\left(\frac{\nu}{c}\right) = 4 \int_0^{+\infty} \left[S_{\text{out}}(x) - \frac{1}{2} S_{\text{out}}(0) \right] \cos\left(2\pi\nu\frac{x}{c}\right) d\left(\frac{x}{c}\right),$$

per poi porre $\tilde{\nu} = \frac{\nu}{c}$ e $\Phi(\tilde{\nu}) = c G(\nu)$, in modo che sia $\Phi(\tilde{\nu}) d\tilde{\nu} = G(\nu) d\nu$, ottenendo

$$\Phi(\tilde{\nu}) = 4 \int_0^{+\infty} \left[S_{\text{out}}(x) - \frac{1}{2} S_{\text{out}}(0) \right] \cos(2\pi\tilde{\nu}x) dx. \quad (8.59)$$

Alla grandezza $\tilde{\nu} = \nu/c = 1/\lambda$ si dà il nome di *numero d'onda*, ed indica il numero di onde di frequenza ν che stanno nell'unità di lunghezza. Tradizionalmente, i numeri d'onda sono misurati in cm^{-1} . La spettroscopia a trasformata di Fourier ha avuto grande sviluppo dopo la diffusione dei calcolatori elettronici, data la necessità di calcolare numericamente un'antitrasformata di Fourier per ottenere lo spettro.

8.6 Forma di riga

E' importante notare che la presenza di un tempo di coerenza finito determina una *non monocromaticità* della radiazione. Se avessimo un'onda incidente rigorosamente monocromatica, con ampiezza complessa $A(t) = A_0 e^{-2\pi i \nu_0 t}$ la (8.28) ci darebbe per la coerenza temporale al primo ordine

$$g^{(1)}(\tau) = \frac{|A_0|^2 e^{2\pi i \nu_0 \tau}}{|A_0|^2}, \quad \text{da cui} \quad \text{Re}[g^{(1)}(\tau)] = \cos(2\pi \nu_0 \tau) \quad (8.60)$$

la cui antitrasformata di Fourier ci dà, per la forma di riga, una delta di Dirac $G(\nu) = \delta(\nu - \nu_0)$. Ma già nel paragrafo 8.4 abbiamo visto due meccanismi che, per una riga isolata emessa da una lampada spettrale, ci portano alle equazioni (8.38) e (8.41), cioè ad una coerenza temporale al primo ordine del tipo

$$g^{(1)}(\tau) = e^{-2\pi i \nu_0 \tau} e^{-\tau/\tau_c}, \quad \text{con} \quad \text{Re}[g^{(1)}(\tau)] = \cos(2\pi \nu_0 \tau) e^{-\tau/\tau_c}$$

dove $\tau_c = \tau_f$ per la (8.38) e $\tau_c = \tau_0$ per la (8.41). Per la densità spettrale di potenza abbiamo allora

$$G(\nu) = 4 \int_0^{+\infty} \cos(2\pi \nu_0 \tau) e^{-\tau/\tau_c} \cos(2\pi \nu \tau) d\tau. \quad (8.61)$$

Introducendo le uguaglianze di Eulero abbiamo

$$\begin{aligned} G(\nu) &= 4 \int_0^{+\infty} e^{-\tau/\tau_c} \frac{e^{2\pi i \nu_0 \tau} + e^{-2\pi i \nu_0 \tau}}{2} \frac{e^{2\pi i \nu \tau} + e^{-2\pi i \nu \tau}}{2} d\tau \\ &= \int_0^{+\infty} e^{-\tau/\tau_c} \left[e^{2\pi i (\nu_0 + \nu) \tau} + e^{2\pi i (\nu_0 - \nu) \tau} + e^{-2\pi i (\nu_0 - \nu) \tau} + e^{-2\pi i (\nu_0 + \nu) \tau} \right] d\tau \\ &= \frac{1}{\frac{1}{\tau_c} - 2\pi i (\nu_0 + \nu)} + \frac{1}{\frac{1}{\tau_c} - 2\pi i (\nu_0 - \nu)} + \frac{1}{\frac{1}{\tau_c} + 2\pi i (\nu_0 + \nu)} + \frac{1}{\frac{1}{\tau_c} + 2\pi i (\nu_0 - \nu)} \\ &= \frac{(2/\tau_c)}{\left(\frac{1}{\tau_c}\right)^2 + 4\pi^2(\nu_0 - \nu)^2} + \frac{(2/\tau_c)}{\left(\frac{1}{\tau_c}\right)^2 + 4\pi^2(\nu_0 + \nu)^2} \\ &\simeq \frac{2}{\tau_c} \frac{1}{\left(\frac{1}{\tau_c}\right)^2 + 4\pi^2(\nu_0 - \nu)^2} = \frac{1}{\pi} \frac{\Delta\nu}{\Delta\nu^2 + (\nu_0 - \nu)^2}, \end{aligned} \quad (8.62)$$

dove nell'approssimazione all'ultima riga abbiamo tenuto conto del fatto che, sperimentalmente, si ha sempre $(1/\tau_c) \ll (\nu + \nu_0)$ e, nella zona spettrale di interesse, $(\nu_0 - \nu)^2 \ll (\nu_0 + \nu)^2$. Quindi la frazione contenente il termine $(\nu_0 + \nu)^2$ al denominatore può essere trascurata rispetto a quella che contiene $(\nu_0 - \nu)^2$. All'ultimo passaggio abbiamo introdotto la semilarghezza a metà altezza $\Delta\nu = 1/(2\pi\tau_c)$. Abbiamo quindi una forma di riga *Lorentziana*. Nel caso della (8.41), con $\tau_c = \tau_0$, si parla di *larghezza naturale* della riga, mentre nel caso $\tau_c = \tau_f$ della (8.38) si parla di *allargamento collisionale* della riga o *allargamento per pressione*.

Come modello classico ulteriormente semplificato rispetto a quanto esposto sopra, ma che porta alla conoscenza di una nuova funzione di interesse, possiamo pensare i fotoni emessi dalla sorgente come treni d'onda di frequenza ν_0 e durata temporale τ_c . Ad ogni passaggio dal beam splitter ogni treno d'onda si divide in due treni d'onda della stessa forma di quello incidente, ma con intensità dimezzata. Sul rivelatore giungono due treni d'onda sovrapposti, dopo che tra loro è stato introdotto un ritardo temporale causato dalla differenza di cammino ottico Δl . C'è interferenza solo se il ritardo temporale è inferiore alla durata τ_c di ciascun treno.

Supponiamo di avere un treno d'onda con ampiezza complessa $A = A_0 e^{2\pi i \nu_0 t}$ per $-\frac{\tau_c}{2} < t < \frac{\tau_c}{2}$, e ampiezza nulla se $t < -\frac{\tau_c}{2}$ oppure $t > \frac{\tau_c}{2}$. Applicando al nostro treno d'onda la trasformata di Fourier, il fatto che il treno sia temporalmente limitato ci fa ottenere, anziché una δ di Dirac $\tilde{A}(\nu) = \delta(\nu_0 - \nu)$,

$$\begin{aligned} \tilde{A}(\nu) &= \int_{-\tau_c/2}^{+\tau_c/2} A_0 e^{2\pi i (\nu_0 - \nu) t} dt = A_0 \left[\frac{e^{2\pi i (\nu_0 - \nu) t}}{2\pi i (\nu_0 - \nu)} \right]_{-\tau_c/2}^{+\tau_c/2} \\ &= \frac{A_0}{\pi(\nu_0 - \nu)} \frac{e^{2\pi i (\nu_0 - \nu) \tau_c/2} - e^{-2\pi i (\nu_0 - \nu) \tau_c/2}}{2i} = \frac{A_0}{\pi(\nu_0 - \nu)} \sin[\pi\tau_c(\nu_0 - \nu)] \\ &= A_0 \tau_c \frac{\sin[\pi\tau_c(\nu_0 - \nu)]}{\pi\tau_c(\nu_0 - \nu)} = A_0 \tau_c \text{sinc}[\tau_c(\nu_0 - \nu)], \end{aligned} \quad (8.63)$$

dove $\text{sinc}(x)$ è la funzione sinc normalizzata, definita da

$$\text{sinc}(x) = \frac{\sin(\pi x)}{\pi x}, \quad \text{con} \quad \text{sinc}(0) = 1, \quad (8.64)$$

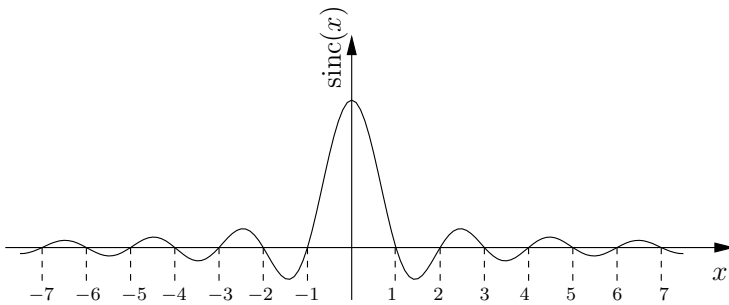


Figura 8.11 La funzione $\text{sinc}(x)$.

e mostrata in fig. 8.11. La funzione $\text{sinc}(x)$ (da *sinus cardinalis*) ha il massimo assoluto, pari a 1, per $x = 0$ (ottenuto come $\lim_{x \rightarrow 0} \sin(\pi x)/\pi x$), e si annulla per ogni altro valore intero di x . E' da notare che in letteratura si trova anche la funzione *sinc non normalizzata*, definita da $\text{sinc}(x) = \sin(x)/x$, che invece si annulla per x uguale a multipli di π diversi da zero. Noi qui, comunque,

useremo solo la sinc normalizzata definita dalla (8.64). Vediamo così dalla (8.63) che $\tilde{A}(\nu)$ è simmetrica attorno al valore $\nu = \nu_0$, al quale ha il massimo assoluto, ed ha due zeri per $\nu = \nu_0 - 1/\tau_c$ e $\nu = \nu_0 + 1/\tau_c$. Come semilarghezza della distribuzione in frequenza, possiamo quindi prendere il valore

$$\Delta\nu = \frac{1}{\tau_c}. \quad (8.65)$$

Per esempio, se prendiamo una sorgente termica a distribuzione spettrale stretta (lampada spettrale), come quelle usate tipicamente in laboratorio, la sua larghezza di banda $\Delta\nu$ è dell'ordine di 10^9 Hz, corrispondente alla tipica vita media di un livello atomico eccitato che è dell'ordine di 10^{-9} s. Questo implica che la differenza massima di cammino ottico a cui osserviamo ancora interferenza è $l_c = c\tau_c \simeq 0.3$ m. Invece una sorgente laser ben stabilizzata può avere una larghezza di riga $\Delta\nu \simeq 10^4$ Hz, cui corrispondono $\tau_c \simeq 10^{-4}$ s, e $l_c = c\tau_c \simeq 3 \times 10^4$ m = 30 km.

8.7 Esperimento di Michelson e Morley

L'interferometro di Michelson è noto per essere stato usato nell'esperimento di Michelson e Morley, eseguito una prima volta dal solo Albert Michelson a Potsdam nel 1881, e poi, in forma più raffinata, insieme a Edward Morley nel 1887 a Cleveland. Scopo dell'esperimento era misurare la velocità della Terra rispetto all'*etere*. Infatti le teorie fisiche della fine del '900 postulavano che, come le onde del mare o di un lago si propagano nell'acqua, o come le onde acustiche hanno bisogno di un mezzo in cui propagarsi, anche le onde luminose si propagassero in un mezzo particolare, detto "etere". Il valore della velocità della luce ottenuto dalle equazioni di Maxwell, $c = 1/\sqrt{\epsilon_0\mu_0}$, quella che oggi chiamiamo *velocità della luce nel vuoto*, era considerato valido solo in un sistema a riposo nell'*etere*. Osservando la propagazione della luce in un sistema di riferimento in moto rispetto all'*etere*, si sarebbe dovuta osservare una composizione delle velocità secondo le regole della fisica classica. Dato l'altissimo valore della velocità della luce, un esperimento per determinare la presenza e le proprietà dell'*etere* avrebbe dovuto richiedere un'altissima precisione e la maggior velocità possibile del sistema dell'apparato sperimentale rispetto all'*etere*. Qui descriveremo l'esperimento per sommi capi, senza scendere nei particolari accorgimenti usati per minimizzare tutte le possibili cause di rumore.

Si pensava che la Terra stessa, viaggiando attorno al sole con una velocità di 30 km/s, viaggiasse con la stessa velocità rispetto all'*etere*. L'interferometro di fig. 8.5 veniva quindi orientato in modo che, per esempio, i cammini 1 e 4 fossero paralleli alla velocità della Terra, e i cammini 2 e 3 perpendicolari. Facciamo l'ipotesi che i percorsi 3 e 4 abbiano esattamente la stessa lunghezza l . Nel percorso 1 tutta la luce emessa dalla sorgente viaggia insieme fino al beam-splitter, e non vengono introdotti ritardi relativi. Allo stesso modo nel percorso 2 la luce ricombinata viaggia di nuovo insieme. La luce riflessa dal beam-splitter percorre due volte, in andata e ritorno, il percorso 3, viaggiando perpendicolarmente alla velocità della terra e quindi non risentendo di effetti di trascinamento. Trascurando gli effetti della piccola variazione di percorso dovuta alla perpendicolarità tra direzione di propagazione della luce e velocità della Terra, ritorna al beam-splitter con un ritardo $\tau_R = 2l/c$. Invece, il fascio trasmesso percorre in andata e ritorno il percorso 4, e, secondo la legge della composizione classica delle velocità dovrebbe subire un ritardo

$$\tau_T = \frac{l}{c+v} + \frac{l}{c-v} = \frac{lc - lv + lc + lv}{c^2 - v^2} = \frac{2l}{c} \frac{1}{1 - v^2/c^2} > \frac{2l}{c} = \tau_R, \quad (8.66)$$

dove v è la velocità della Terra. Michelson e Morley si aspettavano che il ritardo relativo producesse una variazione di 0.04 frange, con un apparato che permetteva una precisione di circa 1/100 di frangia. In particolare poi, una rotazione globale dell'apparato sperimentale avrebbe dovuto portare all'osservazione di uno spostamento delle frange. Spostamenti delle frange avrebbero dovuto essere osservati anche nel corso del giorno, associati alla rotazione terrestre.

Ma nessuno spostamento delle frange viene osservato, e l'esperimento di Michelson e Morley è diventato, probabilmente, il più famoso "esperimento fallito" della fisica. I tentativi di spiegare il risultato dell'esperimento hanno portato alla formulazione della teoria della relatività.

8.8 Interferometro di Young e coerenza spaziale

L'interferometro di Young, realizzato per la prima volta da Thomas Young nel 1802 con lo scopo di determinare se la natura della luce era ondulatoria o corpuscolare, è schematizzato in Fig. 8.12.a). Una sorgente estesa quasi monocromatica (spesso ottenuta ponendo dopo una sorgente estesa termica un filtro ottico che lasci passare radiazione quasi monocromatica), a forma di disco di raggio a , illumina i due forellini P_1 e P_2 , distanti $2d$ tra loro, sullo schermo opaco A , distante r dalla sorgente stessa. Nello

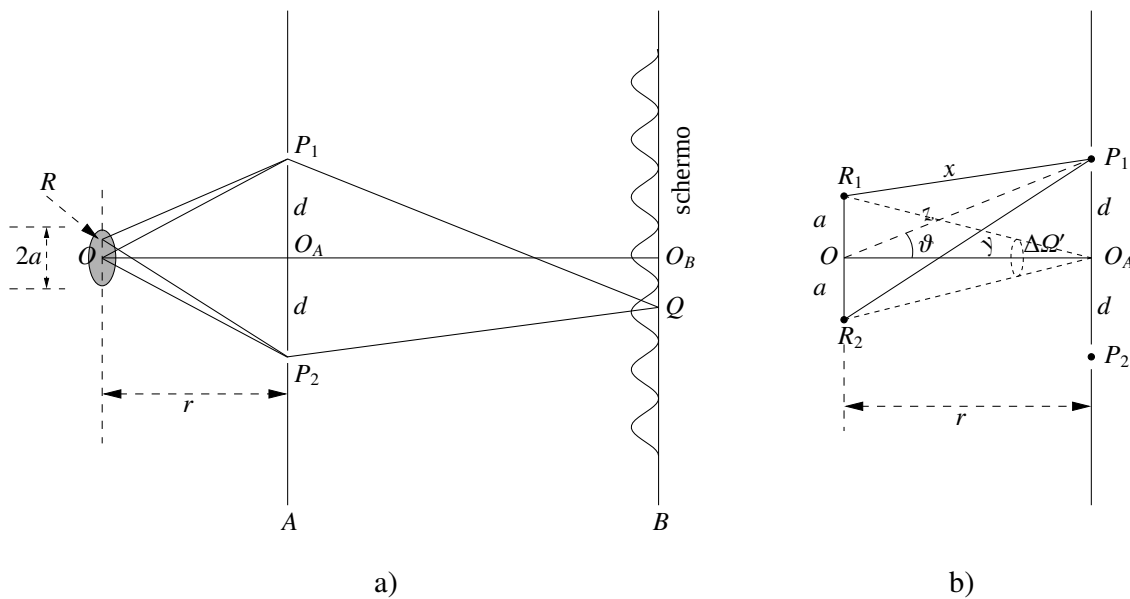


Figura 8.12 Interferometro di Young.

schermo opaco dell'interferometro di Young originale erano presenti in realtà due fenditure, che qui, per chiarezza di ragionamento, abbiamo sostituito con i due forellini (pinholes). Le ampiezze e fasi totali delle *sorgenti secondarie* (in base al principio di Huygens-Fresnel) in P_1 e P_2 sono ottenute per sovrapposizione di tutte le onde parziali emesse dagli elementi di superficie $d\sigma$ in Fig. 8.12, attorno ai punti R della sorgente, tenendo conto dei diversi cammini ottici $\overline{RP_1}$ e $\overline{RP_2}$. L'intensità misurata al punto Q sullo schermo B dipende sia dalla differenza di cammino ottico $\Delta l(Q) = \overline{P_1Q} - \overline{P_2Q}$ che dalla differenza di fase $\Delta\varphi = \varphi(P_1) - \varphi(P_2)$ tra le ampiezze del campo in P_1 e P_2 . Se i diversi punti R della sorgente emettono indipendentemente con fasi casuali (sorgente termica), anche le fasi delle ampiezze totali in P_1 e P_2 fluttueranno a caso. Tuttavia questo non influenza l'intensità in Q se le fluttuazioni in P_1 e P_2 sono sincrone, perché in questo caso la differenza di fase $\Delta\varphi$ rimane costante. In queste condizioni i due fori P_1 e P_2 si comportano come sorgenti coerenti tra loro, che generano una figura di interferenza sullo schermo B . Questa situazione si realizza senz'altro per la radiazione emessa dal centro O della sorgente perché i cammini $\overline{OP_1}$ e $\overline{OP_2}$ sono uguali e tutte le fluttuazioni di fase in O arrivano simultaneamente in P_1 e P_2 . Per ogni punto R della sorgente diverso

da O , invece, si ha una differenza di cammino ottico $\Delta l(R) = \overline{RP_1} - \overline{RP_2}$, che è massima se R si trova sul bordo della sorgente. Nella fig. 8.12.b), che è stata deformata per chiarezza di lettura, i punti R_1 e R_2 , diametralmente opposti, si trovano sul bordo della sorgente. Chiamiamo x la distanza $\overline{R_1P_1}$, y la distanza $\overline{R_2P_1} = \overline{R_1P_2}$, e z la distanza $\overline{OP_1}$. Vogliamo calcolare la differenza di cammino $\Delta l_{\max} = y - x$ nell'approssimazione che sia $a \ll r$. Appliciamo il teorema di Pitagora generalizzato ai due triangoli OR_1P_1 e OR_2P_1 . Dalla figura vediamo che $\cos(\widehat{R_1OP_1}) = \sin \vartheta$, e $\cos(\widehat{R_2OP_1}) = -\sin \vartheta$, quindi abbiamo per x , y e z

$$x = \sqrt{a^2 + z^2 - 2az \sin \vartheta}, \quad y = \sqrt{a^2 + z^2 + 2az \sin \vartheta}, \quad z = \frac{r}{\cos \vartheta}. \quad (8.67)$$

Per ϑ piccolo possiamo usare lo sviluppo al primo ordine della radice quadrata

$$\sqrt{c + \delta} \simeq \sqrt{c} + \frac{\delta}{2\sqrt{c}}, \quad \text{valido se } \delta \ll c, \quad (8.68)$$

ottenendo

$$x \simeq \sqrt{a^2 + z^2} - \frac{az \sin \vartheta}{\sqrt{a^2 + z^2}} \quad \text{e} \quad y \simeq \sqrt{a^2 + z^2} + \frac{az \sin \vartheta}{\sqrt{a^2 + z^2}}, \quad (8.69)$$

da cui

$$\Delta l_{\max} = y - x \simeq \frac{2az \sin \vartheta}{\sqrt{a^2 + z^2}}, \quad \text{che, per } a \ll z, \text{ porta a } \Delta l_{\max} \simeq 2a \sin \vartheta. \quad (8.70)$$

Quando $\Delta l_{\max}$ vale $\lambda/2$ un'onda sferica emessa da R_1 (o da R_2) giunge in P_1 e P_2 in controfase. Questo comporta che, sullo schermo B , i massimi di illuminazione della figura di interferenza generata da un'onda emessa da R_1 coincidano con i minimi di illuminazione della figura di interferenza generata da un'onda emessa da O , e viceversa. Questo porta a una illuminazione uniforme dello schermo e alla scomparsa della figura di interferenza. In realtà la situazione è un po' più complicata perché la radiazione viene emessa non solo da O , R_1 e R_2 , ma da tutti i punti della sorgente compresi tra R_1 e R_2 . Ognuno di questi punti genera onde che giungono in P_1 e P_2 con sfasamenti diversi. Il risultato è che la condizione per avere la cancellazione dell'interferenza non è esattamente $\Delta l_{\max} > \lambda/2$, ma

$$\Delta l_{\max} > \frac{2}{\pi} \lambda \simeq 0.6366 \lambda. \quad (8.71)$$

Quindi avremo illuminazione coerente in P_1 e P_2 se

$$\Delta l_{\max} = 2a \sin \vartheta < \frac{2}{\pi} \lambda, \quad \text{cioè se } \sin \vartheta < \frac{\lambda}{\pi a}. \quad (8.72)$$

Dalla fig. 8.12.b) vediamo che, per ϑ piccolo, $\sin \vartheta \simeq \tan \vartheta \simeq d/r$, per cui possiamo scrivere la condizione di coerenza nella forma

$$d < \lambda \frac{r}{\pi a}. \quad (8.73)$$

Perché si possano osservare le figure di interferenza, bisogna quindi che i due fori P_1 e P_2 si trovino all'interno di un'area di coerenza circolare

$$A_c = \pi d_{\max}^2 = \frac{\pi \lambda^2 r^2}{\pi^2 a^2}, \quad \text{e, ricordando che } A_s = \pi a^2, \quad A_c = \frac{\lambda^2 r^2}{A_s}. \quad (8.74)$$

Dalla sorgente, il cerchio di area A_c è visto entro un angolo solido $\Delta\Omega_c = A_c/r^2$, da cui

$$\Delta\Omega_c = \frac{\lambda^2}{A_s}. \quad (8.75)$$

Riassumendo quanto sopra, possiamo dire che, se abbiamo una sorgente quasi monocromatica di area A_s che emette attorno alla lunghezza d'onda λ , l'area massima A_c che può essere illuminata coerentemente ad una distanza r vale

$$A_c = r^2 \Delta\Omega_c = \frac{\lambda^2 r^2}{A_s} = \frac{\lambda^2}{\Delta\Omega'}, \quad (8.76)$$

dove

$$\Delta\Omega' = \frac{A_s}{r^2} \quad (8.77)$$

è l'angolo solido sotto il quale il punto O_A della fig. 8.12.b) vede l'area della sorgente. Tra $\Delta\Omega'$ e $\Delta\Omega_c$ esiste la relazione

$$\Delta\Omega' = \frac{\lambda^2}{r^2} \frac{1}{\Delta\Omega_c}. \quad (8.78)$$

La superficie della sorgente A_s determina quindi un angolo solido massimo $\Delta\Omega_c = \lambda^2/A_s$ entro al quale il campo di radiazione ha coerenza spaziale. Una volta fissate le dimensioni della sorgente, la superficie di coerenza $A_c = \pi d^2$ aumenta con il quadrato della distanza r dalla sorgente, o, che è lo stesso, in modo inversamente proporzionale all'angolo solido $\Delta\Omega'$. Data la grande distanza dalle stelle, e il conseguente piccolo valore di $\Delta\Omega'$ nella (8.77), la luce di una stella ricevuta da un telescopio è spazialmente coerente su tutta la superficie dell'apertura del telescopio stesso, nonostante il grande diametro della sorgente di radiazione.

8.9 Volume di coerenza

Con una lunghezza di coerenza $\Delta l_c = c \tau_c = c/\Delta\nu$ nella direzione di propagazione della radiazione e la superficie di coerenza $A_c = \lambda^2 r^2/A_s$ abbiamo un volume di coerenza

$$V_c = A_c \Delta l_c = \frac{\lambda^2 r^2 c}{\Delta\nu A_s}. \quad (8.79)$$

La radianza spettrale $L(\nu)$ di una sorgente è definita in modo che $L(\nu) dS d\nu d\Omega$ sia la potenza radiante emessa dall'elemento di superficie dS della sorgente stessa, nell'intervallo spettrale $(\nu, \nu + d\nu)$ all'interno dell'angolo solido $d\Omega$. Le unità di misura per $L(\nu)$ sono quindi $\text{W}/(\text{m}^2 \cdot \text{sterad} \cdot \text{Hz})$. Se $L(\nu)$ è la radianza spettrale della nostra sorgente di superficie A_s , il numero medio di fotoni $\langle n \rangle$ di frequenza compresa nell'intervallo spettrale $(\nu, \nu + \Delta\nu)$ entro l'angolo solido di coerenza $\Delta\Omega_c = \lambda^2/A_s$ durante il tempo di coerenza $\tau_c = 1/\Delta\nu$ sarà

$$\langle n \rangle = \frac{1}{h\nu} A_s L(\nu) \Delta\nu \Delta\Omega_c \tau_c = \frac{1}{h\nu} A_s L(\nu) \Delta\nu \frac{\lambda^2}{A_s} \frac{1}{\Delta\nu} = \frac{L(\nu) \lambda^2}{h\nu}. \quad (8.80)$$

E' possibile calcolare il numero g delle celle dello spazio delle fasi occupato dai fotoni di un fascio luminoso di apertura $\Delta\Omega$ partendo dall'espressione per il numero delle celle dello spazio delle fasi con quantità di moto compresa tra p e $p + \Delta p$, al limite per grandi volumi spaziali V

$$g = \frac{V 4\pi p^2 \Delta p}{h^3}. \quad (8.81)$$

Sostituiamo l'angolo solido 4π , che abbraccia tutto lo spazio, con l'angolo solido $\Delta\Omega$ entro cui vediamo arrivare i fotoni del fascio. Infine, ricordando che per un fotone $p = \frac{h\nu}{c}$, otteniamo

$$g = \frac{V\Delta\Omega_c \nu^2 \Delta\nu}{c^3}. \quad (8.82)$$

Se vogliamo che i fotoni occupino un'unica cella dello spazio delle fasi dobbiamo porre $g = 1$, e questo ci porta alla condizione per il volume

$$V = \frac{c^3}{\Delta\Omega \nu^2 \Delta\nu} = \frac{\lambda^2 c}{\Delta\Omega \Delta\nu}. \quad (8.83)$$

A distanza r dalla sorgente, la sorgente stessa viene vista sotto un angolo solido $\Delta\Omega = A_s/r^2$, e i fotoni vengono visti arrivare entro questo angolo solido. Sostituendo A_s/r^2 al posto di $\Delta\Omega$ nella (8.83) otteniamo

$$V = \frac{\lambda^2 r^2 c}{A_s \Delta\nu}, \quad (8.84)$$

che coincide con il volume di coerenza della (8.79).

Il volume occupato dai fotoni che attraversano una generica superficie S , sufficientemente lontana dalla sorgente e perpendicolare al fascio, durante il tempo t_m di una misura, può essere approssimato dall'espressione

$$V = S c t_m. \quad (8.85)$$

Dalla superficie S la superficie della sorgente A_s viene vista sotto un angolo solido $\Delta\Omega'$. Quindi la (8.82) ci dà un numero di celle dello spazio delle fasi occupato dai fotoni pari a

$$g = \frac{S c t_m \Delta\Omega' \nu^2 \Delta\nu}{c^3} = \frac{S \Delta\Omega'}{\lambda^2} \frac{t_m}{\tau_c}. \quad (8.86)$$

Possiamo definire una grandezza caratteristica del fascio U , detta *estensione del fascio*, come

$$U = \frac{A_s S}{r^2} = A_s \Delta\Omega = S \Delta\Omega', \quad (8.87)$$

dove S è la superficie il cui contorno delimita il fascio ad una distanza r dalla sorgente. L'estensione U è importante perché rimane costante nella propagazione senza perdite di un fascio luminoso sotto le leggi dell'ottica geometrica. Per fotoni che viaggiano entro l'angolo solido $\Delta\Omega_c$ determinato dalla diffrazione alla superficie A_s della sorgente, dato dalla (8.75), l'estensione assume il valore critico U_c

$$U_c = A_s \Delta\Omega_c = A_s \frac{\lambda^2}{A_s} = \lambda^2. \quad (8.88)$$

Se sostituiamo $U_c = \lambda^2$ nell'espressione (8.86) per il numero di celle dello spazio delle fasi otteniamo

$$g = \frac{S \Delta\Omega'}{\lambda^2} \frac{t_m}{\tau_c} = \frac{U}{U_c} \frac{t_m}{\tau_c}, \quad (8.89)$$

dove il termine U/U_c viene definito fattore di coerenza spaziale, ed il termine t_m/τ_c fattore di coerenza temporale. Utilizzando la 8.76, il fattore di coerenza spaziale può anche essere definito

$$\frac{U}{U_c} = \frac{S \Delta\Omega'}{\lambda^2} = \frac{S}{A_c}, \quad (8.90)$$

cioè come rapporto tra la superficie del fascio e la superficie di coerenza, mentre il fattore di coerenza temporale è il rapporto tra il tempo di misura ed il tempo di coerenza. Il valore più basso di g , cioè 1, si ha quando l'area ed il tempo di misura coincidono con l'area ed il tempo di coerenza.

E' da notare che né il fattore di coerenza spaziale, né il fattore di coerenza temporale, possono essere presi come minori di 1. Quindi, per esempio, un tempo di misura minore del tempo di coerenza non compensa una superficie del fascio maggiore dell'area di coerenza.

8.10 Interferometro stellare di Michelson

Un'importante applicazione dell'interferenza è la determinazione dell'angolo solido sotto cui vediamo le stelle. Una stella, data la distanza, ci appare come un punto luminoso e una misura diretta del suo diametro non è possibile. Misure interferometriche permettono però di determinare l'angolo solido $\Delta\Omega$ sotto cui noi vediamo una stella. Se poi la distanza della stella è nota per altra via, da $\Delta\Omega$ possiamo risalire al diametro. Questa è l'idea alla base dell'interferometro stellare usato da Michelson nel 1920 per misurare il *diametro angolare* di Betelgeuse, una supergigante rossa della costellazione di Orione.

La luce proveniente dalla stella incide sui due specchi S_1 e S_2 , posti a distanza d tra loro. I fasci luminosi riflessi da S_1 e S_2 vengono successivamente riflessi, rispettivamente, dagli specchi S_3 e S_4 , e fatti convergere dalla lente L sul rivelatore R , dove, se esistono le condizioni, si

possono osservare figure di interferenza. Le figure di interferenza saranno osservabili solo se S_1 e S_2 si trovano all'interno della stessa area di coerenza. Secondo la 8.76 dobbiamo avere

$$\frac{1}{4}\pi d^2 < A_c = \frac{\lambda^2}{\Delta\Omega},$$

da cui

$$\Delta\Omega < \frac{4}{\pi} \frac{\lambda^2}{d^2}, \quad (8.91)$$

Figura 8.13 Interferometro usato da Michelson per l'interferometria stellare.

dove $\frac{1}{4}\pi d^2$ è l'area del cerchio di diametro d , λ è la lunghezza d'onda della radiazione osservata, e $\Delta\Omega$ è l'angolo solido sotto il quale dalla Terra si vede la stella. Per un valore di d piccolo si osservano figure di interferenza, aumentando poi gradatamente d l'interferenza si attenua fino a scomparire quando si supera un certo valore $d_{\max}$. Questo ci permette di determinare $\Delta\Omega$.

Dalla figura 8.14 possiamo ricavare la relazione tra il diametro angolare α di una stella e l'angolo solido sotto cui la stella viene vista. Nella approssimazione, valida per le stelle, di α molto piccolo, abbiamo infatti

$$\Delta\Omega \simeq \frac{1}{r^2} \pi \left(\frac{a}{2}\right)^2 \simeq \frac{1}{r^2} \pi \left[r^2 \tan\left(\frac{\alpha}{2}\right)\right]^2 \simeq \frac{\pi}{4} \alpha^2, \quad \text{da cui} \quad \alpha \simeq 2 \sqrt{\frac{\Delta\Omega}{\pi}}. \quad (8.92)$$

dove a è il diametro della stella.

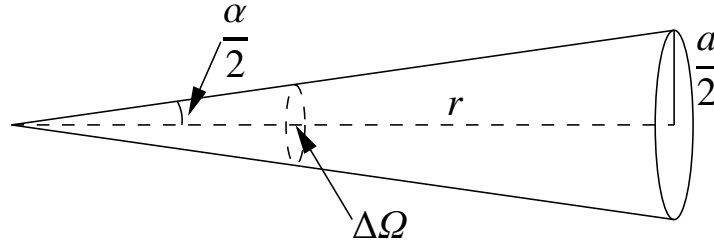


Figura 8.14 Relazione tra angolo solido e diametro angolare di una stella.

8.11 Interferometro stellare di Hanbury-Brown e Twiss

L'interferometro stellare di Hanbury-Brown e Twiss, schematizzato in Fig. 8.15, venne ideato nel 1956. Due fotomoltiplicatori R_1 e R_2 , a distanza d pari a circa 6 m tra loro, venivano puntati verso la stella Sirio. Si osservava una correlazione positiva tra le intensità misurate dai fotomoltiplicatori. Questo nonostante che si misurassero separatamente le *intensità luminose* in due punti diversi, e non la sovrapposizione di due *campi* su un unico rivelatore come nell'interferometro stellare di Michelson. Non era quindi presente alcuna informazione di fase.

Hanbury-Brown e Twiss usarono il segnale di interferenza per determinare l'apertura angolare apparente di Sirio. Il funzionamento dell'apparato può essere spiegato sia in termini di ottica classica, usando la *coerenza temporale al secondo ordine*, che in termini quantistici.

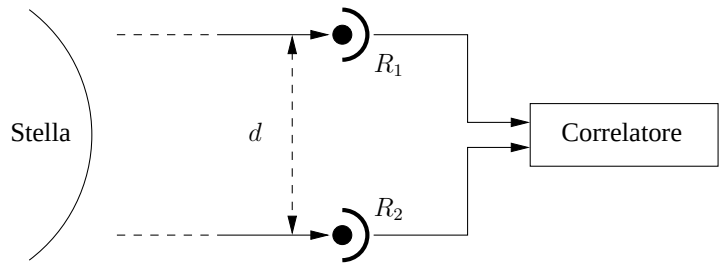


Figura 8.15 Interferometro usato da Hanbury-Brown e Twiss per l'interferometria stellare.

8.11.1 Spiegazione classica

Consideriamo prima la spiegazione classica. Supponiamo di avere un'onda elettromagnetica governata da una statistica stocastica stazionaria, descritta da un campo complesso $A(t)$. L'onda incide sul beam-splitter BS , che supponiamo sia senza perdite e con $|T|^2 = |R|^2 = 1/2$. Poi l'onda trasmessa $A_1(t)$ incide sul rivelatore R_1 , e l'onda riflessa $A_2(t)$ sul rivelatore R_2 . Supponiamo ancora che ai due rivelatori, identici, ci sia un ritardo relativo τ tra le onde. Avremo quindi per le misure effettuate dai due rivelatori all'istante t'

$$\begin{aligned} S_1(t') &= \frac{1}{2} \varepsilon_0 c \overline{|A_1(t')|^2} = \frac{1}{4} \varepsilon_0 c \overline{|A(t)|^2} \\ S_2(t') &= \frac{1}{2} \varepsilon_0 c \overline{|A_2(t')|^2} = \frac{1}{4} \varepsilon_0 c \overline{|A(t + \tau)|^2} \end{aligned} \quad (8.93)$$

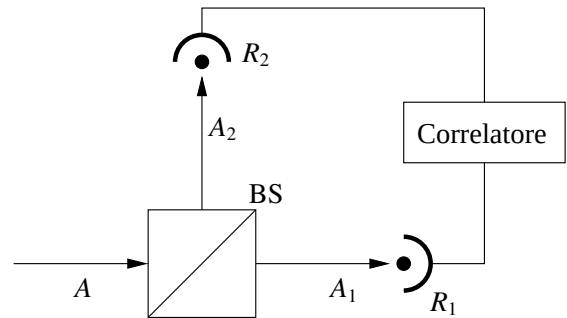


Figura 8.16 Misura di correlazione al secondo ordine.

dove la barra indica la media temporale sul tempo di misura, che questa volta supponiamo *breve rispetto alla scala temporale delle fluttuazioni*. Avremo anche $t' - t = K$, costante che non ci interessa, data la stazionarietà della radiazione. Effettuiamo un gran numero N di coppie di misure simultanee $\{S_1^{(i)} = S_1(t_i), S_2^{(i)} = S_2(t_i)\}$ ai due rivelatori, ogni coppia misurata all'istante t_i , ed inviamo le uscite dei due rivelatori ad un *correlatore*, la cui elettronica moltiplica i segnali tra loro. Per ogni coppia di misure il segnale al correlatore sarà

$$S_{12}^{(i)} = \alpha \overline{A^*(t_i)A(t_i)A^*(t_i + \tau)A(t_i + \tau)}, \quad (8.94)$$

dove α è una costante di proporzionalità, e la barra indica la media temporale sulla durata della misura. A questo punto possiamo definire, in modo analogo alla coerenza temporale al primo ordine della (8.28), la *coerenza temporale al secondo ordine* tra i segnali rivelati da R_1 e quelli rivelati da R_2 come

$$\begin{aligned} g^{(2)}(\tau) &= \frac{\sum_{i=1}^M \overline{A^*(t_i)A(t_i)A^*(t_i + \tau)A(t_i + \tau)}}{\sum_{i=1}^M \left(\overline{A^*(t_i)A(t_i)} \right)^2} \\ &= \frac{\langle \overline{A^*(t)A(t)A^*(t + \tau)A(t + \tau)} \rangle}{\langle \overline{A^*(t)A(t)} \rangle^2} = \frac{\langle I(t)I(t + \tau) \rangle}{\langle I(t) \rangle^2} \end{aligned} \quad (8.95)$$

dove le parentesi ad angolo indicano la media sulle M coppie di misure effettuate e le barre indicano le medie temporali sulle durate delle singole misure. Trattiamo la coerenza al secondo ordine $g^{(2)}(\tau)$ per una sorgente spettrale classica con formalismo analogo a quello usato nel paragrafo 8.4 per la coerenza al primo ordine $g^{(1)}(\tau)$. Facciamo quindi ancora l'ipotesi che il singolo atomo emetta un treno di radiazione con fase costante finché non subisca una collisione. Facciamo poi le ulteriori ipotesi che la durata della collisione sia trascurabile, e che immediatamente dopo la collisione l'atomo riprenda ad emettere con la stessa frequenza ω_0 , ma con fase che ha subito un brusco cambiamento casuale. In questo modo l'ampiezza complessa dell'onda incidente può ancora essere scritta [vedi(8.32)]

$$A(t) = A_0 e^{-i\omega_0 t} \sum_{j=1}^N e^{-i\varphi_j(t)}, \quad (8.96)$$

dove l'indice j corre su tutti gli N atomi della lampada. La funzione di correlazione al secondo ordine che compare al numeratore della (8.95) diventa così

$$\begin{aligned} \overline{A^*(t)A(t)A^*(t + \tau)A(t + \tau)} &= \sum_{j=1}^N \overline{A_j^*(t)A_j(t)A_j^*(t + \tau)A_j(t + \tau)} \\ &+ \sum_{j \neq k} \overline{A_j^*(t)A_k(t)A_j^*(t + \tau)A_k(t + \tau)} \\ &+ \sum_{j \neq k} \overline{A_j^*(t)A_j(t)A_k^*(t + \tau)A_k(t + \tau)}, \end{aligned} \quad (8.97)$$

dove sono stati tenuti solo i termini in cui il campo di ogni atomo è moltiplicato per il proprio complesso coniugato. Tutti gli altri termini si annullano per la casualità delle fasi delle onde emesse dai vari atomi. Ricordando che tutti gli atomi sono equivalenti, come avevamo già fatto per la (8.34), otteniamo

$$\begin{aligned} \overline{A^*(t)A(t)A^*(t+\tau)A(t+\tau)} &= N \overline{A_j^*(t)A_j(t)A_j^*(t+\tau)A_j(t+\tau)} \\ &+ N(N-1) \left\{ \left[\overline{A_j^*(t)A_j(t)} \right]^2 + \left| \overline{A_j^*(t)A_j(t+\tau)} \right|^2 \right\} \end{aligned} \quad (8.98)$$

con j indice di un singolo, generico atomo. Se il numero N degli atomi è molto grande, avremo $N(N-1) \simeq N^2$ e $N^2 \gg N$, cioè, i contributi che coinvolgono coppie di atomi sono molto superiori ai contributi di singolo atomo. Quindi, è una buona approssimazione scrivere

$$\overline{A^*(t)A(t)A^*(t+\tau)A(t+\tau)} \simeq N^2 \left\{ \left[\overline{A_j^*(t)A_j(t)} \right]^2 + \left| \overline{A_j^*(t)A_j(t+\tau)} \right|^2 \right\}. \quad (8.99)$$

Confrontando con l'espressione per $g^{(1)}(\tau)$ della (8.28), possiamo scrivere

$$g^{(2)}(\tau) = 1 + \left| g^{(1)}(\tau) \right|^2 \quad (\text{per } N \gg 1), \quad (8.100)$$

quindi è possibile ricavare informazioni sul modulo di $g^{(1)}(\tau)$ da una misura di $g^{(2)}(\tau)$. Per come è definita, $g^{(2)}(\tau)$ è reale, e, per la stazionarietà ipotizzata per la statistica della radiazione, abbiamo

$$g^{(2)}(-\tau) = g^{(2)}(\tau). \quad (8.101)$$

Inoltre vediamo che $g^{(2)}(0) = 2$, e che $g^{(2)}(\tau) \rightarrow 1$ per $\tau \gg \tau_c$, tempo di correlazione.

Quindi vediamo che, con l'apparato di Fig. 8.15, misureremo correlazione fin tanto che d sarà tale da tenere i due rivelatori entro l'area di coerenza.

8.11.2 Spiegazione quantistica

C'è una spiegazione quantistica dovuta a Ugo Fano (1961). Nella figura 8.17 consideriamo una coppia di fotoni che giungono in coincidenza sui rivelatori R_1 e R_2 , provenienti dai punti a e b della sorgente, nel nostro caso la stella. Ci sono quattro possibilità: i) il punto a ha emesso due fotoni, uno raccolto da R_1 (percorso tratteggiato) e l'altro raccolto da R_2 (percorso punteggiato); ii) il punto b ha emesso due fotoni, ancora una volta raccolti uno da R_1 (percorso punteggiato) e l'altro da R_2 (percorso tratteggiato); iii) il punto a ha emesso il fotone raccolto da R_1 e il punto b quello raccolto da R_2 (i due percorsi tratteggiati); iv) il punto a ha emesso il fotone raccolto da R_2 e il punto b quello raccolto da R_1 (i due percorsi punteggiati). I primi due processi sono distinguibili tra loro e non danno origine a interferenza. Consideriamo invece i processi iii) e iv), chiamando $\langle 1|a \rangle$ l'ampiezza di probabilità

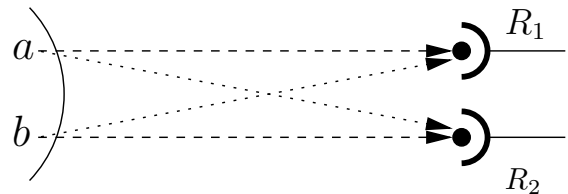


Figura 8.17 Possibili eventi che portano alla rivelazione simultanea di due fotoni provenienti da a e b da parte dei rivelatori 1 e 2.

che R_1 abbia rivelato il fotone emesso da a , $\langle 1|b \rangle$ la probabilità che invece R_1 abbia rivelato il fotone emesso da b , e così via. Poniamo anche, per fissare le idee,

$$|\langle 1|a \rangle|^2 = |\langle 1|b \rangle|^2 = |\langle 2|a \rangle|^2 = |\langle 2|b \rangle|^2 = p.$$

Se avessimo a che fare con particelle distinguibili la rivelazione simultanea di due fotoni su R_1 e R_2 , provenienti uno da a e l'altro da b , sarebbe dovuta o al processo iii) o al processo iv). L'ampiezza di probabilità del processo iii) sarebbe quindi $\langle 1|a \rangle \langle 2|b \rangle$ e la probabilità sarebbe $W_{iii} = |\langle 1|a \rangle|^2 |\langle 2|b \rangle|^2 = p^2$. L'ampiezza di probabilità del processo iv) sarebbe $\langle 1|b \rangle \langle 2|a \rangle$ e la sua probabilità $W_{iv} = |\langle 1|b \rangle|^2 |\langle 2|a \rangle|^2 = p^2$. Ma essendo i fotoni particelle indistinguibili, in particolare bosoni, dobbiamo ragionare in termini di ampiezze di probabilità opportunamente simmetrizzate (dovrebbero essere antisimmetrizzate se avessimo a che fare con fermioni). L'ampiezza di probabilità di ricevere un fotone su ognuno dei due rivelatori, con un fotone proveniente da a e l'altro da b , si scrive quindi

$$\alpha = \frac{1}{\sqrt{2}} (\langle 1|a \rangle \langle 2|b \rangle + \langle 2|a \rangle \langle 1|b \rangle), \quad (8.102)$$

che ci dà una probabilità di rivelazione

$$W = |\alpha|^2 = \frac{1}{2} (|\langle 1|a \rangle|^2 |\langle 2|b \rangle|^2 + |\langle 2|a \rangle|^2 |\langle 1|b \rangle|^2 + 2 \langle 1|a \rangle \langle 2|b \rangle \langle 2|a \rangle \langle 1|b \rangle) = 2p^2, \quad (8.103)$$

dove compare un termine di interferenza costruttiva che la rende superiore alla probabilità dei due eventi separati.

In altri termini, chiamiamo N_1 il numero di fotoni rivelato da R_1 , con numero medio $\bar{N}_1$, e chiamiamo N_2 e $\bar{N}_2$ le corrispondenti quantità per R_2 . Il segnale di correlazione C è

$$C = (N_1 - \bar{N}_1)(N_2 - \bar{N}_2) \quad (8.104)$$

ed il correlatore fornisce la media

$$\bar{C} = \overline{(N_1 - \bar{N}_1)(N_2 - \bar{N}_2)}. \quad (8.105)$$

Se R_1 ed R_2 si trovano nella stessa area di coerenza, definita dalla 8.76, i fotoni che arrivano sui due rivelatori appartengono allo stesso gruppo di celle dello spazio delle fasi, con numero

$$g = \frac{t_m}{\tau_c}, \quad (8.106)$$

dove t_m è il tempo di misura, perché nella (8.89) il fattore di coerenza spaziale U/U_c vale 1. Per le fluttuazioni del numero dei fotoni contenuti in una singola cella $|r\rangle$ dello spazio delle fasi abbiamo dalla (5.78)

$$\langle (\Delta n_r)^2 \rangle = \langle n_r \rangle + \langle n_r \rangle^2.$$

Nel nostro caso, però, i fotoni non appartengono ad una sola cella, ma a g celle dello spazio delle fasi, praticamente equivalenti (le frequenze sono molto vicine). Quindi ci aspettiamo che il numero medio $\bar{N}_1$ di fotoni rivelati da R_1 sia $\bar{N}_1 = g \langle n_r \rangle$. Analogamente sarà $\bar{N}_2 = g \langle n_r \rangle$. Per quel che riguarda le fluttuazioni quadratiche, supponendo che le fluttuazioni di celle diverse non siano correlate, avremo

$$\overline{(\Delta N_1)^2} = \sum \langle (\Delta n_r)^2 \rangle = g \langle (\Delta n_r)^2 \rangle = g \langle n_r \rangle + g \langle n_r \rangle^2 = \bar{N}_1 + \frac{(\bar{N}_1)^2}{g}, \quad (8.107)$$

essendo $(\overline{N_1})^2 = (g\langle n_r \rangle)^2$. La formula per le fluttuazioni di N_2 è analoga. Per le fluttuazioni del numero totale $N = N_1 + N_2$ di fotoni rivelati da R_1 e R_2 avremo, in modo analogo

$$\overline{(\Delta N)^2} = \overline{(N - \overline{N})^2} = \overline{N} + \frac{(\overline{N})^2}{g}. \quad (8.108)$$

Esplicitando N abbiamo

$$\overline{(N_1 + N_2 - \overline{N_1} - \overline{N_2})^2} = \overline{N_1} + \overline{N_2} + \frac{(\overline{N_1} + \overline{N_2})^2}{g}. \quad (8.109)$$

Svolgendo i quadrati a primo e secondo membro otteniamo

$$\overline{(N_1 - \overline{N_1})^2} + \overline{(N_2 - \overline{N_2})^2} + 2\overline{(N_1 - \overline{N_1})(N_2 - \overline{N_2})} = \overline{N_1} + \overline{N_2} + \frac{(\overline{N_1})^2}{g} + \frac{(\overline{N_2})^2}{g} + 2\frac{\overline{N_1}\overline{N_2}}{g}. \quad (8.110)$$

Al secondo membro possiamo sostituire la (8.107) e l'analogo per N_2 , ottenendo

$$\overline{(N_1 - \overline{N_1})^2} + \overline{(N_2 - \overline{N_2})^2} + 2\overline{(N_1 - \overline{N_1})(N_2 - \overline{N_2})} = \overline{(N_1 - \overline{N_1})^2} + \overline{(N_2 - \overline{N_2})^2} + 2\frac{\overline{N_1}\overline{N_2}}{g}, \quad (8.111)$$

da cui, semplificando, otteniamo finalmente per il segnale di correlazione tra i due rivelatori

$$\overline{C} = \overline{(N_1 - \overline{N_1})(N_2 - \overline{N_2})} = \frac{\overline{N_1}\overline{N_2}}{g}. \quad (8.112)$$

Questa correlazione positiva indica che ci aspettiamo di avere, simultaneamente, fluttuazioni dello stesso segno su R_1 e R_2 (*bunching*). Se allontaniamo i rivelatori l'uno dall'altro, ci aspettiamo che i fotoni che giungono su R_1 e R_2 appartengano a celle diverse dello spazio delle fasi, e che la correlazione diminuisca rapidamente. Per un esperimento analogo con fermioni anziché con bosoni, la formula corretta per le fluttuazioni dei numeri di occupazione delle celle dello spazio delle fasi diventa la (5.77) anziché la (5.78), e con conti del tutto analoghi a quelli visti sopra, otteniamo

$$\overline{C} = -\frac{\overline{N_1}\overline{N_2}}{g}, \quad (8.113)$$

ci aspettiamo cioè che a fluttuazioni positive del numero di fermioni rivelati da R_1 corrispondano fluttuazioni negative del numero rivelato da R_2 (*antibunching*). Per particelle classiche, che seguono la statistica di Boltzmann, avremo invece

$$\overline{C} = 0. \quad (8.114)$$

Capitolo 9

Principi di funzionamento del laser

9.1 Coefficienti di Einstein e formula di Planck

Consideriamo un sistema di atomi identici in interazione con un campo di radiazione elettromagnetica. Supponiamo che ogni atomo abbia due livelli energetici $|1\rangle$ e $|2\rangle$, rispettivamente di energia E_1 ed E_2 , con $E_2 > E_1$. Secondo l'impostazione data da Einstein nel 1916 esistono tre processi distinti di interazione tra l'atomo e la radiazione, schematizzati nella Figura 9.1:

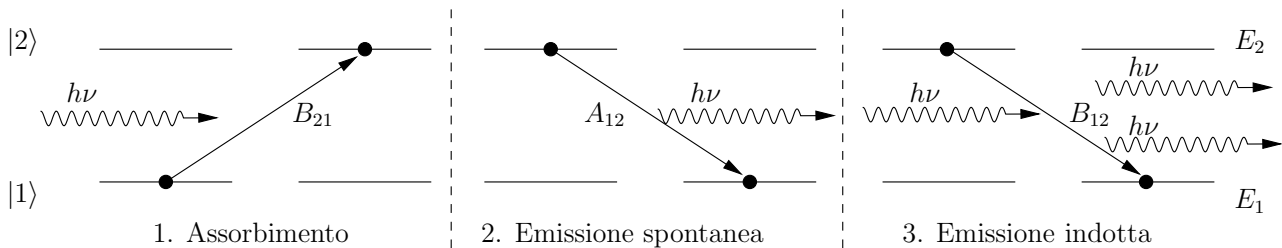


Figura 9.1 Processi di assorbimento, emissione spontanea e emissione indotta, e rispettivi coefficienti di Einstein.

1. Si ha *assorbimento* quando l'atomo si trova inizialmente nello stato di energia più bassa $|1\rangle$ ed assorbe un fotone di energia $h\nu = E_2 - E_1$ (conservazione dell'energia) portandosi allo stato più alto $|2\rangle$. Contemporaneamente scompare un fotone dal campo di radiazione.
2. Se l'atomo si trova inizialmente nel livello più alto $|2\rangle$ può decadere al livello $|1\rangle$ per *emissione spontanea* di un fotone di energia $h\nu$. Si tratta di un fenomeno statistico che genera un decadimento esponenziale della popolazione del livello $|2\rangle$. Si definisce vita media τ del livello $|2\rangle$ il tempo dopo il quale, per effetto dell'emissione spontanea, nel livello $|2\rangle$ è rimasta una frazione $1/e$ degli atomi che vi si trovavano inizialmente. Il fotone è emesso con uguale probabilità in qualunque direzione.
3. Esattamente come un fotone di energia $h\nu$ che incontra un atomo nello stato $|1\rangle$ può essere assorbito, un fotone che incontra un atomo nello stato $|2\rangle$ può stimolare l'emissione di un nuovo fotone, identico al fotone di partenza, portando così l'atomo allo stato più basso $|1\rangle$. Perché

avvenga questo fenomeno, detto *emissione indotta* o *emissione stimolata*, è quindi necessaria la presenza di fotoni “primari”. L’emissione indotta aggiunge un nuovo fotone a quelli già presenti nel campo di radiazione, il nuovo fotone è indistinguibile da quello che ne ha indotto l’emissione. Mentre l’assorbimento e l’emissione spontanea erano fenomeni già noti, l’emissione indotta è stata ipotizzata per la prima volta da Einstein.

Adesso supponiamo che il nostro sistema di atomi sia in equilibrio con della radiazione di corpo nero alla temperatura T . Chiamiamo $u(\nu, T) d\nu$ l’energia di radiazione per unità di volume con frequenza compresa tra ν e $\nu + d\nu$. Nel seguito, per brevità, scriveremo $u(\nu)$ invece di $u(\nu, T)$. Vediamo quanti dei tre processi descritti sopra avvengono nell’intervallo di tempo dt .

1. *Transizioni per assorbimento da $|1\rangle$ a $|2\rangle$* . La frequenza dei processi è proporzionale al numero di atomi N_1 che si trovano nello stato $|1\rangle$ e alla densità di radiazione alla frequenza $\nu = (E_2 - E_1)/h$

$$(dN_{2\leftarrow 1})_{\text{abs}} = B_{21} u(\nu) N_1 dt, \quad (9.1)$$

dove B_{21} è il coefficiente di Einstein per l’assorbimento, pari alla probabilità di transizione per unità di tempo e di densità di radiazione $u(\nu)$. Nei coefficienti di Einstein metteremo come primo indice quello relativo allo stato di arrivo, e come secondo quello relativo allo stato di partenza. Vedremo in seguito che questo è in accordo con la notazione matriciale, e con il normale prodotto righe per colonne delle matrici.

2. *Transizioni per emissione spontanea da $|2\rangle$ a $|1\rangle$* . Il numero di queste transizioni per unità di tempo è proporzionale al numero di atomi che si trova nello stato $|2\rangle$ e non dipende dalla densità di radiazione $u(\nu)$:

$$(dN_{1\leftarrow 2})_{\text{spont}} = A_{12} N_2 dt, \quad (9.2)$$

dove A_{12} è il coefficiente di Einstein per l’emissione spontanea.

3. *Transizioni per emissione indotta da $|2\rangle$ a $|1\rangle$* . Il numero di queste transizioni è proporzionale al numero di atomi N_2 che si trovano nello stato $|2\rangle$ e alla densità di radiazione alla frequenza $\nu = (E_2 - E_1)/h$

$$(dN_{1\leftarrow 2})_{\text{ind}} = B_{12} u(\nu) N_2 dt, \quad (9.3)$$

dove B_{12} è il coefficiente di Einstein per l’emissione indotta.

All’equilibrio N_1 ed N_2 devono essere costanti, quindi nell’intervallo di tempo dt il numero di transizioni da $|1\rangle$ a $|2\rangle$ deve essere uguale al numero complessivo di transizioni da $|2\rangle$ a $|1\rangle$:

$$(dN_{2\leftarrow 1})_{\text{abs}} = (dN_{1\leftarrow 2})_{\text{spont}} + (dN_{1\leftarrow 2})_{\text{ind}}. \quad (9.4)$$

Inserendo le espressioni delle (9.1), (9.2) e (9.3) otteniamo per il rapporto N_2/N_1

$$\frac{N_2}{N_1} = \frac{B_{21} u(\nu)}{A_{12} + B_{12} u(\nu)}. \quad (9.5)$$

D’altra parte, se siamo all’equilibrio termico, il rapporto N_2/N_1 è dato dalla ripartizione di Boltzmann

$$\frac{N_2}{N_1} = \frac{e^{-E_2/kT}}{e^{-E_1/kT}}, \quad (9.6)$$

da cui segue

$$\frac{B_{21} u(\nu)}{A_{12} + B_{12} u(\nu)} = \frac{e^{-E_2/kT}}{e^{-E_1/kT}}, \quad \text{da cui ricaviamo} \quad u(\nu) = \frac{A_{12}}{B_{21} e^{h\nu/kT} - B_{12}}, \quad (9.7)$$

dove abbiamo posto $h\nu = E_2 - E_1$. Per determinare le relazioni tra i coefficienti di Einstein A_{12} , B_{21} e B_{12} consideriamo prima il limite per la temperatura che tende all'infinito. Sappiamo che $\lim_{T \rightarrow \infty} u(\nu) = \infty$, e il denominatore dell'ultima delle (9.7) deve tendere a zero, per cui

$$B_{21} = B_{12}, \quad \text{da cui} \quad u(\nu) = \frac{A_{12}}{B_{21} (e^{h\nu/kT} - 1)}. \quad (9.8)$$

D'altra parte al limite delle basse frequenze, cioè per $h\nu \ll kT$, deve valere la legge di Rayleigh-Jeans, come è verificato sperimentalmente,

$$u(\nu) = \frac{8\pi\nu^2}{c^3} kT. \quad (9.9)$$

Se espandiamo fino al primo ordine l'esponenziale al denominatore della (9.8) otteniamo

$$u(\nu) = \frac{A_{12}}{B_{21} (e^{h\nu/kT} - 1)} \simeq \frac{A_{12}}{B_{21}} \frac{kT}{h\nu}, \quad (9.10)$$

e uguagliando alla (9.9) otteniamo

$$\frac{A_{12}}{B_{21}} = \frac{8\pi h\nu^3}{c^3}, \quad (9.11)$$

da cui otteniamo finalmente la formula di Planck

$$u(\nu) = \frac{8\pi h\nu^3}{c^3} \frac{1}{e^{h\nu/kT} - 1}.$$

Otteniamo poi la relazione tra A_{12} e $B_{21} = B_{12}$

$$A_{12} = \frac{8\pi h\nu^3}{c^3} B_{21}, \quad (9.12)$$

in accordo con la legge di Kirchhoff, secondo cui le probabilità di emissione spontanea e di assorbimento sono proporzionali tra loro.

La derivazione di Einstein della legge di Planck è fondamentale per capire l'interazione radiazione-materia e i principi di funzionamento del laser. Nel paragrafo 10.4 deriveremo i coefficienti di Einstein da considerazioni quantistiche sulle transizioni di dipolo elettrico.

9.2 Propagazione di un fascio di radiazione in un mezzo materiale

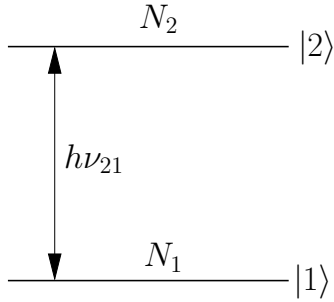
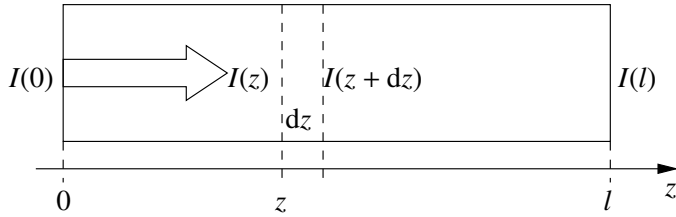


Figura 9.2 Sistema a due livelli.

Supponiamo di avere un fascio di radiazione elettromagnetica monocromatica di frequenza ν (quindi non radiazione di corpo nero!) che attraversi un mezzo materiale inizialmente all'equilibrio termico alla temperatura T , per esempio un gas atomico. Supponiamo che gli atomi del gas abbiano due livelli energetici $|1\rangle$ e $|2\rangle$, di energia rispettivamente E_1 ed E_2 , tali che $E_2 > E_1$ e che sia $h\nu = E_2 - E_1$. Ammettiamo anche che $|1\rangle$ e $|2\rangle$ possano essere degeneri, con degenerazione rispettivamente g_1 e g_2 . In assenza di interazione con il fascio di radiazione l'equilibrio termodinamico impone che i numeri N_1^0 ed N_2^0 di atomi per unità di volume che si trovano rispettivamente nei livelli energetici $|1\rangle$ e $|2\rangle$ stiano nel rapporto di Boltzmann

$$\frac{N_2^0}{N_1^0} = \frac{g_2}{g_1} e^{-(E_2-E_1)/kT}. \quad (9.13)$$

Consideriamo adesso l'interazione del mezzo materiale con il fascio di radiazione. Oltre che dalla frequenza ν , il fascio sarà caratterizzato dalla sua intensità I , definita come la potenza incidente per unità di superficie, e misurata in W/m^2 , e dal flusso di fotoni Φ , definito come il



numero di fotoni incidenti sull'unità di superficie nell'unità di tempo. Dato che il singolo fotone ha energia $h\nu$, le quantità I e Φ sono legate dalla relazione

$$I = \Phi h\nu. \quad (9.14)$$

Figura 9.3 Propagazione di radiazione in un mezzo materiale.

In termini della velocità della luce c , della densità di energia di radiazione per unità di volume u , e del numero di fotoni per unità di volume n , le quantità I e Φ si scrivono

$$I = cu, \quad \Phi = cn. \quad (9.15)$$

Supponiamo che la radiazione si propaghi lungo la direzione z entro un piccolo angolo solido $d\Omega$, come in Fig. 9.3. Passando dalla coordinata z alla coordinata $z + dz$ avremo una variazione dell'intensità del fascio dI e la corrispondente variazione del suo flusso di fotoni $d\Phi$. Queste variazioni sono determinate essenzialmente dai due processi di assorbimento e di emissione indotta. L'emissione spontanea, essendo isotropa, dà infatti un contributo trascurabile (dell'ordine di $d\Omega/4\pi$) all'interno dell'angolo solido $d\Omega$. L'equazione (9.1) può essere riscritta, in termini del flusso di fotoni Φ , nella forma

$$(dN_2)_{\text{abs}} = \sigma_{21} \Phi N_1 dt, \quad (dN_1)_{\text{abs}} = -\sigma_{21} \Phi N_1 dt, \quad (d\Phi)_{\text{abs}} = -\sigma_{21} \Phi N_1 c dt, \quad (9.16)$$

l'ultima relazione si ottiene tenendo conto che N_1 e N_2 sono numeri di atomi per unità di volume, mentre Φ è il numero di fotoni incidenti per unità di superficie e unità di tempo: dalla seconda delle

(9.15) abbiamo $d\Phi = c dn = c dN_1$ (l'emissione di un fotone comporta il passaggio di un atomo dallo stato $|2\rangle$ allo stato $|1\rangle$, l'assorbimento di un fotone il passaggio inverso). La quantità σ_{21} , che è legata al coefficiente di Einstein B_{21} , ha le dimensioni di una superficie ed è detta *sezione d'urto d'assorbimento*. Se consideriamo il numero di assorbimenti al secondo, abbiamo infatti dalle (9.1) e (9.16)

$$\left(\frac{dN_2}{dt}\right)_{\text{abs}} = \sigma_{21} \Phi N_1 = B_{21} u(\nu) N_1, \quad \text{da cui} \quad \sigma_{21} \Phi = B_{21} u(\nu). \quad (9.17)$$

D'altra parte, essendo $u(\nu) d\nu$ l'energia radiante per unità di volume con frequenza compresa tra ν e $\nu + d\nu$, per il nostro fascio possiamo scrivere

$$u(\nu) = \frac{nh\nu}{\Delta\nu}, \quad (9.18)$$

dove $\Delta\nu$ è la larghezza di riga del fascio stesso. Unendo alla seconda delle (9.15) abbiamo

$$\sigma_{21} \Phi = \sigma_{21} nc = B_{21} \frac{nh\nu}{\Delta\nu}, \quad \text{da cui otteniamo} \quad \sigma_{21} = B_{21} \frac{h\nu}{c \Delta\nu}, \quad (9.19)$$

che è la relazione tra la sezione d'urto d'assorbimento e il coefficiente B_{21} di Einstein.

L'equazione (9.3) può essere riscritta, sempre in termini del flusso di fotoni Φ ,

$$(dN_2)_{\text{ind}} = -\sigma_{12} \Phi N_2 dt, \quad (dN_1)_{\text{ind}} = \sigma_{12} \Phi N_2 dt, \quad (d\Phi)_{\text{ind}} = \sigma_{12} \Phi N_2 c dt, \quad (9.20)$$

dove σ_{12} è la *sezione d'urto d'emissione indotta*. Mettendo insieme le relazioni (9.16) e (9.20), abbiamo per la variazione complessiva del flusso di fotoni

$$d\Phi = (d\Phi)_{\text{ind}} + (d\Phi)_{\text{abs}} = (\sigma_{12}N_2 - \sigma_{21}N_1) \Phi c dt = (\sigma_{12}N_2 - \sigma_{21}N_1) \Phi dz, \quad (9.21)$$

nell'ultimo passaggio abbiamo chiamato $dz = c dt$ lo spazio percorso dai fotoni nell'intervallo di tempo dt . Il flusso dei fotoni, passando dalla posizione z alla posizione $z + dz$, subisce quindi una variazione

$$d\Phi = \alpha \Phi(z) dz, \quad (9.22)$$

dove $\alpha = (\sigma_{12}N_2 - \sigma_{21}N_1)$ è il *coefficiente di amplificazione* del mezzo. Se $\alpha < 0$, cioè se $\sigma_{12}N_2 < \sigma_{21}N_1$, il flusso viene smorzato esponenzialmente mentre attraversa il mezzo (legge di Lambert-Beer), questo è ciò che si realizza quando le popolazioni dei livelli energetici sono determinate dall'equilibrio termico. Se invece $\alpha > 0$, cioè $\sigma_{12}N_2 > \sigma_{21}N_1$, il flusso viene amplificato mentre attraversa il mezzo. Questo fenomeno è alla base dell'effetto *maser* (acronimo di *Microwave Amplification by Stimulated Emission of Radiation*) quando la frequenza della transizione tra $|1\rangle$ e $|2\rangle$ è nella regione delle microonde, e dell'effetto *laser* (acronimo di *Light Amplification by Stimulated Emission of Radiation*) quando la frequenza della transizione tra $|1\rangle$ e $|2\rangle$ è nella regione ottica (ma si parla anche di laser operanti nell'infrarosso o nell'ultravioletto). In realtà, sia il maser che il laser vengono usati non come amplificatori, ma come oscillatori. Esattamente come nei circuiti elettronici, un amplificatore di radiazione viene trasformato in oscillatore riinvandogli in ingresso parte della radiazione amplificata in uscita, come vedremo tra poco. Perché il flusso di radiazione sia amplificato attraversando il mezzo è necessario che sia presente un meccanismo che alteri la distribuzione Boltzmanniana. Se le condizioni sono tali per cui il coefficiente α , negativo o positivo che sia, è costante, la (9.22) può essere integrata ottenendo

$$\Phi(z) = \Phi(0) e^{\alpha z}, \quad \text{che porta anche a} \quad I(z) = I(0) e^{\alpha z}. \quad (9.23)$$

Le condizioni di equilibrio fra un sistema a due livelli e la radiazione elettromagnetica con cui interagisce, studiate nel paragrafo 9.1, richiedono che, se le degenerazioni g_1 e g_2 dei livelli $|1\rangle$ e $|2\rangle$ sono uguali, siano uguali anche la sezione d'urto di assorbimento e la sezione d'urto di emissione indotta

$$\sigma_{12} = \sigma_{21} = \sigma \quad \text{se} \quad g_1 = g_2. \quad (9.24)$$

9.3 Equazioni di bilancio per l'effetto laser

9.3.1 Inversione di popolazione

La (9.22) ci dice che, perché l'effetto laser sia possibile, dobbiamo avere

$$\alpha = (\sigma_{12}N_2 - \sigma_{21}N_1) > 0, \quad \text{o, se vale la (9.24), semplicemente} \quad (N_2 - N_1) > 0. \quad (9.25)$$

Per poter funzionare il laser richiede quindi un'*inversione di popolazione* rispetto alla la ripartizione di Boltzmann, visto che il livello $|2\rangle$ ha energia maggiore del livello $|1\rangle$. Per questo è necessaria la presenza di un meccanismo, detto di *pompaggio*, che aumenti la popolazione del livello $|2\rangle$ a spese di altri livelli, spesso a spese del livello fondamentale dell'atomo, che è il più popolato. In altre parole, vengono “pompati” atomi in $|2\rangle$ per rendere N_2 maggiore di N_1 .

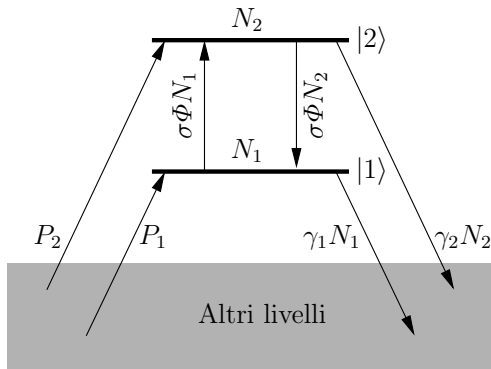


Figura 9.4 Accoppiamento dei livelli laser $|1\rangle$ e $|2\rangle$ con gli altri livelli dell'atomo.

Il processo di pompaggio richiede lavoro da parte dell'esterno. La figura 9.4 rappresenta lo schema generale di un *sistema laser*. Il pompaggio porta atomi nel livello $|2\rangle$ sottraendoli ad altri livelli non rappresentati in figura, e, come effetto collaterale indesiderato, può portare atomi anche nel livello $|1\rangle$. I *termini di sorgente* P_2 e P_1 dicono quanti atomi per unità di volume e unità di tempo vengono portati, rispettivamente, nei due livelli. Questo trasferimento di atomi può essere ottenuto per assorbimento di luce (pompaggio ottico), o per eccitazione da parte di una scarica elettrica, o per altra via. Inoltre i livelli $|1\rangle$ e $|2\rangle$ si spopolano, cioè, perderanno a loro volta atomi verso altri livelli non rappresentati in figura, tramite *processi di*

rilassamento, normalmente non radiativi ma dovuti, per esempio, a collisioni degli atomi tra loro o con le pareti della cella che li contiene. Il numero di atomi persi da ogni livello, per unità di volume e unità di tempo a causa del rilassamento, saranno proporzionali alla popolazione del livello stesso, N_1 o N_2 , tramite i coefficienti di rilassamento γ_1 e γ_2 . Infine l'interazione con la radiazione laser causa i trasferimenti $\sigma\Phi N_1$ e $\sigma\Phi N_2$ da $|1\rangle$ a $|2\rangle$ e viceversa, secondo quanto visto nel paragrafo 9.2. Avremo quindi per il bilancio di N_2

$$\frac{dN_2}{dt} = P_2 - \gamma_2 N_2 - \sigma\Phi(N_2 - N_1), \quad (9.26)$$

e per il bilancio di N_1

$$\frac{dN_1}{dt} = P_1 - \gamma_1 N_1 + \sigma\Phi(N_2 - N_1). \quad (9.27)$$

Facendo, per semplicità, l'ipotesi che sia $\gamma_1 = \gamma_2 = \gamma$, otteniamo per $\Delta N = N_2 - N_1$

$$\frac{d\Delta N}{dt} = (P_2 - P_1) - 2\sigma\Phi\Delta N - \gamma\Delta N. \quad (9.28)$$

Chiamando

$$\Delta N_0 = \frac{P_2 - P_1}{\gamma} \quad (9.29)$$

la differenza di popolazione in assenza di effetto laser otteniamo

$$\frac{d\Delta N}{dt} = \gamma(\Delta N_0 - \Delta N) - 2\sigma\Phi\Delta N, \quad (9.30)$$

dove il primo termine del secondo membro è un termine di rilassamento che tende a riportare la differenza di popolazione al suo valore di equilibrio in assenza di effetto laser ΔN_0 , mentre il secondo termine descrive l'accoppiamento tra il flusso dei fotoni e la differenza di popolazione. Per descrivere con maggior dettaglio l'evoluzione del flusso dei fotoni, dobbiamo fare ipotesi sulla configurazione sperimentale del nostro laser.

9.3.2 Equazioni di bilancio per i fotoni e fattore di qualità

Discutiamo qui le *equazioni di bilancio* del laser, dette anche, all'inglese, *rate equations*. Esistono varie configurazioni di laser, alle quali accenneremo in seguito. La nostra configurazione di base, schematizzata in fig. 9.5, consiste essenzialmente in un tubo che contiene il mezzo attivo, situato tra due specchi che formano una cavità risonante per la radiazione. Dei due specchi, uno è totalmente riflettente (riflettività $R_1 \approx 1.0$), mentre l'altro è solo parzialmente riflettente (normalmente si ha comunque $R_2 \gtrsim 0.95$) per permettere la fuoriuscita di una piccola parte della radiazione. Dei fotoni presenti nella cavità, ci aspettiamo che quelli coinvolti nel processo laser siano quelli assiali perché, riflettendosi avanti e indietro sugli specchi, continuano ad attraversare e riattraversare la cavità, subendo a ogni passaggio un'amplificazione data dalla (9.23), finché non si giunge alla saturazione. Infatti l'amplificazione esponenziale del fascio di radiazione, descritta dalla (9.23), non può ovviamente continuare all'infinito, ma si raggiungono le condizioni in cui il mezzo si comporta da oscillatore anziché da amplificatore.

Come al solito, la transizione laser avviene tra i due livelli $|1\rangle$ e $|2\rangle$, con energie E_1 ed E_2 tali che $E_1 < E_2$, e popolazioni per unità di volume N_1 e N_2 . Chiamiamo n il numero per unità di volume dei fotoni con energia $h\nu = E_2 - E_1$ presenti nella cavità. Il numero dei fotoni n aumenterà per ogni processo di emissione indotta, la cui frequenza è proporzionale sia a n che a N_2 , e diminuirà per ogni assorbimento, la cui frequenza è proporzionale sia a n che a N_1 . Inoltre, indipendentemente da questi processi, ogni fotone non resterà indefinitamente nella cavità, ma avrà un tempo medio di soggiorno τ_s dovuto sia alla probabilità di attraversare lo specchio parzialmente riflettente che a quella di perdersi per processi di diffusione, di diffrazione o di assorbimento. L'equazione di bilancio per il numero di fotoni si scrive così

$$\frac{dn}{dt} = W(N_2 - N_1)n - \frac{n}{\tau_s}, \quad (9.31)$$

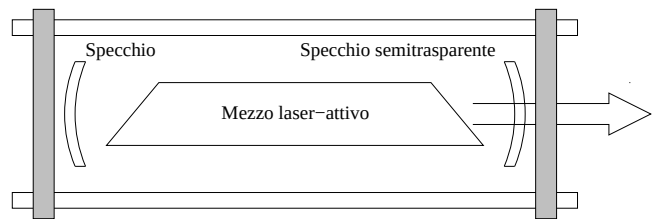


Figura 9.5 Tipica configurazione di un laser.

dove il contributo dell'emissione spontanea è stato trascurato per quanto detto prima. La probabilità W può essere ricavata dalla trattazione di Einstein del paragrafo 9.1. Dalla (9.12) abbiamo infatti

$$B_{21} = \frac{c^3}{8\pi h\nu^3} A_{12} = \frac{c^3}{8\pi h\nu^3 \tau}, \quad (9.32)$$

dove $\tau = 1/A_{12}$ è la vita media del livello $|2\rangle$. Sostituendo, per esempio, nella (9.3), abbiamo

$$WnN_2 = \left(\frac{dN_{12}}{dt} \right)_{\text{ind}} = B_{12} u(\nu) N_2 = \frac{c^3}{8\pi h\nu^3} A_{12} u(\nu) N_2 = \frac{c^3}{8\pi h\nu^3 \tau} u(\nu) N_2. \quad (9.33)$$

Se sostituiamo $u(\nu)$ con la sua espressione (9.18) nella (9.33) e dividiamo a destra e sinistra per il prodotto nN_2 otteniamo

$$W = \frac{c^3}{8\pi h\nu^3 \tau} \frac{h\nu}{\Delta\nu} = \frac{c^3}{8\pi\nu^2 \tau \Delta\nu}. \quad (9.34)$$

Se vogliamo che il numero di fotoni nella cavità non diminuisca deve essere

$$\frac{dn}{dt} = W(N_2 - N_1)n - \frac{n}{\tau_s} = \frac{c^3}{8\pi\nu^2 \tau \Delta\nu} (N_2 - N_1)n - \frac{n}{\tau_s} \geq 0, \quad (9.35)$$

da cui otteniamo la minima *inversione di popolazione* necessaria perché il laser possa operare

$$N_2 - N_1 \geq \Delta N_{\min} = \frac{8\pi\nu^2 \tau \Delta\nu}{c^3 \tau_s}. \quad (9.36)$$

Per avere l'effetto laser dovremo quindi avere una *inversione di popolazione*, cioè, contrariamente a quanto accade all'equilibrio termodinamico, dovrà essere $N_2 > N_1$. Inoltre la differenza di popolazione $\Delta N = N_2 - N_1$ dovrà essere superiore a un certo valore di soglia $\Delta N_{\min}$, fissato dalla (9.36). Un'occhiata alla (9.36) ci dice che l'effetto laser è tanto più difficile da raggiungere quanto più è grande la frequenza ν . Gli atomi laser-attivi saranno tanto più efficaci quanto più è piccola la larghezza di riga $\Delta\nu$. Inoltre dovrà essere il più lungo possibile il tempo di soggiorno dei fotoni nella cavità τ_s . I fotoni vengono persi dalla cavità sia per la non completa riflettività dei due specchi, sia per fenomeni di diffrazione o, comunque, per altre perdite durante il viaggio tra uno specchio e l'altro. Chiamiamo R_1 e R_2 le riflettività dei due specchi, ed F la frazione di fotoni persi in un viaggio tra i due specchi. Chiamiamo $n(0)$ il numero di fotoni nella cavità all'istante $t = 0$, ed escludiamo per il momento gli effetti dovuti al mezzo attivo. Dopo un tempo $t_1 = 2L/c$, dove L è la distanza tra i due specchi, i fotoni "superstiti" avranno fatto un *giro completo* della cavità, cioè percorso due volte la sua lunghezza, e il loro numero sarà

$$n(t_1) = R_1 R_2 (1 - F)^2 n(0), \quad (9.37)$$

dove il fattore $(1 - F)$ compare al quadrato perché F corrisponde alla frazione di fotoni persi percorrendo una sola volta la lunghezza L . Dopo m giri completi, effettuati in un tempo

$$t_m = m \frac{2L}{c}, \quad (9.38)$$

il numero di fotoni superstiti sarà

$$n(t_m) = [R_1 R_2 (1 - F)^2]^m n(0) = [R_1 R_2 (1 - F)^2]^{c t_m / (2L)} n(0). \quad (9.39)$$

In un tempo piccolo da un punto di vista “antropomorfo” avremo comunque un grandissimo numero di viaggi andata e ritorno dei fotoni tra i due specchi. Quindi, con buona approssimazione, possiamo considerare la (9.39) valida anche per un tempo generico t non multiplo intero di t_1 , e scrivere

$$n(t) = n(0) e^{-t/\tau_s}, \quad (9.40)$$

dove τ_s è il tempo di soggiorno dei fotoni nella cavità, pari a

$$\tau_s = -\frac{2L}{c} \frac{1}{\ln [R_1 R_2 (1 - F)^2]}, \quad (9.41)$$

notare che $\tau_s > 0$ perché il logaritmo di una quantità minore di 1 è negativo. Supponiamo adesso di avere un'onda elettromagnetica di frequenza angolare ω che viaggia tra i due specchi della cavità. Dato il tempo di soggiorno finito dei fotoni τ_s , il campo elettrico in un punto della cavità avrà la forma

$$E(t) = E_0 e^{-i\omega t - t/(2\tau_s)}, \quad (9.42)$$

dove il fattore 2 che moltiplica τ_s è dovuto al fatto che τ_s è il tempo di decadimento della potenza (numero di fotoni), proporzionale al quadrato del campo elettrico. In corrispondenza al tempo di soggiorno finito, secondo la (8.62), lo spettro di potenza della radiazione avrà una forma di riga lorentziana di semilarghezza

$$\Delta\nu_s = \frac{1}{2\pi\tau_s}. \quad (9.43)$$

Per ogni sistema risonante, in particolare per una cavità ottica, si definisce il *fattore di qualità* Q come

$$Q = 2\pi \frac{U_{\text{cav}}}{\Delta U_{\text{ciclo}}}, \quad (9.44)$$

dove U_{cav} è l'energia totale immagazzinata nella cavità, e ΔU_{ciclo} è l'energia persa in un ciclo di oscillazione. In altre parole, il fattore di qualità di una cavità è tanto maggiore quanto meno la cavità perde energia. Nel nostro caso

$$U_{\text{cav}} = nh\nu, \quad \Delta U_{\text{ciclo}} = h\nu \left(-\frac{dn}{dt} \right) \frac{1}{\nu} = h\nu \frac{n}{\tau_s} \frac{1}{\nu}, \quad (9.45)$$

dove abbiamo usato la (9.40). Mettendo insieme abbiamo

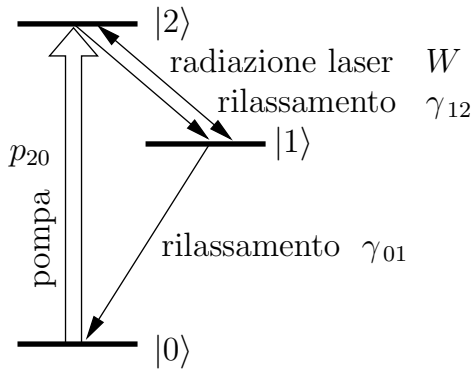
$$Q = 2\pi nh\nu \frac{1}{h\nu} \frac{\tau_s}{n} \nu = 2\pi\nu\tau_s = \frac{\nu}{\Delta\nu_s}, \quad (9.46)$$

dove abbiamo usato la (9.43). Così il fattore Q di una cavità risonante, relativo a un dato modo di oscillazione, può essere interpretato come il rapporto tra la frequenza di risonanza del modo stesso, ν , e la corrispondente larghezza di riga $\Delta\nu_s$.

9.3.3 Equazioni di bilancio per le densità di popolazione dei livelli

Per descrivere completamente la dinamica dell'emissione laser abbiamo bisogno anche, oltre che della (9.31), di equazioni analoghe per le popolazioni dei livelli $|1\rangle$ e $|2\rangle$. E' da notare che non esistono meccanismi di pompaggio per mantenere l'inversione di popolazione tra $|1\rangle$ e $|2\rangle$ che non coinvolgano

almeno un altro livello del mezzo laser-attivo. Questo non esclude la possibilità di *laser a due livelli*, per esempio i *laser chimici*, in cui le molecole laser-attive vengono prodotte direttamente nello stato eccitato. Esistono due schemi possibili di laser a tre livelli, di cui uno è rappresentato nella figura 9.6. Qui è presente un meccanismo di pompaggio che trasferisce gli atomi dal livello fondamentale $|0\rangle$ al livello superiore della transizione laser $|2\rangle$. Esistono poi meccanismi di rilassamento, normalmente non radiativi, che ridistribuiscono gli atomi tra i vari livelli. Per semplicità escluderemo processi di rilassamento che portino atomi da livelli più bassi a livelli più alti, o che riportino atomi dal livello $|2\rangle$ direttamente al livello $|0\rangle$. Questi processi possono facilmente essere inclusi nelle equazioni se necessari, portano solo a conti un po' più lunghi. Per il bilancio della popolazione del livello $|2\rangle$ avremo così



$$\frac{dN_2}{dt} = \underbrace{-Wn(N_2 - N_1)}_{\text{interazione radiazione laser}} + \underbrace{p_{20}N_0}_{\text{pompaggio}} - \underbrace{\gamma_{12}N_2}_{\text{rilassamento}}, \quad (9.47)$$

dove il primo termine descrive la variazione di N_2 dovuta all'interazione con la radiazione coerente del laser, il secondo l'aumento di N_2 a spese della popolazione N_0 del livello fondamentale tramite il meccanismo di pompaggio, e il terzo il decadimento di N_2 dovuto ai processi di rilassamento.

Figura 9.6 Primo schema di laser a tre livelli.

Per il livello $|1\rangle$ avremo, in modo analogo,

$$\frac{dN_1}{dt} = \underbrace{Wn(N_2 - N_1)}_{\text{interaz. radiazione laser}} + \underbrace{\gamma_{12}N_2}_{\text{rilass. da } |2\rangle} - \underbrace{\gamma_{01}N_1}_{\text{rilass. verso } |0\rangle}, \quad (9.48)$$

e per $|0\rangle$

$$\frac{dN_0}{dt} = \underbrace{-p_{20}N_0}_{\text{pompaggio}} + \underbrace{\gamma_{01}N_1}_{\text{rilass. da } |1\rangle}. \quad (9.49)$$

Per semplificare ulteriormente le cose, faremo ancora le ipotesi che: i) la transizione di rilassamento da $|1\rangle$ a $|0\rangle$ sia molto rapida, in modo che N_1 sia praticamente nulla; ii) il pompaggio, il rilassamento e l'emissione laser avvengano in condizioni tali che sia sempre comunque $N_0 \gg N_2$, in modo che N_0 sia praticamente costante. In queste condizioni la (9.47) può essere riscritta

$$\frac{dN_2}{dt} = -WnN_2 + p_{20}N_0 - \gamma_{12}N_2 = -WnN_2 + \gamma_{12}(N_2^{\text{eq}} - N_2), \quad (9.50)$$

dove $N_2^{\text{eq}} = \frac{p_{20}}{\gamma_{12}}N_0$ è il valore di equilibrio di N_2 , in presenza dei soli processi di pompaggio e rilassamento, ma in assenza di interazione con la radiazione laser. Nelle condizioni descritte sopra l'equazione di bilancio per i fotoni (9.35) diventa

$$\frac{dn}{dt} = WN_2n - \frac{n}{\tau_s}. \quad (9.51)$$

Se il laser opera in condizioni stazionarie (*laser continuo*) avremo

$$\frac{dN_2}{dt} = 0, \quad \frac{dn}{dt} = 0, \quad (9.52)$$

in queste condizioni dalla (9.51) otteniamo

$$N_2 = \frac{1}{W\tau_s} = N_2^{\text{cr}}, \quad (9.53)$$

esiste cioè una popolazione critica N_2^{cr} , indipendente da p_{20} , che non aumenta più anche se viene aumentata l'intensità di pompaggio. L'energia fornita in più per il pompaggio va quindi prevalentemente ad aumentare la densità dei fotoni, come si vede risolvendo la (9.50) per n

$$n = \frac{\gamma_{12}}{W} \left(\frac{N_2^{\text{eq}}}{N_2^{\text{cr}}} - 1 \right). \quad (9.54)$$

Un aumento della velocità di pompaggio p_{20} con cui gli atomi vengono trasferiti da $|0\rangle$ a $|2\rangle$ aumenta la popolazione di equilibrio in assenza di effetto laser N_2^{eq} del livello $|2\rangle$, ma non N_2^{cr} . Corrispondentemente, aumenta la densità dei fotoni presenti nella cavità n attraverso la (9.54). Dato che n non può essere negativo, l'emissione laser inizia solo quando N_2^{eq} raggiunge il valore di soglia N_2^{cr} .

9.3.4 Altri possibili schemi di pompaggio

Oltre allo schema di pompaggio rappresentato in Fig. 9.6 per un laser a tre livelli è possibile anche lo schema rappresentato in Fig. 9.7, in cui l'inversione di popolazione viene realizzata tra i livelli $|1\rangle$ e $|0\rangle$. Se la transizione tra livelli $|0\rangle$ e $|1\rangle$ è nella regione spettrale delle microonde, l'emissione spontanea da $|1\rangle$ a $|0\rangle$ può essere un processo non permesso per dipolo, quindi con probabilità sufficientemente bassa per essere trascurata. Inviando sul mezzo radiazione in risonanza con la transizione $|2\rangle \leftrightarrow |0\rangle$ gli atomi vengono pompati dal livello $|0\rangle$ al livello $|2\rangle$. Di qui possono decadere, per rilassamento o per emissione spontanea, verso i due livelli sottostanti. Mentre gli atomi che decadono nel livello $|1\rangle$ vi restano intrappolati, quelli che decadono in $|0\rangle$ possono nuovamente essere pompati in $|2\rangle$. La soluzione stazionaria in presenza di radiazione di pompaggio è quindi quella con tutti gli atomi in $|1\rangle$, con inversione di popolazione per la coppia di livelli $|1\rangle$ e $|0\rangle$. Ovviamente, nel caso reale tale inversione è limitata dalla presenza di altri processi di rilassamento, per deboli che siano. In questo caso il pompaggio ottico può essere realizzato anche con una radiazione di pompa molto debole, come una lampada spettrale.

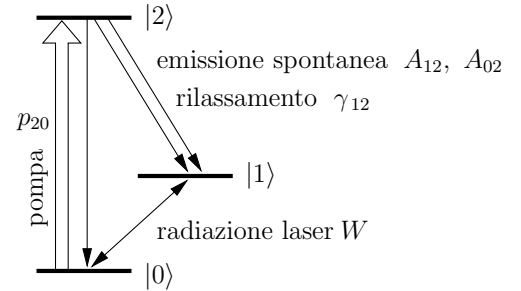


Figura 9.7 Secondo schema di laser a tre livelli.

In Fig. 9.8 è invece rappresentato lo schema di un laser a quattro livelli. Gli atomi vengono trasferiti per pompaggio dal livello fondamentale $|0\rangle$ al livello eccitato $|3\rangle$. Di qui decadono per processi non radiativi con velocità γ_{23} al livello superiore della transizione laser $|2\rangle$. Il livello inferiore della transizione laser $|1\rangle$ è svuotato a sua volta da processi non radiativi con velocità γ_{01} . La differenza di energia tra i livelli $|1\rangle$ e $|0\rangle$ è molto maggiore di kT , così che i livelli $|1\rangle$, $|2\rangle$ e $|3\rangle$ sono praticamente

spopolati all'equilibrio termodinamico. Lo schema laser a quattro livelli è quello più frequente nei laser molecolari a pompaggio ottico. In questo caso il pompaggio avviene normalmente per assorbimento di radiazione emessa da un altro laser, detto *laser di pompa*, a sua volta eccitato da una scarica elettrica.

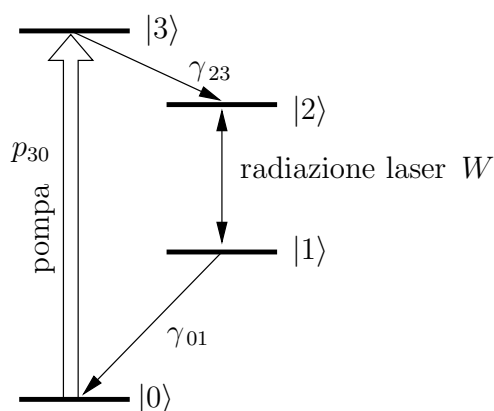


Figura 9.8 Schema di laser a quattro livelli.

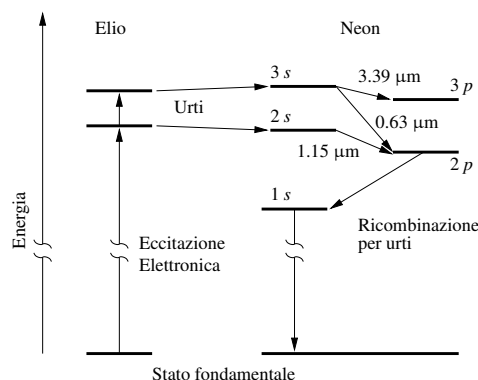


Figura 9.9 Schema di laser a elio-neon.

I laser eccitati da una scarica elettrica hanno il vantaggio di non richiedere l'uso di un laser di pompa per l'eccitazione. Normalmente la scarica avviene in una miscela di gas. Gli elettroni della scarica eccitano gli atomi di uno dei gas in uno o più stati eccitati, da qui l'energia viene trasferita per collisioni agli atomi del gas laser-attivo, che si trovano così nello stato superiore della transizione laser. Spesso è presente un terzo gas, che, per processi collisionali, sottrae l'energia agli atomi del gas laser-attivo che si trovano nello stato inferiore della transizione laser.

La figura 9.9 rappresenta lo schema dei processi di eccitazione e ricombinazione in un laser a elio-neon. Gli atomi di He vengono eccitati negli stati 1S e 3S dagli urti con gli elettroni della scarica. Le energie di eccitazione di questi stati sono in corrispondenza delle energie degli stati $2s$ e $3s$ del Ne, per cui l'energia può facilmente essere trasferita dall'He al Ne per collisioni atomiche. Da qui gli atomi di Ne possono decadere agli stati $3p$ e $2p$ con transizioni che possono essere laser-attive. Dagli stati inferiori delle transizioni laser gli atomi possono decadere allo stato $1s$ mediante emissione di un ulteriore fotone. Finalmente l'atomo di Ne perde l'energia dello stato $1s$ sotto forma di energia cinetica per collisioni con altri atomi, tornando allo stato fondamentale. Ovviamente, lo schema dei livelli energetici dell'elio e del neon rappresentati in Fig. 9.9 è molto semplificato. Si noti anche che

le energie dei livelli fondamentali sono state scelte come energie di riferimento. La riga più nota del laser a HeNe è la riga rossa a circa 633 nm, sulla quale il laser emette normalmente potenze comprese tra 1 e 100 mW.

Uno dei laser eccitati a scarica elettrica più usati è il laser a CO_2 . E' anche uno di laser più efficienti: la potenza in uscita può infatti raggiungere il 20% della potenza di pompaggio.

Anche in questo caso la scarica di eccitazione avviene in una miscela di gas: circa il 10 – 20 % di CO_2 , circa il 10 – 20 % di N_2 , e il resto He. Gli urti con gli elettroni eccitano un modo vibrazionale dell'azoto. Essendo l'azoto una molecola biatomica omonucleare, non può perdere questa energia per emissione di un fotone, così che i suoi stati vibrazionali eccitati sono metastabili. L'energia può però essere trasferita per collisioni a molecole di CO_2 , con efficienza sufficiente a raggiungere l'inversione di popolazione necessaria per l'operazione del laser. Lo stato inferiore della transizione laser viene poi vuotato per trasferimento collisionale di energia agli atomi di elio.

Il laser a CO_2 emette nell'infrarosso, con due bande centrate rispettivamente attorno a 9.4 e 10.6 μm .

9.4 Onde parassiali

9.4.1 Equazione di Helmholtz e approssimazione parassiale

L'equazione delle onde per una qualunque componente $E(\mathbf{r}, t)$ del vettore campo elettrico nel vuoto ha la forma scalare

$$\left(\Delta - \frac{1}{c^2} \frac{\partial^2}{\partial t^2} \right) E(\mathbf{r}, t) = 0. \quad (9.55)$$

Facciamo l'ipotesi che la funzione $E(\mathbf{r}, t)$ possa essere fattorizzata nel prodotto di due funzioni

$$E(\mathbf{r}, t) = \tilde{E}(\mathbf{r}) T(t), \quad (9.56)$$

dove $\tilde{E}(\mathbf{r})$ dipende solo dalla posizione, e $T(t)$ solo dal tempo. In questo caso la (9.55) diventa

$$T(t) \Delta \tilde{E}(\mathbf{r}) - \tilde{E}(\mathbf{r}) \frac{1}{c^2} \frac{\partial^2}{\partial t^2} T(t) = 0, \quad (9.57)$$

che, dividendo entrambi i membri per il prodotto $\tilde{E}(\mathbf{r}) T(t)$, diventa a sua volta

$$\frac{\Delta \tilde{E}(\mathbf{r})}{\tilde{E}(\mathbf{r})} = \frac{1}{c^2 T(t)} \frac{\partial^2}{\partial t^2} T(t), \quad (9.58)$$

dove il primo membro dipende solo dalla posizione $\mathbf{r}$, e il secondo solo dal tempo t . La (9.58) è quindi valida nel caso generale solo se entrambi i membri sono costanti, e uguali ad una stessa costante, che possiamo chiamare $-k^2$ senza perdita di generalità. Abbiamo così il sistema

$$\frac{\Delta \tilde{E}(\mathbf{r})}{\tilde{E}(\mathbf{r})} = -k^2, \quad \frac{1}{c^2 T(t)} \frac{\partial^2}{\partial t^2} T(t) = -k^2. \quad (9.59)$$

Moltiplicata per $\tilde{E}(\mathbf{r})$, la prima equazione diventa

$$(\Delta + k^2) \tilde{E}(\mathbf{r}) = 0, \quad (9.60)$$

che è una forma *indipendente dal tempo* dell'equazione delle onde, detta equazione di Helmholtz, in cui k è il vettore d'onda. La soluzione dell'equazione di Helmholtz dipende dalle condizioni al contorno. La seconda equazione, moltiplicata per $T(t)$, e avendo posto $\omega = kc$, diventa

$$\left(\frac{d^2}{dt^2} + \omega^2 \right) T(t) = 0, \quad (9.61)$$

la cui soluzione è una combinazione lineare di seni e coseni con frequenza angolare ω , dipendente dalle condizioni iniziali.

Nella (9.60) $\tilde{E}(\mathbf{r}) \equiv \tilde{E}(x, y, z)$ è il *fasore* (ampiezza complessa) di un campo che dipende sinusoidalmente dal tempo. Parlando di laser, ci occuperemo di fasci ottici che si propagano prevalentemente lungo una direzione, che sceglieremo come asse z (fasci parassiali). In queste condizioni ci conviene riscrivere la dipendenza spaziale di $\tilde{E}(x, y, z)$ nella forma

$$\tilde{E}(x, y, z) \equiv \tilde{u}(x, y, z) e^{-iz}, \quad (9.62)$$

dove k è il vettore d'onda, e $\tilde{u}$ è un'ampiezza scalare complessa che descrive il profilo trasversale del fascio. Sostituendo nella (9.60) otteniamo la seguente equazione per $\tilde{u}$

$$\frac{\partial^2 \tilde{u}}{\partial x^2} + \frac{\partial^2 \tilde{u}}{\partial y^2} + \frac{\partial^2 \tilde{u}}{\partial z^2} - 2i \frac{\partial \tilde{u}}{\partial z} = 0. \quad (9.63)$$

Una volta separato l'esponentiale e^{-iz} , la dipendenza residua di $\tilde{u}(x, y, z)$ da z è dovuta essenzialmente a effetti di diffrazione. Questa dipendenza, in generale, sarà lenta non solo rispetto alla lunghezza d'onda, ma anche rispetto alle variazioni trasversali dovute alla larghezza finita del fascio. Cioè, matematicamente, ci aspettiamo che sia

$$\left| \frac{\partial^2 \tilde{u}}{\partial z^2} \right| \ll \left| 2k \frac{\partial \tilde{u}}{\partial z} \right|, \quad \left| \frac{\partial^2 \tilde{u}}{\partial z^2} \right| \ll \left| \frac{\partial^2 \tilde{u}}{\partial x^2} \right|, \quad \text{e} \quad \left| \frac{\partial^2 \tilde{u}}{\partial z^2} \right| \ll \left| \frac{\partial^2 \tilde{u}}{\partial y^2} \right|. \quad (9.64)$$

Entro queste approssimazioni possiamo riscrivere la (9.63) in *approssimazione parassiale* trascurando la derivata seconda di $\tilde{u}$ rispetto a z

$$\frac{\partial^2 \tilde{u}}{\partial x^2} + \frac{\partial^2 \tilde{u}}{\partial y^2} - 2i \frac{\partial \tilde{u}}{\partial z} = 0, \quad (9.65)$$

questa è l'*equazione parassiale di Helmholtz* in coordinate cartesiane. La (9.65) viene anche scritta

$$\Delta_t \tilde{u}(s, z) - 2i \frac{\partial \tilde{u}(s, z)}{\partial z} = 0, \quad (9.66)$$

dove Δ_t sta per la parte trasversale (perpendicolare a z) del laplaciano, ed s sta per l'insieme delle due coordinate trasversali, normalmente x e y se usiamo coordinate cartesiane, oppure r e φ se usiamo coordinate cilindriche. In coordinate cilindriche abbiamo

$$\Delta_t \tilde{u}(r, \varphi, z) = \frac{1}{r} \frac{\partial}{\partial r} \left(r \frac{\partial \tilde{u}}{\partial r} \right) + \frac{1}{r^2} \frac{\partial^2 \tilde{u}}{\partial \varphi^2}. \quad (9.67)$$

9.4.2 Fasci gaussiani

Consideriamo un fascio di luce che si propaghi lungo l'asse z di un sistema di coordinate cilindriche (r, φ, z) . Il fascio è detto *gaussiano* se la sua intensità, a z costante, ha un andamento gaussiano con r . L'importanza dei fasci gaussiani è dovuta al fatto che descrivono con ottima approssimazione il comportamento della radiazione emessa da gran parte delle sorgenti laser. Questo, ovviamente, non è il caso più generale di radiazione laser, ma è l'unico che tratteremo in questo contesto. Come esempio semplice di fascio gaussiano, consideriamo un fascio laser nel modo TEM_{00} , che ha simmetria cilindrica attorno all'asse di propagazione z , e quindi è indipendente da φ . Per il modo TEM_{00} facciamo l'ipotesi che $\tilde{u}$ possa essere scritto

$$\tilde{u}(r, z) = \tilde{E}_0 e^{-iP(z)} e^{-ir^2/2q(z)}, \quad (9.68)$$

dove $\tilde{E}_0$ è una costante, k è il numero d'onda, e $P(z)$ e $q(z)$ sono due funzioni complesse di z . Di queste, $P(z)$ determina l'andamento longitudinale di $\tilde{u}(r, z)$, e $q(z)$ ne determina l'andamento trasversale. Per

inserire questa $\tilde{u}(r, z)$ nella (9.66), cominciamo calcolando i due termini

$$\begin{aligned}\frac{\partial \tilde{u}}{\partial z} &= \left(-i \frac{dP}{dz} + \frac{ir^2}{2q^2} \frac{dq}{dz} \right) \tilde{u}, \quad \text{e} \\ \Delta_t \tilde{u} &= -\frac{2i}{q} \tilde{u} - \frac{k^2 r^2}{q^2} \tilde{u},\end{aligned}\tag{9.69}$$

e la (9.66) diventa

$$-\frac{2i}{q} \tilde{u} - \frac{k^2 r^2}{q^2} \tilde{u} - 2i \left(-i \frac{dP}{dz} + \frac{ir^2}{2q^2} \frac{dq}{dz} \right) \tilde{u} = 0.\tag{9.70}$$

Dividendo primo e secondo membro per $\tilde{u}(r, z)$ otteniamo

$$-\frac{2i}{q} - \frac{k^2 r^2}{q^2} - 2k \frac{dP}{dz} + \frac{k^2 r^2}{q^2} \frac{dq}{dz} = 0.\tag{9.71}$$

La (9.71) deve essere valida per ogni valore di r , quindi è equivalente a due equazioni simultanee riguardanti, rispettivamente, la parte che contiene r e quella che non la contiene, cioè

$$\frac{dq}{dz} = 1 \quad \text{e} \quad \frac{dP}{dz} = -\frac{i}{q}.\tag{9.72}$$

La prima delle (9.72) ha come integrale semplicemente

$$q(z) = z + q_0,\tag{9.73}$$

dove q_0 è una costante complessa di integrazione. Se scriviamo $q_0 = z_0 + iz_R$, con z_0 e z_R (lunghezza di Rayleigh) reali, vediamo che la parte reale z_0 corrisponde semplicemente ad uno shift della coordinata z . Quindi, con un'opportuna traslazione dell'origine, possiamo scegliere un sistema di riferimento in cui $z_0 = 0$, e $q_0 = iz_R$ è un immaginario puro. Abbiamo così

$$\frac{1}{q(z)} = \frac{1}{z + iz_R} = \frac{z}{z^2 + z_R^2} - i \frac{z_R}{z^2 + z_R^2}.\tag{9.74}$$

La seconda delle (9.72), moltiplicata per i , ci dà invece

$$iP(z) = \int_0^z \frac{dz'}{q(z')} = \int_0^z \frac{dz'}{z' + iz_R} = \left[\ln(z' + iz_R) \right]_0^z = \ln \left(1 - i \frac{z}{z_R} \right).\tag{9.75}$$

L'argomento (complesso) del logaritmo può essere riscritto, in funzione di ampiezza e fase,

$$1 - i \frac{z}{z_R} = \sqrt{1 + \left(\frac{z}{z_R} \right)^2} e^{i\psi(z)},\tag{9.76}$$

dove l'angolo $\psi(z)$ è detto *sfasamento di Gouy* (Louis Georges Gouy), e vale

$$\psi(z) = -\arctan \left(\frac{z}{z_R} \right).\tag{9.77}$$

Inserendo tutto nella (9.68) otteniamo

$$\tilde{u}(r, z) = \frac{\tilde{E}_0}{\sqrt{1 + \left(\frac{z}{z_R}\right)^2}} e^{i\psi(z)} e^{-kz_R r^2 / 2(z^2 + z_R^2)} e^{-ir^2 / 2(z^2 + z_R^2)}. \quad (9.78)$$

Per interpretare la (9.78) ci conviene definire ancora il *waist* (larghezza minima) del fascio

$$w_0 = \sqrt{\frac{2z_R}{k}} = \sqrt{\frac{\lambda z_R}{\pi}}, \quad (9.79)$$

la *spot size* (dimensione della macchia, o larghezza) del fascio alla coordinata z

$$w(z) = w_0 \sqrt{1 + \left(\frac{z}{z_R}\right)^2} = \sqrt{\frac{2z_R}{k}} \sqrt{1 + \left(\frac{z}{z_R}\right)^2} = \sqrt{\frac{2(z^2 + z_R^2)}{kz_R}}, \quad (9.80)$$

e infine il *raggio di curvatura*

$$R(z) = z \left[1 + \left(\frac{z_R}{z}\right)^2 \right] = \frac{z^2 + z_R^2}{z}. \quad (9.81)$$

Inserendo tutto nella (9.78) otteniamo

$$\tilde{u}(r, z) = \tilde{E}_0 \frac{w_0}{w(z)} e^{-r^2/w^2(z)} e^{-ir^2/2R(z)} e^{i\psi(z)}. \quad (9.82)$$

Dalla (9.62) abbiamo per il fasore completo del campo elettrico del fascio

$$\tilde{E}(r, z) = \tilde{u}(r, z) e^{-iz} = \tilde{E}_0 \frac{w_0}{w(z)} e^{-r^2/w^2(z)} e^{-ir^2/2R(z)} e^{-iz} e^{i\psi(z)}. \quad (9.83)$$

Sul piano $z = 0$ abbiamo

$$w(0) = w_0, \quad \psi(0) = 0, \quad \frac{1}{R(0)} = 0, \quad (9.84)$$

e la (9.83) diventa

$$\tilde{E}(r, 0) = \tilde{E}_0 e^{-r^2/w_0^2}, \quad (9.85)$$

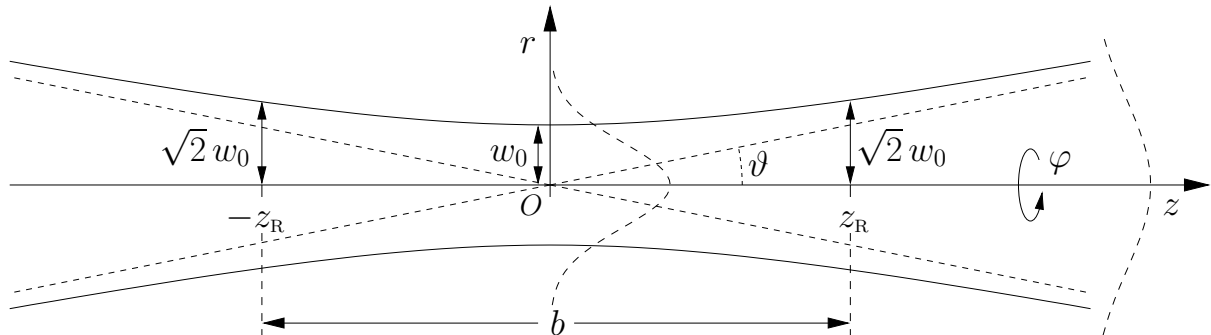


Figura 9.10 Fascio gaussiano

così, per $z = 0$, il campo elettrico ha un massimo per $r = 0$, e, allontanandosi dall'asse z , decresce come e^{-r^2/w_0^2} . Quindi a $r = w_0$ (waist) il campo è diminuito di un fattore $1/e$, cioè $\tilde{E}(w_0, 0) \simeq 0.37 \tilde{E}(0, 0)$, e la potenza del fascio, proporzionale al quadrato del campo, è diminuita di un fattore $1/e^2$, quindi $I(w_0, 0) \simeq 0.14 I(0, 0)$. Per qualunque altro valore di $z \neq 0$, $\tilde{E}(r, z)$ è ancora massimo per $r = 0$, e la sua intensità decresce con r come $e^{-r^2/w^2(z)}$ raggiungendo il valore $\tilde{E}(r, z) = \tilde{E}(0, z)/e$ per $r = w(z)$ (larghezza del fascio). L'espressione (9.80) per la larghezza del fascio $w(z)$ può essere riscritta

$$\frac{w^2(z)}{w_0^2} - \frac{z^2}{z_R^2} = 1 \quad (9.86)$$

che è l'equazione dell'iperboloide di rotazione rappresentato in Fig. 9.10, che, per $z \rightarrow \pm\infty$ tende asintoticamente al cono di semiapertura $\vartheta = w_0/z_R$. La larghezza del fascio è così minima per $z = 0$, e, per $z \rightarrow \pm\infty$, tende all'infinito come $|z|w_0/z_R$. Per $|z| = z_R$ (lunghezza di Rayleigh) la larghezza del fascio vale $w(z_R) = \sqrt{2} w_0$, quindi z_R è un parametro che misura quanto il fascio è collimato. La lunghezza $b = 2z_R$ si chiama *parametro confocale*.

E' da notare che il comportamento del campo lungo l'intero fascio gaussiano dipende solo dalla lunghezza d'onda λ , più uno qualsiasi degli altri parametri, per esempio il waist w_0 . Infatti possiamo scrivere

$$z_R = \frac{\pi}{\lambda} w_0^2, \quad (9.87)$$

$$w(z) = w_0 \sqrt{1 + \left(\frac{z}{z_R}\right)^2} = \frac{1}{\pi w_0} \sqrt{\pi^2 w_0^4 + \lambda^2 z^2}, \quad (9.88)$$

$$R(z) = z + \frac{z_R^2}{z} = z + \frac{\pi^2 w_0^4}{\lambda^2 z} \quad (9.89)$$

$$\psi(z) = -\arctan\left(\frac{z}{z_R}\right) = -\arctan\left(\frac{\lambda z}{\pi w_0^2}\right), \quad (9.90)$$

$$\vartheta = \frac{w_0}{z_R} = \frac{\lambda}{\pi w_0}. \quad (9.91)$$

Mentre sul piano $z = 0$ il campo è in fase in ogni punto, su un piano con $z \neq 0$ si ha una variazione di fase che va come $e^{-ir^2/2R(z)}$. Questo è dovuto al fatto che il fronte d'onda del fascio è sferico, con raggio di curvatura R . Infatti, dalla Fig. 9.11 si vede che la fase del punto $P \equiv (r, \varphi, z)$, con φ non rilevante, non è uguale alla fase del punto $(0, \varphi, z)$, ma a quella del punto $(0, \varphi, z + \Delta z)$, che sta sullo stesso fronte d'onda. Poiché stiamo parlando di fasci parassiali, l'angolo α è piccolo (in altre parole, $r \ll R$), e possiamo approssimare

$$\alpha \simeq \frac{r}{R} \quad (9.92)$$

quindi

$$\Delta z = R - R \cos \alpha \simeq R \frac{1}{2} \frac{r^2}{R^2} = \frac{r^2}{2R}, \quad (9.93)$$

e la fase è

$$-i \left(z + \frac{r^2}{2R(z)} \right). \quad (9.94)$$

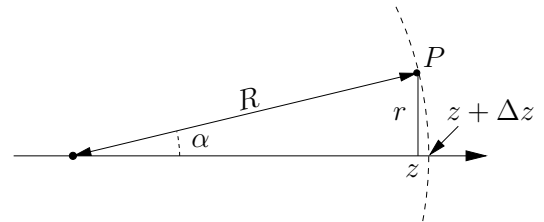


Figura 9.11 Fase del fronte d'onda.

A questa si aggiunge lo sfasamento di Gouy $\psi(z)$, che, secondo la (9.77) varia da $-\pi/2$ a $\pi/2$ per z che varia da $-\infty$ a $+\infty$. Infine, il fattore $w_0/w(z)$ della (9.83) assicura la conservazione dell'energia, diminuendo l'altezza della distribuzione gaussiana man mano che la curva si allarga (gaussiane tratteggiate della Fig. 9.10).

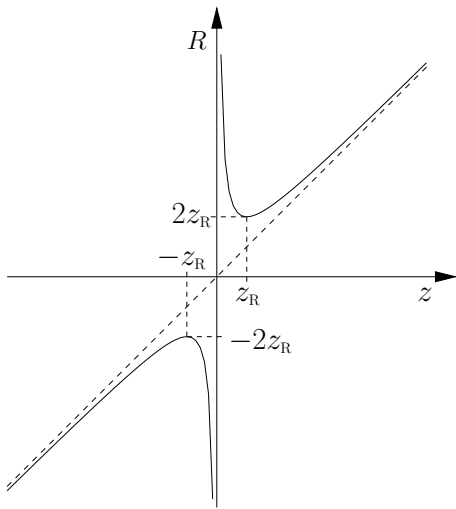


Figura 9.12 Raggio di curvatura del fronte d'onda.

coincidono nel punto $z = 0$, e il sistema è un *risuonatore confocale simmetrico*, che riconsidereremo più avanti.

Per $z > z_R$ il raggio di curvatura comincia a risalire, tendendo asintoticamente a z per $z \rightarrow \infty$, come in Fig. 9.12. Quindi, per $z \gg z_R$, il fascio si comporta come un'onda sferica centrata al centro del waist.

9.5 Stabilità della cavità risonante

Il risuonatore ottico più semplice è costituito da due specchi sferici sistemati uno di fronte all'altro. Se le curvature degli specchi corrispondono a un sistema di focalizzazione periodico stabile, e se le loro dimensioni trasversali sono abbastanza grandi da poter trascurare gli effetti di diffrazione ai bordi

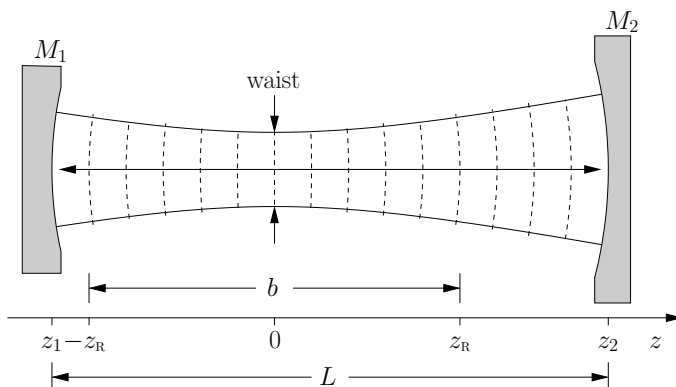


Figura 9.13 Risuonatore stabile a due specchi.

del fascio, gli specchi possono intrappolare un fascio di luce che continuerà a “rimbalzare” da uno specchio all'altro. Consideriamo un fascio gaussiano con un certo waist in $z = 0$, e due specchi sferici M_1 ed M_2 posizionati, rispettivamente, in z_1 e z_2 , dove la loro curvatura si adatta perfettamente al fronte d'onda del fascio, come in Fig. 9.13. Supponiamo che le dimensioni trasversali dei due specchi siano sufficientemente grandi da poter trascurare perdite per diffrazione. In queste condizioni i due specchi riflettono il fascio gaussiano esattamente su sé stesso, invertendone esattamente la curvatura del fronte d'onda e la dire-

A grandi distanze il fascio laser diverge con semiapertura ϑ data dalla (9.91), quindi ha divergenza tanto maggiore quanto più è grande la lunghezza d'onda, e tanto minore quanto più è largo il waist. Il raggio di curvatura del fronte d'onda del fascio è dato dalla (9.89), quindi il fronte d'onda è piano per $z = 0$, dove $R(0) = \pm\infty$, con $R(0^+) = +\infty$ e $R(0^-) = -\infty$. Considerando solo il caso $z > 0$, visto che il caso $z < 0$ è semplicemente antisimmetrico, $R(z)$ diminuisce al crescere di z fino a raggiungere un minimo $R(z_R) = 2z_R$ per $z = z_R$. Quindi, il centro di curvatura del fronte d'onda a $z = z_R$ si trova in $z = -z_R$, e viceversa. Questo caso particolare è molto importante nel disegno di risuonatori laser stabili. Infatti, se poniamo due specchi sferici in $\pm z_R$, ognuno esattamente di raggio di curvaturee $2z_R$, la superficie di ogni specchio si adatta esattamente alla superficie del rispettivo fronte d'onda. Poiché la lunghezza focale f di uno specchio sferico di raggio R è $f = R/2$, i fuochi dei due specchi

coincidono nel punto $z = 0$, e il sistema è un *risuonatore confocale simmetrico*, che riconsidereremo più avanti.

zione. In questo modo i due specchi intrappolano il fascio gaussiano sotto forma di onda stazionaria, formando una cavità risonante, o risuonatore ottico.

Nella pratica, però, si affronta normalmente il problema inverso: dati due specchi sferici concavi M_1 e M_2 , rispettivamente di raggio R_1 e R_2 , posti a una distanza L tra loro, dovremo trovare i parametri w_0 e z_R del fascio gaussiano. La condizione che il fronte d'onda si adatti alle curvature degli specchi porta al seguente sistema di tre equazioni nelle tre incognite z_R , z_1 e z_2

$$\begin{aligned} R(z_1) &= z_1 + \frac{z_R^2}{z_1} = -R_1 \\ R(z_2) &= z_2 + \frac{z_R^2}{z_2} = R_2 \\ z_2 - z_1 &= L \end{aligned} \quad (9.95)$$

e w_0 sarà determinato dalla (9.87). Il segno di R_1 nella prima equazione è dovuto al fatto che la (9.89) attribuisce raggio negativo ai fronti d'onda con concavità verso destra, e positivo ai fronti d'onda con concavità verso sinistra, mentre noi consideriamo sempre positivo il raggio di uno specchio sferico concavo. Prima di risolvere le equazioni, conviene introdurre i *parametri di stabilità* dei due specchi, g_1 e g_2 , definiti come

$$g_1 = 1 - \frac{L}{R_1}, \quad g_2 = 1 - \frac{L}{R_2}, \quad (9.96)$$

in funzione dei quali i parametri del fascio si scrivono

$$z_R = \frac{\sqrt{g_1 g_2 (1 - g_1 g_2)}}{g_1 + g_2 - 2g_1 g_2} L, \quad z_1 = -\frac{g_2 (1 - g_1)}{g_1 + g_2 - 2g_1 g_2} L, \quad z_2 = \frac{g_1 (1 - g_2)}{g_1 + g_2 - 2g_1 g_2} L, \quad (9.97)$$

mentre per il waist abbiamo

$$w_0^2 = \frac{L\lambda}{\pi} \sqrt{\frac{g_1 g_2 (1 - g_1 g_2)}{(g_1 + g_2 - 2g_1 g_2)^2}}. \quad (9.98)$$

Dalle equazioni qui sopra si vede che soluzioni reali e finite per i parametri che caratterizzano il fascio gaussiano esistono solo se è soddisfatta la condizione

$$0 \leq g_1 g_2 \leq 1. \quad (9.99)$$

In altre parole, dati due specchi di raggi R_1 e R_2 posti a distanza L tra loro, si determinano i parametri g_1 e g_2 della (9.96), e se il punto (g_1, g_2) cade nell'area evidenziata del grafico di Fig. 9.14 il risuonatore è stabile e può intrappolare una famiglia di fasci gaussiani con parametri dati dalle (9.97) e (9.98), oltre a fasci corrispondenti a modi superiori. Se invece il punto (g_1, g_2) cade fuori dall'area evidenziata il risuonatore sarà instabile. La Fig. 9.14 indica anche i casi particolari di *risuonatore piano*, con $R_1 = R_2 = \infty$, quindi con $g_1 = g_2 = 1$, *risuonatore confocale*, con $R_1 = R_2 = L$ e $g_1 = g_2 = 0$, e *risuonatore concentrico*, con $R_1 = R_2 = L/2$ e $g_1 = g_2 = -1$. Una volta soddisfatta la condizione di stabilità, i risuonatori possono ancora essere classificati in vari tipi, qui ne ricordiamo due.

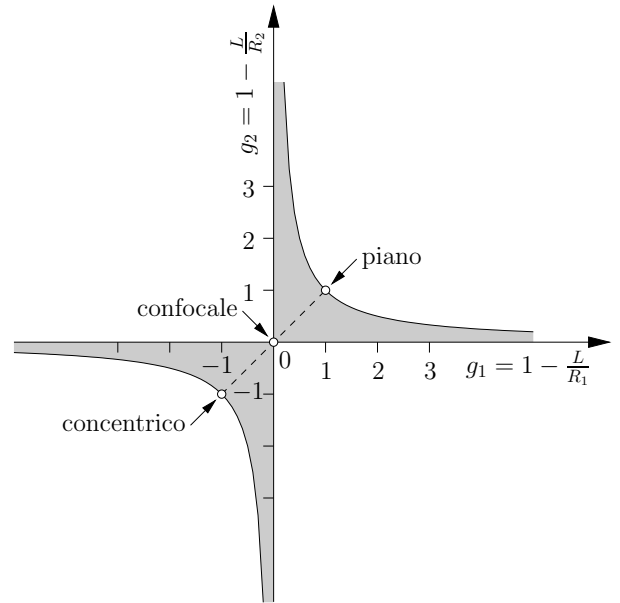


Figura 9.14 Grafico di stabilità per una cavità risonante a due specchi.

Risuonatori simmetrici

Sono la configurazione più semplice, con $R_1 = R_2 = R$ e $g_1 = g_2 = g = 1 - L/R$. Quindi il waist si trova al centro del risuonatore e avremo

$$z_2 = -z_1 = \frac{L}{2}, \quad w_0^2 = \frac{L\lambda}{\pi} \sqrt{\frac{1+g}{4(1-g)}}. \quad (9.100)$$

In Fig. 9.14 i punti rappresentativi di questi risuonatori si trovano, naturalmente, sul segmento tratteggiato che unisce il risuonatore concentrico al risuonatore piano passando per il risuonatore confocale, che sono tre casi particolari.

Risuonatori emisferici

Qui abbiamo uno specchio sferico di raggio R_1 , e uno specchio piano, quindi con $R_2 = \infty$, e i parametri di stabilità sono

$$g_1 = 1 - \frac{L}{R_1}, \quad g_2 = 1, \quad (9.101)$$

e i punti rappresentativi stanno su un segmento orizzontale della Fig. 9.14 che va dall'asse delle g_2 al punto rappresentativo del risuonatore piano. Questo risuonatore è equivalente ad un risuonatore simmetrico di lunghezza doppia, e il waist del fascio si trova sullo specchio piano.

9.6 Instaurazione del regime oscillatorio

Per la statistica dei fotoni, i modi di oscillazione di una cavità laser si comportano come stati di particella singola. Consideriamo, per esempio, una cavità laser limitata da due specchi piani e paralleli a distanza L tra loro, come schematizzato in figura 9.15. Ogni modo della cavità corrisponde ad un'onda stazionaria tra i due specchi. Perché l'onda sia stazionaria deve essere

$$L = r \frac{\lambda_r}{2}, \quad \text{da cui} \quad \lambda_r = \frac{2L}{r}, \quad (9.102)$$

con r numero intero. Le frequenze possibili per i modi della cavità sono quindi

$$\nu_r = \frac{c}{\lambda_r} = r \frac{c}{2L}, \quad (9.103)$$

e a ogni modo corrispondono fotoni di energia $E_r = h\nu_r$. Naturalmente, il laser può emettere la frequenza ν_r solo se questa può eccitare emissione indotta nel mezzo attivo, cioè se ν_r è in risonanza con la transizione laser-attiva, entro la sua larghezza (in altre parole, se ν_r è entro la *curva di guadagno* del laser). Quindi, fissata la lunghezza della cavità, il mezzo laser-attivo, e i parametri che determinano la sua curva di guadagno come, per esempio, la pressione nel caso dei laser a gas, il laser può oscillare solo su una certa gamma di modi. Finché il laser non raggiunge la soglia di oscillazione, i fotoni si comportano come se fossero all'equilibrio termico ad una *temperatura efficace* T_{eff} . Trattandosi di bosoni, ogni modo $|r\rangle$ avrà un numero medio di occupazione

dato dalla (5.68) con il potenziale chimico μ posto uguale a zero, quindi

$$\langle n_r \rangle = \frac{1}{e^{E_r/kT_{\text{eff}}} - 1}. \quad (9.104)$$

E' da notare che la temperatura T_{eff} tende all'infinito man mano che ci si avvicina all'inversione di popolazione. Infatti la temperatura efficace T_{eff} per i due livelli interessati alla transizione laser si ottiene risolvendo per T la relazione di Boltzmann $N_2/N_1 = e^{-(E_2-E_1)/kT}$, che dà

$$T_{\text{eff}} = \frac{E_2 - E_1}{k \ln(N_1/N_2)}. \quad (9.105)$$

Così T_{eff} diverge per $N_2 \rightarrow N_1$, per poi diventare negativa una volta che si è raggiunta l'inversione di popolazione, a cui $N_2 > N_1$ (per un laser, si passa da temperature positive a temperature negative non passando attraverso 0 K, ma passando attraverso l'infinito!). I numeri di occupazione dei modi del laser tenderanno quindi a divergere con T_{eff} . Per le fluttuazioni relative dei numeri di occupazione stessi, man mano che ci si avvicina all'inversione di popolazione otteniamo dalla (5.78)

$$\lim_{N_2 \rightarrow N_1} \frac{\langle (\Delta n_r)^2 \rangle}{\langle n_r \rangle^2} = \lim_{\langle n_r \rangle \rightarrow \infty} \frac{\langle (\Delta n_r)^2 \rangle}{\langle n_r \rangle^2} = \lim_{\langle n_r \rangle \rightarrow \infty} \left(\frac{1}{\langle n_r \rangle} + 1 \right) = 1, \quad (9.106)$$

da cui il numero di fotoni nel modo r ha fluttuazioni dell'ordine del 100% attorno al valor medio $\langle n_r \rangle$. Ponendo $\varepsilon = +1$ e $\mu = 0$ nella (5.68) otteniamo il numero medio di fotoni in uno stato (modo) di energia E_r

$$\langle n_r \rangle = \frac{1}{e^{\beta E_r} - 1}, \quad \text{da cui abbiamo} \quad e^{\beta E_r} = 1 + \frac{1}{\langle n_r \rangle}, \quad \text{e} \quad e^{-\beta E_r} = \frac{\langle n_r \rangle}{\langle n_r \rangle + 1},$$

che, inserita nella (5.61) ci dà per la probabilità di avere n fotoni in un modo con numero medio di occupazione $\langle n_r \rangle$

$$W_r(n) = e^{-n\beta E_r} (1 - e^{-\beta E_r}) = \frac{\langle n_r \rangle^n}{(\langle n_r \rangle + 1)^{n+1}}. \quad (9.107)$$

Appena raggiunta la soglia di oscillazione del laser, il comportamento statistico dei fotoni cambia in modo discontinuo. La probabilità di avere n fotoni nel modo $|r\rangle$ quando il numero medio di occupazione è $\langle n_r \rangle$ è data adesso dalla distribuzione di Poisson

$$W_r(n) = \frac{\langle n_r \rangle^n}{n!} e^{-\langle n_r \rangle}, \quad (9.108)$$

cui corrispondono delle fluttuazioni relative

$$\sqrt{\frac{\langle \Delta n^2 \rangle}{\langle n \rangle^2}} = \frac{1}{\sqrt{\langle n \rangle}}, \quad (9.109)$$

uguali a quelle di un oscillatore classico stabilizzato in ampiezza.

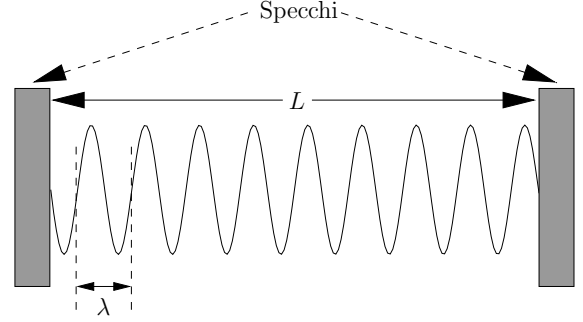


Figura 9.15 Onda stazionaria in una cavità laser a specchi piani.

9.7 Mode-locking

Se non vengono presi particolari accorgimenti, un laser può oscillare contemporaneamente su un numero N abbastanza grande di modi longitudinali. In condizioni normali le fasi dei singoli modi avranno valori casuali, che avranno poi salti casuali dovuti varie cause, quindi, se il laser opera in continua, l'intensità del fascio dipenderà in modo stocastico dal tempo. Se, invece, la fase di ogni modo rimane costante nel tempo, e quindi le differenze di fase tra i vari modi sono costanti nel tempo (*phase-locking*, o *aggancio di fase*), il comportamento è completamente diverso. Infatti il campo complessivo del fascio è dato dalla somma di N frequenze diverse, equispaziate con intervallo di frequenza $\Delta\nu = c/2L$, in base alla (9.103). Pertanto l'intensità ha queste due proprietà, caratteristiche di una serie di Fourier,

1. La forma d'onda è periodica con periodo $\tau_p = 1/\Delta\nu$.
2. Sono presenti impulsi molto intensi quando tutti i modi sono in fase, e questo succede con periodo τ_p . Ognuno di questi impulsi ha una durata $\Delta\tau_p$ che dipende dal numero di modi agganciati, approssimativamente si ha $\Delta\tau_p \approx 1/(N\Delta\nu)$.

In questo modo, in laser con banda di guadagno sufficientemente larga, come laser a stato solido, laser a coloranti organici o laser a semiconduttori, il numero di modi attivi (e agganciati) può essere molto grande, portando a un valore alto di $N\Delta\nu$, e a una conseguente durata $\Delta\tau_p$ dell'impulso molto breve, si può arrivare all'ordine del picosecondo (10^{-12} s) o, addirittura, del femtosecondo (10^{-15} s).

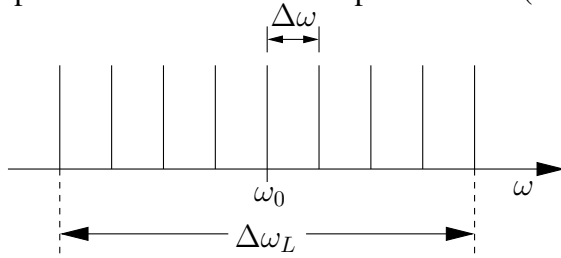


Figura 9.16 9 modi equispaziati, e di uguale intensità, oscillanti contemporaneamente.

Come primo esempio consideriamo il caso rappresentato in Fig. 9.16, con $2n + 1$ (9 nella figura) modi equispaziati di $\Delta\omega$, tutti di uguale ampiezza E_0 . Facciamo anche l'ipotesi che lo sfasamento tra due modi successivi, $\Delta\varphi$, sia costante, cioè, per ogni r con $-n < r \leq n$

$$\varphi_r - \varphi_{r-1} = \Delta\varphi. \quad (9.110)$$

In un qualunque punto del fascio, il campo elettrico complessivo $E(t)$ può essere scritto, a parte una fase dipendente dall'origine dei tempi e dal singolo punto, come

$$E(t) = \sum_{r=-n}^n E_0 e^{i[(\omega_0 + r\Delta\omega)t + r\Delta\varphi]}, \quad (9.111)$$

dove ω_0 è la frequenza del modo centrale e $\Delta\omega$ è la spaziatura tra due modi consecutivi, $\Delta\omega = \pi c/L$, dove L è la lunghezza della cavità. La (9.111) può essere riscritta

$$E(t) = A(t) e^{i\omega_0 t}, \quad (9.112)$$

separando l'onda portante, di frequenza ω_0 , da un'ampiezza dipendente dal tempo $A(t)$:

$$A(t) = \sum_{r=-n}^n E_0 e^{ir(\Delta\omega t + \delta\varphi)}. \quad (9.113)$$

A questo punto ci conviene scegliere una nuova origine dei tempi tale che sia $\Delta\omega t' = \Delta\omega t + \Delta\varphi$, ovvero $t' = t + \Delta\varphi/\Delta\omega$. In funzione di t' la (9.113) diventa

$$A(t') = E_0 \sum_{r=-n}^n e^{ir\Delta\omega t'}, \quad (9.114)$$

dove E_0 moltiplica la somma dei termini tra $-n$ e n della progressione geometrica di ragione $e^{i\Delta\omega t'}$, calcolata in Appendice C, e così abbiamo

$$A(t') = E_0 \sum_{r=-n}^n e^{ir\Delta\omega t'} = E_0 \frac{\sin[(2n+1)\Delta\omega t'/2]}{\sin(\Delta\omega t'/2)}. \quad (9.115)$$

Per interpretare la (9.115) cominciamo osservando che $A(t')$ è periodica nel tempo, con periodo

$$\tau_p = \frac{2\pi}{\Delta\omega} = \frac{1}{\Delta\nu} = 2\frac{L}{c}. \quad (9.116)$$

Il denominatore si annulla ogni volta che $t' = m\tau_p$, con m intero qualunque. Allo stesso valore $t' = m\tau_p$ si annulla anche il numeratore, e per $t' \rightarrow m\tau_p$ la frazione tende al valore finito

$$\lim_{t' \rightarrow 2mL/c} \frac{\sin[(2n+1)\Delta\omega t'/2]}{\sin(\Delta\omega t'/2)} = 2n+1. \quad (9.117)$$

A questi valori di t' l'ampiezza $A(t')$ presenta dei massimi di valore $A(m\tau_p) = (2n+1)E_0$, corrispondenti al fatto che tutti i $2n+1$ modi sono in fase. A questi istanti abbiamo così dei massimi dell'intensità del fascio laser, proporzionali a $(2n+1)^2 E_0^2$. Attorno a ogni massimo l'ampiezza $A(t')$ si annulla al primo zero precedente e al primo zero successivo del numeratore, cioè per $t' = m\tau_p \pm t'_p$,

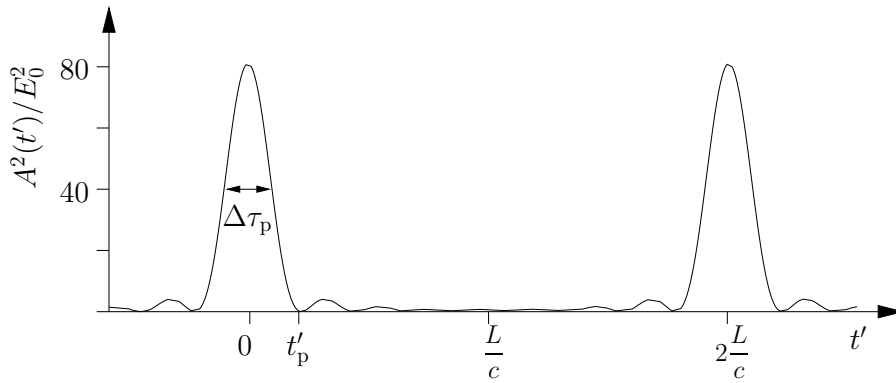


Figura 9.17 Andamento temporale del quadrato dell'ampiezza del campo elettrico nel caso di 9 modi equispaziati di uguale ampiezza agganciati in fase.

con

$$t'_p = \frac{2\pi}{2(n+1)\Delta\omega} = \frac{1}{\Delta\nu_L}, \quad (9.118)$$

dove $\Delta\nu_L = (2n+1)\Delta\omega/(2\pi) = \Delta\omega_L/(2\pi)$ è la distanza in frequenza tra il primo e l'ultimo modo oscillante, vedi Fig. 9.16. A parità di spaziatura tra i modi, i picchi sono così tanto più stretti

quanto è maggiore il numero dei modi oscillanti. La larghezza FWHM a metà altezza del picco è quindi $\Delta\tau_p \approx t'_p = 1/\Delta\nu_L$. Quindi, tanti più modi riusciamo ad agganciare in fase, tanto più brevi sono gli impulsi. La Fig. 9.17 rappresenta, in funzione del tempo, l'intensità di un fascio con 9 modi agganciati in fase. La Fig. 9.17 si spiega facilmente rappresentando le ampiezze sommate nella (9.114) come vettori nel piano complesso. L'ampiezza r -esima corrisponde a un vettore complesso di modulo E_0 che ruota con velocità angolare $r\Delta\omega$. Per $t' = 0$, o t' uguale a un multiplo di τ_p , la (9.114) ci dice che tutti i vettori hanno fase zero e sono sovrapposti come in Fig. 9.18 a). In questo caso, il campo complessivo è $(2n + 1)E_0$. Al crescere di t' il vettore corrispondente a $r = 0$ rimane fisso, mentre i vettori con $r > 0$ ruotano in un verso (per esempio, antiorario), e quelli con $r < 0$ in verso opposto (per esempio, orario), come in Fig. 9.18 b). Quando

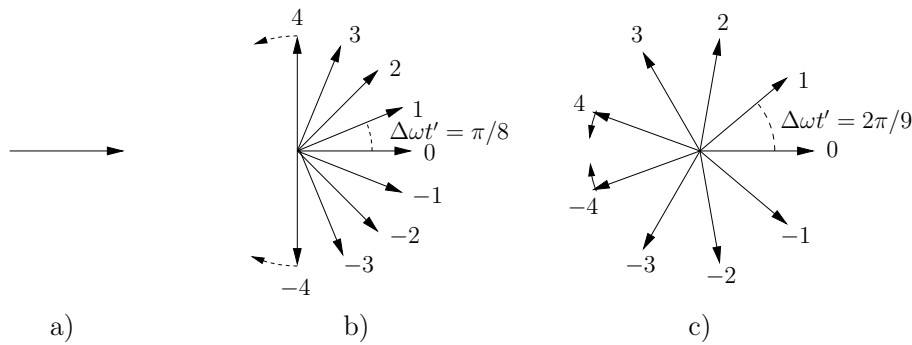


Figura 9.18 Le ampiezze dei nove modi della cavità nel piano complesso: a) $t' = m\tau_p$, le ampiezze sono tutte sovrapposte, l'ampiezza complessiva vale $9E_0$; b) $t' = m\tau_p + \pi/(8\Delta\omega)$, i modi -4 e $+4$ si cancellano; c) $t' = m\tau_p + 2\pi/(9\Delta\omega)$, configurazione simmetrica, l'ampiezza complessiva è nulla.

t' è cresciuto di una quantità $t'_p = 2\pi/[2(n + 1)\Delta\omega]$, la spaziatura angolare tra un modo e l'altro vale $2\pi/[2(n + 1)]$, quindi le ampiezze sono disposte in modo simmetrico, come in Fig. 9.18 c), e il campo complessivo è nullo. Dopo un tempo $\tau_p = 2\pi/\Delta\omega$, il vettore con $r = 1$ ha fatto un giro completo ed è nuovamente sovrapposto a r_0 ,

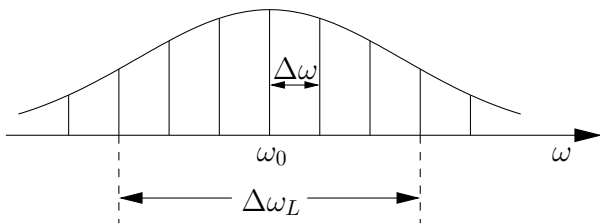


Figura 9.19 Modi equispaziati con ampiezza determinata da una curva di guadagno di profilo gaussiano.

allo stesso modo, il generico vettore r avrà fatto r giri completi (in un verso o nell'altro a seconda del segno), e tutti i vettori sono nuovamente sovrapposti. I conti sviluppati fin qui sono validi nell'ipotesi che tutti i modi agganciati in fase abbiano la stessa ampiezza, come in Fig. 9.16. In realtà, nel caso generale l'ampiezza dei singoli modi è determinata dalla curva di guadagno della transizione laser, come mostrato in Fig. 9.19 per il caso di un profilo

gaussiano. I conti diventano quindi più complicati, ma la trattazione appena fatta ci dà una buona idea qualitativa di quanto accade. Resta il fatto che la durata dell'impulso è inversamente proporzionale a $\Delta\nu_L$, che è dell'ordine di grandezza della larghezza della curva di guadagno. Quindi, come già detto, ci aspettiamo di poter ottenere impulsi molto brevi per laser a stato solido, a semiconduttori o a coloranti organici (*dye lasers*), mentre le curve di guadagno dei laser a gas, molto più strette, non si prestano per questa tecnica.

9.8 Laser a semiconduttore

Mentre in atomi e molecole gli elettroni hanno livelli energetici discreti, la distribuzione dei livelli energetici degli elettroni di un semiconduttore è a bande, come schematizzato in Fig. 9.20. Questa differenza ha conseguenze importanti nei laser che usano semiconduttori come mezzo attivo. La transizione laser, infatti, non avviene tra due livelli energetici ben definiti, ma tra due stati che appartengono a due bande diverse. La Fig. 9.20 mostra che in un semiconduttore per gli elettroni esistono una *banda di conduzione* e una *banda di valenza*, separate da un gap energetico (banda di energie proibite) di ampiezza ΔE_g . Ogni banda è costituita da un grandissimo numero di livelli energetici molto vicini l'uno all'altro, che formano così un quasi-continuo. In base al principio di esclusione di Pauli in ognuno di questi livelli possono trovarsi al massimo due elettroni, con spin opposto. La (5.68), posto $\varepsilon = -1$, ci dà per il numero medio di occupazione di un livello di energia E_r nel caso della statistica di Fermi-Dirac

$$\langle n_r \rangle = \frac{1}{e^{(E_r - \mu)/kT} + 1}, \quad (9.119)$$

dove μ è il potenziale chimico. Al limite $T \rightarrow 0$ il potenziale chimico tende all'energia di Fermi E_F , e, a questo limite, abbiamo per i numeri di occupazione

$$\langle n_r \rangle = \begin{cases} 1 & \text{se } E_r < E_F \\ 0 & \text{se } E_r > E_F \end{cases}. \quad (9.120)$$

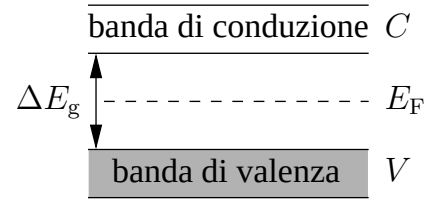


Figura 9.20 Banda di conduzione C, banda di valenza V ed energia di Fermi E_F in un semiconduttore.

La *temperatura di Fermi* T_F è definita come $T_F = E_F/k$. Per tutti i metalli T_F è molto maggiore della temperatura ambiente T , così che il potenziale chimico $\mu(T)$ differisce di molto poco dall'energia di Fermi $E_F = \mu(0)$. Anche nei semiconduttori si tende a confondere potenziale chimico e temperatura di Fermi. Per semiconduttori *intrinseci* (non drogati) l'energia di Fermi E_F si trova all'interno del gap di energia (Fig. 9.20), così che per $T = 0$ la banda di valenza è completamente occupata e la banda di conduzione è completamente vuota. Ma per avere effetto laser abbiamo bisogno di una inversione di popolazione, cioè dobbiamo avere contemporaneamente elettroni nella banda di conduzione e stati elettronici non occupati, o *buche*, nella banda di valenza, in cui gli elettroni della banda di conduzione possano decadere. Questo, non è ottenibile in un semiconduttore intrinseco all'equilibrio termico.

Si possono però realizzare dei semiconduttori *estrinseci*, cioè semiconduttori cui sono state aggiunte impurezze nel reticolo cristallino. In altre parole, si dice che questi semiconduttori sono stati "drogati". Esistono due tipi di *drogaggio*. Nel drogaggio di tipo *n* vengono aggiunti al cristallo atomi detti *donori* che, avendo più elettroni di valenza degli atomi del reticolo, mettono in circolazione un eccesso di elettroni debolmente legati. Questi elettroni, essendo tutti gli stati della banda di valenza occupati, vanno a finire nella banda di conduzione, e si ha un "quasi-livello di Fermi" E_C all'interno della banda di conduzione, come in Fig. 9.21.a). Siamo così riusciti ad ottenere elettroni nella banda di conduzione, ma tutti gli stati della banda di valenza sono ancora occupati: non abbiamo buche in cui gli elettroni della banda di conduzione possano decadere.

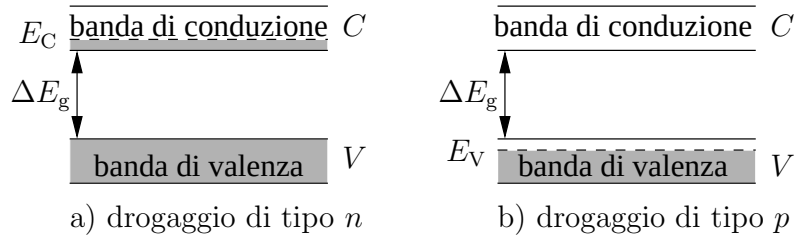


Figura 9.21 Semiconduttori con drogaggio di tipo *n* (a) e tipo *p* (b).

Nel drogaggio di tipo p , invece, vengono aggiunti al semiconduttore atomi detti *accettori* che, avendo orbitali esterni liberi, prelevano elettroni dal cristallo, vuotando così gli stati più alti della banda di valenza, dando origine a buche. Si ha così un “quasi-livello di Fermi” E_V all’interno della banda di valenza, come nella figura 9.21.b).

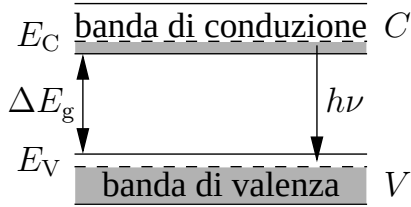


Figura 9.22 Inversione di popolazione e transizione laser in un laser a semiconduttore.

In questo caso abbiamo sì stati liberi, o buche, nella banda di valenza, ma nella banda di conduzione non abbiamo elettroni che vi possano decadere.

Per arrivare all’effetto laser dobbiamo ottenere contemporaneamente elettroni negli stati più bassi della banda di conduzione e buche negli stati più alti della banda di valenza, come indicato nella figura 9.22. In queste condizioni potrebbero essere emessi fotoni con frequenze comprese nell’intervallo

$$\frac{\Delta E_g}{h} < \nu < \frac{E_C - E_V}{h}. \quad (9.121)$$

Fotoni con frequenze in questo intervallo che entrassero nel cristallo sarebbero amplificati per emissione indotta, mentre fotoni di frequenza $\nu > (E_C - E_V)$ sarebbero assorbiti, portando elettroni dalla banda di valenza a stati liberi della banda di conduzione, e, finalmente, fotoni di frequenza $\nu < \Delta E_g$ troverebbero il mezzo trasparente. L’inversione di popolazione in un laser a semiconduttore può essere ottenuta in tre modi:

1. Per eccitazione via pompaggio ottico.
2. Per eccitazione mediante urti con elettroni ad alta energia.
3. Per iniezione di portatori di carica attraverso una giunzione.

Essendo il caso di gran lunga più interessante per le applicazioni pratiche, qui discuteremo soltanto il caso dell’iniezione. Un modo di realizzare questo tipo di laser a semiconduttore consiste nel

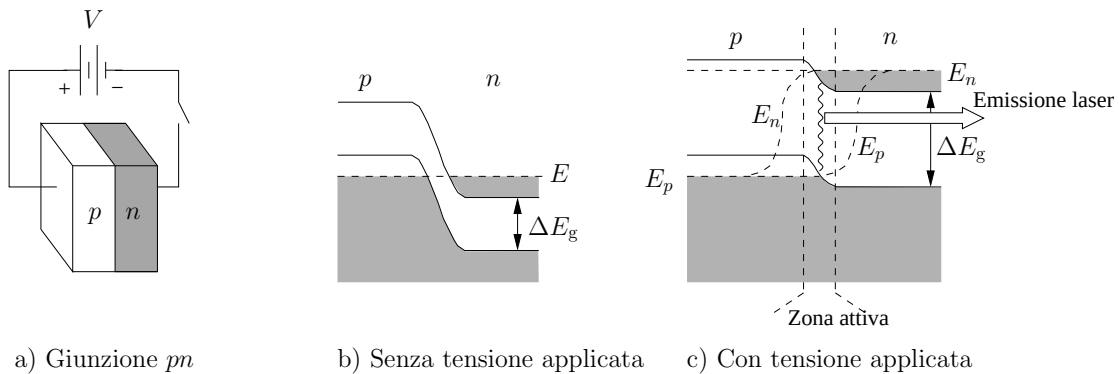


Figura 9.23 Diodo laser.

mettere in sequenza due strati di semiconduttore estrinseci, uno con drogaggio di tipo p ed uno con drogaggio di tipo n , come in Fig. 9.23.a), esattamente come in un diodo a semiconduttore (qui si parla di *laser a diodo*). Se ai capi del diodo non viene applicata tensione, la situazione delle bande

energetiche è quella schematizzata in Fig. 9.23.b). Mentre le bande energetiche della zona p e della zona n sono spostate le une rispetto alle altre, il livello di Fermi ha un valore costante in tutta la zona di transizione pn , corrispondente all'equilibrio termico. A causa del drogaggio il livello di Fermi si trova all'interno della banda di valenza nella zona p , e all'interno della banda di conduzione nella zona n . Alla giunzione si forma una *regione di svuotamento*, dello spessore di circa $0.5\ \mu\text{m}$, in cui gli elettroni si combinano con le buche e non ci sono portatori di carica. Se adesso viene introdotta tra lo strato p e lo strato n una differenza di potenziale V corrispondente al gap di energia, cioè $V = \Delta E_g/e$, nella direzione di conduzione, come in Fig. 9.23.c), le bande energetiche della zona n si innalzano rispetto a quelle della zona p (gli elettroni hanno carica negativa!). Come mostrato in Fig. 9.23.c), l'energia di Fermi E_n dello strato n viene innalzata di una quantità Ve rispetto all'energia di Fermi E_p dello strato p . Esiste così una sottile zona attiva in cui si trovano sia elettroni in banda di conduzione che buche in banda di valenza, con conseguente inversione di popolazione.

Qui è opportuno ricordare che esistono due tipi di banda proibita, o gap energetico per i semiconduttori: il *gap diretto* e il *gap indiretto*. L'energia dei livelli degli elettroni in un cristallo dipendono dai quasimomenti $\mathbf{k}$ associati agli elettroni. Tutte le interazioni tra elettroni, buche, fononi, fotoni e altre particelle o quasi-particelle devono conservare l'energia e il momento complessivo. Un fotone di energia prossima a ΔE_g ha momento trascurabile rispetto al vettore $\mathbf{k}$ dello stato della banda di conduzione in cui si trova l'elettrone. Quindi, se il vettore $\mathbf{k}$ dello stato più basso della banda di conduzione è diverso dal vettore $\mathbf{k}$ dello stato più alto della banda di valenza l'emissione di un fotone non può compensare la variazione di momento di un elettrone che si ricombina con una buca. Questo è per esempio, il caso del silicio, in cui la ricombinazione, non radiativa, può avvenire, per esempio, in difetti del cristallo. Un diodo al silicio può così essere usato come raddrizzatore di corrente, ma non per l'emissione di luce. Questo tipo di gap energetico si dice indiretto. L'emissione di un fotone può invece avvenire se il gap è diretto, cioè se il vettore $\mathbf{k}$ dello stato più basso della banda di conduzione è uguale al vettore $\mathbf{k}$ dello stato più alto della banda di valenza. In questo caso la ricombinazione elettrone-buca può avvenire con emissione di un fotone. Questo è ciò che avviene nei *diodi a emissione luminosa*, o LED (Light Emitting Diode). Se il cristallo viene strutturato in modo da formare una cavità risonante il diodo funziona come laser.

In un laser a semiconduttore il generatore esterno alimenta sia un flusso di elettroni dallo strato n che un flusso di buche dallo strato p sempre verso la giunzione np . Questi flussi rimpiazzano le coppie elettrone-buca man mano che si ricombinano in seguito all'emissione dei fotoni. H. Kroemer e Z. I. Al'fërov (Жорес Иванович Алфёров) hanno proposto indipendentemente un importante miglioramento del laser a semiconduttore. Anziché due, vengono usati tre strati successivi di semiconduttori diversi (Figg. 9.24 e 9.25): uno strato di semiconduttore estrinseco drogato p , uno strato di semiconduttore intrinseco (non drogato) ed infine un uno strato di semiconduttore estrinseco drogato n , avendo così due eterogiunzioni. Gli strati p ed n hanno la stessa ampiezza del gap di energia proibita, mentre l'ampiezza del gap di energia proibita dello strato intrinseco i è minore: $\Delta E_p = \Delta E_n > \Delta E_i$. Questo ha come conseguenza sia il confinamento dell'inversione di popolazione in una zona più ampia che nel laser a singola giunzione, che un ulteriore vantaggio dovuto al fatto che l'indice di rifrazione dello strato i , in cui avviene il processo laser, è maggiore che in p ed in n . Per l'emissione laser si ha così un ef-

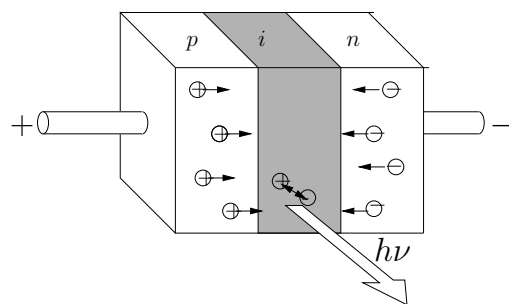


Figura 9.24 Laser a doppia giunzione pin .

fetto guida d'onda. Quando viene applicata la differenza di potenziale alle estremità del “sandwich”, le buche dello strato p e gli elettroni dello strato n possono penetrare nello strato i , dove si forma l'inversione di popolazione ed avvengono i fenomeni di ricombinazione con emissione di fotoni.

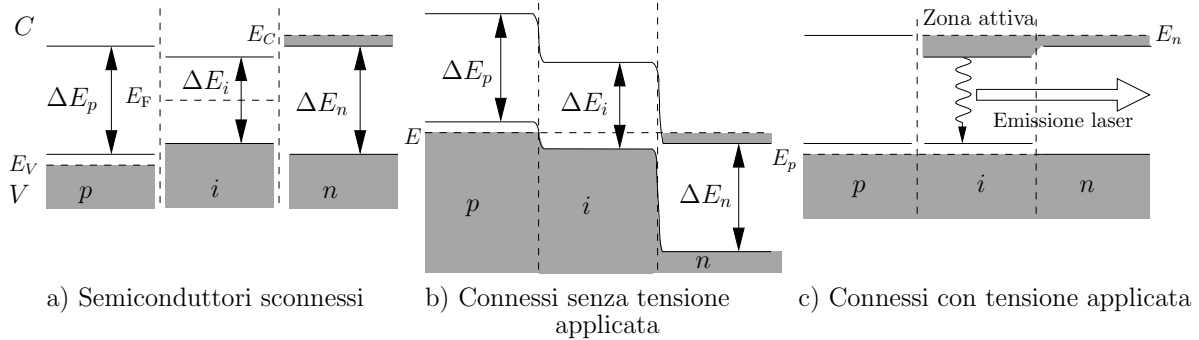


Figura 9.25 Laser a doppia giunzione *pin*. a) Semiconduttori sconnessi, posizione delle energie di Fermi nei tre strati p , i e n . b) Semiconduttori connessi senza tensione applicata: in ogni strato i livelli energetici si spostano in modo che le energie di Fermi siano allineate. c:) La tensione esterna applicata sposta i livelli, rendendo possibile il passaggio degli elettroni in banda di conduzione dalla zona n alla zona i , e delle buche in banda di valenza dalla zona p alla zona i . Nella zona i elettroni e buche si ricombinano, dando luogo all'emissione laser.

9.9 Laser a buca di potenziale e laser a cascata quantica

Nel 1972 Charles H. Henry si rese conto che il “sandwich” di semiconduttori proposto da Kroemer e Alfërov, detto anche *doppia giunzione*, *doppia eterostruttura*, o, appunto, *struttura a sandwich*, riprodotto in Fig. 9.26, oltre a costituire una guida d'onda per le onde luminose, poteva anche essere usato come *buca di potenziale finita* per gli elettroni nella banda di conduzione. Esattamente come nel laser a doppia giunzione *pin* di Kroemer e Alfërov, i semiconduttori sono scelti in modo che il gap di energie proibite del semiconduttore B sia più piccolo del gap del semiconduttore A . Per lo strato A si può usare, per esempio, arseniuro di alluminio AlAs , o arseniuro di alluminio-gallio $\text{Al}_x\text{Ga}_{1-x}\text{As}$, mentre per

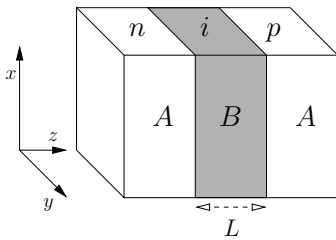


Figura 9.26 Semiconduttori in doppia eterostruttura.

lo strato B si può usare, sempre per esempio, arseniuro di gallio GaAs . Il semiconduttore B è intrinseco (non drogato), mentre i semiconduttori dei due strati A sono drogati uno p e l'altro n . In queste condizioni, una volta allineate le energie di Fermi dei tre strati, gli elettroni che si propagano in direzione $\pm z$ nella banda di conduzione vedono una buca di potenziale di ampiezza L , pari allo spessore dello strato B , e di una certa profondità U_0 , come mostrato in Fig. 9.27. I livelli energetici relativi al moto degli elettroni lungo z coincidono con quelli di una buca di potenziale finita, a parte la differenza, in realtà trascurabile, dovuta al fatto che la massa efficace dell'elettrone, m_{eff} , nel mezzo A è diversa dalla massa efficace nel mezzo B . L'equazione di Schrödinger per la buca di potenziale finita non ha soluzione analitica, perché le equazioni di raccordo tra la funzione d'onda e la sua derivata prima all'interno e all'esterno della buca sono equazioni trascendenti. L'equazione di Schrödinger può essere però risolta grafi-

camente o numericamente. Sappiamo comunque che esiste sempre almeno uno stato legato, e che il numero complessivo N di stati legati di una buca di potenziale finita vale

$$N = \left\lceil \frac{2\sqrt{m_e L^2 U_0}}{\pi \hbar} \right\rceil, \quad (9.122)$$

dove $[x]$ indica, come al solito, la parte intera superiore del numero reale x . Il numero di stati legati dipende quindi dall'ampiezza L e dalla profondità U_0 della buca. Le energie degli stati legati di una buca finita, E_n^{fin} , con $1 \leq n < N$, sono più basse di quelle dei corrispondenti livelli di una buca infinita di uguale ampiezza, E_n^{inf} , che possiamo calcolare analiticamente. Abbiamo cioè

$$E_n^{\text{fin}} < E_n^{\text{inf}} = \frac{n^2 \pi^2 \hbar^2}{2m_{\text{eff}} L^2} \quad \text{per ogni } n < N. \quad (9.123)$$

Ma per N grande ed n piccolo E_n^{inf} costituisce una buona approssimazione per eccesso di E_n^{fin} . La spaziatura tra i livelli dipende così dallo spessore L dello strato di semiconduttore intrinseco B della doppia eterostruttura. Naturalmente il laser può funzionare per decadimento degli elettroni da uno di questi stati a uno stato della banda di valenza (transizione *interbanda*). Ma, se ci sono almeno due livelli nella buca di potenziale, è possibile che, una volta applicata un'opportuna differenza di potenziale tra gli estremi del laser, gli elettroni passano dallo strato A drogato n al più alto di questi due livelli. Si ha così un'inversione di popolazione ed è possibile avere effetto laser all'interno della banda di conduzione in B (transizione *intra-band*), anziché per ricombinazione elettrone-buca (e quindi con salto dell'elettrone dalla banda di conduzione alla banda di valenza), come avviene nei laser a semiconduttore considerati finora. Il livello inferiore della transizione laser può essere vuotato per passaggio dell'elettrone allo strato A drogato p , dove si ricombina con una buca. Un primo vantaggio è che in questo laser a buca di potenziale, o *laser a pozzo quantico* (in inglese *quantum well laser*), la frequenza di emissione viene determinata scegliendo un opportuno spessore L dello strato B , e non si è costretti a cercare materiali semiconduttori con opportuni gap di energie proibite. La 9.123 ci dice che, per avere una separazione in frequenza superiore al THz tra livelli successivi della buca, lo strato B deve avere uno spessore L estremamente sottile, inferiore ai 10 nm. Questo naturalmente significa che la Fig. 9.26 non è assolutamente in scala. Oggi vengono prodotti laser a buca di potenziale operanti dai THz (intra-band) fino all'ultravioletto (interbanda). I laser a buca di potenziale hanno una corrente di soglia molto minore dei corrispondenti laser a doppia eterostruttura convenzionali, e un'efficienza quantica (o rendimento quantico: numero di fotoni emessi/numero di elettroni in ingresso) molto più alta.

Nel 1994 Jerome Faist, Federico Capasso, Deborah Sivco, Carlo Sirtori, Albert Hutchinson e Alfred Cho, dei laboratori Bell, hanno proposto un laser di ancor maggiore efficienza basato su una sequenza di eterogiunzioni di semiconduttori del tipo $ABABA \dots BA$ (questa struttura è detta anche *superreticolo*, in inglese *superlattice*), sempre con $\Delta E_B < \Delta E_A$. Questo laser è detto *laser a cascata quantica*. Come nel laser a buca di potenziale le singole transizioni laser-attive avvengono tra due

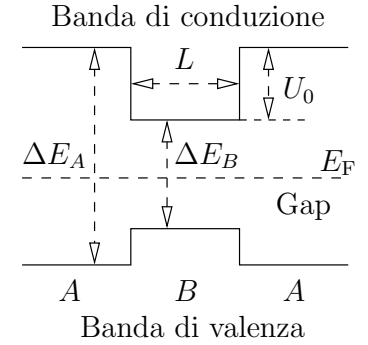


Figura 9.27 Buca di potenziale per gli elettroni della banda di conduzione nel "sandwich" di semiconduttori della Fig. 9.26.

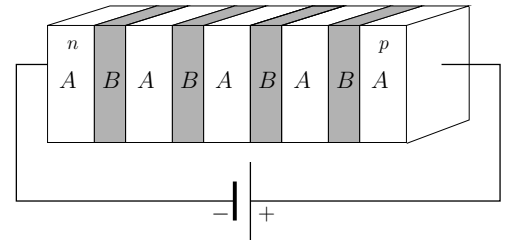


Figura 9.28 Schema di principio di una sequenza di eterogiunzioni (superreticolo) per un laser a cascata quantica.

livelli di una delle buche di potenziale nella banda di conduzione (cioè all'interno dei semiconduttori

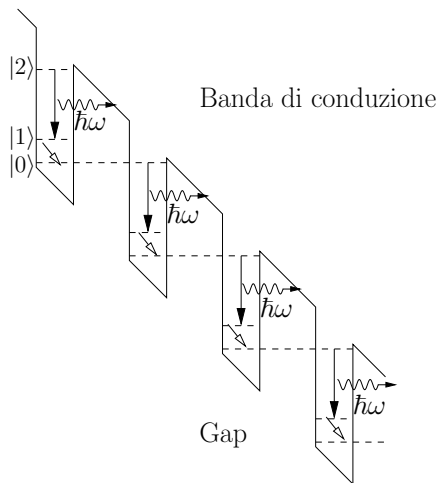


Figura 9.29 Schema di principio di un laser a cascata quantica. Un elettrone viene iniettato nel livello $|2\rangle$ di una buca, per poi passare per emissione indotta al livello $|1\rangle$. Dal livello $|1\rangle$ per transizione fononica passa al livello $|0\rangle$, e da qui, per effetto tunnel, al livello $|2\rangle$ della buca successiva.

che sostituiscono le singole buche della Fig. 9.29, e *zone di iniezione*, che sostituiscono le singole barriere della stessa figura. Ognuna di queste zone è costituita da una successione di buche e barriere

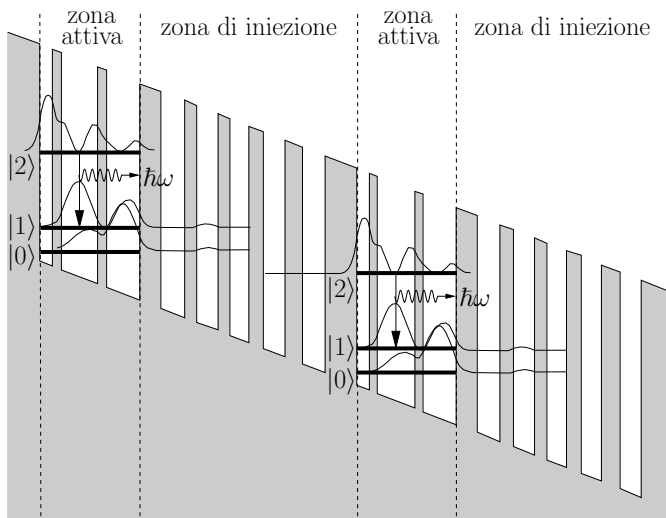


Figura 9.30 Realizzazione pratica di un laser a cascata quantica. Ogni singola zona attiva è costituita da tre buche di potenziale contigue, separate da sottili barriere di potenziale.

re, con le dimensioni scelte in modo che i moduli quadri delle tre funzioni d'onda abbiano i loro massimi ognuno in una diversa delle tre buche. Per aumentare la probabilità di transizione per emis-

di tipo *B*). Applicando una differenza di potenziale tra gli estremi della struttura, come in Fig. 9.28, si spostano i livelli energetici di tutti gli strati, e di conseguenza i fondi delle loro bande di conduzione. Il risultato è la "scala" di buche di potenziale schematizzata in Fig. 9.29. Un elettrone proveniente dallo strato *A* drogato *n* passa allo stato laser-attivo superiore $|2\rangle$ della buca più alta (vedi Fig. 9.29), effettua la transizione laser verso lo stato $|1\rangle$, e questo viene vuotato verso lo stato $|0\rangle$ per emissione di un fonone. Dallo stato $|0\rangle$ l'elettrone passa, per effetto tunnel attraverso uno stato *A* intrinseco, allo stato $|2\rangle$ della buca di potenziale successiva. Il processo viene iterato fino all'ultima buca. Da qui l'elettrone passa all'ultimo stato *A*, drogato *p*, dove si ricombina con una buca (non di potenziale!). Questo è lo schema in linea di principio. La realizzazione pratica è abbastanza più complicata. Bisogna infatti evitare che l'elettrone abbandoni lo stato $|2\rangle$ per emissione di fonone, per effetto tunnel o in qualunque altro modo, prima di effettuare la transizione laser verso lo stato $|1\rangle$. Inoltre bisogna ottenere che lo stato $|1\rangle$ si svuoti con rapidità sufficiente a mantener l'inversione di popolazione. Questo viene ottenuto alternando *zone attive*,

che sostituiscono le singole buche della Fig. 9.29, e *zone di iniezione*, che sostituiscono le singole barriere della stessa figura. Ognuna di queste zone è costituita da una successione di buche e barriere di potenziale di ampiezza opportuna, come mostrato in Fig. 9.30 (adattata dalla figura *Evans82* della Wikipedia in lingua inglese). Gli spessori dei vari strati della singola zona attiva, e quindi le ampiezze delle sue buche e barriere, determinano l'andamento spaziale delle funzioni d'onda degli stati $|1\rangle$, $|2\rangle$ e $|3\rangle$, e di conseguenza le loro probabilità di transizione non radiativa. Le probabilità di transizione non radiativa, infatti, sono fortemente dipendenti dalle sovrapposizioni delle funzioni d'onda e dalle distanze tra i livelli. La figura 9.30 mostra l'andamento dei moduli quadri $|\psi_i|^2$ delle funzioni d'onda degli stati coinvolti nel ciclo laser per una particolare alternanza di zone attive e zone di iniezione. Ogni zona attiva è qui costituita da tre buche di potenziale separate da sottili barriere

sione di fonone dallo stato $|1\rangle$ allo stato $|0\rangle$ si fa in modo che le loro funzioni d'onda abbiano una buona sovrapposizione. Inoltre, la distanza tra i livelli $|1\rangle$ e $|0\rangle$ è in risonanza con la frequenza del fonone longitudinale del reticolo, in modo che il livello $|1\rangle$ venga vuotato rapidamente. Gli spessori degli strati della zona di iniezione sono invece studiati in modo da rendere molto probabile il passaggio per effetto tunnel dell'elettrone nello stato $|0\rangle$ di una buca allo stato $|2\rangle$ della buca successiva, e poco probabile l'uscita da una buca di un elettrone nello stato $|2\rangle$.

Una caratteristica importante del laser a cascata quantica è che la sua efficienza quantica è maggiore di 1, dato che ogni elettrone iniettato genera l'emissione di più fotoni. Al momento l'emissione dei laser a cascata quantica copre la gamma di lunghezza d'onda compresa tra $2.63\text{ }\mu\text{m}$ e $250\text{ }\mu\text{m}$.

Capitolo 10

Interazione radiazione-materia

10.1 Perturbazioni dipendenti dal tempo

10.1.1 Introduzione

Supponiamo di avere un sistema quantistico limitato spazialmente e isolato, descritto da una hamiltoniana indipendente dal tempo $\hat{H}_0$. Gli autostati di $\hat{H}_0$ sono numerabili e formano una base ortonormale e completa $\{|n\rangle\}$ dello spazio di Hilbert relativo al nostro sistema. Ad ogni autostato $|n\rangle$ corrispondono un autovalore dell'energia E_n ed un'autofunzione d'onda $\varphi_n(q_1, \dots, q_f, t)$. Nella rappresentazione di Schrödinger, in cui gli operatori sono costanti mentre gli stati evolvono nel tempo, fintanto che il sistema rimane isolato gli autostati $|n\rangle$ sono stazionari, cioè dipendono dal tempo secondo la relazione $|n(t)\rangle = e^{-iE_n t/\hbar}|n(0)\rangle$.

Adesso supponiamo che il nostro sistema interagisca con un sistema esterno per un certo intervallo di tempo τ . L'hamiltoniana completa del sistema diventa

$$\hat{H} = \hat{H}_0 + \hat{V}(t), \quad \text{con} \quad \hat{V}(t) = \begin{cases} \hat{F}(t) & \text{se } 0 \leq t \leq \tau \\ 0 & \text{se } t < 0 \text{ oppure } t > \tau, \end{cases} \quad (10.1)$$

dove la forma di $\hat{F}(t)$ dipende dalla particolare interazione con il particolare sistema esterno. L'operatore $\hat{V}(t)$ si chiama *hamiltoniana di perturbazione*, mentre l'operatore $\hat{H}_0$ si chiama *hamiltoniana imperturbata*. Qui e in quanto segue faremo l'ipotesi che gli elementi diagonali di $\hat{F}$, e quindi di $\hat{V}$, siano nulli. L'equazione di Schrödinger dipendente dal tempo diventa

$$i\hbar \frac{\partial}{\partial t} |\psi(t)\rangle = [\hat{H}_0 + \hat{V}(t)] |\psi(t)\rangle, \quad (10.2)$$

e non ha soluzioni stazionarie. Ma, naturalmente, gli stati $\{|n\rangle\}$ continuano a costituire una base ortonormale e completa per il nostro spazio di Hilbert. Nei casi più semplici $\hat{F}(t)$ dipende dal tempo attraverso un parametro esterno, come la distanza interatomica durante una collisione, o l'intensità e la frequenza di un campo oscillante nel caso dell'interazione con un'onda elettromagnetica.

10.1.2 Probabilità di transizione

Per risolvere la (10.2) sfruttiamo il fatto che gli autostati dell'hamiltoniana imperturbata continuano a costituire una base completa anche in presenza della perturbazione. Quindi, qualunque sia lo stato

$|\psi(t)\rangle$ del sistema, possiamo sempre scrivere

$$|\psi(t)\rangle = \sum_n a_n(t) e^{-iE_n t/\hbar} |n\rangle, \quad (10.3)$$

dove le $a_n(t)$ sono delle funzioni complesse del tempo da determinare caso per caso, dette *ampiezze di probabilità*. Senza perdita di generalità, da ogni ampiezza abbiamo estratto in forma esplicita il fattore $e^{-iE_n t/\hbar}$ perché così, in assenza di perturbazione, le $a_n(t)$ sono costanti nel tempo. Adesso facciamo l'ulteriore ipotesi che per $t < 0$, cioè prima che venga "accesa" la perturbazione, il sistema si trovi nello stato stazionario $|m\rangle$ con energia E_m . Quindi, per $t < 0$, lo stato iniziale $|\psi_i(t)\rangle$ si scrive, secondo la (10.3),

$$|\psi_i(t)\rangle = e^{-iE_m t/\hbar} |m\rangle, \quad \text{con} \quad a_n(t) = \delta_{mn}. \quad (10.4)$$

Dopo che l'interazione è stata spenta, cioè per $t \geq \tau$, i coefficienti a_n in generale sono diversi da δ_{mn} , ma sono comunque nuovamente indipendenti dal tempo. Ognuno di essi avrà un valore $a_n(\tau)$, dipendente dallo stato iniziale $|m\rangle$, dal tipo di perturbazione $\hat{F}(t)$ applicata, e dal tempo τ durante il quale la perturbazione ha agito. Lo stato finale del sistema si scrive così

$$|\psi_f(t)\rangle = \sum_n a_n(\tau) e^{-iE_n t/\hbar} |n\rangle, \quad t \geq \tau. \quad (10.5)$$

La probabilità di trovare il sistema nello stato $|n\rangle$ vale

$$P_n(\tau) = |a_n(\tau)|^2, \quad (10.6)$$

quindi la quantità $|a_n(\tau)|^2$ è la probabilità che, nel tempo τ , il sistema abbia fatto la transizione dallo stato iniziale $|m\rangle$ allo stato finale $|n\rangle$. Per determinare i coefficienti $a_n(\tau)$ sostituiamo lo sviluppo (10.3) nella (10.2), e svolgiamo i passaggi che seguono

$$\begin{aligned} i\hbar \frac{\partial}{\partial t} \sum_r a_r e^{-iE_r t/\hbar} |r\rangle &= [\hat{H}_0 + \hat{V}(t)] \sum_n a_n e^{-iE_n t/\hbar} |n\rangle \\ i\hbar \sum_r \left[\frac{\partial a_r}{\partial t} - i \frac{E_r}{\hbar} a_r \right] e^{-iE_r t/\hbar} |r\rangle &= \sum_n a_n e^{-iE_n t/\hbar} E_n |n\rangle + \sum_n a_n e^{-iE_n t/\hbar} \hat{V}(t) |n\rangle \\ i\hbar \sum_r e^{-iE_r t/\hbar} \frac{\partial a_r}{\partial t} |r\rangle + \sum_r a_r e^{-iE_r t/\hbar} E_r |r\rangle &= \sum_n a_n e^{-iE_n t/\hbar} E_n |n\rangle + \sum_n a_n e^{-iE_n t/\hbar} \hat{V}(t) |n\rangle \\ i\hbar \sum_r e^{-iE_r t/\hbar} \frac{\partial a_r}{\partial t} |r\rangle &= \sum_n a_n e^{-iE_n t/\hbar} \hat{V}(t) |n\rangle. \end{aligned} \quad (10.7)$$

Nell'ultima delle (10.7) moltiplichiamo primo e secondo membro a sinistra per $-\frac{i}{\hbar} \langle k| e^{iE_k t/\hbar}$, dove $|k\rangle$ è un qualunque autostato di $\hat{H}_0$, ottenendo

$$\frac{\partial a_k}{\partial t} = \frac{1}{i\hbar} \sum_n a_n e^{-i(E_n - E_k)t/\hbar} \langle k| \hat{V}(t) |n\rangle. \quad (10.8)$$

Fino a questo punto non abbiamo fatto approssimazioni. E' possibile, in linea di principio, integrare le equazioni (10.8) e giungere ad una soluzione esatta. Questo si può fare facilmente se ci sono solo due livelli energetici in gioco, e la soluzione così ottenuta è utile per trattare problemi simili, per esempio, a quello dell'inversione dell'ammoniaca. Nel paragrafo 10.5 vedremo invece il caso di un sistema atomico a due livelli in *interazione forte* con il campo elettromagnetico. Ma nel caso generale si preferisce procedere per approssimazioni successive, con un metodo detto *perturbativo*, in cui si pone

$$a_n(\tau) = a_n^{(0)} + a_n^{(1)}(\tau) + a_n^{(2)}(\tau) + \dots, \quad \text{con } a_n^{(0)} = \delta_{mn} \quad \text{per la (10.4).}$$

Perché il metodo sia applicabile è necessario che gli $a_n^{(r)}(\tau)$ decrescano al crescere di r . Tenendo conto dell'ipotesi (10.4), conviene separare il termine inizialmente non nullo a_m dal resto della sommatoria della (10.8), scrivendo

$$\frac{\partial a_k}{\partial t} = \frac{1}{i\hbar} \left[a_m e^{-i(E_m - E_k)t/\hbar} \langle k | \hat{V}(t) | m \rangle + \sum_{n \neq m} a_n e^{-i(E_n - E_k)t/\hbar} \langle k | \hat{V}(t) | n \rangle \right]. \quad (10.9)$$

Se, come ipotizzato, la perturbazione è piccola, gli $a_n(t)$ sono poco diversi dai corrispondenti $a_n(0)$, quindi tutti gli $a_n(t)$ con $n \neq m$ sono piccoli rispetto a $a_m(t) \simeq 1$. Questo ci porta a trascurare, come prima approssimazione, la sommatoria all'interno della parentesi quadra, e ci rimane

$$\frac{\partial a_k}{\partial t} = \frac{1}{i\hbar} e^{-i(E_m - E_k)t/\hbar} \langle k | \hat{V}(t) | m \rangle, \quad (10.10)$$

che può essere integrata, ottenendo per il valore al primo ordine dell'ampiezza $a_k^{(1)}(\tau)$ alla fine dell'interazione

$$a_k^{(1)}(\tau) = \frac{1}{i\hbar} \int_0^\tau dt e^{-i(E_m - E_k)t/\hbar} \langle k | \hat{V}(t) | m \rangle. \quad (10.11)$$

E' da notare che $a_m^{(1)}(\tau) = 0$, perché abbiamo fatto l'ipotesi che $\hat{V}(t)$ sia diagonale su $\{|n\rangle\}$, quindi $\langle m | \hat{V}(t) | m \rangle = 0$. Se si desidera un'approssimazione migliore di quella al primo ordine, il secondo ordine si ottiene non trascurando più la sommatoria della (10.9), ma sostituendo in essa gli $a_n(t)$ con i loro valori al primo ordine $a_n^{(1)}(t)$

$$\frac{\partial a_k}{\partial t} = \frac{1}{i\hbar} \left[a_m^{(0)} e^{-i(E_m - E_k)t/\hbar} \langle k | \hat{V}(t) | m \rangle + \sum_{n \neq m} a_n^{(1)} e^{-i(E_n - E_k)t/\hbar} \langle k | \hat{V}(t) | n \rangle \right] \quad (10.12)$$

Integrando, il primo termine dentro alla parentesi quadra ci ridà l'ampiezza al primo ordine $a_k^{(1)}(\tau)$, mentre il secondo termine (la sommatoria) ci dà l'ampiezza al secondo ordine $a_k^{(2)}(\tau)$

$$a_k^{(2)}(\tau) = \frac{1}{i\hbar} \sum_{n \neq m} \int_0^\tau a_n^{(1)}(t) e^{-i(E_n - E_k)t/\hbar} \langle k | \hat{V}(t) | n \rangle dt. \quad (10.13)$$

Sostituendo il valore (10.11) per $a_n^{(1)}(t)$ nella (10.13) otteniamo

$$a_k^{(2)}(\tau) = -\frac{1}{\hbar^2} \sum_{n \neq m} \int_0^\tau dt e^{-i(E_n - E_k)t/\hbar} \langle k | \hat{V}(t) | n \rangle \int_0^t dt' e^{-i(E_m - E_n)t'/\hbar} \langle n | \hat{V}(t') | m \rangle. \quad (10.14)$$

Le approssimazioni di ordine superiore possono essere ottenute iterativamente, sostituendo le ampiezza calcolate fino all'ordine r nella (10.9) ed integrando per ottenere l'ordine $r + 1$. E' da notare

che la (10.14) descrive effettivamente un processo al secondo ordine, con il sistema che passa prima dallo stato iniziale $|m\rangle$ ad uno stato intermedio $|n\rangle$, per poi passare da $|n\rangle$ allo stato finale $|k\rangle$. Se gli elementi di matrice $\langle f|\hat{V}|i\rangle$ sono piccoli, la (10.11) ci dà un'espressione al primo ordine negli elementi di matrice di $\hat{V}$, mentre la (10.14) è al secondo ordine.

10.1.3 Perturbazioni periodiche

Un caso di particolare interesse è quello delle perturbazioni periodiche nel tempo, che possono essere usate per descrivere l'interazione di un sistema quantistico con un'onda elettromagnetica monocromatica. L'hamiltoniana di perturbazione ha la forma

$$\hat{V}(t) = \begin{cases} \hat{F}^+ e^{-i\omega t} + \hat{F}^- e^{+i\omega t} & \text{se } 0 \leq t \leq \tau \\ 0 & \text{se } t < 0 \text{ oppure } t > \tau, \end{cases} \quad (10.15)$$

dove $\hat{F}^+$ e $\hat{F}^-$ sono matrici costanti nel tempo, e, senza perdita di generalità, supponiamo che sia $\omega > 0$. L'hermiticità dell'hamiltoniana impone che sia

$$\hat{F}^- = (\hat{F}^+)^\dagger, \quad \text{corrispondente a } \langle a|\hat{F}^-|b\rangle = \langle b|\hat{F}^+|a\rangle^* \quad \forall a, b. \quad (10.16)$$

Come sempre, facciamo l'ipotesi che prima dell'accensione della perturbazione ($t < 0$) il sistema si trovi in un certo autostato $|m\rangle$ di $\hat{H}_0$. Se nella (10.11) poniamo $\omega_{km} = (E_k - E_m)/\hbar$ e sostituiamo la (10.15) per $\hat{V}$ nell'integrale a secondo membro, otteniamo

$$\begin{aligned} a_k^{(1)}(\tau) &= \frac{1}{i\hbar} \int_0^\tau dt e^{i\omega_{km}t} \langle k|\hat{F}^+ e^{-i\omega t} + \hat{F}^- e^{+i\omega t}|m\rangle \\ &= \frac{1}{i\hbar} \left[\langle k|\hat{F}^+|m\rangle \int_0^\tau dt e^{i(\omega_{km}-\omega)t} + \langle k|\hat{F}^-|m\rangle \int_0^\tau dt e^{i(\omega_{km}+\omega)t} \right] \\ &= -\frac{1}{\hbar} \left[\langle k|\hat{F}^+|m\rangle \frac{e^{i(\omega_{km}-\omega)\tau} - 1}{\omega_{km} - \omega} + \langle k|\hat{F}^-|m\rangle \frac{e^{i(\omega_{km}+\omega)\tau} - 1}{\omega_{km} + \omega} \right]. \end{aligned} \quad (10.17)$$

La probabilità di transizione al primo ordine si ottiene calcolando il modulo quadro della (10.17). Se la frequenza (angolare) della perturbazione ω è molto diversa da ω_{km} , i due denominatori dell'ultima delle (10.17) sono grandi e, di conseguenza, le ampiezze di probabilità $a_k^{(1)}$ sono piccole e non si osservano transizioni. I casi di interesse sono allora quelli in cui la frequenza della transizione ω è vicina a $\pm\omega_{km}$, cioè $\omega \simeq \omega_{km}$ oppure $\omega \simeq -\omega_{km}$. Consideriamo per esteso solo il primo caso, dato che la trattazione del secondo è del tutto analoga. Qui $E_k > E_m$, cioè lo stato finale $|k\rangle$ ha energia maggiore dello stato iniziale $|m\rangle$, e $\omega \simeq (E_k - E_m)/\hbar$. Essendo $|\omega_{km} - \omega|$ molto piccolo, il primo termine tra parentesi quadra dell'ultima delle (10.17) è dominante rispetto al secondo, che può essere trascurato: questa è l'*approssimazione di onda rotante*. In questa approssimazione possiamo scrivere

$$a_k^{(1)}(\tau) \simeq -\frac{1}{\hbar} \langle k|\hat{F}^+|m\rangle \frac{e^{i(\omega_{km}-\omega)\tau} - 1}{\omega_{km} - \omega}. \quad (10.18)$$

Corrispondentemente avremo per la probabilità di transizione dallo stato iniziale $|m\rangle$ allo stato finale $|k\rangle$

$$\begin{aligned} P_{km}(\tau) &= |a_k^{(1)}(\tau)|^2 \simeq \frac{1}{\hbar^2} |\langle k|\hat{F}^+|m\rangle|^2 \left| \frac{e^{i(\omega_{km}-\omega)\tau} - 1}{\omega_{km} - \omega} \right|^2 \\ &= \frac{2}{\hbar^2} |\langle k|\hat{F}^+|m\rangle|^2 \frac{1 - \cos(\omega_{km} - \omega)\tau}{(\omega_{km} - \omega)^2}. \end{aligned} \quad (10.19)$$

Ponendo per “stenografia”

$$\gamma = \frac{\omega_{km} - \omega}{2} = \frac{E_k - E_m - \hbar\omega}{2\hbar}, \quad (10.20)$$

il numeratore dell’ultima frazione della (10.19) diventa

$$\begin{aligned} 1 - \cos(\omega_{km} - \omega)\tau &= 1 - \cos(2\gamma\tau) = \cos^2(\gamma\tau) + \sin^2(\gamma\tau) - \cos(2\gamma\tau) \\ &= \cos^2(\gamma\tau) + \sin^2(\gamma\tau) - \cos^2(\gamma\tau) + \sin^2(\gamma\tau) \\ &= 2\sin^2(\gamma\tau). \end{aligned} \quad (10.21)$$

Sostituendo nella (10.19) otteniamo

$$\begin{aligned} P_{km} &= \frac{2}{\hbar^2} |\langle k|\hat{F}^+|m\rangle|^2 \frac{2\sin^2(\gamma\tau)}{(\omega_{km} - \omega)^2} = \frac{1}{\hbar^2} |\langle k|\hat{F}^+|m\rangle|^2 \frac{4\sin^2(\gamma\tau)}{4\gamma^2} \\ &= \frac{1}{\hbar^2} |\langle k|\hat{F}^+|m\rangle|^2 \frac{\sin^2(\gamma\tau)}{\gamma^2} = \frac{1}{\hbar^2} |\langle k|\hat{F}^+|m\rangle|^2 \tau^2 \frac{\sin^2\left(\pi \frac{\gamma\tau}{\pi}\right)}{\frac{\pi^2 \gamma^2 \tau^2}{\pi^2}} \\ &= \frac{1}{\hbar^2} |\langle k|\hat{F}^+|m\rangle|^2 \tau^2 \operatorname{sinc}^2\left(\frac{\gamma\tau}{\pi}\right) = \frac{1}{\hbar^2} |\langle k|\hat{F}^+|m\rangle|^2 \tau^2 \operatorname{sinc}^2\left[\frac{(E_k - E_m - \hbar\omega)\tau}{2\pi\hbar}\right] \\ &= \frac{1}{\hbar^2} |\langle k|\hat{F}^+|m\rangle|^2 \tau^2 \operatorname{sinc}^2\left(\frac{E_k - E_m - \hbar\omega}{h} \tau\right), \end{aligned} \quad (10.22)$$

dove la funzione $\operatorname{sinc}^2(x)$, il cui grafico è rappresentato in fig. 10.1, è il quadrato della funzione $\operatorname{sinc}(x) = \sin \pi x / (\pi x)$ che abbiamo incontrato nel paragrafo 8.4. La funzione è simmetrica e sempre positiva, e ha un massimo per $x = 0$, dove vale 1. Si annulla per tutti i valori interi di x , tranne $x = 0$. La regione compresa tra il massimo e il primo zero ha ampiezza

$$\Delta E = \frac{h}{\tau}, \quad (10.23)$$

per cui, se il tempo di interazione τ è molto lungo, l’ampiezza ΔE è molto stretta. Dall’integrale

$$\int_0^{+\infty} \frac{\sin^2(ax)}{x^2} dx = \frac{\pi a}{2} \quad \text{ricaviamo} \quad \int_{-\infty}^{+\infty} \operatorname{sinc}^2(ax) dx = \frac{1}{a}, \quad (10.24)$$

da cui abbiamo

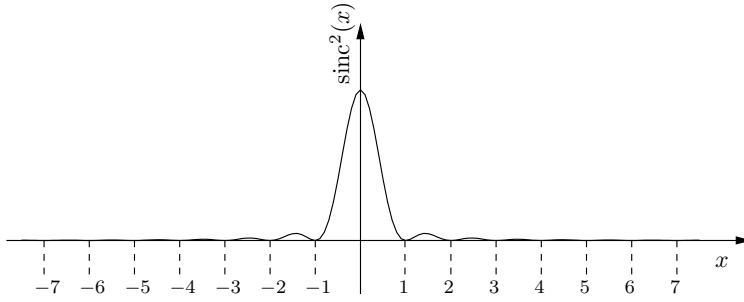


Figura 10.1 La funzione $\text{sinc}^2(x)$.

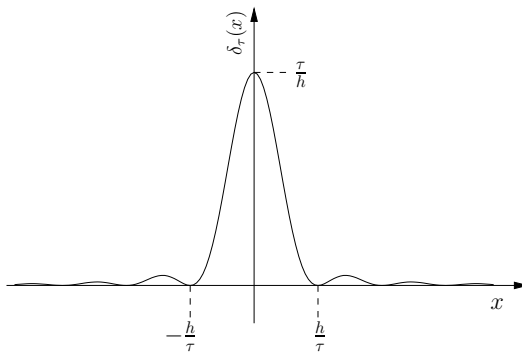


Figura 10.2 La funzione $\delta_\tau(x)$.

$$\frac{\tau}{h} \int_{-\infty}^{+\infty} \text{sinc}^2\left(x \frac{\tau}{h}\right) dx = 1,$$

da cui

$$\lim_{\tau \rightarrow \infty} \frac{\tau}{h} \text{sinc}^2\left(x \frac{\tau}{h}\right) = \delta(x), \quad (10.25)$$

dove $\delta(x)$ è la delta di Dirac. Adesso introduciamo una nuova funzione $\delta_\tau(x)$, definita da

$$\delta_\tau(x) = \frac{\tau}{h} \text{sinc}^2\left(x \frac{\tau}{h}\right), \quad (10.26)$$

il cui grafico è rappresentato in fig. 10.2. Per quanto visto sopra, la funzione $\delta_\tau(x)$ ha le proprietà che

$$\int_{-\infty}^{+\infty} \delta_\tau(x) dx = 1$$

e che

$$\lim_{\tau \rightarrow \infty} \delta_\tau(x) = \delta(x).$$

Invertendo la (10.26) abbiamo

$$\text{sinc}^2\left(x \frac{\tau}{h}\right) = \frac{h}{\tau} \delta_\tau(x), \quad (10.27)$$

che, inserita nella (10.22), ci dà per la probabilità di transizione tra gli stati $|m\rangle$ e $|k\rangle$

$$P_{km} = \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \tau \delta_\tau(E_k - E_m - \hbar\omega). \quad (10.28)$$

Quindi, partendo da uno stato iniziale $|m\rangle$ con energia E_m , si arriva ad uno stato finale $|k\rangle$ di energia $E_k = E_m + \hbar\omega$ entro i limiti della larghezza h/τ della funzione δ_τ . Questa larghezza corrisponde all'indeterminazione tempo-energia, legata alla durata finita dell'interazione. Se dividiamo la (10.28) per τ possiamo ottenere una *probabilità di transizione per unità di tempo* (in inglese *transition rate*)

$$W_{km} = \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \delta_\tau(E_k - E_m - \hbar\omega). \quad (10.29)$$

Ripetendo i conti nell'altro caso di interesse, cioè per $E_k < E_m$ e $\omega \simeq -\omega_{km}$, questa volta è $|\omega_{km} + \omega|$ ad essere molto piccolo rispetto a $|\omega_{km} - \omega|$, e possiamo quindi trascurare il primo termine nella parentesi quadra dell'ultima delle (10.17). Con passaggi del tutto analoghi ai precedenti otteniamo

$$P_{km}(\tau) \simeq \frac{2\pi}{\hbar} |\langle k|\hat{F}^-|m\rangle|^2 \tau \delta_\tau(E_m - E_k - \hbar\omega), \quad (10.30)$$

che ci dà la probabilità di transizione tra gli stati $|m\rangle$ e $|k\rangle$, con $E_k \simeq E_m - \hbar\omega$. E' da notare che, data la (10.16), abbiamo

$$|\langle k|\hat{F}^-|m\rangle|^2 = |\langle k|\hat{F}^+|m\rangle|^2. \quad (10.31)$$

Ad una perturbazione periodica di frequenza angolare ω sono quindi associati quanti di energia di valore $\hbar\omega = h\nu$. Nelle transizioni si ha conservazione dell'energia, entro l'indeterminazione tempo-energia, tenendo conto dell'assorbimento o dell'emissione di un quanto. E' importante notare che siamo giunti a questo risultato trattando come quantizzati i livelli energetici del nostro sistema materiale, ma senza quantizzare esplicitamente la radiazione, che abbiamo trattato classicamente. Questo è il cosiddetto *trattamento semiclassico* dell'interazione radiazione-materia.

Un semplice caso particolare delle perturbazioni periodiche è quello delle *perturbazioni costanti*, o *perturbazioni a gradino*. In questo caso l'hamiltoniana di perturbazione $\hat{V}$, nulla per $t < 0$ e per $t > \tau$, è rappresentata da un operatore costante nel tempo (cioè, con elementi di matrice costanti) $\hat{F}$, durante l'intervallo $0 \leq t \leq \tau$. Questo corrisponde a porre $\omega = 0$ nella (10.15), ottenendo

$$\hat{V}(t) = \begin{cases} \hat{F} & \text{se } 0 \leq t \leq \tau \\ 0 & \text{se } t < 0 \text{ oppure } t > \tau, \end{cases} \quad (10.32)$$

dove $\hat{F}$ è una matrice hermitiana costante con elementi diagonali nulli. Sempre con l'ipotesi che lo stato iniziale del sistema sia l'autostato $|m\rangle$ di $\hat{H}_0$, possiamo portare fuori dall'integrale della (10.11) l'elemento di matrice costante, ottenendo

$$a_k^{(1)}(\tau) = \frac{1}{i\hbar} \langle k|\hat{F}|m\rangle \int_0^\tau e^{i\omega_{km}t} dt = -\frac{1}{\hbar} \langle k|\hat{F}|m\rangle \frac{e^{i\omega_{km}\tau} - 1}{\omega_{km}}, \quad (10.33)$$

dove abbiamo ancora posto $\omega_{km} = (E_k - E_m)/\hbar$. Quindi, una volta spenta la perturbazione, la probabilità che il sistema sia passato dallo stato iniziale $|m\rangle$ allo stato finale $|k\rangle$ è

$$P_{km} = |a_k^{(1)}(\tau)|^2 = \frac{2}{\hbar^2} |\langle k|\hat{F}|m\rangle|^2 \frac{[1 - \cos(\omega_{km}\tau)]}{\omega_{km}^2}. \quad (10.34)$$

Con trattamento analogo a quello delle (10.21) e (10.22), e introducendo la δ_τ , otteniamo

$$P_{km} = \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \tau \delta_\tau(E_k - E_m), \quad W_{km} = \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \delta_\tau(E_k - E_m). \quad (10.35)$$

Anche nella (10.35) la funzione $\delta_\tau(E_k - E_m)$ esprime la conservazione dell'energia con una indeterminazione $\hbar/\tau$, legata alla durata finita dell'interazione. Se abbiamo una hamiltoniana di perturbazione costante nel tempo e di durata infinita, la conservazione dell'energia ci permette di raggiungere solo stati finali $|k\rangle$ che abbiano la stessa energia dello stato iniziale $|m\rangle$.

10.1.4 Transizioni in uno spettro continuo

Un sistema quantistico limitato spazialmente ha uno spettro discreto di autovalori dell'energia, mentre se il sistema non è limitato lo spettro degli autovalori dell'energia è continuo. Un'applicazione importante della teoria delle perturbazioni è il calcolo della probabilità di transizione tra stati dello spettro continuo. Gli stati dello spettro continuo sono quasi sempre degeneri, e questo rende possibile

la transizione tra uno stato iniziale ed uno stato finale della stessa energia dovuto ad una perturbazione costante. Un esempio è costituito da un fenomeno di scattering, in cui una particella cambi la direzione di propagazione conservando l'energia.

Un altro caso di interesse è quando una perturbazione periodica del tipo (10.15) induce una transizione da uno stato iniziale discreto $|m\rangle$ ad uno stato finale $|\alpha\rangle$ dello spettro continuo, con $E_\alpha > E_m$. Naturalmente, l'etichetta α varia con continuità. Analogamente alla (10.3) per il caso dello spettro discreto, uno stato generico $|\psi\rangle$ del sistema nel continuo può essere scritto

$$|\psi(t)\rangle = \int d\alpha a_\alpha(t) e^{-iE_\alpha t/\hbar} |\alpha\rangle, \quad (10.36)$$

dove gli $|\alpha\rangle$ costituiscono una base ortonormale completa nel continuo. La condizione di normalizzazione degli stati impone che sia

$$\int |a_\alpha|^2 d\alpha = 1. \quad (10.37)$$

La (10.17) si riscrive

$$a_\alpha^{(1)}(\tau) \simeq \frac{1}{i\hbar} \langle \alpha | \hat{F}^+ | m \rangle \frac{e^{i(\omega_{\alpha m} - \omega)\tau} - 1}{i(\omega_{\alpha m} - \omega)} \quad (10.38)$$

perché il termine contenente l'elemento di matrice $\langle \alpha | \hat{F}^- | m \rangle$ può essere trascurato, essendo $\omega_{\alpha m} = (E_\alpha - E_m)/\hbar > 0$. Inoltre alla transizione saranno interessati solo gli stati $|\alpha\rangle$ del continuo con energia $E_\alpha \simeq E_m + \hbar\omega$. Per la probabilità di transizione avremo

$$P_{\alpha m}(\tau) \simeq \frac{1}{\hbar^2} |\langle \alpha | \hat{F}^+ | m \rangle|^2 \tau^2 \text{sinc}^2 \left[(E_\alpha - E_m - \hbar\omega) \frac{\tau}{\hbar} \right], \quad (10.39)$$

che possiamo riscrivere

$$P_{\alpha m} = \frac{2\pi}{\hbar} |\langle \alpha | \hat{F}^+ | m \rangle|^2 \tau \delta_\tau(E_\alpha - E_m - \hbar\omega). \quad (10.40)$$

La quantità $dW_{\alpha m} = P_{\alpha m}(\tau) d\alpha/\tau$ esprime la probabilità di transizione per unità di tempo dallo stato iniziale $|m\rangle$ ad uno stato nell'intervallo $(\alpha, \alpha + d\alpha)$. Abbiamo

$$dW_{\alpha m} = \frac{2\pi}{\hbar} |\langle \alpha | \hat{F}^+ | m \rangle|^2 \delta_\tau(E_\alpha - E_m - \hbar\omega) d\alpha. \quad (10.41)$$

Come potevamo attenderci, al limite $\tau \rightarrow \infty$ la probabilità di transizione è zero tranne che per gli stati con energia esattamente uguale a $E_\alpha = E_m + \hbar\omega$.

10.1.5 Regola d'oro di Fermi

Spesso esiste una distribuzione di autostati dell'hamiltoniana imperturbata molto vicini in energia tra loro, e la transizione che ci interessa, indotta da una perturbazione periodica di pulsazione ω , avviene dallo stato iniziale $|m\rangle$ alla distribuzione di possibili stati finali. La situazione è quindi analoga a quanto visto nel paragrafo 10.1.4 sulle transizioni allo spettro continuo. Chiamiamo $\rho(E_k)$ la densità di distribuzione degli stati finali, tale che il numero $dn(E_k)$ di autostati di $\hat{H}_0$ con energia compresa tra E_k ed $E_k + dE_k$ sia

$$dn(E_k) = \rho(E_k) dE_k. \quad (10.42)$$

Facciamo l'ipotesi che la distribuzione $\rho(E_k)$ sia centrata attorno ad un certo valore centrale $\bar{E}_k$, con una certa larghezza $\Delta E_k = \hbar \Delta \omega_k$. Facciamo inoltre l'ipotesi che sia $\rho(E_k)$ che l'elemento di matrice $\langle k|\hat{V}|m\rangle$ varino lentamente attorno allo stato centrale della distribuzione. Quello che ci interessa è la probabilità di transizione *totale* dallo stato iniziale $|m\rangle$ all'insieme degli stati $|k'\rangle$ attorno a $|k\rangle$, che, usando la (10.28), vale

$$\bar{P}_{km} = \sum_{k'} P_{k'm} = \frac{2\pi}{\hbar} \sum_{k'} |\langle k'|\hat{F}|m\rangle|^2 \tau \delta_\tau(E_{k'} - E_m - \hbar\nu). \quad (10.43)$$

Sotto le nostre ipotesi, la somma può essere approssimata da un integrale in dE_k , ottenendo

$$\bar{P}_{km} = \frac{2\pi}{\hbar} \int |\langle k|\hat{F}|m\rangle|^2 \tau \delta_\tau(E_k - E_m - \hbar\omega) \rho(E_k) dE_k. \quad (10.44)$$

Se, invece di una perturbazione periodica, abbiamo una perturbazione costante accesa per un tempo τ , basta porre $\omega = 0$ nella (10.44). Se l'elemento di matrice $\langle k|\hat{F}|m\rangle$ non varia troppo rapidamente può essere portato fuori dall'integrale ottenendo

$$\bar{P}_{km} = \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \tau \int \delta_\tau(E_k - E_m - \hbar\omega) \rho(E_k) dE_k. \quad (10.45)$$

Esistono due casi estremi in cui il calcolo dell'integrale (10.45) è particolarmente facile:

- i. quando la larghezza ΔE_k della funzione di distribuzione $\rho(E_k)$ è molto maggiore della larghezza in energia della δ_τ ;
- ii. quando la larghezza ΔE_k della funzione di distribuzione $\rho(E_k)$ è molto minore della larghezza in energia della δ_τ .

Consideriamo prima il caso (i), schematizzato in Fig. 10.3. La larghezza in energia della δ_τ , cioè $\hbar/\tau$, è molto minore di ΔE_k , quindi

$$\frac{\hbar}{\tau} \ll \Delta E_k, \quad \text{da cui} \quad \tau \gg \frac{\hbar}{\Delta E_k} = \frac{1}{\Delta \nu_k} = \frac{2\pi}{\Delta \omega_k}. \quad (10.46)$$

Quindi, la perturbazione è accesa per un tempo lungo. In questo caso la densità degli stati è praticamente costante su tutto l'intervallo in cui la δ_τ è significativamente diversa da zero, e la $\rho(E_k)$ può essere calcolata sul valore centrale dell'energia $\bar{E}_k$ e portata fuori dall'integrale:

$$\begin{aligned} \bar{P}_{km} &= \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \rho(\bar{E}_k) \tau \int \delta_\tau(E_k - E_m - \hbar\omega) dE_k \\ &= \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \rho(\bar{E}_k) \tau \end{aligned} \quad (10.47)$$

dove abbiamo usato il fatto che l'intervallo di integrazione è sufficientemente largo perché l'integrale della δ_τ sia praticamente 1. Questa espressione per $\bar{P}_{km}$ è conosciuta sotto il nome di *regola d'oro di Fermi*. La probabilità di transizione $\bar{P}_{km}$ della (10.47) è proporzionale alla durata dell'interazione τ , ed è possibile definire una probabilità di transizione per unità di tempo $\bar{W}_{km}$

$$\bar{W}_{km} = \frac{\bar{P}_{km}}{\tau} = \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \rho(\bar{E}_k). \quad (10.48)$$

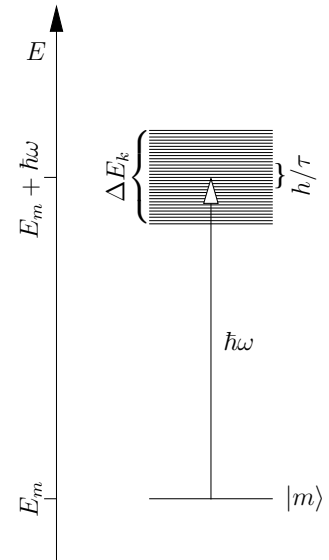
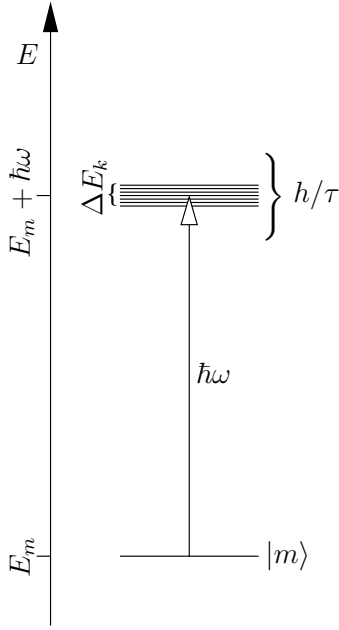


Figura 10.3 Applicabilità della regola d'oro di Fermi: la larghezza relativa all'indeterminazione energia-tempo, $\hbar/\tau$, è molto minore della distribuzione dei livelli ΔE_k



Il secondo caso in cui è facile il calcolo dell'integrale della (10.45) è quando la larghezza della distribuzione degli stati finali è molto stretta rispetto alla larghezza in energia della δ_τ , cioè quando

$$\frac{h}{\tau} \gg \Delta E_k, \quad \text{da cui} \quad \tau \ll \frac{h}{\Delta E_k} = \frac{1}{\Delta \nu_k} = \frac{2\pi}{\Delta \omega_k}. \quad (10.49)$$

In questo caso è la funzione δ_τ che può essere considerata costante, e pari al suo valore massimo $\delta_\tau(0) = \tau/h$, nell'intervallo di energia dove la densità degli stati è apprezzabilmente diversa da zero. La (10.45) si riscrive così

$$\bar{P}_{km} = \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \tau \int \rho(E_k) dE_k = \frac{2\pi}{\hbar} |\langle k|\hat{F}|m\rangle|^2 \frac{\tau^2}{h} N \quad (10.50)$$

dove N è il numero complessivo di stati finali. Ricordando che $h = 2\pi\hbar$ otteniamo finalmente

$$\bar{P}_{km} = \frac{\tau^2}{\hbar^2} |\langle k|\hat{F}|m\rangle|^2 N. \quad (10.51)$$

Figura 10.4 Caso opposto a quello della Fig. 10.3: qui la larghezza relativa all'indeterminazione energia-tempo, h/τ , è molto maggiore della distribuzione dei livelli ΔE_k

Questa espressione per la probabilità di transizione è così proporzionale a τ^2 e, data la condizione (10.49), è valida solo per tempi di interazione τ brevi.

10.2 Perturbazioni periodiche al secondo ordine

In presenza di una perturbazione periodica di frequenza ω del tipo (10.15) le ampiezze di probabilità di transizione al primo ordine dallo stato iniziale $|m\rangle$ ad uno stato finale $|n\rangle$ sono, secondo la (10.18),

$$a_n^{(1)}(t) \simeq \frac{1}{i\hbar} \langle n|\hat{F}^\pm|m\rangle \frac{e^{i(\omega_{nm} \mp \omega)t} - 1}{i(\omega_{nm} \mp \omega)}, \quad (10.52)$$

dove nei simboli $\pm$ e $\mp$ vale il segno superiore se $E_n > E_m$ e la frequenza ω della perturbazione è in prossimità di $(E_n - E_m)/\hbar$, il segno inferiore se $E_n < E_m$ e $\omega \simeq (E_m - E_n)/\hbar$. Abbiamo visto che per ω lontana da questi valori l'ampiezza $a_n^{(1)}(t)$ è trascurabile. Per ottenere l'ampiezza al secondo ordine sostituiamo la (10.52) nella (10.13) ottenendo

$$\begin{aligned} a_k^{(2)}(\tau) &= -\frac{1}{\hbar^2} \sum_{n \neq m} \int_0^\tau \langle n|\hat{F}^\pm|m\rangle \frac{e^{i(\omega_{nm} \mp \omega)t} - 1}{i(\omega_{nm} \mp \omega)} e^{-i(E_n - E_k)t/\hbar} \langle k|\hat{V}(t)|n\rangle dt \\ &= -\frac{1}{\hbar^2} \sum_{n \neq m} \int_0^\tau \langle n|\hat{F}^\pm|m\rangle \frac{e^{i(\omega_{nm} \mp \omega)t} - 1}{i(\omega_{nm} \mp \omega)} e^{i\omega_{kn}t} \langle k|\hat{F}^\pm|n\rangle e^{\mp i\omega t} dt \\ &= \frac{i}{\hbar^2} \sum_{n \neq m} \frac{\langle k|\hat{F}^\pm|n\rangle \langle n|\hat{F}^\pm|m\rangle}{\omega_{nm} \mp \omega} \int_0^\tau [e^{i(\omega_{nm} \mp \omega)t} - 1] e^{i(\omega_{kn} \mp \omega)t} dt, \end{aligned} \quad (10.53)$$

dove va tenuto presente che nell'elemento di matrice $\langle f|\hat{F}^\pm|i\rangle$ compare $\hat{F}^+$ se $E_f > E_i$, altrimenti compare $\hat{F}^-$. Anche negli esponenti all'interno dell'integrale va preso il segno superiore se, rispettivamente, $E_k > E_n$ oppure $E_n > E_m$, il segno inferiore altrimenti. Di questo dobbiamo tener conto sviluppando il prodotto all'interno dell'integrale. Supponendo che sia $E_k > E_m$, possiamo spezzare la sommatoria su n in tre parti, a seconda che l'energia dello stato $|n\rangle$ sia minore di E_m ($E_n < E_m$), compresa tra E_m ed E_k ($E_m < E_n < E_k$) o maggiore di E_k ($E_n > E_k$)

$$\begin{aligned}
a_k^{(2)}(\tau) = & \frac{i}{\hbar^2} \sum_{\substack{n \neq m \\ E_n < E_m}} \frac{\langle k|\hat{F}^+|n\rangle\langle n|\hat{F}^-|m\rangle}{\omega_{nm} + \omega} \int_0^\tau \left[e^{i\omega_{km}t} - e^{i(\omega_{kn}-\omega)t} \right] dt \\
& + \frac{i}{\hbar^2} \sum_{\substack{n \neq m \\ E_m < E_n < E_k}} \frac{\langle k|\hat{F}^+|n\rangle\langle n|\hat{F}^+|m\rangle}{\omega_{nm} - \omega} \int_0^\tau \left[e^{i(\omega_{km}-2\omega)t} - e^{i(\omega_{kn}-\omega)t} \right] dt \\
& + \frac{i}{\hbar^2} \sum_{\substack{n \neq m \\ E_n > E_k}} \frac{\langle k|\hat{F}^-|n\rangle\langle n|\hat{F}^+|m\rangle}{\omega_{nm} - \omega} \int_0^\tau \left[e^{i\omega_{km}t} - e^{i(\omega_{kn}+\omega)t} \right] dt. \quad (10.54)
\end{aligned}$$

Nella prima sommatoria del secondo membro è negativa ω_{nm} , mentre nell'ultimo addendo è negativa ω_{kn} . Il caso $E_k < E_m$ si ottiene dalla (10.54) con opportuni cambiamenti di segno. Per gli integrali che appaiono nella (10.54) abbiamo

$$\begin{aligned}
\int_0^\tau \left[e^{i(\omega_{km}-2\omega)t} - e^{i(\omega_{kn}-\omega)t} \right] dt &= \frac{e^{i(\omega_{km}-2\omega)\tau} - 1}{i(\omega_{km} - 2\omega)} - \frac{e^{i(\omega_{kn}-\omega)\tau} - 1}{i(\omega_{kn} - \omega)}, \\
\int_0^\tau \left[e^{i\omega_{km}t} - e^{i(\omega_{kn}\mp\omega)t} \right] dt &= \frac{e^{i\omega_{km}\tau} - 1}{i\omega_{km}} - \frac{e^{i(\omega_{kn}\mp\omega)\tau} - 1}{i(\omega_{kn} \mp \omega)}, \quad (10.55)
\end{aligned}$$

da cui

$$\begin{aligned}
a_k^{(2)}(\tau) = & \frac{1}{\hbar^2} \sum_{E_n < E_m} \frac{\langle k|\hat{F}^+|n\rangle\langle n|\hat{F}^-|m\rangle}{\omega_{nm} + \omega} \left[\frac{e^{i\omega_{km}\tau} - 1}{\omega_{km}} - \frac{e^{i(\omega_{kn}-\omega)\tau} - 1}{(\omega_{kn} - \omega)} \right] \\
& + \frac{1}{\hbar^2} \sum_{E_m < E_n < E_k} \frac{\langle k|\hat{F}^+|n\rangle\langle n|\hat{F}^+|m\rangle}{\omega_{nm} - \omega} \left[\frac{e^{i(\omega_{km}-2\omega)\tau} - 1}{(\omega_{km} - 2\omega)} - \frac{e^{i(\omega_{kn}-\omega)\tau} - 1}{(\omega_{kn} - \omega)} \right] \\
& + \frac{1}{\hbar^2} \sum_{E_n > E_k} \frac{\langle k|\hat{F}^-|n\rangle\langle n|\hat{F}^+|m\rangle}{\omega_{nm} - \omega} \left[\frac{e^{i\omega_{km}\tau} - 1}{\omega_{km}} - \frac{e^{i(\omega_{kn}+\omega)\tau} - 1}{(\omega_{kn} + \omega)} \right]. \quad (10.56)
\end{aligned}$$

E' da notare che, data la presenza di due addendi in ogni parentesi quadra, nella (10.56) compaiono in realtà 6 sommatorie. Per il calcolo della probabilità di transizione al secondo ordine dobbiamo tener conto delle ampiezza di probabilità sia al primo ordine che al secondo ordine:

$$P_{km}^{(12)} = |a_{km}^{(1)} + a_{km}^{(2)}|^2. \quad (10.57)$$

Lo sviluppo di $P_{km}^{(12)}$ è così decisamente complesso, data l'espressione (10.56) per $a_{km}^{(2)}$. Ma vedremo che considerazioni sulla conservazione dell'energia porteranno a trascurare vari termini dello sviluppo della (10.57), in particolare i “termini di interferenza”, nella maggior parte dei casi di interesse fisico. Cominciamo considerando, separatamente, i moduli quadri dei primi addendi all'interno delle parentesi quadre della (10.56) che, essendo indipendenti da n , possono essere portati fuori dalle sommatorie. Se questi fossero gli unici termini che contribuiscono alla transizione, con procedimento analogo a quello che ci ha portati alla (10.35), otterremmo per la probabilità di transizione per unità di tempo $W_{km}^{(12)} = P_{km}^{(12)}/\tau$

$$W_{km}^{(12)} = \frac{2\pi}{\hbar^3} \left| \sum_{\substack{E_n < E_m \\ E_n > E_k}} \frac{\langle k|\hat{F}^\pm|n\rangle\langle n|\hat{F}^\mp|m\rangle}{\omega_{nm} \pm \omega} \right|^2 \delta_\tau(E_k - E_m) \quad (10.58)$$

come somma dei termini alla prima e alla terza riga della (10.56). La presenza di $\delta_\tau(E_k - E_m)$ ci dice che questa espressione contribuisce significativamente solo se l'energia dello stato finale è molto vicina all'energia dello stato iniziale. Questo è dovuto alla presenza del prodotto di elementi di matrice $\langle k|\hat{F}^\pm|n\rangle\langle n|\hat{F}^\mp|m\rangle$, che comporta l'assorbimento di un quanto di energia $\hbar\omega$ e la sua successiva riemissione, o il processo in ordine inverso, finendo in uno stato con la stessa energia dello stato iniziale.

Consideriamo adesso il primo addendo dentro alla parentesi quadra della seconda riga della (10.56). Se questo fosse l'unico a contribuire avremmo

$$W_{km}^{(12)} = \frac{2\pi}{\hbar^3} \left| \sum_{E_m < E_n < E_k} \frac{\langle k|\hat{F}^+|n\rangle\langle n|\hat{F}^+|m\rangle}{\omega_{nm} - \omega} \right|^2 \delta_\tau(E_k - E_m - 2\hbar\omega). \quad (10.59)$$

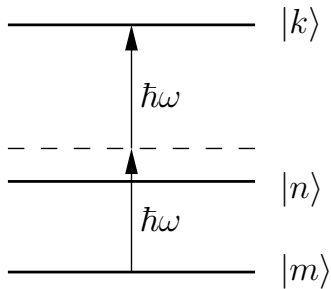


Figura 10.5 Assorbimento di due fotoni.

La funzione δ_τ ci dà la conservazione dell'energia, entro il limite dell'indeterminazione energia-tempo, dallo stato iniziale $|m\rangle$ allo stato finale $|k\rangle$ con l'assorbimento di due quanti di energia $\hbar\omega$. La variazione dell'energia avviene applicando due volte l'operatore $\hat{F}^+$ di interazione con il campo elettromagnetico, per cui sono assorbiti due quanti, ognuno di energia $\hbar\omega$, come schematizzato in fig. 10.5. Questo fenomeno era stato previsto nel 1931 da Maria Göppert-Mayer nella sua tesi di dottorato (Promotion), ed è stato osservato per la prima volta trent'anni dopo, grazie all'introduzione del laser, prima in un cristallo drogato di europio, poi nel vapore di cesio. Il denominatore sotto al prodotto di elementi di matrice nella (10.59) ci dice che al

processo contribuiranno di più gli stati intermedi $|n\rangle$ la cui energia E_n è più vicina al valore $E_m + \hbar\omega$. Nel caso $E_k < E_m$ la (10.59) diventa

$$W_{km}^{(12)} = \frac{2\pi}{\hbar^3} \left| \sum_{E_m > E_n > E_k} \frac{\langle k|\hat{F}^-|n\rangle\langle n|\hat{F}^-|m\rangle}{\omega_{nm} + \omega} \right|^2 \delta_\tau(E_m - E_k - 2\hbar\omega),$$

che corrisponde alla transizione da $|m\rangle$ a $|k\rangle$ mediante l'emissione stimolata di due quanti ognuno di energia $\hbar\omega$.

Passiamo adesso a considerare i secondi addendi all'interno delle parentesi quadre della (10.56). Sempre nell'ipotesi che sia $E_k > E_m$, il termine alla prima riga continua a descrivere l'emissione stimolata di un quanto seguita da un assorbimento, quello alla seconda riga descrive due assorbimenti consecutivi, e quello alla terza riga un assorbimento seguito da un'emissione stimolata. Ma un'occhiata ai denominatori che compaiono uno sotto al prodotto degli elementi di matrice, e l'altro al secondo termine dentro alla parentesi quadra, ci dice che in ognuno di questi processi abbiamo a che fare con due quanti di energia diversa. Il processo diventa quindi particolarmente se abbiamo un'hamiltoniana di perturbazione che, durante l'accensione, contiene due frequenze diverse, del tipo

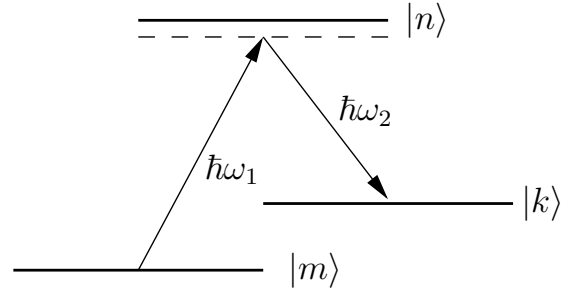


Figura 10.6 Effetto Raman stimolato.

Il processo diventa quindi particolarmente se abbiamo un'hamiltoniana di perturbazione che, durante l'accensione, contiene due frequenze diverse, del tipo

$$\hat{V}(t) = \hat{F}_1^+ e^{-i\omega_1 t} + \hat{F}_1^- e^{i\omega_1 t} + \hat{F}_2^+ e^{-i\omega_2 t} + \hat{F}_2^- e^{i\omega_2 t}, \quad (10.60)$$

con le opportune relazioni tra $\hat{F}_1^+$, $\hat{F}_1^-$, $\hat{F}_2^+$, e $\hat{F}_2^-$ perché $\hat{V}(t)$ sia hermitiano. Se, per esempio, l'unico termine a contribuire alla transizione fosse quello relativo all'ultima riga della (10.56), avremmo

$$W_{km}^{(12)} = \frac{2\pi}{\hbar^3} \left| \sum_{E_n > E_k} \frac{\langle k | \hat{F}_2^- | n \rangle \langle n | \hat{F}_1^+ | m \rangle}{\omega_{nm} - \omega_1} \right|^2 \delta_\tau(E_n - E_k - \hbar\omega_2), \quad (10.61)$$

che corrisponde all'effetto Raman stimolato. L'effetto Raman, osservato prima in liquidi e cristalli, è uno scattering inelastico di un fotone, in cui l'energia del fotone scatterato subisce una variazione, normalmente una diminuzione dovuta al trasferimento di parte dell'energia al sistema "scatterante". Nel caso dell'effetto Raman stimolato, schematizzato in fig. 10.6, l'assorbimento del fotone di frequenza angolare ω_1 è seguito da un'emissione stimolata indotta dalla componente a frequenza ω_2 del campo elettromagnetico. Nella sommatoria della (10.61), gli stati intermedi $|n\rangle$ che contribuiscono di più al processo sono quelli che hanno l'energia E_n contemporaneamente più vicina a $E_m + \hbar\omega_1$ e a $E_k + \hbar\omega_2$.

10.3 Transizioni di dipolo elettrico

La forma completa dell'hamiltoniana di interazione tra un atomo ed un campo elettromagnetico è abbastanza complicata. Per una prima discussione sono sufficienti le approssimazioni che seguono. L'operatore momento di dipolo elettrico $\hat{\mathbf{P}}$ per un atomo neutro può essere scritto

$$\hat{\mathbf{P}} = Ze\hat{\mathbf{R}} - e \sum_{i=1}^Z \hat{\mathbf{r}}_i, \quad (10.62)$$

dove Z è il numero atomico, $\hat{\mathbf{R}}$ è l'operatore posizione del nucleo e gli $\hat{\mathbf{r}}_i$ sono gli operatori posizione dei singoli elettroni nel sistema di riferimento del laboratorio. Trascurando il moto del nucleo e scegliendo un sistema di riferimento con l'origine sul nucleo stesso, che è lo stesso sistema di riferimento

usato per scrivere le funzioni d'onda elettroniche dell'atomo, la (10.62) può essere riscritta

$$\hat{\mathbf{P}} = -e \sum_{i=1}^Z \hat{\mathbf{r}}_i, \quad (10.63)$$

dove, questa volta, gli $\hat{\mathbf{r}}_i$ sono gli operatori posizione degli elettroni rispetto al nucleo. $\hat{\mathbf{P}}$ è un *operatore vettoriale*, con tre componenti $\hat{P}_x$, $\hat{P}_y$ e $\hat{P}_z$

$$\hat{P}_x = -e \sum_{i=1}^Z \hat{x}_i, \quad \hat{P}_y = -e \sum_{i=1}^Z \hat{y}_i, \quad \hat{P}_z = -e \sum_{i=1}^Z \hat{z}_i, \quad (10.64)$$

che si trasformano come le componenti di un vettore per rotazioni degli assi. Per come è definito, $\hat{\mathbf{P}}$ è un operatore dispari, cioè cambia segno per una trasformazione di parità che mandi $\mathbf{r}$ in $-\mathbf{r}$. Di conseguenza, se $|\psi\rangle$ ha una parità definita, $\hat{\mathbf{P}}|\psi\rangle$ ha parità opposta a $|\psi\rangle$. Poiché per un atomo, in generale, gli autostati $|n\rangle$ dell'hamiltoniana imperturbata $\hat{H}_0$ hanno parità definita, avremo sempre

$$\langle n|\hat{\mathbf{P}}|n\rangle = 0. \quad (10.65)$$

In presenza di un campo elettrico esterno uniforme $\mathbf{E}$ l'hamiltoniana di interazione si può scrivere

$$\hat{V} = -\mathbf{E} \cdot \hat{\mathbf{P}}. \quad (10.66)$$

Consideriamo adesso un'onda elettromagnetica piana, monocromatica di frequenza angolare ω , polarizzata linearmente, che si propaghi lungo l'asse z , con il campo elettrico diretto lungo l'asse x ed il campo magnetico diretto lungo l'asse y . Avremo

$$\mathbf{E}(z, t) = \mathbf{E}_0 \cos(kz - \omega t + \varphi), \quad (10.67)$$

dove $k = 2\pi/\lambda$ è il vettore d'onda, o meglio, la componente z del vettore d'onda, che nel nostro caso è l'unica diversa da zero, e φ è una fase. La funzione d'onda degli elettroni dell'atomo è sensibilmente diversa da zero in una regione di dimensioni lineari dell'ordine di grandezza del raggio di Bohr, $a_0 \simeq 5 \times 10^{-11}$ m, mentre la lunghezza d'onda delle radiazioni visibili è dell'ordine di 6×10^{-7} m. Il campo elettrico dell'onda può quindi essere considerato uniforme su tutta l'estensione di interesse delle funzioni d'onda dell'atomo. L'hamiltoniana di interazione tra l'atomo e la radiazione può così essere approssimata dall'espressione

$$\hat{V} = \hat{\mathbf{P}} \cdot \mathbf{E}_0 \cos(\omega t), \quad \text{nel nostro caso} \quad \hat{V} = \hat{P}_x E_0 \cos(\omega t), \quad (10.68)$$

una volta scelta l'origine dei tempi in modo che sia $\varphi = \pi$, per liberarci del segno meno. Per gli elementi di matrice abbiamo

$$\langle k|\hat{V}|m\rangle = \langle k|\hat{P}_x|m\rangle E_0 \cos(\omega t), \quad (10.69)$$

e possono essere diversi da zero solo se la parità di $|k\rangle$ è diversa dalla parità di $|m\rangle$. Come al solito, facciamo l'ipotesi che sia $E_k > E_m$. La (10.69) può essere riscritta

$$\langle k|\hat{V}|m\rangle = \langle k|\hat{P}_x|m\rangle \frac{E_0}{2} [e^{-i\omega t} + e^{+i\omega t}], \quad (10.70)$$

che è un caso particolare della (10.15). Per semplicità, supponiamo di avere un atomo con due soli stati $|1\rangle$ e $|2\rangle$, con $E_2 > E_1$, e che inizialmente l'atomo si trovi nello stato a energia più bassa $|1\rangle$. Con una trattazione del tutto analoga a quella del paragrafo 10.1.3, otteniamo che il termine in $e^{+i\omega t}$ può essere trascurato rispetto al termine in $e^{-i\omega t}$ (approssimazione di *onda rotante*), ed otteniamo per la probabilità di transizione dovuta alla nostra onda monocromatica

$$P_{21}(\tau) = \frac{\pi}{2\hbar} \left| \langle 2 | \hat{P}_x | 1 \rangle \right|^2 E_0^2 \tau \delta_\tau(E_2 - E_1 - h\nu). \quad (10.71)$$

A un'onda piana monocromatica è associata una densità di energia elettrica per unità di volume w_{el}

$$w_{el} = \frac{\epsilon_0}{2} \overline{E^2} = \frac{\epsilon_0}{4} E_0^2, \quad (10.72)$$

e ricordando la relazione tra i campi elettrico e magnetico di un'onda piana $B_y = E_x/c$, vista nel paragrafo 8.1, abbiamo anche una densità di energia magnetica per unità di volume w_{mag}

$$w_{mag} = \frac{1}{2\mu_0} \overline{B^2} = \frac{1}{2\mu_0 c^2} \overline{E^2} = \frac{\epsilon_0}{4} E_0^2, \quad (10.73)$$

quindi, per la densità volumica di energia totale w associata all'onda abbiamo

$$w = w_{el} + w_{mag} = \frac{\epsilon_0}{2} E_0^2. \quad (10.74)$$

Sostituendo nella (10.71) otteniamo

$$P_{21}(\tau) = \frac{\pi}{\hbar \epsilon_0} \left| \langle 2 | \hat{P}_x | 1 \rangle \right|^2 w \tau \delta_\tau(E_2 - E_1 - h\nu), \quad \text{e} \quad W_{21} = \frac{\pi}{\hbar \epsilon_0} \left| \langle 2 | \hat{P}_x | 1 \rangle \right|^2 w \delta_\tau(E_2 - E_1 - h\nu). \quad (10.75)$$

10.4 Assorbimento, emissione stimolata e spontanea

Quanto sviluppato nel paragrafo precedente ci permette di riconsiderare il paragrafo 9.1 sulla derivazione della formula di Planck da parte di Einstein. Nel caso della radiazione di corpo nero, la radiazione non è monocromatica, ma abbiamo una distribuzione spettrale della densità di energia radiante per unità di volume $u(\nu)$. Cominciamo considerando un'onda elettromagnetica che, come nel paragrafo precedente, si propaghi lungo l'asse z e sia polarizzata lungo l'asse x ma che, invece di essere monocromatica, sia caratterizzata da una densità spettrale $u(\nu)$. Per calcolare il contributo della radiazione con frequenza compresa tra ν e $\nu + d\nu$ alla probabilità di transizione da $|1\rangle$ a $|2\rangle$ sostituiamo la densità volumica di energia w con $u(\nu) d\nu$ nella (10.75)

$$dP_{21}(\tau) = \frac{\pi}{\hbar \epsilon_0} \left| \langle 2 | \hat{P}_x | 1 \rangle \right|^2 u(\nu) d\nu \tau \delta_\tau(E_2 - E_1 - h\nu), \quad (10.76)$$

e la probabilità di transizione complessiva è poi data dall'integrale su tutte le frequenze

$$P_{21}(\tau) = \frac{\pi}{\hbar \epsilon_0} \left| \langle 2 | \hat{P}_x | 1 \rangle \right|^2 \tau \int_0^{+\infty} u(\nu) \delta_\tau(E_2 - E_1 - h\nu) d\nu. \quad (10.77)$$

Data la forma della funzione δ_τ , alla transizione contribuiscono in modo apprezzabile solo le frequenze ν prossime al valore $\nu_{21} = (E_2 - E_1)/h$, che annulla l'argomento della δ_τ stessa. Se facciamo

l'ipotesi che la larghezza h/τ sia sufficientemente stretta perché nella regione di frequenze in cui la δ_τ è apprezzabilmente diversa da zero $u(\nu)$ sia praticamente costante, possiamo portare $u(\nu)$ fuori dall'integrale. Nel calcolo dell'integrale ricordiamo che, se $\nu_{21} - h/\tau \gg 0$, l'estremo inferiore di integrazione può essere esteso a $-\infty$, e che $d\nu = dE/h$, ottenendo

$$P_{21}(\tau) = \frac{\pi}{\hbar \varepsilon_0} |\langle 2|\hat{P}_x|1\rangle|^2 u(\nu_{21}) \frac{\tau}{h} = \frac{1}{2\hbar^2 \varepsilon_0} |\langle 2|\hat{P}_x|1\rangle|^2 u(\nu_{21}) \tau, \quad (10.78)$$

che corrisponde a una probabilità di transizione per unità di tempo

$$W_{21} = \frac{P_{21}(\tau)}{\tau} = \frac{1}{2\hbar^2 \varepsilon_0} |\langle 2|\hat{P}_x|1\rangle|^2 u(\nu_{21}). \quad (10.79)$$

Vogliamo confrontare la (10.79) con la (9.1) per il coefficiente di Einstein per l'assorbimento. Per far questo dobbiamo considerare un gas di N atomi (o molecole) per unità di volume. Ad un certo istante, N_1 atomi saranno nello stato $|1\rangle$ e N_2 nello stato $|2\rangle$. Inoltre, le funzioni d'onda elettroniche dei singoli atomi saranno orientate in modo casuale. Se chiamiamo $\mathbf{e}$ il versore di polarizzazione dell'onda elettromagnetica, nella (10.79) dobbiamo sostituire $|\langle 2|\hat{P}_x|1\rangle|^2$ con l'espressione $|\mathbf{e} \cdot \langle 2|\hat{\mathbf{P}}|1\rangle|^2 = |\cos \vartheta \langle 2|\hat{\mathbf{P}}|1\rangle|^2$, dove ϑ è l'angolo tra $\mathbf{e}$ e il vettore $\langle 2|\hat{\mathbf{P}}|1\rangle$. Tenendo presente che, per gli atomi del gas, tutte le orientazioni sono equiprobabili, otteniamo

$$\langle |\cos \vartheta \langle 2|\hat{\mathbf{P}}|1\rangle|^2 \rangle = \frac{1}{3} |\langle 2|\hat{\mathbf{P}}|1\rangle|^2, \quad (10.80)$$

da cui

$$\frac{dN_{2 \leftarrow 1}}{dt} = B_{21} u(\nu_{21}) N_1 = N_1 \frac{1}{6\hbar^2 \varepsilon_0} |\langle 2|\hat{\mathbf{P}}|1\rangle|^2 u(\nu_{21}). \quad (10.81)$$

Come nella (9.1), il termine B_{21} indica il coefficiente di assorbimento per la transizione $|1\rangle \rightarrow |2\rangle$. Otteniamo quindi la seguente relazione tra il coefficiente di assorbimento di Einstein e l'elemento di matrice per la transizione di dipolo elettrico

$$B_{21} = \frac{1}{6\hbar^2 \varepsilon_0} |\langle 2|\hat{\mathbf{P}}|1\rangle|^2. \quad (10.82)$$

La relazione è apparentemente semplice, ma va ricordato che il calcolo degli elementi di matrice di $\hat{\mathbf{P}}$ richiede una conoscenza, analitica o numerica, delle funzioni d'onda degli stati atomici $|1\rangle$ e $|2\rangle$, che, normalmente, non è disponibile. La formula rende tuttavia possibili numerose considerazioni sulle *regole di selezione* per la transizione.

Nel paragrafo 9.1 avevamo ottenuto l'equazione (9.12) per il coefficiente di Einstein per l'emissione spontanea

$$A_{12} = \frac{8\pi h \nu_{21}^3}{c^3} B_{12}. \quad (10.83)$$

Ricordando che $B_{12} = B_{21}$ e utilizzando la (10.82) abbiamo

$$A_{12} = \frac{8\pi h \nu_{21}^3}{c^3} \frac{1}{6\hbar^2 \varepsilon_0} |\langle 2|\hat{\mathbf{P}}|1\rangle|^2 = \frac{8\pi^2 \nu_{21}^3}{3\hbar \varepsilon_0 c^3} |\langle 2|\hat{\mathbf{P}}|1\rangle|^2 = \frac{\omega_{21}^3}{3\pi \hbar \varepsilon_0 c^3} |\langle 2|\hat{\mathbf{P}}|1\rangle|^2. \quad (10.84)$$

A questa espressione si può arrivare anche, e più rigorosamente, dall'elettrodinamica quantistica, come vedremo brevemente nell'Appendice D, ma non direttamente dalla trattazione semiclassica

dell'interazione radiazione-materia. Il coefficiente A_{12} dipende così, oltre che dal modulo quadro dell'elemento di matrice del dipolo tra stato iniziale e stato finale, anche dal cubo della frequenza. Questo si riflette sulla vita media dello stato eccitato $|2\rangle$, che, nel caso di decadimento possibile solo verso lo stato $|1\rangle$, si scrive

$$\tau_2 = \frac{1}{A_{12}}. \quad (10.85)$$

Nel caso, decisamente più frequente in pratica, di più decadimenti possibili verso diversi stati a energia più bassa, avremo

$$\frac{1}{\tau_2} = \sum_i A_{i2}, \quad (10.86)$$

dove la somma corre su tutti gli stati $|i\rangle$ con $E_i < E_2$ verso i quali è possibile il decadimento. Per avere un'idea degli ordini di grandezza in gioco, valutiamo A_{12} per alcuni valori della frequenza ν_{21} . Un ordine di grandezza per l'elemento di matrice del momento di dipolo si ottiene considerando un elettrone distante un raggio di Bohr a_0 dal nucleo:

$$p \simeq ea_0 \simeq 1.6 \times 10^{-19} \text{ C} \times 5 \times 10^{-11} \text{ m} = 8 \times 10^{-30} \text{ Cm}. \quad (10.87)$$

Inserendo questo valore e la (10.84) nella (10.85) otteniamo

$$\tau = \frac{3 \hbar \epsilon_0 c^3}{8\pi^2 \nu_{21}^3} \frac{1}{|\langle 2|\hat{\mathbf{p}}|1\rangle|^2} \simeq \frac{9.1 \times 10^{-22}}{\nu^3} \frac{1}{6.4 \times 10^{-59}} = \frac{1.4 \times 10^{37}}{\nu^3} \text{ s}. \quad (10.88)$$

Se abbiamo a che fare con una transizione nel visibile, con lunghezza d'onda $\lambda \simeq 6000 \text{ \AA} = 600 \text{ nm}$, e quindi frequenza $\nu \simeq 5 \times 10^{14} \text{ Hz}$, otteniamo

$$\tau = \frac{1.4 \times 10^{37}}{1.25 \times 10^{44}} = 1.1 \times 10^{-7} \text{ s}, \quad (10.89)$$

mentre per una transizione nell'infrarosso con $\lambda = 10 \text{ }\mu\text{m}$, e $\nu \simeq 3 \times 10^{13} \text{ Hz}$, abbiamo $\tau = 0.53 \text{ ms}$, e per una frequenza a microonde, con lunghezza d'onda $\lambda = 1 \text{ cm}$ e $\nu \simeq 30 \text{ GHz}$, abbiamo $\tau \simeq 5.3 \times 10^5 \text{ s} \simeq 144 \text{ ore}$.

La (10.84) e le considerazioni fatte qui sopra per gli ordini di grandezza hanno una stretta analogia con i risultati ottenuti per l'emissione di dipolo elettrico nell'elettrodinamica classica. Classicamente, un dipolo oscillante con momento

$$\mathbf{p} = \mathbf{p}_0 \sin \omega t \quad (10.90)$$

emette una potenza, integrata su tutto l'angolo solido, pari a

$$W = \frac{1}{12 \epsilon_0 \pi c^3} p_0^2 \omega^4, \quad (10.91)$$

ed in un intervallo di tempo dt perde quindi una quantità di energia

$$dU_{\text{cl}} = W dt = \frac{1}{12 \epsilon_0 \pi c^3} p_0^2 \omega^4 dt. \quad (10.92)$$

I calcoli quantistici fatti in questo paragrafo ci dicono che un atomo con due livelli che si trovi inizialmente nel livello a energia più alta $|2\rangle$, durante l'intervallo di tempo dt ha probabilità

$$dP = A_{12} dt = \frac{\omega_{21}^3}{3\pi \hbar \epsilon_0 c^3} |\langle 2|\hat{\mathbf{p}}|1\rangle|^2 dt \quad (10.93)$$

di decadere allo stato ad energia più bassa $|1\rangle$. Il decadimento comporta la perdita dell'energia $\hbar\omega$, quindi abbiamo per la perdita di energia

$$dU_{\text{quant}} = \frac{\omega_{21}^4}{3\pi\epsilon_0 c^3} |\langle 2|\hat{\mathbf{p}}|1\rangle|^2 dt, \quad (10.94)$$

Un paragone con la (10.92) ci fa vedere che le due espressioni differiscono per un fattore 4 al denominatore.

Il rapporto tra la perdita di energia per emissione indotta e quella per emissione spontanea, in presenza di una densità spettrale di radiazione $u(\nu)$, vale

$$\frac{B_{12} u(\nu_{21})}{A_{12}} = \frac{c^3 u(\nu_{21})}{8\pi h \nu_{21}^3}. \quad (10.95)$$

Quindi, a parità di densità spettrale di energia radiante presente, il fenomeno dell'emissione indotta è più importante a frequenza più basse. Se abbiamo a che fare con una distribuzione spettrale di energia radiante tipo corpo nero abbiamo

$$u(\nu) = \frac{8\pi h \nu^3}{c^3} \frac{1}{e^{h\nu/kT} - 1},$$

che, sostituita nella (10.95), porta a

$$\frac{B_{12} u(\nu_{21})}{A_{12}} = \frac{1}{e^{h\nu/kT} - 1} = \bar{n}(\nu), \quad (10.96)$$

dove $\bar{n}(\nu)$ è il numero medio di fotoni che si trova in una cella dello spazio delle fasi di energia $h\nu$. Se $\bar{n}(\nu) > 1$ prevale l'emissione indotta, altrimenti prevale l'emissione spontanea. Per vedere in quali condizioni il numero medio di fotoni per cella dello spazio delle fasi è maggiore di 1, risolviamo la disuguaglianza

$$\bar{n}(\nu) = \frac{1}{e^{h\nu/kT} - 1} > 1, \quad (10.97)$$

ottenendo

$$\frac{\nu}{T} < \frac{k}{h} \ln 2 \simeq 1.44 \times 10^{10} \text{ Hz/K}, \quad \text{e} \quad \lambda T > \frac{1}{\ln 2} \frac{hc}{k} \simeq 20\,000 \mu\text{m K} \quad (10.98)$$

Questa relazione vale in generale, e non solo nel caso di radiazione di corpo nero.

10.5 Interazione con campo forte

Supponiamo ancora di avere un sistema atomico a due livelli $|1\rangle$ e $|2\rangle$, con $E_2 > E_1$, che stia interagendo con un campo radiativo così intenso che non sia possibile un trattamento perturbativo. Ricordiamo però che nel paragrafo 10.1.2 avevamo detto che, quando ci sono in gioco due soli livelli, è possibile ottenere la soluzione esatta, e il trattamento perturbativo non è necessario. Come esempio, possiamo pensare ad un sistema a due livelli interagente con radiazione laser. Possiamo scrivere il campo elettrico nella zona di interesse per l'atomo come

$$\mathbf{E} = -\mathbf{E}_0 \cos \omega t = -\frac{E_0}{2} \mathbf{e} \left[e^{-i\omega t} + e^{+i\omega t} \right],$$

dove $\mathbf{e}$ è il versore del campo elettrico dell'onda, e abbiamo scelto l'origine dei tempi in modo che nell'espressione compaia il segno meno che verrà eliminato nella (10.99) qui sotto. L'hamiltoniana di interazione si può scrivere ancora

$$\hat{V} = -\mathbf{E} \cdot \hat{\mathbf{P}} = \mathbf{e} \cdot \hat{\mathbf{P}} \frac{E_0}{2} \left[e^{-i\omega t} + e^{+i\omega t} \right]. \quad (10.99)$$

Il generico stato in cui si trova il sistema può essere scritto, secondo la (10.3),

$$|\psi(t)\rangle = a_1(t)e^{-iE_1t/\hbar}|1\rangle + a_2(t)e^{-iE_2t/\hbar}|2\rangle, \quad (10.100)$$

e, analogamente a quanto fatto nel paragrafo 10.1.2, l'equazione di Schrödinger diventa

$$i\hbar \frac{\partial}{\partial t} \sum_{r=1}^2 a_r(t) e^{-iE_r t/\hbar} |r\rangle = \left[\hat{H}_0 + \hat{V}(t) \right] \sum_{n=1}^2 a_n(t) e^{-iE_n t/\hbar} |n\rangle. \quad (10.101)$$

Seguendo le (10.7) otteniamo per le derivate temporali delle singole ampiezze

$$\frac{\partial a_k}{\partial t} = \frac{1}{i\hbar} \sum_{n=1}^2 a_n e^{-i(E_n - E_k)t/\hbar} \langle k | \hat{V}(t) | n \rangle = \frac{1}{i\hbar} \sum_{n=1}^2 a_n e^{i\omega_{kn}t} \langle k | \hat{V}(t) | n \rangle. \quad (10.102)$$

Gli elementi di matrice di $\hat{V}$ tra i due stati dell'atomo sono

$$\begin{aligned} \langle 2 | \hat{V} | 1 \rangle &= \langle 2 | \mathbf{e} \cdot \hat{\mathbf{P}} | 1 \rangle \frac{E_0}{2} \left[e^{-i\omega t} + e^{+i\omega t} \right] \\ \langle 1 | \hat{V} | 2 \rangle &= \langle 1 | \mathbf{e} \cdot \hat{\mathbf{P}} | 2 \rangle \frac{E_0}{2} \left[e^{+i\omega t} + e^{-i\omega t} \right], \end{aligned} \quad (10.103)$$

con, nel nostro caso, $\omega_{21} > 0$ e $\omega_{12} = -\omega_{21}$. Le fasi degli stati $|1\rangle$ e $|2\rangle$ possono essere scelte in modo che $\langle 2 | \mathbf{e} \cdot \hat{\mathbf{P}} | 1 \rangle$ sia reale, possiamo quindi porre

$$D = \langle 2 | \mathbf{e} \cdot \hat{\mathbf{P}} | 1 \rangle = \langle 1 | \mathbf{e} \cdot \hat{\mathbf{P}} | 2 \rangle, \quad (10.104)$$

mentre $\langle 1 | \mathbf{e} \cdot \hat{\mathbf{P}} | 1 \rangle = \langle 2 | \mathbf{e} \cdot \hat{\mathbf{P}} | 2 \rangle = 0$ per la (10.65). Le equazioni differenziali per $a_1(t)$ e $a_2(t)$ si scrivono quindi

$$\begin{aligned} \frac{\partial a_1(t)}{\partial t} &= \frac{1}{i\hbar} D \frac{E_0}{2} a_2(t) \left[e^{-i(\omega_{21} + \omega)t} + e^{-i(\omega_{21} - \omega)t} \right] \\ \frac{\partial a_2(t)}{\partial t} &= \frac{1}{i\hbar} D \frac{E_0}{2} a_1(t) \left[e^{i(\omega_{21} - \omega)t} + e^{i(\omega_{21} + \omega)t} \right]. \end{aligned} \quad (10.105)$$

Fino a qui non abbiamo fatto approssimazioni. Introducendo adesso l'approssimazione di campo rotante, trascurando cioè i termini contenenti $e^{\pm i(\omega_{21} + \omega)t}$ rispetto a quelli contenenti $e^{\pm i(\omega_{21} - \omega)t}$, le equazioni si riducono a

$$\begin{aligned} \frac{\partial a_1(t)}{\partial t} &= \frac{1}{i\hbar} a_2(t) e^{-i(\omega_{21} - \omega)t} D \frac{E_0}{2} \\ \frac{\partial a_2(t)}{\partial t} &= \frac{1}{i\hbar} a_1(t) e^{i(\omega_{21} - \omega)t} D \frac{E_0}{2}. \end{aligned} \quad (10.106)$$

Introducendo la *frequenza di Rabi*

$$\Omega_R = \frac{DE_0}{\hbar} \quad (10.107)$$

le (10.106) diventano

$$\begin{aligned} \frac{\partial a_1(t)}{\partial t} &= -\frac{i}{2} \Omega_R a_2(t) e^{-i(\omega_{21}-\omega)t} \\ \frac{\partial a_2(t)}{\partial t} &= -\frac{i}{2} \Omega_R a_1(t) e^{i(\omega_{21}-\omega)t}. \end{aligned} \quad (10.108)$$

Adesso introduciamo il cambiamento di variabili

$$\begin{aligned} a_1(t) &= b_1(t) \\ a_2(t) &= b_2(t) e^{i(\omega_{21}-\omega)t}, \end{aligned} \quad (10.109)$$

e la quantità $\delta\omega = \omega - \omega_{21}$. Otteniamo

$$\begin{aligned} \frac{\partial b_1(t)}{\partial t} &= -\frac{i}{2} \Omega_R b_2(t) \\ \frac{\partial b_2(t)}{\partial t} &= -\frac{i}{2} \Omega_R b_1(t) + i\delta\omega b_2(t). \end{aligned} \quad (10.110)$$

Possiamo ottenere un'equazione che coinvolga la sola variabile $b_2(t)$ facendo prima la derivata rispetto al tempo della seconda delle (10.110)

$$\frac{\partial^2 b_2(t)}{\partial t^2} = -\frac{i}{2} \Omega_R \frac{\partial b_1(t)}{\partial t} + i\delta\omega \frac{\partial b_2(t)}{\partial t}, \quad (10.111)$$

e poi sostituendo nel risultato la prima delle (10.110), ottenendo

$$\frac{\partial^2 b_2(t)}{\partial t^2} - i\delta\omega \frac{\partial b_2(t)}{\partial t} + \frac{\Omega_R^2}{4} b_2(t) = 0. \quad (10.112)$$

La (10.112) è un'equazione differenziale lineare, omogenea e a coefficienti costanti, quindi con soluzione generale del tipo

$$b_2(t) = z_1 e^{\alpha_1 t} + z_2 e^{\alpha_2 t}, \quad (10.113)$$

dove z_1 e z_2 sono costanti complesse arbitrarie, da determinare dalle condizioni iniziali del problema, mentre α_1 e α_2 sono le soluzioni dell'equazione algebrica

$$\alpha^2 - i\delta\omega \alpha + \frac{\Omega_R^2}{4} = 0. \quad (10.114)$$

Abbiamo quindi due soluzioni puramente immaginarie

$$\alpha_{1,2} = \frac{i\delta\omega \pm \sqrt{-\delta\omega^2 - \Omega_R^2}}{2}, \quad (10.115)$$

da cui

$$b_2(t) = z_1 e^{i\lambda_1 t} + z_2 e^{i\lambda_2 t} \quad (10.116)$$

con

$$\lambda_1 = -i\alpha_1 = \frac{\delta\omega}{2} + \frac{\sqrt{\Omega_R^2 + \delta\omega^2}}{2}, \quad \lambda_2 = -i\alpha_2 = \frac{\delta\omega}{2} - \frac{\sqrt{\Omega_R^2 + \delta\omega^2}}{2}, \quad (10.117)$$

e la condizione che sia $b_2(0) = 0$ ci dà $z_2 = -z_1$. Per comodità poniamo $z_1 = -i \frac{z}{2}$, da cui

$$b_2(t) = i \frac{z}{2} (e^{i\lambda_1 t} - e^{i\lambda_2 t}) = z e^{i\delta\omega t/2} \sin\left(\frac{\sqrt{\Omega_R^2 + \delta\omega^2}}{2} t\right), \quad (10.118)$$

e da qui otteniamo, attraverso la seconda delle (10.109),

$$a_2(t) = b_2(t) e^{-i\delta\omega t} = z e^{-i\delta\omega t/2} \sin\left(\frac{\sqrt{\Omega_R^2 + \delta\omega^2}}{2} t\right). \quad (10.119)$$

L'ampiezza $a_1(t) = b_1(t)$ si può ottenere inserendo la (10.118), per esempio, nella seconda delle (10.110) riscritta nella forma

$$a_1(t) = b_1(t) = \frac{2}{\Omega_R} \left[\delta\omega b_2(t) + i \frac{\partial b_2(t)}{\partial t} \right]. \quad (10.120)$$

Consideriamo prima il caso particolare di radiazione in risonanza con la transizione atomica, cioè il caso in cui $\delta\omega = \omega - \omega_{21} = 0$, e, di conseguenza, $a_2(t) = b_2(t)$. La (10.118) si riduce a

$$a_2(t) = b_2(t) = z \sin\left(\frac{\Omega_R}{2} t\right), \quad (10.121)$$

e, inserendola nella (10.120), abbiamo

$$a_1(t) = i z \cos\left(\frac{\Omega_R}{2} t\right). \quad (10.122)$$

Se scegliamo la fase in modo che sia $z = i$ otteniamo

$$a_1(t) = \cos\left(\frac{\Omega_R}{2} t\right), \quad e \quad a_2(t) = -i \sin\left(\frac{\Omega_R}{2} t\right), \quad (10.123)$$

che rispettano le condizioni iniziali $a_1(0) = 1$ e $a_2(0) = 0$, oltre alla $|a_1(t)|^2 + |a_2(t)|^2 = 1$. Le probabilità di trovare il sistema nello stato $|1\rangle$ o nello stato $|2\rangle$ sono rispettivamente

$$\begin{aligned} P_1(t) &= |a_1(t)|^2 = \cos^2\left(\frac{\Omega_R}{2} t\right) = \frac{1 + \cos(\Omega_R t)}{2} \\ P_2(t) &= |a_2(t)|^2 = \sin^2\left(\frac{\Omega_R}{2} t\right) = \frac{1 - \cos(\Omega_R t)}{2}. \end{aligned} \quad (10.124)$$

Le (10.124) ci dicono che il campo elettromagnetico guida delle oscillazioni di popolazione, dette *oscillazioni di Rabi*, per cui gli atomi, che per $t = 0$ si trovavano tutti nello stato a energia più bassa $|1\rangle$, dopo un tempo $T_\pi = \pi/\Omega_R$ si trovano tutti nello stato a energia più alta $|2\rangle$, mentre dopo un tempo $T_R = 2T_\pi = 2\pi/\Omega_R$ si ritrovano tutti nello stato $|1\rangle$, e così via. Questo è il motivo per cui abbiamo dato a Ω_R il nome di *frequenza di Rabi*. Se la radiazione elettromagnetica viene applicata per un intervallo di tempo pari a T_π si ha un *impulso* π , che porta tutti gli atomi dallo stato $|1\rangle$ allo stato $|2\rangle$. La condizione per un impulso π è

$$\Omega_R T_\pi = \frac{DE_0}{\hbar} T_\pi = \pi, \quad (10.125)$$

e può essere realizzata agendo sul tempo di durata dell'impulso e sull'intensità del campo elettrico dell'onda.

Nel caso di detuning tra la frequenza dell'onda elettromagnetica e la frequenza di transizione, cioè nel caso di $\delta\omega \neq 0$, $b_2(t)$ conserva la sua forma completa (10.118). Sostituendola nella (10.120) e usando la “stenografia”

$$\Omega'_R = \sqrt{\Omega_R^2 + \delta\omega^2} \quad (10.126)$$

otteniamo $a_1(t)$ con i passaggi che seguono

$$\begin{aligned} a_1(t) &= \frac{2z}{\Omega_R} \left[\delta\omega e^{i\delta\omega t/2} \sin\left(\frac{\Omega'_R}{2} t\right) - \frac{\delta\omega}{2} e^{i\delta\omega t/2} \sin\left(\frac{\Omega'_R}{2} t\right) + i \frac{\Omega'_R}{2} e^{i\delta\omega t/2} \cos\left(\frac{\Omega'_R}{2} t\right) \right] \\ &= \frac{z e^{i\delta\omega t/2}}{\Omega_R} \left[\delta\omega \sin\left(\frac{\Omega'_R}{2} t\right) + i \Omega'_R \cos\left(\frac{\Omega'_R}{2} t\right) \right]. \end{aligned} \quad (10.127)$$

Ricaviamo $|z|^2$ imponendo la condizione di normalizzazione $|a_1(t)|^2 + |a_2(t)|^2 = 1$

$$\begin{aligned} 1 &= |a_1(t)|^2 + |a_2(t)|^2 = \frac{|z|^2}{\Omega_R^2} \left[\delta\omega^2 \sin^2\left(\frac{\Omega'_R}{2} t\right) + \Omega_R'^2 \cos^2\left(\frac{\Omega'_R}{2} t\right) \right] + |z|^2 \sin^2\left(\frac{\Omega'_R}{2} t\right) \\ &= |z|^2 \left[\left(\frac{\delta\omega^2}{\Omega_R^2} + 1 \right) \sin^2\left(\frac{\Omega'_R}{2} t\right) + \frac{\Omega_R'^2}{\Omega_R^2} \cos^2\left(\frac{\Omega'_R}{2} t\right) \right] \\ &= |z|^2 \left[\frac{\delta\omega^2 + \Omega_R^2}{\Omega_R^2} \sin^2\left(\frac{\Omega'_R}{2} t\right) + \frac{\delta\omega^2 + \Omega_R^2}{\Omega_R^2} \cos^2\left(\frac{\Omega'_R}{2} t\right) \right] = |z|^2 \frac{\Omega_R^2 + \delta\omega^2}{\Omega_R^2}, \end{aligned} \quad (10.128)$$

dove abbiamo risostituito la (10.126). Abbiamo così

$$|z|^2 = \frac{\Omega_R^2}{\Omega_R^2 + \delta\omega^2}, \quad \text{ovvero} \quad z = e^{i\varphi} \frac{\Omega_R}{\sqrt{\Omega_R^2 + \delta\omega^2}} \quad (10.129)$$

con φ fase arbitraria. Ponendo $\varphi = -\pi/2$, in modo che sia $e^{i\varphi} = -i$, otteniamo

$$\begin{aligned} a_1(t) &= e^{i\delta\omega t/2} \left[\cos\left(\frac{\sqrt{\Omega_R^2 + \delta\omega^2}}{2} t\right) - i \frac{\delta\omega}{\sqrt{\Omega_R^2 + \delta\omega^2}} \sin\left(\frac{\sqrt{\Omega_R^2 + \delta\omega^2}}{2} t\right) \right] \\ a_2(t) &= -i e^{-i\delta\omega t/2} \frac{\Omega_R}{\sqrt{\Omega_R^2 + \delta\omega^2}} \sin\left(\frac{\sqrt{\Omega_R^2 + \delta\omega^2}}{2} t\right) \end{aligned} \quad (10.130)$$

che rispettano le condizioni $a_1(0) = 1$ e $a_2(0) = 0$. Per la probabilità di trovare il sistema nello stato $|2\rangle$ all'istante t abbiamo

$$P_2(t) = |a_2(t)|^2 = \frac{\Omega_R^2}{\Omega_R^2 + \delta\omega^2} \sin^2 \left(\frac{\sqrt{\Omega_R^2 + \delta\omega^2}}{2} t \right) = \frac{\Omega_R^2}{\Omega_R^2 + \delta\omega^2} \frac{1 - \cos \left(\sqrt{\Omega_R^2 + \delta\omega^2} t \right)}{2}, \quad (10.131)$$

mentre la probabilità di trovarlo nello stato $|1\rangle$ è

$$\begin{aligned} P_1(t) &= |a_1(t)|^2 = \frac{\delta\omega^2}{\Omega_R^2 + \delta\omega^2} + \frac{\Omega_R^2}{\Omega_R^2 + \delta\omega^2} \frac{1 + \cos \left(\sqrt{\Omega_R^2 + \delta\omega^2} t \right)}{2} \\ &= 1 - \frac{\Omega_R^2}{\Omega_R^2 + \delta\omega^2} \frac{1 - \cos \left(\sqrt{\Omega_R^2 + \delta\omega^2} t \right)}{2}. \end{aligned} \quad (10.132)$$

Abbiamo ancora oscillazioni delle popolazioni, alla frequenza $\Omega'_R = \sqrt{\Omega_R^2 + \delta\omega^2} > \Omega_R$. La popolazione dello stato $|2\rangle$, però, risulta sempre inferiore a 1, e ha valore massimo $P_2^{\max} = \Omega_R^2/(\Omega_R^2 + \delta\omega^2)$. L'occupazione dello stato $|1\rangle$ oscilla anche essa alla frequenza Ω'_R senza essere mai nulla, ma con valore minimo $P_1^{\min} = \delta\omega^2/(\Omega_R^2 + \delta\omega^2)$. Un'inversione completa della popolazione non è quindi possibile fuori risonanza, mentre, per $\delta\omega = 0$, le (10.131) e (10.132) si riducono alle (10.124).

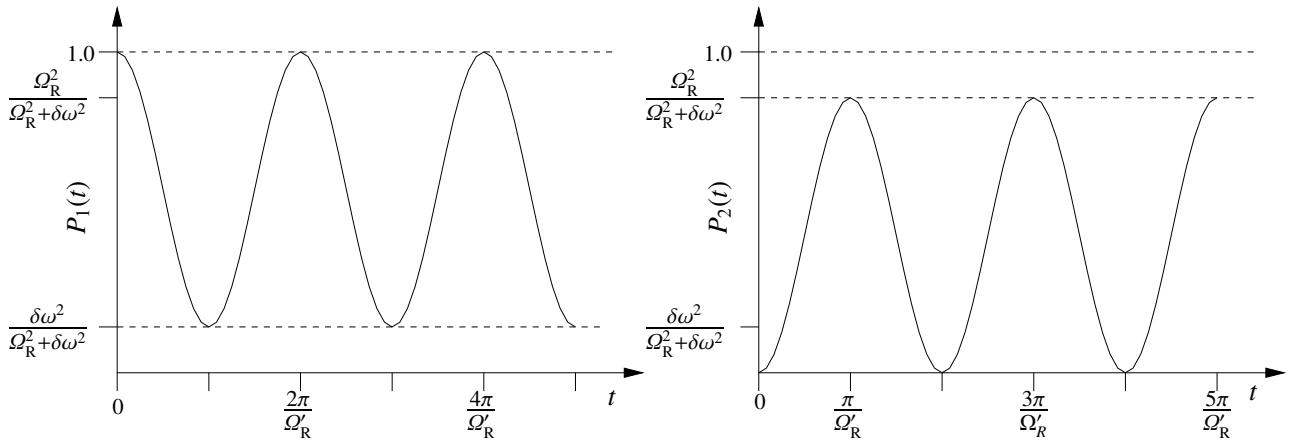


Figura 10.7 Oscillazioni di Rabi delle popolazioni $P_1(t)$ e $P_2(t)$ dei livelli inferiore $|1\rangle$ e superiore $|2\rangle$ della transizione. La frequenza di oscillazione è $\Omega'_R = \sqrt{\Omega_R^2 + \delta\omega^2}$, dove $\delta\omega = \omega - \omega_{21}$ è la differenza tra la frequenza della radiazione e la frequenza di risonanza della transizione. A tempi uguali a multipli dispari di π/Ω'_R gli atomi si trasferiscono completamente dal livello $|1\rangle$ al livello $|2\rangle$ solo se la radiazione è in risonanza, cioè se $\omega = \omega_{21}$ e, di conseguenza, $\Omega'_R = \Omega_R$.

Capitolo 11

Matrice densità

11.1 Matrice densità, popolazioni e coerenze

Consideriamo un atomo isolato, con hamiltoniana imperturbata $\hat{H}_0$, che all'istante t si trovi in un certo stato $|\psi(t)\rangle$. Ordiniamo gli autostati di $\hat{H}_0$ $|1\rangle, |2\rangle, \dots |n\rangle, \dots$, in ordine di energia crescente, in modo, cioè, che per i corrispondenti autovalori dell'energia $E_1, E_2, \dots E_n, \dots$ si abbia $E_n \leq E_{n+1}$ per ogni n . Poiché gli $\{|n\rangle\}$ costituiscono una base ortonormale completa, lo stato $|\psi(t)\rangle$ può essere scritto

$$|\psi(t)\rangle = \sum_n a_n(t) |n\rangle. \quad (11.1)$$

Definiamo *matrice densità dell'atomo singolo* nello stato $|\psi(t)\rangle$ l'operatore di proiezione

$$\hat{\rho} = |\psi(t)\rangle\langle\psi(t)| = \begin{pmatrix} |a_1|^2 & a_1 a_2^* & \dots & a_1 a_n^* & \dots \\ a_2 a_1^* & |a_2|^2 & \dots & a_2 a_n^* & \dots \\ \dots & \dots & \dots & \dots & \dots \\ a_n a_1^* & a_n a_2^* & \dots & |a_n|^2 & \dots \\ \dots & \dots & \dots & \dots & \dots \end{pmatrix}. \quad (11.2)$$

Passiamo adesso a un sistema di N atomi identici, tutti esattamente nello stesso stato $|\psi(t)\rangle$ dato dalla (11.1). Si dice che questo sistema si trova in uno *stato puro*, e la sua matrice densità si scrive esattamente come quella dell'atomo singolo rappresentata nella (11.2). Dalla condizione di normalizzazione dello stato (11.1) abbiamo che

$$\text{Tr}(\hat{\rho}) = 1. \quad (11.3)$$

Se vogliamo calcolare il valor medio di un qualunque operatore di atomo singolo $\hat{A}$ abbiamo

$$\langle A \rangle = \langle \psi(t) | \hat{A} | \psi(t) \rangle = \sum_{n,m} a_n^* \langle n | \hat{A} | m \rangle a_m = \sum_{n,m} \rho_{mn} \langle n | \hat{A} | m \rangle = \text{Tr}(\hat{\rho} \hat{A}). \quad (11.4)$$

L'evoluzione temporale della matrice densità si ottiene partendo dall'equazione di Schrödinger

$$\begin{aligned} \frac{d}{dt} \hat{\rho}(t) &= \left(\frac{d}{dt} |\psi(t)\rangle \right) \langle \psi(t) | + |\psi(t)\rangle \left(\frac{d}{dt} \langle \psi(t) | \right) \\ &= \frac{1}{i\hbar} \hat{H} |\psi(t)\rangle \langle \psi(t) | + \frac{1}{-i\hbar} |\psi(t)\rangle \langle \psi(t) | \hat{H} = \frac{1}{i\hbar} [\hat{H}, \hat{\rho}]. \end{aligned} \quad (11.5)$$

Inoltre abbiamo

$$\hat{\rho}^2 = |\psi(t)\rangle\langle\psi(t)||\psi(t)\rangle\langle\psi(t)| = \hat{\rho} \quad \text{da cui segue che} \quad \text{Tr}(\hat{\rho}^2) = 1. \quad (11.6)$$

Ma, se il nostro sistema si trova in uno stato puro, usare la matrice densità è del tutto equivalente a usare lo stato $|\psi(t)\rangle$, senza vantaggi né svantaggi. Le cose cambiano, e la matrice densità diventa uno strumento importante, se invece il sistema si trova in uno stato *non puro*, cioè se i suoi atomi non si trovano tutti nello stesso stato. Dato il numero di atomi che compongono un sistema macroscopico, non è pensabile conoscere lo stato in cui si trova ogni singolo atomo, e quindi lo stato dell'intero sistema. Come già visto nel paragrafo 2.2, l'unica possibilità è rinunciare alla conoscenza dell'informazione completa, e ricorrere a una trattazione probabilistica. Considereremo tutti gli atomi del sistema come a priori equivalenti, e ci accontenteremo di arrivare a determinare le probabilità p_i per ogni singolo atomo di trovarsi in ognuno di un certo insieme ortonormale e completo di stati

$$|\psi^{(i)}(t)\rangle = \sum_n a_n^{(i)} |n\rangle, \quad (11.7)$$

naturalmente con le condizioni $\sum_i p_i = 1$ e $p_i \geq 0$ per ogni i . Definiamo poi la matrice densità per il sistema in stato non puro come la *media pesata* delle matrici densità $\hat{\rho}^{(i)}$ relative ai singoli stati $|\psi^{(i)}(t)\rangle$, usando come pesi le probabilità p_i :

$$\hat{\rho} = \sum_i p_i \hat{\rho}^{(i)} = \sum_i p_i |\psi^{(i)}(t)\rangle\langle\psi^{(i)}(t)| = \sum_i p_i \begin{pmatrix} |a_1|^2 & a_1 a_2^* & \dots & a_1 a_n^* & \dots \\ a_2 a_1^* & |a_2|^2 & \dots & a_2 a_n^* & \dots \\ \dots & \dots & \dots & \dots & \dots \\ a_n a_1^* & a_n a_2^* & \dots & |a_n|^2 & \dots \\ \dots & \dots & \dots & \dots & \dots \end{pmatrix}. \quad (11.8)$$

Si può verificare che continuano a valere le proprietà

- 1) $\text{Tr}(\hat{\rho}) = \sum_i p_i \text{Tr}(|\psi^{(i)}(t)\rangle\langle\psi^{(i)}(t)|) = \sum_i p_i = 1;$
- 2) $\langle\hat{A}\rangle = \sum_i p_i \langle\psi^{(i)}(t)|\hat{A}|\psi^{(i)}(t)\rangle = \sum_i p_i \text{Tr}(|\psi^{(i)}(t)\rangle\langle\psi^{(i)}(t)|\hat{A}) = \text{Tr}(\hat{\rho}\hat{A});$
- 3) $\frac{d}{dt}\hat{\rho}(t) = \sum_i p_i \frac{d}{dt}(|\psi^{(i)}(t)\rangle\langle\psi^{(i)}(t)|) = \frac{1}{i\hbar} \sum_i p_i [\hat{H}, |\psi^{(i)}(t)\rangle\langle\psi^{(i)}(t)|] = \frac{1}{i\hbar} [\hat{H}, \hat{\rho}];$

mentre, nel caso di stato non puro, in generale abbiamo

$$\hat{\rho}^2 \neq \hat{\rho} \quad \text{e} \quad \text{Tr}(\hat{\rho}^2) < 1. \quad (11.9)$$

Gli elementi diagonali di $\hat{\rho}$ sono chiamati *popolazioni*, mentre gli elementi fuori diagonale sono chiamati *coerenze*. Nel caso di uno stato puro gli elementi diagonali sono le probabilità di occupazione degli stati. Nel caso di uno stato non puro gli elementi diagonali sono le probabilità medie di occupazione, ottenute facendo la media delle probabilità di occupazione corrispondenti ai singoli stati puri $|\psi^{(i)}(t)\rangle$ pesate con le probabilità p_i che l'atomo sia effettivamente nello stato $|\psi^{(i)}(t)\rangle$. Così, se abbiamo N atomi, il numero di atomi che ci aspettiamo di trovare nell'autostato $|n\rangle$ di $\hat{H}_0$ sarà

$$N_n = N\rho_{nn}. \quad (11.10)$$

Una condizione necessaria, ma non sufficiente, perché la coerenza ρ_{nm} tra gli stati $|n\rangle$ e $|m\rangle$ sia diversa da zero è che siano contemporaneamente diverse da zero le occupazioni dei due stati, indipendentemente dal fatto che il sistema si trovi o meno in uno stato puro. Per ogni stato $|\psi^{(i)}(t)\rangle$, di probabilità p_i , la singola ampiezza $a_n^{(i)}(t)$ può essere scritta nella forma

$$a_n^{(i)}(t) = b_n^{(i)}(t) e^{-iE_n t/\hbar}, \quad (11.11)$$

in modo che, finché il singolo atomo può essere considerato isolato, $b_n^{(i)}(t)$ non dipende dal tempo. Per le coerenze delle $\hat{\rho}^{(i)}$ relative ai singoli stati $|\psi^{(i)}(t)\rangle$ abbiamo quindi

$$\rho_{nm}^{(i)} = a_n^{(i)} a_m^{(i)*} = b_n^{(i)} b_m^{(i)*} e^{-i\omega_{nm} t}, \quad \text{e} \quad \rho_{mn}^{(i)} = a_m^{(i)} a_n^{(i)*} = b_m^{(i)} b_n^{(i)*} e^{i\omega_{nm} t}, \quad (11.12)$$

dove abbiamo posto $\omega_{nm} = (E_n - E_m)/\hbar$, mentre per gli elementi fuori diagonale di $\hat{\rho}$ abbiamo

$$\rho_{nm} = \sum_i p_i b_n^{(i)} b_m^{(i)*} e^{i\omega_{nm} t}, \quad \text{e} \quad \rho_{mn} = \sum_i p_i b_m^{(i)} b_n^{(i)*} e^{-i\omega_{nm} t} = \rho_{nm}^*. \quad (11.13)$$

In linea di principio, $b_n^{(i)}$ e $b_m^{(i)}$ sono numeri complessi con fasi qualunque, e le somme della (11.13) possono essere nulle anche se i singoli $b_n^{(i)}$ e $b_m^{(i)}$ sono tutti diversi da zero. Gli elementi di matrice ρ_{nm} , con $n \neq m$, saranno diversi da zero solo se la differenza di fase tra i $b_n^{(i)}$ e i $b_m^{(i)}$ non cambia troppo passando da uno stato $|\psi^{(i)}(t)\rangle$ a un altro $|\psi^{(j)}(t)\rangle$, da qui il nome di *coerenze* a questi elementi di matrice.

Per avere un esempio dell'importanza delle coerenze supponiamo di avere un campione di N atomi, e di voler calcolare il valor medio della componente del momento del dipolo elettrico del campione lungo una direzione definita da un certo versore $\mathbf{e}$, cioè di voler calcolare $\langle \mathbf{e} \cdot \hat{\mathbf{P}} \rangle$. Per semplicità, consideriamo un sistema di atomi a due livelli e mettiamoci nelle condizioni della (10.104), supponendo che, per il singolo atomo, sia

$$D = \langle 2 | \mathbf{e} \cdot \hat{\mathbf{P}} | 1 \rangle = \langle 1 | \mathbf{e} \cdot \hat{\mathbf{P}} | 2 \rangle,$$

per cui possiamo scrivere

$$\mathbf{e} \cdot \hat{\mathbf{P}} = \begin{pmatrix} 0 & D \\ D & 0 \end{pmatrix}. \quad (11.14)$$

Per il valor medio sul sistema abbiamo

$$\langle \mathbf{e} \cdot \hat{\mathbf{P}} \rangle = N \text{Tr}(\hat{\rho} \mathbf{e} \cdot \hat{\mathbf{P}}) = N \text{Tr} \begin{pmatrix} \rho_{12} D & \rho_{11} D \\ \rho_{22} D & \rho_{21} D \end{pmatrix} = N(\rho_{12} D + \rho_{21} D), \quad (11.15)$$

e vediamo che il valor medio della componente lungo $\mathbf{e}$ del dipolo elettrico è determinata dalle coerenze. Ovviamente, al posto di $\mathbf{e}$ possiamo prendere ognuno dei tre versori $\hat{\mathbf{x}}$, $\hat{\mathbf{y}}$ e $\hat{\mathbf{z}}$. Solo se gli atomi si trovano in una sovrapposizione coerente di $|1\rangle$ e $|2\rangle$ il valor medio del dipolo elettrico può essere diverso da zero. E' da notare che, in base alle (11.13), le coerenze hanno una dipendenza temporale del tipo $e^{\pm i\omega_{21} t}$, quindi il valor medio delle componenti del momento di dipolo elettrico, se diverse da zero, hanno un andamento temporale di tipo oscillante.

Se abbiamo un sistema di N atomi all'equilibrio termodinamico in assenza di campi esterni, ci aspettiamo che il valor medio di qualunque componente del momento di dipolo elettrico sia nullo. Questo non accade se il sistema si trova in uno stato puro, perché nella matrice densità corrispondente

a uno stato puro, se le popolazioni di $|n\rangle$ e $|m\rangle$ sono ambedue diverse da zero, le coerenze ρ_{nm} e ρ_{nm} sono necessariamente diverse da zero. Ma è tranquillamente possibile costruire uno stato non puro tale che le popolazioni siano diverse da zero, mentre le coerenze sono nulle. Come esempio, basta pensare a degli stati del tipo

$$|\psi_l\rangle = ae^{-iE_1 t/\hbar}|1\rangle + be^{i\varphi_l}e^{-iE_2 t/\hbar}|2\rangle, \quad (11.16)$$

dove φ_l è una fase casuale compresa tra 0 e 2π . Per le popolazioni avremo $\rho_{11} = |a|^2$ e $\rho_{22} = |b|^2$, mentre per la coerenza ρ_{12} abbiamo

$$\rho_{12} = \frac{1}{2\pi} \int_0^{2\pi} d\varphi_l \langle 1|\psi_l\rangle\langle\psi_l|2\rangle = \frac{ab^*}{2\pi} e^{-i(E_1-E_2)t/\hbar} \int_0^{2\pi} e^{-i\varphi_l} d\varphi_l = 0 \quad (11.17)$$

perché la fase φ_l è casuale.

E' da notare che la distinzione tra popolazioni e coerenze non ha valore assoluto, ma dipende dalla scelta della base per lo spazio di Hilbert. Infatti, essendo $\hat{\rho}$ hermitiana, è sempre possibile trovare una base ortonormale $\{|\chi_i\rangle\}$ in cui $\hat{\rho}$ stessa è diagonale. In questa base esistono solo popolazioni, e tutte le coerenze sono nulle. Va però notato che gli elementi di questa base possono non essere stati stazionari.

11.2 Modello di Feynman-Vernon-Hellwarth

Abbiamo visto che l'evoluzione temporale della matrice densità è data dall'equazione

$$\frac{d}{dt}\hat{\rho} = \frac{1}{i\hbar} [\hat{H}, \hat{\rho}]. \quad (11.18)$$

Per un sistema a due livelli in interazione con un campo elettromagnetico oscillante a frequenza ω abbiamo, riprendendo le notazioni e le ipotesi del paragrafo 10.5, compresa l'approssimazione di onda rotante, ma scrivendo il campo $\mathbf{E}_0 \cos \omega t$ anziché $-\mathbf{E}_0 \cos \omega t$,

$$\hat{H} = \hat{H}_0 + \hat{V} \quad (11.19)$$

con

$$\hat{H}_0 = \begin{pmatrix} E_1 & 0 \\ 0 & E_2 \end{pmatrix} \quad \text{e} \quad \hat{V} = \begin{pmatrix} 0 & -\frac{DE_0}{2}e^{+i\omega t} \\ -\frac{DE_0}{2}e^{-i\omega t} & 0 \end{pmatrix} = \begin{pmatrix} 0 & -\frac{\hbar\Omega_R}{2}e^{+i\omega t} \\ -\frac{\hbar\Omega_R}{2}e^{-i\omega t} & 0 \end{pmatrix}. \quad (11.20)$$

Inserendo nella (11.18) abbiamo

$$\begin{aligned} \frac{d}{dt}\hat{\rho} &= \frac{1}{i\hbar} \begin{pmatrix} E_1 & -\frac{\hbar\Omega_R}{2}e^{+i\omega t} \\ -\frac{\hbar\Omega_R}{2}e^{-i\omega t} & E_2 \end{pmatrix} \begin{pmatrix} \rho_{11} & \rho_{12} \\ \rho_{21} & \rho_{22} \end{pmatrix} - \frac{1}{i\hbar} \begin{pmatrix} \rho_{11} & \rho_{12} \\ \rho_{21} & \rho_{22} \end{pmatrix} \begin{pmatrix} E_1 & -\frac{\hbar\Omega_R}{2}e^{+i\omega t} \\ -\frac{\hbar\Omega_R}{2}e^{-i\omega t} & E_2 \end{pmatrix} \\ &= \begin{pmatrix} -i\frac{\Omega_R}{2}(e^{+i\omega t}\rho_{21} - e^{-i\omega t}\rho_{12}) & i\left[\omega_{21}\rho_{12} + \frac{\Omega_R}{2}(\rho_{11} - \rho_{22})e^{+i\omega t}\right] \\ -i\left[\omega_{21}\rho_{21} + \frac{\Omega_R}{2}(\rho_{11} - \rho_{22})e^{-i\omega t}\right] & i\frac{\Omega_R}{2}(e^{+i\omega t}\rho_{21} - e^{-i\omega t}\rho_{12}) \end{pmatrix}. \end{aligned} \quad (11.21)$$

La matrice densità di un sistema a due livelli ha tre gradi di libertà. Infatti gli elementi diagonali ρ_{11} e ρ_{22} sono due numeri reali, vincolati dalla condizione $\rho_{11} + \rho_{22} = 1$, e basta un solo numero reale, per esempio la differenza $\rho_{11} - \rho_{22}$, per determinarli. Gli elementi fuori diagonale sono numeri complessi, con $\rho_{21} = \rho_{12}^*$, e basta conoscere i due numeri reali $\text{Re}(\rho_{12})$ e $\text{Im}(\rho_{12})$ per determinare automaticamente anche ρ_{21} . I tre numeri reali $(\rho_{11} - \rho_{22})$, $\text{Re}(\rho_{12})$ e $\text{Im}(\rho_{12})$ possono essere combinati con la frequenza ω del campo elettromagnetico per ottenere le tre componenti, reali, di un vettore $\mathbf{R}$, detto *vettore di Feynman-Vernon-Hellwarth*, proposto appunto da Richard P. Feynman, Frank L. Vernon e Robert W. Hellwarth nel 1956, definite da

$$\begin{aligned} R_x &= \rho_{12}e^{-i\omega t} + \rho_{21}e^{+i\omega t} = 2 \text{Re}(\rho_{12}) \cos \omega t - 2 \text{Im}(\rho_{12}) \sin \omega t \\ R_y &= i\rho_{12}e^{-i\omega t} - i\rho_{21}e^{+i\omega t} = 2 \text{Re}(\rho_{12}) \sin \omega t - 2 \text{Im}(\rho_{12}) \cos \omega t \\ R_z &= \rho_{11} - \rho_{22}. \end{aligned} \quad (11.22)$$

Dalla (11.21) otteniamo per le derivate delle componenti di $\mathbf{R}$

$$\begin{aligned} \frac{d}{dt}R_x &= \frac{d\rho_{12}}{dt}e^{-i\omega t} - i\omega\rho_{12}e^{-i\omega t} + \frac{d\rho_{21}}{dt}e^{i\omega t} + i\omega\rho_{21}e^{i\omega t} \\ &= i \left[(\omega_{21} - \omega)\rho_{12} + \frac{\Omega_R}{2}(\rho_{11} - \rho_{22})e^{+i\omega t} \right] e^{-i\omega t} - i \left[(\omega_{21} - \omega)\rho_{21} + \frac{\Omega_R}{2}(\rho_{11} - \rho_{22})e^{-i\omega t} \right] e^{+i\omega t} \\ &= i(\omega_{21} - \omega)(\rho_{12}e^{-i\omega t} - \rho_{21}e^{+i\omega t}) \\ \frac{d}{dt}R_y &= i \frac{d\rho_{12}}{dt} e^{-i\omega t} + \omega\rho_{12}e^{-i\omega t} - i \frac{d\rho_{21}}{dt} e^{i\omega t} + \omega\rho_{21}e^{i\omega t} \\ &= - \left[(\omega_{21} - \omega)\rho_{12} + \frac{\Omega_R}{2}(\rho_{22} - \rho_{11})e^{+i\omega t} \right] e^{-i\omega t} - \left[(\omega_{21} - \omega)\rho_{21} + i \frac{\Omega_R}{2}(\rho_{11} - \rho_{22})e^{-i\omega t} \right] e^{+i\omega t} \\ &= -(\omega_{21} - \omega)(\rho_{12}e^{-i\omega t} + \rho_{21}e^{+i\omega t}) + \Omega_R(\rho_{11} - \rho_{22}) \\ \frac{d}{dt}R_z &= \frac{d\rho_{11}}{dt} - \frac{d\rho_{22}}{dt} = -i\Omega_R(e^{+i\omega t}\rho_{21} - e^{-i\omega t}\rho_{12}) \end{aligned} \quad (11.23)$$

che, confrontando con le (11.22), possono essere riscritte

$$\begin{aligned} \frac{d}{dt}R_x &= (\omega_{21} - \omega)R_y \\ \frac{d}{dt}R_y &= -(\omega_{21} - \omega)R_x + \Omega_R R_z \\ \frac{d}{dt}R_z &= \Omega_R R_y. \end{aligned} \quad (11.24)$$

Se adesso introduciamo il vettore $\mathbf{B}$, di componenti

$$B_x = \Omega_R, \quad B_y = 0, \quad B_z = \delta\omega = (\omega_{21} - \omega) \quad (11.25)$$

le (11.24) sono riassunte dall'equazione

$$\frac{d}{dt}\mathbf{R} = \mathbf{R} \times \mathbf{B} \quad (11.26)$$

che rappresenta un'evoluzione giroscopica del vettore di Feynman-Vernon-Hellwarth $\mathbf{R}$ attorno al vettore $\mathbf{B}$. Se all'istante $t = 0$ il sistema è all'equilibrio termodinamico in assenza di interazione con il campo elettromagnetico, le coerenze sono nulle, quindi $R_x(t = 0) = R_y(t = 0) = 0$, mentre R_z è determinata dalle occupazioni dei livelli secondo la

$$R_z(t = 0) = \frac{1}{N} [N_1^0 - N_2^0], \quad (11.27)$$

dove

$$\begin{aligned} N_1^0 &= N_1(t = 0) = \frac{N}{Z} e^{-E_1/kT} \\ N_2^0 &= N_2(t = 0) = \frac{N}{Z} e^{-E_2/kT}. \end{aligned} \quad (11.28)$$

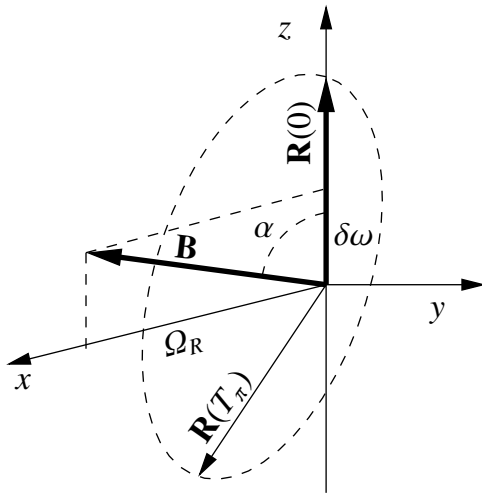


Figura 11.1 Evoluzione del vettore $\mathbf{R}$ di Feynman-Vernon-Hellwarth.

Se la separazione in energia tra i due livelli è nel range ottico, avremo $(E_2 - E_1) \gg kT$, ed è una buona approssimazione porre

$$N_1^0 \simeq N, \quad N_2^0 \simeq 0, \quad \text{da cui} \quad R_z(t = 0) \simeq 1. \quad (11.29)$$

Se a $t = 0$ accendiamo il campo elettromagnetico, il vettore $\mathbf{R}$ parte allineato all'asse z , e comincia a precessere attorno a $\mathbf{B}$ conservando la sua lunghezza. Il moto di precessione cambia progressivamente l'orientazione di $\mathbf{R}$ nello spazio delle coerenze/popolazioni. Per esempio, nel caso di perfetta risonanza tra radiazione e frequenza di transizione da $|1\rangle$ a $|2\rangle$, abbiamo $\omega = \omega_{21}$ e $\delta\omega = 0$, e $\mathbf{B}$ è diretto lungo l'asse x . La precessione di $\mathbf{R}$ avviene così attorno all'asse x nel piano yz . Dopo un impulso di durata $T_\pi = \pi/\Omega_R$, detto impulso π , il vettore $\mathbf{R}$ giace di nuovo sull'asse z , ma è diretto verso il basso. Si è così verificata l'inversione di popolazione descritta nel paragrafo 10.5. Un impulso $\pi/2$, di durata $T_{\pi/2} = \pi/(2\Omega_R)$, allinea invece $\mathbf{R}$ con l'asse y , portando così a una differenza di popolazione uguale a zero e a un massimo del valore delle coerenze. Nel caso di un valore di $\delta\omega \neq 0$ ritroviamo i risultati delle (10.130), (10.131) e (10.132), ottenuti per radiazione fuori risonanza nel paragrafo 10.5 sull'interazione con un campo forte. Infatti, fuori risonanza il vettore $\mathbf{B}$ forma un angolo $\alpha = \arccos(\delta\omega / \sqrt{\Omega_R^2 + \delta\omega^2})$ con l'asse z , e quindi con il vettore $\mathbf{R}$ all'istante $t = 0$, come mostrato in fig. 11.1. Dopo un impulso di durata T_π il vettore $\mathbf{R}$ avrà fatto mezzo giro di precessione attorno a $\mathbf{B}$, e formerà un angolo 2α con l'asse z . La sua proiezione sull'asse z sarà

$$R_z = \cos(2\alpha) = \cos^2 \alpha - \sin^2 \alpha = -\frac{\Omega_R^2 - \delta\omega^2}{\delta\omega^2 + \Omega_R^2}. \quad (11.30)$$

11.3 Rilassamento e equazioni di Bloch ottiche

Il modello di Feynman del paragrafo precedente descrive l'evoluzione delle popolazioni e delle coerenze di un sistema atomico a due livelli sotto l'azione di un campo elettromagnetico oscillante. Partendo da una situazione di equilibrio termodinamico con $R_x = R_y = 0$ e $R_z = 1$, l'azione del campo modifica le popolazioni e crea coerenze non nulle. Supponiamo adesso di spegnere il campo dopo aver raggiunto una situazione di questo tipo. Ci aspettiamo che il sistema ritorni progressivamente a uno stato di equilibrio termodinamico. Le equazioni che abbiamo studiato finora non includono l'interazione tra il sistema e il bagno termico, che è quella che riporta il sistema all'equilibrio. Alle equazioni vanno quindi aggiunti dei termini, detti *termini di rilassamento*, che descrivono, appunto, l'interazione del nostro sistema atomico con il bagno termico. Qui introdurremo i termini di rilassamento in modo fenomenologico, ma ricordiamo che la loro espressione può essere derivata in modo rigoroso formulando una teoria quantistica dell'interazione microscopica tra il nostro sistema e il bagno termico. In modo fenomenologico, i termini di rilassamento si scrivono

$$\begin{aligned}\frac{dR_x}{dt} &= -\gamma_{\perp} R_x \\ \frac{dR_y}{dt} &= -\gamma_{\perp} R_y \\ \frac{dR_z}{dt} &= -\gamma_{\parallel} (R_z - R_z^0)\end{aligned}\tag{11.31}$$

dove R_z^0 è il valore di R_z all'equilibrio termico. Le definizioni di *rilassamento perpendicolare* (o *rilassamento trasversale*), con coefficiente $\gamma_{\perp}$, e *rilassamento parallelo* (o *rilassamento longitudinale*), con coefficiente $\gamma_{\parallel}$, sono legate al fenomeno della risonanza magnetica, in cui questi termini sono stati introdotti per la prima volta da Felix Bloch. I processi fisici che modificano le coerenze, quantificati nel valore di $\gamma_{\perp}$, distruggono le fasi relative dei coefficienti delle espansioni degli stati del sistema in termini di autostati di $\hat{H}_0$, senza coinvolgere scambi di energia. Invece i fenomeni che variano le popolazioni, quantificati nel valore di $\gamma_{\parallel}$, richiedono uno scambio di energia tra il sistema ed il bagno termico. La terza delle (11.31) differisce dalle prime due, oltre che per il diverso coefficiente di rilassamento implicato, per la presenza del termine R_z^0 . Questo corrisponde al fatto che all'equilibrio termico le coerenze sono nulle, mentre esiste una differenza di popolazione diversa da zero data dalla distribuzione di Boltzmann.

Un caso importante di decadimento esponenziale delle popolazioni e delle coerenze è rappresentato dall'emissione spontanea. Se l'emissione spontanea è l'unico processo che agisce sul nostro sistema abbiamo

$$\begin{aligned}\frac{dN_1}{dt} &= A_{12} N_2 \\ \frac{dN_2}{dt} &= -A_{12} N_2\end{aligned}\tag{11.32}$$

che portano all'equazione per R_z (rilassamento longitudinale)

$$\frac{dR_z}{dt} = A_{12}(1 - R_z),\tag{11.33}$$

analoga all'ultima delle (11.31). In questo caso abbiamo così

$$\gamma_{\parallel} = A_{12}. \quad (11.34)$$

Il caso del rilassamento perpendicolare è più delicato. La presenza delle coerenze ρ_{12} e ρ_{21} richiede che sia il livello $|1\rangle$ che il livello $|2\rangle$ siano popolati. Se la popolazione del livello a energia più alta $|2\rangle$ tende a zero, tendono a zero anche ρ_{12} e ρ_{21} indipendentemente dalla popolazione di $|1\rangle$. Se l'unica causa di rilassamento è l'emissione spontanea abbiamo

$$\frac{d|a_2|^2}{dt} = -A_{12}|a_2|^2, \quad \text{che può essere ottenuta da} \quad \frac{da_2}{dt} = -\frac{A_{12}}{2}a_2, \quad (11.35)$$

infatti

$$\frac{d}{dt}(a_2 a_2^*) = \frac{da_2}{dt} a_2^* + \frac{da_2^*}{dt} a_2 = -\frac{A_{12}}{2} a_2 a_2^* - \frac{A_{12}}{2} a_2^* a_2 = -A_{12} |a_2|^2. \quad (11.36)$$

Questo ci porta a porre

$$\frac{d\rho_{12}}{dt} = -\frac{A_{12}}{2}\rho_{12}, \quad \frac{d\rho_{21}}{dt} = -\frac{A_{12}}{2}\rho_{21}, \quad \gamma_{\perp} = \frac{A_{12}}{2}. \quad (11.37)$$

E' importante notare che, nel caso di rilassamento dovuto all'emissione spontanea, il rilassamento trasversale è più lento del rilassamento longitudinale, mentre per rilassamento dovuto all'interazione con il bagno termico tipicamente abbiamo $\gamma_{\perp} \geq \gamma_{\parallel}$.

In presenza di campi elettromagnetici, l'evoluzione del sistema è dovuta all'interazione simultanea del sistema con i campi e con il bagno termico che determina il rilassamento. Una trattazione rigorosa del caso generale è piuttosto complessa. Ma, se la radiazione elettromagnetica non è troppo forte, è una buona approssimazione sommare semplicemente i due contributi all'evoluzione, ottenendo le *equazioni di Bloch ottiche* (da confrontare con le equazioni di Bloch che vedremo nel paragrafo 14.5 per la risonanza magnetica)

$$\begin{aligned} \frac{d}{dt}R_x &= (\omega_{21} - \omega)R_y - \gamma_{\perp}R_x \\ \frac{d}{dt}R_y &= -(\omega_{21} - \omega)R_x + \Omega_R R_z - \gamma_{\perp}R_y \\ \frac{d}{dt}R_z &= -\Omega_R R_y - \gamma_{\parallel}(R_z - R_z^0). \end{aligned} \quad (11.38)$$

La soluzione delle (11.38) permette di determinare l'evoluzione temporale di popolazioni, coerenze e polarizzazione in un esperimento che possa essere approssimato come l'interazione di un sistema a due livelli con un'onda elettromagnetica monocromatica e un bagno termico. Mentre la soluzione transiente, che descrive il passaggio da una generica condizione iniziale allo stato stazionario, è abbastanza complicata, il calcolo della soluzione stazionaria è abbastanza semplice. In condizioni stazionarie dobbiamo avere

$$\frac{d}{dt}R_x = \frac{d}{dt}R_y = \frac{d}{dt}R_z = 0, \quad (11.39)$$

quindi le (11.38) diventano

$$\begin{aligned} R_x &= \frac{\omega_{21} - \omega}{\gamma_{\perp}} R_y \\ R_y &= -\frac{\omega_{21} - \omega}{\gamma_{\perp}} R_x + \frac{\Omega_R}{\gamma_{\perp}} R_z \\ R_z &= R_z^0 - \frac{\Omega_R}{\gamma_{\parallel}} R_y \end{aligned} \quad (11.40)$$

Combinando le prime due delle (11.40) otteniamo R_y in funzione di R_z

$$R_y = \frac{\gamma_{\perp}}{\gamma_{\perp}^2 + (\omega_{21} - \omega)^2} \Omega_R R_z, \quad (11.41)$$

che, inserita nell'ultima delle (11.40), ci dà la soluzione per R_z

$$R_z = \frac{\gamma_{\parallel} [\gamma_{\perp}^2 + (\omega_{21} - \omega)^2]}{\gamma_{\parallel} [\gamma_{\perp}^2 + (\omega_{21} - \omega)^2] + \Omega_R^2 \gamma_{\perp}} R_z^0. \quad (11.42)$$

Questa, a sua volta, ci permette di calcolare la soluzione per R_y

$$R_y = \frac{\gamma_{\parallel} \gamma_{\perp} \Omega_R}{\gamma_{\parallel} [\gamma_{\perp}^2 + (\omega_{21} - \omega)^2] + \Omega_R^2 \gamma_{\perp}} R_z^0, \quad (11.43)$$

e infine, nota R_y , possiamo ricavare R_x dalla prima delle (11.40)

$$R_x = \frac{(\omega_{21} - \omega) \gamma_{\parallel} \Omega_R}{\gamma_{\parallel} [\gamma_{\perp}^2 + (\omega_{21} - \omega)^2] + \Omega_R^2 \gamma_{\perp}} R_z^0. \quad (11.44)$$

Dalla (10.74) abbiamo visto che a un'onda elettromagnetica il cui campo elettrico va come $E_0 \cos \omega t$ è associata una densità di energia per unità di volume

$$w = \varepsilon_0 \frac{E_0^2}{2}, \quad (11.45)$$

e quindi un'intensità, cioè una potenza per unità di superficie

$$I = c w = c \varepsilon_0 \frac{E_0^2}{2}. \quad (11.46)$$

Se introduciamo l'intensità di saturazione I_s (ne vedremo il significato qui sotto), definita come

$$I_s = \frac{\hbar^2}{2} \varepsilon_0 \frac{c \gamma_{\parallel} \gamma_{\perp}}{D^2}, \quad (11.47)$$

abbiamo per il rapporto I/I_s

$$\frac{I}{I_s} = c \varepsilon_0 \frac{E_0^2}{2} \frac{2}{\varepsilon_0 \hbar^2} \frac{D^2}{c \gamma_{\parallel} \gamma_{\perp}} = \frac{\Omega_R^2}{\gamma_{\parallel} \gamma_{\perp}}, \quad (11.48)$$

da cui ricaviamo $\Omega_R^2 = \gamma_{\parallel}\gamma_{\perp}I/I_s$. Sostituendo nelle (11.42-11.44) ottenendo

$$\begin{aligned} R_x &= \Omega_R \frac{\omega_{21} - \omega}{(\omega_{21} - \omega)^2 + \gamma_{\perp}^2 \left(1 + \frac{I}{I_s}\right)} R_z^0 \\ R_y &= \Omega_R \frac{\gamma_{\perp}}{(\omega_{21} - \omega)^2 + \gamma_{\perp}^2 \left(1 + \frac{I}{I_s}\right)} R_z^0 \\ R_z &= \frac{(\omega_{21} - \omega)^2 + \gamma_{\perp}^2}{(\omega_{21} - \omega)^2 + \gamma_{\perp}^2 \left(1 + \frac{I}{I_s}\right)} R_z^0 = \left[1 - \frac{\gamma_{\perp}^2 (I/I_s)}{(\omega_{21} - \omega)^2 + \gamma_{\perp}^2 \left(1 + \frac{I}{I_s}\right)} \right] R_z^0. \end{aligned} \quad (11.49)$$

Le dipendenze di R_x e di R_y dalla frequenza della radiazione ω hanno così le forme, rispettivamente, di una lorentziana di dispersione e di una lorentziana di assorbimento, ambedue centrate attorno alla frequenza di risonanza ω_{21} e ambedue di larghezza $\delta\omega = \gamma_{\perp} \sqrt{1 + I/I_s}$. Invece R_z ha la forma di una lorentziana di assorbimento rovesciata, sempre centrata attorno a ω_{21} e sempre di larghezza $\delta\omega$. L'espressione per R_z ci chiarisce il significato dell'intensità di saturazione I_s . Consideriamo radiazione alla frequenza di risonanza $\omega = \omega_{21}$. Per $I = 0$, cioè in assenza di campo, abbiamo $R_z = R_z^0$, mentre al limite $I \rightarrow \infty$ l'interazione sistema-campo elettromagnetico tende ad uguagliare le popolazioni, portando a $R_z = 0$. Se invece applichiamo radiazione di intensità pari a I_s l'ultima delle (11.49) ci dà $R_z^{\text{sat}} = R_z^0/2$. L'intensità di saturazione I_s corrisponde così al valore dell'intensità a cui la differenza di popolazione è esattamente intermedia tra i due casi estremi. Per frequenze ottiche $\hbar\omega \gg kT$, e in assenza di radiazione solo lo stato più basso è popolato: $N_1^0 \simeq N$, $N_2^0 \simeq 0$, e $R_z^0 = 1$, così che per $I = I_s$ abbiamo

$$R_z^{\text{sat}} = \frac{R_z^0}{2} = \frac{1}{2} = \frac{N_1^{\text{sat}} - N_2^{\text{sat}}}{N}, \quad \text{da cui} \quad N_1^{\text{sat}} = \frac{3}{4}, \quad N_2^{\text{sat}} = \frac{1}{4}, \quad N_1^{\text{sat}} - N_2^{\text{sat}} = \frac{N}{2}. \quad (11.50)$$

Nel caso più generale abbiamo $0 < N_2^0 < N_1^0 < N$, e in risonanza con $I = I_s$ troviamo

$$R_z^{\text{sat}} = \frac{R_z^0}{2} = \frac{N_1^0 - N_2^0}{2N} = \frac{N_1^{\text{sat}} - N_2^{\text{sat}}}{N}, \quad \text{da cui} \quad N_1^{\text{sat}} - N_2^{\text{sat}} = \frac{N_1^0 - N_2^0}{2}. \quad (11.51)$$

Se consideriamo anche radiazione fuori risonanza l'espressione per R_z diventa

$$R_z = \frac{(\omega_{21} - \omega)^2 + \gamma_{\perp}^2}{(\omega_{21} - \omega)^2 + \gamma_{\perp}^2 \left(1 + \frac{I}{I_s}\right)} \frac{N_1^0 - N_2^0}{N} \quad (11.52)$$

da cui otteniamo per la differenza di popolazione

$$N_1 - N_2 = N R_z = \frac{(\omega_{21} - \omega)^2 + \gamma_{\perp}^2}{(\omega_{21} - \omega)^2 + \gamma_{\perp}^2 \left(1 + \frac{I}{I_s}\right)} (N_1^0 - N_2^0). \quad (11.53)$$

In caso di risonanza ($\omega = \omega_{21}$) quest'ultima espressione si riduce a

$$N_1 - N_2 = \frac{N_1^0 - N_2^0}{1 + I/I_s}. \quad (11.54)$$

Per le popolazioni dei due livelli abbiamo nel caso generale

$$\begin{aligned} N_1 &= \frac{N}{2} + \frac{(\omega_{21} - \omega)^2 + \gamma_\perp^2}{(\omega_{21} - \omega)^2 + \gamma_\perp^2 (1 + I/I_s)} \frac{(N_1^0 - N_2^0)}{2} \\ N_2 &= \frac{N}{2} - \frac{(\omega_{21} - \omega)^2 + \gamma_\perp^2}{(\omega_{21} - \omega)^2 + \gamma_\perp^2 (1 + I/I_s)} \frac{(N_1^0 - N_2^0)}{2}, \end{aligned} \quad (11.55)$$

e, in caso di risonanza

$$N_1 = \frac{1}{2} \left(N + \frac{N_1^0 - N_2^0}{1 + I/I_s} \right), \quad N_2 = \frac{1}{2} \left(N - \frac{N_1^0 - N_2^0}{1 + I/I_s} \right). \quad (11.56)$$

11.4 Equazioni di rate e forme di riga

Consideriamo un sistema a due livelli in presenza di un'onda elettromagnetica. Escludendo per il momento l'interazione con il bagno termico per la differenza di popolazione abbiamo

$$\frac{d}{dt}(N_1 - N_2) = -2W_{21}(N_1 - N_2), \quad (11.57)$$

dove abbiamo ricordato che $W_{21} = W_{12}$. Se adesso “accendiamo” anche l'interazione con il bagno termico abbiamo visto che la terza equazione di Bloch ottica si scrive

$$\frac{d}{dt}R_z = \frac{1}{N} \frac{d}{dt}(N_1 - N_2) = -\Omega_R R_y - \gamma_\parallel (R_z - R_z^0), \quad (11.58)$$

dove il primo termine dell'ultimo membro descrive l'interazione con il campo di radiazione. La soluzione stazionaria dell'equazione di Bloch per R_y è data dalla (11.41)

$$R_y = \frac{\gamma_\perp}{\gamma_\perp^2 + (\omega_{21} - \omega)^2} \Omega_R R_z,$$

che, sostituita nella (11.58), ci dà

$$\frac{d}{dt}R_z = -\Omega_R^2 \frac{\gamma_\perp}{\gamma_\perp^2 + (\omega_{21} - \omega)^2} R_z - \gamma_\parallel (R_z - R_z^0). \quad (11.59)$$

Separando la parte che riguarda l'interazione con la radiazione dalla parte contenente $\gamma_\parallel$ che descrive la variazione delle popolazioni per effetto dell'interazione con il bagno termico, la (11.59) diventa

$$\left. \frac{d}{dt}(N_1 - N_2) \right|_{\text{rad}} = -\Omega_R^2 \frac{\gamma_\perp}{\gamma_\perp^2 + (\omega_{21} - \omega)^2} (N_1 - N_2) \quad (11.60)$$

che, paragonata alla (11.57) ci dà

$$W_{21} = \frac{\Omega_R^2}{2} \frac{\gamma_{\perp}}{\gamma_{\perp}^2 + (\omega_{21} - \omega)^2} = \frac{D^2 E_0^2}{2\hbar^2} \frac{\gamma_{\perp}}{\gamma_{\perp}^2 + (\omega_{21} - \omega)^2}, \quad (11.61)$$

che corrisponde a una probabilità di transizione con dipendenza lorentziana dalla differenza tra la frequenza della radiazione ω e la frequenza di transizione ω_{21} . Tutti i conti svolti in questo paragrafo presuppongono che R_x ed R_y evolvano così rapidamente da raggiungere lo stato stazionario prima di R_z . Questo implica che sia $\gamma_{\perp} \gg \gamma_{\parallel}$, condizione normalmente verificata nell'interazione con il bagno termico.

11.5 Rate equations e bilancio dettagliato

Supponiamo di avere un sistema di N atomi, descritto da una matrice densità $\hat{\rho}$, con $\hat{H}_0$ hamiltoniana imperturbata del singolo atomo. La popolazione dell'autostato $|k\rangle$ di $\hat{H}_0$ che misuriamo all'istante t sarà

$$N_k = N\rho_{kk}(t). \quad (11.62)$$

Il valore di N_k al passare del tempo viene modificato dalle interazioni degli atomi tra di loro e degli atomi con l'ambiente esterno. Se per queste interazioni è possibile calcolare una *probabilità di transizione per unità di tempo* W_{km} dallo stato $|m\rangle$ allo stato $|k\rangle$, per ogni stato $|k\rangle$ avremo

$$\frac{d}{dt}\rho_{kk} = \sum_{m \neq k} W_{km}\rho_{mm} - \rho_{kk} \sum_{m \neq k} W_{mk}, \quad (11.63)$$

dove la prima sommatoria tiene conto della probabilità che l'atomo sia inizialmente in uno stato $|m\rangle$ diverso da $|k\rangle$ e passi in $|k\rangle$, mentre la seconda tiene conto della probabilità che l'atomo fosse inizialmente in $|k\rangle$ e da lì passi a un qualunque altro stato.

La condizione di normalizzazione della matrice densità implica

$$\sum_k \rho_{kk} = 1, \quad \text{da cui} \quad \sum_k \frac{d}{dt}\rho_{kk} = 0 \quad (11.64)$$

che, combinata con la (11.63), ci dà

$$\sum_k \sum_{m \neq k} W_{km}\rho_{mm} - \sum_k \rho_{kk} \sum_{m \neq k} W_{mk} = 0. \quad (11.65)$$

Se il numero totale di atomi è costante vale la (11.62), quindi, moltiplicando primo e secondo membro per N , otteniamo

$$\sum_k \sum_{m \neq k} W_{km}N_m - \sum_k N_k \sum_{m \neq k} W_{mk} = 0. \quad (11.66)$$

Se siamo in uno stato stazionario, come, per esempio, all'equilibrio termodinamico, non solo la popolazione totale, ma anche la popolazione di ogni singolo livello deve essere costante, quindi dobbiamo avere per ogni stato $|k\rangle$

$$\sum_{m \neq k} W_{km}N_m = N_k \sum_{m \neq k} W_{mk}. \quad (11.67)$$

Questo costituisce un sistema di equazioni, dette *rate equations*, che, note le W_{km} , permette di calcolare le popolazioni dei singoli livelli, cioè gli elementi diagonali della matrice densità, trascurando le coerenze.

In condizione di equilibrio termodinamico normalmente è valida la relazione più ristretta della (11.67), detta *equazione del bilancio dettagliato*,

$$W_{mk}N_k = W_{km}N_m, \quad (11.68)$$

che implica un equilibrio all'interno di ogni coppia di livelli $|m\rangle$ e $|k\rangle$ interessati dalla perturbazione: la velocità con cui gli atomi passano da $|m\rangle$ a $|k\rangle$ è uguale a quella con cui passano da $|k\rangle$ a $|m\rangle$. In altre parole, all'equilibrio ogni processo elementare è bilanciato dal suo processo inverso.

11.6 Accoppiamento atomo-bagno termico

Consideriamo un sistema di atomi a due livelli che interagisca con un bagno termico in assenza di radiazione. Il bagno termico può indurre transizioni tra i livelli atomici $|1\rangle$ e $|2\rangle$. L'equazione di rate per l'evoluzione della popolazione del livello atomico eccitato $|2\rangle$ si scrive

$$\frac{dN_2}{dt} = W_{21}N_1 - W_{12}N_2 \quad (11.69)$$

che ha come soluzione stazionaria

$$\frac{N_1^0}{N_2^0} = \frac{W_{12}}{W_{21}}. \quad (11.70)$$

Se il sistema atomi più bagno termico si trova all'equilibrio termodinamico ad una temperatura T , e supponiamo che le condizioni siano tali che sia soddisfatta la statistica di Maxwell-Boltzmann, deve anche essere

$$\frac{N_1^0}{N_2^0} = e^{(E_2-E_1)/kT}, \quad \text{da cui otteniamo} \quad \frac{W_{12}}{W_{21}} = e^{(E_2-E_1)/kT} = e^{\Delta E/kT}. \quad (11.71)$$

Questa relazione implica probabilità diverse, e dipendenti dalla temperatura del bagno termico, per i due versi della transizione atomica tra $|1\rangle$ e $|2\rangle$. A prima vista questo sembra in contrasto con quanto abbiamo calcolato trattando le transizioni tra stati atomici dovute ad interazioni con la radiazione elettromagnetica, o, più in generale, ad hamiltoniane di perturbazione dipendenti dal tempo. Infatti avevamo sempre trovato probabilità uguali, e indipendenti dalla temperatura, per le transizioni nei due versi, dato che $|\langle 1|\hat{V}|2\rangle|^2 = |\langle 2|\hat{V}|1\rangle|^2$.

La cosa però si spiega ricordando che adesso abbiamo a che fare con un sistema isolato costituito dall'insieme "sistema atomico + bagno termico", e che in questo sistema isolato deve conservarsi l'energia. Il bagno termico, per ipotesi, ha una capacità termica molto maggiore di quella del sistema atomico, così che, per interazioni con questo, subisce variazioni di energia trascurabili rispetto alla sua energia totale. Le proprietà del bagno termico non vengono quindi alterate sensibilmente negli scambi di energia con gli atomi del nostro sistema.

Perché ci sia conservazione dell'energia dobbiamo considerare non, per esempio, la W_{12} usata sopra, ma una W_{12}^{tot} che dia la probabilità per unità di tempo che contemporaneamente:

- i. l'atomo passi dallo stato ad energia più alta $|2\rangle$ a quello ad energia più bassa $|1\rangle$ cedendo l'energia $\Delta E = E_2 - E_1$, evento che ha probabilità $W_{12}^{\text{at}} = W_{21}^{\text{at}}$, e

- ii. il bagno termico assorba questa energia ΔE , passando da un certo stato del bagno $|\alpha\rangle$ a un altro stato di energia superiore $|\beta\rangle$, con $E_\beta - E_\alpha = \Delta E$. Questo evento avrà probabilità $W_{\beta\alpha}^{\text{res}}$.

Quella che noi osserviamo non è quindi semplicemente la probabilità di transizione atomica W_{12}^{at} , ma una probabilità di transizione “complessiva” $W_{12}^{\text{tot}} = W_{12}^{\text{at}} W_{\beta\alpha}^{\text{res}}$ (dove il suffisso “res” sta per *reservoir*, bagno termico). Essendo $W_{12}^{\text{at}} = W_{21}^{\text{at}}$ abbiamo

$$\frac{W_{12}^{\text{tot}}}{W_{21}^{\text{tot}}} = \frac{W_{12}^{\text{at}}}{W_{21}^{\text{at}}} \frac{W_{\beta\alpha}^{\text{res}}}{W_{\alpha\beta}^{\text{res}}} = \frac{W_{\beta\alpha}^{\text{res}}}{W_{\alpha\beta}^{\text{res}}}. \quad (11.72)$$

Mentre gli stati $|1\rangle$ e $|2\rangle$ sono stati atomici singoli, gli stati $|\alpha\rangle$ e $|\beta\rangle$ corrispondono in realtà a molteplicità di stati, data la complessità del bagno termico, che è un sistema macroscopico. La probabilità di passare da un singolo stato di $|\alpha\rangle$ a uno qualunque degli stati $|\beta\rangle$ è così uguale al modulo quadro dell’elemento di matrice dell’hamiltoniana di perturbazione per il numero degli stati finali $N^{\text{res}}(\beta)$, viceversa per il passaggio da $|\beta\rangle$ a $|\alpha\rangle$. Abbiamo quindi

$$\frac{W_{12}^{\text{tot}}}{W_{21}^{\text{tot}}} = \frac{W_{\beta\alpha}^{\text{res}}}{W_{\alpha\beta}^{\text{res}}} = \frac{N^{\text{res}}(\beta)}{N^{\text{res}}(\alpha)}. \quad (11.73)$$

Facendo il logaritmo della (11.73) otteniamo

$$\ln\left(\frac{W_{12}^{\text{tot}}}{W_{21}^{\text{tot}}}\right) = \ln[N^{\text{res}}(\beta)] - \ln[N^{\text{res}}(\alpha)] = \frac{1}{k}[S^{\text{res}}(\beta) - S^{\text{res}}(\alpha)] \quad (11.74)$$

dove con S^{res} abbiamo indicato l’entropia del bagno termico, e abbiamo sfruttato il fatto che

$$S(\alpha) = k \ln[N^{\text{res}}(\alpha)]. \quad (11.75)$$

Infatti, quando diciamo che il bagno termico si trova nello stato $|\alpha\rangle$, in realtà non sappiamo in quale di $N^{\text{res}}(\alpha)$ stati effettivamente si trovi. Abbiamo quindi un’informazione mancante relativa a $N^{\text{res}}(\alpha)$ casi equiprobabili. Mentre il bagno termico assorbe l’energia ΔE passando dallo stato $|\alpha\rangle$ allo stato $|\beta\rangle$ il suo volume e la sua temperatura restano praticamente costanti, per cui abbiamo

$$S^{\text{res}}(\beta) - S^{\text{res}}(\alpha) = \frac{\Delta E}{T}, \quad (11.76)$$

quindi

$$\ln\left(\frac{W_{21}^{\text{tot}}}{W_{12}^{\text{tot}}}\right) = \frac{\Delta E}{kT}, \quad \text{da cui} \quad \frac{W_{21}^{\text{tot}}}{W_{12}^{\text{tot}}} = e^{\Delta E/kT} \quad (11.77)$$

che dimostra la (11.71). In altre parole, il bagno termico fornisce al sistema atomico l’informazione sulla temperatura a causa delle transizioni all’interno del bagno termico stesso, ed alla variazione di informazione collegata, che necessariamente accompagnano le transizioni all’interno del sistema atomico.

11.7 Allargamento omogeneo

Nelle formule che descrivono l'interazione tra la radiazione elettromagnetica e la materia compare una *forma di riga* $\phi(\omega - \omega_{21})$, caratterizzata da una certa *larghezza di riga* $\Delta\omega = \gamma_{\perp}$, che dice come varia la probabilità dell'assorbimento in funzione della frequenza. La larghezza di riga può avere origini diverse, e può essere di tipo *omogeneo* o *non omogeneo*.

Si parla di allargamento omogeneo quando si ha la stessa forma di riga (inclusi centro e larghezza) per ogni atomo (o molecola) del sistema. Nella maggior parte dei casi l'allargamento omogeneo è dovuto alla durata finita del tempo di interazione tra la radiazione elettromagnetica e gli atomi assorbenti. Il tempo di interazione può essere limitato ad un valore T da processi di vario genere, come, per esempio, l'emissione spontanea o la durata finita dell'impulso di radiazione. Un'applicazione diretta del principio di indeterminazione di Heisenberg $\Delta E \Delta t > \hbar/2$ ci porta ad una larghezza della riga di assorbimento

$$\Delta\omega = \frac{\Delta E}{\hbar} = \frac{1}{2T}. \quad (11.78)$$

Un altro meccanismo di allargamento omogeneo, nel caso di un gas, è l'allargamento per pressione, dovuto dalle collisioni interatomiche che modificano i processi di interazione tra la radiazione elettromagnetica e gli atomi. Possiamo supporre che ad ogni collisione la fase relativa dell'interazione radiazione-atomo sia modificata bruscamente e cominci una nuova evoluzione coerente, con fase non correlata a quella di prima dell'urto. L'allargamento omogeneo è così determinato dal tempo medio tra due collisioni T_c , dato da

$$T_c = \frac{\ell}{v}, \quad (11.79)$$

dove ℓ è il cammino libero medio tra due collisioni e $v = \sqrt{\frac{8kT}{\pi\mu}}$, con $\mu = m_1 m_2 / (m_1 + m_2)$, la velocità relativa media tra due atomi di massa m_1 e m_2 che collidono. Il libero cammino medio per urti tra sfere rigide di raggio a , con n sfere per unità di volume, vale

$$\ell = \frac{1}{4\sqrt{2}\pi n a^2} \quad (11.80)$$

dove il fattore $\sqrt{2}$ appare quando viene fatta una integrazione sulla distribuzione delle velocità di Maxwell-Boltzmann. Otteniamo così

$$T_c = \frac{1}{4\sqrt{2}\pi n a^2} \sqrt{\frac{\pi\mu}{8kT}}. \quad (11.81)$$

Questa formula può essere scritta in funzione della pressione del gas. Dall'equazione dei gas

$$pV = NkT \quad (11.82)$$

otteniamo per il numero di atomi per unità di volume n

$$n = \frac{N}{V} = \frac{p}{kT}, \quad (11.83)$$

che, sostituito nella (11.81) ci dà

$$T_c = \frac{1}{4\sqrt{2}\pi a^2} \sqrt{\frac{\pi\mu}{8kT}} \frac{kT}{p} = \sqrt{\frac{\mu kT}{\pi}} \frac{1}{16a^2 p}, \quad \text{da cui} \quad \Delta\omega = \frac{1}{2T_c} = \sqrt{\frac{\pi}{\mu kT}} 8 a^2 p. \quad (11.84)$$

Sperimentalmente si osserva che, per diversi sistemi atomici e molecolari, l'allargamento per pressione è dato con buona approssimazione dall'espressione

$$\Delta\omega \simeq 8.357 \frac{\text{MHz}}{\text{Pa}} p \quad (11.85)$$

11.8 Allargamento disomogeneo: effetto Doppler

L'allargamento per effetto Doppler è dovuto alla componente della velocità dei singoli atomi lungo la direzione di propagazione della radiazione elettromagnetica. Consideriamo un'onda elettromagnetica di frequenza ν_0 che si propaghi lungo l'asse z del nostro sistema di riferimento, ed un atomo di massa m che si muova con velocità v_z opposta alla direzione di propagazione dell'onda. Per l'effetto Doppler, nel proprio sistema di quiete l'atomo vede un'onda di frequenza

$$\nu = \nu_0 \sqrt{\frac{1 + v_z/c}{1 - v_z/c}} \simeq \nu_0 \left(1 + \frac{v_z}{c}\right). \quad (11.86)$$

Perché l'onda venga assorbita bisogna che non ν_0 , ma ν sia uguale alla frequenza ν_{21} di separazione tra il livello $|1\rangle$ e il livello $|2\rangle$. Pertanto la radiazione sarà assorbita solo dagli atomi che hanno velocità lungo l'asse z

$$v_z = c \left(\frac{\nu_{21}}{\nu_0} - 1 \right) = c \frac{\nu_{21} - \nu_0}{\nu_0} \simeq c \frac{\nu_{21} - \nu_0}{\nu_{21}}, \quad (11.87)$$

dove abbiamo tenuto conto del fatto che $\nu_{21} \simeq \nu_0$. La probabilità che la componente z della velocità dell'atomo sia nell'intervallo $(v_z, v_z + dv_z)$ è

$$P(v_z) dv_z = \sqrt{\frac{m}{2\pi kT}} e^{-mv_z^2/2kT} dv_z, \quad (11.88)$$

sostituendo la (11.87), da cui ricaviamo anche che

$$dv_z = -c \frac{d\nu_0}{\nu_{21}},$$

otteniamo per la probabilità che un atomo possa assorbire radiazione nell'intervallo di frequenza $(\nu_0, \nu_0 + d\nu_0)$

$$\phi(\nu_0) d\nu_0 = \sqrt{\frac{m}{2\pi kT}} \exp\left[-\frac{m}{2kT} c^2 \frac{(\nu_{21} - \nu_0)^2}{\nu_{21}^2}\right] \frac{c}{\nu_{21}} d\nu_0. \quad (11.89)$$

Il valore massimo dell'esponenziale si ha per $\nu_0 = \nu_{21}$. La *larghezza Doppler* $\Delta\nu_D$ è definita in modo che per $\nu_0 = \nu_{21} \pm \Delta\nu_D$ l'esponenziale valga $1/2$. Abbiamo quindi

$$\frac{m}{2kT} c^2 \frac{\Delta\nu_D^2}{\nu_{21}^2} = \ln 2 \quad \text{da cui} \quad \Delta\nu_D = \frac{\nu_{21}}{c} \sqrt{\frac{2kT}{m} \ln 2}. \quad (11.90)$$

La (11.89) può così essere riscritta

$$\phi(\nu_0) = \sqrt{\frac{\ln 2}{\pi}} \frac{1}{\Delta\nu_D} e^{-(\nu_0 - \nu_{21})^2 / \Delta\nu_D^2}, \quad (11.91)$$

che ci dà una forma di riga Gaussiana.

11.9 Regole di selezione: considerazioni di simmetria

In questo paragrafo introdurremo alcuni concetti di simmetria per l'hamiltoniana imperturbata di un atomo. Per semplicità, consideriamo l'atomo di idrogeno col suo unico elettrone. L'equazione di Schrödinger indipendente dal tempo si scrive

$$\hat{H}_0(\mathbf{r})\psi_n(\mathbf{r}) = \left[-\frac{\hbar^2}{2m}\Delta + V(\mathbf{r}) \right] \psi_n(\mathbf{r}) = E_n\psi_n(\mathbf{r}). \quad (11.92)$$

Il potenziale generato dal nucleo è invariante per parità, cioè per trasformazioni che mandino $\mathbf{r}$ in $-\mathbf{r}$:

$$V(\mathbf{r}) = V(-\mathbf{r}). \quad (11.93)$$

D'altra parte, anche il laplaciano che compare nell'espressione dell'operatore relativo all'energia cinetica è invariante per parità perché

$$\frac{\partial^2}{\partial(-x_i)^2} = \frac{\partial^2}{\partial x_i^2}, \quad \text{per } x_i = x, y, z. \quad (11.94)$$

Quindi, in assenza di campi esterni, tutta l'hamiltoniana è invariante per parità.

Indipendentemente da questo, se effettuiamo il cambiamento di variabili $\mathbf{r} \rightarrow -\mathbf{r}$, da

$$\hat{H}_0(\mathbf{r})\psi_n(\mathbf{r}) = E_n\psi_n(\mathbf{r}) \quad \text{otteniamo} \quad \hat{H}_0(-\mathbf{r})\psi_n(-\mathbf{r}) = E_n\psi_n(-\mathbf{r}). \quad (11.95)$$

Ma dall'invarianza per parità abbiamo che $\hat{H}_0(-\mathbf{r}) = \hat{H}_0(\mathbf{r})$, quindi

$$\hat{H}_0(\mathbf{r})\psi_n(-\mathbf{r}) = E_n\psi_n(-\mathbf{r}), \quad (11.96)$$

che ci dice che $\psi_n(-\mathbf{r})$ e $\psi_n(\mathbf{r})$ sono autofunzioni di $\hat{H}_0$ corrispondenti allo stesso autovalore. Se siamo in assenza di degenerazione all'autovalore E_n corrisponde un solo autostato. Naturalmente sappiamo che se $\psi_n(\mathbf{r})$ è un'autofunzione di un operatore corrispondente ad un certo autostato $|n\rangle$, anche la funzione $\alpha\psi_n(\mathbf{r})$, con α numero complesso qualunque, è un'autofunzione corrispondente allo stato $|n\rangle$. Quindi deve essere

$$\psi_n(-\mathbf{r}) = \alpha\psi_n(\mathbf{r}). \quad (11.97)$$

Se in ambedue i membri sostituiamo $\mathbf{r}$ con $-\mathbf{r}$ otteniamo

$$\psi_n(\mathbf{r}) = \alpha\psi_n(-\mathbf{r}). \quad (11.98)$$

D'altra parte, se sostituiamo la (11.97) nella (11.98) abbiamo

$$\psi_n(-\mathbf{r}) = \alpha\psi_n(\mathbf{r}) = \alpha^2\psi_n(-\mathbf{r}) \quad (11.99)$$

da cui

$$\alpha^2 = 1, \quad \alpha = \pm 1, \quad \psi_n(-\mathbf{r}) = \pm\psi_n(\mathbf{r}). \quad (11.100)$$

In altre parole, in assenza di degenerazione, le autofunzioni di $\hat{H}_0$ hanno parità definita: o sono pari, o sono dispari.

Un'altra proprietà di $\hat{H}_0$ è la sua invarianza per rotazioni attorno a qualunque asse passante per il nucleo, per esempio, attorno all'asse z . Per vedere le conseguenze di questa simmetria cominciamo scrivendo l'equazione di Schrödinger indipendente dal tempo in coordinate sferiche

$$\hat{H}_0(r, \vartheta, \varphi) \psi_n(r, \vartheta, \varphi) = E_n \psi_n(r, \vartheta, \varphi), \quad (11.101)$$

e notiamo che, se ruotiamo il sistema di riferimento attorno all'asse z di un angolo arbitrario $-\varphi_1$, l'equazione diventa

$$\hat{H}_0(r, \vartheta, \varphi + \varphi_1) \psi_n(r, \vartheta, \varphi + \varphi_1) = E_n \psi_n(r, \vartheta, \varphi + \varphi_1). \quad (11.102)$$

D'altra parte per l'invarianza rotazionale dell'hamiltoniana abbiamo

$$\hat{H}_0(r, \vartheta, \varphi) = \hat{H}_0(r, \vartheta, \varphi + \varphi_1), \quad (11.103)$$

che, introdotta nella (11.102) ci dà

$$\hat{H}_0(r, \vartheta, \varphi) \psi_n(r, \vartheta, \varphi + \varphi_1) = E_n \psi_n(r, \vartheta, \varphi + \varphi_1). \quad (11.104)$$

Quindi $\psi_n(r, \vartheta, \varphi)$ e $\psi_n(r, \vartheta, \varphi + \varphi_1)$ devono essere autofunzioni corrispondenti allo stesso autovalore, e, in assenza di degenerazione, deve essere

$$\psi_n(r, \vartheta, \varphi + \varphi_1) = \alpha_{\varphi_1} \psi_n(r, \vartheta, \varphi), \quad (11.105)$$

con il fattore di proporzionalità α_{φ_1} che può dipendere dall'angolo φ_1 . Se facciamo una rotazione di φ_2 anziché di φ_1 abbiamo

$$\psi_n(r, \vartheta, \varphi + \varphi_2) = \alpha_{\varphi_2} \psi_n(r, \vartheta, \varphi). \quad (11.106)$$

Adesso sostituiamo $\varphi + \varphi_2$ al posto di φ nella (11.105), ottenendo

$$\psi_n(r, \vartheta, \varphi + \varphi_1 + \varphi_2) = \alpha_{\varphi_1} \psi_n(r, \vartheta, \varphi + \varphi_2), \quad (11.107)$$

che, in base alla (11.106), può essere riscritta

$$\psi_n(r, \vartheta, \varphi + \varphi_1 + \varphi_2) = \alpha_{\varphi_1} \alpha_{\varphi_2} \psi_n(r, \vartheta, \varphi), \quad (11.108)$$

da cui si ricava

$$\alpha_{\varphi_1 + \varphi_2} = \alpha_{\varphi_1} \alpha_{\varphi_2}. \quad (11.109)$$

Quest'ultima relazione è verificata solo se

$$\alpha_\varphi = e^{im\varphi}, \quad (11.110)$$

con m parametro da determinare. Ogni funzione deve essere trasformata in sé stessa da una rotazione di un angolo uguale a 2π . Abbiamo quindi la condizione

$$e^{im2\pi} = 1, \quad (11.111)$$

che implica che m sia un numero intero, positivo o negativo. Questa relazione ci permette anche di determinare la dipendenza delle ψ_n dalla coordinata sferica φ . Infatti abbiamo

$$\psi_n(r, \vartheta, 0 + \varphi) = \alpha_\varphi \psi_n(r, \vartheta, 0) = e^{im\varphi} \psi_n(r, \vartheta, 0), \quad (11.112)$$

con m intero.

11.10 Regole di selezione: transizioni di dipolo elettrico

Come abbiamo visto nel paragrafo 10.3 l'operatore momento di dipolo elettrico per un atomo può essere scritto

$$\hat{\mathbf{P}} = -e \sum_{i=1}^Z \hat{\mathbf{r}}_i, \quad \text{con componenti} \quad \hat{P}_x = -e \sum_{i=1}^Z \hat{x}_i, \quad \hat{P}_y = -e \sum_{i=1}^Z \hat{y}_i, \quad \hat{P}_z = -e \sum_{i=1}^Z \hat{z}_i. \quad (11.113)$$

Limitandoci, per semplicità di calcolo, all'atomo di idrogeno, abbiamo

$$\hat{\mathbf{P}} = -e \hat{\mathbf{r}}, \quad \text{con componenti} \quad \hat{P}_x = -e \hat{x}, \quad \hat{P}_y = -e \hat{y}, \quad \hat{P}_z = -e \hat{z}, \quad (11.114)$$

dove $\hat{\mathbf{r}} \equiv (\hat{x}, \hat{y}, \hat{z})$ è l'operatore posizione dell'elettrone nel sistema di riferimento del nucleo. Se vogliamo calcolare il valor medio di $\hat{\mathbf{P}}$ su un qualunque autostato $|n\rangle$ di $\hat{H}_0$ abbiamo

$$\langle n | \hat{\mathbf{P}} | n \rangle = -e \int_{-\infty}^{+\infty} dx \int_{-\infty}^{+\infty} dy \int_{-\infty}^{+\infty} dz \psi_n^*(\mathbf{r}) \mathbf{r} \psi_n(\mathbf{r}). \quad (11.115)$$

Se nell'integrale effettuiamo il cambiamento di variabili

$$x \rightarrow -x, \quad y \rightarrow -y, \quad z \rightarrow -z,$$

il valore dell'integrale stesso rimane invariato, quindi

$$\langle n | \hat{\mathbf{P}} | n \rangle = -e \int_{+\infty}^{-\infty} -dx \int_{+\infty}^{-\infty} -dy \int_{+\infty}^{-\infty} -dz \psi_n^*(-\mathbf{r}) (-\mathbf{r}) \psi_n(-\mathbf{r}). \quad (11.116)$$

L'integrale triplo rimane inalterato anche se adesso per ognuno degli integrali scambiamo i limiti di integrazione e cambiamo il segno del differenziale

$$\langle n | \hat{\mathbf{P}} | n \rangle = -e \int_{-\infty}^{+\infty} dx \int_{-\infty}^{+\infty} dy \int_{-\infty}^{+\infty} dz \psi_n^*(-\mathbf{r}) (-\mathbf{r}) \psi_n(-\mathbf{r}). \quad (11.117)$$

Se supponiamo di non avere degenerazione, continua ad essere valida per la funzione d'onda la proprietà (11.100), quindi $\psi_n^*(\mathbf{r})\psi_n(\mathbf{r})$ è invariante per parità, e $\psi_n^*(-\mathbf{r})\psi_n(-\mathbf{r}) = \psi_n^*(\mathbf{r})\psi_n(\mathbf{r})$:

$$\begin{aligned} \langle n | \hat{\mathbf{P}} | n \rangle &= -e \int_{-\infty}^{+\infty} dx \int_{-\infty}^{+\infty} dy \int_{-\infty}^{+\infty} dz \psi_n^*(-\mathbf{r}) (-\mathbf{r}) \psi_n(-\mathbf{r}) \\ &= -e \int_{-\infty}^{+\infty} dx \int_{-\infty}^{+\infty} dy \int_{-\infty}^{+\infty} dz \psi_n^*(\mathbf{r}) (-\mathbf{r}) \psi_n(\mathbf{r}) \\ &= +e \int_{-\infty}^{+\infty} dx \int_{-\infty}^{+\infty} dy \int_{-\infty}^{+\infty} dz \psi_n^*(\mathbf{r}) (\mathbf{r}) \psi_n(\mathbf{r}), \end{aligned} \quad (11.118)$$

che, paragonato alla (11.115), porta a

$$\langle n | \hat{\mathbf{P}} | n \rangle = -\langle n | \hat{\mathbf{P}} | n \rangle = 0. \quad (11.119)$$

Come avevamo già detto, $\hat{\mathbf{P}}$ ha valor medio nullo su qualunque autostato non degenero di $\hat{H}_0$. Se abbiamo a che fare con un atomo diverso dall'atomo di idrogeno, rimane il fatto che gli autostati non

degeneri dell'hamiltoniana imperturbata hanno parità definita. I conti appena fatti possono quindi essere ripetuti con funzioni d'onda che dipendono dalle coordinate di tutti gli elettroni dell'atomo (opportunamente antisimmetrizzate), e con il momento di dipolo scritto nella forma (11.113), riottenendo il risultato che il valor medio di $\hat{\mathbf{P}}$ è nullo su qualunque stato stazionario dell'atomo. Una cosa molto importante da notare è che, usando proprietà di simmetria, siamo arrivati a dire che il valore dell'integrale è zero senza doverlo calcolare esplicitamente.

Per gli elementi di matrice di $\hat{\mathbf{P}}$ tra due autostati diversi di $\hat{H}_0$, per esempio $|m\rangle$ e $|n\rangle$, il risultato precedente ci dice che

$$\langle n|\hat{\mathbf{P}}|m\rangle = -e \int_{-\infty}^{+\infty} dx \int_{-\infty}^{+\infty} dy \int_{-\infty}^{+\infty} dz \psi_n^*(\mathbf{r}) \mathbf{r} \psi_m(\mathbf{r}). \quad (11.120)$$

può essere (ma non necessariamente è) diverso da zero solo se gli stati $|m\rangle$ e $|n\rangle$ hanno parità diversa. Per andare oltre conviene considerare separatamente le singole componenti del dipolo ed usare le coordinate sferiche, con l'asse z preso come asse di quantizzazione. Etichettiamo esplicitamente gli autostati come $|n, l, m\rangle$, e le funzioni d'onda come $\psi_{n,l,m}(r, \vartheta, \varphi)$. Per la componente $\hat{P}_z$, abbiamo

$$\langle n', l', m'|\hat{P}_z|n, l, m\rangle = -e \int \psi_{n',l',m'}^*(r, \vartheta, \varphi) (r \cos \vartheta) \psi_{n,l,m}(r, \vartheta, \varphi) d\tau, \quad (11.121)$$

dove z è stato scritto come $r \cos \vartheta$, e $d\tau$ è l'elemento di volume. Appliciamo una rotazione di un angolo arbitrario φ_0 attorno all'asse z , che lascia inalterata l'hamiltoniana. Poiché z non viene cambiato dalla rotazione, la (11.112) ci dice che la trasformazione dell'elemento di matrice è

$$\langle n', l', m'|\hat{P}_z|n, l, m\rangle \rightarrow \langle n', l', m'|\hat{P}_z|n, l, m\rangle e^{i(m-m')\varphi_0}. \quad (11.122)$$

D'altra parte l'elemento di matrice deve restare invariato, quindi, se $\langle n, l, m|\hat{P}_z|n', l', m'\rangle \neq 0$, dobbiamo avere

$$e^{i(m-m')\varphi_0} = 1 \quad (11.123)$$

qualunque sia il valore di φ_0 . Questo è possibile solo se $m' = m$. Questa è una *regola di selezione*: se il campo elettrico dell'onda elettromagnetica è diretto lungo l'asse z , possono avvenire transizioni di dipolo elettrico solo tra stati con $m' = m$. In questo caso si parla di radiazione polarizzata π (il valore fonetico di “ π ” è “p”, per polarizzazione “parallela” all'asse di quantizzazione).

Anziché trattare direttamente le componenti $\hat{P}_x$ e $\hat{P}_y$, adesso conviene considerare le loro combinazioni $\hat{P}_+$ e $\hat{P}_-$ definite da

$$\begin{aligned} \hat{P}_+ &= \hat{P}_x + i\hat{P}_y \\ \hat{P}_- &= \hat{P}_x - i\hat{P}_y. \end{aligned} \quad (11.124)$$

Esprimendo x e y in coordinate polari, cioè $x = r \sin \vartheta \cos \varphi$ e $y = r \sin \vartheta \sin \varphi$, abbiamo per gli elementi di matrice

$$\langle n', l', m'|\hat{P}_\pm|n, l, m\rangle = -e \int \psi_{n',l',m'}^*(r, \vartheta, \varphi) r \sin \vartheta e^{\pm i\varphi} \psi_{n,l,m}(r, \vartheta, \varphi) d\tau. \quad (11.125)$$

Se, di nuovo, facciamo una rotazione di un angolo arbitrario φ_0 attorno all'asse z , che lascia inalterata l'hamiltoniana, abbiamo la trasformazione

$$\langle n, l, m|\hat{P}_\pm|n', l', m'\rangle \rightarrow \langle n, l, m|\hat{P}_\pm|n', l', m'\rangle e^{i(m-m'\pm 1)\varphi_0}. \quad (11.126)$$

Imponendo nuovamente che l'elemento di matrice resti invariato, vediamo che deve essere

$$m' = m + 1 \quad \text{perché possa essere} \quad \langle n, l, m | \hat{P}_+ | n', l', m' \rangle \neq 0$$

$$m' = m - 1 \quad \text{perché possa essere} \quad \langle n, l, m | \hat{P}_- | n', l', m' \rangle \neq 0. \quad (11.127)$$

Una seconda regola di selezione ci dice quindi che, per radiazione polarizzata perpendicolarmente all'asse di quantizzazione, si possono avere solo transizioni con $m' = m \pm 1$. In questo caso si parla di luce polarizzata σ (il valore fonetico di σ è “s”, “perpendicolare” in tedesco si dice “senkrecht”, quindi, polarizzazione “perpendicolare” all'asse di quantizzazione). In caso di radiazione σ polarizzata circolarmente, si chiama polarizzazione σ^+ quella che induce transizioni con $m' = m + 1$, polarizzazione σ^- quella che induce transizioni con $m' = m - 1$.

Con considerazioni analoghe, ma più complicate dal punto di vista matematico, possiamo ricavare la regola di selezione per il momento angolare orbitale dell'elettrone:

$$l' = l \pm 1. \quad (11.128)$$

Regole di selezione più generali possono essere ottenute, in maniera più elegante, dal teorema di Wigner-Eckart, per il quale rimandiamo ad un testo di meccanica quantistica.

Capitolo 12

Superconduttività

12.1 Un richiamo di magnetismo classico

Il fenomeno della superconduttività non può essere interpretato semplicemente come il caso limite di un materiale a resistività nulla dell'elettromagnetismo classico. La presenza dell'*effetto Meissner*, che vedremo nel seguito, ci dice che un materiale superconduttore reale è un materiale perfettamente diamagnetico. Per questo ci conviene cominciare il capitolo con due richiami di magnetismo, oltre che uno di termodinamica.

Cominciamo dal problema di una sfera di permeabilità magnetica relativa μ_r e raggio R , inserita in un campo magnetico esterno uniforme $\mathbf{B}_a$, come in Fig. 12.1. Con $\mathbf{B}_a$ indicheremo sempre il campo applicato dallo sperimentatore, per esempio regolando la corrente che passa in un solenoide. A $\mathbf{B}_a$ si aggiungerà, eventualmente, il campo dovuto alla presenza di materiali con $\mu_r \neq 1$. Controlliamo se esiste una soluzione in cui: i) il campo magnetico $\mathbf{B}^{(i)}$ all'interno della sfera è uniforme e proporzionale a $\mathbf{B}_a$, ii) la magnetizzazione della sfera, proporzionale a $\mathbf{B}^{(i)}$, di conseguenza è uniforme, e iii) il campo all'esterno $\mathbf{B}^{(e)}$ è la sovrapposizione di $\mathbf{B}_a$ e di un campo $\mathbf{B}^{(m)}$ generato dalla magnetizzazione della sfera. Sappiamo che $\mathbf{B}^{(m)}$ sarà uguale al campo generato da un dipolo magnetico $\mathbf{m} = \alpha \mathbf{B}_a$ situato al centro O della sfera stessa. Controlliamo quindi se è possibile che sia

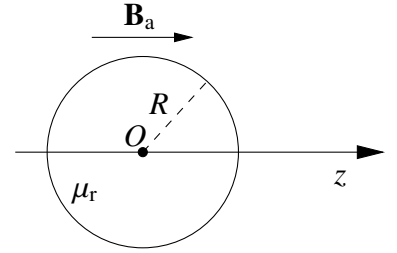


Figura 12.1 Sfera con permeabilità magnetica μ_r in campo magnetico esterno uniforme.

$$\mathbf{B}^{(i)} = k \mathbf{B}_a \quad (12.1)$$

$$\mathbf{B}^{(e)} = \mathbf{B}_a + \mathbf{B}^{(m)}, \quad (12.2)$$

dove α e k sono costanti di proporzionalità da determinare. In coordinate sferiche, con origine in O e z come asse polare, abbiamo per $\mathbf{B}^{(m)}$ e le sue componenti

$$\mathbf{B}^{(m)} = \alpha B_a \frac{\mu_0}{4\pi} \left[\left(3 \frac{\hat{z} \cdot \mathbf{r}}{r^5} \right) \mathbf{r} - \frac{\hat{z}}{r^3} \right], \quad (12.3)$$

$$B_r^{(m)} = \alpha B_a \frac{\mu_0}{4\pi} \frac{2 \cos \vartheta}{r^3}, \quad (12.4)$$

$$B_\vartheta^{(m)} = \alpha B_a \frac{\mu_0}{4\pi} \frac{\sin \vartheta}{r^3}, \quad (12.5)$$

$$B_\varphi^{(m)} = 0, \quad (12.6)$$

dove $\hat{\mathbf{z}}$ è il versore dell'asse z ed r il modulo della distanza da O . I valori di α e k si determinano dalle condizioni di continuità della componente perpendicolare di $\mathbf{B}$ e della componente parallela del campo ausiliare $\mathbf{H}$ in ogni punto della superficie della sfera. Abbiamo (vedi Fig. 12.2)

$$B_{\perp}^{(i)} = B_{\perp}^{(e)} \Rightarrow k B_a \cos \vartheta = B_a \cos \vartheta + B_r^{(m)} = B_a \cos \vartheta + \alpha B_a \frac{\mu_0}{4\pi} \frac{2 \cos \vartheta}{R^3} \quad (12.7)$$

$$H_{\parallel}^{(i)} = H_{\parallel}^{(e)} \Rightarrow k \frac{B_a}{\mu_0 \mu_r} \sin \vartheta = \frac{B_a}{\mu_0} \sin \vartheta - \frac{B_{\vartheta}^{(m)}}{\mu_0} = \frac{B_a}{\mu_0} \sin \vartheta - \alpha B_a \frac{1}{4\pi} \frac{\sin \vartheta}{R^3}, \quad (12.8)$$

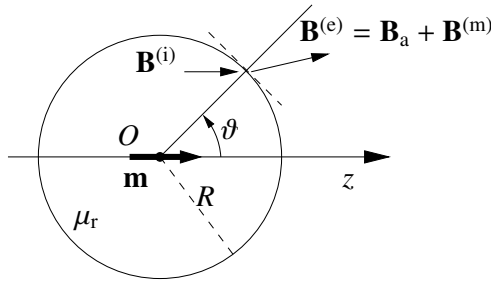


Figura 12.2 Raccordo dei campi $\mathbf{B}$ e $\mathbf{H}$ tra dentro e fuori la sfera.

il segno meno al secondo membro della (12.8) è dovuto all'orientazione della componente ϑ di $\mathbf{B}^{(m)}$. Semplificando, otteniamo il sistema di due equazioni nelle due incognite k ed α

$$k = 1 + \frac{2\alpha\mu_0}{4\pi R^3} \quad (12.9)$$

$$\frac{k}{\mu_r} = 1 - \frac{\alpha\mu_0}{4\pi R^3}, \quad (12.10)$$

che ha soluzioni

$$k = \frac{3\mu_r}{\mu_r + 2} \quad \text{e} \quad \alpha = \frac{4\pi R^3}{\mu_0} \left(\frac{\mu_r - 1}{\mu_r + 2} \right). \quad (12.11)$$

Da qui otteniamo per il campo all'interno della sfera, per il suo momento di dipolo magnetico e per il campo ausiliare $\mathbf{H}$ fuori e dentro alla sfera

$$\mathbf{B}^{(i)} = \frac{3\mu_r}{\mu_r + 2} \mathbf{B}_a, \quad \mathbf{m} = \alpha \mathbf{B}_a = \frac{4\pi R^3}{\mu_0} \left(\frac{\mu_r - 1}{\mu_r + 2} \right) \mathbf{B}_a, \quad \mathbf{H}^{(e)} = \frac{\mathbf{B}^{(e)}}{\mu_0}, \quad \mathbf{H}^{(i)} = \frac{\mathbf{B}^{(i)}}{\mu_0 \mu_r}. \quad (12.12)$$

Poiché sono rispettate la condizione che sia $\mathbf{B} = \mathbf{B}_a$ all'infinito, e le condizioni di raccordo alla superficie della sfera, le (12.12) sono l'unica soluzione del problema. All'interno della sfera la magnetizzazione $\mathbf{M}$ sarà

$$\mathbf{M} = \chi_m \mathbf{H}^{(i)} = (\mu_r - 1) \mathbf{H}^{(i)} = (\mu_r - 1) \frac{\mathbf{B}^{(i)}}{\mu_0 \mu_r} = 3 \left(\frac{\mu_r - 1}{\mu_r + 2} \right) \frac{\mathbf{B}_a}{\mu_0} = \frac{\mathbf{m}}{V}, \quad (12.13)$$

dove $V = 4\pi R^3/3$ è il volume della sfera. Un altro caso che ci interesserà è quello di un cilindro infinito di permeabilità magnetica μ_r immerso in un campo magnetico uniforme $\mathbf{B}_a$, parallelo al suo asse. In questo caso è facile ricavare dalla continuità di $H_{\parallel} = H$ alla superficie laterale del cilindro che il campo interno sarà $\mathbf{B}^{(i)} = \mu_r \mathbf{B}_a$.

Inoltre, è possibile dimostrare che in un qualunque ellissoide di materiale omogeneo e isotropo, con permeabilità magnetica μ_r e posto in un campo applicato $\mathbf{B}_a$ uniforme e parallelo a uno degli assi, il campo magnetico all'interno $\mathbf{B}^{(i)}$ è uniforme e proporzionale a $\mathbf{B}_a$. La dimostrazione richiede però l'uso di coordinate ellissoidali ed esula dagli scopi del corso. Qui ci limitiamo a notare che la sfera e il cilindro che abbiamo appena trattato sono casi particolari dell'ellissoide.

12.2 Lavoro di magnetizzazione

Consideriamo un cilindro non ferromagnetico di raggio R e lunghezza l , e quindi volume $V = \pi R^2 l$, inizialmente in assenza di campo magnetico. Vogliamo calcolare il lavoro fatto sul campione quando lo si magnetizza, accendendo un campo magnetico applicato $\mathbf{B}_a$ parallelo al suo asse. Supponiamo che valga la condizione $l \gg R$, in modo che, all'interno del cilindro $\mathbf{B}$, $\mathbf{H}$ e la magnetizzazione $\mathbf{M}$ possano essere considerati uniformi, con $\mathbf{H}$ continuo alla superficie laterale. Chiameremo $\mathbf{B}^{(i)}$, $\mathbf{H}^{(i)}$ e $\mathbf{M}^{(i)}$ i valori finali all'interno del campione. Il calcolo può essere effettuato in due modi, ed è utile considerarli entrambi,

12.2.1 Lavoro “meccanico” sul campione

Consideriamo un campo magnetico $\mathbf{B}(x)$ diretto lungo l'asse z di un sistema cartesiano. Poniamo il nostro campione cilindrico nel campo, con l'asse parallelo a $\hat{z}$ sulla retta $(x, y = 0)$. Supponiamo anche che $\mathbf{B}(x)$ vari molto lentamente con x , da un valore massimo $\mathbf{B}_a$ per $x = 0$ fino a annullarsi all'infinito. Per variazione lenta intendiamo che il campo possa essere considerato con buona approssimazione uniforme su tutto il volume del nostro cilindro. Per prima cosa calcoliamo il lavoro che il campo fa sul cilindro quando questo trasla di una quantità dx .

Partiamo dalle forze di Lorentz che agiscono lungo x sui lati di una spiretta rettangolare centrata in $(x, 0, 0)$, di lati dx e dy , e percorsa da una corrente I , come in fig. 12.3. Avremo

$$\begin{aligned} f(x) &= -I dy B(x) \\ f(x + dx) &= I dy \left(B(x) + \frac{\partial B}{\partial x} dx \right), \end{aligned}$$

per cui la forza complessiva lungo l'asse x vale

$$F(x) = I dx dy \frac{\partial B}{\partial x} = dm \frac{\partial B}{\partial x},$$

dove dm è il momento magnetico della spiretta, diretto lungo l'asse z . Se il gradiente di $\mathbf{B}$ cambia sufficientemente lentamente, la forza complessiva sul cilindro può essere scritta

$$F_x = m \frac{\partial B}{\partial x},$$

dove m è il momento magnetico del cilindro. Il lavoro fatto dal campo quando il cilindro si sposta di dx vale così

$$dL = F_x dx = m \frac{\partial B}{\partial x} dx = m dB.$$

Questa espressione può essere generalizzata in

$$dL = \mathbf{m} \cdot d\mathbf{B}, \quad (12.14)$$

che corrisponde al lavoro che le forze del campo fanno sul cilindro quando questo passa da una posizione in cui il campo vale $\mathbf{B}$ a una in cui il campo vale $\mathbf{B} + d\mathbf{B}$. Il lavoro fatto dal campo è l'opposto

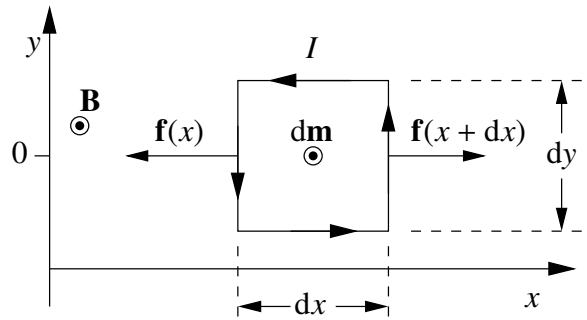


Figura 12.3 Forze di Lorentz sui lati di una spiretta rettangolare di lati dx e dy in presenza di un campo magnetico $\mathbf{B}(x)$ diretto lungo z .

del *lavoro di magnetizzazione*, definito come il lavoro che deve essere fatto *contro le forze del campo* per spostare il campione da dove il campo vale $\mathbf{B}$ a dove vale $\mathbf{B} + d\mathbf{B}$, mentre, conseguentemente, il suo momento di dipolo passa da $\mathbf{m} = V\mathbf{M}$ a $\mathbf{m} + d\mathbf{m} = V(\mathbf{M} + d\mathbf{M})$. In particolare, con le nostre ipotesi, il lavoro di magnetizzazione per portare il cilindro da zero alla magnetizzazione che si ha in presenza del campo applicato B_a è il lavoro necessario per portare il cilindro dall'infinito all'origine delle coordinate:

$$L = - \int_0^{B_a} \mathbf{m} \cdot d\mathbf{B} \quad (12.15)$$

12.2.2 Lavoro fatto dal generatore

Consideriamo un campione cilindrico che occupi interamente l'interno di un solenoide di lunghezza l , con N spire di raggio R . Vogliamo calcolare il lavoro del generatore quando la corrente del solenoide passa da 0 al valore finale I_a , trascurando l'energia dissipata per effetto Joule.

Aumentiamo gradualmente la corrente che circola nelle spire, partendo da $I = 0$ fino a giungere al valore finale I_a . Durante la trasformazione a ogni valore di I corrisponderà un valore $B^{\text{app}}(I) = \mu_0 IN/l$ del campo applicato, e un valore $B^{\text{int}}(I) = B^{\text{app}}(I) + \mu_0 M^{\text{camp}}(I)$ del campo all'interno del campione, dove $M^{\text{camp}}(I)$ è la magnetizzazione del campione. Oltre che compensare l'effetto Joule, il generatore deve fare lavoro perché, quando varia B^{int} , compare una contro-forza elettromotrice $\mathcal{E}$ sul solenoide, e anche questa deve essere compensata. Quando la corrente vale I il lavoro necessario per contrastare $\mathcal{E}$ per un tempo dt vale

$$dL_{\text{tot}} = -\mathcal{E}I dt. \quad (12.16)$$

Esprimiamo $\mathcal{E}$ e I in termini di B^{int} e di B^{app} . Abbiamo

$$\mathcal{E} = -\frac{d\Phi}{dt} = -N\pi R^2 \frac{dB^{\text{int}}(I)}{dt}, \quad \text{e} \quad B^{\text{app}}(I) = \mu_0 \frac{N}{l} I, \quad \text{da cui} \quad I = \frac{l}{N\mu_0} B^{\text{app}}. \quad (12.17)$$

Sostituendo nella (12.16)

$$dL_{\text{tot}} = N\pi R^2 \frac{dB^{\text{int}}}{dt} \frac{l}{N\mu_0} B^{\text{app}} dt = \pi R^2 l \frac{B^{\text{app}}}{\mu_0} dB^{\text{int}}.$$

Ricordando che $\pi R^2 l = V$, volume del campione (e dell'interno del solenoide), otteniamo

$$dL_{\text{tot}} = V \frac{B^{\text{app}}}{\mu_0} dB^{\text{int}}. \quad (12.18)$$

Come abbiamo visto, il campo $\mathbf{B}^{\text{int}}$ può essere scritto nella forma $\mathbf{B}^{\text{int}} = \mathbf{B}^{\text{app}} + \mu_0 \mathbf{M}^{\text{camp}}$, e ricordando che nel solenoide $\mathbf{B}^{\text{int}}$, $\mathbf{B}^{\text{app}}$ e $\mathbf{M}^{\text{camp}}$ sono paralleli, abbiamo

$$dL_{\text{tot}} = V \frac{B^{\text{app}}}{\mu_0} (dB^{\text{app}} + \mu_0 dM^{\text{camp}}) = d\left(V \frac{B^{\text{app}^2}}{2\mu_0}\right) + V \mathbf{B}^{\text{app}} \cdot d\mathbf{M}. \quad (12.19)$$

dove $d\mathbf{M}$ è il differenziale di $\mathbf{M}^{\text{camp}}$. Integrando, e ricordando che $B^{\text{app}}(I_a) = B_a$, otteniamo

$$L_{\text{tot}} = \int_0^{I_a} d\left(V \frac{B^{\text{app}^2}}{2\mu_0}\right) + V \int_0^{I_a} \mathbf{B}^{\text{app}} \cdot d\mathbf{M} = V \frac{B_a^2}{2\mu_0} + \int_0^{I_a} \mathbf{B}^{\text{app}} \cdot d\mathbf{m}, \quad (12.20)$$

Il primo termine all'ultimo membro è l'energia che si avrebbe all'interno del solenoide vuoto con corrente I_a , il secondo termine è il lavoro in più fatto a causa della presenza del cilindro nel solenoide, mentre $\mathbf{m} = V\mathbf{M}^{\text{camp}}$ è il momento magnetico del cilindro stesso.

Alternativamente possiamo supporre di partire ancora in assenza di corrente, e quindi di campi magnetici e di magnetizzazione, ma questa volta con il campione fuori dal solenoide e a distanza infinita, per poi arrivare alla stessa situazione finale, in cui il campione riempie il solenoide, nel quale scorre la corrente I_a . In questo caso il lavoro necessario per passare dalla situazione iniziale alla situazione finale può essere scomposto in tre parti

1. Portare da zero al valore finale I_a la corrente del solenoide, così portando da zero a $\mathbf{B}_a$ il campo magnetico applicato.
2. Portare il campione dentro al solenoide, mantenendo costante la corrente I_a mentre il campione si muove e si magnetizza. Questo processo comporta lavoro per due motivi
 - (a) A corrente I_a fissa nel solenoide, bisogna fare lavoro sul campione per portarlo dall'infinito all'interno del solenoide stesso, muovendolo contro le forze esercitate dal campo del solenoide. Durante questo processo il campione si magnetizza.
 - (b) E' necessario lavoro per mantenere costante la corrente I_a nelle spire del solenoide, contrastando la forza elettromotrice indotta dal campione che si avvicina ed entra.

Il lavoro necessario per il passo (1) è semplicemente

$$L^{(1)} = V \frac{B_a^2}{2\mu_0}, \quad (12.21)$$

che corrisponde all'energia immagazzinata nel solenoide vuoto percorso dalla corrente I_a . Il lavoro necessario per il processo (2a) non è altro che quello dato dall'eq. (12.15)

$$L^{(2a)} = - \int_0^{B_a} \mathbf{m} \cdot d\mathbf{B},$$

dove $\mathbf{B}$ è, istante per istante, il campo alla posizione del campione. Nel processo (2b) il generatore deve compiere lavoro per mantenere costante la corrente I_a nelle spire mentre all'interno del solenoide B^{int} cambia dal valore iniziale B_a , quando il campione è a distanza infinita, al valore finale $B^{(i)}$, quando il campione è completamente inserito nel solenoide. Alla fine avremo

$$\mathbf{B}^{(i)} = \mu_0 (\mathbf{H}^{(i)} + \mathbf{M}^{(i)}) = \mathbf{B}_a + \mu_0 \mathbf{M}^{(i)}.$$

Abbiamo quindi per questo contributo al lavoro

$$\begin{aligned} L^{(2b)} &= - \int_{\text{inizio}}^{\text{fine}} I_a \mathcal{E} dt = I_a \int_{\text{inizio}}^{\text{fine}} \frac{d\Phi}{dt} dt = I_a \pi R^2 N \int_{B_a}^{B_a + \mu_0 M^{(i)}} dB^{\text{int}} \\ &= I_a \frac{N}{l} \pi R^2 \mu_0 M^{(i)} = B_a V M^{(i)} = \mathbf{B}_a \cdot \mathbf{m}. \end{aligned} \quad (12.22)$$

dove $\mathbf{m} = V\mathbf{M}^{(i)}$ è il momento magnetico finale. Il lavoro totale fatto dal generatore sarà quindi

$$L_{\text{tot}} = L^{(1)} + L^{(2a)} + L^{(2b)} = V \frac{B_a^2}{2\mu_0} - \int_0^{B_a} \mathbf{m} \cdot d\mathbf{B} + \mathbf{B}_a \cdot \mathbf{m}. \quad (12.23)$$

Un'integrazione per parti mostra che

$$\int_0^{I_a} \mathbf{B}^{\text{app}} \cdot d\mathbf{m} = \mathbf{B}_a \cdot \mathbf{m} - \int_0^{B_a} \mathbf{m} \cdot d\mathbf{B},$$

per cui la (12.23) coincide con la (12.20).

12.3 Un richiamo di termodinamica

Nel caso che un sistema termodinamico interagisca con un campo magnetico, nel primo principio della termodinamica

$$\delta L = \delta Q - dU \quad (12.24)$$

dobbiamo tenere conto del fatto che il lavoro δL è costituito, oltre che dal lavoro meccanico, dal lavoro magnetico associato a variazioni $d\mathbf{B}$ del campo esterno. La (12.24) può essere riscritta, nel caso di una trasformazione reversibile, in cui $dS = \delta Q/T$,

$$\delta L = T dS - dU. \quad (12.25)$$

Introduciamo l'*energia libera* (o *energia libera di Helmholtz*) F , definita da

$$F = U - TS, \quad \text{con} \quad dF = dU - T dS - S dT. \quad (12.26)$$

F è considerata una funzione di V e T . Per una trasformazione reversibile isoterma, in cui $dT = 0$, abbiamo che

$$\delta L = -dF \quad (12.27)$$

Quindi nei processi isotermi reversibili (come quelli che avvengono in un sistema in contatto con un bagno termico) l'energia libera ha un ruolo analogo a quello dell'energia potenziale per un sistema isolato in meccanica. Infatti la (12.27) ci dice che il lavoro fatto dal sistema è uguale alla sua diminuzione di energia libera. Per processi isotermi non reversibili, il secondo principio della termodinamica

$$dS \geq \frac{\delta Q}{T} \quad (12.28)$$

implica che il lavoro ottenuto è minore della diminuzione di energia libera. L'energia libera, essendo costruita partendo da funzioni di stato, è una funzione di stato.

In una qualunque trasformazione reversibile vale la relazione $dU = \delta Q - \delta L = T dS - \delta L$. In assenza di lavoro di magnetizzazione avremmo $\delta L = p dV$, mentre in sua presenza abbiamo, dalla (12.14), $\delta L = p dV + \mathbf{m} \cdot d\mathbf{B}$. Il differenziale dell'energia libera può così essere riscritto

$$\begin{aligned} dF &= T dS - p dV - \mathbf{m} \cdot d\mathbf{B} - T dS - S dT \\ &= -p dV - \mathbf{m} \cdot d\mathbf{B} - S dT. \end{aligned} \quad (12.29)$$

Essendo F una funzione di stato, la (12.29) è valida anche per una trasformazione non reversibile.

12.4 Aspetti sperimentali della superconduttività

Heike Kammerlingh Onnes (1853 – 1926), dopo aver realizzato la liquefazione dell'elio nel 1908, ha iniziato uno studio sistematico della dipendenza della resistività elettrica dei metalli dalla temperatura. Nel 1911, studiando la resistività di un campione di mercurio, ha notato che scendendo sotto a circa 4 K la resistenza crollava bruscamente a un valore che Onnes non era più in grado di distinguere dallo zero. A questo fenomeno Onnes ha dato il nome di *superconduttività*, e alla temperatura al di sotto della quale il fenomeno si presenta il nome di *temperatura critica* T_c . Un valore più preciso della temperatura critica del mercurio è $T_c = 4.2$ K. I superconduttori convenzionali, o di tipo I (i *superconduttori ad alta temperatura*, o di tipo II, scoperti più recentemente, sono un caso differente) hanno temperature critiche normalmente comprese tra meno di 1 K e circa 20 K.

Se il flusso Φ_B del campo magnetico attraverso una spira conduttrice cambia, per la legge di Faraday-Neumann compare una forza elettromotrice indotta $\mathcal{E}$ che induce nella spira stessa una corrente elettrica che si oppone alla variazione di Φ_B (la legge di Lenz). Se la spira ha resistenza R e induttanza L , dopo un brusco cambiamento di Φ_B la corrente indotta decade nel tempo secondo la legge

$$I(t) = I_0 e^{-(R/L)t}. \quad (12.30)$$

La misura dell'andamento temporale di $I(t)$ permette di determinare valori di R molto più piccoli di quelli misurabili con metodi potenziometrici. Nel 1956 S.C. Collins ha osservato la corrente indotta in una spira superconduttrice scorrere per circa 2 anni e mezzo senza decadimento apprezzabile, fissando così un limite superiore alla resistività del campione pari a $10^{-21} \Omega \text{ cm}$.

A temperature inferiori a T_c la superconduttività può essere distrutta, con ritorno alla conducibilità normale, applicando al campione un campo magnetico esterno superiore a un certo valore critico B_c . Il campo critico B_c dipende approssimativamente dalla temperatura secondo la legge

$$B_c \simeq B_0 \left[1 - \left(\frac{T}{T_c} \right)^2 \right], \quad (12.31)$$

dove B_0 è il campo critico allo zero assoluto. L'andamento è riportato in Fig. 12.4. Spesso è conveniente usare le grandezze adimensionali $t = T/T_c$ e $b = B_c(T)/B_0$, con le quali l'equazione (12.31) diventa

$$b \simeq 1 - t^2. \quad (12.32)$$

In realtà, la dipendenza di b da t è rappresentata con maggiore precisione da uno sviluppo in serie di Taylor in cui il coefficiente di t è nullo e il coefficiente di t^2 differisce da -1 per qualche percento. La superconduttività di un filo o di una lamina viene distrutta quando la corrente che vi circola raggiunge un certo valore critico I_c . Per campioni sufficientemente spessi da rendere trascurabili gli effetti ai bordi, la corrente critica è quella che genera un campo pari a B_c alla superficie del campione (regola di R. H. Silsbee). Tuttavia, campioni sufficientemente piccoli perché gli effetti ai bordi non siano

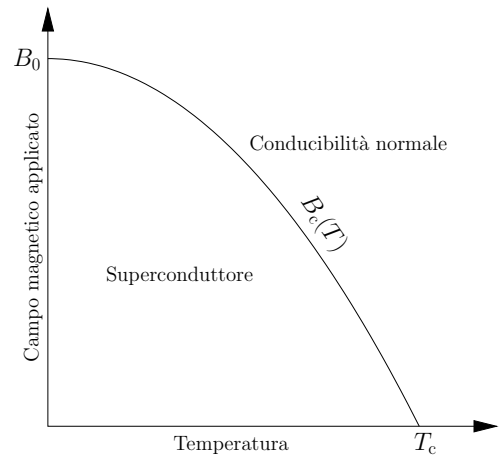


Figura 12.4 Dipendenza del campo critico B_c dalla temperatura in un superconduttore. Al di sotto della curva si ha la superconduttività, al di sopra la conducibilità normale.

trascurabili restano superconduttori anche quando in essi circolano correnti molto più alte di quella determinata dalla regola di Silsbee.

C'è però una differenza molto importante tra un *superconduttore* (materiale osservato sperimentalmente) e un *conduttore perfetto* (materiale ideale per cui viene ipotizzata una resistività nulla). Supponiamo di avere un campione conduttore semplicemente connesso che diventi “conduttore perfetto” (cioè la cui resistività si annulla senza “effetti collaterali”) al di sotto di una certa temperatura critica T_c . Supponiamo inoltre di cominciare un esperimento con il campione a $T < T_c$ e in assenza di campo magnetico. Se a un certo istante accendiamo il campo magnetico vengono indotte delle correnti superficiali che impediscono al flusso magnetico di penetrare nel campione. A resistività nulla, queste correnti superficiali possono durare indefinitamente.

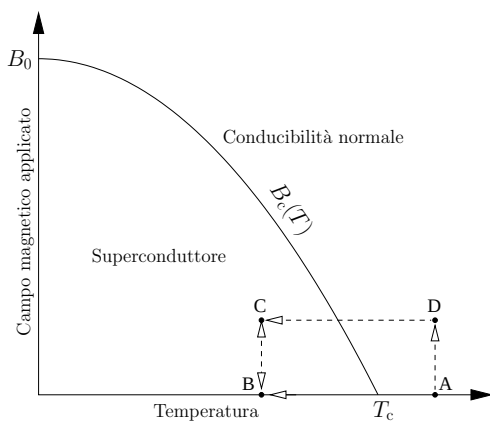


Figura 12.5 Dipendenza dello stato finale dal percorso della transizione per un conduttore perfetto.

in cui $T > T_c$, e quindi si ha conducibilità normale, e non c'è campo magnetico esterno, e uno stesso stato finale B, in cui $T < T_c$, quindi a resistività è nulla. Sia in A che in B il campo magnetico applicato è nullo.

- i) Percorso $A \rightarrow B \rightarrow C \rightarrow B$. In A il campione è “normale”, $\mathbf{B} = 0$ e $\Phi_B = 0$. In B è ancora $\mathbf{B} = 0$, $\Phi_B = 0$, ma la resistività del campione nulla. Nel passaggio da B a C si accende il campo magnetico, ma, poiché nel tratto $B \rightarrow C$ la resistività del campione è sempre nulla, il flusso magnetico Φ_B che lo attraversa rimane nullo. Con il tratto $C \rightarrow B$ si torna allo stato B sempre mantenendo nullo Φ_B .
- ii) Percorso $A \rightarrow D \rightarrow C \rightarrow B$. In A il campione è “normale” e $\Phi_B = 0$. In D il campione è ancora “normale”, ma $\Phi_B \neq 0$. Nel tratto $D \rightarrow C$ il campo magnetico esterno è costante, in D il campione è “normale”, ma quando si attraversa la curva $B = B_c(T)$ la resistenza si annulla e Φ_B viene “congelato”. Nel tratto $C \rightarrow B$ il campo esterno viene spento, ma, essendo il campione sempre conduttore perfetto, Φ_B rimane congelato e si conserva.

Tuttavia l'esperimento di W. Meissner e R. Ochsenfeld del 1933 ha dimostrato che, quando un materiale diventa *superconduttore* il flusso magnetico viene sempre espulso dal suo interno (effetto Meissner). Quindi, se abbiamo a che fare con un *superconduttore* reale anziché con un *conduttore perfetto*, negli stati B e C della Fig. 12.5 il flusso interno Φ_B è sempre nullo, indipendentemente dalle trasformazioni precedenti subite dal campione. Questa è una caratteristica fondamentale di tutti

Adesso supponiamo invece che lo stesso campione sia inizialmente a una temperatura $T > T_c$, quindi un conduttore normale, e immerso in un campo magnetico applicato $B_a \neq 0$. All'interno del campione ci sarà un certo flusso $\Phi_B \neq 0$. Ad un certo istante raffreddiamo il campione sotto T_c : la resistenza si annulla, ma non compaiono correnti indotte e Φ_B non cambia. Successivamente spegniamo il campo applicato. A questo punto compaiono correnti superficiali indotte che mantengono inalterato il flusso interno Φ_B indefinitamente. Questo comporta che lo stato del campione non sia determinato solo dai valori della temperatura e del campo magnetico applicato, ma che dipenda anche dalle trasformazioni subite precedentemente. Consideriamo, per esempio, due percorsi diversi tra uno stesso stato iniziale A, rappresentato in Fig. 12.5,

i superconduttori reali di dimensioni grandi rispetto a una lunghezza caratteristica che definiremo in seguito.

12.5 Campo magnetico critico

Nel paragrafo 12.2 abbiamo visto che, se teniamo un campione a temperatura costante T in un campo magnetico esterno costante $\mathbf{B}_a$, possono avvenire solo trasformazioni con

$$\Delta F \leq 0,$$

cioè, lo stato di equilibrio stabile è quello per cui l'energia libera di Helmholtz F è minima. Il problema dell'equilibrio tra la fase normale e quella superconduttrice del metallo è quindi del tutto analogo a una normale transizione di fase. Per semplicità, in quanto segue supporremo sempre di avere un campione cilindrico di volume V , lungo rispetto al proprio raggio, con l'asse parallelo al campo magnetico. In presenza di un campo magnetico $\mathbf{B}$ alla temperatura T , l'energia libera del campione nella fase superconduttrice si scriverà $F_s(V, T, B)$, mentre quella nella fase normale si scriverà semplicemente $F_n(V, T)$, perché qui abbiamo con ottima approssimazione $\mu_r = 1$ e il campione non risente del campo magnetico.

La condizione di equilibrio tra fase superconduttrice e fase normale a una data temperatura $T < T_c$ si realizza al campo magnetico critico $B_c(T)$, al quale, quindi, si ha

$$F_s(V, T, B_c) = F_n(V, T). \quad (12.33)$$

Invece in condizioni tali che sia $F_n < F_s$ è stabile la fase normale, mentre se $F_s < F_n$ è stabile la fase superconduttrice. Al di sotto della temperatura critica T_c avremo $F_s(V, T, 0) < F_n(V, T)$, visto che, in assenza di campo magnetico, il campione è stabile nella fase superconduttrice.

La condizione $F_s(V, T, B_c) = F_n(V, T)$ vale su ogni punto della curva rappresentata in Fig. 12.4. Se ci muoviamo lungo la curva facendo un piccolo spostamento le due energie libere rimangono uguali, quindi avremo

$$F_s(V, T + dT, B_c + dB_c) = F_n(V, T + dT),$$

da cui, sottraendo la (12.33), otteniamo

$$dF_s = dF_n. \quad (12.34)$$

Dalla (12.29), tenendo conto del fatto che il volume del campione rimane praticamente immutato, abbiamo

$$dF_s = -V\mathbf{M}_c^{(i)} \cdot d\mathbf{B}_c - S_s dT \quad \text{e} \quad dF_n = -S_n dT, \quad (12.35)$$

dove $\mathbf{M}_c^{(i)}$ è la magnetizzazione del campione superconduttore in presenza del campo critico, e con S_s ed S_n abbiamo indicato l'entropia nelle due fasi. La (12.34) diventa

$$-V\mathbf{M}_c^{(i)} \cdot d\mathbf{B}_c - S_s dT = -S_n dT, \quad (12.36)$$

da cui otteniamo

$$S_n - S_s = V\mathbf{M}_c^{(i)} \cdot \frac{d\mathbf{B}_c}{dT}. \quad (12.37)$$

Nel caso di un campione a forma di cilindro allungato con l'asse parallelo al campo, per la (12.45) abbiamo

$$S_n - S_s = -V \frac{B_c}{\mu_0} \frac{dB_c}{dT}, \quad (12.38)$$

analoga all'equazione di Clausius-Clapeyron. Poiché $dB_c/dT < 0$, deve essere $S_n > S_s$: lo stato superconduttore ha un'entropia più bassa, e corrisponde a un maggiore ordine. Associato alla differenza di entropia si ha un calore latente

$$Q_L = T(S_n - S_s) \quad (12.39)$$

che deve essere assorbito nella transizione dallo stato superconduttore a quello normale. Si noti che $S_n - S_s = 0$ quando la transizione avviene a campo magnetico nullo (cioè per $T = T_c$), quindi in questo caso non c'è calore latente associato alla transizione di fase. Per $T \rightarrow 0$ il terzo principio della termodinamica ci dice che anche $(S_n - S_s) \rightarrow 0$, quindi il calore latente deve tendere a zero anche per $T \rightarrow 0$. Inoltre, poiché B_c non tende a zero per $T \rightarrow 0$, l'equazione (12.38) ci dice che deve tendere a zero la derivata dB_c/dT . Tutto questo è in accordo con le osservazioni sperimentali, come mostrato nel grafico della Fig. 12.4.

Se adesso poniamo il campione a campo zero a una temperatura $T < T_c$, e accendiamo lentamente il campo magnetico esterno mantenendo la temperatura costante, per la (12.29), con considerazioni analoghe a quelle usate per la (12.35), abbiamo

$$dF_s = -V\mathbf{M}^{(i)} \cdot d\mathbf{B}_a = V \frac{B_a}{\mu_0} dB_a, \quad \text{e} \quad dF_n = 0.$$

Quando raggiungiamo il campo critico B_c abbiamo

$$F_s(V, T, B_c) = F_s(T, V, 0) + \frac{V}{\mu_0} \int_0^{B_c} B_a dB_a = F_s(V, T, 0) + V \frac{B_c^2}{2\mu_0},$$

d'altra parte, al campo critico abbiamo che $F_s(V, T, B_c) = F_n(V, T)$, quindi, in assenza di campo, la differenza tra l'energia libera dello stato normale e quella dello stato superconduttore vale

$$\Delta F = F_n(V, T) - F_s(V, T, 0) = V \frac{B_c^2}{2\mu_0}. \quad (12.40)$$

Sempre dalla Fig. 12.4 e dalla legge approssimata (12.31) vediamo che il valore massimo del campo critico si ha a $T=0$. Questo valore massimo di B_c è dell'ordine di qualche 10^{-2} Tesla, che corrisponde a una differenza di energia libera dell'ordine di 10^{-8} eV ($\simeq 10^{-27}$ J) per atomo del campione. Come termine di paragone, si ricordi che l'energia di Fermi per gli elettroni di conduzione di un normale metallo è tra i 10 e i 20 eV ($\simeq 10^{-18}$ J).

12.6 Diamagnetismo perfetto

L'espulsione del campo magnetico dall'interno di un superconduttore è dovuta alla comparsa di correnti elettriche sulla sua superficie, distribuite in modo tale da generare un campo magnetico che cancelli il campo applicato ovunque all'interno del superconduttore stesso. Una descrizione formale dei campi $\mathbf{B}$ e $\mathbf{H}$, e della magnetizzazione $\mathbf{M}$, all'interno di un superconduttore macroscopico in presenza di un campo applicato $\mathbf{B}_a$ è quindi

- 1.a) interno: $\mathbf{B}^{(i)} = \mathbf{H}^{(i)} = \mathbf{M}^{(i)} = 0$;
- 1.b) superficie: $\mathbf{K}_c \neq 0$ e $\mathbf{K}_m = 0$, dove $\mathbf{K}_c$ e $\mathbf{K}_m$ sono rispettivamente le densità superficiali di corrente di *conduzione* e di *magnetizzazione*;
- 1.c) esterno: $\mathbf{B}^{(e)} = \mathbf{B}_a + \mathbf{B}_s$, dove $\mathbf{B}_a$ è il campo applicato e $\mathbf{B}_s$ il campo generato da $\mathbf{K}_c$.

Sia all'interno che all'esterno del superconduttore il campo generato dalle correnti superficiali $\mathbf{B}_s$ si somma al campo applicato $\mathbf{B}_a$, dando un campo complessivo nullo all'interno e tangente alla superficie del superconduttore immediatamente all'esterno, come mostrato in Fig. 12.6 nel caso di un superconduttore sferico. La descrizione del punto 1) è formalmente corretta, ma spesso è considerato più conveniente considerare il superconduttore in presenza di un campo esterno come un materiale in cui si abbia

2.a) interno: $\mathbf{B}^{(i)} = 0$, $\mathbf{H}^{(i)} \neq 0$, $\mathbf{M}^{(i)} \neq 0$;

2.b) superficie: $\mathbf{K}_c = 0$ e $\mathbf{K}_m \neq 0$;

2.c) esterno: $\mathbf{B}^{(e)} = \mathbf{B}_a + \mathbf{B}_s$, dove adesso $\mathbf{B}_s$ è considerato generato da $\mathbf{K}_m$.

I campi $\mathbf{B}_a$, $\mathbf{B}_s$, $\mathbf{B}^{(i)}$ e $\mathbf{B}^{(e)}$, che possono essere misurati, sono ovviamente gli stessi ai punti 1) e 2). Le interpretazioni 1) e 2) attribuiscono invece valori diversi ad $\mathbf{H}$, cosa che non meraviglia troppo perché è solo un vettore ausiliario, ma anche a $\mathbf{M}$, $\mathbf{K}_c$ e $\mathbf{K}_m$, che siamo abituati a considerare quantità con un significato fisico ben preciso, e un'ambiguità sul loro valore può sorprendere. Ma nel paragrafo 12.7 vedremo che in un superconduttore possono esistere solo correnti superficiali e non ha senso distinguere tra corrente di conduzione e di magnetizzazione. Ha senso fisico solo la densità complessiva di corrente superficiale $\mathbf{K}$, e la somma $\mathbf{K} = \mathbf{K}_c + \mathbf{K}_m$ è la stessa ai punti 1) e 2). Per quel che riguarda la magnetizzazione, sappiamo dall'elettromagnetismo classico che il campo magnetico generato da $\mathbf{M}$ è uguale al campo generato dalla densità volumica di corrente di magnetizzazione $\mathbf{J}_m = \nabla \times \mathbf{M}$, più il campo generato dalla densità superficiale di corrente di magnetizzazione $\mathbf{K}_m = \mathbf{M} \times \hat{\mathbf{n}}$, dove $\hat{\mathbf{n}}$ è il versore normale uscente dalla superficie che delimita il materiale magnetizzato. Nel nostro caso le due interpretazioni danno entrambi una magnetizzazione uniforme, quindi è comunque $\mathbf{J}_m = \nabla \times \mathbf{M} \equiv 0$. Resta il campo generato dalla densità di corrente superficiale, che è lo stesso se consideriamo $\mathbf{K}$ come corrente sia di magnetizzazione che di conduzione.

Seguendo l'interpretazione del punto 2) e ricordando le relazioni

$$\mathbf{B} = \mu_0(\mathbf{H} + \mathbf{M}) \quad \text{e} \quad \mathbf{M} = \chi_m \mathbf{H},$$

otteniamo dalla condizione $\mathbf{B}^{(i)} = 0$

$$\mathbf{M}^{(i)} = -\mathbf{H}^{(i)},$$

da cui segue che, per un superconduttore, si ha

$$\begin{aligned} \chi_m &= -1, \\ \mu_r &= \chi_m + 1 = 0. \end{aligned} \tag{12.41}$$

Il superconduttore si comporta quindi come un materiale perfettamente diamagnetico. D'altra parte la relazione $\mu_r = 0$, unita alla condizione $\mathbf{B}^{(i)} = 0$, implica che $\mathbf{H}^{(i)} = \mathbf{B}^{(i)} / \mu_0 \mu_r$ non abbia senso fisico, e di conseguenza non lo abbia $\mathbf{M}$. Le interpretazioni 1) e 2) non sono quindi contraddittorie. Se

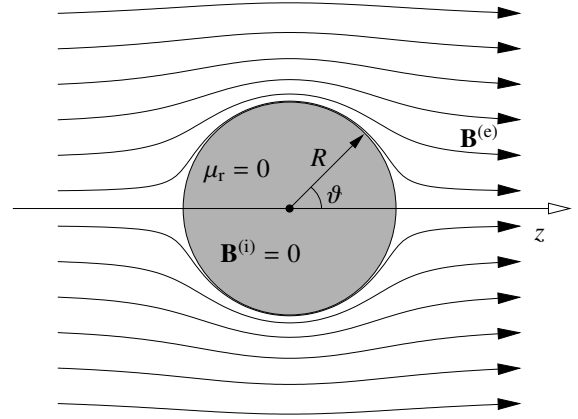


Figura 12.6 Sfera perfettamente diamagnetica. Essendo il campo interno $\mathbf{B}^{(i)} = 0$, le linee di forza del campo esterno $\mathbf{B}^{(e)}$ sono tangenti alla superficie.

seguiamo l'interpretazione 2) e consideriamo una sfera superconduttrice di raggio R immersa in un campo applicato $\mathbf{B}_a$ (Fig. 12.6), e inseriamo le (12.41) nella (12.13) otteniamo

$$\mathbf{H}^{(i)} = -\mathbf{M}^{(i)} = \frac{3}{2} \frac{\mathbf{B}_a}{\mu_0}. \quad (12.42)$$

Dalla condizione di continuità della componente parallela di $\mathbf{H}$ all'equatore della sfera ($\vartheta = \pi/2$ in Fig. 12.6) otteniamo

$$\mathbf{H}_{\text{eq}}^{(e)} = \mathbf{H}_{\text{eq}}^{(i)} = \frac{3}{2} \frac{\mathbf{B}_a}{\mu_0}, \quad \text{da cui} \quad \mathbf{B}_{\text{eq}}^{(e)} = \mu_0 \mathbf{H}_{\text{eq}}^{(e)} = \frac{3}{2} \mathbf{B}_a, \quad (12.43)$$

mentre dalla condizione di continuità della componente perpendicolare di $\mathbf{B}$ ai poli ($\vartheta = 0$ e $\vartheta = \pi$ in Fig. 12.6) della sfera otteniamo

$$\mathbf{B}_{\text{pol}}^{(e)} = \mathbf{B}_{\text{pol}}^{(i)} = 0. \quad (12.44)$$

Il modulo del campo esterno $B^{(e)}$ è quindi massimo all'equatore, e ha due minimi ai poli.

Se invece di una sfera consideriamo un cilindro superconduttore, con l'asse parallelo a $\mathbf{B}_a$ e sufficientemente lungo da poter essere considerato infinito, la sua magnetizzazione non genererà campo all'esterno del cilindro stesso, quindi avremo $\mathbf{B}^{(e)} = \mathbf{B}_a$. La condizione di continuità della componente parallela di $\mathbf{H}$ alla superficie laterale del cilindro impone che sia

$$\mathbf{H}^{(i)} = \mathbf{H}^{(e)} = \frac{\mathbf{B}_a}{\mu_0}, \quad \text{da cui} \quad \mathbf{M}^{(i)} = -\frac{\mathbf{B}_a}{\mu_0}. \quad (12.45)$$

A temperature $T < T_c$ l'intero cilindro infinito rimane superconduttore fino a che il campo applicato non supera il valore critico B_c . Appena il modulo di $\mathbf{B}_a$ supera il valore B_c l'intero cilindro passa alla conducibilità normale.

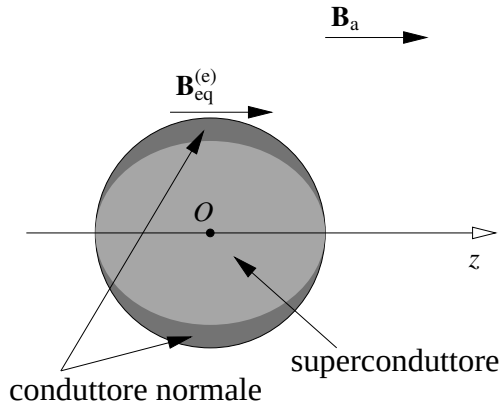


Figura 12.7 Interpretazione non possibile del caso $B_{\text{eq}}^{(e)} > B_c > B_a$.

La situazione è un po' più complicata per una sfera superconduttrice, data la non uniformità di $\mathbf{B}^{(e)}$ alla sua superficie. Il problema si pone quando abbiamo $B_{\text{eq}}^{(e)} \geq B_c > B_a$, cioè quando il campo applicato è ancora inferiore al campo critico B_c , ma il campo esterno all'equatore della sfera raggiunge il valore critico. Ipotizzare che una porzione della sfera in prossimità dell'equatore passi alla conducibilità normale (parte più scura in Fig. 12.7) porta a una contraddizione. Infatti il confine tra la parte a conduzione normale e la parte superconduttrice si ha per $B^{(e)} = B_c$, ma adesso, nella parte diventata normale, avremmo $B^{(e)} < B_c$. Non è possibile immaginare una divisione macroscopica del campione in una parte normale e una superconduttrice che permetta una distribuzione del campo con $B \geq B_c$ nella prima, $B < B_c$ nella seconda,

e $B = B_c$ sulle superfici di separazione.

L'unica spiegazione possibile è postulare, come fatto separatamente da R. Peierls e da F. London nel 1936, che, una volta raggiunta la condizione $B_a \geq \frac{2}{3} B_c$, l'intera sfera si divida in una moltitudine di regioni mesoscopiche (cioè grandi rispetto alle dimensioni atomiche, ma ancora piccole rispetto alle dimensioni macroscopiche) alternativamente normali e superconduttrici, con $B = B_c$ nelle prime

e $B = 0$ nelle seconde. La distribuzione di queste regioni è tale che la magnetizzazione macroscopica per unità di volume cambia linearmente da

$$M^{(i)} = -\frac{B_c}{\mu_0} = -\frac{3}{2} \frac{B_a}{\mu_0} \quad \text{per} \quad B_a = \frac{2}{3} B_c$$

a

$$M^{(i)} = 0 \quad \text{per} \quad B_a = B_c.$$

Quindi, per $\frac{2}{3}B_c \leq B_a \leq B_c$, avremo all'interno della sfera superconduttrice

$$\begin{aligned} M^{(i)} &= -\frac{3}{\mu_0} (B_c - B_a) \\ H^{(i)} &= \frac{B_c}{\mu_0} \\ B^{(i)} &= B_c - \mu_0 M = 3B_a - 2B_c. \end{aligned} \quad (12.46)$$

La figura 12.8 mostra la dipendenza della magnetizzazione interna per una sfera e per un cilindro superconduttori dal campo magnetico applicato B_a . Notate che, secondo la (12.15), il lavoro di magnetizzazione per unità di volume per portare il campione da campo nullo al campo critico è dato, indipendentemente dalla forma del campione, da

$$\frac{L}{V} = - \int_0^{B_c} \mathbf{M}(B_a) \cdot d\mathbf{B}_a = \frac{1}{2} \frac{B_c^2}{\mu_0},$$

pari alla densità di energia magnetica per unità di volume al campo critico.

Nell'intervallo di campo $\frac{2}{3}B_c \leq B_a \leq B_c$, in cui il campione sferico non è né interamente superconduttore né interamente normale, si dice che il campione si trova nello *stato intermedio*. Qualunque sia la forma del campione, lo stato intermedio cessa a $B_a = B_c$. Il valore del campo applicato B_{im} , a cui lo stato intermedio inizia, dipende invece dalla particolare forma del campione. Per un campione a forma di ellissoide generico, con uno degli assi parallelo a $\mathbf{B}_a$, le (12.46) si riscrivono

$$\begin{aligned} M^{(i)} &= -\frac{1}{\mu_0 D} (B_c - B_a) \\ H^{(i)} &= \frac{B_c}{\mu_0} \\ B^{(i)} &= B_c - \frac{1}{D} (B_c - B_a), \end{aligned} \quad (12.47)$$

dove D prende il nome di *coefficiente di demagnetizzazione* relativo all'asse dell'ellissoide parallelo a $\mathbf{B}_a$. Per una sfera si ha $D = \frac{1}{3}$, mentre nei due casi limite di cilindro infinito con l'asse parallelo a $\mathbf{B}_a$ e cilindro infinito con l'asse perpendicolare a $\mathbf{B}_a$ si ha, rispettivamente, $D = 0$ e $D = \frac{1}{2}$. Il campo a cui comincia lo stato intermedio vale $B_{im} = (1 - D) B_c$. Abbiamo quindi $B_{im} = \frac{2}{3} B_c$ per una sfera, $B_{im} = \frac{1}{2} B_c$ per un cilindro infinito trasversale, e $B_{im} = B_c$ per un cilindro infinito longitudinale. Per quest'ultimo, quindi, non esiste stato intermedio.

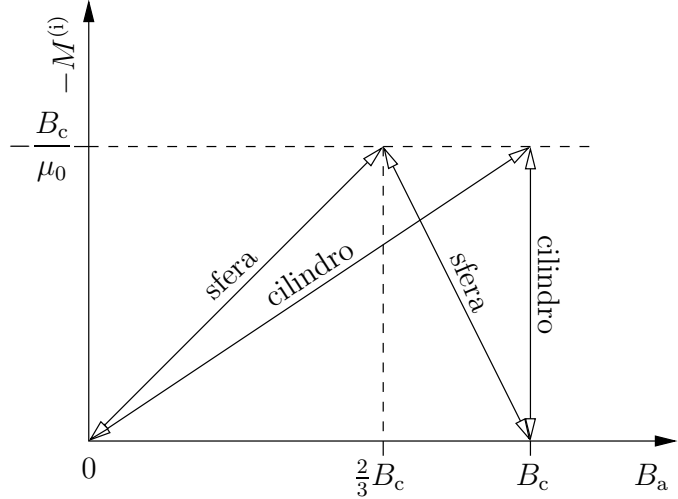


Figura 12.8 Dipendenza della magnetizzazione dal campo applicato per un superconduttore sferico e uno cilindrico con l'asse parallelo al campo.

12.7 Corrente nei superconduttori, trattazione classica

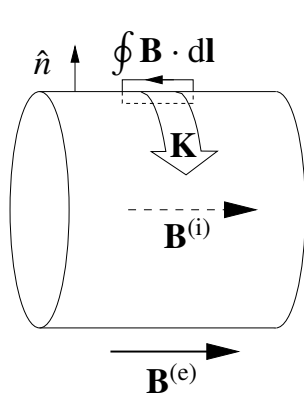
Passiamo adesso alla distribuzione di corrente in un superconduttore partendo dalla quarta equazione di Maxwell, che, in assenza di campi elettrici variabili nel tempo, si riduce a

$$\nabla \times \mathbf{B} = \mu_0 \mathbf{J}, \quad (12.48)$$

dove

$$\mathbf{J} = \mathbf{J}_c + \nabla \times \mathbf{M}$$

è la densità di corrente macroscopica complessiva, somma delle densità volumiche di corrente di conduzione (corrente libera) $\mathbf{J}_c$ e di magnetizzazione $\mathbf{J}_m = \nabla \times \mathbf{M}$. Nell'elettromagnetismo classico la magnetizzazione è originata dalle correnti elettriche microscopiche legate ai moti orbitali degli elettroni e dagli spin degli elettroni e dei nuclei, qui tralascieremo questi ultimi. In un materiale non magnetizzato le correnti microscopiche, sommandosi, si cancellano dando origine a correnti macroscopiche nulle. Invece, come abbiamo già ricordato nel paragrafo 12.6, in un materiale magnetizzato le correnti microscopiche si combinano dando origine a correnti macroscopiche di magnetizzazione, con densità volumica $\mathbf{J}_m = \nabla \times \mathbf{M}$ e densità superficiale $\mathbf{K}_m = \mathbf{M} \times \hat{\mathbf{n}}$. Queste correnti, legate a moti orbitali, non incontrano resistenza nel mezzo. In caso di magnetizzazione uniforme $\nabla \times \mathbf{M} = 0$, e abbiamo solo la corrente di magnetizzazione superficiale. Per prima cosa dimostriamo che in un conduttore normale le correnti di magnetizzazione sono le uniche correnti superficiali che possono esistere. Poiché a una densità di corrente superficiale non nulla $\mathbf{K}$ corrisponde una densità di corrente volumica *infinita*, $\mathbf{K} \neq 0$ implica una discontinuità della componente tangenziale di $\mathbf{B}$ tra l'esterno e l'interno del conduttore. Dalla circuitazione di $\mathbf{B}$ sul perimetro di un rettangolo con due lati infinitesimi perpendicolari alla superficie del conduttore, e che la attraversano, come indicato in Fig. 12.9, otteniamo



$$\mathbf{K} = \frac{1}{\mu_0} \hat{\mathbf{n}} \times (\mathbf{B}^{(e)} - \mathbf{B}^{(i)}), \quad (12.49)$$

dove $\hat{\mathbf{n}}$ è il versore normale alla superficie del conduttore, diretto verso l'esterno. Se all'esterno c'è il vuoto e il mezzo è omogeneo con permeabilità magnetica relativa μ_r , la continuità della componente parallela di $\mathbf{H} = \mathbf{B}/\mu_0\mu_r$ ci dice che $\hat{\mathbf{n}} \times \mathbf{B}^{(e)} = (\hat{\mathbf{n}} \times \mathbf{B}^{(i)})/\mu_r$, da cui ricaviamo

$$\mathbf{K} = \frac{1}{\mu_0} \hat{\mathbf{n}} \times \left(\frac{\mathbf{B}^{(i)}}{\mu_r} - \mathbf{B}^{(i)} \right) = \hat{\mathbf{n}} \times \mathbf{B}^{(i)} \left(\frac{1 - \mu_r}{\mu_0 \mu_r} \right), \quad (12.50)$$

in funzione del solo campo interno $\mathbf{B}^{(i)}$. Notare che la (12.50) può essere riscritta

$$\begin{aligned} \mathbf{K} &= \hat{\mathbf{n}} \times \mathbf{B}^{(i)} \left(\frac{1 - \mu_r}{\mu_0 \mu_r} \right) = -\hat{\mathbf{n}} \times \mathbf{B}^{(i)} \left(\frac{\chi_m}{\mu_0 \mu_r} \right) \\ &= -\hat{\mathbf{n}} \times (\chi_m \mathbf{H}^{(i)}) = \mathbf{M}^{(i)} \times \hat{\mathbf{n}}, \end{aligned} \quad (12.51)$$

da cui vediamo che la corrente superficiale coincide con la corrente di magnetizzazione. D'altra parte, in presenza di resistività, una corrente superficiale di conduzione incontrerebbe una resistenza infinita perché circolerebbe in uno strato di spessore nullo.

Figura 12.9 Corrente superficiale su un conduttore normale (corrente di magnetizzazione).

Consideriamo ancora un conduttore normale non semplicemente connesso, per esempio un toro, come in Fig. 12.10. Come mostrato in figura, sulla sua superficie possiamo disegnare linee chiuse sia di tipo l_1 , possibili anche sulla superficie di un corpo semplicemente connesso, che di tipo l_2 , che non possono essere ridotte con continuità a un punto senza attraversare la superficie. La carica totale per unità di tempo I che attraversa una di queste linee è data dal flusso di $\mathbf{K}$ attraverso la linea stessa

$$I = \oint \mathbf{K} \cdot (\hat{\mathbf{n}} \times d\mathbf{l}), \quad (12.52)$$

dove $\hat{\mathbf{n}}$ è il versore normale alla superficie, quindi $\hat{\mathbf{n}} \times d\mathbf{l}$ giace sulla superficie stessa ed è perpendicolare al percorso di integrazione. L'integrale è esteso alla linea chiusa, e nel caso di una linea del tipo l_1 , $\hat{\mathbf{n}} \times d\mathbf{l}$ sarà uscente per un percorso antiorario ed entrante per un percorso orario. Dalla (12.50) abbiamo, qualunque sia il tipo di curva,

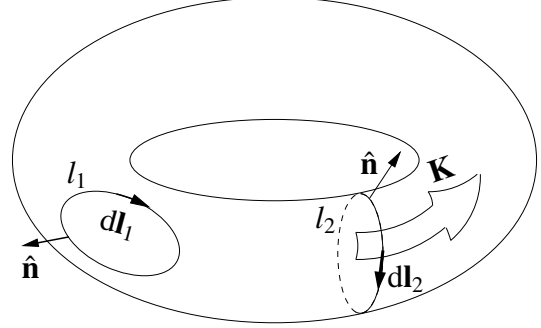


Figura 12.10 Corrente superficiale su un conduttore a forma di toro.

$$\begin{aligned} I &= \oint \mathbf{K} \cdot (\hat{\mathbf{n}} \times d\mathbf{l}) \\ &= \oint \left(\frac{1 - \mu_r}{\mu_0 \mu_r} \right) (\hat{\mathbf{n}} \times \mathbf{B}^{(i)}) \cdot (\hat{\mathbf{n}} \times d\mathbf{l}) \\ &= \left(\frac{1 - \mu_r}{\mu_0 \mu_r} \right) \oint \mathbf{B}^{(i)} \cdot d\mathbf{l} = 0, \end{aligned} \quad (12.53)$$

data l'irrotazionalità di $\mathbf{B}$ in assenza di campi elettrici variabili nel tempo. Quindi dalla superficie delimitata da una linea chiusa del tipo l_1 non entra né esce carica, e la linea l_2 non può essere attraversata da corrente netta.

Invece nel caso di un superconduttore abbiamo visto che è $\mathbf{B}^{(i)} \equiv 0$, quindi da una parte la (12.49) diventa semplicemente

$$\mathbf{K} = \hat{\mathbf{n}} \times \frac{\mathbf{B}^{(e)}}{\mu_0}, \quad (12.54)$$

mentre, dall'altra, la (12.50) diventa indeterminata, perché sono nulli sia $\mathbf{B}^{(i)}$ che μ_r . Cade quindi la (12.53), e non esiste più un vincolo a priori per la carica superficiale per unità di tempo che può attraversare le linee chiuse della Fig. 12.10. Inoltre, da $\mathbf{B}^{(i)} \equiv 0$ segue che $\nabla \times \mathbf{B}^{(i)} = 0$, quindi, per la (12.48), anche $\mathbf{J} = 0$. La densità di corrente è quindi nulla ovunque all'interno di un superconduttore, e in esso possono esistere solo correnti superficiali. Ma, essendo la resistività nulla, la corrente superficiale può essere di conduzione.

In queste condizioni, come anticipato nel paragrafo 12.6, non ha più senso distinguere la corrente in corrente di conduzione e corrente di magnetizzazione. Questo porta alla perdita di senso fisico delle definizioni sia di $\mathbf{M}$, legata alla sola corrente di magnetizzazione, che di $\mathbf{H}$, il cui rotore è uguale alla densità di corrente di sola conduzione. Ma è ancora possibile, come abbiamo già visto, continuare a usare $\mathbf{M}$ e $\mathbf{H}$ come vettori ausiliari, affetti sì da ambiguità, ambiguità che però si risolvono usando coerentemente solo una delle possibili interpretazioni.

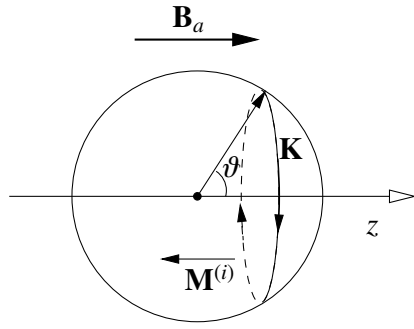


Figura 12.11 Corrente superficiale su una sfera superconduttrice.

circolano correnti. In questo dominio “esterno” esiste un potenziale magnetico scalare φ tale che il campo $\mathbf{B}$ è dato dalla relazione $\mathbf{B} = -\nabla\varphi$.

Per le condizioni al contorno sulla continuità della componente normale di $\mathbf{B}$, alla superficie di un superconduttore abbiamo

$$B_{\perp} = -\frac{\partial\varphi}{\partial n} = 0,$$

dove $\hat{n}$ è la normale alla superficie. Ma dalla teoria del potenziale sappiamo che se $\partial\varphi/\partial n$ è nullo sia alla superficie di un corpo semplicemente connesso che all'infinito, φ è costante su tutto lo spazio fuori dal corpo. Quindi $\mathbf{B}$ deve essere nullo, come le correnti superficiali che lo generano. D'altra parte, queste correnti dovrebbero scorrere in modo da non generare campo all'interno del superconduttore.

Invece, in presenza di un campo applicato esterno $\mathbf{B}_a$, avremo delle correnti superficiali tali da annullare il campo stesso all'interno del superconduttore. In questo caso $\mathbf{K}$ può essere calcolata a partire dal vettore ausiliario $\mathbf{M}^{(i)}$ tramite la relazione

$$\mathbf{K} = \mathbf{M}^{(i)} \times \hat{n}.$$

Nei casi semplici considerati, cioè sfera superconduttrice e cilindro infinito superconduttore con l'asse parallelo a $\mathbf{B}_a$, per le correnti superficiali abbiamo dalle (12.42) e (12.45)

$$\text{sfera: } |\mathbf{K}| = \frac{3}{2} \frac{B_a}{\mu_0} \sin\vartheta, \quad \text{cilindro: } |\mathbf{K}| = \frac{B_a}{\mu_0}.$$

Il caso della sfera è rappresentato in Fig. 12.11. Le correnti scorrono in modo da cancellare $\mathbf{B}_a$ all'interno dei superconduttori.

Le cose cambiano per un superconduttore non semplicemente connesso. Adesso, infatti, anche lo spazio esterno al superconduttore non è più semplicemente connesso, il potenziale scalare φ diventa polidromo e non è più possibile ottenere

Figura 12.12 Toro in stato di conducibilità normale ($T > T_c$), in presenza di campo magnetico applicato $\mathbf{B}_a$.

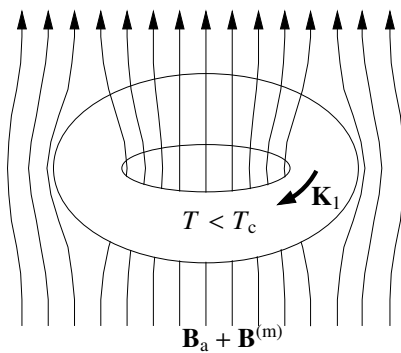


Figura 12.13 La temperatura scende sotto al valore critico e il toro diventa superconduttore. Compare una corrente superficiale $\mathbf{K}_1$ che genera un campo $\mathbf{B}^{(m)}$ che annulla $\mathbf{B}_a$ all'interno del materiale (effetto Meissner).

$\mathbf{B}$ da un unico potenziale scalare. Non è quindi più valida la dimostrazione appena fatta.

Consideriamo, come controesempio, un superconduttore a forma di toro, come nelle figure 12.12 - 12.14. Inizialmente teniamo il toro a una temperatura $T > T_c$, cioè nel suo stato di conducibilità normale, e accendiamo un campo magnetico esterno uniforme $\mathbf{B}_a$, parallelo all'asse del toro stesso. Poiché a $T > T_c$ il materiale ha $\mu_r = 1$, la presenza del toro non perturba il campo magnetico, e il campo complessivo è uguale ovunque al solo $\mathbf{B}_a$, come in Fig. 12.12. Partendo da questa situazione, raffreddiamo il toro fino a una temperatura inferiore a T_c . Il materiale diventa così superconduttore e, per effetto Meissner, il campo magnetico viene espulso. Questo è dovuto alla comparsa di una corrente superficiale $\mathbf{K}_1$ che genera un suo campo magnetico $\mathbf{B}^{(m)}$. Il campo $\mathbf{B}^{(m)}$, sommandosi a $\mathbf{B}_a$, dà un campo nullo all'interno del materiale, e un campo complessivo $\mathbf{B}^{(e)}$ fuori da esso, come in Fig. 12.13. Finalmente, sempre mantenendo la temperatura sotto al valore critico, spegniamo il campo applicato $\mathbf{B}_a$. Poiché il materiale rimane superconduttore, il flusso attraverso la cavità del toro non può cambiare. Sulla sua superficie rimane una distribuzione di corrente superficiale $\mathbf{K}_2$ che mantiene il flusso costante, come abbiamo già visto all'inizio di questo capitolo. Abbiamo così sul toro una corrente superficiale anche in assenza di campi esterni.

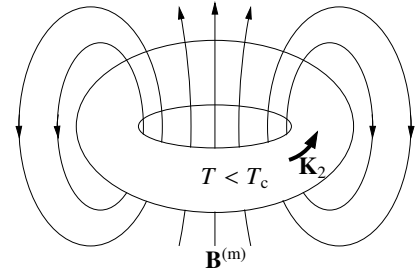


Figura 12.14 Viene spento il campo magnetico applicato $\mathbf{B}_a$. Rimane una corrente superficiale $\mathbf{K}_2$ che mantiene costante il flusso di $\mathbf{B}$ attraverso il toro.

12.8 Corrente nei superconduttori, trattazione quantistica

12.8.1 Equazione di continuità per la funzione d'onda

L'equazione di Schrödinger dipendente dal tempo per una particella di carica q e massa m , in presenza di un campo elettromagnetico con potenziale elettrico V e potenziale magnetico vettore $\mathbf{A}$, si scrive (vedi Appendice E)

$$i\hbar \frac{\partial \psi}{\partial t} = \hat{H}\psi = \frac{1}{2m} (-i\hbar \nabla - q\mathbf{A})^2 \psi + qV\psi. \quad (12.55)$$

Supponiamo di avere una particella in uno stato descritto dalla funzione d'onda $\psi(\mathbf{r}, t)$. La probabilità di trovarla all'istante t dentro a un volume infinitesimo d^3x centrato attorno a un punto di vettore posizione $\mathbf{r}$ vale $|\psi(\mathbf{r}, t)|^2 d^3x$. Supponiamo poi che la $\psi(\mathbf{r}, t)$ vari nel tempo. La condizione di normalizzazione $\int |\psi(\mathbf{r}, t)|^2 d^3x = 1$ impone che a una diminuzione di $|\psi|^2$ in una certa parte dello spazio corrisponda un aumento da qualche altra parte. Ci aspettiamo così che compaia un qualche tipo di "flusso di probabilità" nella zona intermedia. Vogliamo vedere se è possibile definire una *corrente di probabilità* che fluisca da dove $|\psi|^2$ diminuisce a dove $|\psi|^2$ aumenta. In altre parole, cerchiamo per il modulo quadro della funzione d'onda un'equazione analoga all'equazione di continuità $\nabla \cdot \mathbf{J} = -\partial \rho / \partial t$, che lega la densità volumica di carica elettrica ρ alla densità di corrente $\mathbf{J}$. Scritta la densità di probabilità $P(\mathbf{r}, t)$ come

$$P(\mathbf{r}, t) = \psi(\mathbf{r}, t) \psi^*(\mathbf{r}, t), \quad (12.56)$$

cerchiamo una *densità di corrente di probabilità* $\mathbf{W}(\mathbf{r}, t)$ tale che valga la relazione

$$\frac{\partial P}{\partial t} = -\nabla \cdot \mathbf{W}. \quad (12.57)$$

Cominciamo derivando la (12.56) rispetto al tempo, ottenendo

$$\frac{\partial P}{\partial t} = \psi^* \frac{\partial \psi}{\partial t} + \psi \frac{\partial \psi^*}{\partial t}. \quad (12.58)$$

Le derivate di ψ e ψ^* sono date dalla (12.55) e dalla sua complessa coniugata. Sostituite nella (12.58) ci danno

$$\begin{aligned} \frac{\partial P}{\partial t} &= -\frac{i}{\hbar} \left[\psi^* \frac{1}{2m} (-i\hbar \nabla - q\mathbf{A})^2 \psi + qV\psi^*\psi - \psi \frac{1}{2m} (i\hbar \nabla - q\mathbf{A})^2 \psi^* - qV\psi\psi^* \right] \\ &= -\frac{i}{2m\hbar} \left(-\hbar^2 \psi^* \Delta \psi + i\hbar q |\psi|^2 \nabla \cdot \mathbf{A} + i\hbar q \psi^* \mathbf{A} \cdot \nabla \psi + q^2 A^2 |\psi|^2 \right. \\ &\quad \left. + \hbar^2 \psi \Delta \psi^* + i\hbar q |\psi|^2 \nabla \cdot \mathbf{A} + i\hbar q \psi \mathbf{A} \cdot \nabla \psi^* - q^2 A^2 |\psi|^2 \right) \\ &= \frac{1}{2m} [\psi^* (i\hbar \Delta \psi + q\mathbf{A} \cdot \nabla \psi + q\psi \nabla \cdot \mathbf{A}) \\ &\quad + \psi (-i\hbar \Delta \psi^* + q\mathbf{A} \cdot \nabla \psi^* + q\psi^* \nabla \cdot \mathbf{A})] \\ &= \frac{1}{2m} \nabla \cdot [\psi^* (i\hbar \nabla + q\mathbf{A}) \psi + \psi (-i\hbar \nabla + q\mathbf{A}) \psi^*]. \end{aligned} \quad (12.59)$$

Trattandosi della somma di un'espressione con la propria complessa coniugata, complessivamente abbiamo a che fare con un'espressione reale. Abbiamo quindi per la densità di corrente di probabilità

$$\mathbf{W} = \frac{1}{2} \left[\psi^* \left(\frac{-i\hbar \nabla - q\mathbf{A}}{m} \right) \psi + \psi \left(\frac{i\hbar \nabla - q\mathbf{A}}{m} \right) \psi^* \right], \quad (12.60)$$

e la (12.57) può essere riscritta

$$\frac{\partial P}{\partial t} = -\nabla \cdot \left[\frac{1}{2m} \psi^* (-i\hbar \nabla - q\mathbf{A}) \psi + \frac{1}{2m} \psi (i\hbar \nabla - q\mathbf{A}) \psi^* \right]. \quad (12.61)$$

12.8.2 Densità di corrente in un superconduttore

Come abbiamo ricordato più sopra, se facciamo una misura per vedere se la particella è all'interno di un volumetto τ centrato attorno al punto di vettore posizione $\mathbf{r}$, la probabilità di trovarcela all'istante t è $|\psi(\mathbf{r}, t)|^2 \tau$. Quello che stiamo considerando è un volume τ sufficientemente piccolo perché in esso la ψ sia praticamente costante, ma, soddisfatta questa condizione, non lo facciamo tendere a zero. Ben inteso, una volta effettuata la misura, o abbiamo trovato la particella in τ , o non ce l'abbiamo trovata. Supponiamo però di avere non una, ma un "grande" numero N di particelle, tutte nello stesso stato quantistico, e quindi tutte rappresentate dalla stessa funzione d'onda $\psi(\mathbf{r}, t)$. Per "grande" intendiamo che valga la condizione $N|\psi(\mathbf{r}, t)|^2 \tau \gg 1$. In queste condizioni la legge dei grandi numeri ci dice che il numero di particelle che troveremo in τ all'istante t è, con ottima approssimazione, $N|\psi(\mathbf{r}, t)|^2 \tau$. Il prodotto $N|\psi(\mathbf{r}, t)|^2$, legato a una grandezza "microscopica" quantistica come la funzione d'onda ψ , acquista così un significato macroscopico, e corrisponde alla densità volumica delle particelle.

A questo punto va fatta una precisazione. Nei superconduttori, come nei normali metalli, i portatori di carica sono elettroni che, in quanto fermioni, non possono trovarsi tutti nello stesso stato. Vedremo però nel paragrafo 12.9 che l'interazione con il reticolo cristallino dà origine a una piccola interazione attrattiva tra gli elettroni. Questa attrazione porta alla formazione di coppie legate di elettroni, dette coppie di Cooper, che possono essere trattate come *quasiparticelle*. Queste quasiparticelle

hanno spin totale intero, e si comportano quindi come bosoni. Niente impedisce che più coppie di Cooper si trovino nello stesso stato. Quindi la (12.55) va interpretata come l'equazione di Schrödinger non per un elettrone, ma per una coppia di Cooper. La carica q di una di queste quasiparticelle vale $-2e$, mentre m , la sua *massa efficace*, non è esattamente $2m_e$ (con m_e massa dell'elettrone), perché vanno considerati effetti dovuti alle interazioni col reticolo. La quantità N sarà il numero di quasiparticelle, cioè la metà del numero degli elettroni accoppiati.

Il numero di coppie di Cooper per unità di volume sarà così $N|\psi(\mathbf{r}, t)|^2$, e, poiché ogni coppia ha carica q , la densità di carica $\rho_p(\mathbf{r})$ nel punto $\mathbf{r}$, relativa ai soli portatori (trascurando cioè il contributo degli ioni del reticolo), sarà

$$\rho_p(\mathbf{r}, t) = Nq|\psi(\mathbf{r}, t)|^2. \quad (12.62)$$

Per la (12.60) la densità di corrente $\mathbf{J}(\mathbf{r}, t)$ si scriverà

$$\mathbf{J}(\mathbf{r}, t) = Nq \mathbf{W} = \frac{1}{2} Nq \left[\psi^* \left(\frac{-i\hbar \nabla - q\mathbf{A}}{m} \right) \psi + \psi \left(\frac{i\hbar \nabla - q\mathbf{A}}{m} \right) \psi^* \right]. \quad (12.63)$$

Le grandezze macroscopiche classiche ρ_p e $\mathbf{J}$ sono così legate alla funzione d'onda ψ .

In prossimità dello zero assoluto possiamo supporre che ogni coppia di Cooper si trovi nello stato fondamentale dell'hamiltoniana di quasiparticella singola, con energia E_0 . Questo stato è rappresentato dalla funzione d'onda $\psi_0(\mathbf{r}, t)$. La densità di carica dei portatori sarà così $\rho_p(\mathbf{r}, t) = Nq |\psi_0(\mathbf{r}, t)|^2$. Nel nostro superconduttore, globalmente neutro, ρ_p è sovrapposta alla densità di carica degli ioni atomici del reticolo ρ_{ret} , che possiamo supporre uniforme. Quindi anche ρ_p deve essere uniforme, in modo che la densità di carica complessiva $\rho_{\text{tot}} = \rho_p + \rho_{\text{ret}}$ sia uniformemente nulla su tutto il campione. In caso contrario comparirebbero campi elettrici che farebbero muovere i portatori di carica, e questi si redistribuirebbero fino a rendere ρ_{tot} uniformemente nulla. Questo impone che la $|\psi_0(\mathbf{r}, t)|^2$ sia costante su tutto il volume del superconduttore V , e la condizione di normalizzazione

$$\int_V |\psi_0(\mathbf{r}, t)|^2 d^3x = 1 \quad \text{impone che sia} \quad |\psi_0(\mathbf{r}, t)|^2 \equiv \frac{1}{V}. \quad (12.64)$$

La funzione d'onda $\psi_0(\mathbf{r}, t)$ deve quindi avere la forma

$$\psi_0(\mathbf{r}, t) = \frac{1}{\sqrt{V}} e^{i\vartheta(\mathbf{r})} e^{-iE_0 t/\hbar}, \quad (12.65)$$

dove $\vartheta(\mathbf{r})$ è una funzione reale delle coordinate, che, per il momento, non conosciamo.

Se sostituiamo la ψ_0 della (12.65) al posto della ψ nella (12.63) otteniamo, per la densità di corrente elettrica dovuta alle quasiparticelle nello stato fondamentale

$$\begin{aligned} \mathbf{J}(\mathbf{r}) &= \frac{Nq}{2V} \left[e^{-i\vartheta} \left(\frac{-i\hbar \nabla - q\mathbf{A}}{m} \right) e^{i\vartheta} + e^{i\vartheta} \left(\frac{i\hbar \nabla - q\mathbf{A}}{m} \right) e^{-i\vartheta} \right] \\ &= \frac{nq}{2} \left(-\frac{i\hbar}{m} e^{-i\vartheta} \nabla e^{i\vartheta} - \frac{q}{m} \mathbf{A} + \frac{i\hbar}{m} e^{i\vartheta} \nabla e^{-i\vartheta} - \frac{q}{m} \mathbf{A} \right) \\ &= \frac{nq}{m} (\hbar \nabla \vartheta - q\mathbf{A}), \end{aligned} \quad (12.66)$$

dove abbiamo introdotto la densità volumica delle coppie di Cooper $n = N/V$, e abbiamo sfruttato la relazione $\nabla e^{f(\mathbf{r})} = e^{f(\mathbf{r})} \nabla f(\mathbf{r})$, valida per qualunque funzione continua e derivabile $f(\mathbf{r})$ delle coordinate spaziali.

Questa espressione per la densità di corrente nei superconduttori ci permette di interpretare l'effetto Meissner. Calcolando il rotore di entrambi i membri della (12.66) otteniamo la seconda equazione di London (fratelli Fritz e Heinz London, 1935)

$$\nabla \times \mathbf{J} = -\frac{nq^2}{m} \nabla \times \mathbf{A} = -\frac{nq^2}{m} \mathbf{B}, \quad (12.67)$$

perché $\nabla \times \nabla \vartheta = 0$. Scrivendo la $\mathbf{J}$ al primo membro in funzione della quarta equazione di Maxwell nel caso statico abbiamo

$$\frac{1}{\mu_0} \nabla \times \nabla \times \mathbf{B} = -\frac{nq^2}{m} \mathbf{B}. \quad (12.68)$$

Dall'identità $\nabla \times \nabla \times \mathbf{f} = -\Delta \mathbf{f} + \nabla(\nabla \cdot \mathbf{f})$, con $\mathbf{f}(\mathbf{r})$ funzione vettoriale qualunque, otteniamo che $\nabla \times \nabla \times \mathbf{B} = -\Delta \mathbf{B}$, essendo nulla la divergenza di $\mathbf{B}$. Abbiamo così

$$\Delta \mathbf{B} = \frac{\mu_0 n q^2}{m} \mathbf{B} = \frac{1}{\lambda^2} \mathbf{B}, \quad (12.69)$$

con

$$\lambda = \frac{1}{q} \sqrt{\frac{m}{\mu_0 n}}. \quad (12.70)$$

Proviamo a stimare l'ordine di grandezza di λ , che ha le dimensioni di una lunghezza. La costante μ_0 vale $4\pi \times 10^{-7} \text{ N/A}^2$. Per un metallo come il piombo abbiamo circa $3 \times 10^{26} \text{ atomi/m}^3$, supponendo che ognuno contribuisca con un elettrone alla conduzione abbiamo $n \simeq 1.5 \times 10^{26} \text{ coppie di Cooper per m}^3$. La carica dell'elettrone vale $e = -1.602 \times 10^{-19} \text{ C}$, e la sua massa vale $m_e = 9.109 \times 10^{-31} \text{ kg}$. Poniamo quindi $q = 2e$ e, approssimativamente, $m = 2m_e$. Alla fine otteniamo

$$\lambda \simeq 3 \times 10^{-7} \text{ m}. \quad (12.71)$$

Torniamo alla (12.69). Nel caso unidimensionale, con il piano $x = 0$ che separa il vuoto ($x < 0$) dal materiale superconduttore ($x \geq 0$), l'equazione diventa

$$\frac{\partial^2 \mathbf{B}}{\partial x^2} = \frac{1}{\lambda^2} \mathbf{B}, \quad \text{con soluzione del tipo } \mathbf{B} = \mathbf{B}_0 e^{\pm x/\lambda}. \quad (12.72)$$

Questo significa che il campo magnetico deve diminuire esponenzialmente dalla superficie verso l'interno del materiale (un aumento esponenziale porterebbe a un'esplosione!). Dato l'ordine di grandezza di λ ottenuto nella (12.71), $\mathbf{B}$ penetrerà nel materiale solo per uno strato molto sottile. All'interno del materiale, a profondità molto maggiori di λ , avremo $\mathbf{B} = 0$. Questa è la spiegazione dell'effetto Meissner. Come abbiamo già ricordato, $\mathbf{B}$ uniformemente nullo all'interno del materiale implica che vi sia nullo anche $\nabla \times \mathbf{B}$, e quindi anche $\mathbf{J}$, che penetrerà nel materiale solo per un sottile strato superficiale dell'ordine di λ . Quanto detto finora vale per superconduttori sia semplicemente connessi che non. Per un superconduttore semplicemente connesso non c'è niente da aggiungere: un campo applicato $\mathbf{B}_a$ presente prima che il materiale venga raffreddato sotto la temperatura critica viene espulso, un campo applicato quando il materiale è già superconduttore non vi entra.

12.8.3 Quantizzazione del flusso attraverso un anello superconduttore

Le cose cambiano per un superconduttore non semplicemente connesso, come il toro della Fig. 12.15. Supponiamo che il raggio della sezione del toro sia $a \gg \lambda$. Effettuiamo la sequenza di processi schematizzata nelle figure 12.12-12.14. Quanto detto nel paragrafo precedente resta valido, il campo inizialmente applicato $\mathbf{B}_a$ viene espulso dal materiale, penetrando solo in uno strato di profondità dell'ordine di λ , nel quale sono localizzate anche le correnti superficiali.

Il campo magnetico è però presente nel “buco della cianbella”. Quando togliamo $\mathbf{B}_a$ (Fig. 12.14) ci troviamo con una distribuzione superficiale di corrente $\mathbf{K}$ che “congela” il flusso del campo magnetico intrappolato nel “buco”. Siamo quindi in assenza di campo applicato, ma con $\mathbf{A} \neq 0$. All'interno del toro, a distanze molto maggiori di λ dalla superficie, la densità di corrente $\mathbf{J}$ è nulla, e, per la (12.66), avremo

$$\hbar \nabla \vartheta = q \mathbf{A}. \quad (12.73)$$

Adesso calcoliamo la circuitazione di $\mathbf{A}$ lungo la circonferenza l , di raggio r , giacente sul piano di simmetria perpendicolare all'asse del toro, e coassiale al toro stesso. Il raggio r è tale che la circonferenza sia a distanza molto maggiore di λ dalla superficie, come in Fig. 12.15. Per la (12.73) avremo

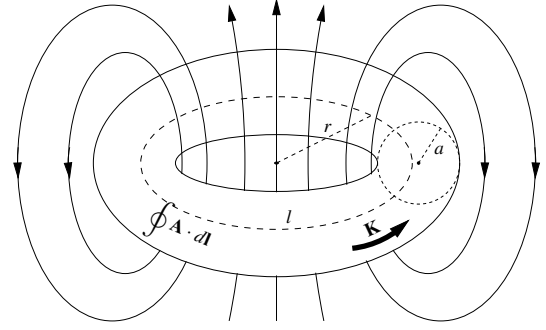


Figura 12.15 Circuitazione di $\mathbf{A}$ all'interno del toro superconduttore.

$$\hbar \oint \nabla \vartheta \cdot d\mathbf{l} = q \oint \mathbf{A} \cdot d\mathbf{l} = q \int \mathbf{B} \cdot d\mathbf{S} = q \Phi, \quad (12.74)$$

dove Φ è il flusso di $\mathbf{B}$ attraverso il “buco” del toro. Va notato che $\mathbf{B}$ non penetra nel superconduttore, quindi la circuitazione di $\mathbf{A}$ non deve variare se cambio il raggio r della circonferenza, purché la circonferenza rimanga all'interno del toro, sempre a distanza molto maggiore di λ dalla superficie. Con la nostra geometria avremo quindi, per le componenti in coordinate cilindriche di $\mathbf{A}$ lungo l , $A_\varphi = \Phi/(2\pi r)$, $A_z = A_r = 0$. Il rotore di $\mathbf{A}$ è quindi nullo, in accordo con il fatto che $\mathbf{B}$ non penetra per più di λ nel superconduttore. Tornando alla (12.74), abbiamo

$$\oint \nabla \vartheta \cdot d\mathbf{l} = \frac{q}{\hbar} \Phi. \quad (12.75)$$

L'integrale di linea del gradiente di ϑ tra due punti qualunque, chiamiamoli A e B , vale

$$\int_A^B \nabla \vartheta \cdot d\mathbf{l} = \vartheta(B) - \vartheta(A), \quad (12.76)$$

quindi potremmo aspettarci che l'integrale tenda a zero se facciamo tenere A a B . Ma mentre questo è senz'altro vero per un percorso di integrazione chiuso all'interno di un superconduttore semplicemente connesso, non è necessariamente vero all'interno del nostro toro. L'unica condizione che possiamo imporre è che la funzione d'onda (12.65) sia monodroma, cioè abbiamo uno e un sol valore in ogni punto, mentre non è necessariamente monodroma ϑ . Se percorriamo la circonferenza l tornando al punto di partenza dobbiamo ritrovare lo stesso valore per la (12.65), ma questo si realizza anche se ϑ cambia di $2\pi n$, con n intero qualunque. Qualunque sia il comportamento di ϑ lungo la circonferenza

(o lungo una qualunque linea chiusa all'interno del toro, a profondità maggiore di λ), una volta fatto un giro completo dobbiamo avere

$$\oint \nabla \vartheta \cdot d\mathbf{l} = 2\pi n, \quad (12.77)$$

che, inserita nella (12.75), ci dà

$$\Phi = n \frac{2\pi\hbar}{q}. \quad (12.78)$$

Il flusso intrappolato deve quindi essere un multiplo intero di $2\pi\hbar/q$. Questo risultato era stato previsto da London nel 1950. London aveva però posto q uguale alla carica elettronica e , ottenendo un *quanto di flusso* (a volte detto *flussone*) pari a $\Phi_0 = 2\pi\hbar/e \simeq 4 \times 10^{-15} \text{ T}\cdot\text{m}^2$. Nel 1961 Deaver e Fairbank a Stanford, e contemporaneamente Doll e Näbauer a Herrsching, sono riusciti a osservare sperimentalmente che il flusso intrappolato variava effettivamente a scatti, ma a scatti pari alla metà di quanto previsto da London, cioè

$$\Phi_0 = \frac{\pi\hbar}{e}, \quad (12.79)$$

in accordo con la teoria secondo cui la superconduttività è dovuta alle coppie di Cooper.

12.9 Coppie di Cooper

Un elettrone di conduzione in un metallo si comporta con buona approssimazione come una particella libera. Infatti, se è vero che c'è una interazione repulsiva tra gli elettroni, è anche vero che c'è una interazione attrattiva tra gli elettroni e gli ioni positivi del reticolo, nel quale gli elettroni di conduzione sono immersi. Gli ioni del reticolo che circondano un elettrone di conduzione lo schermano così dall'interazione con gli altri elettroni di conduzione. Questa considerazione è alla base sia del modello di Drude, in cui gli elettroni di conduzione in un metallo vengono trattati come un gas perfetto classico, che della trattazione quantistica degli elettroni di conduzione come un gas perfetto di Fermi.

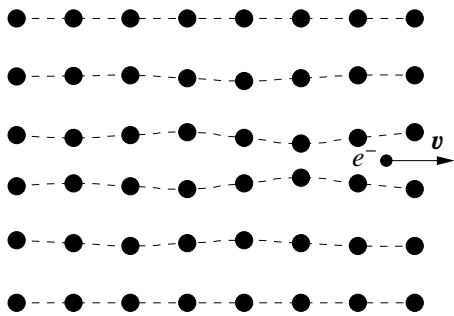


Figura 12.16 Deformazione del reticolo degli ioni positivi da parte di un elettrone.

Lo stato fondamentale degli elettroni trattati come un gas perfetto (quindi senza interazioni tra gli elettroni stessi) di Fermi a 0 K corrisponde ad avere tutti i livelli di particella singola occupati fino all'energia di Fermi E_F . Se per caso compare una sia pur debolissima forza attrattiva tra una coppia di elettroni, questi possono formare uno stato legato. A sua volta la formazione dello stato legato, con energia negativa, porta a uno stato complessivo di energia inferiore di quella dello stato fondamentale degli elettroni liberi. Vediamo, anche se in modo molto qualitativo, come può nascere un'interazione attrattiva tra gli elettroni di un metallo. La teoria di Cooper è basata sull'ipotesi che un elettrone in moto,

attraendo gli ioni positivi del reticolo, lasci dietro di sé una traccia consistente in una distorsione del reticolo stesso, come schematizzato in Fig. 12.16. Al momento del passaggio dell'elettrone gli ioni del reticolo ricevono un impulso che, dato il rapporto tra le masse di ioni e elettrone, porta a uno spostamento apprezzabile, e quindi ad una polarizzazione del reticolo, solo dopo che l'elettrone è passato. La deformazione si propaga dando origine a un fonone. La concentrazione di ioni positivi della deformazione genera una forza attrattiva su di un secondo elettrone. Questo modello può essere

descritto anche in modo quantistico, interpretando le deformazioni del reticolo come sovrapposizioni di fononi, che l'elettrone assorbe e riemette continuamente attraverso la sua interazione con il reticolo. Per il principio di indeterminazione energia-tempo i fononi virtuali possono vivere solo durante un tempo $\tau = 2\pi/\omega$, dove ω è la frequenza angolare del fonone. Partiamo da un gas di Fermi di elettroni non interagenti alla temperatura di 0 K. Tutti gli stati di particella singola con vettore d'onda $\mathbf{k}$ corrispondente a un'energia minore dell'energia di Fermi, cioè con $|\mathbf{k}| < k_F$, dove $k_F = \sqrt{2mE_F/\hbar}$, sono occupati, mentre tutti gli stati con $|\mathbf{k}| > k_F$ sono vuoti. Adesso aggiungiamo al sistema due elettroni di vettori d'onda rispettivamente $\mathbf{k}_1$ e $\mathbf{k}_2$. Facciamo le ipotesi che le energie corrispondenti $E(\mathbf{k}_1)$ e $E(\mathbf{k}_2)$ siano ambedue superiori a E_F , ma all'interno di un range $E_F < E < E_F + \hbar\omega_D$, dove ω_D è la frequenza massima per i fononi nel reticolo, detta frequenza di Debye. Supponiamo che questi due elettroni siano accoppiati dall'interazione attrattiva descritta sopra. I due elettroni aggiuntivi interagiscono scambiando tra loro un fonone di vettore d'onda $\mathbf{q}$, passando, rispettivamente, ai vettori d'onda

$$\mathbf{k}'_1 = \mathbf{k}_1 - \mathbf{q} \quad \text{e} \quad \mathbf{k}'_2 = \mathbf{k}_2 + \mathbf{q}.$$

Per la conservazione della quantità di moto totale dovrà essere

$$\mathbf{k}_1 + \mathbf{k}_2 = \mathbf{k}'_1 + \mathbf{k}'_2 = \mathbf{K}.$$

D'altra parte $\mathbf{k}_1$ e $\mathbf{k}_2$, così come $\mathbf{k}'_1$ e $\mathbf{k}'_2$, dovranno trovarsi, nello spazio dei vettori d'onda, all'interno di un guscio sferico di raggio interno pari a k_F e raggio esterno pari a $k_F + \delta k$, con $E(k_F + \delta k) = E_F + \hbar\omega_D$. Nella Fig. 12.17, che va intesa come simmetrica per rotazioni attorno a $\mathbf{K} = \mathbf{k}_1 + \mathbf{k}_2$, abbiamo tracciato due gusci sferici, ognuno di raggio interno k_F e raggio esterno $k_F + \delta k$, centrati rispettivamente attorno alla "coda" di $\mathbf{k}_1$ e alla punta di $\mathbf{k}_2$. Dalla figura vediamo che, una volta fissata la somma $\mathbf{K}$ dei vettori d'onda della coppia di elettroni, tutte le possibili coppie di vettori d'onda $(\mathbf{k}_1, \mathbf{k}_2)$ devono avere la punta di $\mathbf{k}_1$, coincidente con la "coda" di $\mathbf{k}_2$, all'interno del volume evidenziato in grigio. Questo volume è direttamente legato al numero di processi di scambio di fonone che abbassano l'energia complessiva del gas di elettroni. Pertanto il numero di interazioni attrattive tra coppie di elettroni diventa massimo quando diventa massimo questo volume, cioè per $\mathbf{K} = 0$, quando i due gusci sferici si sovrappongono esattamente. Dobbiamo quindi avere

$$\mathbf{k}_1 = -\mathbf{k}_2. \quad (12.80)$$

Nel seguito ci occuperemo quindi di coppie di elettroni con vettore d'onda opposto (coppie di Cooper). E' possibile dimostrare che i termini di scambio riducono il valore assoluto dell'energia di interazione per coppie con spin parallelo, per cui possiamo limitare la nostra attenzione alle coppie di elettroni con vettore d'onda opposto e spin antiparallelo.

L'ipotesi fondamentale della teoria BCS è quindi che *nello stato fondamentale del gas di elettroni in un superconduttore a 0 K tutti gli stati elettronici in un sottile guscio sferico nello spazio reciproco dei vettori d'onda siano occupati con la maggior densità possibile da coppie di elettroni di quantità di moto opposta e spin opposto*. L'energia di questo stato è più bassa dell'energia del metallo nello stato normale per la quantità (12.40)

$$\Delta E = V \frac{B_c^2}{2\mu_0}. \quad (12.81)$$

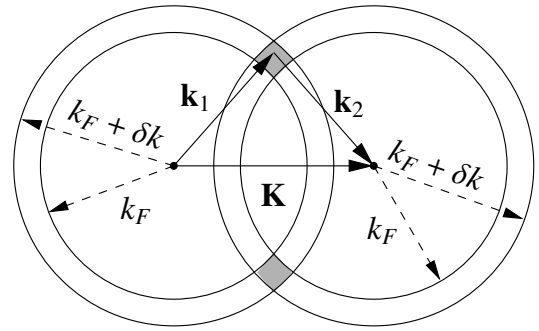


Figura 12.17 Visualizzazione dei vettori d'onda delle coppie di elettroni interagenti nello spazio reciproco dei numeri d'onda.

L'equazione di Schrödinger per la funzione d'onda $\psi(\mathbf{r}_1, \mathbf{r}_2)$ della coppia di Cooper può essere scritta

$$\left[-\frac{\hbar^2}{2m}(\Delta_1 + \Delta_2) + V(\mathbf{r}_1, \mathbf{r}_2) \right] \psi(\mathbf{r}_1, \mathbf{r}_2) = (E + 2E_F) \psi(\mathbf{r}_1, \mathbf{r}_2), \quad (12.82)$$

dove $V(\mathbf{r}_1, \mathbf{r}_2)$ è il potenziale di interazione per scambio di fononi e $2E_F$ è l'energia della coppia di elettroni in assenza di interazione ($V(\mathbf{r}_1, \mathbf{r}_2) = 0$).

Capitolo 13

Raffreddamento magnetico

13.1 Introduzione

Come abbiamo visto nel paragrafo 12.2, è possibile effettuare lavoro su di un corpo cambiando il valore del campo magnetico ad esso applicato. Questo fatto può essere sfruttato per raffreddare particolari campioni paramagnetici fino a temperature molto basse, al di sotto del mK.

Il procedimento presenta analogie con il raffreddamento “meccanico” di un gas. Nel raffreddamento meccanico si ha una prima fase in cui si effettua lavoro sul gas comprimendolo, mentre è in contatto con un bagno termico. L’apporto di energia dovuto al lavoro fatto sul gas viene così ceduto al bagno termico, ed il gas rimane a temperatura costante. In una seconda fase viene tolto il contatto tra il gas e il bagno termico, lasciando il gas termicamente isolato. A questo punto il gas fa un’espansione adiabatica quasi statica (con buona approssimazione reversibile), nella quale compie lavoro a spese della propria energia interna, raffreddandosi. Analogamente, nel raffreddamento magnetico, mentre il campione si trova in contatto con un bagno termico (normalmente elio liquido) gli viene applicato un campo magnetico che induce una magnetizzazione, compiendo così il lavoro dato dalla (12.15). Il campione non si riscalda perché il calore viene asportato dal bagno termico. Poi il campione viene isolato termicamente, ed il campo esterno viene azzerato adiabaticamente.

I materiali nei quali il raffreddamento magnetico è più evidente sono sali paramagnetici, come il CMN (nitrato di cerio-magnesio), e leghe di gadolinio. Nei sali paramagnetici il paramagnetismo è dovuto allo spin degli elettroni periferici degli atomi paramagnetici.

13.2 Energia

Poiché gli elettroni hanno spin $\frac{1}{2}$, ci sono solo due orientazioni possibili rispetto al campo magnetico $\mathbf{B}$, che supponiamo orientato lungo l’asse z . Le due orientazioni corrispondono ai due stati stazionari di particella singola $|\uparrow\rangle$ e $|\downarrow\rangle$, in cui lo spin è rispettivamente parallelo e antiparallelo al campo. La componente m_z del momento magnetico dell’elettrone $\mathbf{m}$ ha così due valori possibili

$$m_z = \mp g\mu_B \frac{S_z}{\hbar}, \text{ dove } \mu_B = \frac{e\hbar}{2m_e} \simeq 9.27 \times 10^{-24} \text{ J/T è il magnetone di Bohr, } S_z = \pm \frac{1}{2}\hbar,$$

e $g = 2.002319 \dots$ è il fattore g dell’elettrone. La discordanza tra i segni è dovuta alla carica negativa dell’elettrone. Per l’energia di particella singola $U = -\mathbf{m} \cdot \mathbf{B}$, tenendo conto che $g \simeq 2$, abbiamo

quindi i due autovalori

$$E_{\uparrow} = \mu_B B \quad \text{e} \quad E_{\downarrow} = -\mu_B B.$$

Per trattare l'insieme degli elettroni che danno origine al paramagnetismo possiamo usare la statistica classica (Boltzmann) perché, avendo a che fare con un solido in cui gli atomi paramagnetici sono presenti come impurità, la distanza media tra gli elettroni responsabili del paramagnetismo è grande, le sovrapposizioni tra le funzioni d'onda sono trascurabili, e non facciamo errori apprezzabili se li consideriamo come distinguibili. Se abbiamo un totale di N elettroni, il numero di elettroni con spin rispettivamente parallelo e antiparallelo al campo sarà

$$N_{\uparrow} = \frac{N}{Z_{\text{sing}}} e^{-\mu_B B/kT}, \quad N_{\downarrow} = \frac{N}{Z_{\text{sing}}} e^{\mu_B B/kT}, \quad (13.1)$$

dove Z_{sing} è la funzione di partizione per la particella singola

$$Z_{\text{sing}} = e^{-\mu_B B/kT} + e^{\mu_B B/kT} = 2 \cosh\left(\frac{\mu_B B}{kT}\right). \quad (13.2)$$

Dalle relazioni (3.33) abbiamo per il moltiplicatore di Lagrange Ω_{sing} , relativo alla normalizzazione della matrice densità per la particella singola,

$$\Omega_{\text{sing}} = -\ln Z_{\text{sing}} = -\ln[2 \cosh(\beta\mu_B B)], \quad (13.3)$$

dove $\beta = 1/kT$. Per l'energia totale U del campione abbiamo

$$\begin{aligned} U &= N_{\uparrow}\mu_B B - N_{\downarrow}\mu_B B = \mu_B B \frac{N}{Z} (e^{-\mu_B B/kT} - e^{\mu_B B/kT}) \\ &= -\mu_B B N \tanh\left(\frac{\mu_B B}{kT}\right), \end{aligned} \quad (13.4)$$

che può anche essere ottenuta da Ω_{sing} nel modo seguente

$$\begin{aligned} U &= N U_{\text{sing}} = N \frac{\partial \Omega_{\text{sing}}}{\partial \beta} = -N \frac{\partial}{\partial \beta} \ln[2 \cosh(\beta\mu_B B)] \\ &= -N \frac{1}{2 \cosh(\beta\mu_B B)} 2 \sinh(\beta\mu_B B) \mu_B B = -\mu_B B N \tanh\left(\frac{\mu_B B}{kT}\right), \end{aligned} \quad (13.5)$$

dove $U_{\text{sing}} = \partial \Omega_{\text{sing}} / \partial \beta$ è l'energia media della particella singola.

Per il contributo della magnetizzazione alla capacità termica abbiamo

$$C = \left. \frac{\partial U}{\partial T} \right|_{N,p},$$

da cui

$$C = -\mu_B B N \frac{1}{\cosh^2\left(\frac{\mu_B B}{kT}\right)} \left(-\frac{\mu_B B}{kT^2}\right) = N \frac{\mu_B^2 B^2}{kT^2} \frac{1}{\cosh^2\left(\frac{\mu_B B}{kT}\right)} \quad (13.6)$$

La dipendenza di $C/(Nk)$ dal parametro $\mu_B B/kT$ è mostrata in fig. 13.1. Il calore specifico tende a zero sia per $B \rightarrow 0$, corrispondente a $\frac{\mu_B B}{kT} \rightarrow 0$, che per $T \rightarrow 0$, corrispondente a $\frac{\mu_B B}{kT} \rightarrow \infty$, in accordo con il terzo principio della termodinamica.

Se $\mu_B B \ll kT$ possiamo approssimare $\cosh(\mu_B B/kT) \simeq 1$, da cui $Z \simeq 2$ e per $N_{\uparrow, \downarrow}$ abbiamo

$$N_{\uparrow} = \frac{N}{2} \left(1 - \frac{\mu_B B}{kT} \right) \quad (13.7)$$

$$N_{\downarrow} = \frac{N}{2} \left(1 + \frac{\mu_B B}{kT} \right). \quad (13.8)$$

Per la magnetizzazione abbiamo così

$$M_z = -\mu_B N_{\uparrow} + \mu_B N_{\downarrow} = \frac{N}{2} \left(2 \frac{\mu_B^2 B}{kT} \right) = N \frac{\mu_B^2 B}{kT} = \chi B, \quad (13.9)$$

dove

$$\chi = \frac{\mu_B^2}{kT} \quad (13.10)$$

è, a campi deboli, la suscettività magnetica del materiale.

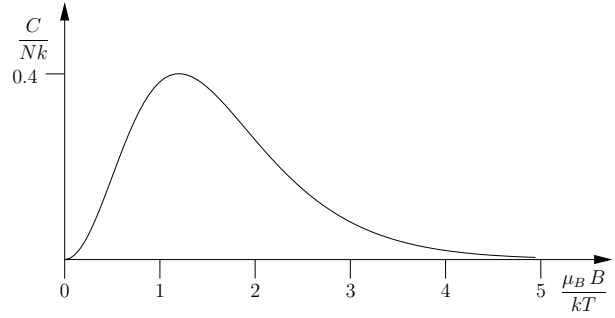


Figura 13.1 Dipendenza del calore specifico dal parametro $\mu_B B/kT$.

13.3 Entropia e raffreddamento magnetico

Con le nostre ipotesi abbiamo per l'entropia

$$S = N S_{\text{sing}} = N \left(-k \Omega_{\text{sing}} + k \beta U_{\text{sing}} \right) = -k N \Omega_{\text{sing}} + k \beta U,$$

dove S_{sing} è l'entropia per la particella singola. Infatti, avendo supposto che non ci sia interazione tra i singoli elettroni che contribuiscono al paramagnetismo, l'informazione mancante totale sarà semplicemente la somma delle informazioni mancanti relative ai singoli elettroni. Inserendo le (13.3) e (13.4), e sostituendo ovunque β con $1/(kT)$, otteniamo

$$\begin{aligned} S &= kN \ln \left[2 \cosh \left(\frac{\mu_B B}{kT} \right) \right] - kN \frac{\mu_B B}{kT} \tanh \left(\frac{\mu_B B}{kT} \right) \\ &= kN \left\{ \ln 2 + \ln \left[\cosh \left(\frac{\mu_B B}{kT} \right) \right] - \frac{\mu_B B}{kT} \tanh \left(\frac{\mu_B B}{kT} \right) \right\}. \end{aligned} \quad (13.11)$$

Introducendo la variabile $x = \frac{\mu_B B}{kT}$ abbiamo

$$\lim_{T \rightarrow \infty} S = \lim_{x \rightarrow 0} S = N k \ln 2, \quad (13.12)$$

e

$$\begin{aligned}
\lim_{T \rightarrow 0} S &= \lim_{x \rightarrow \infty} S = kN \lim_{x \rightarrow \infty} [\ln 2 + \ln(\cosh x) - x \tanh x] \\
&= kN \lim_{x \rightarrow \infty} \left[\ln 2 + \ln \left(\frac{e^x + e^{-x}}{2} \right) - x \frac{e^x - e^{-x}}{e^x + e^{-x}} \right] \\
&= kN \lim_{x \rightarrow \infty} (\ln 2 + x - \ln 2 - x) = 0,
\end{aligned} \tag{13.13}$$

avendo tenuto conto del fatto che $\lim_{x \rightarrow \infty} e^{-x} = 0$. L'andamento di $S/(Nk)$ in funzione del parametro $kT/(\mu_B B)$ è mostrato in figura 13.2.

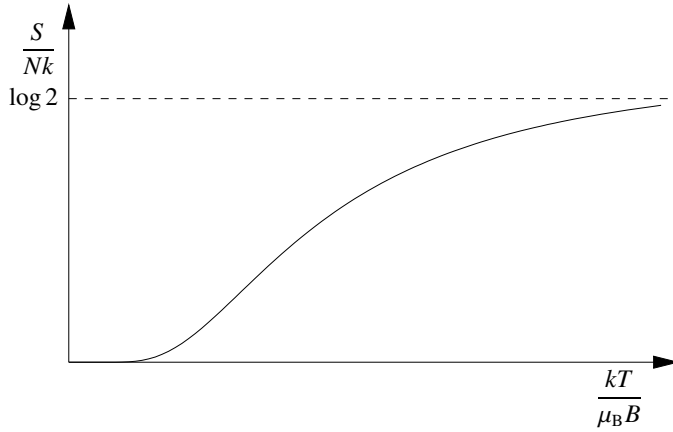


Figura 13.2 Dipendenza dell'entropia da $kT/(\mu_B B)$.

pre in contatto con il bagno termico, quindi a temperatura costante, il campo magnetico sul campione viene aumentato fino al valore B_f , corrispondente allo stato b .

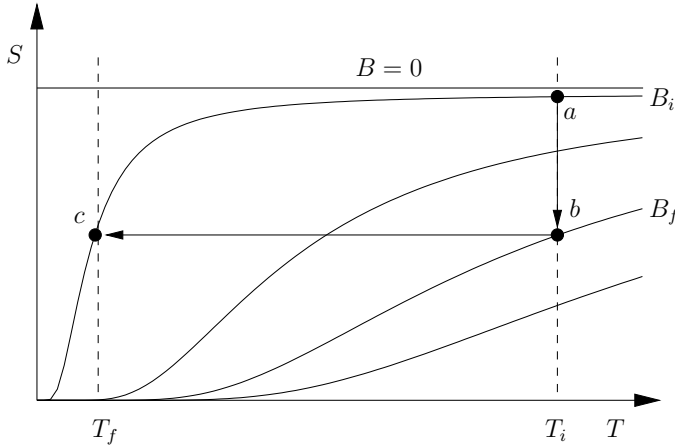


Figura 13.3 Raffreddamento magnetico nel piano ST .

E' da notare che, in assenza di campo magnetico, cioè per $B = 0$, il parametro $kT/(\mu_B B)$ diverge per ogni valore di T diverso da 0. Quindi, per $B = 0$, abbiamo $S = Nk \ln 2$ per ogni valore di T , tranne $T = 0$.

Passiamo adesso a considerare il raffreddamento magnetico. Nella prima fase il campione viene messo in contatto termico con un bagno di elio liquido, alla temperatura T_i , in presenza di un campo magnetico B_i , che non sarà mai rigorosamente nullo. Questo stato iniziale corrisponde al punto a nel grafico di figura 13.3. Sem-

Nella seconda fase il campione viene isolato termicamente, ed il campo applicato viene riportato lentamente, in un processo *adiabatico*, al valore iniziale B_i . In questo processo, rappresentato dalla linea orizzontale $\overline{bc}$ in Fig. 13.3, l'entropia rimane costante. Dalla (13.13) vediamo che l'entropia dipende dalla temperatura e dal campo magnetico solo attraverso il rapporto $x = \frac{\mu_B B}{kT}$, quindi la relazione $S_c = S_b$ implica che sia

$$\frac{\mu_B B_c}{kT_c} = \frac{\mu_B B_b}{kT_b}, \quad \text{ovvero} \quad \frac{B_i}{T_f} = \frac{B_f}{T_i},$$

da cui ricaviamo per la temperatura finale $T_f = T_i \frac{B_i}{B_f}$. Quindi, apparentemente, partendo da e tornando a campo nullo, cioè $B_i = 0$, passando per un qualunque valore $B_f \neq 0$, dovrebbe essere possibile raggiungere lo zero assoluto! Questo in realtà non è possibile perché non è possibile realizzare un

campo rigorosamente nullo. Il processo di raffreddamento isothermico può essere iterato per raggiungere temperature sempre più basse in più stadi. Ma anche a questo c'è un limite, fissato dalla necessità di trovare, ad ogni stadio, un bagno termico, a temperatura ogni volta più bassa, per la trasformazione isoterma $\overline{ab}$. Ovviamente non sono disponibili bagni termici a temperatura arbitrariamente bassa, che, oltre a tutto, renderebbero inutile il raffreddamento magnetico!

Capitolo 14

Risonanza magnetica

14.1 Descrizione classica

Come modello classico di partenza consideriamo una distribuzione rigida di carica elettrica vincolata a ruotare attorno a un proprio asse di inerzia con velocità angolare ω costante, come in Fig. 14.1. Il sistema avrà così un momento magnetico μ e un momento angolare $\mathbf{L}$ proporzionali tra loro. Possiamo quindi scrivere $\mu = \gamma \mathbf{L}$, dove γ è una costante chiamata *rapporto giromagnetico* del sistema (vedremo più avanti la differenza tra *rapporto giromagnetico* e *fattore giromagnetico*). Se sia la carica totale q che la massa totale m sono distribuite uniformemente sul corpo, si dimostra che

$$\gamma = \frac{q}{2m}. \quad (14.1)$$

Introduciamo adesso un campo magnetico uniforme $\mathbf{B}$ che formi un angolo ϑ con μ (e quindi con $\mathbf{L}$), come in Fig. 14.2. Il nostro sistema sarà sottoposto a un momento meccanico $\Gamma = \mu \times \mathbf{B}$, che porta a un'equazione di evoluzione temporale per il momento angolare $\mathbf{L}$

$$\frac{d\mathbf{L}}{dt} = \mu \times \mathbf{B}. \quad (14.2)$$

Data la proporzionalità ipotizzata tra μ e $\mathbf{L}$ possiamo scrivere

$$\frac{d\mu}{dt} = \mu \times (\gamma \mathbf{B}). \quad (14.3)$$

Questa equazione, che è valida indipendentemente dal fatto che $\mathbf{B}$ sia o meno costante nel tempo, ci dice che, istante per istante, le variazioni di μ sono perpendicolari sia a μ stesso che a $\mathbf{B}$. In altre parole, se nella Fig. 14.2 consideriamo fissa la “coda” del vettore μ , la sua punta si muoverà in modo da uscire dalla pagina con l'angolo ϑ che resta costante. Si ha, cioè, un moto di precessione e, se $\mathbf{B}$ è costante nel tempo, il vettore μ descriverà un cono. Una tecnica che rende più facile la soluzione della (14.3) è l'introduzione di un sistema di riferimento rotante $x'y'z'$. Consideriamo una funzione vettoriale del tempo $\mathbf{f}(t)$ che, nel nostro sistema di riferimento rotante avrà componenti $f_{x'}(t)$, $f_{y'}(t)$ e $f_{z'}(t)$. Supponiamo che il sistema cartesiano $x'y'z'$ abbia l'origine

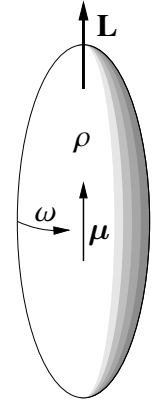


Figura 14.1 Distribuzione di carica rotante.

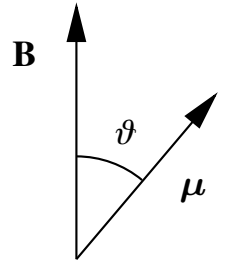


Figura 14.2 Carica rotante in campo magnetico.

coincidente con quella del sistema xyz , inerziale, del laboratorio, ma ruoti rigidamente, con velocità angolare istantanea $\boldsymbol{\omega}$, attorno a un asse passante per l'origine comune. Per l'evoluzione temporale dei tre versori del sistema rotante abbiamo

$$\frac{d\hat{\mathbf{x}}'}{dt} = \boldsymbol{\omega} \times \hat{\mathbf{x}}', \quad \frac{d\hat{\mathbf{y}}'}{dt} = \boldsymbol{\omega} \times \hat{\mathbf{y}}', \quad \frac{d\hat{\mathbf{z}}'}{dt} = \boldsymbol{\omega} \times \hat{\mathbf{z}}', \quad (14.4)$$

quindi l'equazione dell'evoluzione temporale di $\mathbf{f} = \hat{\mathbf{x}}' f_{x'} + \hat{\mathbf{y}}' f_{y'} + \hat{\mathbf{z}}' f_{z'}$ si può scrivere

$$\begin{aligned} \frac{d\mathbf{f}}{dt} &= \hat{\mathbf{x}}' \frac{df_{x'}}{dt} + f_{x'} \frac{d\hat{\mathbf{x}}'}{dt} + \hat{\mathbf{y}}' \frac{df_{y'}}{dt} + f_{y'} \frac{d\hat{\mathbf{y}}'}{dt} + \hat{\mathbf{z}}' \frac{df_{z'}}{dt} + f_{z'} \frac{d\hat{\mathbf{z}}'}{dt} \\ &= \hat{\mathbf{x}}' \frac{df_{x'}}{dt} + \hat{\mathbf{y}}' \frac{df_{y'}}{dt} + \hat{\mathbf{z}}' \frac{df_{z'}}{dt} + \boldsymbol{\omega} \times (\hat{\mathbf{x}}' f_{x'} + \hat{\mathbf{y}}' f_{y'} + \hat{\mathbf{z}}' f_{z'}) \\ &= \frac{\delta \mathbf{f}}{\delta t} + \boldsymbol{\omega} \times \mathbf{f}, \end{aligned} \quad (14.5)$$

dove $\delta \mathbf{f} / \delta t$ indica la derivata di $\mathbf{f}$ nel sistema di riferimento rotante $x'y'z'$. Uguagliando il secondo membro della (14.5) al secondo membro della (14.3) otteniamo la relazione

$$\frac{\delta \boldsymbol{\mu}}{\delta t} + \boldsymbol{\omega} \times \boldsymbol{\mu} = \boldsymbol{\mu} \times \gamma \mathbf{B}, \quad \text{da cui} \quad \frac{\delta \boldsymbol{\mu}}{\delta t} = \boldsymbol{\mu} \times (\gamma \mathbf{B} + \boldsymbol{\omega}). \quad (14.6)$$

In altre parole, l'equazione dell'evoluzione temporale di $\boldsymbol{\mu}$ in un sistema di riferimento rotante con frequenza $\boldsymbol{\omega}$ è la stessa che nel sistema del laboratorio, a condizione di sostituire il campo magnetico applicato $\mathbf{B}$ con un campo magnetico efficace $\mathbf{B}_{\text{eff}}$ dato da

$$\mathbf{B}_{\text{eff}} = \mathbf{B} + \frac{\boldsymbol{\omega}}{\gamma}. \quad (14.7)$$

Così, un modo facile per risolvere il moto di $\boldsymbol{\mu}$ in un campo magnetico statico $\mathbf{B} = B_0 \hat{\mathbf{z}}$ è porsi in un sistema rotante a velocità angolare $-\boldsymbol{\omega}_L$, con $\boldsymbol{\omega}_L = \gamma B_0 \hat{\mathbf{z}}$. Nel sistema rotante abbiamo infatti $\mathbf{B}_{\text{eff}} = (B_0 - \omega_L / \gamma) \hat{\mathbf{z}} = 0$, da cui $\delta \boldsymbol{\mu} / \delta t = 0$, e $\boldsymbol{\mu}$ è costante nel tempo. Quindi, nel sistema del laboratorio $\boldsymbol{\mu}$ precede a velocità angolare $-\boldsymbol{\omega}_L$. La velocità angolare $\omega_L = \gamma B_0$ è detta *frequenza di Larmor*. La componente μ_z di $\boldsymbol{\mu}$ è costante nel tempo, mentre le componenti μ_x e μ_y oscillano alla frequenza di Larmor. Con un'opportuna scelta dell'origine dei tempi possiamo avere, per esempio,

$$\mu_x = \mu \sin \vartheta \cos \omega_L t, \quad \text{e} \quad \mu_y = -\mu \sin \vartheta \sin \omega_L t. \quad (14.8)$$

14.2 Descrizione quantistica

Consideriamo una particella dotata di spin $\hat{\mathbf{S}}$, quindi con momento angolare $\hbar \hat{\mathbf{S}}$, posta in un campo magnetico statico $\mathbf{B}_0$ diretto lungo l'asse z . L'hamiltoniana classica $H = -\boldsymbol{\mu} \cdot \mathbf{B}_0$ diventa

$$\hat{H}_0 = -\gamma \hbar \hat{\mathbf{S}} \cdot \mathbf{B}_0, \quad \text{con autovalori} \quad E_m = -\gamma \hbar B_0 m, \quad (14.9)$$

dove γ è il rapporto giromagnetico e gli m sono gli autovalori della componente $\hat{S}_z$ di $\hat{\mathbf{S}}$. Invece del rapporto giromagnetico γ (rapporto tra momento magnetico e spin della particella, quindi misurata

in $\text{s}^{-1}\text{T}^{-1}$, o, equivalentemente, in $\text{C}\cdot\text{kg}^{-1}$ nel sistema internazionale SI), spesso si preferisce usare il *fattore giromagnetico* g , quantità adimensionale definita come

$$g = -\frac{\gamma\hbar}{\mu_B}, \quad \text{dove} \quad \mu_B = \frac{e\hbar}{2m_e} = 9.27400915(23) \times 10^{-24} \text{ JT}^{-1} \quad (14.10)$$

è il magnetone di Bohr, mentre m_e ed e sono, rispettivamente, la massa e il valore assoluto della carica dell'elettrone. Il segno meno nella definizione di g tiene conto del fatto che la carica dell'elettrone è negativa. L'hamiltoniana diventa così

$$\hat{H}_0 = g\mu_B B_0 \hat{S}_z. \quad (14.11)$$

Se la nostra particella è un elettrone isolato avremo $g \simeq 2.002\,332$, che può essere approssimato con 2 ai nostri fini. Sempre per un elettrone isolato abbiamo per il rapporto giromagnetico

$$\gamma_e = -1.760\,859\,770(44) \times 10^{11} \text{ rad s}^{-1} \text{ T}^{-1}. \quad (14.12)$$

Se, in vece di un elettrone, stiamo trattando un protone, un neutrone o un nucleo atomico, si preferisce scrivere

$$g = \frac{\gamma\hbar}{\mu_N}, \quad \text{dove} \quad \mu_N = \frac{e\hbar}{2m_p} = 5.05078324(13) \times 10^{-27} \text{ JT}^{-1} \quad (14.13)$$

è il magnetone nucleare, ed m_p è la massa del protone. Per il fattore giromagnetico di un protone abbiamo

$$\gamma_p = 2.675\,222\,099(70) \times 10^8 \text{ rad s}^{-1} \text{ T}^{-1}. \quad (14.14)$$

Tornando, per fissare le idee, al caso dell'elettrone, scegliamo un base in cui $\hat{S}_z$ è diagonale, con autovalori $\pm\frac{1}{2}$. In questa stessa base anche $\hat{H}_0$ è diagonale con autovalori

$$E_{+\frac{1}{2}} = \frac{1}{2} g\mu_B B_0, \quad E_{-\frac{1}{2}} = -\frac{1}{2} g\mu_B B_0, \quad (14.15)$$

che possono essere riscritti

$$E_{\pm\frac{1}{2}} = \pm\frac{1}{2} g\mu_B B_0 = \pm\frac{1}{2} \hbar\omega_0, \quad \text{con} \quad \omega_0 = \frac{g\mu_B}{\hbar} B_0 = -\gamma B_0. \quad (14.16)$$

Usando il rapporto giromagnetico della (14.12) abbiamo $\omega_0 \simeq -1.76 \times 10^{11} B_0 \text{ rad s}^{-1}$, dove B_0 è espresso in Tesla. La dipendenza temporale degli autostati dell'hamiltoniana è

$$\left|\frac{1}{2}, t\right\rangle = \left|\frac{1}{2}, 0\right\rangle e^{-i\omega_0 t/2} \quad \text{e} \quad \left|-\frac{1}{2}, t\right\rangle = \left|-\frac{1}{2}, 0\right\rangle e^{i\omega_0 t/2}. \quad (14.17)$$

Supponiamo che la nostra particella all'istante $t = 0$ si trovi, anziché in un autostato dell'hamiltoniana, nell'autostato, non stazionario, $|m_x = \frac{1}{2}\rangle$ di $\hat{S}_x$. Dato che abbiamo

$$|\psi, 0\rangle = |m_x = \frac{1}{2}\rangle = \frac{1}{\sqrt{2}}(|\frac{1}{2}\rangle + |-\frac{1}{2}\rangle) \quad (14.18)$$

l'evoluzione temporale successiva dello stato sarà

$$|\psi, t\rangle = \frac{1}{\sqrt{2}}(|\frac{1}{2}, 0\rangle e^{-i\omega_0 t/2} + |-\frac{1}{2}, 0\rangle e^{i\omega_0 t/2}). \quad (14.19)$$

L'evoluzione temporale del valore di aspettazione della componente x del momento magnetico della particella, $\hat{\mu}_x = \gamma\hbar\hat{S}_x$, può essere calcolato tramite il valor medio sullo stato $|\psi, t\rangle$ della matrice di Pauli corrispondente a $\hat{S}_x$

$$\hat{S}_x = \begin{pmatrix} 0 & \frac{1}{2} \\ \frac{1}{2} & 0 \end{pmatrix} = \frac{1}{2} \left(\left| \frac{1}{2} \right\rangle \left\langle -\frac{1}{2} \right| + \left| -\frac{1}{2} \right\rangle \left\langle \frac{1}{2} \right| \right). \quad (14.20)$$

Avremo

$$\langle \hat{\mu}_x(t) \rangle = \gamma\hbar \langle \psi, t | \hat{S}_x | \psi, t \rangle = \frac{\gamma\hbar}{4} (e^{i\omega_0 t} + e^{-i\omega_0 t}) = \frac{\gamma\hbar}{2} \cos \omega_0 t, \quad (14.21)$$

e analogamente possiamo ottenere per $\langle \hat{\mu}_y(t) \rangle$

$$\langle \hat{\mu}_y(t) \rangle = -\frac{\gamma\hbar}{2} \sin \omega_0 t, \quad (14.22)$$

in accordo con la precessione classica (14.8).

Più in generale, supponiamo di avere una particella di momento angolare qualunque $\hat{\mathbf{J}}$ e con momento magnetico $\hat{\boldsymbol{\mu}} = \gamma\hbar\hat{\mathbf{J}}$, situata in un campo magnetico $\mathbf{B}$. Per tener conto di campi magnetici sia statici che variabili nel tempo (questo è il motivo per cui abbiamo scritto $\mathbf{B}$ anziché $\mathbf{B}_0$) ci mettiamo in un sistema di riferimento il cui asse z coincida, istante per istante, con la direzione di $\mathbf{B}$. Così possiamo scrivere l'hamiltoniana nella forma

$$\hat{H} = -\gamma\hbar B \hat{J}_z. \quad (14.23)$$

Le relazioni di commutazione tra le componenti del momento angolare si ottengono permutando ciclicamente gli indici x, y e z nell'equazione

$$[\hat{J}_x, \hat{J}_y] = i\hat{J}_z. \quad (14.24)$$

Quindi, nella rappresentazione di Heisenberg, avremo

$$\frac{d\hat{J}_x}{dt} = \frac{i}{\hbar} [\hat{H}, \hat{J}_x] = -\gamma B i [\hat{J}_z, \hat{J}_x] = \gamma B \hat{J}_y, \quad (14.25)$$

e, analogamente, possiamo ottenere

$$\frac{d\hat{J}_y}{dt} = -\gamma B \hat{J}_x \quad \text{e} \quad \frac{d\hat{J}_z}{dt} = 0. \quad (14.26)$$

Le (14.25) e (14.26) possono essere compattate nell'espressione

$$\frac{d\hat{\mathbf{J}}}{dt} = \hat{\mathbf{J}} \times \gamma\mathbf{B}. \quad (14.27)$$

Dato che l'operatore momento magnetico della particella $\hat{\boldsymbol{\mu}}$ è legato all'operatore momento angolare dalla relazione $\hat{\boldsymbol{\mu}} = \gamma\hbar\hat{\mathbf{J}}$, per il suo valore d'aspettazione abbiamo

$$\frac{d}{dt} \langle \hat{\boldsymbol{\mu}} \rangle = \langle \hat{\boldsymbol{\mu}} \rangle \times \gamma\mathbf{B}, \quad (14.28)$$

che coincide con l'equazione classica (14.3). Quindi il valore di aspettazione del momento magnetico obbedisce all'equazione del moto classica. In un punto del campione di vettore posizione $\mathbf{r}$ possiamo definire la grandezza locale mesoscopica *magnetizzazione* $\hat{\mathbf{M}}(\mathbf{r})$ come

$$\hat{\mathbf{M}}(\mathbf{r}) = \frac{1}{\Delta V} \sum_{k=1}^N \hat{\boldsymbol{\mu}}_k, \quad (14.29)$$

dove ΔV è un volume *mesoscopico* centrato attorno alla posizione $\mathbf{r}$. Mesoscopico significa che ΔV è ancora molto piccolo da un punto di vista macroscopico (ordine di grandezza dei μm), ma tale da contenere un numero N di particelle sufficientemente grande da essere trattato statisticamente. Ogni particella contenuta in ΔV ha un momento magnetico $\hat{\boldsymbol{\mu}}_k$, con $k = 1, 2, \dots, N$. Se le particelle non interagiscono tra loro l'equazione (14.28) sarà valida anche sostituendo $\hat{\mathbf{M}}(\mathbf{r})$ per $\langle \hat{\boldsymbol{\mu}} \rangle$. Dato che, in pratica, quello che misureremo sarà sempre la magnetizzazione del campione, l'evoluzione temporale del nostro risultato sperimentale sarà descritta correttamente dall'equazione classica (14.3), sempre a condizione che i momenti magnetici delle singole particelle non interagiscano tra loro.

E' importante ricordare che la (14.28) è valida anche per un campo magnetico $\mathbf{B}$ dipendente dal tempo, non solo per un campo statico.

14.3 Campo magnetico rotante

Supponiamo che le nostre particelle, dotate di spin e momento magnetico, si trovino in presenza di un campo magnetico stazionario $\mathbf{B}_0$ diretto lungo l'asse z . In condizioni di equilibrio termodinamico, non esistendo una direzione privilegiata perpendicolare a z , le componenti x e y dei momenti magnetici $\hat{\boldsymbol{\mu}}_k$ delle singole particelle si annullano l'una con l'altra. La magnetizzazione mesoscopica osservata, pari al valore di aspettazione della (14.29),

$$\mathbf{M} = \left\langle \frac{1}{\Delta V} \sum_k \hat{\boldsymbol{\mu}}_k \right\rangle, \quad (14.30)$$

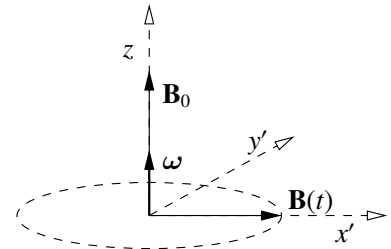


Figura 14.3 Campo magnetico rotante.

avrà quindi componente nulla sul piano xy . Sempre all'equilibrio termodinamico, la componente di $\mathbf{M}$ parallela a $\mathbf{B}_0$ sarà invece diversa da zero perché particelle con diversa componente z di $\hat{\boldsymbol{\mu}}$ hanno energia diversa, ed $\mathbf{M}$ sarà orientata lungo z . Naturalmente, perché questa magnetizzazione sia osservabile è necessario che μB_0 non sia trascurabile rispetto a kT .

Aggiungiamo adesso un secondo campo magnetico $\mathbf{B}(t)$ rotante nel piano perpendicolare a z con velocità angolare $\boldsymbol{\omega} = \omega_z \hat{\mathbf{z}}$, come in Fig.14.3. Cioè aggiungiamo a $\mathbf{B}_0$ un campo del tipo

$$\mathbf{B}(t) = \hat{\mathbf{x}} B_1 \cos \omega_z t + \hat{\mathbf{y}} B_1 \sin \omega_z t, \quad (14.31)$$

per quanto detto nei paragrafi precedenti, nel sistema del laboratorio l'evoluzione temporale della magnetizzazione osservata $\mathbf{M}$ sarà data dall'equazione

$$\frac{d\mathbf{M}}{dt} = \mathbf{M} \times \gamma [\mathbf{B}_0 + \mathbf{B}(t)]. \quad (14.32)$$

Adesso mettiamoci in un sistema di riferimento rotante $x'y'z$ solidale con $\mathbf{B}(t)$, cioè con l'asse z lungo il campo statico $\mathbf{B}_0$, l'asse x' allineato istante per istante con $\mathbf{B}(t)$, e l'asse y' perpendicolare a z e a x' . In questo sistema di riferimento sia $\mathbf{B}_0$ che $\mathbf{B}(t)$ sono costanti nel tempo, e, in base alla (14.6), l'equazione del moto per la magnetizzazione si scrive

$$\frac{\delta \mathbf{M}}{\delta t} = \mathbf{M} \times \gamma \left[\mathbf{B}_0 + \mathbf{B}_1 + \frac{\boldsymbol{\omega}}{\gamma} \right], \quad (14.33)$$

dove $\mathbf{B}_1 = \hat{\mathbf{x}}' B_1$. In questo sistema di riferimento, quindi, la magnetizzazione esegue il moto di precessione rappresentato dalla circonferenza (vista in prospettiva, ellissi!) tratteggiata di Fig. 14.4 attorno a un *campo magnetico efficace* costante $\mathbf{B}_{\text{eff}}$, con componente z uguale a $B_0 + \omega/\gamma$, e componente x' uguale a B_1 . Questa precessione, sempre osservata in $x'y'z$, avviene a velocità angolare $\boldsymbol{\omega}^* = -\gamma \mathbf{B}_{\text{eff}}$, con

$$\omega^* = |\boldsymbol{\omega}^*| = \gamma \sqrt{\left(B_0 + \frac{\omega}{\gamma} \right)^2 + B_1^2}. \quad (14.34)$$

Figura 14.4 Precessione della magnetizzazione $\mathbf{M}$ attorno al campo efficace $\mathbf{B}_{\text{eff}}$ osservata nel sistema di riferimento rotante.

Quindi la magnetizzazione, inizialmente allineata lungo z , torna ad allinearsi periodicamente con quella direzione, con periodo $T = 2\pi/\omega^*$. Vediamo adesso che cosa succede se $\boldsymbol{\omega}$ è antiparallela a $\mathbf{B}_0$, cioè se la rotazione di $\mathbf{B}(t)$ avviene in senso orario, nell'ipotesi che il rapporto giromagnetico γ sia positivo (ipotesi valida per magnetismo nucleare, ovviamente non per magnetismo elettronico). In questo caso ω_z è negativa, e, per comodità, definiamo la quantità positiva ω

$$\omega = -\omega_z, \quad (14.35)$$

$$\text{di valore} \quad \mathbf{B}_{\text{eff}} = \mathbf{B}_0 + \frac{\boldsymbol{\omega}}{\gamma} + \mathbf{B}_1, \quad (14.36)$$

in modo che il campo efficace si scriva

$$\mathbf{B}_{\text{eff}} = \hat{\mathbf{z}} \left(B_0 - \frac{\omega}{\gamma} \right) + \hat{\mathbf{x}}' B_1. \quad (14.37)$$

Se variamo ω mantenendo costanti i valori di B_0 e B_1 , la componente x' di $\mathbf{B}_{\text{eff}}$ rimane costante, mentre la sua componente z diminuisce al crescere di ω , fino ad annullarsi e cambiare di segno. Abbiamo risonanza quando $\omega = \omega_0 = \gamma B_0$. In queste condizioni la componente z di $\mathbf{B}_{\text{eff}}$, data dalla quantità tra parentesi nella (14.37), è nulla, e $\mathbf{B}_{\text{eff}}$ ha solo la componente x' , pari a B_1 . La precessione della magnetizzazione $\mathbf{M}$ avviene così nel piano $y'z$, come in Fig. 14.5, con frequenza $\omega_1 = -\gamma B_1$. Se il campo rotante viene acceso all'istante $t = 0$, dopo un intervallo di tempo $T_{\frac{\pi}{2}}$, dato da

$$T_{\frac{\pi}{2}} = \frac{\pi}{2\omega_1} = \frac{\pi}{2\gamma B_1}, \quad (14.38)$$

la magnetizzazione ha ruotato di 90° ed è allineata con l'asse y' . Se a questo punto il campo rotante viene spento, trascurando effetti di rilassamento la magnetizzazione continua a precedere restando

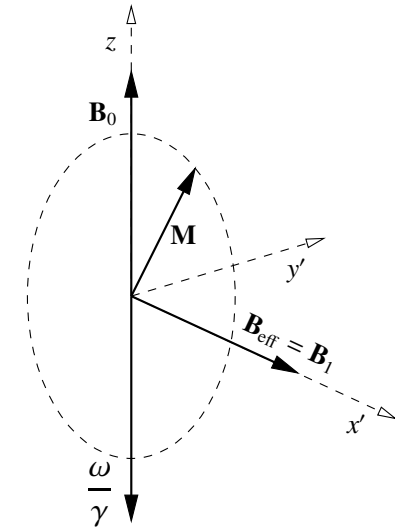


Figura 14.5 Precessione della magnetizzazione in condizioni di risonanza: $\omega = \gamma B_0$.

perpendicolare a $\mathbf{B}_0$. In questo caso si parla di un *impulso* $\pi/2$. Se invece il campo rotante viene spento dopo un intervallo di tempo T_π

$$T_\pi = \frac{\pi}{\omega_1} = \frac{\pi}{\gamma B_1} \quad (14.39)$$

la magnetizzazione si inverte, e si parla di un *impulso* π . Infine, dopo un *impulso* 2π , di durata $T_{2\pi} = 2T_\pi$, la magnetizzazione è tornata alla direzione originale.

14.4 Campo oscillante

Sperimentalmente è molto più facile aggiungere al campo statico $\mathbf{B}_0$ un campo perpendicolare $\mathbf{B}(t)$ oscillante piuttosto che un campo rotante. Basta infatti aggiungere un solenoide, o, più semplicemente, una coppia di bobine di Helmholtz con l'asse perpendicolare a $\mathbf{B}_0$ e alimentare con una corrente oscillante alla frequenza opportuna. Se l'asse del solenoide è diretto lungo $\hat{\mathbf{x}}$ avremo, con un'opportuna scelta dell'origine dei tempi, un campo oscillante del tipo

$$\mathbf{B}(t) = \hat{\mathbf{x}} B_1 \cos \omega t. \quad (14.40)$$

Un campo oscillante può essere considerato la sovrapposizione di due campi controrotanti, $\mathbf{B}_L(t)$ rotante in senso antiorario, e $\mathbf{B}_R(t)$ rotante in senso orario, con

$$\begin{aligned} \mathbf{B}_L(t) &= \hat{\mathbf{x}} \frac{B_1}{2} \cos \omega t + \hat{\mathbf{y}} \frac{B_1}{2} \sin \omega t \\ \mathbf{B}_R(t) &= \hat{\mathbf{x}} \frac{B_1}{2} \cos \omega t - \hat{\mathbf{y}} \frac{B_1}{2} \sin \omega t, \end{aligned} \quad (14.41)$$

dove $\mathbf{B}_L$ e $\mathbf{B}_R$ differiscono soltanto per la sostituzione di $+\omega$ con $-\omega$. E' da notare che la (14.31) ci dice che, allo stesso modo, un campo rotante può essere considerato come la sovrapposizione di due campi oscillanti in quadratura, perpendicolari tra loro e di uguale intensità.

Torniamo al nostro campo oscillante. In presenza di una magnetizzazione del mezzo, una delle due componenti controrotanti ruoterà nello stesso verso del moto di precessione, mentre l'altra ruoterà in verso opposto. Supponiamo che la frequenza di oscillazione di $\mathbf{B}(t)$ sia esattamente ω_0 , e mettiamoci nel sistema di riferimento rotante che ha $\mathbf{B}_R$ come asse x' (sempre nell'ipotesi di $\gamma > 0$). In questo sistema di riferimento la (14.33) diventa

$$\begin{aligned} \frac{\delta \mathbf{M}}{\delta t} &= \mathbf{M} \times \gamma \left[\mathbf{B}_0 + \hat{\mathbf{x}}' \frac{B_1}{2} - \frac{\omega_0}{\gamma} + \hat{\mathbf{x}}' \frac{B_1}{2} \cos 2\omega_0 t + \hat{\mathbf{y}}' \frac{B_1}{2} \sin 2\omega_0 t \right] \\ &= \mathbf{M} \times \gamma \left[\hat{\mathbf{x}}' \frac{B_1}{2} + \hat{\mathbf{x}}' \frac{B_1}{2} \cos 2\omega_0 t + \hat{\mathbf{y}}' \frac{B_1}{2} \sin 2\omega_0 t \right] \end{aligned} \quad (14.42)$$

e la precessione, osservata nel sistema rotante, dovrebbe avvenire attorno a un campo magnetico efficace dato dalla somma di un campo statico diretto lungo $\hat{\mathbf{x}}'$ e pari a $\frac{1}{2}B_1$, e di un campo, sempre di intensità $\frac{1}{2}B_1$, rotante a frequenza $2\omega_0$ nel piano $x'y'$. Nelle normali condizioni sperimentali di $B_1 \ll B_0$ la frequenza del campo rotante $2\omega_0 = 2\gamma B_0$ è molto maggiore della frequenza di precessione della magnetizzazione attorno alla componente "statica" del campo efficace, che vale $\frac{1}{2}\omega_1 = \frac{1}{2}\gamma B_1$. In

prima approssimazione, quindi, la precessione osservata nel sistema $x'y'z$ avviene attorno a un campo efficace medio pari alla sola componente “corotante” di $\mathbf{B}(t)$, mentre la rapida rotazione della componente “controrotante” viene mediata a zero durante una singola precessione. Questa è la cosiddetta *approssimazione di onda rotante*, o *approssimazione di campo rotante*. Con un campo oscillante di intensità $B_1 \ll B_0$ e frequenza prossima alla risonanza, quindi, tutto avviene come se fossimo in presenza della sola *componente corotante* del campo, di intensità $\frac{1}{2}B_1$. Al crescere dell'intensità del campo oscillante, quando la condizione $B_1 \ll B_0$ non è più verificata, la presenza della componente controrotante del campo oscillante non può più essere trascurata, e porta a uno spostamento della frequenza ω a cui si osserva la risonanza, fenomeno detto *effetto Bloch-Siegert*. Nel resto del capitolo faremo sempre l'ipotesi di essere nelle condizioni $B_1 \ll B_0$, in cui l'approssimazione di campo rotante è valida, e l'effetto Bloch-Siegert è trascurabile.

14.5 Rilassamento e equazioni di Bloch

Finora abbiamo considerato le singole particelle dotate di momento magnetico come interagenti solo con i campi. In realtà esistono sia le interazioni delle particelle tra loro, che le interazioni con l'ambiente in cui sono immerse, che possiamo considerare come *bagno termico*. In assenza di campo statico ci aspettiamo una magnetizzazione complessiva nulla, visto che tutte le possibili orientazioni per i momenti magnetici sono equiprobabili. Quando introduciamo il campo statico $\mathbf{B}_0$, in assenza di bagno termico non dovrebbe comparire una magnetizzazione, perché tutti i momenti comincerebbero semplicemente a precedere attorno a $\mathbf{B}_0$, mantenendo quindi costante la loro componente lungo il campo. In presenza di bagno termico ci aspettiamo che venga raggiunto l'equilibrio boltzmanniano con un'equazione del tipo

$$\frac{dM_z}{dt} = \frac{M_0 - M_z}{T_1}, \quad (14.43)$$

dove T_1 è una costante di tempo caratteristica, detta *tempo di rilassamento longitudinale*, M_z è la componente z della magnetizzazione $\mathbf{M}$ (l'unica componente non nulla nelle nostre condizioni), e M_0 è il valore di M_z all'equilibrio boltzmanniano in presenza di $\mathbf{B}_0$. Se il nostro sistema può essere descritto mediante una suscettività magnetica statica χ avremo

$$M_0 = \frac{\chi}{\mu_0} B_0. \quad (14.44)$$

Se combiniamo la (14.43) con la (14.3) mediata su tutti i momenti magnetici del sistema otteniamo

$$\frac{dM_z}{dt} = \gamma (\mathbf{M} \times \mathbf{B})_z + \frac{M_0 - M_z}{T_1}. \quad (14.45)$$

All'equilibrio termico in presenza di un campo statico diretto lungo z , le componenti x e y della magnetizzazione, se inizialmente non nulle per un qualunque motivo, tenderanno ad annullarsi. Possiamo descrivere questo fenomeno mediante due ulteriori equazioni

$$\begin{aligned} \frac{dM_x}{dt} &= \gamma (\mathbf{M} \times \mathbf{B})_x - \frac{M_x}{T_2} \\ \frac{dM_y}{dt} &= \gamma (\mathbf{M} \times \mathbf{B})_y - \frac{M_y}{T_2}, \end{aligned} \quad (14.46)$$

dove abbiamo introdotto un'ulteriore costante di tempo caratteristica T_2 , detta *tempo di rilassamento trasversale*, che, per motivi di simmetria di rotazione attorno all'asse z , abbiamo supposto uguale per le direzioni x e y . Possiamo giustificare il fatto che ci aspettiamo che T_2 possa essere diverso da T_1 ricordando che, in presenza di un campo statico, il rilassamento trasversale conserva l'energia del sistema, mentre il rilassamento longitudinale implica uno scambio di energia tra sistema e bagno termico. Invece l'ipotesi che il rilassamento avvenga in forma esponenziale è una nostra ipotesi arbitraria, o postulato, giustificata a posteriori se le nostre equazioni descrivono correttamente i risultati sperimentali. E' bene però notare che esistono casi in cui il rilassamento non avviene in modo esponenziale. Le equazioni (14.45) e (14.46) sono state introdotte da Felix Bloch nel 1946, e sono dette *equazioni di Bloch*.

14.6 Equazioni di Bloch a campi oscillanti deboli

Partendo dalle condizioni di equilibrio in presenza di $\mathbf{B}_0$ e dei fenomeni di rilassamento, introduciamo un campo rotante definito dalla (14.31) con la condizione $B_1 \ll B_0$. In base a quanto visto nel paragrafo 14.4 quanto osserveremo sarà analogo a quanto si osserva nel caso, più facile da realizzare sperimentalmente, dell'introduzione di un campo oscillante di intensità $2B_1$, per il quale abbiamo visto che siamo autorizzati a considerare la sola componente corotante. Supponiamo quindi di avere solo il campo rotante della (14.31) e mettiamoci nel sistema di riferimento con l'asse z parallelo a $\mathbf{B}_0$ e l'asse x' parallelo a $\mathbf{B}_1$, e definiamo la quantità $b_0 = B_0 - \omega/\gamma$. Le equazioni di Bloch diventano

$$\begin{aligned}\frac{\delta M_z}{\delta t} &= -\gamma M_{y'} B_1 + \frac{M_0 - M_z}{T_1} \\ \frac{\delta M_{x'}}{\delta t} &= +\gamma M_{y'} b_0 - \frac{M_{x'}}{T_2} \\ \frac{\delta M_{y'}}{\delta t} &= +\gamma [M_z B_1 - M_{x'} b_0] - \frac{M_{y'}}{T_2} .\end{aligned}\quad (14.47)$$

Poiché M_x ed M_y devono tendere a zero per B_1 che tende a zero, vediamo dalla prima delle (14.47) che il valore di M_z allo stato stazionario può differire da M_0 solo per termini proporzionale a B_1^2 . Possiamo quindi approssimare M_z con M_0 nella terza delle (14.47). Se poi introduciamo la combinazione $M_+ = M_{x'} + iM_{y'}$ otteniamo dalle ultime due delle (14.47)

$$\frac{\delta M_+}{\delta t} = -\alpha M_+ + i\gamma M_0 B_1 , \quad \text{dove} \quad \alpha = \frac{1}{T_2} + i\gamma b_0 . \quad (14.48)$$

La soluzione dell'equazione differenziale (14.48) è

$$M_+ = A e^{-\alpha t} + \frac{i\gamma M_0 B_1}{1/T_2 + i\gamma b_0} . \quad (14.49)$$

Se trascuriamo il transiente che va esponenzialmente a zero al passare del tempo, sostituiamo $M_0 = \chi/\mu_0 B_0$ e definiamo $\omega_0 = \gamma B_0$, otteniamo, separando parte reale e parte immaginaria della (14.49)

$$\begin{aligned} M_{x'} &= \frac{\chi \omega_0 T_2 B_1}{\mu_0} \frac{(\omega_0 - \omega) T_2}{1 + (\omega - \omega_0)^2 T_2^2} \\ M_{y'} &= \frac{\chi \omega_0 T_2 B_1}{\mu_0} \frac{1}{1 + (\omega - \omega_0)^2 T_2^2}. \end{aligned} \quad (14.50)$$

Le (14.50) ci dicono che, all'equilibrio in presenza di campo statico, campo rotante e rilassamento, la magnetizzazione trasversale è costante nel sistema rotante $x'y'z$. Inserendo la seconda delle (14.50) nella prima delle (14.47) vediamo che, sempre in presenza di rilassamento, esiste una soluzione stazionaria anche per M_z , data da

$$M_z = M_0 - \gamma \frac{\chi \omega_0 T_2 T_1 B_1^2}{\mu_0} \frac{1}{1 + (\omega - \omega_0)^2 T_2^2}, \quad (14.51)$$

che, come previsto, differisce da M_0 solo per una quantità proporzionale a B_1^2 . La magnetizzazione globale, in condizioni stazionarie, è così costante nel sistema $x'y'z$, e quindi ruota a frequenza ω nel sistema di laboratorio. Pertanto, se aggiungiamo una bobina orientata lungo l'asse x (o l'asse y) del laboratorio, la magnetizzazione vi indurrà una forza elettromotrice $\mathcal{E}$, che può essere misurata.

14.7 Misura del tempo di rilassamento trasversale T_2

Partiamo da un sistema di spin in presenza di un campo statico $\mathbf{B}_0$, che indurrà una magnetizzazione $\mathbf{M}_0 = \chi/\mu_0 \mathbf{B}_0$ parallela a sé stesso, e dei fenomeni di rilassamento. All'istante $t = 0$ applichiamo un impulso $\pi/2$ che ruoti la magnetizzazione fino a portarla sul piano xy . Dopo l'impulso ci troveremo

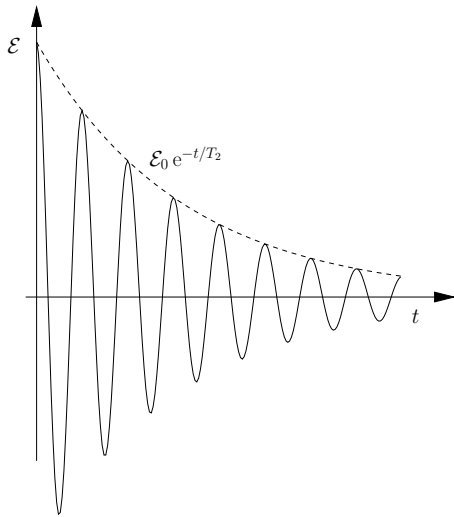


Figura 14.6 Smorzamento della f.e.m. indotta dalla magnetizzazione trasversale in una bobina dopo un impulso $\pi/2$.

quindi una magnetizzazione trasversale $M_\perp$, uguale in modulo alla magnetizzazione longitudinale originale M_0 , che comincerà a precedere attorno a $\mathbf{B}_0$ con frequenza ω_0 nel piano xy . Quindi, se mettiamo una bobina con l'asse diretto, per esempio, lungo x , misuriamo una forza elettromotrice indotta $\mathcal{E}$ ai suoi capi. Naturalmente, in presenza di fenomeni di rilassamento gli impulsi $\pi/2$, π e 2π definiti dalle equazioni (14.38) e (14.39) sono efficaci solo se i tempi $T_{\pi/2}$, T_π e $T_{2\pi}$ sono molto minori sia di T_1 che di T_2 , in modo che il sistema non sia alterato sensibilmente dal rilassamento durante l'impulso. Nel seguito supporremo che questa condizione sia verificata. Se $\mathbf{B}_0$ fosse perfettamente omogeneo su tutto il volume occupato dal campione, dovremmo osservare nella bobina un segnale smorzato esponenzialmente con costante di tempo pari a T_2 per effetto del rilassamento trasversale, come in Fig. 14.6. In realtà, si osserva praticamente sempre un decadimento molto più rapido, caratterizzato da una costante di tempo $T_2^* \ll T_2$. Questo è dovuto al fatto che, in

generale, $\mathbf{B}_0$ non è uniforme su tutto il campione. La disomogeneità di $\mathbf{B}_0$ fa sì che la precessione dei singoli spin avvenga a frequenza non uniforme, ma con uno spettro di valori centrato attorno a ω_0 , dipendente dal valore effettivo del campo nel sito in cui si trova il singolo spin. Quindi, alla distruzione della magnetizzazione trasversale per effetto del rilassamento, si aggiunge la diminuzione dovuta al progressivo sfasamento degli spin tra loro, che avviene con una certa costante temporale T_2' . Per questo osserviamo la magnetizzazione decadere con una costante temporale effettiva T_2^* data da

$$\frac{1}{T_2^*} = \frac{1}{T_2} + \frac{1}{T_2'}, \quad \text{da cui} \quad T_2^* = \frac{T_2 T_2'}{T_2 + T_2'}, \quad \text{in generale con} \quad T_2' \ll T_2. \quad (14.52)$$

Se la larghezza della curva di distribuzione di B_0 è ΔB , la larghezza della distribuzione delle frequenze di precessione di Larmor è $\Delta\omega = |\gamma|\Delta B$ (mettiamo il valore assoluto perché γ può essere sia positivo che negativo), e la costante di tempo per l'apertura del ventaglio vale

$$T_2' = \frac{1}{\Delta\omega} = \frac{1}{|\gamma|\Delta B}. \quad (14.53)$$

Nonostante la presenza della disuniformità del campo statico, è possibile risalire al valore di T_2 tramite la tecnica dell'*eco di spin*, introdotta da Erwin Hahn e descritta in un suo articolo del 1950. Questa tecnica, descritta schematicamente in Fig. 14.7, è basata sul fatto che, mentre il rilassamento è un processo irreversibile, è reversibile la precessione degli spin attorno al campo magnetico statico. Si inizia con un impulso $\pi/2$, che, nel sistema rotante $x'y'z$ con il campo $\mathbf{B}_1(t)$ allineato lungo x' , ruota i momenti magnetici inizialmente allineati lungo z fino a portarli

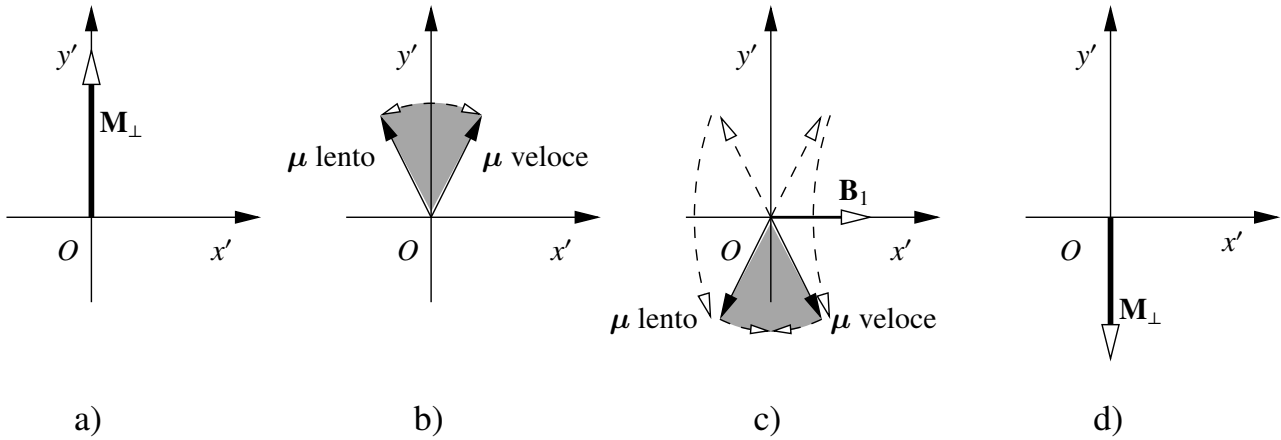


Figura 14.7 Eco di spin. a) Dopo un impulso $\pi/2$ la magnetizzazione, inizialmente orientata lungo z , si trova allineata lungo l'asse y' del sistema rotante alla frequenza di Larmor. b) A causa della non uniformità del campo statico gli spin, sempre osservati nel sistema rotante, cominciano ad aprirsi a ventaglio. c) All'istante t' viene applicato un impulso π , durante il quale gli spin effettuano una rotazione di π radianti attorno al campo $\mathbf{B}_1(t)$, diretto lungo x' . In questo modo, adesso nel sistema del laboratorio gli spin che precedono velocemente inseguono gli spin che precedono lentamente, e, all'istante $2t'$, la magnetizzazione si ricompone, come in d).

ad allinearsi lungo y' . Quindi, dopo l'impulso, la magnetizzazione complessiva $\mathbf{M}$ è diretta lungo y' come in Fig. 14.7.a). Avremmo quindi $M_z = 0$ e $\mathbf{M} = \mathbf{M}_\perp$. A questo punto, se $\mathbf{B}_0$ fosse perfettamente uniforme, nel sistema di laboratorio $\mathbf{M}_\perp$ comincerebbe a precedere attorno a z alla frequenza di Larmor, mentre nel sistema rotante resterebbe costantemente allineata lungo y' . Durante la precessione il rilassamento trasversale smorzerebbe progressivamente $\mathbf{M}_\perp$ e quello longitudinale ricostruirebbe M_z .

Ma come abbiamo visto, in realtà i singoli spin precedono a velocità angolari leggermente diverse a causa della non uniformità di $\mathbf{B}_0$, aprendosi gradatamente a ventaglio. Questa apertura a ventaglio è rappresentata in Fig. 14.7.b): nel sistema del laboratorio l'asse y' del sistema rotante alla frequenza di Larmor media ω_0 lascia indietro gli spin più lenti, mentre sono gli spin più veloci a lasciare indietro l'asse y' . Man mano che il “ventaglio” si apre, la magnetizzazione trasversale complessiva, data dalla somma vettoriale dei singoli momenti magnetici, diminuisce. Vediamo adesso cosa succede se all'istante $t = t'$, con $t' \gg T_\pi$, viene applicato al campione un impulso π . Durante questo impulso gli spin ruotano di 180° attorno al campo $\mathbf{B}_1(t)$, diretto lungo l'asse x' , come in Fig. 14.7.c). In questo modo dopo l'impulso gli spin lenti si trovano, nel moto di precessione, davanti agli spin veloci che li inseguono, chiudendo gradatamente il ventaglio, e raggiungendoli all'istante $2t'$. Viene così ricostruita la magnetizzazione trasversale originaria $\mathbf{M}_\perp$, ovviamente smorzata dal simultaneo rilassamento trasversale, come in Fig. 14.7.d). A $t = 2t'$ il ventaglio appena chiuso comincia a riaprirsi, e il procedimento viene iterato applicando un nuovo impulso π a $t = 3t'$, che porta a una ricostituzione di $\mathbf{M}_\perp$ a $t = 4t'$. E così via, come rappresentato in Fig. 14.8. In altre parole, dopo il primo impulso $\pi/2$ a $t = 0$, vengono applicati impulsi π a tempi corrispondenti a multipli dispari di un certo intervallo temporale $t' \gg T_\pi$, quindi, ogni volta che $t = (2n - 1)t'$, con $n \geq 1$ e intero, il ventaglio della distribuzione degli spin comincia a richiudersi e la magnetizzazione trasversale $M_\perp$ aumenta fino a un massimo per $t = 2nt'$,

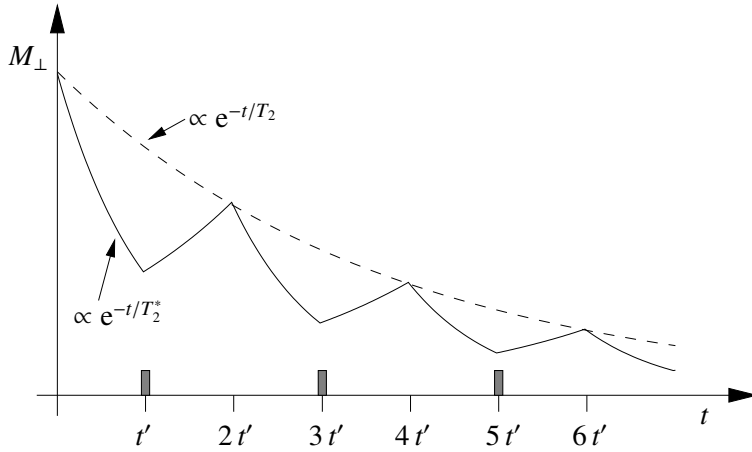


Figura 14.8 Echi di spin. Dopo il primo impulso $\pi/2$ a $t = 0$, vengono applicati impulsi π a multipli dispari di un certo intervallo temporale $t' \gg T_\pi$. Ai multipli pari di t' si ha il riallineamento della magnetizzazione trasversale $\mathbf{M}_\perp$.

quando il ventaglio, dopo essersi richiuso, comincia a riaprirsi. $M_\perp$ così decade fino all'applicazione del successivo impulso π a $t = (2n + 1)t'$. I massimi relativi di $M_\perp$, che si verificano agli istanti $t = 2nt'$, giacciono su un'esponenziale decrescente del tipo

$$M_\perp^{\max} = M_0 e^{-t/T_2}, \quad (14.54)$$

dove T_2 è la costante di rilassamento trasversale. Per ogni $n \geq 0$, nell'intervallo di tempo $2nt' < t < (2n + 1)t'$ la magnetizzazione trasversale diminuisce con la costante di tempo $T_2^* < T_2$, introdotta nella (14.52). Ovviamente, perché la tecnica dell'eco di spin sia applicabile è necessario che sia $T_2^* \ll T_2$. Nel caso della risonanza di spin elettronico ESR (Electron Spin Resonance), detta anche EPR (Electron Paramagnetic Resonance), in presenza di un campo statico di 1 T abbiamo, secondo la (14.12), $\omega_0 = 1.76 \times 10^{11} \text{ rad s}^{-1}$, corrispondente a una frequenza di 28 GHz. Il tempo di rilassamento trasversale è dell'ordine di

$$T_2 \simeq 1 \mu\text{s}, \quad (14.55)$$

mentre, ipotizzando una disomogeneità del campo statico ΔB dell'ordine di 10^{-3} T (10 G), corrispondente a $\Delta\omega = 3 \times 10^7 \text{ rad/s}$, abbiamo per T_2^*

$$T_2^* \simeq T_2' \simeq 10^{-2} \mu\text{s}. \quad (14.56)$$

14.8 Spettroscopia NMR

Nella spettroscopia a risonanza magnetica nucleare viene misurata la frequenza di risonanza di *nuclei attivi* (nuclei con momento magnetico, e quindi spin nucleare, non nullo). I nuclei più usati sono l'idrogeno (^1H), il ^{13}C e il ^{15}N , anche se, in linea di principio, nuclei di isotopi di molti altri elementi possono essere, e spesso vengono, usati.

Le misure possono essere effettuate in due modi: i) a onda continua, e ii) a trasformata di Fourier.

- i) **Spettroscopia NMR a onda continua.** Storicamente, questa è stata la prima a essere usata. In questa tecnica la frequenza del campo oscillante viene mantenuta costante, mentre l'intensità del campo magnetico "statico" B_0 viene incrementato passo passo. Contemporaneamente viene registrato l'assorbimento da parte del campione. Poiché la frequenza di risonanza di un nucleo attivo è proporzionale a B_0 , si ha un aumento dell'assorbimento ogni volta che la frequenza di risonanza di uno dei nuclei presenti nel campione diventa uguale alla frequenza (costante, come abbiamo detto) del campo oscillante.
- ii) **Spettroscopia NMR a trasformata di Fourier.** Un progresso importantissimo della spettroscopia NMR è stato portato dall'idea di utilizzare brevi impulsi con frequenza al centro dello spettro NMR. Come abbiamo visto nel paragrafo 8.6, in particolare dalle (8.63) e (8.64), un impulso quadrato di una data *frequenza portante* ν contiene in realtà tutta una gamma di frequenze, centrata attorno alla frequenza portante stessa, con larghezza inversamente proporzionale alla durata dell'impulso. Poiché l'intero spettro NMR è relativamente stretto, è sufficiente applicare impulsi a radiofrequenza di durata compresa tra il μs e il ms per eccitare l'intero spettro. Quindi, un impulso di questo tipo è simultaneamente in risonanza con tutti i nuclei attivi presenti nel campione. Per quanto visto nel paragrafo 14.3, l'impulso ruota così il vettore magnetizzazione dalla sua posizione di equilibrio lungo $\mathbf{B}_0$. La componente trasversale di $\mathbf{M}$ precede così attorno $\mathbf{B}_0$ alle frequenze di risonanza dei vari nuclei, generando un segnale che induce una forza elettromotrice nella bobina di misura. Una trasformata di Fourier del segnale ci rivela le frequenze, e quindi i nuclei, presenti nel campione. Il metodo quindi richiede l'elaborazione del segnale al computer tramite metodi di trasformata di Fourier veloce.

L'informazione che si può trarre dall'NMR è molto maggiore di quanto potrebbe sembrare dalla trattazione appena sviluppata. Mentre da quanto detto sopra potrebbe sembrare che tutti i nuclei dello stesso tipo risuonino alla stessa frequenza, in realtà la situazione è più complessa. Ogni nucleo, infatti, è circondato da elettroni con i quali il suo momento di dipolo magnetico interagisce. Questa interazione è dovuta alla presenza di campi magnetici originati sia dal moto orbitale degli elettroni che dai momenti magnetici associati agli spin elettronici. Il "moto orbitale" degli elettroni produce una "schermatura" del nucleo, che fa sì che il campo magnetico visto dal nucleo si minore del campo statico applicato B_0 . Questo comporta una diminuzione del valore della frequenza del campo oscillante a cui si osserva la risonanza, detta *spostamento chimico* (chemical shift). In questo modo i segnali di NMR sono sensibili, per esempio, alla molecola in cui il nucleo si trova. A meno che la simmetria delle funzioni d'onda molecolari sia molto vicina a una simmetria sferica (inducendo uno spostamento chimico isotropo), l'effetto di schermo dipenderà dall'orientazione della molecola rispetto al campo esterno $\mathbf{B}_0$.

Il caso più noto di chemical shift si osserva nell'alcol etilico $\text{CH}_3\text{CH}_2\text{OH}$, per il quale si notano tre righe diverse corrispondenti alla risonanza del protone, con intensità nei rapporti 3:2:1. Le frequenze

di queste tre righe corrispondono ai diversi shift chimici subiti dai tre “tipi” di protoni: quelli del gruppo CH_3 , del gruppo CH_2 e del gruppo OH .

Un'applicazione importante della NMR è la tomografia a risonanza magnetica (Magnetic Resonance Imaging, MRI, oppure Magnetic Resonance Tomography, MRT). In questa tecnica viene usato quasi esclusivamente l'idrogeno come fonte di segnale. L'immagine tridimensionale viene costruita usando gradienti lineari del campo statico, in modo che la frequenza a cui si osserva la risonanza dipenda dalla posizione del protone. Ovviamente il segnale deve essere elaborato da un computer per ricostruire l'immagine.

14.9 Spettroscopia EPR

La Spettroscopia EPR viene impiegata per individuare e analizzare specie chimiche contenenti uno o più elettroni non accoppiati (chiamate specie paramagnetiche). Queste specie includono: radicali liberi, ioni di metalli di transizione, difetti in cristalli, molecole in stato elettronico di tripletto fondamentale (per esempio l'ossigeno molecolare O_2) o stato di tripletto indotto per fotoeccitazione. A differenza dalla spettroscopia NMR, qui sono gli spin elettronici a essere eccitati dal campo oscillante, al posto degli spin nucleari.

Il primo a osservare il fenomeno della risonanza paramagnetica elettronica è stato il fisico russo Evgenij Zavojskij (Евгений Константинович Завойский) nel 1944: egli notò che un cristallo di CuCl_2 esposto a un campo magnetico statico di 4 mT assorbiva radiazione elettromagnetica a 133 MHz.

La maggioranza delle misure EPR viene effettuata in campi magnetici di circa 0.35 T, che corrispondono a frequenze di risonanza di spin nella regione delle microonde, tra 9 e 10 GHz. Anche qui, come nel caso dell'NMR, si potrebbe pensare che, osservando il segnale di risonanza dello spin degli elettroni, lo spettro EPR consista di un'unica riga. In realtà il momento magnetico di un elettrone non accoppiato interagisce, per esempio, con i momenti magnetici dei nuclei circostanti. Questo porta a nuovi livelli energetici e a spettri a molte righe. In questo caso, la separazione delle righe dello spettro EPR indica il grado di interazione tra l'elettrone non accoppiato e i nuclei “perturbatori”. La costante di accoppiamento iperfine di un nucleo è legata alla separazione tra le righe.

I due meccanismi più comuni con cui e elettroni e nuclei interagiscono tra loro sono l'interazione di contatto di Fermi e l'interazione dipolo-dipolo.

Appendice A

Derivazione alternativa dell'equazione di Fokker-Planck

Qui forniamo derivazione alternativa dell'equazione di Fokker-Planck rispetto a quella vista nel paragrafo 7.4. Questo per evitare l'“ineleganza” matematica dello sviluppo in Δx visto nella (7.56), in cui si eravamo partiti da un termine del tipo $x + \Delta x - \Delta x$. La maggiore eleganza matematica si paga però con una dimostrazione più lunga e forse meno intuitiva.

Cominciamo considerando una funzione $f(x)$ della variabile stocastica x , che sia una funzione reale liscia e a supporto compatto, ma per il resto arbitraria. Cioè $f(x)$ è arbitraria sotto le condizioni di essere continua e infinitamente derivabile, e di tendere a zero con tutte le sue derivate per $x \rightarrow \pm\infty$, in modo da essere integrabile tra $-\infty$ e $+\infty$. Se all'inizio il valore di x era x_0 la nostra funzione valeva $f(x_0)$, mentre all'istante t il suo valor medio sarà

$$\bar{f}(t, x_0) = \int_{-\infty}^{+\infty} f(x) P(x, t|x_0) dx. \quad (\text{A.1})$$

La derivata temporale del valor medio di f sarà

$$\frac{\partial \bar{f}(t, x_0)}{\partial t} = \int_{-\infty}^{+\infty} f(x) \frac{\partial P(x, t|x_0)}{\partial t} dx. \quad (\text{A.2})$$

Scriviamo la derivata all'interno dell'integrale della (A.2) come limite di rapporto incrementale

$$\frac{\partial P(x, t|x_0)}{\partial t} = \lim_{\Delta t \rightarrow 0} \frac{P(x, t + \Delta t|x_0) - P(x, t|x_0)}{\Delta t}, \quad (\text{A.3})$$

e portiamo il limite, che non coinvolge la variabile di integrazione, fuori dall'integrale stesso

$$\int_{-\infty}^{+\infty} f(x) \frac{\partial P(x, t|x_0)}{\partial t} dx = \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \int_{-\infty}^{+\infty} dx f(x) [P(x, t + \Delta t|x_0) - P(x, t|x_0)]. \quad (\text{A.4})$$

Se applichiamo l'equazione di Smoluchowski, o Chapman-Kolmogorov, (7.21) a $P(x, t + \Delta t|x_0)$, la (A.4) può essere riscritta

$$\begin{aligned} & \int_{-\infty}^{+\infty} f(x) \frac{\partial P(x, t|x_0)}{\partial t} dx = \\ &= \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \left[\int_{-\infty}^{+\infty} dx f(x) \int_{-\infty}^{+\infty} dx' P(x, \Delta t|x') P(x', t|x_0) - \int_{-\infty}^{+\infty} f(x) P(x, t|x_0) dx \right]. \end{aligned} \quad (\text{A.5})$$

Nell'ultimo integrale tra parentesi quadre al secondo membro della (A.5) possiamo chiamare x' , anziché x , la variabile di integrazione. Possiamo poi moltiplicare questo integrale per l'integrale $\int_{-\infty}^{+\infty} P(x, \Delta t | x') dx$, che è uguale a 1 per la condizione di normalizzazione della densità probabilità condizionata. Con queste modifiche la (A.5) diventa

$$\begin{aligned} \int_{-\infty}^{+\infty} f(x) \frac{\partial P(x, t | x_0)}{\partial t} dx &= \\ &= \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \left[\int_{-\infty}^{+\infty} dx f(x) \int_{-\infty}^{+\infty} dx' P(x, \Delta t | x') P(x', t | x_0) - \right. \\ &\quad \left. - \int_{-\infty}^{+\infty} dx' f(x') P(x', t | x_0) \int_{-\infty}^{+\infty} dx P(x, \Delta t | x') \right]. \quad (\text{A.6}) \end{aligned}$$

Nel primo doppio integrale a secondo membro possiamo scambiare l'ordine di integrazione, e nel secondo possiamo spostare $f(x')$ dall'integrale esterno all'integrale interno, ottenendo

$$\begin{aligned} \int_{-\infty}^{+\infty} f(x) \frac{\partial P(x, t | x_0)}{\partial t} dx &= \\ &= \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \left[\int_{-\infty}^{+\infty} dx' P(x', t | x_0) \int_{-\infty}^{+\infty} dx f(x) P(x, \Delta t | x') - \right. \\ &\quad \left. - \int_{-\infty}^{+\infty} dx' P(x', t | x_0) \int_{-\infty}^{+\infty} dx f(x') P(x, \Delta t | x') \right] \\ &= \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \int_{-\infty}^{+\infty} dx' P(x', t | x_0) \int_{-\infty}^{+\infty} dx P(x, \Delta t | x') [f(x) - f(x')] . \quad (\text{A.7}) \end{aligned}$$

La condizione che x vari nel tempo con continuità impone che, al limite $\Delta t \rightarrow 0$, la densità di probabilità condizionata $P(x, \Delta t | x')$ tenda rapidamente a zero per x che si allontana da x' . All'integrale interno, in dx , dell'ultimo membro della (A.7) contribuiranno quindi solo i valori di x molto vicini a x' , e la $f(x)$ può essere sviluppata in serie di Taylor attorno al valore $f(x')$:

$$\begin{aligned} \int_{-\infty}^{+\infty} f(x) \frac{\partial P(x, t | x_0)}{\partial t} dx &= \\ &= \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \int_{-\infty}^{+\infty} dx' P(x', t | x_0) \int_{-\infty}^{+\infty} dx P(x, \Delta t | x') \sum_{n=1}^{\infty} \frac{\partial^n f(x)}{\partial x^n} \Big|_{x=x'} \frac{(x-x')^n}{n!} \\ &= \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \int_{-\infty}^{+\infty} dx' P(x', t | x_0) \sum_{n=1}^{\infty} \frac{\partial^n f(x')}{\partial x'^n} \int_{-\infty}^{+\infty} dx P(x, \Delta t | x') \frac{(x-x')^n}{n!}, \quad (\text{A.8}) \end{aligned}$$

dove le derivate di f sono state portate fuori dall'integrale interno perché, essendo calcolate nel punto x' , sono indipendenti dalla variabile di integrazione x . Inoltre il termine di ordine zero dello sviluppo di $f(x)$, pari a $f(x')$, si è cancellato con il termine $-f(x')$ all'interno della parentesi quadra. A questo punto ricordiamo i *coefficienti di Kramers-Moyal* di ordine n della variabile x' , già visti nella (7.61),

$$D_n(x') = \lim_{\Delta t \rightarrow 0} \frac{1}{\Delta t} \int_{-\infty}^{+\infty} dx (x - x')^n P(x, \Delta t | x') = \lim_{\Delta t \rightarrow 0} \frac{\langle (x - x')^n \rangle}{\Delta t}. \quad (\text{A.9})$$

Il coefficiente D_1 è la velocità con cui il valor medio della variabile stocastica x si allontana dal valore iniziale x' (deriva, o *drift*), D_2 è la velocità con cui cresce la varianza di x , e così via. Introducendo la (A.9) nella (A.8) abbiamo

$$\begin{aligned} \int_{-\infty}^{+\infty} f(x) \frac{\partial P(x, t|x_0)}{\partial t} dx &= \int_{-\infty}^{+\infty} dx' P(x', t|x_0) \sum_{n=1}^{\infty} \frac{D_n(x')}{n!} \frac{\partial^n f(x')}{\partial x'^n} \\ &= \sum_{n=1}^{\infty} \int_{-\infty}^{+\infty} dx' P(x', t|x_0) \frac{D_n(x')}{n!} \frac{\partial^n f(x')}{\partial x'^n}. \end{aligned} \quad (\text{A.10})$$

Gli addendi della sommatoria all'ultimo membro della (A.10) sono integrabili per parti, ognuno tante volte quanto è l'ordine della derivata di $f(x')$ che contiene. Dalle proprietà di $f(x')$, e in particolare dal fatto che le sue derivate di ogni ordine tendono a zero per $x' \rightarrow \pm\infty$, segue che n integrazioni per parti del singolo addendo trasferiscono l'operazione di derivazione n -sima rispetto a x' dalla $f(x')$ al resto dell'integrando

$$\int_{-\infty}^{+\infty} dx' P(x', t|x_0) \frac{D_n(x')}{n!} \frac{\partial^n f(x')}{\partial x'^n} = \frac{(-1)^n}{n!} \int_{-\infty}^{+\infty} dx' f(x') \frac{\partial^n}{\partial x'^n} [D_n(x') P(x', t|x_0)]. \quad (\text{A.11})$$

Sostituendo la (A.11) negli addendi della sommatoria all'ultimo membro della (A.10) otteniamo

$$\int_{-\infty}^{+\infty} f(x) \frac{\partial P(x, t|x_0)}{\partial t} dx = \sum_{n=1}^{\infty} \frac{(-1)^n}{n!} \int_{-\infty}^{+\infty} dx' f(x') \frac{\partial^n}{\partial x'^n} [D_n(x') P(x', t|x_0)]. \quad (\text{A.12})$$

Chiamando x anziché x' la variabile di integrazione al secondo membro, e sottraendo poi il secondo membro dal primo, l'equazione diventa

$$\int_{-\infty}^{+\infty} dx f(x) \left\{ \frac{\partial P(x, t|x_0)}{\partial t} - \sum_{n=1}^{\infty} \frac{(-1)^n}{n!} \frac{\partial^n}{\partial x^n} [D_n(x) P(x, t|x_0)] \right\} = 0, \quad (\text{A.13})$$

e questa, data l'arbitrarietà di $f(x)$, implica che sia

$$\frac{\partial P(x, t|x_0)}{\partial t} = \sum_{n=1}^{\infty} \frac{(-1)^n}{n!} \frac{\partial^n}{\partial x^n} [D_n(x) P(x, t|x_0)], \quad (\text{A.14})$$

che è, come sappiamo, l'equazione di Kramers-Moyal. Questa si riduce all'equazione di Fokker-Planck se possiamo trascurare i coefficienti $D_n(x)$ con $n > 2$.

Appendice B

Teorema centrale del limite

Il *teorema centrale del limite* afferma che, se S_n è la somma di n variabili aleatorie indipendenti, $S_n = \sum_{j=1}^n x_j$, sotto opportune condizioni la funzione di distribuzione di S_n tende ad una gaussiana al crescere di n . Condizioni frequentemente incontrate, e sotto cui il teorema centrale del limite è valido, sono: 1) che le singole x_j , pur indipendenti tra loro, siano governate da un'unica funzione di distribuzione normalizzata di probabilità $p(x)$, e 2) che esistano, e siano finiti, sia il valore di aspettazione μ che la varianza σ^2 , con

$$\mu = \int_{-\infty}^{+\infty} x_j p(x_j) dx_j \quad \text{e} \quad \sigma^2 = \int_{-\infty}^{+\infty} (x_j - \mu)^2 p(x_j) dx_j, \quad (\text{B.1})$$

con μ e σ^2 ovviamente indipendenti da j . In queste condizioni il valore di aspettazione di S_n tende a $n\mu$ e la sua varianza a $n\sigma^2$, indipendentemente dalla forma analitica di $p(x)$.

Piuttosto che lavorare direttamente su S_n conviene definire una variabile aleatoria *standardizzata* Z_n come

$$Z_n = \frac{S_n - n\mu}{\sigma \sqrt{n}} = \frac{\sum_{j=1}^n x_j - n\mu}{\sigma \sqrt{n}}, \quad (\text{B.2})$$

la cui funzione di distribuzione, in condizioni di validità del teorema centrale del limite, converge ad una gaussiana centrata sullo zero con varianza 1

$$\lim_{n \rightarrow \infty} W(Z_n) = \frac{1}{\sqrt{2\pi}} e^{-Z_n^2/2}. \quad (\text{B.3})$$

Segue una delle possibili dimostrazioni.

Consideriamo una variabile aleatoria x , descritta da una densità di probabilità $p(x)$. La probabilità di avere un valore nell'intervallo $(x, x + dx)$ vale, cioè, $p(x) dx$. Chiamiamo μ il valor medio di x e σ^2 la sua varianza, in formule

$$\mu = \bar{x} = \int_{-\infty}^{+\infty} x p(x) dx \quad \text{e} \quad \sigma^2 = \overline{(x - \mu)^2} = \int_{-\infty}^{+\infty} (x - \mu)^2 p(x) dx, \quad (\text{B.4})$$

dove con le barre abbiamo indicato i valori medi aspettati su un grande numero di campionamenti. Possiamo definire una nuova variabile y legata ad x dalla relazione

$$y = \frac{x - \mu}{\sigma} \quad \text{in modo che} \quad \bar{y} = 0 \quad \text{e} \quad \overline{y^2} = 1. \quad (\text{B.5})$$

Essendo costruita partendo da una variabile aleatoria, anche la y sarà una variabile aleatoria. La densità di probabilità $w(y)$, tale che $w(y) dy$ sia la probabilità che y assuma un valore compreso tra y e $y + dy$, sarà

$$w(y) = \sigma p(\sigma y + \mu), \quad \text{con} \quad \int_{-\infty}^{+\infty} w(y) dy = 1, \quad \int_{-\infty}^{+\infty} y w(y) dy = 0 \quad \text{e} \quad \int_{-\infty}^{+\infty} y^2 w(y) dy = 1. \quad (\text{B.6})$$

Per una variabile aleatoria y descritta da una densità di probabilità $w(y)$, si definisce la *funzione caratteristica* $\varphi_y(t)$ come il valore di aspettazione di $e^{-2\pi i t y}$, dove $t \in \mathbb{R}$ è l'argomento della funzione caratteristica stessa. In altre parole, $\varphi_y(t)$ è la trasformata di Fourier di $w(y)$

$$\varphi_y(t) = \int_{-\infty}^{+\infty} e^{-2\pi i t y} w(y) dy. \quad (\text{B.7})$$

Matematicamente, la funzione caratteristica gode di alcune proprietà che la rendono spesso vantaggiosa rispetto alla densità di probabilità. Per esempio, supponiamo di fare coppie di campionamenti, cui corrispondono coppie di valori x_1 e x_2 , con i corrispondenti y_1 ed y_2 . Definiamo $z = y_1 + y_2$ e cerchiamo la densità di probabilità $W(z)$ che governa la distribuzione di z . Avremo

$$W(z) = \int_{-\infty}^{+\infty} w(y_1) w(z - y_1) dy_1, \quad (\text{B.8})$$

cioè $W(z)$ è il prodotto di convoluzione di $w(y)$ con sé stessa. Se definiamo una variabile aleatoria $z_2 = \frac{1}{\sqrt{2}}(y_1 + y_2)$, per la sua distribuzione di probabilità avremo

$$W(z_2) = C \int_{-\infty}^{+\infty} w(y_1) w(\sqrt{2}z_2 - y_1) dy_1, \quad (\text{B.9})$$

con C costante da determinare dalla condizione di normalizzazione su $W(z_2)$. Deve essere

$$\begin{aligned} 1 &= \int_{-\infty}^{+\infty} W(z_2) dz_2 = C \int_{-\infty}^{+\infty} dz_2 \int_{-\infty}^{+\infty} w(y_1) w(\sqrt{2}z_2 - y_1) dy_1 \\ &= C \int_{-\infty}^{+\infty} dy_1 w(y_1) \int_{-\infty}^{+\infty} dz_2 w(\sqrt{2}z_2 - y_1). \end{aligned} \quad (\text{B.10})$$

Nel secondo integrale possiamo porre $\xi = \sqrt{2}z_2 - y_1$, con $d\xi = \sqrt{2}dz_2$ perché y_1 viene mantenuto costante durante l'integrazione in dz_2 , ottenendo

$$1 = \frac{C}{\sqrt{2}} \int_{-\infty}^{+\infty} dy_1 w(y_1) \int_{-\infty}^{+\infty} d\xi w(\xi) = \frac{C}{\sqrt{2}} \quad (\text{B.11})$$

per la normalizzazione della distribuzione di probabilità w . Abbiamo quindi $C = \sqrt{2}$ e

$$W(z_2) = \sqrt{2} \int_{-\infty}^{+\infty} w(y_1) w(\sqrt{2}z_2 - y_1) dy_1. \quad (\text{B.12})$$

Per ottenere la funzione caratteristica $\varphi_{z_2}(t)$ corrispondente alla distribuzione W dobbiamo farne la trasformata di Fourier

$$\begin{aligned}
 \varphi_{z_2}(t) &= \tilde{W}(t) = \sqrt{2} \int_{-\infty}^{+\infty} dy_1 \int_{-\infty}^{+\infty} dz_2 e^{-2\pi i t z_2} w(y_1) w(\sqrt{2}z_2 - y_1) \quad \text{sostituendo } \eta = \sqrt{2}z_2 \\
 &= \sqrt{2} \frac{1}{\sqrt{2}} \int_{-\infty}^{+\infty} dy_1 w(y_1) \int_{-\infty}^{+\infty} d\eta e^{-2\pi i t \frac{\eta}{\sqrt{2}}} w(\eta - y_1) \\
 &= \int_{-\infty}^{+\infty} dy_1 w(y_1) \int_{-\infty}^{+\infty} d\eta e^{-2\pi i \frac{t}{\sqrt{2}} \eta} w(\eta - y_1)
 \end{aligned} \tag{B.13}$$

Nel secondo integrale possiamo sostituire $\xi = \eta - y_1$, con $d\xi = d\eta$ (y_1 è considerato fisso nell'integrazione in $d\eta$) e $\eta = \xi + y_1$, ottenendo

$$\begin{aligned}
 \varphi_{z_2}(t) &= \int_{-\infty}^{+\infty} dy_1 w(y_1) \int_{-\infty}^{+\infty} d\xi e^{-2\pi i \frac{t}{\sqrt{2}} (\xi + y_1)} w(\xi) \\
 &= \int_{-\infty}^{+\infty} dy_1 e^{-2\pi i \frac{t}{\sqrt{2}} y_1} w(y_1) \int_{-\infty}^{+\infty} d\xi e^{-2\pi i \frac{t}{\sqrt{2}} \xi} w(\xi) \\
 &= \left[\tilde{w}\left(\frac{t}{\sqrt{2}}\right) \right]^2 = \left[\varphi_y\left(\frac{t}{\sqrt{2}}\right) \right]^2
 \end{aligned} \tag{B.14}$$

Questa è in realtà una conseguenza del *teorema di convoluzione*, secondo il quale la trasformata di Fourier del prodotto di convoluzione di due funzioni è uguale al prodotto delle trasformate di Fourier delle due funzioni. Più in generale, con procedimento analogo, si può dimostrare che se definiamo la variabile aleatoria

$$z_n = \frac{\sum_{i=1}^n x_i - n\mu}{\sigma \sqrt{n}} = \sum_{i=1}^n \frac{y_i}{\sqrt{n}}, \tag{B.15}$$

supponendo di fare un numero grande n di campionamenti della variabile, ottenendo i valori $x_1, x_2 \dots x_n$ (e i corrispondenti $y_1, y_2 \dots y_n$), avremo per la funzione caratteristica φ_{z_n} di z_n

$$\varphi_{z_n}(t) = \left[\varphi_y\left(\frac{t}{\sqrt{n}}\right) \right]^n. \tag{B.16}$$

Adesso sviluppiamo la funzione caratteristica $\varphi_y(t)$ in serie di Taylor attorno a $t = 0$. Abbiamo per le (B.6)

$$\begin{aligned}
 \varphi_y(0) &= \int_{-\infty}^{+\infty} w(y) dy = 1, \quad \frac{\partial}{\partial t} \varphi_y(t) \Big|_{t=0} = -2\pi i \int_{-\infty}^{+\infty} y w(y) dy = 0 \\
 \frac{\partial^2}{\partial t^2} \varphi_y(t) \Big|_{t=0} &= -4\pi^2 \int_{-\infty}^{+\infty} y^2 w(y) dy = -4\pi^2,
 \end{aligned} \tag{B.17}$$

quindi possiamo scrivere

$$\varphi_y(t) \simeq 1 - 2\pi^2 t^2 + O(t^2). \tag{B.18}$$

Abbiamo così l'approssimazione

$$\varphi_{z_n}(t) = \left[\varphi_y\left(\frac{t}{\sqrt{n}}\right) \right]^n \simeq \left(1 - \frac{2\pi^2 t^2}{n} \right)^n, \quad (\text{B.19})$$

di cui vogliamo calcolare il limite per $n \rightarrow \infty$. E' più facile calcolare il limite del logaritmo della funzione, anziché il limite della funzione stessa. Abbiamo

$$\lim_{n \rightarrow \infty} \ln \left(1 - \frac{2\pi^2 t^2}{n} \right)^n = \lim_{n \rightarrow \infty} n \ln \left(1 - \frac{2\pi^2 t^2}{n} \right). \quad (\text{B.20})$$

Usando lo sviluppo in serie di Taylor, valido per $x < 1$,

$$\ln(1-x) = \sum_{n=1}^{\infty} \frac{(-x)^n}{n}, \quad \text{abbiamo} \quad \ln \left(1 - \frac{2\pi^2 t^2}{n} \right) = \sum_{m=1}^{\infty} \frac{(-2\pi^2 t^2)^m}{m}, \quad (\text{B.21})$$

che, introdotto nella (B.20), ci dà

$$\begin{aligned} \lim_{n \rightarrow \infty} \ln \left(1 - \frac{2\pi^2 t^2}{n} \right)^n &= \lim_{n \rightarrow \infty} n \left[-\frac{2\pi^2 t^2}{n} + \frac{1}{2} \frac{(2\pi^2 t^2)^2}{n^2} - \dots \right] \\ &= \lim_{n \rightarrow \infty} \left[-2\pi^2 t^2 + \frac{1}{2} \frac{(2\pi^2 t^2)^2}{n} - \dots \right]. \end{aligned} \quad (\text{B.22})$$

Al limite $n \rightarrow \infty$ scompaiono tutti i termini con potenze di n maggiori di zero al denominatore, lasciando

$$\lim_{n \rightarrow \infty} \ln \left(1 - \frac{2\pi^2 t^2}{n} \right)^n = -2\pi^2 t^2, \quad \text{da cui} \quad \lim_{n \rightarrow \infty} \varphi_{z_n}(t) = e^{-2\pi^2 t^2}. \quad (\text{B.23})$$

Per passare alla distribuzione di probabilità $W(z_n)$ dobbiamo fare l'antitrasformata di Fourier della sua funzione caratteristica $\varphi_{z_n}(t)$. La trasformata di Fourier di una gaussiana è una gaussiana, e

$$\text{se } f(x) = e^{-\alpha x^2} \quad \text{abbiamo} \quad \tilde{f}(\xi) = \sqrt{\frac{\pi}{\alpha}} e^{-\beta \xi^2}, \quad \text{dove } \beta = \frac{\pi^2}{\alpha}. \quad (\text{B.24})$$

Abbiamo quindi

$$\lim_{n \rightarrow \infty} W(z_n) = A e^{-z_n^2/2} = \frac{1}{\sqrt{2\pi}} e^{-z_n^2/2} \quad (\text{B.25})$$

dove il parametro A è stato determinato dalla condizione di normalizzazione della distribuzione di probabilità $W(z_n)$. Questo è, appunto, il teorema centrale del limite.

Appendice C

Progressione geometrica

Per ricavare la (9.115), cominciamo considerando la somma S_n dei termini tra l'ordine 0 e l'ordine n , compresi, di una progressione geometrica di ragione q , con q numero complesso qualunque, cioè

$$S_n = \sum_{r=0}^n q^r = 1 + q + q^2 + \cdots + q^n. \quad (\text{C.1})$$

Moltiplichiamo poi S_n per q , ottenendo

$$qS_n = q + q^2 + q^3 + \cdots + q^{n+1}, \quad (\text{C.2})$$

e se togliamo la (C.1) dalla (C.2) ci resta

$$S_n(q - 1) = q^{n+1} - 1, \quad (\text{C.3})$$

perché i primi n addendi di qS_n si cancellano con gli ultimi n addendi di S_n . Quindi possiamo ricavare il valore di S_n come

$$S_n = \sum_{r=0}^n q^r = \frac{q^{n+1} - 1}{q - 1}. \quad (\text{C.4})$$

Adesso calcoliamo la somma comprendente indici negativi

$$\begin{aligned} \sum_{r=-n}^n q^r &= \frac{1}{q^n} \sum_{r=0}^{2n} q^r = \frac{1}{q^n} \frac{q^{2n+1} - 1}{q - 1} = \frac{1}{q^n} \frac{q^{(2n+1)/2} q^{(2n+1)/2} - 1}{q^{1/2} q^{1/2} - 1} \\ &= \frac{1}{q^n} \frac{q^{(2n+1)/2} [q^{(2n+1)/2} - q^{-(2n+1)/2}]}{q^{1/2} (q^{1/2} - q^{-1/2})} = \frac{1}{q^n} q^n \frac{q^{(2n+1)/2} - q^{-(2n+1)/2}}{q^{1/2} - q^{-1/2}} \\ &= \frac{q^{(2n+1)/2} - q^{-(2n+1)/2}}{q^{1/2} - q^{-1/2}}. \end{aligned} \quad (\text{C.5})$$

Infine, prendiamo il caso particolare $q = e^{i\alpha}$ e otteniamo

$$\sum_{r=-n}^n e^{ir\alpha} = \frac{e^{i(2n+1)\alpha/2} - e^{-i(2n+1)\alpha/2}}{e^{i\alpha/2} - e^{-i\alpha/2}} = \frac{\sin[(2n+1)\alpha/2]}{\sin(\alpha/2)}, \quad (\text{C.6})$$

che porta alla (9.115).

Appendice D

Emissione spontanea e seconda quantizzazione

Nel paragrafo 10.4 abbiamo ricavato i coefficienti di assorbimento e di emissione indotta B di Einstein dalla trattazione semiclassica delle transizioni di dipolo elettrico vista nel paragrafo 10.3. Nello stesso paragrafo il coefficiente di emissione spontanea A è invece stato ottenuto partendo da B , e usando le considerazioni euristiche di Einstein sulla radiazione di corpo nero. Questo perché il coefficiente A non può essere ricavato da una trattazione semiclassica. In questa appendice accenneremo al calcolo di A a partire dalla seconda quantizzazione.

La seconda quantizzazione, introdotta da Dirac nel 1927, tratta il campo elettromagnetico come un insieme di bosoni (fotoni) non interagenti tra loro. I modi del campo, etichettati da un vettore d'onda $\mathbf{k}$ e dalla polarizzazione μ , costituiscono gli stati di particella singola $|\mathbf{k}, \mu\rangle$ per i fotoni. Gli stati del campo elettromagnetico sono così degli stati di Fock (5.43) del tipo

$$|\{n_{\mathbf{k},\mu}\}\rangle = |n_{\mathbf{k}_1,\mu_1}\rangle |n_{\mathbf{k}_2,\mu_2}\rangle |n_{\mathbf{k}_3,\mu_3}\rangle \cdots = \prod_{\mathbf{k},\mu} |n_{\mathbf{k},\mu}\rangle, \quad (\text{D.1})$$

dove gli $n_{\mathbf{k},\mu}$ sono i numeri di occupazione degli stati di particella singola, e il prodotto è un prodotto tensore che corre su tutti i possibili stati di particella singola. Un fotone nello stato $|\mathbf{k}, \mu\rangle$ ha energia $\hbar\omega_k$, con $\omega_k = ck$.

La quantizzazione di un singolo modo del campo elettromagnetico nel vuoto è analoga alla quantizzazione di un oscillatore armonico. E' possibile scrivere un operatore di creazione $\hat{a}_{\mathbf{k},\mu}^\dagger$ e un operatore di distruzione $\hat{a}_{\mathbf{k},\mu}$, analoghi agli operatori di salita e discesa dell'oscillatore armonico, tali che valgono le relazioni

$$\hat{a}_{\mathbf{k},\mu} |\mathbf{k}, \mu\rangle = \sqrt{n_{\mathbf{k},\mu}} |\mathbf{k}, \mu - 1\rangle, \quad \hat{a}_{\mathbf{k},\mu}^\dagger |\mathbf{k}, \mu\rangle = \sqrt{n_{\mathbf{k},\mu} + 1} |\mathbf{k}, \mu + 1\rangle. \quad (\text{D.2})$$

L'operatore di distruzione toglie quindi un fotone dal modo, mentre l'operatore di creazione ne aggiunge uno. Se il modo è vuoto abbiamo $\hat{a}_{\mathbf{k},\mu} |0_{\mathbf{k},\mu}\rangle = 0$, come deve essere perché non esistono numeri negativi di fotoni. L'operatore "numero di fotoni" nel singolo modo si scrive $\hat{N}_{\mathbf{k},\mu} = \hat{a}_{\mathbf{k},\mu}^\dagger \hat{a}_{\mathbf{k},\mu}$, infatti possiamo controllare che

$$\hat{N}_{\mathbf{k},\mu} |\mathbf{k}, \mu\rangle = \hat{a}_{\mathbf{k},\mu}^\dagger \hat{a}_{\mathbf{k},\mu} |\mathbf{k}, \mu\rangle = n_{\mathbf{k},\mu} |\mathbf{k}, \mu\rangle. \quad (\text{D.3})$$

L'hamiltoniana relativa al singolo modo si scrive

$$\hat{H}_{\mathbf{k},\mu} = \frac{1}{2} \hbar\omega_k (\hat{a}_{\mathbf{k},\mu}^\dagger \hat{a}_{\mathbf{k},\mu} + \hat{a}_{\mathbf{k},\mu} \hat{a}_{\mathbf{k},\mu}^\dagger), \quad (\text{D.4})$$

i suoi autostati sono gli stati $|\mathbf{k}, \mu\rangle$, e per gli autovalori abbiamo

$$\hat{H}_{\mathbf{k}, \mu} |\mathbf{k}, \mu\rangle = \left(n_{\mathbf{k}, \mu} + \frac{1}{2}\right) \hbar \omega_k |\mathbf{k}, \mu\rangle. \quad (\text{D.5})$$

La quantità $\frac{1}{2} \hbar \omega_k$ è l'energia dello stato fondamentale ($n_{\mathbf{k}, \mu} = 0$) del modo. L'hamiltoniana complessiva del campo elettromagnetico è

$$\hat{H}_{\text{phot}} = \sum_{\mathbf{k}, \mu} \hat{H}_{\mathbf{k}, \mu}, \quad (\text{D.6})$$

i suoi autostati sono gli stati di Fock (D.1), e abbiamo per gli autovalori

$$\hat{H}_{\text{phot}} |\{n_{\mathbf{k}, \mu}\}\rangle = \left[\sum_{\mathbf{k}, \mu} \left(n_{\mathbf{k}, \mu} + \frac{1}{2}\right) \hbar \omega_k \right] |\{n_{\mathbf{k}, \mu}\}\rangle. \quad (\text{D.7})$$

Consideriamo adesso l'interazione del campo elettromagnetico con un atomo che, se isolato, è descritto da un'hamiltoniana $\hat{H}_{\text{at}}$ con autostati numerabili $|\text{at}, m\rangle$ e autovalori U_m (usiamo il simbolo U perché il simbolo E ci servirà per il campo elettrico),

$$\hat{H}_{\text{at}} |\text{at}, m\rangle = U_m |\text{at}, m\rangle. \quad (\text{D.8})$$

Dell'atomo a noi interesseranno solo lo stato fondamentale $|\text{at}, g\rangle$ di energia U_g e uno stato eccitato $|\text{at}, e\rangle$ di energia U_e . L'hamiltoniana complessiva del sistema "atomo più campo elettromagnetico" si scrive

$$\hat{H}_{\text{tot}} = \hat{H}_{\text{at}} + \hat{H}_{\text{phot}} + \hat{H}_{\text{int}}, \quad (\text{D.9})$$

dove l'operatore $\hat{H}_{\text{int}}$ descrive l'interazione atomo-radiazione. Una base per lo spazio di Hilbert del sistema è costituita dagli autostati dell'hamiltoniana del sistema imperturbato $\hat{H}_{\text{at}} + \hat{H}_{\text{phot}}$

$$|m, \{n_{\mathbf{k}, \mu}\}\rangle = |\text{at}, m\rangle |\{n_{\mathbf{k}, \mu}\}\rangle, \quad \text{con autovalori} \quad U_m + \sum_{\mathbf{k}, \mu} \left(n_{\mathbf{k}, \mu} + \frac{1}{2}\right) \hbar \omega_k. \quad (\text{D.10})$$

Il contributo all'hamiltoniana di interazione di un singolo modo del campo con $\mathbf{k}$ parallelo all'asse z , polarizzazione lungo x , e vicino alla risonanza con la transizione atomica, cioè con $\omega_k = \omega_r \simeq (U_e - U_g)/\hbar$, si scrive, seguendo le (10.66) e (10.68), nella forma

$$\hat{H}_{\text{int}}^{(\text{res}, x)} = -\hat{\mathbf{E}}^{(\text{res})} \cdot \hat{\mathbf{P}} = -\hat{E}_x^{(\text{res})} \hat{P}_x \quad (\text{D.11})$$

dove qui, oltre al momento di dipolo atomico, anche il campo elettrico associato al modo risonante, $\hat{\mathbf{E}}^{(\text{res})}$, è un operatore. Questo operatore può essere scritto in termini di operatori di creazione e distruzione nella forma

$$\hat{E}_x^{(\text{res})} = \frac{1}{\sqrt{2}} \sqrt{\frac{\hbar \omega_r}{\epsilon_0 V}} \left(\hat{a}_x^{(\text{res})} e^{ikz} + \hat{a}_x^{(\text{res})\dagger} e^{-ikz} \right), \quad (\text{D.12})$$

dove V è il volume in cui consideriamo confinato il sistema, mentre $\hat{a}_x^{(\text{res})}$ e $\hat{a}_x^{(\text{res})\dagger}$ sono gli operatori di distruzione e creazione del nostro modo risonante con polarizzazione x . È importante notare che gli stati di Fock (D.1) non sono autostati dell'operatore campo elettrico. Come nel paragrafo 10.3

facciamo l'ipotesi che le dimensioni dell'atomo siano molto piccole rispetto alla lunghezza d'onda dei nostri fotoni, in modo che $e^{ikz} \simeq 1$, e approssimiamo la (D.12) con

$$\hat{E}_x^{(\text{res})} = \frac{1}{\sqrt{2}} \sqrt{\frac{\hbar\omega_r}{\varepsilon_0 V}} (\hat{a}_x^{(\text{res})} + \hat{a}_x^{(\text{res})\dagger}). \quad (\text{D.13})$$

Il contributo all'hamiltoniana di interazione diventa

$$\hat{H}_{\text{int}}^{(\text{res},x)} = -\frac{1}{\sqrt{2}} \sqrt{\frac{\hbar\omega_r}{\varepsilon_0 V}} (\hat{a}_x^{(\text{res})} + \hat{a}_x^{(\text{res})\dagger}) \hat{P}_x. \quad (\text{D.14})$$

Si può controllare che $\hat{H}_{\text{int}}^{(\text{res},x)}$ ha elementi diagonali nulli sulla base (D.10), mentre possono essere diversi da zero gli elementi di matrice tra coppie di stati del tipo

$$|g; n_r + 1, \{n'_{\mathbf{k},\mu}\}\rangle \quad \text{e} \quad |e; n_r, \{n'_{\mathbf{k},\mu}\}\rangle, \quad \text{dove} \quad |m; n_r, \{n'_{\mathbf{k},\mu}\}\rangle = |\text{at}, m\rangle |n_r\rangle \prod'_{\mathbf{k},\mu} |n_{\mathbf{k},\mu}\rangle, \quad (\text{D.15})$$

qui n_r è il numero di occupazione del modo risonante, $\{n'_{\mathbf{k},\mu}\}$ è l'insieme dei numeri di occupazione di tutti gli altri modi, uguali modo per modo per $|g; n_r + 1, \{n'_{\mathbf{k},\mu}\}\rangle$ e $|e; n_r, \{n'_{\mathbf{k},\mu}\}\rangle$, e il simbolo $\prod'_{\mathbf{k},\mu}$ indica il prodotto tensore su tutti i modi e tutte le polarizzazioni del campo tranne quello risonante. La differenza di energia tra i due stati vale

$$\Delta U = U_e - (U_g + \hbar\omega_r) = U_e - U_g - \hbar\omega_r, \quad (\text{D.16})$$

e l'elemento fuori diagonale di $\hat{H}_{\text{int}}^{(\text{res},x)}$ vale

$$\langle g; n_r + 1, \{n'_{\mathbf{k},\mu}\} | \hat{H}_{\text{int}}^{(\text{res},x)} | e; n_r, \{n'_{\mathbf{k},\mu}\} \rangle = -\langle \text{at}, g | \hat{P}_x | \text{at}, e \rangle \langle n_r + 1 | \hat{E}_x^{(\text{res})} | n_r \rangle. \quad (\text{D.17})$$

Inserendo la (D.13) nell'elemento di matrice di $\hat{E}_x^{(\text{res})}$ abbiamo

$$\langle n_r + 1 | \hat{E}_x^{(\text{res})} | n_r \rangle = \sqrt{\frac{n_r + 1}{2}} \sqrt{\frac{\hbar\omega_r}{\varepsilon_0 V}}, \quad (\text{D.18})$$

e il modulo quadro del nostro elemento di matrice complessivo vale così

$$\left| \langle g; n_r + 1, \{n'_{\mathbf{k},\mu}\} | \hat{H}_{\text{int}}^{(\text{res},x)} | e; n_r, \{n'_{\mathbf{k},\mu}\} \rangle \right|^2 = \frac{(n_r + 1) \hbar\omega_r}{2\varepsilon_0 V} \left| \langle \text{at}, g | \hat{P}_x | \text{at}, e \rangle \right|^2. \quad (\text{D.19})$$

La probabilità di transizione per unità di tempo W_{ge} dallo stato $|e; n_r, \{n'_{\mathbf{k},\mu}\}\rangle$ allo stato $|g; n_r + 1, \{n'_{\mathbf{k},\mu}\}\rangle$ può essere ottenuta dalle (10.29) e (10.35)

$$\begin{aligned} W_{ge} &= \frac{2\pi}{\hbar} \left| \langle g; n_r + 1, \{n'_{\mathbf{k},\mu}\} | \hat{H}_{\text{int}}^{(\text{res})} | e; n_r, \{n'_{\mathbf{k},\mu}\} \rangle \right|^2 \delta_\tau(U_e - U_g - \hbar\omega_r) \\ &= \frac{\pi}{\hbar} \frac{(n_r + 1) \hbar\omega_r}{\varepsilon_0 V} \left| \langle \text{at}, g | \hat{\mathbf{P}} | \text{at}, e \rangle \right|^2 \delta_\tau(U_e - U_g - \hbar\omega_r). \end{aligned} \quad (\text{D.20})$$

Se il numero di fotoni n_r è molto grande abbiamo $n_r + 1 \simeq n_r$, la quantità $n_r \hbar \omega / V$ non è altro che l'energia della radiazione per unità di volume w , e la (D.20) diventa

$$W_{ge} \simeq \frac{\pi}{\hbar \varepsilon_0} |\langle \text{at}, g | \hat{\mathbf{P}} | \text{at}, e \rangle|^2 w \delta_\tau(U_e - U_g - \hbar \omega_r), \quad (\text{D.21})$$

che coincide con la (10.75) ottenuta dal trattamento semiclassico. Ma l'approssimazione semiclassica non vale per piccoli numeri di fotoni. In particolare, in assenza totale di fotoni abbiamo $w = 0$, quindi è nullo il secondo membro della (D.21) e, per la trattazione semiclassica, lo stato atomico eccitato $|\text{at}, e\rangle$ non dovrebbe decadere. Ma per $n_r = 0$ non è nullo l'elemento di matrice (D.19), e questo spiega l'emissione spontanea da $|\text{at}, e\rangle$.

Per arrivare all'espressione corretta del coefficiente A di decadimento spontaneo dobbiamo ancora notare che, se l'atomo interagisce con radiazione distribuita in modo isotropo, i fotoni possono propagarsi con direzioni e polarizzazioni casuali, e con considerazioni analoghe a quelle che hanno portato alla (10.80), la (D.19) può essere riscritta

$$|\langle g; n_r + 1, \{n'_{\mathbf{k}, \mu}\} | \hat{H}_{\text{int}}^{(\text{res})} | e; n_r, \{n'_{\mathbf{k}, \mu}\} \rangle|^2 = \frac{(n_r + 1) \hbar \omega_r}{6 \varepsilon_0 V} |\langle \text{at}, g | \hat{\mathbf{P}} | \text{at}, e \rangle|^2, \quad (\text{D.22})$$

in termini del modulo quadro dell'elemento di matrice dell'operatore vettoriale $\hat{\mathbf{P}}$ anziché dell'elemento di matrice della sola componente $\hat{P}_x$. Per $n_r = 0$ abbiamo così

$$W_{ge}|_{n_r=0} = \frac{2\pi}{\hbar} \frac{\hbar \omega_r}{6 \varepsilon_0 V} |\langle \text{at}, g | \hat{\mathbf{P}} | \text{at}, e \rangle|^2 \delta_\tau(U_e - U_g - \hbar \omega_r), \quad (\text{D.23})$$

e l'atomo può passare allo stato fondamentale emettendo un fotone. Ma in quale modo viene emesso il fotone? La $\delta_\tau(U_e - U_g - \hbar \omega_r)$ che compare nelle (D.20) e (D.23) ci dà un certo range di frequenze possibili. Come abbiamo visto nella (8.81) il numero di celle dello spazio delle fasi con quantità di moto compresa tra p e $p + \Delta p$, al limite per grandi volumi spaziali V , è

$$g = \frac{V 4\pi p^2 \Delta p}{h^3} \quad (\text{D.24})$$

da cui, ricordando che per un fotone è $p = \hbar \omega / c$, e che un fotone ha due stati di polarizzazione possibili, otteniamo per il numero Δm di modi di frequenza compresa tra ω e $\omega + \Delta \omega$

$$\Delta m = \frac{V \omega^2 \Delta \omega}{\pi^2 c^3}, \quad (\text{D.25})$$

cui corrisponde una densità di stati di particella singola per unità di energia

$$\rho = \frac{\Delta m}{\Delta U} = \frac{\Delta m}{\hbar \Delta \omega} = \frac{V \omega^2}{\pi^2 \hbar c^3}. \quad (\text{D.26})$$

A questo punto possiamo applicare la regola d'oro di Fermi (10.45) per ottenere

$$\begin{aligned} W_{ge} &= \frac{2\pi}{\hbar} |\langle \text{fin} | \hat{H}_{\text{int}} | \text{in} \rangle|^2 \rho = \frac{2\pi}{\hbar} \frac{\hbar \omega_r}{6 \varepsilon_0 V} |\langle \text{at}, g | \hat{\mathbf{P}} | \text{at}, e \rangle|^2 \frac{V \omega_r^2}{\pi^2 \hbar c^3} \\ &= \frac{\omega_r^3}{3\pi \hbar \varepsilon_0 c^3} |\langle \text{at}, g | \hat{\mathbf{P}} | \text{at}, e \rangle|^2, \end{aligned} \quad (\text{D.27})$$

che coincide con il valore di A_{ge} dato dalla (10.84).

Appendice E

Carica in campo elettromagnetico

L'hamiltoniana che compare nell'equazione quantistica (12.55) è ottenuta dall'hamiltoniana classica per un particella di massa m e carica q , interagente con un campo elettromagnetico di potenziale elettrico V e potenziale magnetico vettore $\mathbf{A}$. Partiamo cioè dalla

$$H(\mathbf{p}, \mathbf{r}) = \frac{1}{2m}(\mathbf{p} - q\mathbf{A})^2 + qV, \quad (\text{E.1})$$

ed effettuiamo la sostituzione standard $\mathbf{p} \rightarrow -i\hbar\nabla$. A sua volta la (E.1) è ottenuta modificando prima la lagrangiana, poi, di conseguenza, l'hamiltoniana della particella in modo da tener conto della forza dovuta al campo magnetico, $q\mathbf{v} \times \mathbf{B}$, che non fa lavoro, dipende dalla velocità e non può essere semplicemente derivata da un potenziale. Qui non ricaveremo la (E.1), ma ci limiteremo a far vedere che da essa si ottiene la forza di Lorentz tramite le equazioni di Hamilton. Cominciamo riscrivendo l'equazione nella forma

$$H(\mathbf{p}, \mathbf{r}) = \frac{1}{2m}[(p_x - qA_x)^2 + (p_y - qA_y)^2 + (p_z - qA_z)^2] + qV \quad (\text{E.2})$$

e calcoliamo $\dot{x}$ dalla prima delle equazioni canoniche (2.3)

$$\dot{x} = \frac{\partial H}{\partial p_x} = \frac{1}{m}(p_x - qA_x), \quad \text{da cui} \quad m\dot{x} = p_x - qA_x. \quad (\text{E.3})$$

Qui vediamo che la componente x dell'impulso (quantità di moto) canonico, p_x non è più uguale alla componente x dell'impulso cinetico, $m\dot{x}$, ma abbiamo $p_x = m\dot{x} + qA_x$. Questo vale anche per le componenti y e z dell'impulso, e il primo addendo a secondo membro della (E.1) è ancora l'energia cinetica della particella. Analogamente, in meccanica quantistica $(-i\hbar\nabla - q\mathbf{A})^2/2m$ è l'operatore energia cinetica. Derivando la (E.3) rispetto al tempo otteniamo

$$m\ddot{x} = m \frac{d\dot{x}}{dt} = \dot{p}_x - q \frac{dA_x}{dt} = \dot{p}_x - q \frac{\partial A_x}{\partial t} - q\dot{x} \frac{\partial A_x}{\partial x} - q\dot{y} \frac{\partial A_x}{\partial y} - q\dot{z} \frac{\partial A_x}{\partial z}. \quad (\text{E.4})$$

L'ultimo membro della (E.4) contiene $\dot{p}_x$, che possiamo ottenere dalla seconda delle equazioni del moto canoniche (2.3)

$$\dot{p}_x = -\frac{\partial H}{\partial x} = \frac{q}{m} \left[(p_x - qA_x) \frac{\partial A_x}{\partial x} + (p_y - qA_y) \frac{\partial A_y}{\partial x} + (p_z - qA_z) \frac{\partial A_z}{\partial x} \right] - q \frac{\partial V}{\partial x}. \quad (\text{E.5})$$

Sostituendo $m\dot{x} = p_x - qA_x$, e le formule analoghe per le componenti y e z , la (E.5) diventa

$$\dot{p}_x = q \left(\dot{x} \frac{\partial A_x}{\partial x} + \dot{y} \frac{\partial A_y}{\partial x} + \dot{z} \frac{\partial A_z}{\partial x} \right) - q \frac{\partial V}{\partial x}. \quad (\text{E.6})$$

Se inseriamo la (E.6) nella (E.4) i termini contenenti $\dot{x} \partial A_x / \partial x$ si cancellano, e resta

$$\begin{aligned} m\ddot{x} &= q \left(\dot{y} \frac{\partial A_y}{\partial x} + \dot{z} \frac{\partial A_z}{\partial x} - \frac{\partial A_x}{\partial t} - \dot{y} \frac{\partial A_x}{\partial y} - \dot{z} \frac{\partial A_x}{\partial z} - \frac{\partial V}{\partial x} \right) \\ &= q \left[\dot{y} \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) - \dot{z} \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) - \frac{\partial A_x}{\partial t} - \frac{\partial V}{\partial x} \right] \\ &= q (\dot{y} B_z - \dot{z} B_y + E_x) = q (\mathbf{v} \times \mathbf{B} + \mathbf{E})_x, \end{aligned} \quad (\text{E.7})$$

che è la componente x della forza di Lorentz. Ricordiamo che, in presenza di un campo magnetico dipendente dal tempo, abbiamo $\mathbf{E} = -\nabla V - \partial \mathbf{A} / \partial t$. Ripetendo i conti qui sopra due volte, effettuando ogni volta la permutazione ciclica $x \rightarrow y \rightarrow z \rightarrow x$, si ottengono le altre due componenti della forza di Lorentz.

Bibliografia

- [1] E. Arimondo, *Lezioni di Struttura della Materia*, Edizioni ETS, Pisa, 2005.
- [2] M. Born and E. Wolf, *Principles of Optics*, Pergamon Press, Oxford, New York, Seoul, Tokyo, 1993.
- [3] S. Chandrasekhar, *Stochastic Problems in Physics and Astronomy*, Rev. Mod. Phys. **15** 1–89 (1943) disponibile in rete su http://rmp.aps.org/abstract/RMP/v15/i1/p1_1.
- [4] W.T. Coffey, Yu.P. Kalmykov, and J.P. Waldron, *The Langevin Equation*, World Scientific Series in Contemporary Chemical Physics, Vol. 14, World Scientific Publishing Co., Singapore 2005.
- [5] W. Demtröder, *Laser Spectroscopy*, Springer-Verlag, Berlin, Heidelberg, New York, 1996.
- [6] U. Fano, *Quantum theory of interference effects in the mixing of light from phase independent sources*, American Journal of Physics **29**, 539–545, 1961.
- [7] R.P. Feynman, F.L. Vernon, Jr., and R.W. Hellwarth, *Geometrical Representation of the Schrödinger Equation for Solving Maser Problems*, Journal of Applied Physics, **28**, 49–52, 1956.
- [8] R.P. Feynman, R.B. Leighton, and M. Sands, *The Feynman Lectures on Physics, Volume III, Quantum Mechanics*, Addison-Wesley Publishing Company, Menlo-Park, London, Amsterdam, Don Mills, Sydney, 2006.
- [9] C.W. Gardiner, *Handbook of Stochastic Methods for Physics, Chemistry and the Natural Sciences*, Springer-Verlag, Berlin, Heidelberg, New York, 2004.
- [10] H. Haken e H. C. Wolf, *Fisica Atomica e Quantistica*, edizione italiana a cura di G. Moruzzi, Bollati Boringhieri, Torino, 1990.
- [11] K. Huang, *Statistical Mechanics*, John Wiley and Sons, New York, Chichester, Brisbane, Toronto, Singapore, 1987.
- [12] A. Katz, *Principles of Statistical Mechanics*, W. H. Freeman and Company, San Francisco and London, 1967.
- [13] Charles Kittel, *Elementary Statistical Physics*, Dover Publications Inc., 2004.
- [14] F. K. Kneubühl und M. W. Sigrist, *Laser*, Teubner Verlag, Stuttgart, 1988.

- [15] L. D. Landau and E. M. Lifshitz, *Electrodynamics of Continuous Media*, Volume 8 of *Course of Theoretical Physics*, Chapter VI, Pergamon Press, Oxford, London, New York, Paris, 1960.
- [16] R. Loudon, *The Quantum Theory of Light*, Oxford University Press, Oxford and New York, 2008.
- [17] E. A. Lynton, *Superconductivity*, Methuen & Co. Ltd., London, 1964.
- [18] R.K. Pathria, *Statistical Mechanics*, Elsevier, Amsterdam, Boston, Heidelberg, London, New York, Oxford, Paris, San Diego, San Francisco, Singapore, Sydney, Tokyo, 2006.
- [19] F. Reif, *Fundamentals of Statistical and Thermal Physics*, McGraw-Hill Kogakusha, Tokyo, 1965.
- [20] H. Risken, *The Fokker-Planck Equation*, Springer-Verlag, Berlin, Heidelberg, New York, London, Paris, Tokyo, 1989.
- [21] C. E. Shannon, *A Mathematical Theory of Communication*, The Bell System Technical Journal, **28** 379–423 (1948).
- [22] C. P. Slichter, *Principles of Magnetic Resonance*, Harper & Row, New York, Evanston and London, 1964.
- [23] O. Svelto, *Principles of Lasers*, Springer, New York, 2010.
- [24] J. von Neumann, *Mathematische Grundlagen der Quantenmechanik*, Springer-Verlag, Berlin, 1955. *Mathematical Foundations of Quantum Mechanics*, Princeton University Press, Princeton, 1996.
- [25] Wikipedia

Indice analitico

- Accoppiamento atomo-bagno termico, 177
Aleatoria, passeggiata, 47
Alfërov, Žores Ivanovič, 135
Allargamento collisionale, 96
Allargamento Doppler, 180
Allargamento omogeneo, 179
Allargamento per pressione, 96
Ampiezze di probabilità, 142
Antibunching, 107
Approssimazione di campo rotante, 224
Approssimazione di onda rotante, 144, 224
Approssimazione parassiale dell'equazione di Helmholtz, 121
Assorbente, muro (random walk), 54
Assorbimento, 155
Assorbimento (radiazione), 109
Assorbimento, cella di, 94
Autocorrelazione, funzione di, 63, 64
- Bachelier, Louis, 70
Beam-splitter, matrice del, 85
Bianco, rumore, 77
Bilancio dettagliato, 176, 177
Binomiale, coefficiente, 48
bit, 3
Bloch, equazioni di, 225
Bloch, equazioni ottiche di, 171
Bloch-Siegert, effetto, 224
Bose-Einstein, statistica di, 39
Brown, Robert, 70
Browniano, moto, 70
Buca di potenziale, laser a, 136
Bunching, 107
- Calcolo delle variazioni, 20
Campi complessi, formalismo dei, 83
Campo forte, interazione con, 158
Campo magnetico critico (superconduttività), 193, 195
Campo rotante, approssimazione di, 224
Canoniche, coordinate, 1, 9
Capasso, Federico, 137
Caratteristica, funzione, 236
Carica in campo elettromagnetico, hamiltoniana, 245
Cascata quantica, laser a, 136, 137
Catena markoviana, 61
Cavità risonante, fattore di qualità, 115
Cavità risonante, stabilità, 126
Cella di assorbimento, 94
Centrale, teorema del limite, 235
Chapman-Kolmogorov, equazione di, 62
Chemical shift, 229
Chimico, potenziale, 42
Chimico, spostamento, 229
Cho, Alfred, 137
Coefficiente binomiale, 48
Coefficiente di demagnetizzazione, 199
Coefficienti di Kramers-Moyal, 67, 232
Coerenza temporale al primo ordine, 88
Coerenza temporale al secondo ordine, 103, 104
Coerenze, 166
Collisionale, allargamento, 96
Condizionata, densità di probabilità, 60
Convoluzione, teorema di, 237
Cooper, coppie di, 204
Coordinate canoniche, 1, 9
Coppie di Cooper, 204
Correlazione, funzione di, 63
Correlazione, teorema di, 64
Corrente di superconduttività, 200, 203
Costanti, perturbazioni, 147
Curva di guadagno, 128
- Delta di Dirac, 84

- Densità di corrente in un superconduttore, 204
 Densità di probabilità al primo ordine, 59
 Densità di probabilità al secondo ordine, 59
 Densità di probabilità condizionata, 60
 Densità spettrale, 62
 Densità spettrale di energia, 62
 Densità spettrale di potenza, 63, 92
 Deriva, equazione di, 68
 Derivazione alternativa dell'equazione di Fokker-Planck, 231
 Diamagnetismo perfetto (superconduttori), 196
 Diffusione, equazione di, 68, 69
 Diodi a emissione luminosa, 135
 Dipendenti dal tempo, perturbazioni, 141
 Dirac, delta di, 84
 Diretto, gap, 135
 Doppia eterostruttura, 136
 Doppia giunzione, 136
 Doppler, allargamento, 180
 Drift, equazione di, 68
 Drude, Paul, 82
 Dulong e Petit, legge di, 32

 Eco di spin, 227
 Effetto Bloch-Siegert, 224
 Effetto Meissner, 194, 206
 Effetto Raman stimolato, 153
 Efficienza quantica, 137
 Einstein, Albert, 70, 109
 Emissione indotta, 110
 Emissione spontanea, 109, 155
 Emissione spontanea e seconda quantizzazione, 241
 Emissione stimolata, 110, 155
 Energia di Fermi, 133
 Energia libera, 192
 Energia libera di Helmholtz, 192
 Energia, densità spettrale di, 62
 Entropia di Shannon, 7
 Entropia di von Neumann, 7
 Entropia e informazione mancante, 30
 Entropia per l'insieme macrocanonico, 34
 EPR, spettroscopia, 230
 Equazione del bilancio dettagliato, 177
 Equazione di Chapman-Kolmogorov, 62
 Equazione di continuità per la funzione d'onda, 203
 Equazione di deriva, 68
 Equazione di diffusione, 68, 69
 Equazione di diffusione, soluzione fondamentale, 54
 Equazione di drift, 68
 Equazione di Fokker-Planck, 66
 Equazione di Helmholtz, 121
 Equazione di Helmholtz, approssimazione parasiale, 121
 Equazione di Kramers-Moyal, 67
 Equazione di Langevin, 71
 Equazione di Smoluchowski, 62
 Equazioni di bilancio per i fotoni, 115
 Equazioni di bilancio per le popolazioni, 117
 Equazioni di Bloch, 225
 Equazioni di Bloch ottiche, 171
 Equazioni di Hamilton, 9
 Equilibrio termodinamico, 167, 170, 176, 177
 Equipartizione dell'energia, teorema di, 31
 Esperimento di Michelson e Morley, 97
 Esponenziale di un operatore, 14
 Estensione del fascio, 101
 Estrinseci, semiconduttori, 133

 Faist, Jerome, 137
 Fano, Ugo, 105
 Fasci gaussiani, 122
 Fascio, estensione del, 101
 Fattore di qualità di un cavità risonante Q , 115
 Fattore giromagnetico, 217, 219
 Fenomeni di interferenza, 83
 Fermi, energia di, 133
 Fermi, regola d'oro di, 148
 Fermi, temperatura di, 133
 Fermi-Dirac, statistica di, 39
 Feynman-Vernon-Hellwarth, modello di, 168
 Flusso magnetico, quantizzazione nei superconduttori, 203
 Flusso, quanto di, 208
 Flussone, 208
 Fluttuazione quadratica media, 38
 Fluttuazioni dei numeri di occupazione, 43
 Fluttuazioni in statistica classica, 38

- Fock, stati di, 40, 241
- Fock, Vladimir Alexandrovič, 40
- Fokker-Planck, derivazione alternativa, 231
- Fokker-Planck, equazione di, 66
- Forma di riga, 95, 179
- Formalismo dei campi complessi, 83
- Fotoconteggi, 78
- Fotomoltiplicatore, 78
- Fotoni, equazioni di bilancio per i, 115
- Fourier, spettroscopia a trasformata di, 91
- Fourier, trasformata di, 62
- Frequenza di Larmor, 218
- Frequenza di Rabi, 160, 162
- Funzione caratteristica, 236
- Funzione ceiling, 3
- Funzione di autocorrelazione, 63, 64
- Funzione di correlazione, 63
- Funzione di partizione, 25
- Funzione floor, 3
- Gap diretto, 135
- Gap indiretto, 135
- Gas perfetti, legge dei, 29
- Gas perfetto classico, 28
- Gas perfetto classico: insieme macrocanonico, 36
- Gaussiana, trasformata di Fourier della, 238
- Gaussiani, fasci, 122
- Geometrica, progressione, 239
- Geometrica, serie, 239
- Giromagnetico, fattore, 217, 219
- Giromagnetico, rapporto, 217
- Gouy, Louis Georges, 123
- Gran canonico, insieme, 33
- Grandezza “numero delle particelle”, 34
- Hahn, Erwin, 227
- Hamilton, equazioni di, 9
- Hamiltoniana di perturbazione, 141
- Hamiltoniana di una carica in campo elettromagnetico, 245
- Hamiltoniana imperturbata, 141
- Hanbury-Brown e Twiss, interferometro di, 103
- Heisenberg, rappresentazione di, 11
- Helmholtz, equazione di, 121
- Henry, Charles H., 136
- Hutchinson, Albert, 137
- Imperturbata, hamiltoniana, 141
- Impulso 2π , 223
- Impulso π , 223
- Impulso $\pi/2$, 223
- Indiretto, gap, 135
- Indotta, emissione, 110
- Informazione disponibile in statistica classica, 20
- Informazione mancante, 2, 10
- Ingenhousz, Jan, 70
- Insieme gran canonico, 33
- Insieme macrocanonico, 33
- Insieme macrocanonico, entropia, 34
- Insieme macrocanonico: gas perfetto classico, 36
- Insieme macrocanonico: statistica quantistica, 39
- Insieme statistico, 59
- Installazione del regime oscillatorio, 128
- Intensità di saturazione, 173
- Interazione con campo forte, 158
- Interazione radiazione-materia, 141
- Interbanda, transizioni, 137
- Interferenza, fenomeni di, 83
- Interferometro di Hanbury-Brown e Twiss, 103
- Interferometro di Mach-Zehnder, 86
- Interferometro di Michelson, 89
- Interferometro di Young, 98
- Interferometro stellare di Michelson, 102
- Intrabanda, transizioni, 137
- Intrinseci, semiconduttori, 133
- Inversione di popolazione, 114
- Johnson, John B., 80
- Johnson, rumore, 80
- Kirchhoff, Gustav, 111
- Kramers-Moyal, coefficienti di, 67, 232
- Kramers-Moyal, equazione di, 67
- Kroemer, Herbert, 135
- Lagrange, Joseph-Louis, 19
- Lagrange, moltiplicatori di, 19
- Langevin, equazione di, 71
- Langevin, Paul, 71
- Larghezza naturale, 96
- Larmor, frequenza di, 218
- Laser a buca di potenziale, 136
- Laser a cascata quantica, 136, 137

- Laser a CO₂, 120
 Laser a elio-neon, 120
 Laser a quattro livelli, 119
 Laser a semiconduttore, 133
 Laser, curva di guadagno, 128
 Laser, modi di oscillazione, 128
 Lavoro di magnetizzazione, 189
 Lavoro di magnetizzazione “meccanico”, 189
 Lavoro di magnetizzazione effettuato dal generatore, 190
 LED, 135
 Legge dei gas perfetti, 29
 Legge di Dulong e Petit, 32
 Legge di Rayleigh-Jeans, 111
 Logaritmo di un operatore, 18

 Mach, Ludwig, 86
 Mach-Zehnder, interferometro di, 86
 Macrocanonico, insieme, 33
 Magnetica, risonanza, 217
 Magnetico, raffreddamento, 211
 Magnetizzazione, lavoro di, 189
 Markov, Andrej Andrejevič, 61
 Markoviani, processi, 61
 Massima informazione mancante, 18
 Massimo condizionato, 19
 Matrice del beam-splitter, 85
 Matrice densità, 14, 165
 Meissner, effetto, 194, 206
 Metodo perturbativo, 143
 Michelson e Morley, esperimento di, 97
 Michelson, interferometro di, 89
 Michelson, interferometro stellare di, 102
 Mode-locking, 130
 Modello di Feynman-Vernon-Hellwarth, 168
 Modi di oscillazione di un laser, 128
 Moltiplicatori di Lagrange, 19
 Moltiplicatori di Lagrange in statistica classica, 20
 Momenti di una variabile aleatoria, 56
 Moto browniano, 70
 Muro assorbente (random walk), 54
 Muro riflettente (random walk), 54

 Naturale, larghezza, 96
 NMR, spettroscopia, 229

 Numeri di occupazione, 42
 Numeri di occupazione, fluttuazione dei, 43
 Numero d'onda, 95
 Numero delle particelle, grandezza fisica, 34
 Numero di Avogadro, 1
 Numero variabile di particelle identiche, 33
 Nyquist, Harry, 80
 Nyquist, teorema di, 80

 Occupazione, numeri di, 42
 Omogeneo, allargamento, 179
 Onda rotante, approssimazione di, 144, 224
 Onda, numero di, 95
 Onde parassiali, 121
 Onnes, H. Kammerlingh, 193
 Oscillazioni di Rabi, 162

 Parametro di stabilità, 127
 Parassiale, approssimazione dell'equazione di Helmholtz, 121
 Parassiali, onde, 121
 Parentesi di Poisson, 10
 Parseval, teorema di, 62
 Particelle identiche in fisica classica, 22
 Partizione, funzione di, 25
 Passeggiata aleatoria, 47
 Periodiche, perturbazioni, 144
 Perturbativo, metodo, 143
 Perturbazione, hamiltoniana di, 141
 Perturbazioni costanti, 147
 Perturbazioni dipendenti dal tempo, 141
 Perturbazioni periodiche, 144
 Perturbazioni periodiche al secondo ordine, 150
 Poisson, distribuzione di, 37
 Poisson, parentesi di, 10
 Popolazione, inversione di, 114
 Popolazioni, 166
 Popolazioni, equazioni di bilancio per le, 117
 Potenza, densità spettrale di, 63, 92
 Potenziale chimico, 42
 Poynting, vettore di, 84
 Pressione, allargamento per, 96
 Primo ordine, coerenza temporale al, 88
 Probabilità a priori, 17
 Probabilità di transizione, 141
 Probabilità, ampiezze di, 142

- Processi markoviani, 61
- Processi stocastici, 59
- Processo stocastico stazionario, 60
- Progressione geometrica, 239
- Punti di sella, 20
- Q, fattore, 115
- Quantica, efficienza, 137
- Quantico, rendimento, 137
- Quantizzazione del flusso magnetico, 203
- Quanto di flusso, 208
- Quantum cascade laser, 136
- Quantum well laser, 136
- Quasiparticella, 204
- Rabi, frequenza di, 160, 162
- Rabi, oscillazioni di, 162
- Raffreddamento magnetico, 211
- Raman, effetto, 153
- Random walk, 47
- Random walk, limite continuo, 53
- Random walk, limite per grandi N, 51
- Rapporto giromagnetico, 217
- Rappresentazione di Schrödinger, 141
- Rate equations, 176, 177
- Rate equations per i fotoni, 115
- Rate equations per le popolazioni, 117
- Rayleigh-Jeans, legge di, 111
- Regione di svuotamento, 135
- Regola d'oro di Fermi, 148
- Regole di selezione, 181
- Rendimento quantico, 137
- Reservoir, 178
- Riflettente, muro (random walk), 54
- Riga, forma di, 95
- Rilassamento, 171
- Rilassamento (risonanza magnetica), 224
- Rilassamento longitudinale, 171
- Rilassamento longitudinale, tempo di, 224
- Rilassamento parallelo, 171
- Rilassamento perpendicolare, 171
- Rilassamento trasversale, 171
- Rilassamento trasversale, tempo di, 225
- Rinormalizzazione dell'informazione mancante, 7
- Risonanza magnetica, 217
- risuonatore concentrico, 127
- Risuonatore confocale, 127
- Risuonatore piano, 127
- Risuonatori semisferici, 128
- Risuonatori simmetrici, 128
- Risuonatori stabili a due specchi, 126
- Rumore $1/f$, 77
- Rumore bianco, 77
- Rumore Johnson, 80
- Rumore rosa, 77
- Sandwich, struttura a, 136
- Saturazione, intensità di, 173
- Scelta nel continuo, 5
- Scelte equiprobabili, 2
- Schottky, Walter H., 75
- Schrödinger, rappresentazione di, 11, 141
- Seconda quantizzazione ed emissione spontanea, 241
- Secondo ordine, coerenza temporale al, 103, 104
- Secondo ordine, perturbazioni periodiche al, 150
- Selezione, regole di, 181
- Semiconduttori estrinseci, 133
- Semiconduttori intrinseci, 133
- Semisferici, risuonatori, 128
- Serie geometrica, 239
- Sfera in campo magnetico, 187
- Shannon, entropia di, 7
- Shot noise, 75
- Simmetrici, risuonatori, 128
- sinc, 96, 145
- sinc non normalizzata, 96
- sinc normalizzata, 96
- sinc^2 , 145
- Sirtori, Carlo, 137
- Sivco, Deborah, 137
- Smoluchowski, equazione di, 62
- Soluzione fondamentale dell'equazione di diffusione, 54
- Spettrale, densità, 62
- Spettroscopia a trasformata di Fourier, 91
- Spettroscopia EPR, 230
- Spettroscopia NMR, 229
- Spin, eco di, 227
- Spontanea, emissione, 109, 241
- Spostamento chimico, 229

- Stabilità della cavità risonante, 126
Stabilità, parametro di, 127
Stati di Fock, 40, 241
Statistica classica, informazione disponibile, 20
Statistica di Bose-Einstein, 39
Statistica di Fermi-Dirac, 39
Statistico, insieme, 59
Stato intermedio (superconduttività), 199
Stato puro, 165
Stazionario, processo stocastico, 60
Stimolata, emissione, 110
Stirling, approssimazione di, 5
Stocastici, processi, 59
Struttura a sandwich, 136
Superconduttività, 187
Superconduttività, corrente di, 200, 203
Superconduttività, teoria microscopica, 208
Superconduttore, densità di corrente, 204
Superlattice, 137
Superreticolo, 137
Svuotamento, regione di, 135
- Temperatura critica (superconduttività), 193
Temperatura di Fermi, 133
Tempo di rilassamento longitudinale, 224
Tempo di rilassamento trasversale, 225
Teorema centrale del limite, 235
Teorema di convoluzione, 237
Teorema di correlazione, 64
Teorema di equipartizione dell'energia, 31
Teorema di Nyquist, 80
Teorema di Parseval, 62
Teorema di Wiener-Khinchin, 65, 81
Teoria microscopica della superconduttività, 208
Transizione, probabilità di, 141
Transizioni a due fotoni, 152
Transizioni di dipolo elettrico, 153
Transizioni in uno spettro continuo, 147
Transizioni interbanda, 137
Transizioni intrabanda, 137
Trasformata di Fourier, 62
Trasformata di Fourier della gaussiana, 238
Trattamento semiclassico dell'interazione radiazione-materia, 147
- Variabile aleatoria, momenti di una, 56
Variazioni, calcolo delle, 20
Vettore di Poynting, 84
von Neumann, entropia di, 7
Wiener-Khinchin, teorema di, 64, 65, 81
Young, interferometro di, 98
Zavojskij, Evgenij Konstantinovič, 230
Zehnder, Ludwig, 86