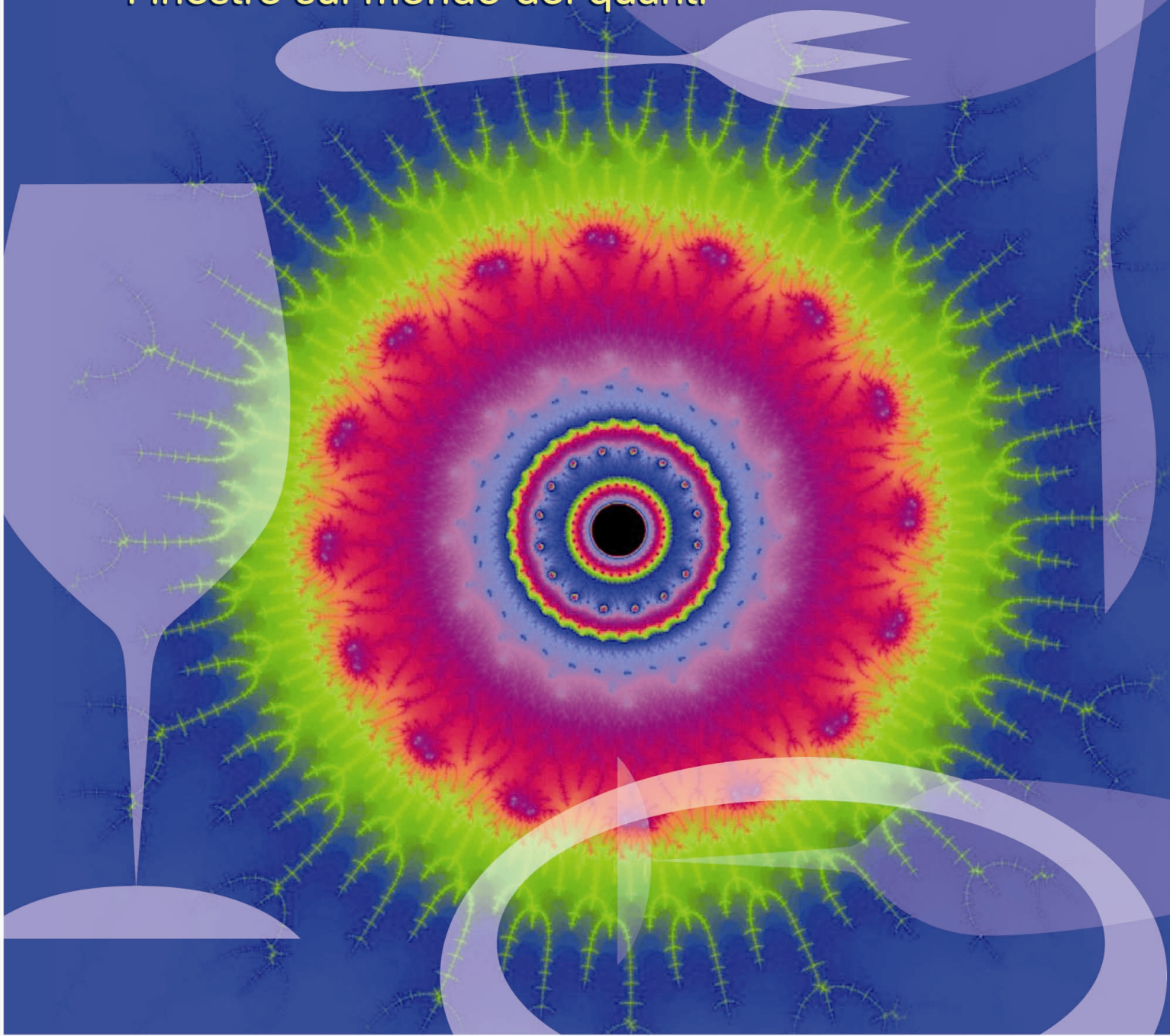


Attilio Rigamonti
Andrei Varlamov

Magico caleidoscopio della fisica

Fisica, poesia e armonia nei fenomeni naturali
La fisica del sabato sera
Da "Homo erectus" a "Cuoco sapiens"
Finestre sul mondo dei quanti



Attilio Rigamonti e Andrei Varlamov

MAGICO CALEIDOSCOPIO DELLA FISICA

Fisica, poesia e armonia nei fenomeni naturali
La fisica del Sabato sera
Da “Homo erectus” a “cuoco sapiens”
Finestre sul mondo dei quanti

2007

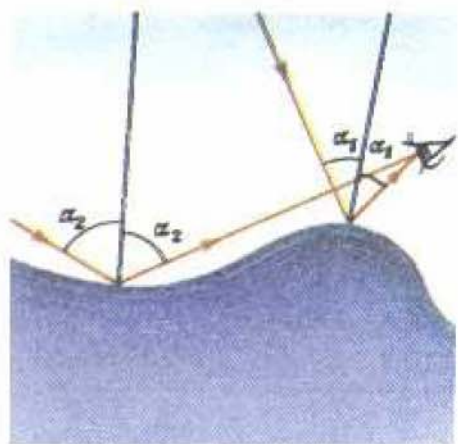


Fig. 1.

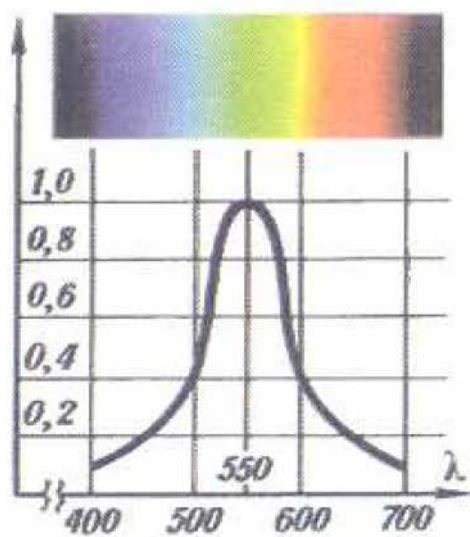


Fig. 2. 1

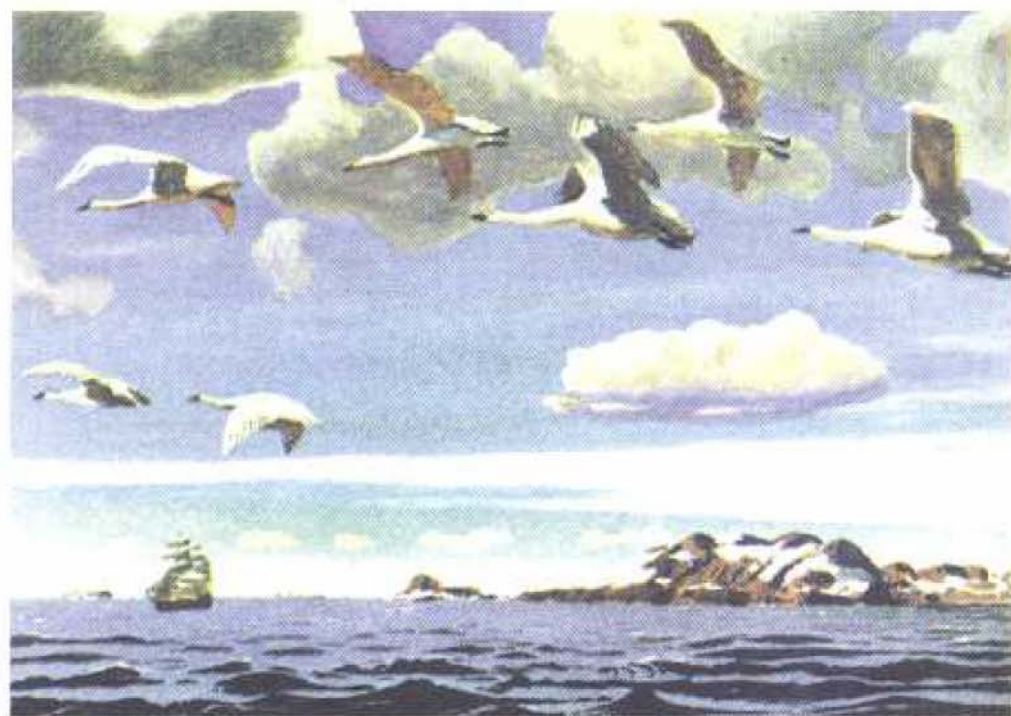


Fig. 3.

In memoria di Lev Aslamazov, preclaro fisico e divulgatore, con commosso ricordo.

Prefazione

Quale bimbo non ha guardato con meraviglia l'interno di un caleidoscopio, nel quale frammenti di vetri colorati e di specchi componevano una cangiante varietà di immagini? Un caleidoscopio della fisica si propone analoghi effetti: che attraverso trame e frammenti si possano cogliere meraviglianti fenomeni, si possa osservare come nella stupefacente realtà della natura, bellezza e armonia impensate nascano dalla astrazione matematica, sia nel mondo macroscopico sia a livello dei microscopici quanti.

L'intento principale di questo testo è quello di avvicinare agli aspetti della stupenda architettura costituita dalla realtà del mondo le persone non familiari con la disciplina fisica. I criteri ispiratori sono i seguenti: in un linguaggio piano e la forma la più agevole possibile, il tentativo di sollecitare il lettore a guardare con occhi da "praticante ricercatore" un caleidoscopio di fenomeni fisici, di aiutarlo verso l'esplorazione degli eventi, con aspirazione ad apprendere e desiderio di acquisire conoscenza, attitudini che secondo Einstein costituiscono caratteristica del saggio, a differenza della erudizione in sé stessa.

La speranza degli autori è che coloro che si lasceranno guidare in una lettura meditante possano un giorno sorprendersi a seguire incuriositi il serpeggiare di un fiume e la struttura delle sue rive, a soffermarsi nella osservazione del rosso fuoco di un tramonto, a studiare il battito d'ala degli uccelli, ad ascoltare il bollire di una teiera o il suono di un violino, oppure a meditare sulle disposizioni degli elettroni negli atomi o ancora a stupirsi come con la superconduttività l'uomo abbia imparato a controllare la natura a un livello delicato e bellissimo, quello quantistico. In questo percorso sugli aspetti quasi fiabeschi della natura non mancano cenni alle più

importanti realizzazioni applicative (come lo SQUID, la risonanza magnetica nucleare, il metodo SNIF, le nanoscienze e i qubit), richiami ad aspetti artistici o di divertissement, come la dettagliata osservazione di una lampada magica o le connessioni tra fisica e vino (sino a stimare con procedura matematica la quantità ottimale di questa bevanda che adorna la nostra tavola). Una sezione è dedicata alla fisica in cucina: si potrà scoprire come sia diffuso il suo ruolo e come si possano individuare gioiosi aspetti dell'esecuzione culinaria, aumentando nel contempo la qualità dei prodotti.

Gli autori si augurano che i colleghi che appartengono alla universale carovana della fisica, legati dalla magica arte galileiana di intendersi tra di loro con equazioni e numeri e in tal modo capaci di cogliere la bellezza architettonica del reale, leggano questo libro con piacere, trovandovi qualche inusitato richiamo. Per gli insegnanti di fisica potrà essere una fonte di spunti didattici. Tuttavia, per sua natura, il testo è soprattutto destinato ad avvicinare alla stupefacente architettura della realtà, come si è detto, le persone non professionalmente coinvolte con la disciplina fisica. Peraltro, sarebbe riduttivo considerarlo un testo di divulgazione scientifica. Si tratta piuttosto di "lezioni" impartite in una forma particolare: accanto alla massima riduzione dei termini del problema si privilegia il ricorso all'intuito, a considerazioni di scala, ad analogie e a visualizzazioni del contenuto fisico delle equazioni e degli ordini di grandezza. Per la comprensione della maggior parte degli argomenti è necessaria soltanto della matematica elementare e la conoscenza dei principi più generali della fisica, quali si apprendono in una scuola media superiore. Pur tuttavia, non è un libro "facile": infatti nulla è banalizzato e si è mantenuto rigore concettuale, l'obiettivo essendo quello di insegnare qualcosa, anche se in una maniera particolare.

Gli autori di questo libro sono fisici professionisti che hanno alternato ricerche su temi specialistici e piuttosto complessi a pause, spesso dopo serate conviviali, per coltivare la possibilità di vedere al di là degli argomenti più tecnici sui quali operavano, di non perdere gli aspetti ludici della ricerca scientifica e la capacità di cogliere i gratificanti momenti legati alla osservazione e comprensione dei fenomeni naturali. Ritengono di poter aggiungere a questa breve prefazione il commento di Einstein: "L'uomo per il quale non è più familiare il sentimento del mistero, che ha perso la facoltà di meravigliarsi davanti alla creazione, è come un uomo morto, i suoi occhi sono spenti. . . Nessuno si può sottrarre a un sentimento di riverente commozione contemplando i misteri della stupenda struttura della realtà".

Infine, nell'elaborare sugli argomenti che si riportano in questo testo gli autori si sono divertiti: sommessamente si augurano che analogo piacere possa ricevere il lettore che senta in sè lo stimolo a osservare e a comprendere.

Gli autori esprimono il più caloroso apprezzamento, purtroppo postumo, a Lev Aslamazov, ispiratore dell'intento generale che sottende questo loro testo, autore di alcuni dei precedenti articoli divulgativi dai quali sono successivamente nati capitoli come "Fiumi, meandri e laghi", "Il pendolo di Foucault e la legge di Beer", "Tracce sulla sabbia" e "Al suono del violino". Alla Sua memoria dedicano questo loro lavoro. Molti colleghi e amici hanno direttamente o indirettamente collaborato, attraverso articoli su specifici argomenti o discussioni, alla stesura di questo testo. Si ringraziano, in particolare, Giuseppe Balestrino, Alexander Budzin, Carlo Camerlingo, Yuri Galperin. Gli autori sono debitori a Giuseppe Angilella per la segnalazione dell'articolo sulla frattura degli spaghetti e a Maurizio Corti per le immagini NMR.

Stefania Tentoni ha garantito una apprezzata assistenza tecnico - computazionale.

Una parte del presente testo riprende alcuni argomenti presenti in "Fisica che meraviglia" di Aslamazov e Varlamov. Pertanto gli autori sono grati a Loredana Bove e Lucia Reggiani, che avevano provveduto alla traduzione dal russo e a una prima redazione di quel testo. Si segnala, inoltre, che nel 2004 la traduzione in lingua inglese (per la casa editrice World Scientific, con il titolo "The Wonders of Physics") aveva già comportato l'inclusione di alcuni argomenti riportati ora nel presente libro. Gli autori ringraziano pertanto l'editore scientifico di quel testo, A.A. Abrikosov (jr.), per avere contribuito con la sua assistenza al miglioramento di alcune parti.

Infine gli autori manifestano il loro ringraziamento ai gestori della trattoria "La Barcela", Lorenza e Tino Alberico per l'ospitalità loro concessa nel corso di diverse serate, nelle quali parti di questo libro sono state elaborate o presentate al Circolo culturale "La Barcela", di Pavia.

(Pavia, 2007).

Indice

Prefazione	vii
Parte I	<u>Fisica, poesia e armonia nei fenomeni naturali</u>
Capitolo 1	Fiumi, meandri e laghi. 5
Capitolo 2	Un interfono oceanico. 15
Capitolo 3	Nell'azzurro dello spazio. 25
Capitolo 4	Il pendolo di Foucault e la legge di Beer. 39
Capitolo 5	Le maree e il freno lunare. 49
Capitolo 6	Bolle e gocce. 55
Capitolo 7	Il microfono ad acqua. 69
Capitolo 8	Tracce sulla sabbia. 75
Parte II	<u>“Divertissement”: la fisica del Sabato sera</u>
Capitolo 9	Conversando nel viaggio in treno. 91
Capitolo 10	Al suono del violino. 99
Capitolo 11	Calici sonori e calici silenti. 107
Capitolo 12	Osservando la lampada magica. 113
Capitolo 13	Nunc est bibendum. 123

Parte III	<u>Da “Homo erectus” a “cuoco sapiens”</u>	
Capitolo 14	La Signora prepara il tè.	145
Capitolo 15	Porchetta alle micro-onde?	159
Capitolo 16	Ab ovo.	169
Capitolo 17	Il tacchino nel giorno del ringraziamento.	177
Capitolo 18	Pasta, spaghetti e fisica.	183
Capitolo 19	La fisica di un buon caffè.	195
Capitolo 20	Gelato all’azoto liquido e cucina “molecolare”.	209
Parte IV	<u>Finestre sul mondo dei quanti</u>	
Capitolo 21	Il principio di indeterminazione.	217
Capitolo 22	Fenomeni meraviglianti.	229
Capitolo 23	Palle di neve e bollicine nell’elio liquido.	237
Capitolo 24	L’affascinante mondo dei superconduttori.	247
Capitolo 25	Visione NMR del nostro interno.	277
Capitolo 26	Nanoscienza e futuro del computer.	289

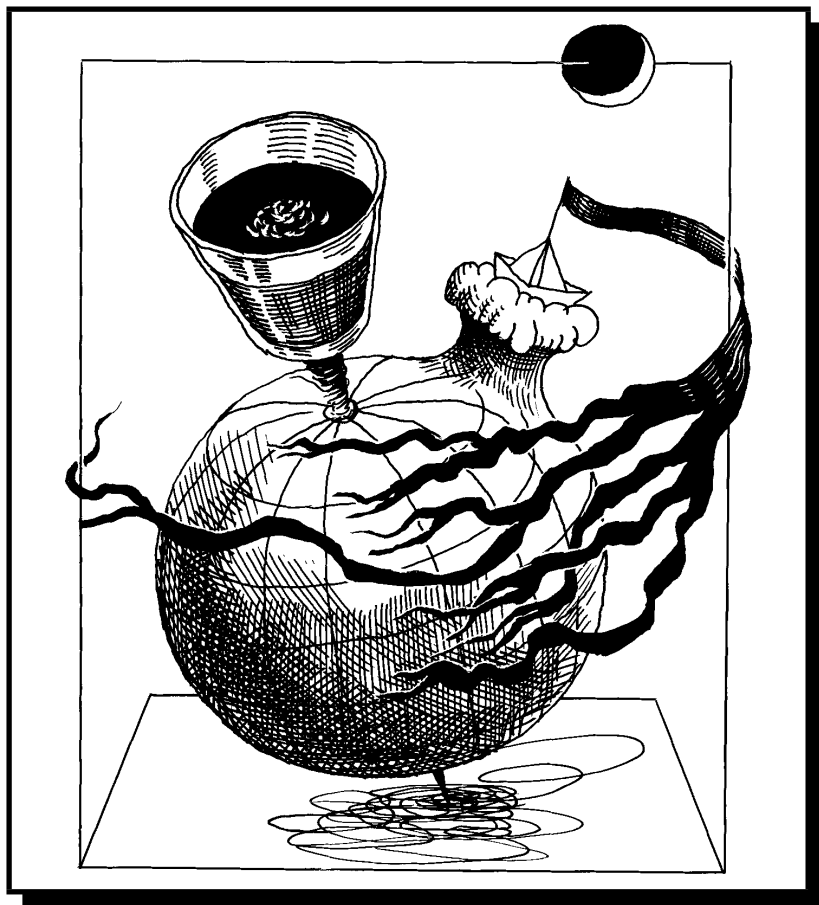
PARTE I

Fisica, poesia e armonia nei fenomeni naturali.

In questa prima parte i lettori capiranno perché i fiumi scorrono a meandri, come cambiano i loro letti e come si formano le rive; in che cosa consista un interessante esempio di guida d'onda esistente in natura, il canale sonoro subacqueo che si forma negli oceani ed è usato dalle balene e dai sottomarini per comunicare alle grandi distanze; cosa si debba “vedere” in un quadro o nell'azzurro del cielo e nel rosso del sole al tramonto o nelle bianche nuvole che solcano i cieli; cosa determina la frequenza del battito d'ala degli uccelli, perché il mare è azzurro o bianco di schiuma.

Si scriverà degli oceani e dei venti, di come la rotazione terrestre influisca su di loro, di come la luna determini la periodicità delle maree e l'allungamento della durata del giorno. Descriveremo come una bolla di sapone con i suoi cangianti colori, che ci ha affascinato da bambini, possa insegnarci molta fisica; come si comportano le gocce nel miscelarsi e quali forme assumano.

In breve, guarderemo ad alcuni fenomeni che avvengono ogni giorno nel mondo che ci circonda e che osservati con occhi meno distratti possono essere fonte di meraviglia e anche di serenità. Sarete sorpresi nel vedere come fenomeni apparentemente molto strani e complessi possano spesso essere spiegati e meglio apprezzati con una fisica relativamente semplice.



Capitolo 1

Fiumi, meandri e laghi.

Quante volte abbiamo percorso un sentiero che si snodava sulla riva di un ruscello o di un fiume e ci siamo chiesti perché il sentiero, sulla riva, procedesse a meandri? Naturalmente si incontrano anche tratti di un fiume in cui esso avanza in linea retta o quasi, ma fiumi del tutto privi di curve non esistono. Pur se il fiume scorre in una pianura, di solito scende dolcemente a zig-zag e spesso le curve si ripetono ad intervalli regolari (si veda la Fig. 1.1). Come spiegarne i meandri ?



Fig. 1.1: Esempi del procedere a meandri dei fiumi.

L'idrodinamica, il campo della fisica che studia il moto dei liquidi, è una disciplina di carattere assai avanzato. Tuttavia, nell'applicazione a oggetti naturali così complessi come un ruscello o un fiume, l'idrodinamica non è sempre in grado di spiegare tutti gli aspetti del loro moto. Peraltro, come

vedremo, a molti interrogativi si può dare una risposta.

Il problema delle cause della formazione delle tortuosità è stato affrontato per la prima volta da **Einstein**. Nella sua nota “Causa della formazione delle tortuosità nei letti dei fiumi e cosiddetta legge di Beer”, presentata al congresso dell’Accademia delle Scienze Prussiana nel 1926, Einstein, lasciando per qualche tempo la complessa fisica-matematica degli spazi-tempo curvi della relatività generale, analizzò il movimento dei vortici dell’acqua in un bicchiere e quello in un letto di un fiume. Del resto, sin da ragazzo Einstein mostrava una naturale tendenza alla attenta e stupita osservazione dei fenomeni naturali. **Ernestina Marangoni** (alla quale Einstein indirizzò alcune lettere conservate alla Università di Pavia) in un articolo lo ricorda in Pavia, quando giovinetto osservava pensoso gli anelli di fumo che dalla sigaretta salgono in lenti vortici ruotanti, cercando il perché del fenomeno (un problema che fu successivamente affrontato e risolto da **Richard Feynman**, nel contesto dei vortici in superfluidi e in superconduttori).

Ritornando ai fiumi e ai vortici nel bicchiere, l’analogia tra i due casi ha permesso di spiegare la tendenza dei fiumi ad acquisire una forma tortuosa.

Cerchiamo di capire questo fenomeno, sia pure qualitativamente, iniziando con il guardare ai granelli di tè in un bicchiere.

1.1 Granelli di tè nel bicchiere

Prendete un bicchiere contenente grani di tè, mescolate con un cucchiaino, quindi rimuovetelo. L’acqua si ferma gradatamente, mentre i granelli si raccolgono al centro del fondo del bicchiere. Cosa spinge i granelli a raggrupparsi al centro? Per rispondere a questa domanda spiegheremo inizialmente quale forma assume la superficie dell’acqua che si muove nel bicchiere.

L’osservazione ci indica che la superficie del liquido prende una forma concava. È facile capire il perché. Affinché abbia luogo un movimento rotatorio, la risultante delle forze agenti su ogni particella deve determinare un’accelerazione centripeta. Immaginiamo di individuare un cubetto di massa Δm all’interno del liquido, ad una distanza r dall’asse di rotazione (si veda la Fig. 1.2, *a*).

Con una velocità angolare di rotazione ω l’accelerazione centripeta del cubetto è uguale a $\omega^2 r$. Questa accelerazione è legata alla differenza di pressione esercitata radialmente sulle facce laterali (destra e sinistra) del

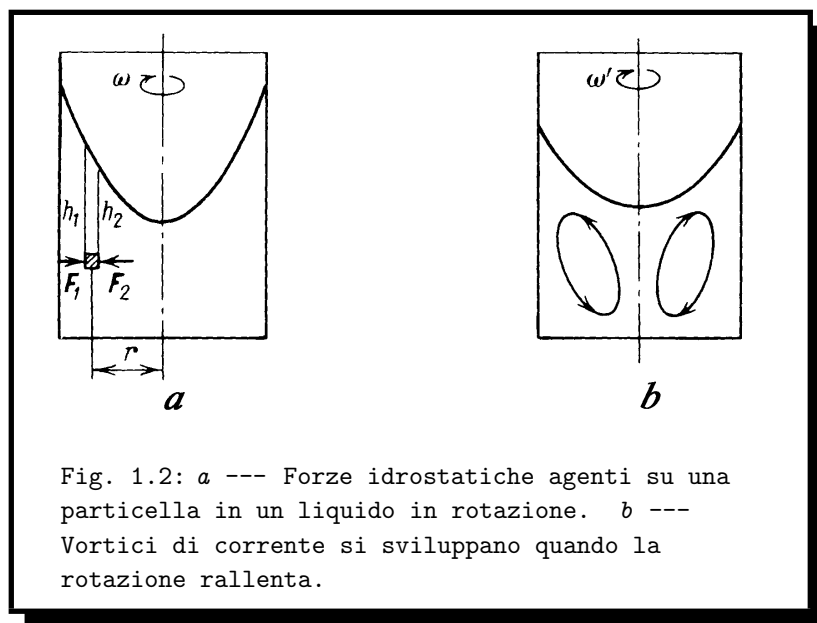
cubo. Scriveremo di conseguenza

$$m\omega^2 r = F_1 - F_2 = (P_1 - P_2) \Delta S, \quad (1.1)$$

dove ΔS è l'area della faccia laterale del cubo. Le pressioni P_1 e P_2 sono legate alle distanze h_1 e h_2 dalla superficie libera del liquido:

$$P_1 = \rho g h_1 \quad \text{e} \quad P_2 = \rho g h_2, \quad (1.2)$$

dove ρ è la densità e g l'accelerazione di gravità. Poiché la forza F_1 deve essere maggiore della forza F_2 , anche h_1 deve essere maggiore di h_2 , ossia la superficie libera del liquido durante la rotazione deve incurvarsi, come è indicato nella figura 1.2. Quanto maggiore è la velocità di rotazione, tanto più si incurva la superficie, in modo da aumentare la forza necessaria ad equilibrare la forza centrifuga nel moto rotatorio.



Si può calcolare la forma della superficie libera del liquido. Tale superficie è un paraboloide, ossia una superficie la cui sezione con il piano del foglio è una parabola.

Quando mescoliamo il tè con il cucchiaino, manteniamo la rotazione nel liquido. Non appena si estrae il cucchiaino dal bicchiere, a seguito

dell'attrito tra i singoli strati del liquido (il fenomeno della viscosità) e dell'attrito sulle pareti e sul fondo del bicchiere, l'energia cinetica associata alla rotazione viene dissipata in calore e il moto del liquido gradualmente si arresta.

Quando la frequenza di rotazione diminuisce, la superficie tende a divenire piana. In questo caso all'interno del liquido si formano vortici, la cui direzione è indicata nella Fig. 1.2, *b*. La formazione di tali vortici è collegata al rallentamento non uniforme del liquido sul fondo del bicchiere e sulla superficie libera. In profondità, infatti, a seguito del forte attrito sul fondo il liquido rallenta più velocemente che sulla superficie. Per questo, particelle di liquido che si trovano a distanze uguali dall'asse di rotazione hanno differenti velocità: quanto più sono vicine al fondo del bicchiere, tanto minore è la loro velocità. La risultante delle forze della pressione laterale, agenti su queste particelle è ancora la stessa. Questa forza non può produrre la necessaria accelerazione centripeta lungo tutta la profondità (come nel caso di rotazione di tutta la massa del liquido con una stessa velocità angolare). Sulla superficie la velocità angolare è più alta e le particelle d'acqua tendono verso le pareti del bicchiere. Sul fondo la velocità angolare è minore e la risultante delle forze di pressione spinge l'acqua a spostarsi verso il centro.

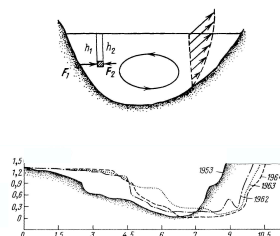
Ora è chiaro perché i granelli di tè tendono a raccogliersi verso il centro del fondo del bicchiere: essi sono sollecitati dai vortici che si formano durante il rallentamento. Questa descrizione, benché molto semplificata, riflette correttamente l'essenza del fenomeno.

1.2 Come cambiano i letti dei fiumi

Analizziamo ora le caratteristiche del movimento dell'acqua di un fiume quando effettua una svolta. In questo caso abbiamo un moto simile a quello dell'acqua nel bicchiere: la superficie dell'acqua s'inclina nel verso della svolta, in modo che la differenza delle forze di pressione determina l'accelerazione centripeta. Nella Fig. 1.3 è mostrata la sezione verticale del fiume alla svolta. Quindi, come nel caso del bicchiere con grani di tè, la velocità dell'acqua sul fondo, a causa dell'attrito, è inferiore alla velocità sulla superficie del fiume (la distribuzione della velocità in funzione della profondità è indicata in figura).

Pertanto, sulla superficie, la risultante delle forze non è in grado di sostenere il movimento delle particelle lungo la circonferenza ad alta velocità

Fig. 1.3: Sezione del letto di un fiume alla svolta. Forze idrostatiche, vortici di corrente e distribuzione della velocità. Nel corso del tempo si produce l'alterazione del letto del fiume.



e l'acqua si sposta verso la riva lontana (dal centro della svolta). Sul fondo, al contrario, la velocità è piccola e l'acqua si dirige verso la riva vicina (al centro della svolta). Quindi, in aggiunta alla corrente principale, si sovrappone questa circolazione di acqua. Nella figura 1.3 è indicata la direzione della circolazione nel piano della sezione del fiume.

Tale circolazione produce un'erosione della riva lontana dal centro della svolta, mentre sulla riva vicina gradualmente si deposita uno strato di terreno sempre maggiore (ricordate i granelli di tè nel bicchiere!). Pertanto la forma dell'alveo subisce nel tempo notevoli variazioni. Si osservi che una analoga circolazione può sussistere anche in caso di corrente rettilinea del fiume a seguito della rotazione terrestre, della quale ci occuperemo successivamente (al capitolo 4 relativo al pendolo di Foucault e la legge di Beer). In conseguenza di tale rotazione i fiumi dell'emisfero settentrionale erodono prevalentemente la riva destra, mentre quelli dell'emisfero meridionale la sinistra.

È ora interessante capire come varia la velocità del flusso lungo una sezione, da riva a riva. Nei tratti rettilinei dell'alveo la velocità massima si ha al centro del fiume. Alla svolta del letto la linea della corrente corrispondente alla velocità massima si sposta verso la riva lontana dal centro della svolta. Questo fenomeno si produce perché è più difficile far ruotare le particelle d'acqua che corrono veloci, poiché in questo caso è necessario creare una maggiore accelerazione centripeta. Tuttavia, laddove la velocità della corrente è maggiore compare anche una circolazione di acqua più marcata, quindi anche una maggior erosione del terreno. Ecco perché nell'alveo del fiume il punto in cui il movimento è più rapido è, in genere, anche il più

profondo. Questa regola è ben nota ai barcaioli.



L'erosione del terreno sulla riva lontana e il suo depositarsi sulla riva vicina, determinano un graduale spostamento di tutto l'alveo del fiume rispetto al centro della svolta e allo stesso tempo l'aumento della sua sinuosità. Cosicché, una piccola sinuosità iniziale formatasi per una ragione casuale (per esempio, a seguito di una frana, la caduta di un albero e così via), con il tempo è destinata ad aumentare. Lo scorrere rettilineo di un fiume lungo una pianura è instabile. Nella Fig. 1.4 è schematizzato il processo di erosione di una riva alla svolta del fiume e il depositarsi di terreno sulle riva opposta. Con il passare degli anni, se questo trasporto di terreno da una riva all'altra diviene cospicuo è possibile che un tratto di fiume venga per così dire isolato dal ramo principale. L'ansa isolata diviene un lago, come è schematicamente mostrato nella terza parte della Fig. 1.4.

1.3 Quali forme hanno i meandri?

L'andamento del letto di un fiume è fortemente determinato dal rilievo del terreno sul quale scorre. In una zona dove il terreno non è uniforme il fiume procede a zig-zag per evitare le asperità e per riempire le depressioni e sceglie la via con la pendenza massima. Come scorrono i fiumi nelle pianure? Come influisce sulla forma dell'alveo l'instabilità dello scorrimento rettilineo del fiume rispetto alle sinuosità? È tale instabilità che causa l'aumento della lunghezza del fiume, così che questi inizia a procedere a meandri? È naturale pensare che in un caso ideale (un sito completamente piano e omogeneo) debba comparire una curva periodica. Qual è la sua forma?

I geologi hanno avanzato l'ipotesi che i letti dei fiumi che scorrono attraverso pianure alle svolte debbano prendere la forma di una particolare

curva. Prendete un'asta d'acciaio e tentate di comprimerla lungo la sua lunghezza (come tentate di diminuire la distanza tra le sue estremità). L'asta si incurva (Fig.1.5).

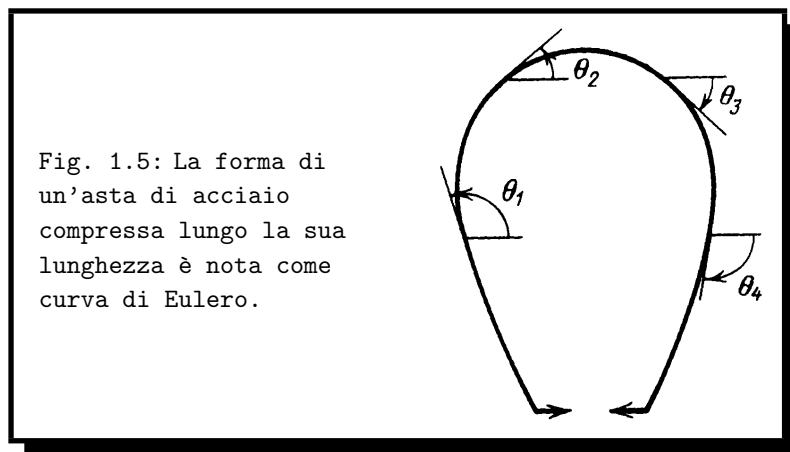
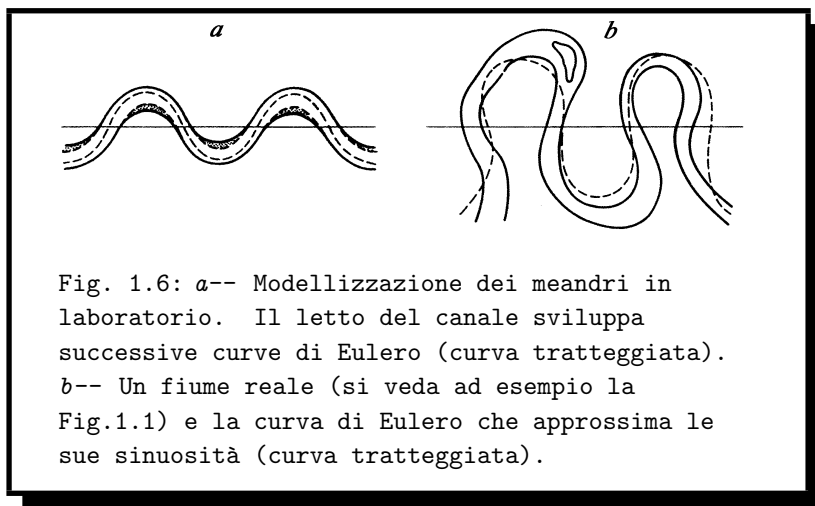


Fig. 1.5: La forma di un'asta di acciaio compressa lungo la sua lunghezza è nota come curva di Eulero.

La curva che descrive matematicamente la forma assunta dall'asta viene chiamata curva di Eulero (dal nome dello scienziato, **Leonhard Euler** (1707-1783), che ha esaminato teoricamente questo problema). Questa curva possiede una particolare proprietà: di tutte le possibili curve a lunghezza assegnata che uniscono le estremità, la curva di Eulero è in media quella a curvatura minore. Se si misurano le deviazioni angolari θ (si veda la figura 1.5), a distanze uguali lungo la lunghezza e si sommano i quadrati di tali deviazioni angolari, lungo la curva di Eulero questa somma ha un minimo. La forma "economica" della curvatura di Eulero è servita da base per lo studio della forma dei letti dei fiumi. I geologi hanno simulato il processo che porta al cambiamento del letto di un fiume in un canale artificiale, sistemato in un mezzo omogeneo costituito da piccole particelle poco connesse e perciò facilmente soggette all'erosione. Molto presto il canale rettilineo iniziava ad incurvarsi e la forma delle svolte appariva descritta dalla curva di Eulero (Fig.1.6).

Ovviamente in situazioni reali non si può osservare una simile perfezione della forma dei letti dei fiumi (per esempio, a causa dell'eterogeneità del terreno). Tuttavia in pianura i fiumi di solito si incurvano e formano una struttura periodica. In figura 1.6 è raffigurato il letto di un fiume reale; la linea tratteggiata indica le curve di Eulero più vicine alla sua forma.



Le svolte periodiche del letto del fiume si chiamano meandri (termine legato al nome greco del fiume turco “Meandro”, noto per le sue anse, oggi Grande Menderes). Si indicano come meandri anche le svolte periodiche delle correnti oceaniche e dei torrenti che si formano sulla superficie dei ghiacciai. In tutti questi fenomeni che si verificano in ambienti omogenei, processi casuali portano alla formazione di una struttura periodica: sebbene le cause che provocano le svolte possano essere diverse, la forma delle curve risultanti è strettamente simile.

1.4 Fiumi dai laghi.

Per quanti fiumi si immettano in un lago, di regola uno soltanto ne sfocia. Il fenomeno può essere spiegato come segue. L’acqua scorre lungo l’alveo più profondo, mentre le altre possibili fonti giacciono a livello più alto della superficie del lago. È poco probabile che i possibili alvei dei fiumi in prossimità della sorgente si trovino alla stessa altezza. Se al lago affluisce molta acqua, da esso possono talvolta uscire anche due emissari. Tuttavia questa situazione è instabile ed è eventualmente possibile osservarla solo in laghi di formazione relativamente recente. Con il tempo il fiume con l’alveo più profondo, in cui la velocità della corrente è maggiore, produrrà maggiore erosione e ciò influirà sull’aumento della portata e sulla riduzione

del livello d'acqua del lago. L'emissione dell'acqua attraverso il fiume meno profondo diminuirà ed esso si coprirà di melma. È per questo che si può dire che “sopravvive” il fiume più profondo tra gli emissari.

Affinché dal lago possano uscire contemporaneamente due fiumi, è necessario che i loro letti si trovino esattamente alla stessa altezza; in questo caso si dice che si ha una *biforcazione* (questo termine oggi è ampiamente utilizzato dai matematici per indicare il raddoppiamento del numero delle soluzioni delle equazioni). Tuttavia in natura la biforcazione è un fenomeno raro e di solito dal lago fuoriesce soltanto un fiume.

Fenomeni analoghi hanno luogo anche nei fiumi. È noto che i fiumi tendono spontaneamente a confluire mentre lo sdoppiamento di un fiume, un chiaro esempio di biforcazione, si osserva relativamente di rado. Un fiume, in ogni luogo, scorre seguendo la curva di inclinazione massima ed è poco probabile che in qualche punto si verifichi uno sdoppiamento di questa curva. Nel delta di un fiume la situazione si può modificare e potete intuire facilmente la ragione.

Capitolo 2

Un interfono oceanico.

In tempi relativamente recenti, circa quarant'anni fa, gli scienziati hanno scoperto un fenomeno sorprendente: onde sonore che si propagavano nell'oceano a migliaia di chilometri dalla loro fonte. In uno degli esperimenti di maggior successo, il rumore di un'esplosione provocata presso le rive australiane ha percorso metà del globo terrestre ed è stato registrato da un altro gruppo di ricercatori presso le isole Bermuda, ad una distanza di 19.600 *km* dall'Australia. Qual' è il meccanismo della propagazione del suono a distanze così grandi?

Per rispondere a questa domanda ricordiamo che la propagazione di onde sonore a “grande” distanza si verifica in qualche misura anche nella vita di tutti i giorni. Infatti, seduti in cucina per la colazione, spesso possiamo sentire un piacevole tintinnio, che un burlone ha chiamato “la canzone dei tubi di conduttura dell'acqua”. A volte si riesce ad interrompere questa “musica condominiale” aprendo il rubinetto dell'acqua nel proprio appartamento. La maggior parte delle persone dopo di ciò ritorna alla colazione interrotta, senza prestare particolare attenzione alla fisica del fenomeno avvenuto. E invece vale la pena di riflettere. Perché il suono, provocato dal getto dell'acqua in un rubinetto in un appartamento collegato da un solo circuito distributore, disturba gli inquilini di tutto lo stabile, anche a distanze più grandi di quelle alle quali si propagherebbe normalmente un suono o un fischio? Questo peculiare canto delle condutture dell'acqua si trasmette invece ovunque, dal primo all'ultimo piano.

Vi sono due cause da tenere in considerazione. La prima è l'azione insonorizzante delle pareti e dei solai che riflettono e assorbono suoni ordinari, mentre l'onda sonora che si propaga nella conduttura dell'acqua passa

da un piano all'altro senza ostacoli. La seconda causa merita particolare attenzione. Nel suono, per esempio dovuto a un fischietto, l'onda acustica si propaga nello spazio in tutte le direzioni e il cosiddetto fronte d'onda ha forma sferica. L'area di questa sfera cresce con il quadrato della distanza dalla sorgente e l'intensità dell'onda sonora, ossia l'energia che attraversa l'unità di area del fronte ondulatorio nell'unità di tempo, diminuisce mano a mano che ci si allontana. Invece l'onda sonora che si forma e si propaga all'interno di una condotta dell'acqua, è "unidimensionale": riflettendo continuamente sulle pareti del tubo essa non si disperde in tutte le direzioni dello spazio, ma si propaga in una sola direzione, lungo l'asse del tubo, senza allargamento del fronte dell'onda. Pertanto l'intensità del suono cambia poco con la distanza dalla sorgente. In questo senso il tubo della condotta dell'acqua rappresenta una guida d'onda acustica: un canale in cui le onde sonore si propagano quasi senza indebolirsi.

Un esempio di guida d'onda acustica è rappresentato da tubi di comunicazione (interfoni) mediante i quali sin dall'antichità e fino a tempi relativamente recenti, sulle navi venivano trasmessi i comandi dal capitano alla sala macchine. Si osservi che lo smorzamento del suono all'interno della guida d'onda è talmente piccolo che, se si riuscisse a fare un tubo lungo 700 km , esso potrebbe servire da particolare "telefono" per la trasmissione di una conversazione, per esempio, da Milano a Napoli. Tuttavia conversare con un tale "telefono" sarebbe piuttosto difficoltoso, poiché l'interlocutore sentirebbe ciò che viene detto oltre mezz'ora dopo.

Sottolineiamo che la riflessione dell'onda che si propaga nella guida lungo la superficie della stessa ha una caratteristica peculiare: l'onda non si propaga in tutte le direzioni ma solo nella direzione assegnata.

Gli esempi riportati fanno pensare che anche la propagazione del suono nell'oceano a enormi distanze sia determinata da un meccanismo analogo. Come si forma questa gigantesca guida d'onda all'interno dell'oceano stesso e che cosa, in questo caso, agisce da superficie-limite riflettente?

Si potrebbe ritenere che come limite superiore possa intervenire la superficie del mare, che riflette abbastanza bene il suono. La relazione tra l'intensità dell'onda sonora riflessa e quella che attraversa l'interfaccia di due mezzi, dipende dalle densità di questi mezzi e dalle velocità del suono in ognuno di essi. Se i mezzi si differenziano di molto (per esempio, per l'acqua e l'aria le densità si differenziano quasi di mille volte, mentre le velocità del suono differiscono all'incirca di 4,5 volte), perfino nel caso in cui l'onda incida perpendicolarmente attraverso l'acqua sulla superficie

dell'aria, praticamente quasi tutta l'onda si rifletterà ritornando all'indietro nell'acqua: l'intensità sonora trasferita all'aria è solo lo 0.01% di quella incidente. In caso di incidenza obliqua, l'onda si rifletterà ancora più marcatamente. Tuttavia la superficie dell'oceano è raramente uniforme a causa della presenza pressoché continua di onde. Ciò porta alla dispersione caotica delle onde sonore e quindi alla distruzione del carattere a guida d'onda della loro propagazione sulla superficie libera del mare.

Le cose non vanno meglio per quanto riguarda la riflessione sul fondo dell'oceano. La densità del fondo in genere si trova nell'intervallo $1.2 - 2.0 \text{ g/cm}^3$ e la velocità di propagazione del suono in essa è del 2–3% inferiore a quella nell'acqua. Per questo, a differenza dell'interfaccia "acqua-aria", una notevole frazione di energia si trasferisce dall'acqua e viene assorbita dal terreno. Pertanto il fondo del mare riflette il suono debolmente e non può costituire il limite inferiore della guida d'onda.

I limiti della guida d'onda nel mare devono così trovarsi da qualche parte, tra il fondo e la superficie. E sono stati individuati. Questi limiti sono strati d'acqua a determinate profondità dell'oceano. Come avviene dunque la riflessione delle onde sonore dalle "pareti" del canale sonoro subacqueo (CSS)? Per rispondere a questa domanda dovremo esaminare come avviene in generale la propagazione del suono nell'oceano.

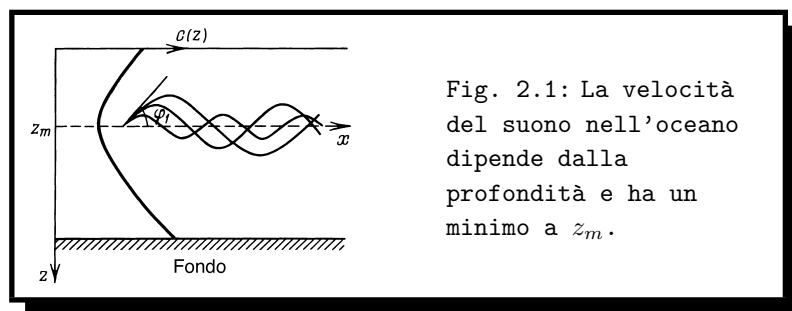
2.1 Propagazione del suono in acqua

È noto che la velocità del suono nel mare oscilla da 1.450 m/s a 1.540 m/s , dipendendo dalla temperatura dell'acqua, dalla salinità, dal valore della pressione idrostatica e da altri fattori. L'aumento della pressione idrostatica $P(z)$ produce un aumento della velocità del suono di circa 1.6 m/s per ogni cento metri di profondità z .

Con l'aumento della temperatura la velocità del suono aumenta. Nell'oceano la temperatura dell'acqua, di regola, diminuisce notevolmente man mano che ci si allontana dagli strati superficiali relativamente caldi verso maggiori profondità, dove si stabilizza ad un valore praticamente costante. Il risultato di questi due meccanismi porta alla dipendenza della velocità del suono $c(z)$ dalla profondità nella forma indicata nella Fig. 2.1.

Infatti, nei pressi della superficie è la rapida riduzione della temperatura ad esercitare l'influenza predominante e in questa zona la velocità del suono diminuisce all'aumentare della profondità. Quanto più aumenta la

profondità, lentamente la temperatura si riduce, mentre la pressione idrostatica continua a crescere. Ad una certa profondità l'influenza di questi due fattori si equilibra e a questo livello la velocità del suono risulta minima. All'ulteriore aumento della profondità corrisponde un aumento della velocità a seguito dell'aumento della pressione idrostatica.



2.2 Propagazione della luce in acqua

Il profilo della velocità di propagazione al variare della profondità influisce sul carattere della propagazione. Per comprendere come avviene la propagazione delle onde acustiche nell'oceano, ricordiamo come si propaga un raggio di luce in un mezzo costituito da una successione di lamelle piane e parallele con diversi indici di rifrazione. Quindi approssimeremo il nostro risultato al caso di un mezzo il cui indice di rifrazione varia in modo graduale $n_0, n_1, \dots, n_k$ e $n_0 < n_1 < \dots < n_k$ (Fig. 2.2).

Il raggio che passa dalla regione superiore alla lamella 1 con angolo di incidenza α_0 dopo la rifrazione forma un angolo α_1 con la normale alla lamella 1; con questo angolo incide sulla lamella 2 e forma un angolo α_2 con la normale alla lamella 3, si rifrange di nuovo e così via. In accordo alla legge di rifrazione si ha

$$\frac{\sin \alpha_0}{\sin \alpha_1} = \frac{n_1}{n_0}, \quad \frac{\sin \alpha_1}{\sin \alpha_2} = \frac{n_2}{n_1}, \dots, \quad \frac{\sin \alpha_{k-1}}{\sin \alpha_k} = \frac{n_k}{n_{k-1}}.$$

Poiché l'indice di rifrazione di un mezzo è inversamente proporzionale alla velocità di propagazione della luce in esso, riscriviamo tutte le relazioni

nella forma

$$\frac{\sin \alpha_0}{\sin \alpha_1} = \frac{c_0}{c_1}, \quad \frac{\sin \alpha_1}{\sin \alpha_2} = \frac{c_1}{c_2}, \dots, \quad \frac{\sin \alpha_{k-1}}{\sin \alpha_k} = \frac{c_{k-1}}{c_k}.$$

Moltiplicando in successione queste uguaglianze si ottiene la relazione

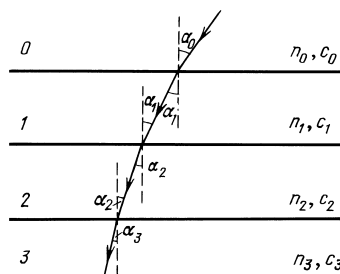
$$\frac{\sin \alpha_0}{\sin \alpha_k} = \frac{c_0}{c_k}.$$

Facendo tendere a zero lo spessore delle lamelle e il loro numero all'infinito, otteniamo la legge di rifrazione generalizzata. Tale legge descrive l'andamento del raggio luminoso nel mezzo con indice di rifrazione lentamente variabile (legge generalizzata di **Willebrord Snell**):

$$c(z) \sin \alpha(0) = c(0) \sin \alpha(z). \quad (2.1)$$

dove $c(0)$ è la velocità della luce nel punto in cui il raggio entra nel mezzo e $c(z)$ è la velocità della luce ad una distanza z . Con il passaggio al limite, la linea spezzata, che indica il percorso del raggio, si trasforma in una curva continua. Quindi nella propagazione di un raggio luminoso in un mezzo otticamente non omogeneo, all'aumentare della velocità della luce (diminuzione dell'indice di rifrazione) il raggio si allontana sempre più dalla verticale e "si stringe" all'interfaccia.

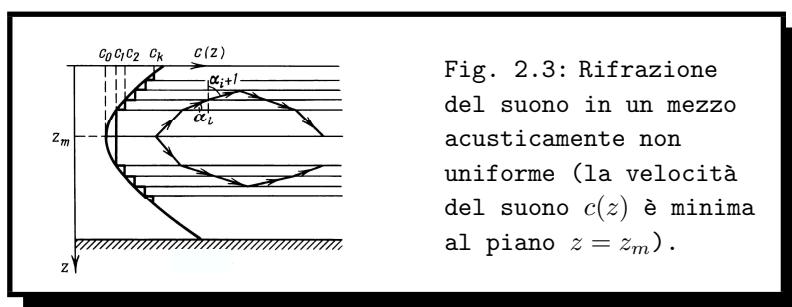
Fig. 2.2: Un mezzo otticamente non omogeneo può essere modellizzato con un con una sovrapposizione di lamine di vetro a indici di rifrazione diversi.



Conoscendo come cambia la velocità della luce in un mezzo, utilizzando la legge generalizzata di Snell, possiamo calcolare l'andamento di qualsiasi raggio che si propaga in un mezzo non omogeneo. In tal modo si può descrivere l'incurvamento delle onde acustiche in un mezzo non omogeneo, dove la velocità in gioco è quella del suono.

2.3 Guida d'onda subacquea

Ora torniamo al problema della propagazione del suono in un canale sonoro subacqueo. Immaginiamo che la sorgente del suono si trovi ad una profondità z_m , corrispondente al minimo della velocità (Fig. 2.3). Come si comportano le onde acustiche? Il raggio che viaggia lungo l'orizzontale $z = z_m$ sarà rettilineo. I raggi che formano un certo angolo rispetto a questa linea incurveranno. Per analogia questo fenomeno si chiama rifrazione del suono. Poiché sia sopra sia sotto il livello z_m la velocità del suono cresce, le onde acustiche incurveranno in direzione dell'orizzontale a $z = z_m$. Ad un certo momento il raggio diventerà “parallelo” a questa orizzontale e dopo essersi “riflesso” ritornerà verso questa stessa orizzontale (Fig. 2.3). Così

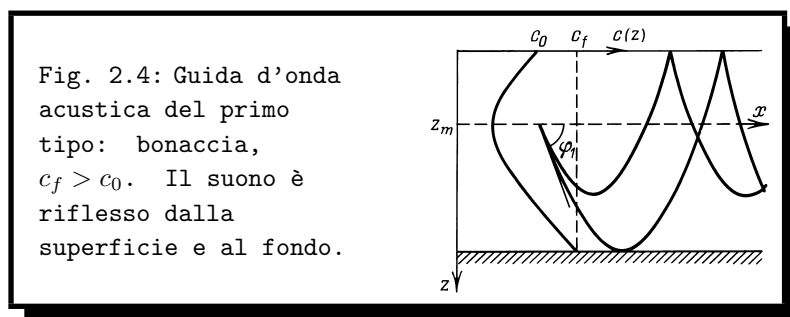


la rifrazione del suono nell'oceano fa sì che una parte dell'energia sonora emessa dalla sorgente possa propagarsi senza uscire in superficie né arrivare sul fondo. Ciò significa che in questo mezzo si realizza un meccanismo di propagazione del suono in guida d'onda: un canale sonoro subacqueo. Il ruolo di “pareti” di questa guida d'onda è svolto dagli strati di acqua alle profondità alle quali avviene l'inversione del raggio sonoro.

Il livello di profondità z_m alla quale la velocità del suono è minima è definito asse del canale sonoro subacqueo. Di solito il valore di z_m è tra i 1.000 e i 1.200 m; tuttavia alle latitudini tropicali, dove l'acqua si riscalda fino a grandi profondità, l'asse del CSS può scendere anche fino a 2.000 m. Al contrario, ad alte latitudini l'effetto della temperatura sulla velocità del suono influisce soltanto sullo strato vicino alla superficie e l'asse del CSS si trova a una profondità di 200 – 500 m. Alle latitudini polari esso arriva ancora più vicino alla superficie.

Nell'oceano possono esistere due diversi tipi di CSS. Il canale del primo

tipo si forma quando la velocità del suono sulla superficie dell'acqua (c_0) è minore che sul fondo (c_f). Questo caso di solito ha luogo nei mari molto profondi, dove la pressione raggiunge centinaia di atmosfere. Come abbiamo già detto, il suono che dall'acqua tenta di passare all'aria, si riflette alla superficie dell'interfaccia e, se la superficie è liscia (bonaccia), essa agisce da limite superiore della guida d'onda ed il canale occupa tutto lo strato dell'acqua, dalla superficie fino al fondo (si veda la figura 2.4).



Vediamo ora quale parte dei raggi sonori viene “imprigionata” nel CSS in questo primo caso. Riscriviamo, per questo, la legge di Snell nella forma

$$c(z) \cos \varphi_1 = c_1 \cos \varphi(z).$$

dove φ_1 e $\varphi(z)$ sono gli angoli formati dal raggio sonoro con il piano orizzontale rispettivamente alle profondità z_1 e z . Questi angoli si chiamano angoli di scivolamento (è manifesto che $\varphi_1 = \frac{\pi}{2} - \alpha_1$, $\varphi(z) = \frac{\pi}{2} - \alpha(z)$).

Se la sorgente del suono si trova sull'asse del CSS, allora $c_1 = c_m$ e il canale cattura i raggi per i quali l'angolo di scivolamento rispetto al fondo è uguale a $\varphi(z) = 0$. Così tutti i raggi emessi dalla sorgente sotto angoli di scivolamento che soddisfano la condizione

$$\cos \varphi_1 \geq \frac{c_m}{c_f},$$

cadono nel CSS (si veda la Fig. 2.4).

In caso di superficie irregolare dell'acqua, le onde acustiche si disperderanno su di essa. I raggi che lasciano la superficie con angoli di scivolamento sufficientemente grandi raggiungeranno il fondo e saranno assorbiti. Tuttavia, anche in questo caso il canale può intrappolare tutte quelle onde acustiche che raggiungono la superficie ondata (Fig. 2.5). In questo caso il

canale si estende dalla superficie fino alla profondità z_k , che viene determinata dalla condizione $c(z_k) = c_0$. Pertanto questo canale intrappola tutte le onde acustiche con angoli di scivolamento

$$\varphi_1 \leq \arccos \frac{c_m}{c_0}.$$

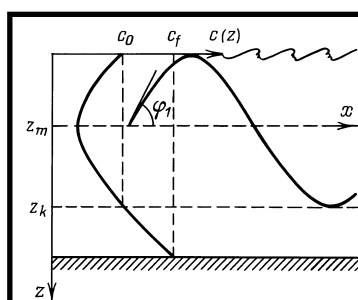


Fig. 2.5: Un'altra guida d'onda acustica del primo tipo: mare mosso, $c_f > c_0$. Il suono è riflesso dalla superficie ma non raggiunge il fondo.

Il canale del secondo tipo è caratteristico delle zone poco profonde e si forma quando la velocità del suono vicino alla superficie è maggiore che sul fondo (Fig. 2.6). Il CSS occupa lo strato d'acqua contenuto tra il fondo e la profondità z_k , tale che $c(z_k) = c_f$. Ossia si tratta quasi di un canale del primo tipo rivoltato.

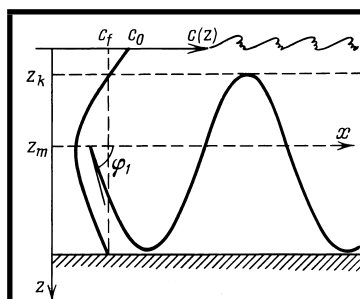
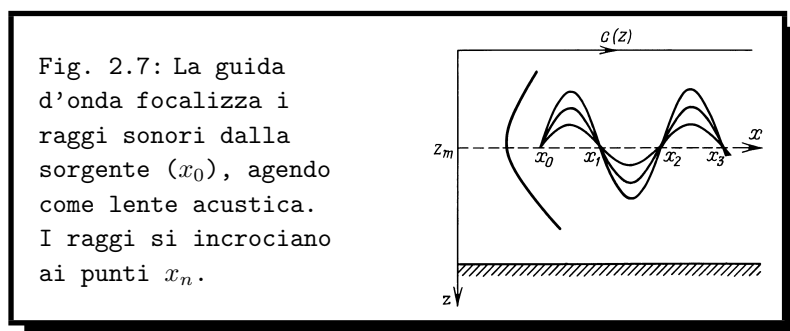


Fig. 2.6: Guida d'onda acustica del secondo tipo. Quando $c_f < c_0$ il suono riflesso dal fondo non raggiunge la superficie.

Se la sorgente del suono si trova vicino all'asse del CSS, nel punto di arrivo del segnale confluiscono, di regola, molte onde acustiche. Il tempo di attraversamento è diverso e risulta massimo per il raggio assiale. La velocità di propagazione del suono a questa profondità è minima. L'intensità di un

segnale costituito da un breve impulso cresce dall'inizio della ricezione alla fine. Infatti la differenza tra i tempi di arrivo degli impulsi per le diverse onde diminuisce ed essi cominciano a sovrapporsi uno all'altro, determinando l'aumento dell'intensità. Per ultima arriva l'onda acustica che si propaga lungo l'asse del CSS (con angolo di scivolamento nullo), dopo di che il segnale si interrompe.

In alcuni casi si verificano particolari dipendenze della velocità del suono a seconda della profondità del CSS, analoghe, per esempio, a quella di una lente: se la sorgente è situata sull'asse del CSS, le onde uscenti sotto diversi angoli convergeranno periodicamente e contemporaneamente sull'asse del canale, in un punto detto fuoco del CSS. Così, nella situazione in cui la velocità del suono varia con la profondità con andamento di tipo parabolico, vale a dire $c(z) = c_m(1 + \frac{1}{2}b^2(z - z_m)^2)$, i fuochi delle onde acustiche emesse sotto angoli piccoli si troveranno nei punti $x_n = x_0 + \pi n / b$, dove $n = 1, 2, \dots$ e b è un coefficiente avente le dimensioni dell'inverso di una lunghezza (m^{-1}) (Fig. 2.7). Il suddetto profilo della curva $c(z)$ è vicino alla distribuzione reale della velocità del suono in CSS profondi. Deviazioni dall'andamento parabolico portano allo spostamento dei fuochi sull'asse del CSS.



2.4 Applicazioni?

È possibile forzare il suono che si propaga lungo il CSS a compiere sott'acqua il percorso di tutto il globo terrestre, fino al punto di partenza? Ciò è impossibile. Come primo ostacolo, si hanno continenti e i notevoli gradienti nella profondità degli oceani. Tuttavia questo non è l'unico motivo. Le onde

acustiche che si propagano nel CSS in realtà si differenziano da quelle che si propagano nelle condutture dell'acqua e dal caso dell'interfono. Come si è già detto, nella propagazione in queste guide d'onda, le onde acustiche hanno carattere unidimensionale e l'area del fronte d'onda è costante a qualsiasi distanza dalla sorgente del suono. Di conseguenza l'intensità del suono, se non si tiene conto delle perdite per dissipazione in forma di calore, è costante su qualsiasi sezione del tubo. Nel CSS, invece, l'onda non si propaga lungo una retta, ma in tutte le direzioni sul piano $z = z_m$. Perciò il fronte d'onda è costituito dalla superficie di un cilindro e l'intensità del suono diminuisce all'allontanarsi dalla sorgente in modo proporzionale a $1/R$, dove R è la distanza dalla sorgente. Semplici calcoli vi consentiranno di esprimere tale dipendenza, a confronto con quella che caratterizza la diminuzione dell'intensità dell'onda acustica sferica in uno spazio tridimensionale.

Un'altra causa di indebolimento del suono è costituita dallo smorzamento dell'onda acustica durante il suo propagarsi. L'energia si trasforma in calore a causa dell'attrito dovuto alla viscosità dell'acqua e a causa di altri processi irreversibili. Oltre a ciò, l'onda acustica si disperde nell'oceano per varie eterogeneità che possono essere particelle in sospensione, bolle d'aria, plancton e anche bolle gassose dovute ai pesci.

Infine osserviamo che il canale subacqueo sonoro descritto non è l'unico esempio di guida d'onda esistente in natura. Per esempio la radiodiffusione a distanza da radiostazioni terrestri è possibile soltanto grazie alla propagazione di onde radio nell'atmosfera attraverso gigantesche guide di onda. In determinati casi nell'atmosfera possono formarsi canali tipo guida d'onda anche per onde elettromagnetiche dello spettro visibile. Possono così prodursi miraggi a grande distanza in virtù dei quali nel centro del deserto si può vedere un lago con una superficie brillante e chiara o ritenere che nel centro dell'oceano appaiano macchie e righe che simulano gli edifici di una città.

Capitolo 3

Nell' azzurro dello spazio.

Gli artisti hanno un acuto spirito d'osservazione e per questo in genere il mondo rappresentato dai pittori realisti nei loro paesaggi è molto ricco e dettagliato e particolari fenomeni o eventi naturali appaiono più evidenti. Tuttavia anche un grande pittore a volte ha bisogno di comprendere più intimamente la complessa essenza dei fenomeni che hanno luogo in natura e che realisticamente egli raffigura sulle tele. In altri casi si può studiare il mondo circostante meglio sui quadri di un buon pittore paesaggista che in natura, perché nel paesaggio è quasi come il pittore fissasse il tempo, rafforzando intuitivamente le cose più sostanziali e tralasciando l'effimero. Guardiamo assieme il quadro di **Arcadi Rylov** "*Nello spazio azzurro*" (la sua riproduzione è sulla seconda pagina della copertina interna). Si vedono bianche nubi e uccelli volteggianti nell'azzurro del cielo. Allo stesso modo, dolcemente, sulle onde turchine dell'oceano scivola una barca a vela. Ammirando questa eccellente tela si ha la sensazione di essere partecipi di una festa della natura. Ora iniziamo a guardare il quadro con gli scrutatori occhi di un fisico.

3.1 Brezze e venti.

Prima di tutto: da dove il pittore ha ritratto il paesaggio? Da uno scoglio sulla riva o dal bordo di una nave? Probabilmente da una nave, poiché in primo piano nel quadro non si vede la risacca, la distribuzione delle onde è simmetrica e non è alterata dalla presenza di una riva vicina. Cerchiamo ora di valutare la velocità del vento. Non siamo i primi a considerare il problema della valutazione della velocità del vento in base all'entità delle onde o ad

altri fenomeni. Già nel 1806, l'ammiraglio inglese **Francis Beaufort** mise a punto una scala in 12 gradi per la valutazione approssimativa della velocità del vento in base alla sua azione su oggetti terrestri e in base al moto ondoso in mare aperto (si veda la Tabella 3.1 alle pagini 36–37). Questa scala è stata riconosciuta dalla Organizzazione Mondiale Meteorologica ed è usata tuttora.

Dal quadro possiamo osservare che il moto ondoso sull'acqua è leggero e sulle creste delle onde si formano occasionalmente dei cavalloni. In base alla scala di Beaufort ciò corrisponde ad un vento debole, con velocità di circa 5 m/s .

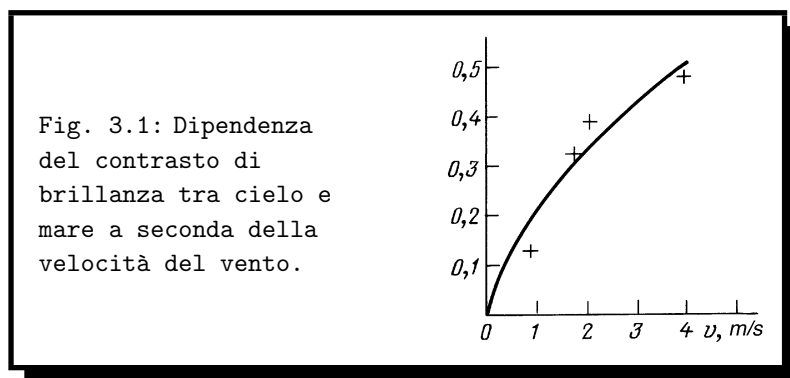
Notiamo che si può valutare la velocità del vento non solo in base alla scala di Beaufort, ma anche dal contrasto tra la luminosità del cielo e del mare. Osservando l'orizzonte nel mare aperto di regola si vede una divisione netta tra mare e cielo. Soltanto in caso di bonaccia la loro luminosità diventa uguale. Il contrasto, in questo caso, sfuma e il mare e il cielo si fondono in un tutt'uno. Questo fenomeno in natura si nota molto raramente, la bonaccia deve essere praticamente assoluta, 0 gradi sulla scala Beaufort. Anche per una gentile brezza sulla superficie del mare comparirebbero le onde.

Il coefficiente di riflessione della luce da una superficie "rugosa" non è uguale all'unità: per questo si verifica un contrasto tra la luminosità del cielo e del mare, che può essere valutato sperimentalmente.

La dipendenza del contrasto tra la luminosità del cielo e quella del mare dalla velocità del vento, è stata misurata durante una delle spedizioni della nave di ricerca scientifica "Dmitri Mendeleev". Nel grafico riportato in Fig. 3.1 le crocette indicano i risultati delle misure, la linea intera indica la dipendenza teorica trovata da **A.V. Bialko** e **V.N. Pelevin**.

Ancora dal quadro, perché i cavalloni sono bianchi e si differenziano così fortemente dal colore azzurro-verde del mare?

Il colore del mare è determinato da molti fattori, i più importanti essendo la posizione del sole, il colore del cielo, il rilievo della superficie e la profondità. Se la profondità del mare non è marcata, la presenza o l'assenza di alghe o di particelle in sospensione diventa anch'essa importante. Tutti questi fattori influiscono sulla riflessione della luce alla superficie, sull'assorbimento e la diffusione in profondità. Pertanto, una spiegazione univoca del perché il mare abbia un determinato colore è difficile. Tuttavia qualcosa possiamo dedurre. Così, per esempio, si può dedurre la ragione per la quale il colore delle onde vicine al pittore è molto più scuro



dello sfondo del mare, mentre all'orizzonte esso diventa più luminoso.

L'entità della riflessione dell'onda luminosa, in caso di incidenza alla superficie di separazione tra due mezzi con caratteristiche ottiche diverse, è determinata dall'angolo di incidenza α e dall'indice di rifrazione. Quantitativamente il coefficiente di riflessione viene definito come rapporto tra l'intensità della luce riflessa e quella incidente. A sua volta l'intensità è data dal valore medio rispetto al tempo del flusso di luce attraverso la unità di superficie dell'area perpendicolare alla direzione di propagazione.

Il coefficiente di riflessione dipende all'angolo di incidenza. Per rilevare questa dipendenza notate come si riflettono i raggi della luce dalla superficie di un tavolo levigato. In questo caso il mezzo otticamente più denso è lo strato trasparente della lacca. Vedrete che in caso di raggi perpendicolari tutto il flusso luminoso viene riflesso mentre diminuendo l'angolo di incidenza una parte sempre maggiore del flusso di luce penetra nel mezzo con maggiore densità ottica e una parte sempre minore si riflette all'interfaccia tra i due mezzi. Dedurrete così con facilità come il coefficiente di riflessione diminuisca con il diminuire dell'angolo di incidenza.

Consideriamo ora la raffigurazione schematica dell'onda, riportata nella figura nella seconda pagina della copertina. Si osservi come gli angoli di incidenza α_1 e α_2 dei raggi che giungono all'occhio di chi guarda l'onda frontalmente o nella regione al di là della cresta sono diversi e $\alpha_2 > \alpha_1$. Per questo dalle regioni lontane del mare giunge all'osservatore più luce riflessa e il fronte anteriore dell'onda risulta più scuro della superficie liscia dietro di essa. In caso di molte onde sulla superficie del mare, l'angolo α , in genere, cambia a seconda se fermiamo la attenzione sulla cresta dell'onda

o sull'avvallamento vicino o lontano dal fronte della stessa. Tuttavia, allontanandosi dal fronte anteriore, la dimensione angolare delle creste scure diminuisce rapidamente, mentre l'angolo α_2 rimane comunque maggiore dell'angolo α_1 . Quando lo sguardo si allontana verso l'orizzonte, l'insieme delle onde è come se venisse mediato e l'osservatore non coglie più gli avvallamenti. Per questo la regione del mare vicina all'orizzonte appare più luminosa che quella in primo piano.

Ora possiamo rendere ragione del perché l'acqua sulla cresta delle onde appare di colore bianco. L'acqua scossa contiene molte bolle d'aria, che si muovono continuamente, cambiano forma e dimensioni e collidono l'una con l'altra. Gli angoli di riflessione cambiano da punto a punto e nel tempo. Per questo sulla schiuma delle creste delle onde i raggi solari si riflettono quasi completamente ed essa ha colore quasi totalmente bianco.

3.2 Il colore del cielo.

Come si è detto, sul colore del mare molto influisce il colore del cielo. Se il primo non è facile da prevedere compiutamente, viceversa il colore del cielo può essere giustificato quantitativamente in base a rigorose leggi fisiche.

È intuitivo che il colore del cielo è determinato dalla diffusione dei raggi solari nell'atmosfera. Perché la diffusione della luce solare, il cui spettro è continuo, ossia contiene tutte le lunghezze d'onda dal rosso al violetto (nel campo del visibile) determina invece il colore celeste e azzurro del cielo, mentre il sole appare piuttosto con tonalità marcatamente di colore giallo?

Per chiarire questo aspetto del mondo fisico, così presente nel nostro quotidiano, ci baseremo sulla legge di Rayleigh per la diffusione della luce.

Nel 1898 il fisico inglese **John Strutt Rayleigh** elaborò la teoria della diffusione della luce da particelle le cui dimensioni siano notevolmente inferiori alla lunghezza d'onda. La legge da lui trovata (che successivamente ha avuto una conferma anche dalle teorie microscopiche basate sulla meccanica quantistica, della quale parleremo nella quarta parte del libro) stabilisce che *l'intensità della luce diffusa è inversamente proporzionale alla quarta potenza della lunghezza d'onda*.

Per spiegare il colore del cielo Rayleigh utilizzò la sua legge nei riguardi della diffusione della luce solare da parte della atmosfera di molecole di azoto e di ossigeno che circonda la terra (per questo a volte si parla di *"legge del cielo azzurro"*).

Cerchiamo di capire qualitativamente il contenuto di questa legge. La luce è costituita da onde elettromagnetiche, le particelle emettenti o assorbenti le onde sono gli atomi o le molecole, con i loro nuclei e i loro elettroni, dotati di carica elettrica. Nel campo di un'onda elettromagnetica, le particelle cariche sono sollecitate a muoversi e si può assumere che il loro movimento avvenga in base alla legge armonica: $x(t) = A_0 \sin \omega t$, dove A_0 è l'ampiezza delle oscillazioni e ω è la frequenza angolare dell'onda luminosa. Con questo movimento le cariche subiscono una accelerazione $a = x''_{tt} = -\omega^2 A_0 \sin \omega t$. Le cariche elettriche che si muovono di moto accelerato diventano esse stesse sorgenti di radiazione elettromagnetica, le cosiddette onde secondarie. L'ampiezza di una onda secondaria è legata all'accelerazione della particella che la emette. (Infatti particelle cariche che si muovono di moto uniforme creano una corrente elettrica, ma non emettono onde elettromagnetiche).

L'intensità di radiazione delle onde secondarie diviene così proporzionale al quadrato dell'ampiezza del moto degli elettroni (il movimento dei nuclei pesanti può essere trascurato) e di conseguenza alla quarta potenza della frequenza ($I \propto a^2 \propto (x''_t)^2 \propto \omega^4$).

Ora, nei riguardi del colore del cielo, si osservi che il rapporto della lunghezza dell'onda di luce rossa e di quella dell'onda blu è all'incirca uguale a $650 \text{ nm} / 450 \text{ nm} = 1.44$ (1 nm (nanometro) = 10^{-9} m). Elevando questo numero alla quarta potenza, otterremo il valore 4.3. Quindi, in accordo alla legge di Rayleigh, l'intensità della luce azzurra diffusa nell'atmosfera è maggiore dell'intensità della luce rossa all'incirca di quattro volte e uno strato d'aria dello spessore di alcune decine di chilometri acquista una colorazione della luce diffusa con forte predominanza di azzurro-blu.

D'altra parte la luce solare che arriva direttamente fino a noi come trasmessa attraverso la "cortina" dell'atmosfera è in gran parte priva delle onde corte del suo spettro. Per questo il sole, che vediamo attraverso i raggi che hanno attraversato l'atmosfera, acquista una debole sfumatura gialla. Questa sfumatura può diventare più marcata e assumere toni tra l'arancio e il rosso nel caso del sole al tramonto, quando cioè i raggi solari devono compiere un maggiore cammino nell'atmosfera e perdono in larga misura le loro componenti violette.

Questa considerazione vi permette di comprendere perchè nella eclissi di luna prodottasi nella notte del 3 Marzo 2007 milioni di persone hanno potuto ammirare il fenomeno della *luna rossa*. Infatti sulla direttrice luna-terra-sole pervenivano sulla prima i raggi solari che avevano percorso un

largo tratto della atmosfera terrestre. D'altra parte, in conseguenza del fenomeno della rifrazione in tale strato di gas (molecole di azoto e di ossigeno) i raggi solari subivano un incurvamento, per così dire "rimediando" all'ombra della notte che si sarebbe prodotta sulla luna in caso di propagazione in linea rigorosamente retta dei raggi di luce provenienti dal sole.

Notiamo che nella legge di Rayleigh si ipotizza che la lunghezza d'onda della luce superi notevolmente le dimensioni delle particelle diffondenti e tuttavia questa dimensione caratteristica non compare nell'espressione dell'intensità. Rayleigh aveva inizialmente ipotizzato che il colore del cielo fosse determinato dalla diffusione della luce solare da parte di particelle piccolissime che vagano nell'atmosfera. Più tardi egli giunse correttamente alla convinzione che i raggi solari sono diffusi dalle molecole stesse del gas che compongono l'aria. Dopo dieci anni, nel 1908, il fisico teorico polacco **Marian Smolukhovskij** avanzò l'idea che i diffusori dovessero in realtà essere rappresentati da oggetti del tutto inaspettati: dalle disomogeneità della densità delle particelle. Basandosi su quest'ipotesi, Smolukhovskij riuscì a spiegare il fenomeno noto già da prima delle sue ricerche, quello dell' *opalescenza critica*, cioè della marcata diffusione della luce in un liquido vicino al punto critico, quando sta trasformandosi con continuità in vapore. Infine **Einstein**, nel 1910, elaborò la teoria quantitativa della diffusione molecolare della luce, basata sull'ipotesi di Smolukhovskij. Per i gas, l'intensità della luce diffusa, calcolata in base alla formula di Einstein, coincide esattamente con il risultato ottenuto precedentemente da Rayleigh. Quindi sembrava che il fenomeno della diffusione fosse compiutamente descritto, almeno in base alla fisica classica. Rimaneva da chiarire quale fosse l'origine della disomogeneità della densità dell'aria.

L'aria si trova in uno stato di equilibrio termodinamico e anche se soffia il vento le disomogeneità legate a questo movimento hanno dimensioni che superano di molte volte la lunghezza d'onda della luce e non possono quindi influire in nessun modo sulla diffusione. Per comprendere la natura delle disomogeneità cerchiamo di capire meglio il concetto di equilibrio termodinamico. Per semplificare ci riferiremo a volumi macroscopici di gas in un recipiente chiuso.

La fisica della materia studia i sistemi costituiti da un numero enorme di particelle, perciò l'unico modo per descriverne le proprietà è quello statistico. L'approccio statistico significa che non si studia il singolo stato di ogni molecola, ma si calcolano i valori medi delle grandezze fisiche di tutto

il sistema in generale. Tuttavia non è affatto necessario che i singoli valori locali o istantanei siano uguali ai valori medi. In una porzione macroscopica di gas è certo più probabile lo stato in corrispondenza al quale le molecole sono distribuite in media con la stessa densità in tutto il volume, ma in virtù dell'agitazione termica delle molecole si ha una probabilità diversa da zero che la concentrazione delle molecole in una data parte del recipiente e durante un certo intervallo di tempo sia inferiore o superiore alla concentrazione media. Teoricamente è anche possibile uno stato in cui tutte le molecole si raccolgano in una metà del recipiente, mentre l'altra metà resti assolutamente vuota. Tuttavia la probabilità che ciò si verifichi è talmente piccola che non vi è alcuna speranza che si realizzi neanche sul tempo di vita dell'Universo, che è calcolato essere all'incirca di 10^{10} anni.

Sono peraltro possibili più piccoli discostamenti delle grandezze fisiche dai valori medi, discostamenti che i fisici chiamano *fluttuazioni*. Non solo tali discostamenti sono possibili, ma si verificano costantemente grazie al movimento termico delle molecole. Le fluttuazioni rispetto ai valori medi fanno sì che in alcune zone la densità del gas aumenti, mentre in altre diminuisce. Ciò influisce sull'indice di rifrazione della luce, determinando così quelle disomogeneità responsabili della diffusione della stessa.

Se ora torniamo allo studio della diffusione della luce solare nella atmosfera, tutte le osservazioni fatte per il volume del gas in un recipiente rimangono valide. Oltre a ciò, l'aria è una miscela di gas diversi e le differenze del moto termico per molecole di gas diversi determina un'ulteriore possibilità di disomogeneità dell'indice di rifrazione legata alle fluttuazioni.

La misura caratteristica della "disomogeneità" dell'indice di rifrazione (o disomogeneità della densità) dipende dalla temperatura. Per gli strati d'atmosfera in cui avviene la fonte maggiore della diffusione della luce solare le dimensioni di queste disomogeneità sono molto minori della lunghezza d'onda della luce visibile, ma superano le dimensioni delle molecole di gas che compongono l'aria. Per questo la diffusione della luce avviene a seguito delle disomogeneità locali piuttosto che direttamente dalle molecole, come aveva ipotizzato Rayleigh.

Perché vediamo il cielo azzurro e non violetto, sebbene la legge di Rayleigh ipotizzi una prevalenza di luce violetta? Ciò è dovuto a due motivi. In primo luogo nella luce solare vi è comunque un minor numero di componenti violette che azzurre. In secondo luogo dobbiamo tenere conto anche del nostro "strumento di registrazione", cioè l'occhio umano. L'acutezza della ricezione visiva dipende fortemente dalla lunghezza d'onda della luce.

Nella figura sulla seconda pagina della copertina interna è riportata la curva sperimentale che caratterizza questa dipendenza. Da essa si vede come l'occhio umano sia maggiormente sensibile alle componenti spettrali sul giallo e quindi reagisca più debolmente ai raggi violetti che a quelli verde-azzurri. Pertanto la luce solare diffusa e percepita dall'occhio umano, in pratica, non ha componente violetta.

3.3 Bianche nubi e loro forma.

Perché nel cielo azzurro vediamo spesso nuvole bianche, in genere comportanti una più marcata diffusione della luce? La giustificazione di questo fatto è un po' più difficile e articolata, in quanto richiede la distinzione tra sorgenti di luce *coerenti* (o pilotate tutte con la stessa fase) e *incoerenti* o indipendenti. Le nuvole sono composte da piccolissime gocce di acqua o di cristalli di ghiaccio, le cui dimensioni superano peraltro notevolmente la lunghezza d'onda della luce visibile. Per questo, per la diffusione della luce solare attraverso le particelle che compongono le nuvole, la legge di Rayleigh va integrata. Le componenti rosse, a lunghezza d'onda maggiore, producono un effetto collettivo di emissione in fase dalle diverse sorgenti che si estende su dimensioni più grandi che non per le componenti azzurre a lunghezza d'onda più piccola. Così si attua una specie di compensazione all'interno della goccia e la diffusione della luce avviene all'incirca con la stessa intensità per tutte le lunghezze d'onda.

Quando ci capiterà di osservare bianche nubi che si rincorrono e si accavallano nel cielo, così come le vediamo nel quadro, ora potremo pensare a questi delicati meccanismi fisici che ci consentono di beneficiare di tanta varietà e bellezza della natura.

Possiamo ora proseguire l'analisi dei fenomeni riportati nel quadro fermando la nostra attenzione sulla forma delle nubi. Nel quadro si vede che sebbene la parte superiore delle nuvole abbia gibbosità (queste nuvole sono dette cumuli) il limite inferiore è ben delineato. Da che cosa dipende? Le nuvole cumuliformi (a differenza di quelle stratiformi) si formano a seguito della ascesa convettiva dalla superficie della terra di strati d'aria calda, saturi d'umidità. Quanto più si sale dal mare (così come dai continenti) tanto più la temperatura scende. È intuitivo che ad altezze molto minori del raggio della terra e della distanza dalla riva più vicina, le superfici a temperatura costante (isoterme) sono piani paralleli alla superficie del mare.

In questo caso, nei pressi della superficie la temperatura diminuisce abbastanza rapidamente, all'incirca $1^\circ C$ ogni $100\ m$ (in genere la dipendenza dalla temperatura dell'aria non è lineare con l'altezza; tuttavia, fino ad altezze di alcuni chilometri, il valore riportato è sostanzialmente corretto).

Torniamo al flusso d'aria in salita. Non appena si raggiunge l'altezza alla quale la temperatura dell'aria corrisponde al punto di condensazione, il vapore acqueo in essa contenuto inizia a condensarsi in piccolissime goccioline. Il limite inferiore della nuvola è definito da questa isoterma. Tenuto conto che le dimensioni della nuvola orizzontalmente sono di centinaia o di migliaia di metri, osserviamo che la spiegazione proposta conduce ad un limite inferiore della nuvola abbastanza netto. Infatti la condensazione del vapore acqueo avviene entro alcune decine di metri in altezza, un tratto molto minore della dimensione orizzontale della nuvola.

A conferma di quanto detto, in secondo piano nel quadro, vediamo un'intera serie di nuvole con limiti inferiori piani alla stessa altezza sul livello del mare.

Dopo la formazione del limite inferiore della nuvola, l'aria non si ferma ma raffreddandosi continua a salire. Il vapore acqueo nella nuvola si condensa trasformandosi non più in gocce d'acqua ma in piccoli cristalli di ghiaccio di cui, di solito, è composto il limite superiore della nuvola cumuliforme. Perdendo la propria umidità e raffreddandosi, l'aria arresta la sua salita e inizia il movimento di ritorno verso il basso. Essa scende sfiorando la nube. A causa dei flussi convettivi descritti, sul limite superiore delle nuvole cumuliformi si formano anche caratteristiche creste. Poiché l'aria che si raffredda scende verso il basso, queste nuvole, di regola, non formano una massa continua ma sono intervallate da tratti di cielo azzurro.

3.4 Uccelli e battiti d'ala.

In primo piano nel quadro si vede uno stormo di uccelli bianchi. Valutiamo la frequenza dei battiti d'ala di un uccello di medie dimensioni (massa $m \approx 10\ kg$, con una superficie alare $S \approx 1\ m^2$), in caso di volo senza planata. Assumiamo una velocità media dell'ala $\bar{v}$. Quindi, in un tempo Δt , con movimento dell'ala verso il basso, l'uccello comunica alla massa d'aria $\Delta m = \rho S \bar{v} \Delta t$ (ρ è la densità dell'aria) un impulso $\Delta \bar{p} = \rho S \bar{v}^2 \Delta t$. Affinché l'uccello si mantenga in volo, quest'impulso deve compensare l'azione della forza di gravità: $\Delta \bar{p} = m g \Delta t$. Pertanto si ha

$$m g = \rho S \bar{v}^2,$$

e quindi, per la velocità media dell'ala, abbiamo $\bar{v} = \sqrt{m g / \rho S}$. Questa velocità può essere collegata alla frequenza ν del battito delle ali e alla lunghezza dell'ala L dalla relazione

$$\bar{v} = \omega L = 2\pi \nu L.$$

Considerando che $L \sim \sqrt{S}$, ricaviamo

$$\nu \approx \frac{1}{2\pi S} \sqrt{\frac{m g}{\rho}} \approx 1 \text{ s}^{-1}. \quad (3.1)$$

Così, in accordo alla nostra valutazione, l'uccello deve compiere all'incirca un battito d'ala al secondo, il che costituisce, come ordine di grandezza, una stima ragionevole.

È interessante studiare più dettagliatamente la formula ottenuta. Facciamo l'ipotesi che tutti gli uccelli, indipendentemente dalle dimensioni e dalla specie, abbiano la stessa forma del corpo. Quindi l'area dell'ala può essere espressa in termini della massa dell'uccello secondo la relazione $S \propto m^{2/3}$. Inserendo questa espressione nell'equazione da noi trovata, la frequenza dei battiti delle ali diviene

$$\nu \propto \frac{1}{m^{1/6}}. \quad (3.2)$$

In questo modo vediamo che diminuendo le dimensioni dell'uccello, la frequenza dei suoi battiti d'ala cresce, come del resto ci sarà capitato di osservare. L'ipotesi di una forma unica per tutti gli uccelli è ovviamente approssimativa: gli uccelli grandi hanno anche ali più grandi di quelle degli uccelli più piccini. Questo fatto, peraltro, rafforza la conclusione alla quale siamo pervenuti.

Notiamo che la formula da noi ricavata (senza il fattore 2π) poteva facilmente essere ottenuta anche con il metodo chiamato dai fisici delle *leggi di scala*. È intuitivo che la frequenza dei battiti debba dipendere soltanto dal peso dell'uccello, dall'area delle sue ali S e dalla densità dell'aria circostante ρ . Scrivendo $\nu = \rho^\alpha S^\beta (m g)^\gamma$ e confrontando le dimensioni

delle grandezze fisiche a sinistra e a destra di questa uguaglianza, troviamo che $\alpha = -\gamma = -1/2$ and $\beta = -1$, ossia

$$\nu \sim \frac{1}{S} \sqrt{\frac{mg}{\rho}}.$$

Non pensiamo di avere esaurito l'analisi di tutti i fenomeni naturali, né di aver dato tutte le risposte possibili, nei riguardi quanto avviene nello “*spazio azzurro del cielo*”. Ad un osservatore accorto, anche nello stesso quadro che ci ha dato lo spunto di avvio, possono rivelarsi altri fenomeni, perfino più interessanti. E non vi è la necessità di limitarsi al quadro. Guardatevi attorno e osservate il mondo che vi circonda, magari in una mattina in cui non vi sentite del tutto piacevolmente “rinati” alla vita. Scoprirete altri eventi e fenomeni e sorgeranno in voi molti interrogativi, che arricchiranno la vostra giornata e il vostro spirito.

Tabella 3.1: SCALA DI BEAUFORT.

Gradi di Beaufort	Definizione della forza del vento	Velocità media del vento ad un'altezza di 10 m		Azione del vento
		(nodi)	(miglia/ora)	
0	Bonaccia	0-1	< 1.15	Mare liscio a specchio.
1	Calmo	1-3	1.2-3.5	Sul mare un leggero increspamento; sulle creste assenza di schiuma. Altezza dell'onda fino a 0.1 m.
2	Leggero	4-6	4.6-6.9	Sul mare onde corte; altezza massima fino a 0.3 m.
3	Debole	7-10	8.0-12	Ondosità leggera sull'acqua: raramente si formano piccoli cavalloni. Altezza delle onde fino a 0.6 m.
4	Moderato	11-16	13-18	Cavalloni bianchi sul mare si vedono in molte zone. Altezza massima delle onde fino a 1.5 m.
5	Freschetto	17-21	20-24	Ovunque sono visibili cavalloni bianchi. Altezza media delle onde 2 m.
6	Intenso	22-27	25-31	Creste bianche e spumose occupano un'area notevole; si forma polvere d'acqua. Altezza massima delle onde 4 m.
7	Forte	28-33	32-38	Le creste delle onde sono rotte dal vento. Altezza massima delle onde 5.5 m.
8	Molto forte	34-40	39-46	Ondosità del mare molto forte. Altezza massima delle onde fino a 7.5 m.

Tabella 3.1: SCALA DI BEAUFORT (*continuazione*).

Gradi di Beaufort	Definizione della forza del vento	Velocità media del vento ad un'altezza di 10 m		Azione del vento
		(nodi)	(miglia/ora)	
9	Tempesta	41–47	47–54	Ondosità del mare molto forte. Altezza massima delle onde fino a 10 m.
10	Forte tempesta	48–55	55–63	La superficie del mare é bianca per la schiuma. Forte fracasso simile a colpi. Onde molto alte, fino a 12.5 m.
11	Tempesta violenta	56–65	64–75	Sul mare solo onde alte fino a 16 m. Grandi navi a tratti scompaiono dalla vista.
> 12	Uragano	> 65	> 75	L'aria é piena di vapori e di schiuma, il mare é completamente bianco; la visibilità é marcatamente ridotta.

Capitolo 4

Il pendolo di Foucault e la legge di Beer.

Nella fotografia riportata nella figura 4.1 è mostrato l'interno del Pantheon di Parigi con un ingrandimento dell'area di oscillazione del pendolo lanciato da **Jean Bernard Foucault** nel 1851. L'oscillazione del pendolo ci consentirà di discutere di alcuni importanti aspetti del moto terrestre e delle sue conseguenze.



Fig. 4.1: Interno del Pantheon di Parigi ove fu lanciato da Foucault il famoso pendolo.

Oltre al moto di oscillazione, mentre il pendolo continua ad oscillare

si ha anche una rotazione del piano delle oscillazioni. Questa osservazione in realtà era stata fatta per la prima volta da un italiano, **Vincenzo Viviani**, nel 1661 in Firenze. Lo scienziato francese Foucault ripeté l'esperimento, dandone anche la giustificazione teorica che Viviani non aveva potuto formulare per la mancanza degli strumenti matematici a quel tempo. L'esperimento di Foucault si tenne nell'enorme sala del Pantheon, la massa della sfera era di 28 kg e la lunghezza della corda di 67 m. Da allora questo tipo di pendolo è chiamato pendolo di Foucault. Il periodo di rotazione del piano di oscillazione del pendolo è di 24 ore al polo. Alle nostre latitudini, tenendo conto della componente della forza di Coriolis, come chiariremo nel seguito, il piano di oscillazione del pendolo descrive circa 10 gradi all'ora. In generale tale periodo risulta $T = 24 \text{ ore} / \sin \alpha$, dove α è l'angolo di latitudine.

Come spiegare il movimento congiunto di oscillazione e rotazione del pendolo?

Se in un sistema di riferimento terrestre si realizzassero le più semplici leggi di Newton, come sappiamo dai corsi elementari di fisica, il pendolo oscillerebbe in un piano che si manterrebbe costante nel tempo. In realtà nel sistema di riferimento terrestre, le leggi di Newton devono essere corrette, e ciò è possibile introducendo speciali forze, *le forze d'inerzia*.

4.1 Pseudo-forze nei sistemi ruotanti.

Le forze d'inerzia, o pseudo-forze, devono essere introdotte in qualsiasi sistema di riferimento accelerato rispetto al Sole (più precisamente rispetto al cosiddetto sistema delle stelle fisse). Questi sistemi sono anche detti non-inerziali, a differenza dei sistemi inerziali che si muovono rispetto al Sole e alle stelle fisse di moto rettilineo uniforme.

La Terra, a rigor di termini non è un *sistema inerziale*, poiché essa ruota intorno al Sole e al proprio asse. Tuttavia, poiché le accelerazioni legate a questi movimenti possono essere in genere trascurate, sulla terra vengono in genere essere utilizzate le leggi di Newton in forma diretta.

La rotazione del pendolo di Foucault si spiega grazie all'azione di una particolare pseudo-forza caratteristica dei sistemi non-inerziali, la forza di **Gaspard Gustave Coriolis**. Parliamone più dettagliatamente.

Consideriamo un esempio di sistema rotante in cui si manifestano chiaramente le forze d'inerzia. Immaginate che un uomo si trovi su una giostra

(indicheremo la velocità angolare con ω e il raggio con r). Supponiamo che egli voglia muoversi da una sedia all'altra (sfidando i pericoli che egli corre in tal modo e anche le disposizioni di legge!) (vedi Fig. 4.2). Per far ciò egli si muove, nel sistema di riferimento della giostra, lungo una circonferenza, con una certa velocità v_0 , per esempio nel senso della rotazione.



Fig. 4.2: Forze di inerzia in un sistema di riferimento ruotante.

Esaminiamo inizialmente il movimento dell'uomo in un sistema di riferimento fisso. La velocità totale è la somma della velocità periferica della giostra e della velocità lineare:

$$v = \omega r + v_0.$$

L'accelerazione centripeta viene definita dalla relazione

$$a_{cp} = \frac{v^2}{r} = \frac{v_0^2}{r} + \omega^2 r + 2 v_0 \omega.$$

In base alla seconda legge di Newton avremo

$$m a_{cp} = Q,$$

dove Q è la componente della forza di reazione agente sull'uomo da parte della sedia della giostra.

Ora esaminiamo questo stesso movimento nel sistema di riferimento (ruotante) della giostra. In esso la velocità è uguale a v_0 e l'accelerazione

centripeta è $a'_{cp} = \frac{v_0^2}{r}$. Utilizzando le due relazioni si può scrivere

$$m a'_{cp} = \frac{m v_0^2}{r} = Q - m \omega^2 r - 2m v_0 \omega.$$

Come si vede, se vogliamo utilizzare le legge di Newton anche nel sistema rotante è necessario introdurre la forza d'inerzia

$$F_{in} = -(m \omega^2 r + 2m v_0 \omega) = -(F_{cf} + F_{Cor}),$$

dove il segno meno indica che questa forza è diretta verso il centro di rotazione. Nel sistema di riferimento non-inerziale l'equazione di moto sarebbe

$$m a'_{cp} = Q + F_{in} = Q - (F_{cf} + F_{Cor}).$$

È come se la forza d'inerzia spingesse l'uomo dal centro verso l'esterno, mentre questi gira sulla giostra. L'espressione "come se" non è qui a caso. Nel sistema ruotante non compaiono in realtà nuove forze di interazione tra i corpi. Come in precedenza, da parte della sedia agisce sull'uomo la stessa forza di reazione, che ha la stessa componente Q , nel piano orizzontale diretta verso il centro di rotazione. Tuttavia se la forza Q determina un'accelerazione centripeta totale a_{cp} , nel sistema rotante l'accelerazione diminuirebbe. Per questo è necessario introdurre la forza d'inerzia F_{in} che compensa parzialmente la forza Q .

Nel nostro caso la forza d'inerzia è composta da due termini, corrispondenti ai due addendi nell'espressione di F_{in} . Il primo è costituito dalla forza centrifuga d'inerzia F_{cf} . Tale forza è tanto più grande tanto più veloce è la rotazione e tanto più lontano dal centro si trova il corpo. La seconda forza si chiama forza di Coriolis F_{Cor} (dal nome dello scienziato francese che l'ha considerata per primo). Questa forza deve essere introdotta soltanto quando il moto del corpo è descritto in un sistema rotante. Essa non dipende dalla posizione del corpo, ma dalla velocità del moto e dalla velocità angolare di rotazione del sistema di riferimento.

Se il corpo in un sistema rotante si muove non lungo la circonferenza, ma, per esempio, lungo il raggio (si veda la Fig. 4.2) anche in questo caso sarà necessario introdurre la forza di Coriolis. Essa però non sarà diretta lungo il raggio, ma perpendicolarmente ad esso. Comunque, con qualsiasi movimento in un sistema rotante, la forza di Coriolis è diretta perpendicolarmente all'asse di rotazione e alla velocità del corpo. Desta meraviglia che

in caso di movimento in un sistema rotante la forza d'inerzia non soltanto allontana il corpo dal centro ma anche agisca come se lo spingesse tangenzialmente.

Occorre sottolineare che l'origine della forza di Coriolis è la stessa di tutte le pseudo-forze: non è legata all'interazione diretta dei corpi, ma solo alla natura non-inerziale del sistema di riferimento.

Ecco un esempio chiarificatore. Immaginate che al Polo sia sistemato un cannone che spara lungo un meridiano (il Polo è stato preso per semplicità). Il bersaglio si trova sullo stesso meridiano. Arriverà il proiettile a colpire il bersaglio? Se si guarda la traiettoria del proiettile lateralmente, utilizzando un sistema di riferimento inerziale (relativo al sole) la situazione è chiara: la traiettoria del proiettile giace nel piano iniziale del meridiano, mentre la Terra ruota. Per questo il proiettile non colpirà mai il bersaglio. Come spiegare lo stesso fenomeno nel sistema di riferimento relativo alla terra? Come spiegare l'effetto dell'allontanarsi del bersaglio dal piano di traiettoria del proiettile?

Per spiegarlo occorre introdurre la forza di Coriolis, diretta perpendicolarmente alla velocità del corpo e all'asse di rotazione. Ciò fatto anche nel sistema di riferimento terrestre diventa manifesta la ragione per la quale il proiettile è spinto fuori dal piano del meridiano e non colpisce il bersaglio.

In modo analogo si spiega la rotazione del piano delle oscillazioni del pendolo di cui parlavamo all'inizio di questa sezione. Nel sistema inerziale del sole il piano delle oscillazioni del pendolo rimane invariato mentre la Terra ruota. Relativamente alla Terra, il piano delle oscillazioni invece ruota (si veda l'illustrazione in Fig. 4.3).

Se per maggiore semplicità, immaginiamo che il pendolo oscilli al Polo, il piano delle oscillazioni compie un giro completo in un giorno, come si è detto all'inizio di questo capitolo. Nel sistema di riferimento terrestre il fenomeno si può spiegare appunto introducendo la forza di Coriolis.

4.2 Conseguenze sulla terra.

La forza di Coriolis associata alla rotazione della Terra produce una serie di importanti effetti. Prima di parlare di questi, esaminiamo più dettagliatamente il problema della direzione della forza di Coriolis. È stato già detto che essa è perpendicolare all'asse di rotazione e alla velocità del corpo. In linea di principio sono possibili due direzioni della forza, indicate nella

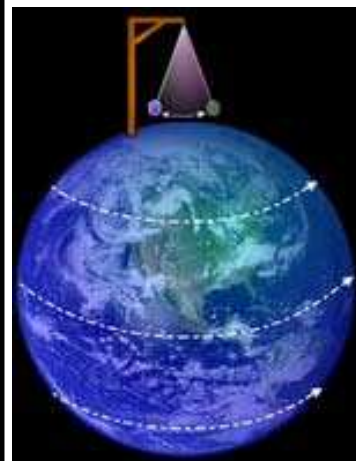


Fig. 4.3: Un pendolo oscillante al polo mantiene il suo piano di oscillazione rispetto alle stelle fisse, mentre la terra ruota. In un sistema di riferimento solidale con la terra l'osservatore vedrebbe il pendolo compiere una rotazione su 360 gradi in ventiquattro ore.

Fig. 4.4. Ricordiamo che una situazione analoga si ha nella definizione della forza di **Lorentz** che agisce su una carica in movimento in un campo magnetico. Come è noto dalla fisica dei corsi liceali, questa forza è perpendicolare alla velocità del corpo e alla direzione del campo magnetico. Tuttavia, per definire univocamente la sua direzione, bisogna utilizzare *la regola della mano sinistra*.

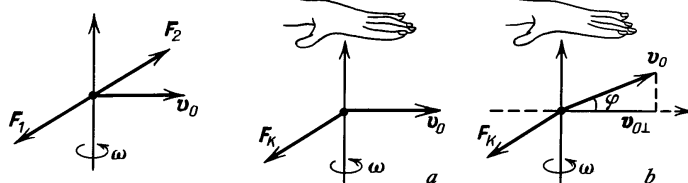


Fig. 4.4: Le due possibilità per la direzione della forza di Coriolis. La convenzione fissa la direzione con la regola della mano sinistra.

La direzione della forza di Coriolis può essere dedotta per mezzo di una

regola analoga. Ciò é illustrato nella Fig. 4.4, *a*.

Prima di tutto scegliamo una determinata direzione dell'asse di rotazione: se si guarda al corpo rotante in questa direzione, la rotazione deve avvenire in senso orario. Ora poniamo la mano sinistra in modo che la direzione delle quattro dita allungate coincida con la direzione della velocità e la direzione dell'asse di rotazione trafori idealmente il palmo. Il pollice ripiegato con un angolo di 90° mostra la direzione della forza di Coriolis. Le due possibilità di definizione della direzione della forza di Coriolis o della forza di Lorentz corrispondono a due tipi di simmetria che s'incontrano in natura: simmetria sinistra e simmetria destra. Per indicare il tipo di simmetria appropriata, bisogna riferirsi ad un "modello", come quello della mano. Naturalmente in realtà non vi è nessun speciale ruolo della mano sinistra. Semplicemente, in questo modo si possono formulare le regole per trovare la direzione della forza.

Così abbiamo esaminato dettagliatamente il problema della forza di Coriolis per il caso in cui la velocità del corpo in un sistema rotante è perpendicolare all'asse di rotazione. In questo caso il modulo della forza è uguale a $2m\omega v_0$ mentre la direzione è definita dalla regola della mano sinistra. Cosa succede in un caso generale?

Se la velocità del corpo v_0 forma un angolo arbitrario con l'asse di rotazione (Fig. 4.4, *b*), nel valutare la forza di Coriolis bisogna tener conto soltanto della proiezione della velocità sul piano perpendicolare all'asse di rotazione. Quindi l'entità della forza di Coriolis è data dall'espressione

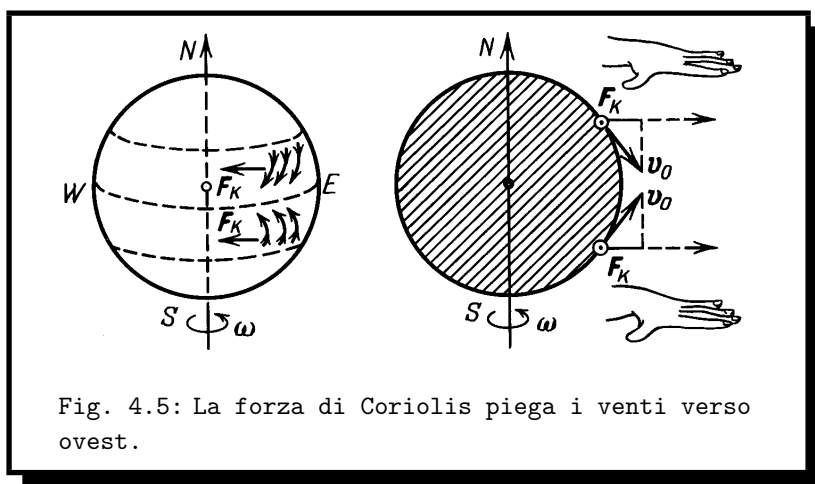
$$F_{\text{Cor}} = 2m\omega v_{\perp} = 2m\omega v_0 \cos \phi.$$

La direzione di questa forza è definita dalla stessa regola della mano sinistra, ma in questo caso le quattro dita allungate non devono essere dirette nella direzione della velocità del corpo, ma lungo la perpendicolare all'asse di rotazione della componente della velocità, come è dimostrato nella Fig. 4.4, *b*.

Ora conosciamo quanto è necessario sulla forza di Coriolis: sia come trovare il suo modulo, sia come definire la direzione. Forti di queste conoscenze, passiamo alla spiegazione di altri effetti interessanti.

È noto che i venti che spirano dai tropici all'equatore piegano ad ovest. Questo effetto è spiegato dalla Fig. 4.5. Inizialmente convincetevi che ciò accade per l'emisfero settentrionale, dove i venti spirano da nord a sud. Mettete la mano sinistra in modo che l'asse di rotazione della Terra en-

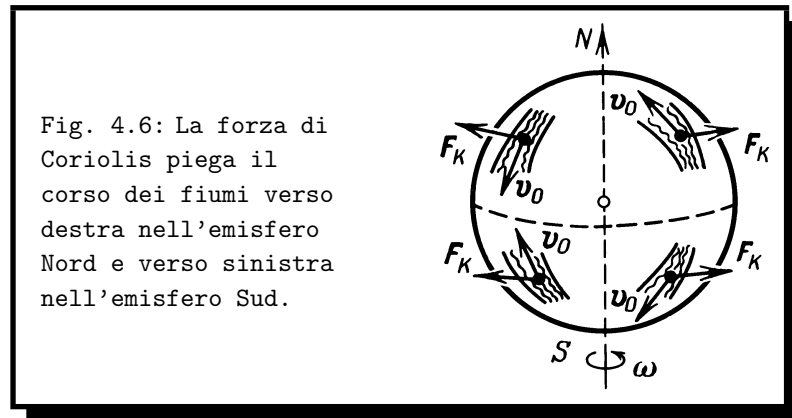
tri nel palmo e dirigete le quattro dita allungate lungo la perpendicolare all'asse di rotazione. Vedrete che la forza di Coriolis è diretta verso di voi perpendicolarmente al disegno e, di conseguenza, ad ovest. Nell'emisfero meridionale i venti spirano al contrario, da sud a nord. Ma né la direzione dell'asse di rotazione, né la direzione della perpendicolare a esso cambiano, di conseguenza non cambia neanche la direzione della forza di Coriolis. Quindi questa forza, in ambedue i casi, è rivolta ad ovest.



La figura 4.6 illustra la legge di **Beer**: nei fiumi dell'emisfero settentrionale la riva destra è più aspra ed erosa della sinistra (nell'emisfero meridionale succede il contrario). In questo caso l'azione della forza di Coriolis porta l'acqua verso la riva destra. A causa dell'attrito la velocità della corrente è maggiore sulla superficie che sul fondo; conseguentemente è maggiore anche la forza di Coriolis. Ne consegue una circolazione dell'acqua e il terreno della riva destra viene eroso mentre sulla riva sinistra vi si deposita. Questo fenomeno è analogo all'erosione della riva alla svolta di un fiume, di cui si è parlato a proposito dei meandri dei fiumi (vedi capitolo 1).

Si potrà notare questo effetto se il fiume scorre lungo un parallelo? Che succede se il fiume taglia l'equatore? Cercate di meditare da soli su tali situazioni, così come sul fatto che la forza di Coriolis porta alla deviazione verso est dei corpi in caduta libera.

Nel 1833 il fisico tedesco **Ferdinand Raich** ha eseguito degli esperimenti molto precisi nella miniera di Friburgo e ha constatato che con la



caduta libera di corpi da un'altezza di 158 m la loro deviazione era in media (in base a 106 esperimenti) di 28.3 mm . Questa é stata una delle prime dimostrazioni sperimentali della teoria di Coriolis.

Capitolo 5

Le maree e il freno lunare.

Sino dai tempi più remoti gli uomini hanno osservato con meraviglia e stupore e subito il fascino delle maree. In alcune regioni il fenomeno è particolarmente marcato. Sulle coste della Normandia, nell'area di Mont Saint Michel, distese aree vengono giornalmente invase e svuotate dall'acqua di mare. Dopo una bassa marea, al momento che il flusso del mare cambia verso, non è difficile vedere come un impetuoso fiume ritorni verso la riva e ricopra campi a perdita d'occhio, con qualche pericolo per turisti e osservatori disattenti che troppo si siano inoltrati nelle zone che la bassa marea aveva precedentemente portato all'asciutto. La variazione dell'altezza dell'acqua di mare varia da punto a punto di quella regione ma può arrivare anche a 5 metri. Un fenomeno naturale di grande bellezza e di impressionante imponenza (si vedano le fotografie in figura 5.1).



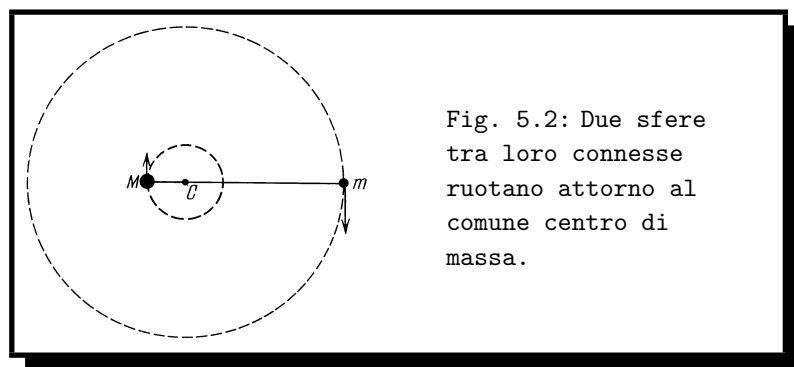
Fig. 5.1: Scorci di alta e bassa marea nei pressi di Mont Saint Michel.

Sino dalle prime loro stupite osservazioni i nostri antenati correttamente riconobbero nella Luna la causa delle maree. La Luna attira l'acqua degli oceani e, un po' ingenuamente, si potrebbe da ciò dedurre che nell'oceano si possa formare una "gobba" d'acqua. Questa gobba rimane di fronte alla Luna, mentre la Terra ruota intorno al suo asse. Raggiungendo la riva, l'acqua che si alza provoca l'alta marea e, allontanandosi, la bassa marea. Questa ingenua interpretazione conduce a una inconsistenza: se così fosse le alte maree dovrebbero osservarsi ogni 24 ore, mentre invece esse si ripetono ogni 12 ore.

La prima giustificazione corretta delle maree è stata formulata da **Newton**, subito dopo la sua scoperta della gravitazione universale.

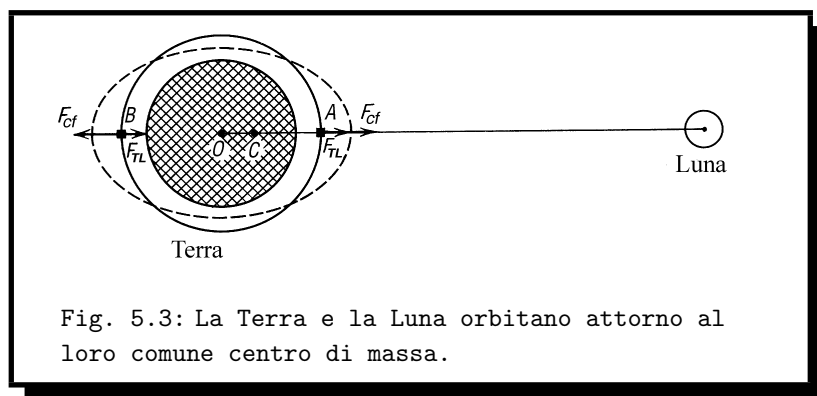
Per la discussione del fenomeno delle maree utilizziamo il concetto di forza d'inerzia. Come abbiamo già detto nel precedente capitolo, anche nel sistema di riferimento ruotante si possono usare le leggi di Newton, purché alle forze di interazione tra i corpi si aggiungano le forze d'inerzia.

La Terra ruota intorno al suo asse, intorno al Sole e intorno ... alla Luna. Sebbene di solito ci si dimentichi di questo, proprio ciò permette di fornire una corretta giustificazione delle maree. Immaginate due piccole sfere, una pesante e una leggera, unite da un filo, che si trovano su di un piano orizzontale liscio (Fig. 5.2).



Queste sferette ruotano insieme, su se stesse, e attorno a un centro comune che è il centro di massa. Alla stessa maniera la Terra e la Luna, attraendosi l'un l'altra, in base alla legge di gravitazione universale, ruotano nello spazio attorno al centro delle due loro masse (si veda la Fig. 5.3).

In ragione della grande massa della Terra, questo centro si trova all'inter-



no del globo terrestre, ma è purtuttavia spostato rispetto al centro O . Le velocità angolari ω della Terra e della Luna intorno al punto C sono evidentemente uguali.

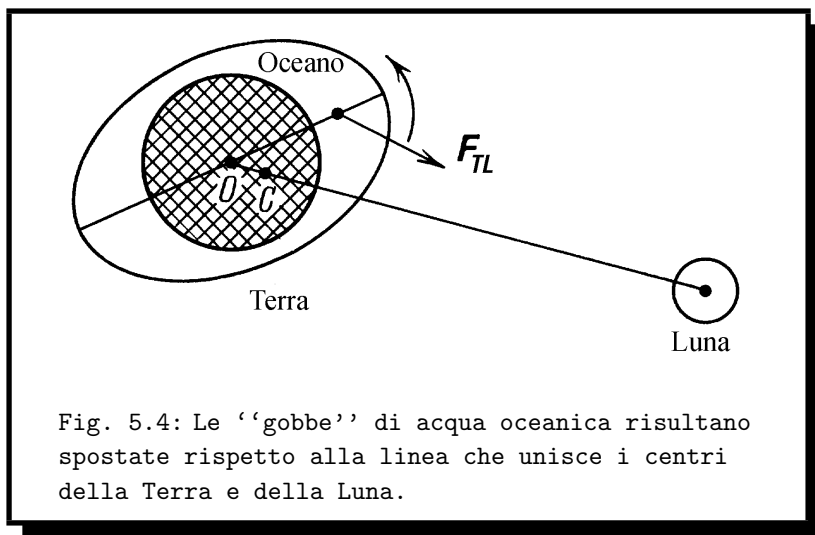
Passiamo ora al sistema di riferimento che ruota con una velocità angolare ω . In esso, poiché questo sistema non è inerziale, su ogni elemento della massa agisce non soltanto la forza di gravitazione ma anche la forza centrifuga che è tanto maggiore quanto più la particella è lontana dal centro di rotazione. Per semplicità immaginiamo che l'acqua copra con uno strato uniforme tutta la superficie della Terra. Può questo strato d'acqua trovarsi in equilibrio? Evidentemente no. La Luna turberà questo equilibrio, poiché compariranno forze aggiuntive di attrazione e forze centrifughe d'inerzia. Sulla superficie rivolta alla Luna la forza gravitazionale verso la Luna e la forza centrifuga si sommano e pertanto si forma una gobba d'acqua A (Fig. 5.3). Sulla superficie lontana dell'acqua si verifica una situazione analoga. Qui aumenta la forza d'inerzia (cresce la distanza dal centro di rotazione C), mentre la forza gravitazionale verso la Luna diminuisce. La risultante di queste forze è ancora sulla direttrice verso il centro della Terra e determina la formazione di un'altra gobba B . All'equilibrio corrisponde l'immagine mostrata dalla linea tratteggiata nella Fig. 5.3.

Questa descrizione del fenomeno delle maree è molto semplificata. In particolare non tiene conto della non uniformità della distribuzione dell'acqua degli oceani sulla superficie della Terra, dell'effetto delle forze gravitazionali verso il Sole e di molti altri fattori, capaci di alterare il quadro descritto. Comunque la nostra spiegazione fornisce la risposta alla domanda più rilevante. Poiché le gobbe sono praticamente immobili rispetto al sis-

tema delle stelle fisse e la Terra ruota intorno al suo asse, le alte maree e le basse maree devono essere rilevate per due volte, nel corso delle 24 ore.

Ed è ora il momento opportuno per spiegare il principio d'azione del freno lunare. In realtà le gobbe non si trovano esattamente sulla linea che unisce il centro della Terra e quello della Luna (per semplificare è stato mostrato in tal modo in Fig. 5.3) ma sono un po' spostate (Fig. 5.4). Ciò si verifica per la seguente ragione. La Terra, ruotando intorno al proprio asse, a seguito dell'attrito trascina con sé l'acqua dell'oceano. Per questo, nella formazione delle gobbe delle maree devono essere attratte ogni volta nuove masse d'acqua. La deformazione, tuttavia, ritarda rispetto alla forza che la provoca (la forza crea un'accelerazione e deve passare un certo tempo perché le particelle acquistino velocità e si spostino di una certa distanza). Per questo il punto di massimo sollevamento dell'acqua (cima della gobba) e il punto sulla linea dei centri, dove sull'acqua agisce la massima forza di gravità verso la Luna, non coincidono.

La gobba si forma con un certo ritardo e si forma in direzione della rotazione della Terra. In questo caso, come si vede nella figura, la forza di gravità della Luna non è più diretta verso il centro della Terra e crea un momento che frena la sua rotazione. La durata dei giorni aumenta ogni giorno! Il primo a capirlo fu il grande fisico inglese **Lord Kelvin**.



Il freno lunare funziona senza guasti già da molti millenni ed è capace di variare sensibilmente la durata del giorno terrestre. Gli studiosi hanno notato in coralli pietrificati, che si trovano nell'oceano da 400 milioni di anni, delle strutture chiamate anelli "giornalieri". Quando sono stati contati gli anelli giornalieri è risultato che l'anno solare, corrispondente al periodo in cui la Terra compie un giro completo intorno al Sole (questo evidentemente non è cambiato) era inizialmente di 395 giorni. Significa quindi che un giorno a quel tempo durava 22 ore!

Il freno lunare continua a funzionare anche oggi giorno, aumentando la durata dei giorni. Poiché il tempo di rotazione della Terra intorno al suo asse è confrontabile con quello di rotazione della Luna intorno alla Terra verrà un giorno che questa azione frenante arresterà la rotazione. La Terra volgerà verso la Luna sempre la stessa parte, alla stessa maniera in cui la Luna è rivolta verso la Terra solo da un lato. La durata dei giorni sulla Terra aumenterà progressivamente e, a seguito di questo, potrà cambiare il clima. Infatti, quando la Terra rimarrà rivolta al Sole con una sola faccia, dall'altra parte, in corrispondenza, ci sarà una lunga notte. L'aria calda dalla parte riscaldata tenderà con enorme velocità a portarsi verso la parte fredda, sulla Terra cominceranno a spirare forti venti, si alzeranno tempeste di sabbia.

Tuttavia ciò non succederà molto presto e i fisici del futuro escogiteranno certamente qualcosa per evitare questa catastrofe.

Capitolo 6

Bolle e gocce.

Le manifestazioni della tensione superficiale di un liquido, in natura e nella tecnica, sono sorprendentemente varie. La tensione superficiale fa sì che l'acqua tenda ad assumere la forma di goccia; grazie alla tensione superficiale si possono produrre bolle di sapone e si può scrivere con la penna. La tensione superficiale gioca un ruolo importante nella fisiologia del nostro organismo e viene anche utilizzata nella tecnologia spaziale.

Perché la superficie del liquido si comporta come una pellicola elastica tesa?

Le molecole situate in uno strato di liquido nei pressi della superficie libera si trovano in condizioni particolari. Esse hanno dei vicini a loro simili soltanto da un lato della superficie, a differenza delle molecole che si trovano all'interno del liquido le quali, invece, sono circondate da tutti i lati da molecole simili. Poiché l'interazione delle molecole a distanze non troppo piccole è attrattiva, l'energia potenziale di ciascuna molecola è negativa. Il suo valore assoluto, in prima approssimazione può essere considerato proporzionale al numero delle molecole più vicine. Pertanto le molecole che si trovano nello strato superficiale (per le quali quindi il numero delle vicine è minore) hanno un'energia potenziale maggiore che le molecole all'interno del liquido. Un altro fattore responsabile dell'aumento dell'energia potenziale delle molecole nello strato superficiale è costituito dalla diminuzione della concentrazione quando ci si avvicina alla superficie libera.

D'altra parte le molecole del liquido si trovano in un moto continuo di agitazione termica: alcune fuoriescono dalla superficie e altre, al contrario, dallo stato gassoso ritornano nel liquido. Si può così parlare di energia potenziale dello strato superficiale di liquido in media in eccesso rispetto a

quella dell'interno. In altre parole, per estrarre una molecola dall'interno del liquido e portarla sulla superficie, forze esterne devono compiere un lavoro positivo. L'eccesso di energia potenziale per unità di superficie, rispetto all'energia potenziale che le molecole avrebbero nell'interno del liquido, si chiama coefficiente di tensione superficiale σ e rappresenta numericamente tale lavoro.

È noto che di tutti i possibili stati del sistema lo stato d'equilibrio è quello in cui l'energia è minima. In particolare, anche la superficie del liquido tende a prendere la forma per la quale la sua energia superficiale, nelle date condizioni, sia minima. Per questo il liquido è caratterizzato dalla tensione superficiale, cioè dalla tendenza a minimizzare la sua superficie libera, condizioni che si ottiene in corrispondenza alla forma sferica.

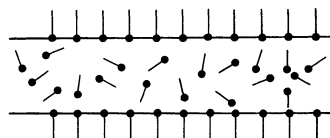
6.1 Bolle di sapone

“Fate una bolla di sapone e guardatela: potreste passare tutta la vita a studiarla senza smettere di ricevere lezioni di fisica”, scriveva il fisico inglese lord **Kelvin**. In particolare, il film di sapone è un bellissimo oggetto per lo studio della tensione superficiale. La forza di gravità gioca in questo sistema un ruolo trascurabile poiché le pellicole di sapone sono molto sottili e la loro massa è quindi assai piccola. Il ruolo principale è in effetti giocato dalla tensione superficiale, che fa sì che la pellicola mantenga una forma tale che la sua area sia la minore possibile, nelle condizioni fissate.

Perché proprio bolle di sapone? Perché non studiare una pellicola di acqua distillata, il cui coefficiente di tensione superficiale è di alcune volte più elevato del coefficiente di tensione superficiale della soluzione di sapone? La ragione di riferirsi al film della soluzione non consiste nel valore del coefficiente di tensione superficiale, ma nella struttura della pellicola. Il sapone è ricco delle cosiddette sostanze *superficiali attive*, molecole lunghe le cui estremità hanno per così dire un rapporto con l'acqua diverso: un'estremità si unisce facilmente con la molecola dell'acqua, l'altra non è in alcun modo influenzata. Per questo la pellicola di sapone ha una struttura complessa: la soluzione che la costituisce è come “armata” di uno steccato di molecole (sistematiche in maniera molto ordinata) della sostanza attiva che compone il sapone (Fig. 6.1).

Torniamo alle bolle di sapone. A tutti è capitato di osservare questi meravigliosi oggetti. Le bolle hanno una forma sferica e possono librarsi

Fig. 6.1: La stabilità di un film di acqua saponata viene garantita dalla presenza di particolari molecole organiche alla superficie.



liberamente in aria, per lungo tempo. La pressione all'interno della bolla è superiore alla pressione atmosferica. Il contributo in eccesso è determinato dal fatto che la pellicola di sapone, cercando di ridurre la propria superficie, comprime l'aria all'interno della bolla. Quanto minore è il suo raggio R tanto maggiore risulta la pressione in eccesso all'interno della bolla. Definiamo questa pressione in eccesso ΔP_{sf} . Supponiamo ora che la tensione superficiale della pellicola sia diminuita di un certo valore e, a seguito di ciò, il suo raggio sia aumentato di una quantità $\delta R \ll R$ (Fig. 6.2).

In questo caso la superficie esterna aumenta di

$$\delta S = 4\pi (R + \delta R)^2 - 4\pi R^2 \approx 8\pi R \delta R$$

($S = 4\pi R^2$ è la superficie della sfera). Di conseguenza aumenta anche l'energia di superficie della bolla:

$$\delta E = \sigma (2 \delta S) = 16\pi \sigma R \delta R, \quad (6.1)$$

(si può trascurare in questa stima la variazione del coefficiente di tensione superficiale) avendo fatto comparire un fattore due per tenere conto che la bolla di sapone ha due superfici, una esterna e una interna. Aumentando il raggio di una quantità δR l'area di entrambe cresce di $8\pi R \delta R$.

L'aumento dell'energia di superficie si è prodotto a spese del lavoro dell'aria in essa compressa. Tenendo conto che la pressione della bolla in presenza di una variazione piccola del volume δV si può assumere costante, possiamo scrivere

$$\delta A_{aria} = \Delta P_{sf} \delta V = \delta E.$$

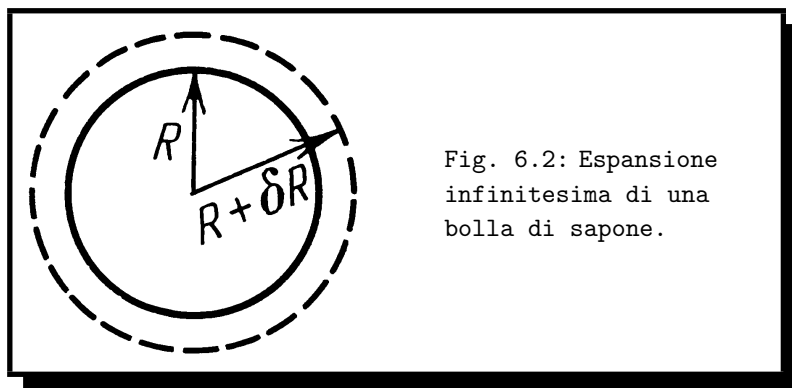


Fig. 6.2: Espansione infinitesima di una bolla di sapone.

La variazione del volume della bolla può essere legata al volume della sfera a parete sottile (Fig. 6.2) :

$$\delta V = \frac{4\pi}{3}(R + \delta R)^3 - \frac{4\pi}{3} R^3 \approx 4\pi R^2 \delta R,$$

dalla quale ricaviamo per δE

$$\delta E = 4\pi R^2 \Delta P_{sf} \delta R.$$

Confrontando questa espressione con la precedente variazione di energia otteniamo che la pressione in eccesso prodotta dalla tensione superficiale all'interno della bolla è

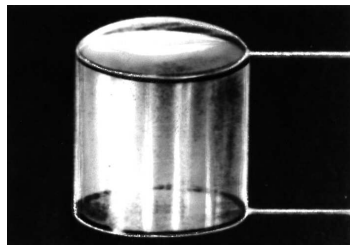
$$\delta P_{sf} = \frac{4\sigma}{R} = \frac{4\sigma'}{R} = 2\sigma' \rho. \quad (6.2)$$

Con $\sigma' = 2\sigma$ abbiamo indicato il doppio del coefficiente di tensione superficiale del liquido. E evidente che se si valutasse la pressione in eccesso legata a un'ordinaria superficie curva (per esempio, all'interno di una goccia di liquido), essa sarebbe definita dall'espressione $\delta P_{sf} = 2\sigma / R$. La grandezza ρ che entra nella (6.2) costituisce la *curvatura* della sfera: $\rho = 1/R$.

Siamo arrivati all'importante conclusione che la pressione in eccesso è proporzionale alla curvatura della sfera. Tuttavia alla bolla di sapone si possono dare anche altre forme, oltre che quella sferica. Se si sistema la sfera tra due anelli, la si può tendere finché non acquista la forma di un cilindro con “cappelli” che sono porzioni di sfera (Fig. 6.3).

Quanto vale la pressione in eccesso all'interno di tale bolla? Sulla superficie cilindrica la curvatura è diversa nelle diverse direzioni. Lungo un

Fig. 6.3: Si riesce a realizzare una bolla di sapone a forma cilindrica con l'aiuto di anelli metallici.



asse parallelo all'asse del cilindro la curvatura è uguale a zero, mentre nella sezione perpendicolare all'asse del cilindro la curvatura è uguale a $1/R$, dove R è il raggio del cilindro. Quale valore dobbiamo dunque dare a ρ ? La differenza di pressione nei diversi lati di una qualsiasi superficie è determinata dalla curvatura media. Tracciamo dei piani attraverso la normale alla superficie, nel punto A (Fig. 6.4). Le sezioni della superficie cilindrica tagliate da questi piani possono essere circonferenze, ellissi o rette. È facile vedere che le curvature di queste sezioni nel punto A sono diverse: la sezione trasversale, la circonferenza, possiede la curvatura massima, mentre la retta (sezione longitudinale) ha la curvatura minima, uguale a zero. La curvatura media $\bar{\rho}$ si definisce come semisomma della curvatura massima e minima delle sezioni normali:

$$\bar{\rho} = \frac{\rho_{\max} + \rho_{\min}}{2}.$$

Questa definizione non è limitata al cilindro: per questa via si può definire anche la curvatura, in un dato punto, di una qualsiasi superficie. Su una superficie cilindrica, in un qualsiasi punto, la curvatura massima è $\rho_{\max} = 1/R$, dove R è il raggio della sezione trasversale del cilindro, mentre $\rho_{\min} = 0$. Per questo la curvatura media del cilindro è $\bar{\rho} = 1/2R$, mentre la pressione in eccesso all'interno della bolla cilindrica vale

$$\Delta P_{\text{cil}} = \frac{\sigma'}{R}.$$

Quindi, nella bolla cilindrica la pressione in eccesso è la stessa che nella sfera di raggio doppio del raggio del cilindro. Per questo il raggio di curvatura dei "cappelli" nella bolla cilindrica sarà due volte maggiore del raggio del cilindro ed essi non sono quindi semisfere.

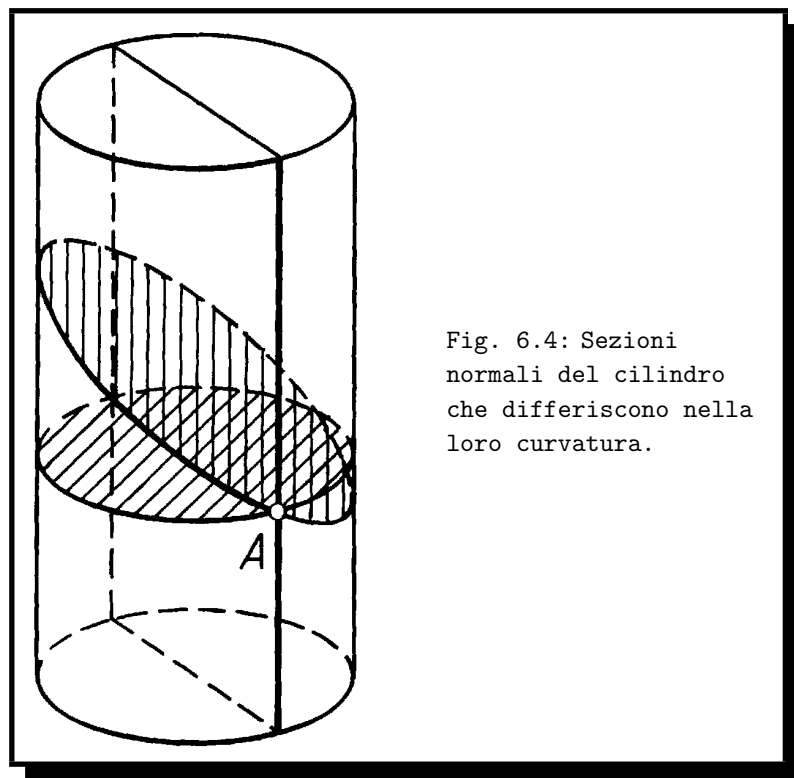


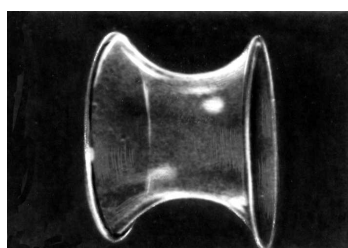
Fig. 6.4: Sezioni normali del cilindro che differiscono nella loro curvatura.

E se volessimo eliminare la pressione in eccesso in questa bolla, costringendola, per esempio, a rompere i “cappelli”? Poiché all'interno della nuova “bolla” non vi è nessuna pressione in eccesso, la sua superficie non dovrebbe avere curvatura. Eppure, le pareti cilindriche della “bolla” si piegano verso l'interno prendendo la forma della *catenoide* (questa superficie si può ottenere ruotando intorno a un asse la curva data da una catena appesa per le estremità: la catenaria). Perché?

Guardate questa superficie (Fig. 6.5). Nella parte stretta vi è un'intercetta. È facile vedere che questa intercetta è contemporaneamente convessa e concava. La sua sezione trasversale è una circonferenza; quella longitudinale una linea. La curvatura verso l'interno deve aumentare la pressione all'interno della bolla, mentre verso l'esterno deve ridurla (la pressione sotto la superficie concava è maggiore che quella al di sopra di essa). Nel caso della catenoide queste curvature sono uguali in modulo e poiché hanno

direzioni opposte, la curvatura media è uguale a zero. Di conseguenza, all'interno di tale bolla non vi è pressione in eccesso. Esiste una grande varietà di superfici che sembrano curve in tutte le direzioni, nonostante la loro curvatura media sia uguale a zero. A queste superfici non è associata alcuna pressione. Per ottenerle è sufficiente prendere una qualsiasi cornice a filo, piegarla e immergerla in acqua saponata. Estraeendo la cornice, si possono vedere superfici diverse con curvatura media uguale a zero, la cui forma dipende dalla forma della cornice. Tuttavia la catenoide è l'unica superficie di rotazione, a parte il piano, con curvatura media nulla.

Fig. 6.5: Lasciato libero il film di acqua saponata assume la forma della catenoide. Tale superficie ha curvatura media nulla.



Uno dei problemi di *geometria differenziale* è costituito dalla ricerca di superfici con curvatura media uguale a zero, su curve spaziali chiuse. Esiste un teorema che afferma che l'area di simili superfici rispetto alle altre possibili, sulla stessa curva, è quella minima. Ora il motivo è manifesto.

Le bolle di sapone possono unirsi le une alle altre, formando una sorta di catena. Nonostante la disposizione dei film appaia caotica, viene rispettata la legge che vuole che i film si intersecano tra di loro soltanto sotto gli stessi angoli (si veda la Fig. 6.6).

Esaminiamo, ad esempio, due bolle che si trovano in contatto tra loro e aventi un diaframma comune (Fig. 6.7). Le pressioni in eccesso (rispetto a quella atmosferica) all'interno delle bolle sono diverse e sono definite dalla formula di **Laplace** (si veda l'Eq. (6.2)):

$$\Delta P_1 = \frac{2\sigma'}{R_1}, \quad \Delta P_2 = \frac{2\sigma'}{R_2}.$$

Pertanto il diaframma deve essere tale da determinare una pressione aggiuntiva all'interno delle bolle. Di conseguenza esso deve avere una certa

curvatura. Il raggio R_3 della curvatura del diaframma è dato dalla relazione

$$\frac{2\sigma'}{R_3} = \frac{2\sigma'}{R_2} - \frac{2\sigma'}{R_1},$$

da cui

$$R_3 = \frac{R_1 R_2}{R_1 - R_2}.$$

La figura 6.7 mostra una sezione delle bolle con il piano che passa attraverso i loro centri. I punti A e B rappresentano i punti di intersezione con il piano del disegno della circonferenza dove si incontrano le due bolle. In un punto qualsiasi di questa circonferenza si incontrano tre film. Poiché la loro tensione superficiale è la stessa, le forze associate possono “equilibrarsi” a vicenda quando gli angoli, sotto i quali i film si intersecano, sono uguali fra loro. Di conseguenza, ogni angolo è uguale a 120° .

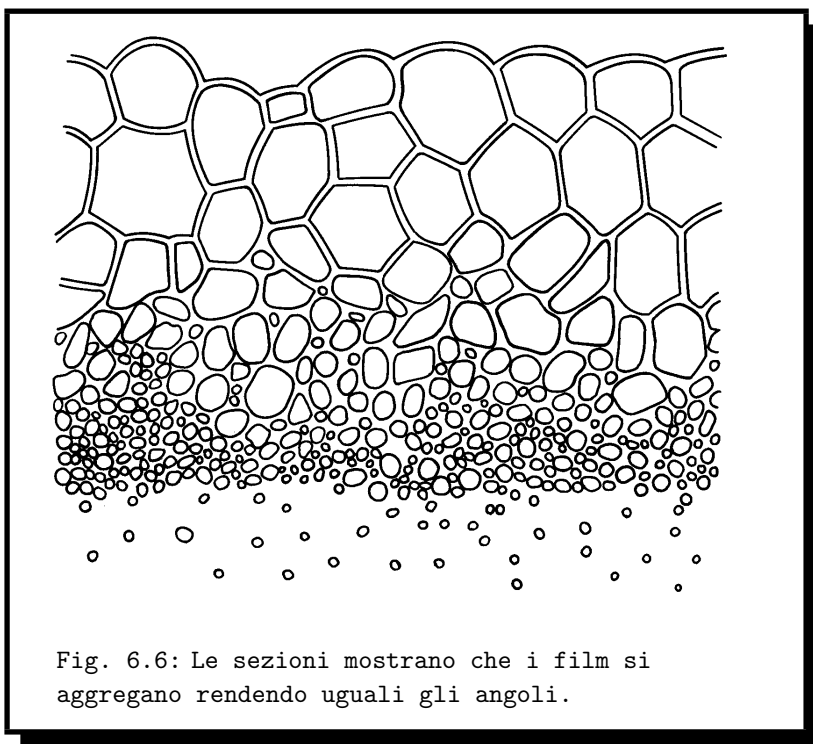
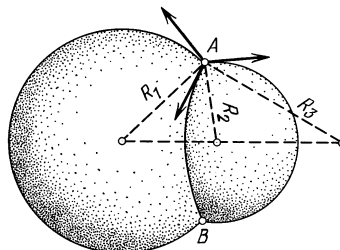


Fig. 6.6: Le sezioni mostrano che i film si aggregano rendendo uguali gli angoli.

Fig. 6.7: L'angolo tra le tangenti al contatto delle bolle è 120° .



6.2 Cosa fanno le gocce

Un problema complesso è costituito dalla forma che assumono gocce di liquidi. Al tentativo della tensione superficiale di ridurre la superficie del liquido, si contrappongono altre forze. Per esempio, le gocce non sono quasi mai sferiche, sebbene la sfera abbia la superficie minore tra tutte le forme geometriche di pari volume. Quando la goccia si stabilizza su una superficie orizzontale immobile essa appare schiacciata. Anche la goccia che cade nell'aria ha una forma complessa. Solo in assenza di gravità la goccia assumerebbe una forma perfettamente sferica.

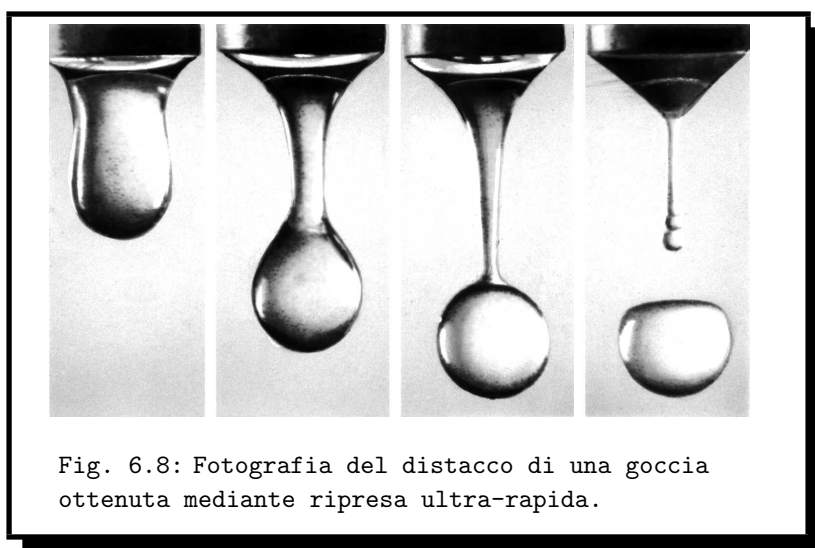
Lo scienziato belga **J. Plato** riuscì per primo ad eliminare l'azione della forza di gravità nello studio della tensione superficiale dei liquidi. Plato propose semplicemente di equilibrare la forza di gravità con la forza di Archimede. Egli sistemò il liquido in esame (olio) in una soluzione avente la stessa densità e, come scrive il suo biografo, "con meraviglia vide che l'olio prendeva una forma sferica; immediatamente adottò la sua norma di "meravigliarsi" quando lo desiderava e questo fenomeno gli servì da oggetto per lunghe considerazioni".

Plato utilizzò il suo metodo per studiare diversi tipi di fenomeni. Per esempio, studiò il processo di formazione e di distacco di una goccia di liquido all'estremità di una cannuccia.

Di solito, per quanto lentamente venga fatta ingrandire la goccia, essa si distacca dalla cannuccia così velocemente che l'occhio non riesce a seguire i dettagli di questo processo. Plato sistemò l'estremità del tubicino in un liquido la cui densità era di poco inferiore alla densità della goccia. L'azione della forza di gravità in questo caso è debole, si può quindi produrre una

goccia molto grande e osservare comodamente come essa si distacca dalla cannuccia.

La figura 6.8 mostra i diversi stadi del processo di formazione e distacco di una goccia (le fotografie sono state ottenute con tecnica moderna, mediante una cine-ripresa accelerata). Tentiamo di spiegare questo fenomeno. Se l'accrescimento è lento, si può ritenere che in ogni istante di tempo la goccia si trovi in equilibrio. A un dato volume della goccia, la forma è determinata dalla condizione che la somma dell'energia di superficie e dell'energia potenziale gravitazionale della goccia sia minima.



La tensione superficiale provoca la riduzione della superficie della goccia, e tende a farle assumere una forma sferica. La forza di gravità, al contrario, tende a sistemare il centro delle masse della goccia quanto più in basso possibile. A seguito di ciò la goccia risulta “stirata”.

Quanto più grande è la goccia, tanto maggiore è il ruolo che gioca l'energia potenziale dovuta alla forza di gravità. Al crescere delle dimensioni della goccia la massa principale si raccoglie in basso e si forma una strozzatura (seconda fotografia della figura 6.8). La forza associata alla tensione superficiale ha una direzione verticale rispetto alla tangente ed equilibra la forza di gravità agente sulla goccia. Ora per la bolla è sufficiente ingrandirsi anche di poco perché le forze di tensione superficiale non

riescano più ad equilibrare la forza di gravità. La strozzatura della goccia si restringe rapidamente (terza fotografia della figura 6.8) e di conseguenza la goccia si distacca (quarta fotografia). In questo caso si separa una piccola gocciolina, che segue nella caduta quella più grande. La seconda gocciolina si forma in ogni caso (essa viene chiamata sfera di Plato). Solitamente non notiamo questa gocciolina secondaria poiché il processo di distacco della goccia è molto rapido.

Non entreremo nei dettagli della formazione della sfera di Plato: si tratta di un problema particolarmente complesso. Cerchiamo, però, di spiegare la forma delle gocce che cadono. Le istantanee delle gocce mostrano che quelle piccole sono quasi sferiche, mentre quelle più grosse sono simili a un panino. Valutiamo dapprima il raggio della goccia, all'istante in cui sta per perdere la sua sfericità.

In caso di moto uniforme della goccia la forza di gravità che agisce, per esempio, sulla colonnina centrale AB (Fig. 6.9), deve essere equilibrata dalle forze di tensione superficiale. Perciò è necessario che i raggi della curvatura della goccia nei punti A e B siano diversi. Effettivamente la tensione superficiale crea una pressione in eccesso determinata dalla formula di Laplace $\Delta P_L = \sigma' / R$.

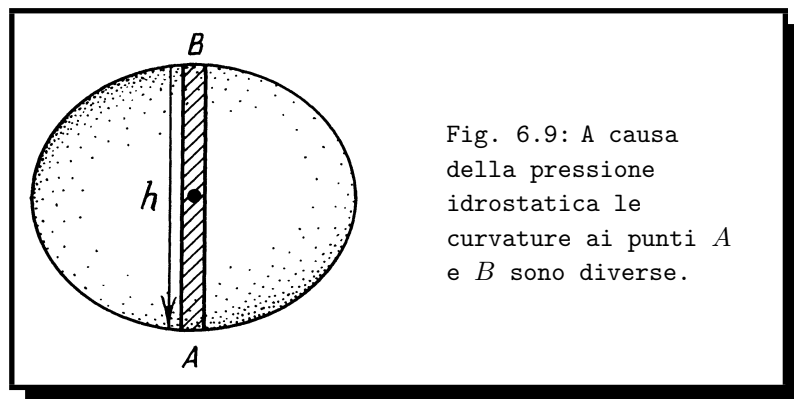
Se la curvatura della superficie della goccia nel punto A è maggiore che nel punto B , la differenza delle pressioni può equilibrare la pressione idrostatica del liquido:

$$\rho g h = \frac{\sigma'}{R_A} - \frac{\sigma'}{R_B}$$

(ρ è qui la massa specifica del liquido).

Quanto vale la differenza tra R_A e R_B ? Per gocce piccole, con raggio dell'ordine di 1μ ($10^{-6} m$), si ha $\rho g h \approx 2 \cdot 10^{-2} Pa$, mentre $\Delta P_L = \sigma' / R \approx 1.6 \cdot 10^5 Pa$! In questo caso la pressione idrostatica è talmente piccola, rispetto a quella di Laplace, che si la può tranquillamente trascurare. Questa goccia può essere considerata sferica.

Diversa è la situazione con una goccia di diametro, ad esempio, di $4 mm$. Per essa, la pressione idrostatica diviene all'incirca pari a $\rho g h \approx 40 Pa$, e quella laplaciana $\Delta P_L = 78 Pa$. Queste grandezze sono dello stesso ordine e le perturbazioni alla sfericità per una simile goccia sono sostanziali.



Supponendo che $R_B = R_A + \delta R$ e $R_A + R_B = h = 4 \text{ mm}$, troveremo

$$\delta R \sim h \left[\sqrt{\left(\frac{\Delta P_L}{\rho g h} \right)^2 + 1} - \frac{\Delta P_L}{\rho g h} \right] \approx 1 \text{ mm}.$$

La differenza dei raggi di curvatura nei punti A e B in questo caso risulta già dell'ordine della dimensione della goccia.

La nostra valutazione indica per quale dimensione delle gocce ci si può attendere una perturbazione della sfericità. Tuttavia la forma della goccia che cade risulta in realtà diversa da quella che appare nell'esperimento (nella fotografia le gocce sono schiacciate verso il basso).

Perché ? Perché noi abbiamo implicitamente assunto uguali la pressione dell'aria sulla goccia e sotto di essa. In caso di movimento lento della goccia è effettivamente così. Ma se la goccia si muove nell'aria con velocità sufficientemente alta, l'aria non fa in tempo ad avvolgere uniformemente la goccia: davanti ad essa si forma una regione di pressione maggiore, mentre dietro di essa si ha una regione di bassa pressione (lì si formano i vortici). La differenza delle due pressioni può essere maggiore della pressione idrostatica e la pressione laplaciana ora deve compensare questa differenza. In questo caso, il termine

$$\frac{\sigma'}{R_A} - \frac{\sigma'}{R_B}$$

diventa negativo e, di conseguenza, il raggio R_A sarà maggiore di R_B .

Avete talvolta osservato gocce molto grandi? In condizioni normali non

ne esistono. E ciò non è casuale: le gocce di grandi dimensioni sono instabili e si disintegrano in piccole. La conservazione della forma della goccia su una superficie è determinata dalla tensione superficiale. Quando la pressione idrostatica supera quella laplaciana la goccia si disperde e si frantuma in gocce più piccole. Si può valutare la dimensione massima possibile di una goccia con la diseuguaglianza $\rho g h \gg \frac{\sigma'}{R_A}$, dove $h \sim R$. Si ha

$$R_{\max} \sim \sqrt{\frac{\sigma'}{\rho g}}.$$

Per l'acqua, ad esempio, come ordine di grandezza della massima dimensione possibile si ha $R_{\max} \approx 0.3 \text{ cm}$. Ecco perché sulle foglie degli alberi e su altre superfici scarsamente ricopribili da acqua, non troveremo gocce di dimensioni elevate.

Capitolo 7

Il microfono ad acqua.

Tutti conoscono il microfono. Tuttavia pochi probabilmente conoscono l'esistenza di un microfono ad acqua. Mediante un getto d'acqua si possono amplificare i suoni! Questo insolito strumento è stato costruito da **Alexander Bell**, più noto come l'inventore del telefono (anche se pare ormai assodato che il primo ad idearlo sia stato l'italiano **Antonio Meucci**).

Descriviamo, innanzi tutto, le proprietà del getto d'acqua. Se sul fondo di un recipiente riempito di acqua si pratica un piccolo foro, si può notare che la vena che ne fuoriesce è composta di due parti, diverse per le loro proprietà. La parte superiore è trasparente e tanto immobile da sembrare di vetro. Quanto più si allontana dalla sorgente tanto più la vena diventa sottile e nel punto in cui si riduce maggiormente, ha inizio una parte inferiore opaca e dalla forma variabile. Di primo acchito questa parte del getto, come anche quella superiore, sembra ininterrotta. Tuttavia qualche volta si riesce a passare attraverso di essa velocemente un dito, senza bagnarlo. Il fisico francese **Felix Savar** ha studiato attentamente le proprietà di getti di liquido ed è arrivato alla conclusione che nella parte più stretta essi non sono continui, ma si frantumano in realtà in singole goccioline. Ora, dopo più di cento anni, è facile rendersene conto fotografando un getto con un flash o osservandolo con luce stroboscopica (Fig. 7.1). Ai tempi di Savar, tuttavia, le proprietà del getto potevano essere studiate solo mediante un lampo di luce.

Osserviamo l'immagine della parte inferiore di un getto. Essa è composta da singole goccioline, piccole e grandi alternate. Come si vede, nel processo di caduta le gocce grandi "pulsano", cambiando in sequenza la loro forma da ellissoide allungato in direzione orizzontale (bolle 1 e 2), ad

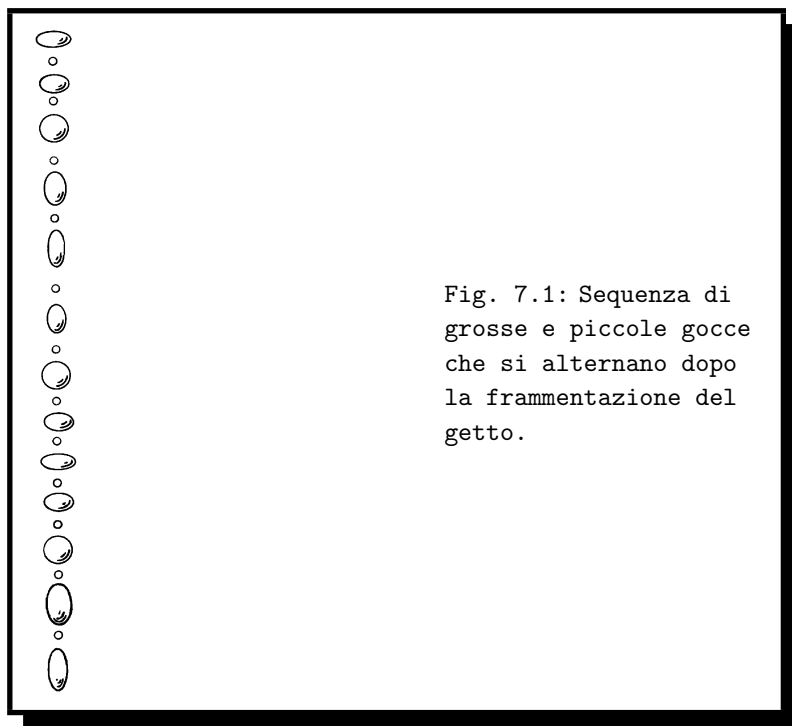


Fig. 7.1: Sequenza di grosse e piccole gocce che si alternano dopo la frammentazione del getto.

una sfera (3), quindi trasformandosi in ellissoidi allungati in direzione verticale (4, 5, 6). Poi la goccia ritorna di nuovo alla forma sferica (7) e così via. Ogni goccia, pulsando rapidamente nel corso della propria caduta^a, produce effetti diversi all'occhio umano. Per questo si crea l'immagine di un getto torbido, con ispessimenti nelle regioni in cui le gocce ellissoidali sono allungate orizzontalmente e restringimenti laddove esse sono allungate verticalmente.

Presentandosi dapprima come protuberanze, man mano che si allon-

^aNon è difficile valutare la frequenza della variazione di forma della goccia. Come vedremo nel capitolo dedicato a calici sonori o silenti, per le bolle d'aria in un liquido si ha per la frequenza

$$\nu \sim \sigma^{1/2} \rho^{-1/2} r^{-3/2}.$$

Usando i valori $\sigma = 0.07 \text{ N/m}$, $\rho = 10^3 \text{ kg/m}^3$, $r = 3 \cdot 10^{-3} \text{ m}$, ricaviamo per la frequenza $\nu \approx 50 \text{ Hz}$. Si osservi che il susseguirsi delle immagini sullo schermo in cui è proiettato un film è di 24 riquadri al secondo: per il nostro occhio ciò è sufficiente per creare l'illusione di un movimento senza discontinuità.

tanano dal foro le gocce appaiono sempre più distinte finché non si distaccano definitivamente. Le protuberanze si susseguono l'un l'altra in modo che producono un debole suono. La frequenza caratteristica della vibrazione sonora favorisce una più rapida disintegrazione del getto in singole gocce. Pertanto il suono trasforma la parte del getto prima trasparente in torbida. Un'altra importante scoperta di Savar fu che i suoni producono un effetto più marcato sulla zona superiore del getto, quella trasparente. Se vicino ad esso si producono delle vibrazioni sonore di determinata frequenza, la parte trasparente di questo, appunto, immediatamente diviene torbida. Savar lo spiegò così: le goccioline nelle quali si divide il getto d'acqua nella sua parte inferiore, vengono in tal caso prodotte anche più in alto, vicino al foro di uscita dello zampillo.

Il fisico inglese **John Tindal** ripeté l'esperienza di Savar. Riuscì a creare un getto (vena, secondo la sua espressione), la cui parte trasparente raggiungeva una lunghezza di circa 90 piedi (27.4 m). Sotto l'azione del suono di una canna d'organo questa vena trasparente diveniva torbida, disintegrandosi in una elevatissima quantità di gocce d'acqua. In uno dei suoi articoli Tindal descrive così lo studio della caduta del getto in una piscina : "Quando un getto d'acqua cade su una superficie liquida posta sopra il punto critico (il punto del passaggio dalla parte trasparente del getto a quella torbida) e la pressione non è molto forte, essa entra nel liquido in silenzio; ma, quando questa superficie incontra il getto al di sotto del punto critico, in quel momento stesso si sente un mormorio e compaiono una moltitudine di bollicine. Nel primo caso non soltanto non vi è una marcata polverizzazione del liquido, ma esso si raccoglie a mucchio intorno alla base della vena, nella sporgenza dove il moto è opposto al getto".

Le proprietà del getto sono state utilizzate da Bell per la costruzione del microfono ad acqua, che è riportato nella figura 7.2. L'apparecchio è costituito da un tubicino metallico dove è saldato un bocchettone, su cui è montato un imbuto. L'estremità inferiore di questo tubicino è sistemata su un supporto, mentre quella superiore è chiusa da una membrana di gomma elastica, fissata al tubicino con un filo. Come sappiamo dall'esperimento di Tindal, la parte inferiore del getto, disintegrandosi in tante goccioline, produce un suono. Se invece è la parte superiore, intera, che penetra nell'acqua, essa vi "sfocia" con continuità. Un esperimento simile si può fare anche con un pezzo di cartone. Se si avvicina un foglio di cartone su cui cade un getto d'acqua verso la sorgente, i colpi delle gocce si sentiranno sempre più debolmente e quando sarà raggiunto il punto di intervallo non si sentiranno

affatto. La membrana nel microfono di Bell gioca lo stesso ruolo di un foglio di cartone. Tuttavia, grazie al risuonatore, ossia al tubicino e al megafono, ogni piccolo colpo delle gocce viene più facilmente avvertito. Quindi, le gocce che cadono sulla membrana di gomma producono l'impressione di deboli colpi di martello su un'incudine.

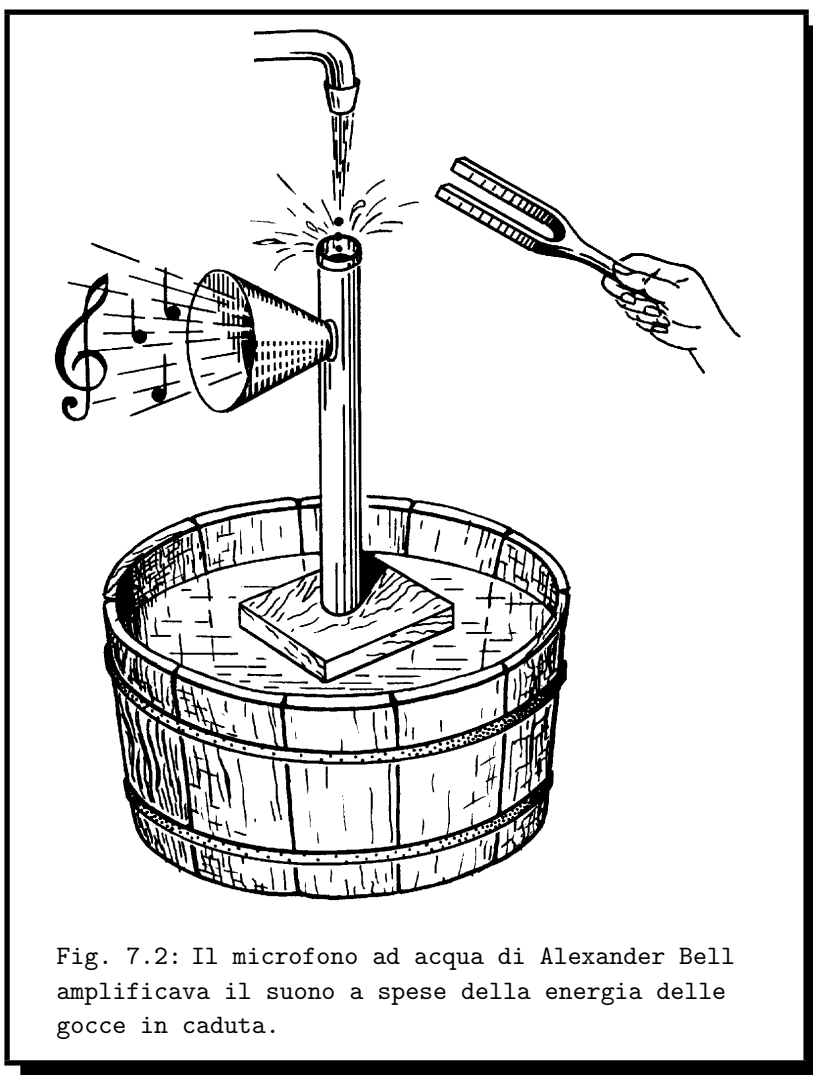


Fig. 7.2: Il microfono ad acqua di Alexander Bell amplificava il suono a spese della energia delle gocce in caduta.

Mediante il microfono ad acqua è possibile, per esempio, convincersi della sensibilità di un getto ai suoni, come descritto da Savar e da Tindal. Così, se si avvicina un diapason, il getto si disintegra immediatamente in gocce che, percuotendo la membrana, iniziano a “cantare” fortemente. Nell’amplificazione di questo suono inizialmente debole, a spese dell’energia del getto che cade, si manifesta l’azione del microfono ad acqua. Se al posto del diapason si accosta al getto un orologio, tutti i presenti nella stanza ne sentiranno il ticchettio. Il noto divulgatore scientifico **Donat**, alla fine del XIX secolo, affermava che adattando ad un tubo di vetro in cui scorre acqua un imbuto, egli aveva tentato di trasmettere i suoni della sua voce attraverso il getto d’acqua. Nell’esperimento il getto si mise realmente a “parlare”, ma in modo così confuso e rozzo e con voce così sgradevole, che i presenti fuggirono.

Leggendo queste righe possiamo rallegrarci del fatto che la principale invenzione di Bell (o di Meucci) - il telefono con il microfono elettrico - non abbia questo inconveniente.

Capitolo 8

Tracce sulla sabbia.

Ecco quanto disse su questo argomento **Reynolds**, noto per i suoi studi di idrodinamica, nella relazione alla riunione dell'Associazione Britannica nel 1885: "Quando il piede preme sulla sabbia, sul tratto compatto dopo il ritiro della marea, la regione che si trova attorno al piede diventa immediatamente asciutta. La pressione del piede rende la sabbia meno compatta e quanto più tale pressione è forte, tanta più acqua fuoriesce... Essa rende la sabbia asciutta finché dal basso non arriva altra acqua".

Perché a seguito della pressione aumenta lo spazio tra i granelli e l'acqua non è più sufficiente per riempirlo? Non a caso questo interrogativo interessava gli scienziati del XIX secolo. La risposta a questa domanda è legata alla "struttura" della sostanza. Cerchiamo di capire questo processo.

8.1 Compattamento di sfere.

Si può riempire con sfere solide tutto lo spazio? Naturalmente no, tra di esse rimangono comunque spazi liberi. La frazione di spazio occupata dalle sfere definisce la *densità del loro compattamento*. Quanto più compattamente sono disposte le sfere, tanto meno spazio libero vi sarà tra di esse e tanto maggiore sarà la densità. Quando si raggiunge la densità massima di sfere solide simili? La risposta a questa domanda fornirà la chiave alla soluzione del "mistero" delle tracce sulla sabbia.

Studiamo inizialmente il caso più semplice: la disposizione di cerchi uguali su di un piano. Si può raggiungere una disposizione compatta inscrivendo i cerchi in un mosaico di poligoni regolari che copre tutto il piano. Esistono soltanto tre modi per costruire questi mosaici: usare triangoli re-

golari, quadrati o esagoni regolari. La figura 8.1 mostra disposizioni di cerchi in un mosaico quadrato e in uno esagonale. Anche “ad occhio” si vede che il secondo modo (*b*) è più “economico”. Una semplice valutazione numerica mostra che in questo caso il piano è riempito da cerchi al 90,7 %, mentre nel primo caso (*a*) soltanto al 78 % circa. La disposizione esagonale sul piano è quindi più compatta. Evidentemente, proprio per questo la utilizzano le api nel costruire i favi.

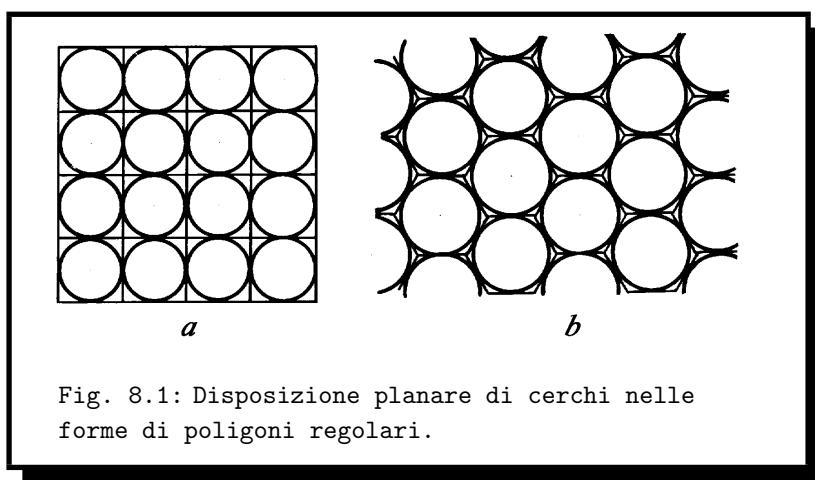
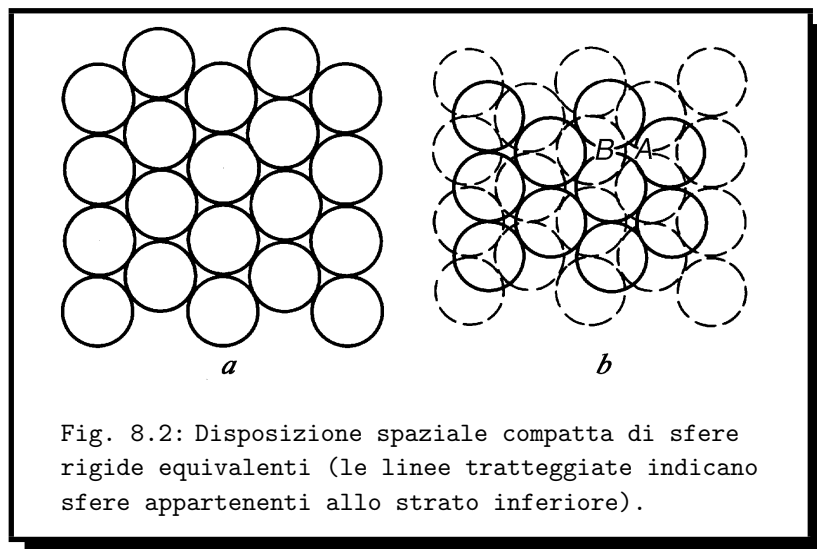


Fig. 8.1: Disposizione planare di cerchi nelle forme di poligoni regolari.

Un compattamento di sfere nello spazio si può ottenere nel seguente modo. Sistemiamo su un piano il primo strato di sfere, secondo la configurazione esagonale che abbiamo compreso essere la più compatta. Quindi su di esse mettiamo un secondo strato ordinato allo stesso modo. Se ogni sfera dello strato superiore si sistema esattamente sopra la corrispondente dello strato inferiore, così che tutte le sfere risultino verticalmente inscritte in poligoni regolari, rimarrà inutilizzato troppo spazio. Con un simile metodo di impacchettamento delle sfere, si riempirà infatti soltanto il 52% dello spazio a disposizione.

È intuitivo che si possono disporre le sfere in modo più compatto. Per far ciò è sufficiente sistemare una sfera superiore nella cavità formata da tre sfere vicine inferiori. In questo caso, però, le sfere superiori non possono riempire tutte le cavità, uno dei due incavi vicini rimane sempre libero (Fig. 8.2).

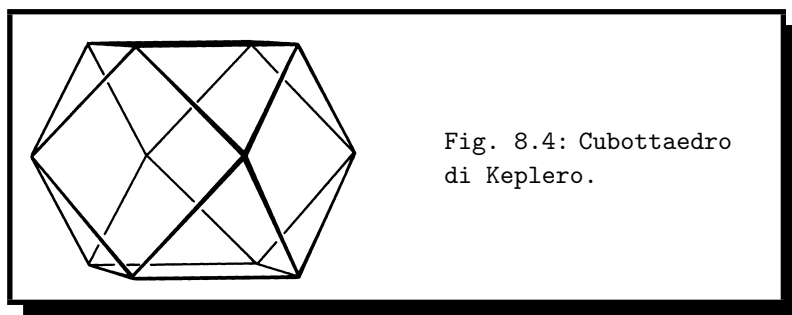
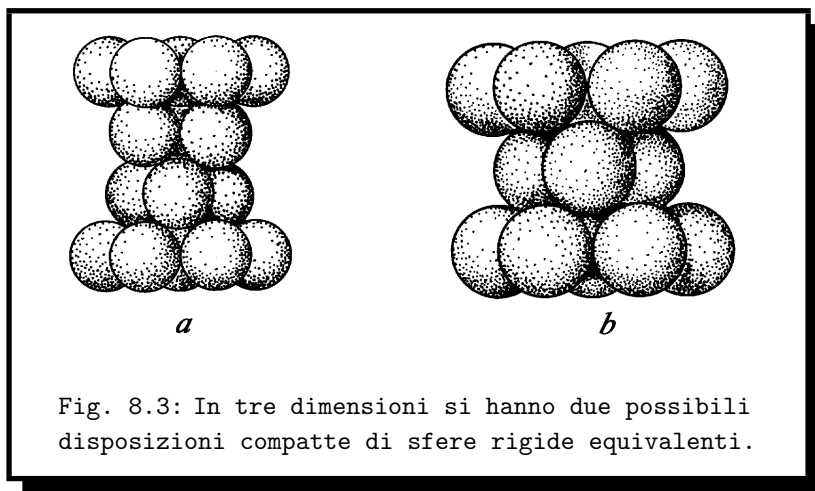
Per questo quando sistemeremo il terzo strato di sfere lo potremo fare



in due modi: si possono mettere le sfere del terzo strato sugli incavi del primo strato che le sfere del secondo hanno lasciato liberi (il centro di una di queste sfere nella figura 8.2, *b* si trova nel punto *A*), oppure in corrispondenza delle sfere del primo strato (in questo caso il centro di una delle sfere del terzo strato si troverà nel punto *B*). Per gli strati successivi l'ordine di sistemazione delle sfere si ripete allo stesso modo. Otteniamo così i due modi di compattamento mostrati nella figura 8.3. In ambedue i casi le sfere riempiono circa il 74% dello spazio.

E' facile dimostrare che in una disposizione di questo genere ogni sfera è in contatto con 12 sfere contigue. I punti di contatto definiscono un poliedro a 14 spigoli. I suoi spigoli sono quadrati alternati a triangoli equilateri. Con il secondo tipo di disposizione (Fig. 8.3, *b*) si ottiene il poliedro mostrato nella figura 8.4.

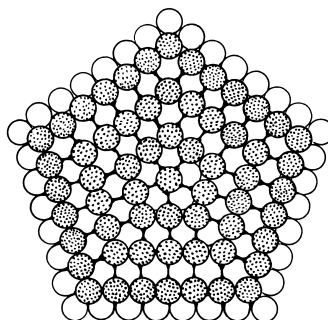
Fino ad ora abbiamo considerato soltanto quei modi di compattamento delle sfere secondo i quali esse vengono inscritte in "celle" periodiche. Si può ottenere compattamento ignorando questa condizione? Uno dei metodi è mostrato nella figura 8.5. Le sfere di ogni piano sono situate sui lati di pentagoni regolari. In ogni pentagono le sfere contigue si toccano l'un l'altra ma le sfere dei diversi pentagoni di un piano sono separate. I lati dei pentagoni degli strati alternantisi contengono alternativamente un



numero pari e dispari di sfere. Il coefficiente di riempimento dello spazio per una tale struttura è pari al 72%, solo di poco inferiore al caso del compattamento mostrato in figura. Si possono compattare le sfere senza seguire la struttura reticolare ottenendo il coefficiente di riempimento del 74%? Oppure esistono metodi di impacchettamento ancora più compatti? La risposta a questo quesito non è definitiva, neppure oggi.

Torniamo alle orme sulla sabbia. Ora sappiamo che esistono particolari disposizioni delle sfere che comportano poco spazio libero tra di esse. Se si perturba questa disposizione, per esempio portando fuori le sfere di uno strato dalle cavità dell'altro strato, gli spazi liberi tra le sfere aumenteranno. Tuttavia nessuno impacca i granelli di sabbia in modo speciale. Come può,

Fig. 8.5:
Compattamento
pentagonale di sfere.



allora, realizzarsi un compattamento dei granelli di sabbia?

Eseguiamo un semplice esperimento. Avete bisogno di riempire un recipiente di semola, in modo tale che ne contenga la massima quantità. Che cosa fate? Scuotete leggermente il recipiente oppure lo picchiate ottenendo l'effetto desiderato. Anche dopo una forte compressione della semola nel recipiente, si può, usando nuovamente tale metodo, trovare posto per un'altra quantità di semola.

Negli anni 1950, lo studioso inglese **Scott** intraprese lo studio scientifico di questo problema. Riempì di sferette recipienti di diverse dimensioni. Quando i recipienti vengono riempiti senza scuotimenti, così che le sfere si dispongano a caso, sperimentalmente si nota la seguente dipendenza della densità dal numero N di sfere

$$\rho_1 = 0.6 - \frac{0.37}{\sqrt[3]{N}}.$$

Se il numero di sfere è molto alto (negli esperimenti N raggiungeva alcune migliaia) si vede che la densità di compattamento tende a diventare costante, pari al 60% dello spazio. Se il recipiente viene scosso mentre lo si riempie il compattamento aumenta:

$$\rho_2 = 0.64 - \frac{0.33}{\sqrt[3]{N}}.$$

Per la verità, anche in questo caso la densità è notevolmente inferiore al 74% che corrisponde alla disposizione regolare delle sfere.

Vale la pena di ripensare agli esperimenti idealmente condotti. Perché la correzione è inversamente proporzionale a $\sqrt[3]{N}$? Le sfere situate presso le pareti del recipiente si trovano in una particolare posizione in confronto alle sfere all'interno e influiscono sul compattamento. Il loro contributo è proporzionale al rapporto dell'area della superficie del recipiente ($\sim R^2$) diviso il volume del recipiente ($\sim R^3$) e decresce quindi in modo inversamente proporzionale alla dimensione del sistema (R). Come volume del sistema intendiamo il volume totale dello spazio, occupato dalle sfere e dalle zone vuote. La dimensione R è $\sqrt[3]{N}$, poiché il volume del sistema è proporzionale al numero totale delle sfere. Questo tipo di calcoli si usa spesso in fisica, quando è necessario tener conto degli effetti di superficie.

Esperimenti più accurati confermano l'osservazione empirica e mostrano che scuotendo un mezzo granulare si può ottenere un maggiore compattamento. In ogni caso, perché ciò si verifica? Il fatto è che ad una posizione di equilibrio corrisponde il minimo dell'energia potenziale. Una sfera può giacere stabilmente in una cavità, ma scivolerà se si trova sul vertice di una gibbosità. Una cosa più o meno simile succede anche nel caso discusso. Le sfere a seguito dello scuotimento scivolano negli interstizi liberi, la densità aumenta, mentre il volume totale del sistema diminuisce. Di conseguenza il livello di riempimento del recipiente diminuisce e si abbassa il centro di massa, riducendo l'energia potenziale.

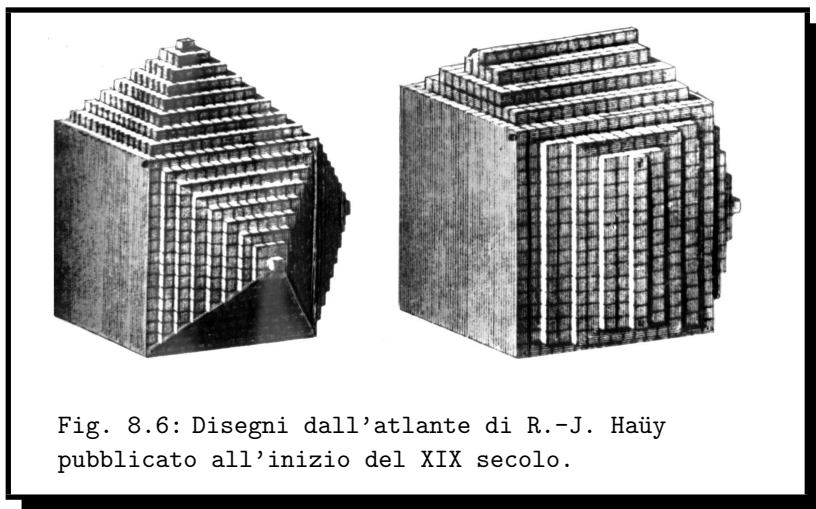
Ora, infine, possiamo comprendere che cosa si verifica con la sabbia. Il movimento dell'acqua scuote la sabbia così che si raggiunge un denso compattamento dei granelli. Premendo la sabbia con il piede, alteriamo questo compattamento e aumentiamo gli spazi vuoti. L'acqua dagli strati superiori della sabbia scenderà, allora, in profondità, riempiendo questi nuovi interstizi. Per questo la sabbia si "asciuga". Allorché si toglie il piede la deformazione scompare, il compattamento si ripristina e l'acqua eliminata dagli interstizi nuovamente ridottisi riempie l'orma lasciata dal piede. Può succedere che dopo una forte pressione il compattamento più elevato non si ripristini. L'orma diventerà in tal caso di nuovo "bagnata" soltanto quando l'acqua salirà dagli strati inferiori e riempirà gli interstizi.

Si può osservare che le proprietà dei mezzi costituiti da molti grani erano empiricamente note ai fachiri indiani. Uno dei loro trucchi consisteva nell'introdurre ripetutamente un coltello sottile in un recipiente con l'apertura stretta e riempito fino all'orlo di riso. Ad un certo momento il coltello si incagliava nel riso e tirando verso l'alto il coltello si poteva alzare il recipiente. È chiaro che il trucco consisteva nello scuotere il re-

ciante mentre lo si riempiva di riso, ottenendo così un marcato compattamento. L'inserimento del coltello turbava tale compattamento e aumentava il volume dello spazio tra i granelli di riso. Poiché il volume del recipiente non cambia, aumentano necessariamente le forze di pressione sui grani e di conseguenza anche l'attrito tra loro e il coltello. Ad un certo punto, nell'esibizione del fachimiro, tale attrito risultava sufficiente ad impedire l'estrazione del coltello dal riso.

8.2 Ordine a corto e a lungo raggio.

Gli atomi di cui sono composti i corpi non possono ovviamente essere considerati come sfere solide. Nonostante ciò, una tale semplificazione può essere di grande aiuto per capire la struttura dei corpi. L'approccio geometrico ai solidi è stato usato per la prima volta già nel 1611 dallo scienziato tedesco **Keplero**, che avanzò l'ipotesi che la forma esagonale dei fiocchi di neve fosse legata al compattamento delle sfere. Lo scienziato russo **M.V. Lomonosov** nel 1760 ha dato la prima rappresentazione del compattamento cubico e ha spiegato in questo modo la forma dei poligoni cristallini. L'abate francese **R.-J. Haüy** nel 1783 ha notato che qualsiasi cristallo è composto di tante parti ripetute (figura 8.6). Egli giustificò la forma geo-



metrica regolare dei cristalli con il fatto che essi sono composti di piccoli “mattoncini” simili. Infine, nel 1824 lo studioso tedesco **A. I. Seeber** ipotizzò un modello di cristallo costituito da piccole sfere posizionate con regolarità, interagenti come gli atomi. L’ammassamento denso di queste sfere corrisponde al minimo dell’energia potenziale associata alla loro interazione.

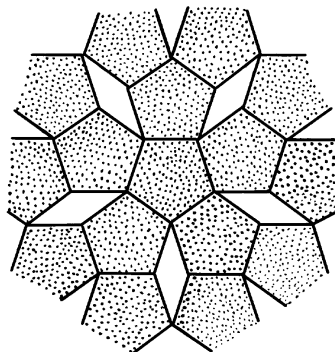
Della descrizione della struttura dei cristalli si interessa una scienza particolare: la *cristallografia*. Attualmente la disposizione periodica degli atomi nei cristalli è un fatto dimostrato in maniera conclusiva. I microscopi elettronici ci permettono di osservarlo, in pratica con i nostri occhi. Nel mondo atomico sussiste una tendenza all’ammassamento compatto. Circa 35 elementi cristallizzano in modo tale che i loro atomi nello spazio si sistemano come le sfere mostrate nella figura 8.3. I centri degli atomi (più precisamente i nuclei) formano il cosiddetto *reticolo cristallino*, che è composto di unità che periodicamente si ripetono nello spazio. I reticoli più semplici, che possono essere ottenuti con uno spostamento periodico nello spazio di un solo atomo, si chiamano *reticoli di Bravais* (dal nome dell’ufficiale della marina francese **Auguste Bravais** che ha esposto per la prima volta, nel XIX secolo, la teoria dei reticoli).

Di reticoli di Bravais non ne esistono poi molti, soltanto 14 tipi diversi. La limitazione nel numero dei reticoli bravaisiani è legata al fatto che non tutti gli elementi di simmetria possono realizzarsi in reticoli periodici. Un pentagono regolare può, per esempio, essere ruotato attorno all’asse che passa per il centro, così che per 5 volte coinciderà con se stesso. In questo caso si dice che si ha un’asse di simmetria del 5 ordine. Tuttavia in un reticolo di Bravais, che deve ripetersi con periodicità spaziale, non può sussistere quest’asse di simmetria. Ciò comporterebbe infatti la possibilità di realizzare con continuità un piano mediante pentagoni regolari, e ciò non è possibile (si veda la figura 8.7)!

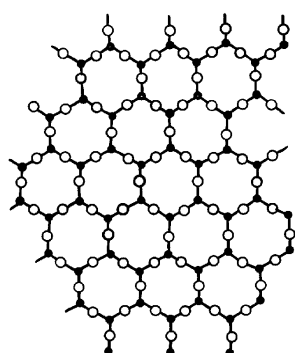
In altre parole, a meno di difetti strutturali o di impurezze, un cristallo deve essere costruito mediante ripetizione regolare, nello spazio, di unità strutturali identiche. Questa proprietà si chiama *simmetria di traslazione*. Si può anche affermare che nei cristalli sussiste un *ordine a lungo raggio*. Questa è la peculiarità più caratteristica dei cristalli, che li distingue da tutti gli altri corpi.

Vi è un’altra classe di materiali solidi, non meno importanti: gli *amorfi*. In tali materiali non si ha ordine a lungo raggio. In uno stato analogo agli amorfi, dal punto di vista strutturale, si trovano i liquidi. Tuttavia anche

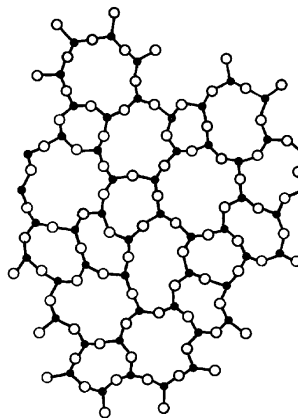
Fig. 8.7: Non è possibile ricoprire con continuità un piano con pentagoni regolari.



un solido può essere di tipo amorfo. L'esempio più semplice è il normale vetro. Nella figura 8.8 è mostrata la struttura di un vetro e di un cristallo di quarzo, che ha la stessa composizione chimica del vetro.



a



b

Fig. 8.8: Struttura del quarzo (*a*) e del vetro (*b*).

Il quarzo è un solido cristallino, mentre il vetro è un solido amorfo. Benché, manifestamente, manchi nel vetro un ordine a lungo raggio, ciò

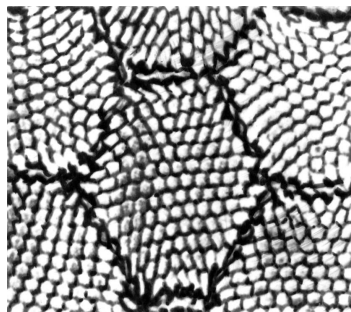
non significa che nella disposizione degli atomi sussista un disordine totale. Una certa struttura nella disposizione degli atomi contigui più vicini, come si vede dalla figura, si conserva anche nel vetro. Si dice che nei corpi amorfi si realizza *ordine a corto raggio*. I materiali amorfi recentemente hanno avuto largo impiego nella tecnica. Le leghe metalliche amorfe (*vetri metallici*) hanno proprietà uniche. Esse si ottengono raffreddando rapidamente il metallo liquido, con una velocità di raffreddamento dell'ordine di alcune migliaia di gradi al secondo. Ciò si può ottenere, per esempio, spruzzando piccole gocce di metallo sulla superficie di un disco freddo che ruota velocemente. La goccia si "spande" su uno strato molto sottile (dello spessore di alcuni micrometri) e una buona dissipazione di calore permette al metallo di raffreddarsi così rapidamente che i suoi atomi non riescono ad accomodarsi nel modo corretto. Ne consegue che le leghe amorfe presentano un'alta durezza, un'alta resistenza alla corrosione e una combinazione ottimale delle proprietà elettriche e magnetiche. Il campo d'impiego di questi materiali va rapidamente espandendosi.

Esperimenti interessanti per spiegare la struttura dei materiali amorfi sono stati condotti nel 1959 dallo scienziato inglese **J. Bernal**. Sfere di plastilina sono state sistemate e pressate disordinatamente in una zolla. Quando successivamente sono state studiate si è osservato che i poliedri che ne descrivevano la forma geometrica erano soprattutto pentaedri. Questi esperimenti sono stati condotti anche con palline di piombo sferiche. Se prima di essere pressate si sistemano le palline in modo molto ordinato e compatto, a seguito della pressione si ottengono dei rombododecaedri quasi perfetti. Se invece vengono ammassate casualmente si ottengono dei corpi irregolari a quattordici facce. In questo caso si trovano delle facce quadrangolari, pentagonali ed esagonali, ma prevalgono i pentagoni.

Nella moderna tecnologia spesso è necessario raggruppare compattamente i componenti di un manufatto. La figura 8.9 è una fotografia della sezione di un cavo superconduttore, costituito da molte anime superconduttrici, racchiuse in una guaina di rame. Inizialmente le anime hanno una forma cilindrica, ma dopo la pressione si trasformano in prismi esagonali. Quanto più compattamente e regolarmente si impacchettano le anime, tanto più regolari sono gli esagoni visibili nella sezione. Ciò indica la qualità di preparazione del cavo. Quando la densità dell'ammassamento viene alterata, nella sezione compaiono pentagoni.

La simmetria del 5° ordine è molto diffusa nella materia biologica. La figura 8.10 riproduce una fotografia al microscopio elettronico di colonne di

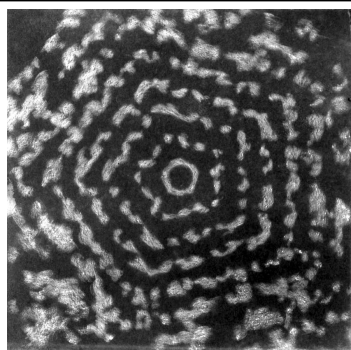
Fig. 8.9: Sezione di un cavo superconduttore di alta qualità. A seguito della compressione i cavi cilindrici divengono esagonali.



particelle virali. Come si vede, si ha una completa analogia con l'ammassamento pentagonale delle sfere, mostrato nella Fig. 8.5.

I paleontologi usano la presenza degli assi di ordine 5 negli oggetti fossili per provare la loro origine biologica (e non geologica).

Fig. 8.10: Fotografia al microscopio elettronico che evidenzia la simmetria di ordine cinque di una colonia di virus.



PARTE II

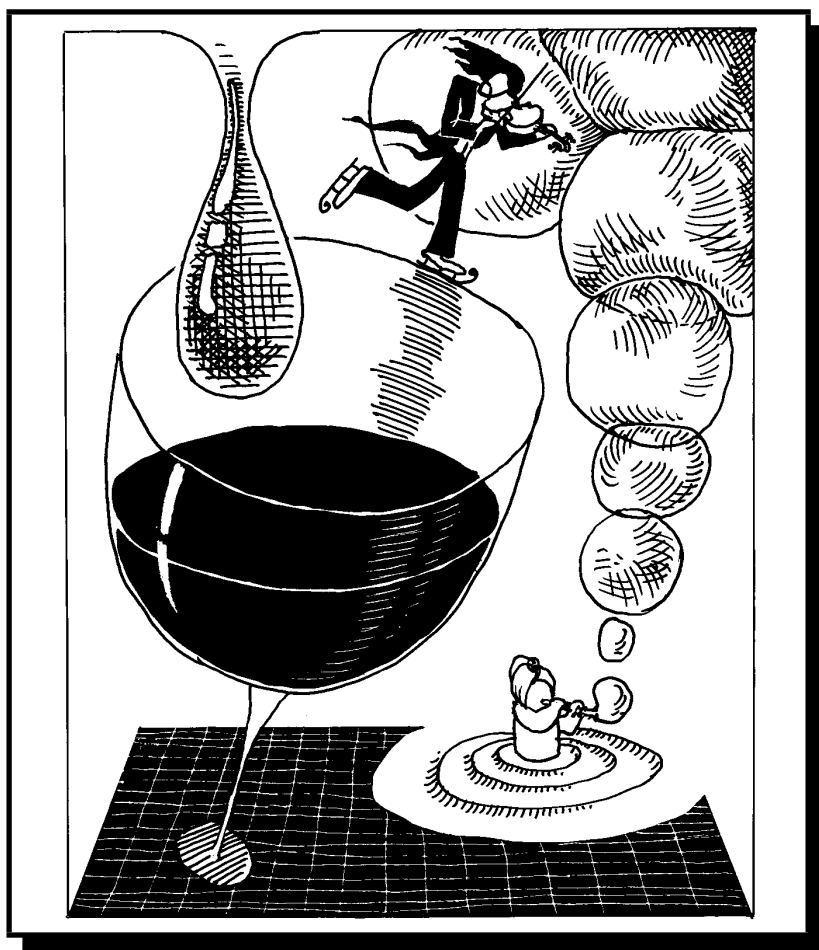
“Divertissement”: la fisica del Sabato sera

In questa parte i lettori saranno idealmente invitati a un modo forse non usuale di trascorrere qualche ora di un fine settimana: discutendo di fenomeni fisici che per loro natura trovano collocazione in serate conviviali o che ad eventi che le accompagnano fanno riferimento.

In molti salotti potete imbattervi in una sorta di “lampada magica” nella quale l’alternarsi di miscele e demiscelazioni di gocce multicolori su uno sfondo di luce diffusa crea fenomeni molto belli e complessi, che cercheremo di illustrarvi in termini di principi fisici.

Forse non avete mai osservato come la fisica giochi un ruolo di rilievo nella produzione del vino, nel frizzare degli spumanti, nel tintinnio di calici bene auguranti. Inoltre questa disciplina ha fornito un metodo scientifico per la certificazione della provenienza dei vini da questa o quella area e per provare o meno la mancanza di manipolazione. E ancora, quali sono le ragioni dello struggente e melodioso suono di un violino? Perché lo scorrere dell’archetto sulle corde produce onde sonore?

In breve, si mostrerà come anche in argomenti apparentemente frivoli o nella conversazione in un viaggio in treno si possa fare della fisica del tutto “rispettabile” e come sia facile cogliere interessanti aspetti e ricavare piacevoli commenti da eventi, oggetti o situazioni comuni nella vita di tutti i giorni.



Capitolo 9

Conversando nel viaggio in treno.

Recentemente ci si trovava in un gruppetto di viaggiatori, in una carrozza dell' espresso Venezia-Napoli. Il treno viaggiava veloce (ad una velocità di circa 150km/h) e attraverso il finestrino comparivano paesaggi quasi presi a prestito dalle tele dei maestri del Rinascimento. In somiglianza con questi quadri il paesaggio era collinare; il treno ora correva su un ponte ora entrava in una galleria. In una di queste gallerie, particolarmente lunga, nei pressi di Firenze, inaspettatamente ai viaggiatori si sono “tappate” le orecchie, come succede in un aereo al momento del decollo o dell'atterraggio. Nello scompartimento tutti muovevano la testa, cercando di liberarsi da questa spiacevole sensazione. Quando il treno fuoriuscì dall'angusto tunnel, il disturbo scomparve spontaneamente.

Alcuni dei viaggiatori, considerato che beneficiavano della presenza di un fisico, decisero di chiedere lumi e prese così l'avvio una piacevole discussione che intrattenne i viaggiatori e della quale ora desideriamo rendervi partecipi.

L'effetto sui timpani delle nostre orecchie era manifestamente legato alla differenza di pressione che si era prodotta. Di primo acchito sembrava evidente che al passaggio del treno la pressione dell'aria compresa tra le pareti della galleria dovesse aumentare. Tuttavia, quanto più ci si addentrava nell'esame del fenomeno tanto più questa considerazione si evidenziava poco valida. Approfondiamo l'argomento.

Consideriamo un treno con un'area della sezione trasversale S_t , che si muove ad una velocità costante v_t in una lunga galleria con un'area della sezione trasversale S_0 . Passiamo al sistema di riferimento solidale con il treno. Riterremo la corrente dell'aria stazionaria, mentre trascureremo la viscosità dell'aria. In questo caso il moto delle pareti della galleria rispetto

al treno può essere trascurato: in virtù dell'assenza di viscosità esso non influisce sulla corrente d'aria. Inoltre considereremo il treno abbastanza lungo da poter ignorare gli effetti agli estremi anteriore e posteriore dello stesso, mentre la pressione dell'aria nella galleria la considereremo costante.

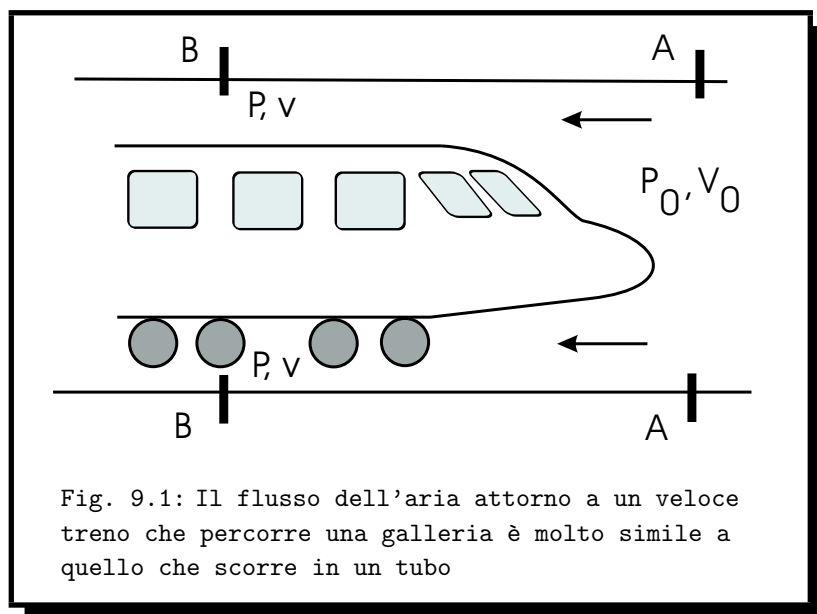
Cosicché, trascurando i dettagli meno rilevanti, siamo passati dal movimento del treno reale ad un modello fisico semplificato, che possiamo ora descrivere con il ricorso al linguaggio universale delle equazioni. Abbiamo un tubo (la galleria), al cui interno si trova un cilindro (il treno) con gli estremi aerodinamici. Attraverso il tubo scorre aria: lontano dal treno la pressione dell'aria p_0 è uguale a quella atmosferica, mentre la velocità del flusso d'aria è uguale in modulo e opposta in direzione alla velocità $\vec{v}_t$ con la quale, fino al passaggio nel nostro sistema di riferimento, si muoveva il treno. Esaminiamo una certa sezione, lontana dalle estremità del treno (affinché le nostre ipotesi non siano invalidate). Indichiamo la pressione dell'aria in questa sezione con p e la velocità del flusso d'aria con v . Queste grandezze possono essere collegate a v_t e p_0 mediante l'equazione di **Jacob Bernoulli**:

$$p + \frac{\rho v^2}{2} = p_0 + \frac{\rho v_t^2}{2}, \quad (9.1)$$

dove ρ è la densità dell'aria. Si osservi che abbiamo scritto l'equazione di Bernoulli nella forma strettamente valida per liquidi incompressibili.

Ciò impone ulteriori limitazioni, che vale la pena di considerare più dettagliatamente. Può, nel problema in esame, la densità dell'aria essere considerata costante? La domanda non è delle più semplici e una risposta rigorosa può essere data soltanto con un esame molto più complesso, sulla base dell'equazione di Bernoulli relativa a un liquido comprimibile (detta integrale di Bernoulli). Riporteremo soltanto le ipotesi fisiche più significative. Se la velocità del flusso non è troppo alta (cosa ciò significhi, sarà chiarito più oltre) e poichè abbiamo trascurato la viscosità dell'aria, si può trascurare la variazione della temperatura. La densità cambierà pertanto allo stesso modo della pressione. Ci aspettiamo di ottenere una piccola variazione della pressione, che sarà determinata dalla densità dell'aria. Quest'ultima può essere ricavata dall'equazione di **Mendeleev-Klaeperton**, utilizzata per l'aria lontana dal treno:

$$\rho = \frac{m}{V} = \frac{p_0 \mu}{RT}, \quad (9.2)$$



dove $\mu = 29$ è la massa molecolare dell'aria, T è la temperatura assoluta e R la costante dei gas.

Non è difficile intuire il limite ove cade la ipotesi sulla costanza della densità: se la velocità del flusso d'aria in una qualsiasi delle sezioni del tubo diventa confrontabile con la velocità quadratica media del moto termico delle molecole, allora non si potrà parlare di costanza della densità nelle diverse sezioni. In realtà, proprio questa velocità determina il tempo caratteristico di stabilizzazione per la densità media del gas in un volume macroscopico. Pertanto, con velocità del flusso elevata la densità non potrà stabilizzarsi nel passaggio da una sezione ad un'altra. Tuttavia, come vedremo più avanti, questo limite può risultare insufficiente. Per il momento considereremo la densità costante.

Nell'equazione compaiono due grandezze incognite ρ e v ed è quindi necessario disporre di un'altra relazione. Essa ci viene fornita dalla condizione di invariabilità della massa dell'aria che attraversa una qualunque sezione nell'unità di tempo:

$$\rho v_t S_0 = \rho v (S_0 - S_t). \quad (9.3)$$

(nota anche come condizione di continuità del flusso). Utilizzando tale

relazione e sostituendo l'espressione per la densità, ricaviamo la pressione dell'aria nella galleria:

$$p = p_0 \left[1 - \frac{\mu v_t^2}{2RT} \left(\frac{S_0^2}{(S_0 - S_t)^2} - 1 \right) \right]. \quad (9.4)$$

La quantità $\frac{\mu v_t^2}{2RT}$, che compare in questa espressione, evidentemente è adimensionale. Di conseguenza, la grandezza $\sqrt{RT/\mu}$ ha come dimensioni quelle della velocità. In essa, si può subito ravvisare la velocità quadratica media del moto termico delle molecole. Tuttavia, per noi sarà importante un'altra caratteristica del gas: la velocità di propagazione del suono. Tale velocità, come quella del moto termico delle molecole, è determinata dalla temperatura e dalla massa molecolare del gas. Inoltre il suo valore numerico dipende anche dal cosiddetto indice della curva adiabatica γ , caratteristico di ogni gas (per l'aria $\gamma = 1.41$):

$$v_s = \sqrt{\gamma \frac{RT}{\mu}}. \quad (9.5)$$

In condizioni normali $v_s = 335 \text{ m/s}$ (ossia circa 1200 km/h). Utilizzando la espressione per v_s potremo riscrivere la pressione in una forma agevole

$$p = p_0 \left[1 - \frac{\gamma}{2} \frac{v_t^2}{v_s^2} \left(\frac{S_0^2}{(S_0 - S_t)^2} - 1 \right) \right]. \quad (9.6)$$

Ora è giunto il momento di fermarsi e riflettere. Abbiamo calcolato la pressione nei pressi della superficie del treno all'interno della galleria. Tuttavia le nostre orecchie hanno sofferto non tanto per la pressione, ma per la sua variazione, quando il treno viaggia in galleria rispetto alla pressione p_0 durante il moto all'esterno della stessa. La differenza di pressione può essere determinata dalla equazione (9.6). La variazione relativa della pressione si scrive

$$\frac{\Delta p}{p_0} = \frac{p - p_0}{p_0} = -\frac{\gamma}{2} \left(\frac{v_t}{v_s} \right)^2 \left(\frac{S_0^2}{(S_0 - S_t)^2} - 1 \right). \quad (9.7)$$

Notiamo da tale relazione che all'entrata nella galleria la pressione diminuisce (anziché aumentare, come ingenuamente si poteva ritenere). Ora valutiamo l'entità dell'effetto.

Per una galleria ferroviaria si può assumere all'incirca $S_t = S_0/4$. Ponendo $v_t = 150 \text{ km/h}$ e $v_s = 1200 \text{ km/h}$ si ricava

$$\frac{\Delta p}{p_0} = -\frac{1.41}{2} \left(\frac{1}{8}\right)^2 \left(\left(\frac{4}{3}\right)^2 - 1\right) \approx 1\%.$$

Come ci si poteva attendere, la variazione della pressione risulta piccola. Di conseguenza, l'ipotesi fatta è accettabile: la variazione della densità dell'aria può essere effettivamente trascurata. La piccola variazione relativa di pressione non significa che l'organismo umano non sia in grado di avvertirla^a. Effettivamente se ricordiamo che $p_0 = 10^5 \text{ N/m}^2$ e consideriamo per l'area del timpano $\sigma \sim 1 \text{ cm}^2$, la forza in eccesso agente su di esso dall'interno vale $\Delta F = \Delta p_0 \cdot \sigma \sim 0.1 \text{ N}$, che non è difficile da avvertire (all'incirca come quella associata a dieci grammi peso).

Sembrerebbe che l'effetto sia stato giustificato. Tuttavia, qualcosa nella formula ottenuta prese a preoccupare i viaggiatori.

Infatti si può osservare che dalla Eq. (9.7) consegue che con una condizione usuale per i comuni treni cioè $\frac{v_t}{v_s} \ll 1$ (questo rapporto delle velocità si incontra costantemente in aerodinamica ed è chiamato numero di **Mach**) in una galleria molto angusta, quando $S_t \sim S_0$, la grandezza $|\Delta p|$ può diventare dell'ordine di p_0 ed anche superare questo limite!

È evidente che in questo caso debbono cadere alcune condizioni assunte precedentemente. Dove? Sembrerebbe che la condizione $v_t \ll v_s$ ci difenda del tutto da impreviste non validità! Cercammo di chiarirci ulteriormente le idee.

Supponiamo che effettivamente Δp , calcolato in base alla Eq. (9.7) sia dell'ordine di p_0 . Ciò significa che

$$\frac{v_t}{v_s} \left(\frac{S_0}{S_0 - S_t} \right) \sim 1,$$

e quindi

$$v_t S_0 \sim v_s (S_0 - S_t).$$

^aVale la pena di ricordare che in biofisica esiste la cosiddetta legge di **Weber-Chefner**, in accordo alla quale qualsiasi variazione dell'azione esterna viene recepita dagli organi sensoriali soltanto quando la variazione relativa supera una determinata soglia. In secondo luogo, se la galleria è abbastanza lunga, l'organismo ha il tempo di adattarsi alle nuove condizioni e la sensazione spiacevole scompare per ripresentarsi solo all'uscita.

Confrontando l'ultima eguaglianza con la condizione di continuità (9.3) vediamo le conseguenze: $|\Delta p|$ diventa dell'ordine di p_0 quando la velocità di scorrimento della flusso d'aria tra il rivestimento del treno e le pareti della galleria risulta dell'ordine della velocità del suono. In questo caso, come abbiamo detto, è l'intera analisi che viene invalidata. Di conseguenza, la condizione per l'utilizzazione dell'espressione (9.7) non è solo $v_t \ll v_s$, ma una più restrittiva:

$$v_t \ll v_s \left(\frac{S_0 - S_t}{S_0} \right).$$

Per i treni e le gallerie ordinarie questa condizione si realizza praticamente in ogni caso.

Ciò nonostante la nostra ricerca sui limiti di applicabilità della formula ottenuta non deve essere considerata solo come una preoccupazione di rigore matematico. Il fisico deve sempre individuare "la regione operativa" del risultato ottenuto. Inoltre vi è un aspetto pratico da tener presente.

Negli ultimi anni sono state discusse possibilità di realizzare mezzi di trasporto basati su principi innovativi, anche di tipo ferroviario. Già molti anni fa è stata proposta l'idea di un treno che si muove per effetto della levitazione diamagnetica. Si tratta di un vagone, all'interno del quale viene trasportato un potente magnete superconduttore. Grazie al campo magnetico creato dal magnete, il vagone si solleva sul binario metallico e la resistenza al moto viene determinata soltanto dalle proprietà aerodinamiche del mezzo stesso. In Giappone e in Cina sono già stati realizzati modelli di questo tipo capaci di sviluppare velocità sino a 540 km/h , quasi la metà della velocità del suono. Ne discuteremo a proposito della superconduttività, nella quarta parte del testo.

Il passo successivo nello sviluppo di questo mezzo di trasporto è l'idea di sistemare il vagone in un tubo ermetico e di ridurre al suo interno la pressione atmosferica. Ecco quindi che i progettisti probabilmente si troveranno di fronte al problema da noi esaminato per il treno nella galleria con condizioni ancora più complesse, quando si creeranno le condizioni $v_t \sim v_s$ e $S_0 - S_t \ll S_0$.

Il moto dell'aria rarefatta è di carattere turbolento, la temperatura varia notevolmente e oggi la scienza non conosce ancora le risposte ai molti problemi che sorgono in simili circostanze. La nostra analisi vi permette di capire alcune delle difficoltà che gli scienziati incontrano sulla loro strada.

Infine vi proponiamo di riflettere su altri interrogativi che possono sorg-

ere in un vostro prossimo viaggio in treno:

- Perché il rumore prodotto dal treno in movimento aumenta bruscamente quando il treno entra in una galleria?
- Ultimamente su molte strade ferrate viene sistemato il cosiddetto “binario di velluto”, le rotaie vengono cioè saldate senza lasciare gioco tra di esse. Che cosa bisogna fare per bilanciare gli effetti di dilatazione termica quando cambia la temperatura?
- Perché quando due treni in corsa in direzioni opposte si incrociano i loro finestrini subiscono una sorta di onda d’urto. Dove è diretta la forza che agisce in questo caso sui vetri: verso l’interno o verso l’esterno del vagone?

Che bel viaggio, con questi pensieri!

Capitolo 10

Al suono del violino.

A tutti è capitato di apprezzare lo struggente e melodioso suono di un violino. L'evoluzione storica di questo strumento musicale vede come protagonista la città di Cremona, adagiata sulla riva sinistra del fiume Po, tra distese di campi coltivati e lunghi filari di pioppi che accompagnano i corsi d'acqua, qua e là raggruppandosi in piccoli boschi.

In questa città, alla metà del secolo XVI, inizia con **Andrea Amati** il cammino della celeberrima liuteria che ha visto nei secoli successivi il susseguirsi di dinastie come quelle dei Bergonzi, dei Guarneri, dei Ruggeri e degli Stradivari, la cui fama percorse l'Europa tutta. Ancora oggi Cremona, con le sue oltre 130 botteghe liutarie costituisce il punto di riferimento per eccellenza nella liuteria internazionale. In particolare Cremona custodisce nel museo Stradivariano preziosissimi strumenti storici della collezione dei violini Stradivari (si vedano le fotografie in Fig. 10.1).

Quali sono i fenomeni fisici che controllano il fascino delle dolcissime note di un violino? Sarebbe troppo lungo e anche difficile spiegare compiutamente tutti i complicati fenomeni relativi alla generazione dei suoni nel violino e alle particolarità di risonanza che si realizzano in questo singolare strumento. Tuttavia siamo in grado di illustrare almeno la ragione per la quale compaiono le oscillazioni della corda quando viene sfregata uniformemente dall'archetto.

Nel movimento di un corpo in un mezzo qualsiasi compaiono forze di resistenza tendenti a rallentare il moto stesso. Il moto di un solido sulla superficie di un altro è ostacolato da forze di attrito. Nel caso di un corpo in moto in un liquido o in un gas compaiono l'attrito viscoso e la resistenza aerodinamica.

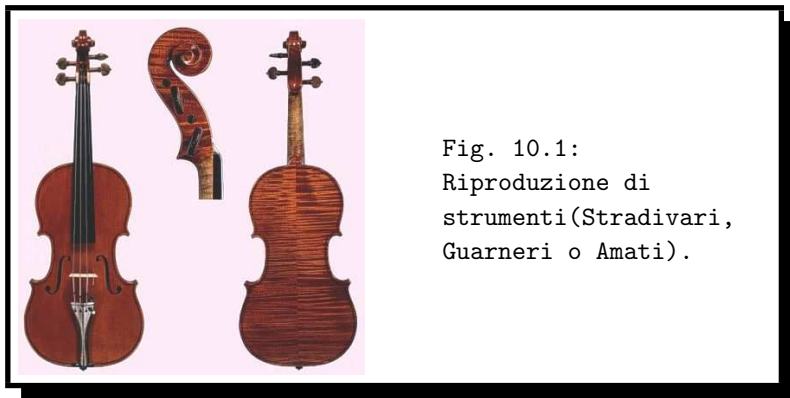


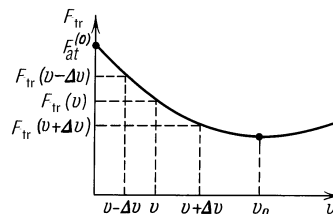
Fig. 10.1:
Riproduzione di
strumenti(Stradivari,
Guarneri o Amati).

L'interazione del corpo con un mezzo è un processo piuttosto complesso. Al termine di tale processo, dopo un tempo sufficientemente lungo generalmente l'energia del corpo viene trasformata in calore. In alcuni casi l'interazione con il mezzo può giocare anche un ruolo opposto, ossia aumentare l'energia di un sistema. In tali casi di regola si formano delle oscillazioni. Come esempio si può pensare all'attrito tra un armadio che viene spostato e il pavimento. Una situazione analoga si produce nel caso della forza d'attrito tra l'archetto del violino e la corda, il che provoca le oscillazioni della stessa. Come vedremo tra poco, la causa della comparsa delle oscillazioni è costituita dalla dipendenza della forza d'attrito dalla velocità. In particolare, le oscillazioni compaiono quando la forza d'attrito diminuisce all'aumentare della velocità.

Esaminiamo il processo di formazione delle oscillazioni meccaniche mentre ascoltiamo il dolce suono di un violino, provocato dal movimento dell'archetto. Tra l'archetto e la corda si generano evidentemente delle forze di attrito. Si può distinguere tra forze d'attrito statico e forze d'attrito dinamico o radente. Il primo tipo di forza compare tra corpi in contatto, immobili uno rispetto all'altro, mentre le forze di attrito dinamico riguardano il caso di scivolamento di un corpo sulla superficie di un altro. La forza d'attrito statico, come è ovvio, può assumere qualsiasi valore, a seconda della forza esterna: dallo zero fino al massimo F_{at}^0 . In questo caso la forza di attrito statico essendo uguale in modulo e opposta in direzione alla forza esterna.

La forza d'attrito radente dipende, invece, dal materiale di cui sono costituiti i corpi, dallo stato delle superfici a contatto e anche dalla velocità relativa di questi corpi. Di quest'ultimo aspetto parleremo più dettagliata-

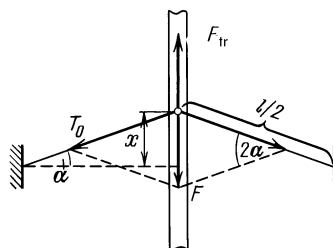
Fig. 10.2: Tipica dipendenza dell'attrito dalla velocità relativa.



mente. La dipendenza della forza d'attrito radente dalla velocità è diversa per corpi diversi. Spesso aumentando la velocità di scivolamento si ha una iniziale diminuzione della forza d'attrito radente, che per altro, ad un certo punto comincia a crescere. Nella figura 10.2 è rappresentata la dipendenza del modulo della forza d'attrito dalla velocità. La forza d'attrito tra l'archetto del violino e la corda ha proprio questo carattere. Il tratto verticale a $v = 0$ corrisponde alla forza d'attrito statico. Se la velocità relativa tra la corda e l'archetto si trova nel tratto di curva discendente $0 < v < v_0$, all'aumento della velocità di una quantità Δv corrisponde la diminuzione della forza d'attrito e al contrario, diminuendo la velocità la variazione della forza d'attrito è positiva (si veda la Fig. 10.2). Come ora vedrete, proprio grazie a questa particolarità l'energia della corda può crescere a spese delle forze d'attrito.

Durante il movimento iniziale dell'archetto, la corda si piega insieme ad esso. In questo caso, la forza d'attrito statico viene equilibrata dalle forze di tensione della corda (Fig.10.3).

Fig. 10.3: Quando la corda segue l'archetto senza scivolamento l'attrito statico compensa la risultante delle due forze di tensione.



La risultante F delle forze di tensione è proporzionale alla deviazione x

della corda dalla posizione d'equilibrio:

$$F = 2T_0 \sin \alpha \approx \frac{4T_0}{l}x,$$

dove l è la lunghezza della corda e T_0 è la forza di tensione, che in caso di piccole deviazioni può essere considerata costante. Per questo, nel movimento della corda insieme all'archetto, la forza F cresce e nel momento in cui essa diventa uguale alla forza d'attrito statico massima F_{at}^0 inizia lo scivolamento.

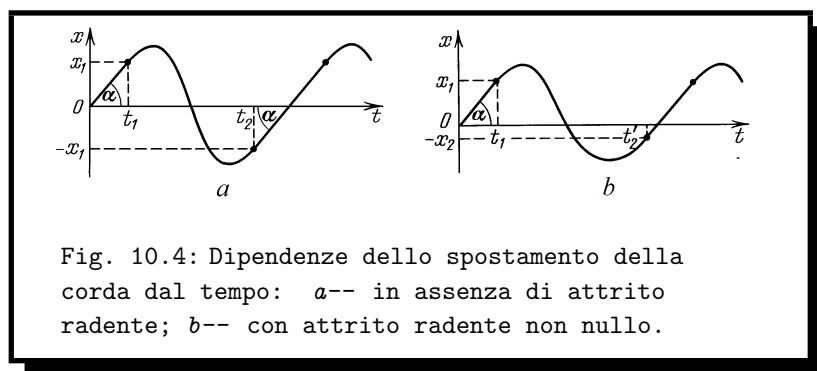
Per semplicità considereremo che al momento d'inizio dello scivolamento il cambiamento della forza d'attrito avvenga bruscamente, cioè tale forza verrà assunta diminuire dal valore massimo della forza d'attrito statico fino al valore della forza d'attrito radente, in genere assai minore. In altre parole, dopo l'inizio dello scivolamento il movimento della corda può essere considerato quasi libero.

Al momento del distacco la velocità della corda è uguale alla velocità dell'archetto e inizialmente la corda continua a spostarsi nella direzione del movimento dell'archetto. La risultante delle forze di tensione, però, non essendo compensata da nessun'altra forza, frenerà il movimento della corda, rallentandolo. Ad un certo momento la velocità della corda si annullerà e la corda inizierà a muoversi all'indietro. Dopo la deviazione massima dalla posizione d'equilibrio nel senso opposto a quello iniziale, la corda si muoverà di nuovo nella direzione del movimento dell'archetto. L'archetto invece continua a muoversi uniformemente con una velocità u . Ad un dato istante le velocità dell'archetto e della corda diverranno eguali e non essendoci più slittamento compare a quel punto la forza d'attrito statico, uguale alla risultante delle forze di tensione.

Con l'ulteriore movimento della corda fino alla posizione di equilibrio, le forze di tensione diminuiscono e quindi diminuisce la forza d'attrito statico. Quando la corda supera la posizione d'equilibrio, il processo si ripete.

Il grafico relativo alla dipendenza della deviazione della corda dal tempo è riportato nella figura 10.4, *a*. Il movimento della corda è periodico e in ogni periodo si hanno due intervalli diversi. Per esempio, nell'intervallo $0 < t < t_1$ la corda si muove con l'archetto ad una velocità costante, così la deviazione x dipende linearmente dal tempo ($\tan \alpha \propto u$). Al tempo t_1 avviene il distacco e per $t_1 < t < t_2$ la variazione di x con il tempo ha andamento sinusoidale. Al tempo t_2 , quando la tangente alla sinusoide ha la stessa inclinazione α del tratto rettilineo iniziale (condizione d'uguaglianza

delle velocità), la corda segue di nuovo il movimento dell'archetto.



La fig.10.4, *a* corrisponde al caso ideale in cui la forza d'attrito radente è assente e non ci sono perdite d'energia nel cammino libero della corda. Il lavoro totale della forza d'attrito statico lungo i tratti lineari, in questo caso, è uguale a zero, poiché quando x è negativo si compie un lavoro negativo (la forza d'attrito è diretta contro il movimento) e con $x > 0$ si compie un lavoro uguale in modulo, ma positivo.

Cosa accade invece nel caso in cui la forza dell'attrito radente è diversa da zero? L'attrito dissipa energia. Il movimento della corda in caso di sfregamento è descritto dal grafico riportato nella figura 10.4, *b*. Per spostamenti negativi questa curva è più smussata rispetto alla parte relativa allo spostamento positivo. Per questo l'archetto riafferra la corda nel caso di uno spostamento negativo x_2 più piccolo in modulo dello spostamento positivo x_1 corrispondente al distacco. Di conseguenza, la forza d'attrito statico durante il contatto della corda con l'archetto compie nel periodo un lavoro positivo dato da

$$A = \frac{k(x_1^2 - x_2^2)}{2}.$$

dove $k = 4T/l$ è il coefficiente di proporzionalità tra la forza d'attrito statico e lo spostamento della corda. Questo lavoro compensa la dissipazione d'energia della corda dovuta alle forze d'attrito radente. Le oscillazioni della corda non si smorzano nel tempo.

In genere, per compensare l'energia della corda mediante le forze d'attrito non è necessario che ci sia accoppiamento della corda con l'archetto. È suf-

ficiente che la velocità relativa v dell'archetto e della corda durante le oscillazioni di quest'ultima si trovi nell'intervallo del tratto discendente della curva dell'attrito radente al variare della velocità. Analizziamo più dettagliatamente il fenomeno dell'eccitazione delle oscillazioni della corda in questo caso.

Supponiamo che l'archetto si muova con velocità costante u e che la corda sia spostata dalla posizione d'equilibrio nel punto x_0 , così che la risultante $F(x_0)$ delle forze elastiche equilibri la forza d'attrito radente $F_{at}(u)$. Se la corda si sposta casualmente nella direzione del movimento dell'archetto, la velocità relativa diminuirà. Così la forza d'attrito cresce (la velocità relativa corrisponde al tratto discendente!) e la corda si sposta ulteriormente. La forza elastica ad un certo punto supererà la forza d'attrito (la forza elastica è proporzionale al valore assoluto dello spostamento, mentre la forza d'attrito radente non può superare il valore massimo della forza d'attrito statico) e la corda inizierà a muoversi nel senso inverso. Essa attraverserà la posizione d'equilibrio, se ne discosterà di nuovo, si fermerà ancora e così via. Quindi si innescheranno le oscillazioni della corda.

E' importante che queste oscillazioni non si smorzino nel tempo. Nel movimento della corda con una velocità Δv nella direzione dell' archetto la forza d'attrito compie un lavoro positivo, mentre nel movimento inverso il lavoro è negativo. Tuttavia la velocità relativa $v_1 = u - \Delta v$ nel primo caso è minore della velocità $v_2 = u + \Delta v$ nel secondo caso e di conseguenza la forza d'attrito $F_{at}(u - \Delta v)$ al contrario è maggiore di $F_{at}(u + \Delta v)$. Quindi il lavoro positivo delle forze d'attrito durante il movimento della corda nella direzione dell'archetto è maggiore del lavoro negativo nella direzione inversa e in generale le forze d'attrito compiono un lavoro positivo. L'ampiezza delle oscillazioni aumenterà, ma solo fino ad un certo limite. Con $v > v_0$ (si veda la Fig.10.2) la velocità v supera il minimo e quindi il lavoro negativo della forza d'attrito diventa maggiore di quello positivo. L'energia, e quindi anche l'ampiezza delle oscillazioni, diminuiscono.

In conclusione, si stabilisce un'ampiezza tale che il lavoro totale delle forze d'attrito è uguale a zero (ossia, più precisamente, questo lavoro compensa la dissipazione d'energia dovuta alla resistenza dell' aria, al carattere anelastico delle deformazioni e così via). Con questa ampiezza costante si avranno oscillazioni della corda che non si smorzano.

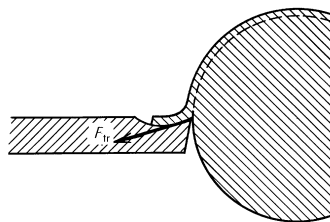
Il generarsi delle oscillazioni sonore durante il moto di un corpo solido sulla superficie di un altro è molto frequente. L'attrito del cardine di una porta può provocare il cigolio della stessa. Scricchiolano le assi, le scarpe...

Le dita delle mani possono produrre un fruscio se le si fa scorrere su una superficie liscia. In questo caso i fenomeni che si verificano hanno molto in comune con l'eccitazione delle oscillazioni della corda del violino. All'inizio non c'è scivolamento e compare una deformazione elastica. Quindi avviene il distacco e si innescano le oscillazioni del corpo. Le oscillazioni non si smorzano, il lavoro delle forze d'attrito (grazie alla loro caratteristica) fornisce l'energia necessaria.

Cambiando la dipendenza delle forze d'attrito dalla velocità il cigolio scompare. È noto, per esempio, che è necessario oliare le superfici ove avvengono sfregamenti. La forza d'attrito viscoso (per piccole velocità) è proporzionale alla velocità. Non si realizzano più le condizioni necessarie per provocare oscillazioni. Al contrario, quando si vogliono provocare oscillazioni, le superfici vengono trattate in modo particolare, così da ottenere una marcata diminuzione delle forze d'attrito all'aumentare della velocità. Proprio per questo, per esempio, l'archetto del violino viene strofinato con la pece.

La conoscenza delle leggi sull'attrito aiuta a risolvere importanti problemi pratici. Durante la lavorazione del metallo sul tornio a volte si verifica una vibrazione della lama. Queste oscillazioni sono provocate dalle forze d'attrito tra la lama e il truciolo metallico che scivola sulla superficie durante la tornitura (Fig. 10.5). La forza d'attrito dipende dalla velocità del truciolo (velocità di lavorazione) e per una serie di metalli ad alta qualità può diminuire bruscamente a certe velocità. Per questa via si possono spiegare le oscillazioni dell'utensile. Per evitare queste vibrazioni si ricorre, ad esempio, ad una speciale affilatura della lama così che il truciolo non possa scivolare. Si elimina così la causa della formazione delle oscillazioni.

Fig. 10.5: Le vibrazioni della lama dello strumento possono essere eliminate con una opportuna scelta dell'angolo di taglio.



Capitolo 11

Calici sonori e calici silenti.

Non è una novità che con calici di vetro si possano produrre suoni. Toni musicali possono essere ottenuti in un modo tutto particolare, come voi stessi potrete sperimentare. Immergete un dito nell'acqua e passatelo con cura sul bordo di un calice, mantenendolo bagnato. All'inizio il calice produrrà un suono stridulo, poi, quando l'orlo sarà asciutto, i suoni diventeranno più melodici. Cambiando la pressione del dito si può cambiare anche il tono del suono prodotto. L'altezza del tono dipende anche dalle dimensioni, dallo spessore delle pareti del calice e infine dalla quantità di liquido in esso contenuto.

Non tutti i calici possono produrre suoni piacevoli e la ricerca di un calice idoneo può essere lunga. Meglio suonano melodiosamente (e non scricchiolano) i calici molto sottili a forma di paraboloide di rotazione, su uno stelo lungo e affusolato. Il tono può cambiare aggiungendo acqua: quanta più acqua vi è nel calice, tanto più basso è il suono. Quando il livello dell'acqua raggiunge la metà del calice, sulla superficie si formano delle onde, prodotte dalle vibrazioni delle pareti del calice. Le onde più marcate si verificano nel punto dove si trovava il dito.

Sulla base del fenomeno descritto, lo scienziato **Benjamin Franklin** (che, in particolare, ha scoperto l'elettricità atmosferica) aveva realizzato uno strumento musicale molto originale, che ricorda quello descritto da Hoffmann nella favola "Il piccino Zaches, soprannominato Zinnober". Ad un asse veniva fissata una serie di tazze ben levigate, forate al centro e sistemate a distanza uguale l'una dall'altra. Sotto la cassa in cui era contenuto tale asse era fissato un pedale (come nella macchina da cucire) che lo faceva ruotare. Al tocco delle dita bagnate dell'esecutore sulle tazze rotanti,

i suoni si rafforzavano fino al fortissimo o diminuivano fino al sussurro.

Certo oggi è difficile ipotizzare un largo uso di uno strumento musicale siffatto. Tuttavia chi lo aveva ascoltato assicurava che l'armonia dei suoni aveva un effetto sorprendente, sia sull'esecutore sia sugli ascoltatori. Nel 1763 Franklin regalò il suo strumento all'inglese Miss Davis che per alcuni anni lo fece conoscere in molti paesi d'Europa. In seguito questo straordinario strumento scomparve senza lasciar traccia. È possibile che lo stesso Ernest Teodor Amadeus Hoffmann sia venuto a sapere dell'esistenza di questo strumento e che lo abbia utilizzato ne "Il piccino Zaches".

Dal momento che ci intratteniamo con i calici, possiamo osservare come di solito non siano opportuni brindisi mediante l'urto di calici che contengono spumanti. Alla base di questa inopportunità vi è un preciso motivo fisico: il suono prodotto dall'urto di calici - anche di cristallo - pieni di spumante è sordo. Perché i calici di spumante non tintinnano, sono pressoché silenti?

La melodiosità e la sfumatura "cristallina" del suono è data dalle oscillazioni prodotte nella cassa di risonanza, ossia in questo caso il calice, sia ad alta frequenza ($\nu \sim 10 - 20 \text{ KHz}$) sia ultrasonore ($\nu > 20 \text{ KHz}$). Nell'urto di calici vuoti o contenenti bevande non gassate, si producono oscillazioni che risuonano anche abbastanza a lungo. Pertanto si pensa subito che la causa dell'assenza di suono in caso di urti tra calici contenenti spumanti sia da associare alle bolle di anidride carbonica, che si liberano copiosamente quando si stappa la bottiglia. È possibile che esse provochino un forte smorzamento delle oscillazioni ad onde corte nel calice, così come le fluttuazioni della densità delle molecole nell'atmosfera disperdono fortemente le radiazioni della luce solare (ritornate al capitolo "Nell'azzurro dello spazio")?

Perfino per suoni a frequenze al limite superiore della percezione dell'orecchio umano ($\nu \sim 20 \text{ KHz}$), la lunghezza dell'onda nell'acqua, $\lambda = c/\nu \sim 10 \text{ cm}$ ($c = 1450 \text{ m/s}$ è la velocità del suono nell'acqua), supera di molto le dimensioni delle bolle di anidride carbonica nello spumante ($r_0 \sim 1 \text{ mm}$). Pertanto la diffusione delle onde sonore in una forma analoga alla diffusione Rayleigh della luce sembrerebbe giocare un ruolo irrilevante. Tuttavia cerchiamo di analizzare criticamente il ruolo dell'onda a lunghezza minima $\lambda_{\min}$ corrispondente alla massima frequenza percepibile. Per semplicità non considereremo la forma reale del calice, ma lo immagineremo come una scatola rettangolare. Supponiamo che in essa si insedi un'onda sonora piana, costituita da un'onda di compressione e decompressione.

La pressione in eccesso nel mezzo, in caso di propagazione dell'onda piana, può essere scritta nella forma

$$P_e(x, t) = P_0 \cos\left(\frac{2\pi x}{\lambda} - \omega t\right), \quad (11.1)$$

dove P_0 è l'ampiezza delle oscillazioni pressorie, ω è la frequenza del suono, λ è la lunghezza dell'onda sonora corrispondente; x è la coordinata del punto in esame lungo la direzione di propagazione dell'onda. Poiché anche $\lambda_{\min}$ supera le dimensioni del calice, per le onde sonore la funzione $P_e(x, t_0)$ (cioè il *campo di pressione*) nel dato istante t_0 , nei limiti del volume del calice, cambia lentamente al variare di x , ossia il primo addendo nell'argomento del coseno non è importante (poiché $x \ll \lambda_{\min}$). Il ruolo principale nella variazione della pressione in eccesso all'interno del calice, quindi, è giocato dal secondo addendo nell'argomento del coseno. Pertanto si può ritenere che all'interno del calice si stabilisca un campo di sovra-pressione praticamente omogeneo (non dipendente da x) e che cambia rapidamente col tempo:

$$P_e(t) = P_0 \cos \omega t, \quad (11.2)$$

La pressione totale nel liquido all'interno del calice è determinata dalla somma di $P_e(t)$ e della pressione atmosferica:

$$P_e(t) = P_{\text{atm}} + P_0 \cos \omega t.$$

Rimane ancora un passo per capire le cause del rapido smorzamento del suono nei calici di spumante. Il liquido saturo di gas è un “*mezzo acusticamente non lineare*”. Dietro questa espressione scientifica si nasconde quanto segue. La solubilità del gas nel liquido dipende dalla pressione: quanto maggiore è la pressione tanto più gas viene “intrappolato” nell'unità di volume di liquido. Come abbiamo visto, con l'emissione del suono dai calici all'interno si forma un campo alternato di pressione. Quando la pressione nel liquido diventa più bassa di quella atmosferica, in esso si produce l'emissione di bolle. Naturalmente, l'emissione di gas cambia la semplice legge armonica della variazione della pressione col tempo da noi considerata. Questo intendiamo quando definiamo un liquido saturo di gas come un mezzo acustico non lineare. Infatti nell'emissione di gas si dissipa l'energia delle oscillazioni ed esse si smorzano rapidamente.

Quando si fanno collidere i calici, si sviluppano inizialmente oscillazioni a diverse frequenze. In virtù del meccanismo descritto, le oscillazioni ad alta

frequenza si smorzano molto più rapidamente di quelle a bassa frequenza e quindi percepiamo soltanto un suono sordo, privo della sua sfumatura cristallina ad alta frequenza.

Le bolle in un liquido possono non soltanto smorzare le onde sonore, ma al contrario, in determinate condizioni, anche emettere suoni. Recentemente è stato scoperto che sotto l'azione di un'intensa radiazione laser piccole bolle d'aria in acqua generano onde sonore. Questo effetto è determinato dall'arrivo del raggio laser sulla superficie della bolla dalla quale, a seguito del fenomeno della riflessione totale, esso può essere riflesso. A seguito di questa specie di urto la bollicina vibra per un certo tempo (fino allo smorzamento delle oscillazioni) provocando onde sonore nel mezzo circostante. Valutiamone la frequenza. Esiste tutta una serie di importanti fenomeni che, nonostante l'apparente differenza, sono descritti dalla stessa equazione alla quale del resto abbiamo già fatto ricorso: l'equazione dell'oscillatore armonico. Tale equazione descrive processi oscillatori di diverso tipo: oscillazioni di un peso su una molla, di atomi nelle molecole e nei cristalli, della carica sulle facce di un condensatore in un circuito LC e così via. Tutti questi fenomeni hanno in comune la presenza di una forza detta elastica, che dipende linearmente dallo spostamento. Tale forza tende a riportare il sistema in posizione di equilibrio dopo che esso se ne è allontanato a causa di un'azione esterna. Anche una bollicina d'aria in un liquido costituisce appunto un sistema oscillatorio. La frequenza propria delle sue oscillazioni può essere valutata in base alla formula per la frequenza delle oscillazioni del peso sulla molla, cercando di individuare cosa giocherà il ruolo di costante elastica in quel caso.

Come primo candidato per il ruolo di costante elastica si ha il coefficiente di tensione superficiale del liquido σ : esso ha le stesse dimensioni della costante elastica, cioè forza per unità di lunghezza. In luogo della massa del peso nella formula per la frequenza propria delle oscillazioni dovrà esser tenuta in conto la massa di liquido coinvolta nelle oscillazioni della bolla. Questa grandezza è dell'ordine del volume della bolla moltiplicato per la densità dell'acqua: $m \sim \rho r_0^3$. Quindi la frequenza propria delle oscillazioni della bolla d'aria nell'acqua sarà descritta dall'equazione

$$\nu_1 \sim \sqrt{\frac{k_1}{m}} \sim \frac{\sigma^{\frac{1}{2}}}{\rho^{\frac{1}{2}} r_0^{\frac{3}{2}}}.$$

Tuttavia questa non è l'unica soluzione possibile. Vi è un importante

parametro che non è ancora intervenuto: la pressione dell'aria nella bolla. Se tale pressione P_0 viene moltiplicata per il raggio della bolla, si otterrà ugualmente una grandezza che ha la dimensione di una costante elastica (N/m): $k_2 \sim P_0 r_0$. Sostituendo questa nuova espressione nella formula per la frequenza propria delle oscillazioni, otterremo una frequenza completamente diversa:

$$\nu_2 \sim \sqrt{\frac{k_2}{m}} \sim \frac{P_0^{\frac{1}{2}}}{\rho^{\frac{1}{2}} r_0}.$$

Tra le due frequenze trovate qual'è la vera? Entrambe, nel senso che esse corrispondono ai diversi tipi di oscillazioni della bolla d'aria. Il primo tipo è costituito dalle oscillazioni che la bolla compie dopo la sua iniziale compressione (a seguito dell'urto con il raggio laser). Durante queste oscillazioni la forma della bolla cambia e con questa anche l'area della superficie, mentre il volume rimane invariato. In questo caso la forza elastica viene determinata effettivamente dal coefficiente di tensione superficiale^a. Tuttavia è possibile anche un altro tipo di oscillazione. Se si comprime uniformemente da tutti i lati la bollicina d'aria nel liquido e poi la si lascia libera, essa oscillerà grazie alle forze di pressione. Appunto a queste oscillazioni radiali, con variazione del volume, corrisponde la seconda delle frequenze ν_2 da noi trovate.

Nel caso di eccitazione delle oscillazioni con il raggio laser l'azione esterna non è simmetrica e il tipo di oscillazione delle bollicine sarà più vicino al primo caso. Se si conoscono le dimensioni delle bollicine, in base alla frequenza del suono da esse generate si potrà giudicare il tipo di oscillazioni prodotte. Nelle esperienze alle quali abbiamo fatto riferimento questa frequenza era di $3 \cdot 10^4 \text{ Hz}$, mentre non si conosce con precisione le dimensioni delle piccolissime bolle d'aria nell'acqua: si può dire che sono dell'ordine di frazioni di millimetro. Sostituendo nelle formule corrispondenti $\nu_0 = 3 \cdot 10^4 \text{ Hz}$, $\sigma = 0.07 \text{ N/m}$, $P_0 = 10^5 \text{ Pa}$, $\rho = 10^3 \text{ kg/m}^3$ troviamo

^aOccorre sottolineare che le oscillazioni, nel corso delle quali il volume della bolla rimane invariato, possono essere dei più svariati tipi: dalla normale deformazione periodica della bolla, alla sua trasformazione a vibrazioni quasi nella forma di una ciambella. Le frequenze in questi casi possono peraltro differire solo di poco, come ordine di grandezza rimanendo date da

$$\nu_1 \sim \frac{\sigma^{\frac{1}{2}}}{\rho^{\frac{1}{2}} r_0^{\frac{3}{2}}}.$$

che le dimensioni caratteristiche delle bolle generanti il suono durante le oscillazioni del primo o del secondo tipo sono rispettivamente

$$r_1 \sim \frac{\sigma^{\frac{1}{3}}}{\rho^{\frac{1}{3}} \nu_0^{\frac{2}{3}}} = 0.05 \text{ mm};$$
$$r_2 \sim \frac{P_0^{\frac{1}{2}}}{\rho^{\frac{1}{2}} \nu_0} = 0.3 \text{ mm}.$$

Come si vede queste dimensioni non sono molto diverse e non è possibile discriminare in base ad esse il tipo di oscillazione effettivamente prodottosi nelle condizioni dell'esperimento in esame. Comunque le dimensioni delle bolle che abbiamo ottenuto nelle nostre stime sono poco diverse da quelle che si presentano nell'esperienza quotidiana.

Capitolo 12

Osservando la lampada magica.

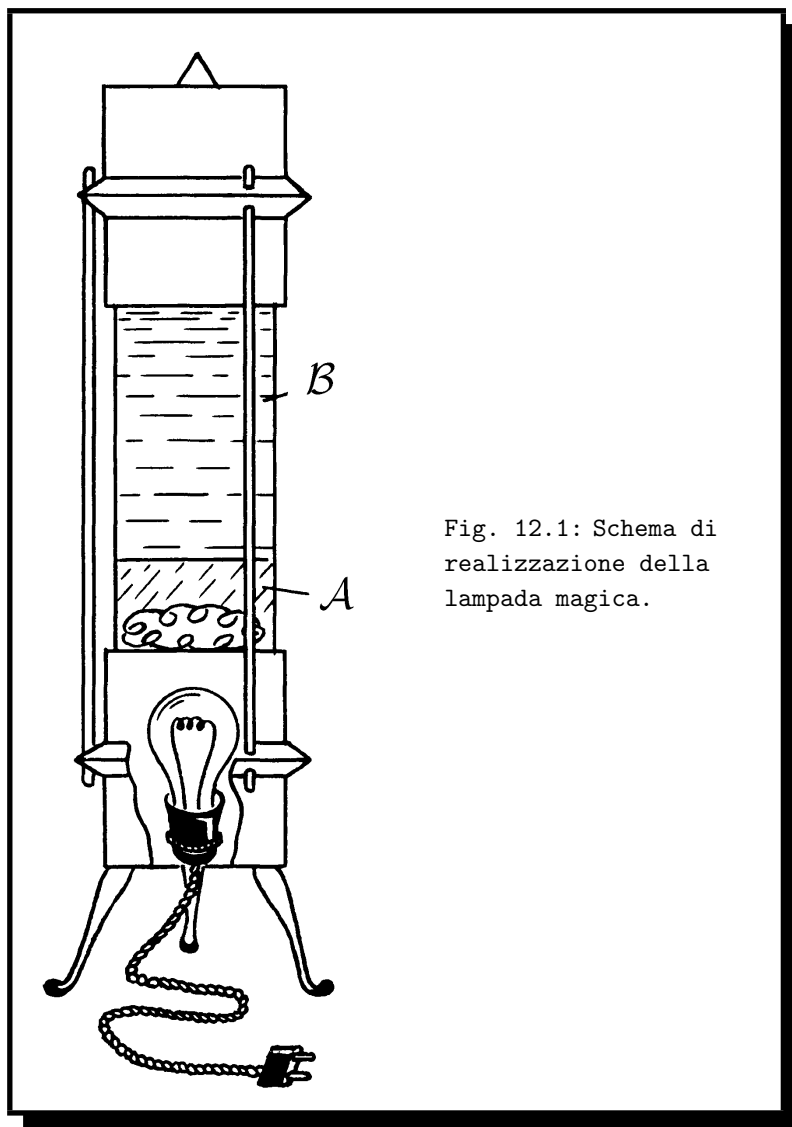
Le fotografie della terza pagina della copertina non sono state realizzate su Soljaris, né scattate da una navicella cosmica che entra nelle tenebrose profondità dell'atmosfera di Giove o dall'oblò di un batiscafo che si è avvicinato ad un vulcano subacqueo in eruzione. Esse riproducono una lampada detta "Lava-lamp", che non è infrequente trovare sui tavoli di un salotto o a fianco di un televisore. Rilassiamoci, un Sabato sera, e dopo aver inserito la spina nella presa di alimentazione osserviamo la lampada con attenzione: scopriremo che essa cela molti fenomeni, belli e complessi. È più conveniente effettuare le osservazioni al buio, quando la lampada magica costituisce l'unica fonte di luce.

Il meccanismo della lampada è piuttosto semplice. Si tratta di un bulbo cilindrico trasparente, alla cui base, sotto un fondo di vetro, è montata una lampadina elettrica che oltre ad illuminare ha la funzione di elemento riscaldante. Il vetro sul fondo è rivestito con un filtro colorato e lungo il suo perimetro è disposta una spirale metallica (Fig. 12.1).

Il bulbo è riempito per circa $1/6$ del suo volume di una sostanza quale cera o ad essa simile (che in seguito chiameremo *sostanza A*) mentre, quasi fino al limite superiore, è aggiunto un liquido trasparente (che chiameremo *sostanza B*). In base a quale considerazioni vengono scelte queste sostanze e quali proprietà esse debbano avere, lo spiegheremo più avanti, studiando i fenomeni che avvengono nella lampada.

I fenomeni che hanno luogo nella lampada, possono essere divisi in alcune fasi. Chiameremo convenzionalmente la prima "*fase di quiete e di accumulo di forze*".

La sostanza *A* è amorfa, ossia non ha una struttura microscopica per-



fettamente regolare. All'aumentare della temperatura essa si ammorbidisce e passa gradualmente allo stato liquido. Rammentiamo l'importante differenza tra il passaggio allo stato liquido di una sostanza cristallina e di una amorfa. Per un materiale cristallino questo passaggio avviene ad una

definita temperatura e richiede un apporto di energia: il *calore di fusione* necessario per disarticolare i legami chimici responsabili della struttura cristallina. Per quanto riguarda invece l'amorfo, lo stato solido e quello liquido non presentano differenze sostanziali, come abbiamo già avuto occasione di illustrare (paragrafo 8.2). Cosicché, aumentando la temperatura la viscosità della sostanza amorfa diminuisce ed essa diventa sempre più fluida.

La lampadina illumina dal basso l'interno del cilindro, con una luce rosso-verde dovuta al filtro, e come abbiamo detto agisce anche da fonte di calore. Sul fondo, accanto alla lampada, si forma una "macchia bollente" (una specie di regione di alta temperatura). In questa regione la sostanza $\mathcal{A}$ inizia ad ammorbidirsi, mentre né la crosta superiore, né tanto meno il liquido $\mathcal{B}$ aumentano di temperatura, per il momento restando relativamente freddi. A seguito del riscaldamento una parte maggiore della sostanza viene liquefatta e la crosta solida diventa più sottile. A causa della dilatazione termica, il volume degli strati inferiori della sostanza $\mathcal{A}$ che vengono liquefatti tende a crescere; la pressione sotto la crosta aumenta sino a che il liquido $\mathcal{A}$ rompe la crosta solida e si proietta verso l'alto, producendo bolle. Sul fondo si potrebbe dire che una sorta di vulcano è entrato nella fase di eruzione. "La fase di quiete e accumulo di forze" è completata, e viene sostituita dalla "fase d'attività vulcanica" (si veda la Fig.2 sulla terza pagina della copertina).

Le sostanze $\mathcal{A}$ e $\mathcal{B}$ vengono scelte in modo tale che la densità della prima nella fase liquida sia un po' inferiore alla densità della sostanza $\mathcal{B}$ ancora fredda. Per questo porzioni di sostanza $\mathcal{A}$ salgono verso l'alto, una dopo l'altra^a. Durante il cammino esse si raffreddano nel liquido $\mathcal{B}$ e raggiungendo la superficie solidificano, prendendo le forme più bizzarre. Raffreddandosi, la densità della sostanza $\mathcal{A}$ diventa leggermente maggiore di quella del liquido $\mathcal{B}$ e "i frammenti" iniziano lentamente a scendere. Tuttavia alcuni di loro rimangono a lungo sospesi sulla superficie. La causa della fluttuazione dei piccoli frammenti sulla superficie può essere la tensione superficiale. Infatti il liquido $\mathcal{B}$ non "bagna" la sostanza $\mathcal{A}$ e pertanto la forza di tensione superficiale agente sui frammenti semisommersi è diretta verso l'alto e tende ad espellerli dal liquido. Grazie a un effetto di questo genere gli idrometri si

^aIn questo senso il dispositivo può essere considerato come uno sviluppo del noto esperimento di **Darling**, in cui una goccia di anilina si riscalda in un grande bicchiere d'acqua. All'incirca a $70^\circ C$ la densità dell'anilina diventa inferiore a quella dell'acqua e la goccia sale verso l'alto.

mantengono sulla superficie dell'acqua e un ago d'acciaio cosparso di grasso galleggia. Inoltre la pressione in eccesso nella parte inferiore del recipiente, sotto la crosta, è già diminuita, i margini della crepa si sono smussati e attraverso questo cratere continuano ad uscire a bassa velocità le successive porzioni di sostanza $\mathcal{A}$ fluidificata. Tuttavia, ora esse non si staccano dal fondo, ma lentamente si distendono dal cratere sotto forma di getto allungato verso l'alto. La superficie di questo getto, venendo a contatto con il liquido freddo $\mathcal{B}$, solidifica formando una specie di tubo. Guardando questo tubo forse vi meraviglierete: esso è sottile ed è riempito di... liquido $\mathcal{B}$. Il fatto è che, quando il getto di sostanza fluidificata $\mathcal{A}$ esce dal cratere e tende verso l'alto, ad un certo momento non dispone più di sostanza per un'ulteriore crescita.

All'interno del getto si crea così una specie di rarefazione e in qualche parte sul limite del tubo compare una frattura nella quale penetra il liquido freddo $\mathcal{B}$. La parte superiore del getto continua il suo moto verso l'alto. Così il liquido $\mathcal{B}$ riempie il tubo all'interno, raffreddando e formando le sue pareti, dopo di che esse solidificano. Nella parte inferiore continua intanto il processo di fluidificazione e un'ulteriore massa di sostanza $\mathcal{A}$ liquida esce dal cratere. Essa sale verso l'alto all'interno del tubicino formatosi e arrivata all'estremità superiore, grazie alla sua massa ancora riscaldata, lo allunga. Con ogni nuova porzione di sostanza $\mathcal{A}$ il tubicino si allunga, formando una figura ondulata che cresce verso l'alto (si veda la Fig. 3 sulla terza pagina della copertina). Accanto ad essa, allungando i frammenti prodotti dalla attività "vulcanica", dopo un certo tempo possono ancora crescere uno o più steli. Tali steli si intrecciano bizzarramente alla maniera degli steli di piante esotiche, tra zolle che costellano il fondo e frammenti che continuano a scendere a seguito del riscaldamento del liquido $\mathcal{B}$. Il quadro si blocca per un certo tempo. Questa fase può essere chiamata "*la fase del bosco di pietra*". Se in questo momento si spegne la lampada, il "bosco pietrificato" rimane invariato: il sistema non potrà tornare allo stato iniziale. Tuttavia, nonostante la varietà degli eventi avvenuti, non siamo ancora giunti al regime di utilizzo. Quindi lasciamo la lampada accesa e continuiamo la nostra osservazione. Col tempo il liquido $\mathcal{B}$ si riscalda, i frammenti che si trovano sul fondo iniziano a fluidificare mentre i tubi rivolti verso l'alto si depositano gradualmente verso il fondo. Tuttavia, tra quelli che erano i frammenti non si vedono gocce oblunghe: gradatamente esse assumono tutte forma sferica. In condizioni normali lo schiacciamento delle bolle su una superficie avviene in virtù della forza di gravità. La gravità si oppone

alle forze di tensione superficiale, che cercano di dare alla bolla una forma sferica, per la quale a volume fissato la superficie è minima. Nella lampada, oltre alla forza di gravità e di tensione superficiale, agisce sulla goccia la forza di Archimede, che compensa quasi totalmente la forza di gravità. Per questo la goccia risulta quasi in stato di imponderabilità e pertanto essa può prendere una forma sferica. Per una goccia la forma sferica è quella più naturale dal punto di vista dell'equilibrio energetico. Per due o più gocce vicine e che si toccano l'un l'altra sarebbe più "utile" fondersi insieme: la superficie di un'unica sfera, di dimensioni maggiori, è più piccola della superficie totale di più sfere con la stessa massa complessiva e di conseguenza l'energia di superficie di un'unica bolla è minore di quella delle bolle componenti. Osservando la lampada noterete che coesistono ancora alcune gocce quasi sferiche della sostanza $\mathcal{A}$ che sembrano non volere fondersi in una sola. Probabilmente più di una volta avete osservato come gocce di mercurio e di acqua si miscelano quasi istantaneamente su una superficie che non si bagna. Sorgerà spontanea la domanda: da cosa dipende il tempo di miscelazione di due gocce?

Da tempo i fisici hanno studiato questo problema. Non si tratta di un quesito peregrino. Al contrario, esso ha un significato pratico di grande importanza per la comprensione dei processi fisici che hanno luogo, per esempio, nella metallurgia delle polveri. In tal caso i grani metallici pressati a seguito della lavorazione termica sinterizzano in sostanze che possiedono proprietà particolari. Nel 1944, il fisico russo **Ja. I. Frenkel'** propose un modello di questo processo, sul quale si basava il suo lavoro d'avanguardia che poneva i fondamenti della metallurgia delle polveri. L'idea alla base di quel lavoro ci permetterà di valutare il tempo di miscelazione.

Poniamo che due gocce di liquido simili vengano a contatto. Nel punto di contatto si forma un istmo (Fig. 12.2), che cresce gradualmente all'aumentare della fusione delle gocce. Per valutare il tempo di fusione, risulta conveniente fare ricorso a valutazioni di bilancio energetico. All' "attivo" di un sistema di due bolle dobbiamo porre un'energia ΔE_s uguale alla differenza delle energie di superficie dello stato iniziale e finale (ossia di due singole gocce con raggi r_0 e di un'unica goccia "comune" di raggio r):

$$\Delta E_s = 8\pi\sigma r_0^2 - 4\pi\sigma r^2.$$

Poiché con la fusione il volume totale non cambia, avremo $\frac{4\pi}{3}r^3 = 2 \cdot \frac{4\pi}{3}r_0^3$, da cui $r = r_0\sqrt[3]{2}$. Quindi

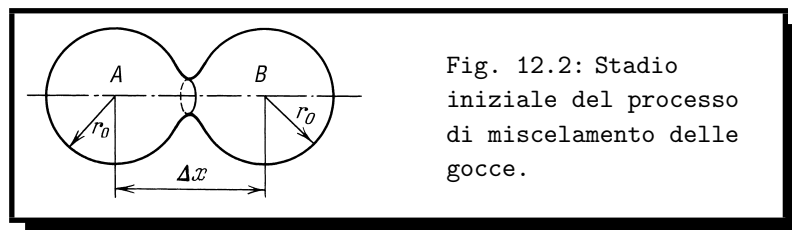


Fig. 12.2: Stadio iniziale del processo di miscelamento delle gocce.

$$\Delta E_s = 4\pi\sigma(2 - 2^{\frac{2}{3}})r_0^2. \quad (12.1)$$

In accordo all'idea di Frenkel' questa energia deve essere dissipata come lavoro contro le forze d'attrito viscoso che sorgono nel processo di mescolamento delle gocce. Calcoleremo l'ordine di grandezza di questo lavoro. Per la forza d'attrito viscoso utilizziamo l'espressione di **Stokes**, valida per il caso di una sfera di raggio R che si muove con velocità $\vec{v}$ nel liquido di viscosità η : $\vec{F} = -6\pi\eta R\vec{v}$ ^b.

Assumeremo che la viscosità della sostanza di cui sono costituite le gocce sia assai maggiore della viscosità dell'ambiente circostante e per questo nella formula di Stokes inseriremo η_A . In luogo di R porremo r_0 . Questa grandezza caratterizza anche le proporzioni del mescolamento della massa di liquido: $\Delta x \sim r_0$. Quindi, per il lavoro delle forze d'attrito viscoso avremo

$$\Delta A \sim 6\pi\eta_A r_0^2 v.$$

È intuitivo che quanto più rapidamente le gocce si miscelano, tanto maggiore è l'energia necessaria per raggiungere lo stato finale (a causa dell'aumento delle forze d'attrito viscoso). Tuttavia la riserva d'energia è limitata, pari a ΔE_s (12.1). Da questa considerazione viene determinato il tempo di fusione delle gocce τ_F (il cosiddetto *tempo di fusione di Frenkel'*).

^bLa formula di Stokes è stata dedotta per il moto di una sfera in un liquido viscoso. Tuttavia anche nel caso in esame la forza dell'attrito viscoso può comunque dipendere soltanto dalla viscosità, dalla dimensione della goccia e dalla velocità di svolgimento del processo. Per questo, in accordo a considerazioni di dimensionalità, con queste grandezze non riusciamo ad ottenere altra espressione che abbia le dimensioni di una forza che non sia quella di Stokes (una differenza può sussistere nel coefficiente di proporzionalità; tuttavia noi ricorriamo alla formula di Stokes soltanto per valutazioni di ordini di grandezza).

Stimando la velocità del processo come data da $v \sim r_0 / \tau_F$ otteniamo

$$\Delta A \sim \frac{6\pi \eta_A r_0^3}{\tau_F} \sim 4\pi \sigma (2 - 2^{\frac{2}{3}}) r_0^2,$$

da cui

$$\tau_F \sim \frac{r_0 \eta_A}{\sigma}. \quad (12.2)$$

Per le gocce d'acqua con $r_0 \sim 1 \text{ cm}$, $\sigma \sim 0.1 \text{ N/m}$ e $\eta \sim 10^{-3} \text{ kg/(m} \cdot \text{s)}$ questo tempo è all'incirca $\sim 10^{-4} \text{ s}$. Tuttavia per la glicerina, notevolmente più viscosa (a circa 20° C $\sigma_{gl} \sim 0.01 \text{ N/m}$ e $\eta_{gl} \sim 1 \text{ kg/(m} \cdot \text{s)}$) il tempo corrispondente è vicino al secondo. Per liquidi diversi, al variare della viscosità e della tensione superficiale, τ_F può cambiare in intervallo molto ampio di valori.

È importante osservare che, grazie alla forte dipendenza della densità dalla temperatura, il tempo τ_F può variare sensibilmente anche per uno stesso liquido. Così, la viscosità della glicerina variando la temperatura da 20° a 30° C diminuisce di 2,5 volte. La tensione superficiale dipende dalla temperatura molto più debolmente (nell'intervallo di temperature indicato σ_{gl} diminuisce soltanto di alcuni percento). Per questo si può considerare la dipendenza del tempo di fusione di Frenkel' dalla temperatura come determinata essenzialmente dalla dipendenza della viscosità dalla temperatura.

Torniamo ora alle sfere che giacciono sul fondo della lampada. Fino a quando la temperatura del liquido $\mathcal{B}$ non è salita sensibilmente, la viscosità della sostanza amorfa $\mathcal{A}$ è ancora alta. Proprio per questa ragione le sfere non si mescolano. Allo stesso modo due sferette di cera non si mescolano se si portano in contatto a temperatura ambiente e si comprimono. Tuttavia, quando si riscaldano la viscosità della cera diminuisce marcatamente e le sfere si mescolano abbastanza rapidamente. Sottolineiamo anche l'importante ruolo della superficie delle sfere: se una non è liscia ma piuttosto rugosa sarà più difficile ottenere “un ponte” tra le sfere.

La miscelazione delle gocce è necessaria perché la lampada continui a funzionare ed è previsto uno speciale meccanismo di “travaso” della sostanza $\mathcal{A}$ dalle singole gocce alla massa principale già fluidificatasi. Si tratta della spirale che corre lungo il perimetro del fondo della lampada. Essa è ad alta temperatura e quando le gocce della sostanza $\mathcal{A}$ entrano in contatto con essa si riscaldano, la viscosità diminuisce e le gocce si mescolano con facilità con la massa principale.

Così, sul fondo del recipiente si è formata un'unica massa liquida di sostanza $\mathcal{A}$. A causa del riscaldamento, che continua, essa non può rimanere in quiete. Inizia la “fase delle protuberanze”.

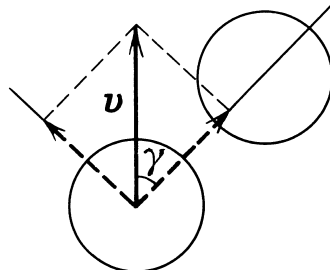
Una protuberanza che si stacca dalla superficie sotto l'azione della forza di espulsione sale lentamente verso l'alto (si veda la figura 5 sulla terza pagina della copertina), assumendo gradualmente la forma di una sfera. Dopo essere salita nella parte superiore della lampada, dove il liquido $\mathcal{B}$ (a causa della bassa conducibilità termica) fino a questo momento non si è riscaldato, questa sfera si raffredda un po' (rimanendo comunque liquida) e lentamente scende verso il basso. Tuttavia, come abbiamo già illustrato, non è agevole per la sfera fondersi con la superficie e così essa ballonzola abbastanza a lungo sulla stessa, spostandosi gradatamente verso la periferia. Qui la spirale “scova” la sua superficie e quella che era la protuberanza completa il suo viaggio, ritornando all'ambiente che la ha generata.

La lampadina alla base del cilindro continua a riscaldare il sistema e il processo di generazione di protuberanze reitera. Staccandosi dalla superficie, le protuberanze lasciano sospese “tra cielo e terra” singole gocce che non riescono a decidere se gettarsi all'inseguimento della protuberanza o ritornare nell'ambiente naturale. A questo punto nel cilindro si trovano già una decina di sfere liquide, alcune delle quali salgono, mentre altre scendono. Prende così avvio la “fase delle collisioni e delle catastrofi”. Proprio questa fase, più lunga e spettacolare, viene considerata dai costruttori la più proficua ai fini estetici della lampada magica.

Le sfere si scontrano, cambiano la direzione del proprio moto e non riusciamo facilmente a seguire la loro fusione nel corso delle collisioni. Come abbiamo già spiegato, da un punto di vista energetico è più conveniente per le sfere miscelarsi. Per far ciò è necessario del tempo. Il tempo che “hanno perduto” è il tempo di collisione t . Se τ_F supera notevolmente t , le sfere non riescono a fondersi e si disperdono. Da che cosa è determinato il tempo di collisione? Nella lampada avvengono soprattutto urti obliqui (Fig. 12.3), nel corso dei quali le sfere si deformano facilmente e scivolano una sull'altra. In questo caso il tempo caratteristico di collisione è $t \sim r_0/v$. La velocità delle sfere v è solo di alcuni centimetri al secondo, i raggi sono di alcuni centimetri. Pertanto $t \sim 1\text{ s.}$ e in questo tempo le sfere non riescono a miscelarsi. Ecco che esse devono vagabondare per la lampada, rimanendo per un certo tempo al fondo o quasi pendendo dall'alto, urtandosi ma senza potersi fondere.

La “fase delle collisioni e delle catastrofi” dura a lungo: 5 - 7 ore. Dopo

Fig. 12.3: La maggior parte delle collisioni di gocce della sostanza $\mathcal{A}$ avvengono in forma di urti obliqui poco violenti.



questo tempo, le istruzioni raccomandano di spegnere la lampada. Tuttavia se la temperatura dell'aria circostante è abbastanza alta, questa fase può non essere l'ultima. Dopo che nella lampada si stabilisce una distribuzione stazionaria della temperatura (cioè tutto il liquido $\mathcal{B}$ si riscalda), le densità delle sostanze $\mathcal{A}$ e $\mathcal{B}$ praticamente si eguagliano. Tutta la sostanza $\mathcal{A}$ si raccoglie in una grossa bolla, che plana al fondo. Con il tempo questa sfera, a causa del contatto con le pareti del cilindro si raffredda un poco, la sua densità aumenta lievemente ed essa lentamente scende sul fondo. Toccato il fondo, la sfera riceve una quantità supplementare di calore e ritorna al posto precedente. Qui essa si ferma fino a che non si raffredda di nuovo, dopo di che il processo si ripete. Questa fase non prevista dalle istruzioni, può essere chiamata “fase della grande sfera” (si veda la figura sulla terza pagina della copertina).

Ora, dopo aver osservato e compreso molti aspetti del funzionamento della lampada, consideriamo questo fenomeno nel suo insieme. Si pone la domanda: perché compaiono questi processi di formazione, collisione e scomparsa delle sfere che si alternano l'un l'altro e si ripetono? È manifesto che tutta la “sorgente” del processo consiste nella differenza di temperatura tra l'estremità superiore e quella inferiore della lampada (“riscaldatore” e “raffreddatore”). Se supponiamo che il flusso di calore si propaghi in virtù della conducibilità termica del liquido $\mathcal{B}$, la sua temperatura cambierà lievemente con l'altezza e nel sistema non accadrà nulla di insolito. La comparsa di sfere, così come anche la conduzione convettiva, è la conseguenza delle instabilità che compaiono in condizioni particolari nei sistemi in cui la differenza di temperatura provoca una propagazione di flussi di calore sui contorni. Dello studio delle leggi generali che riguardano il comportamento di questi sistemi si interessa una nuova scienza, che si sta sviluppando rap-

idamente: *la sinergetica.*

Capitolo 13

Nunc est bibendum.

Dopo aver osservato e studiato la lampada magica, un Sabato sera beneficieremo di un momento conviviale, accompagnando la cena con una bottiglia di vino. Anche in questo caso, la fisica mostra la sua onnipresenza. Cominciamo con un accenno al ruolo di questa disciplina nelle tecniche di vinificazione.

13.1 La fisica nella vinificazione

Le origini del vino si perdono nei tempi nei quali sono nate forme di vita associativa. Resti di vinaccioli sono stati rinvenuti in caverne preistoriche. Da millenni avanti Cristo si conoscevano i frutti della vite, pianta presumibilmente originaria dall'India. Dopo la scoperta della possibilità di ottenere la bevanda, al vino venne progressivamente attribuito un ruolo quasi sacrale. Esso divenne parte essenziale nel Sacramento dell'Eucarestia nelle pratiche religiose del Cristianesimo e anche nel rito ebreo del Kiddush la benedizione cerimoniale si attua sul vino e sul pane.

Storicamente è abbastanza certo che in Mesopotamia e in Egitto siano state sviluppate per la prima volta tecniche di vinificazione. Le migrazioni degli ariani dall'India, verso il 2500 A.C., determinarono il trasferimento in occidente della pratica della vinificazione. Certa è, intorno all'anno 2000 A.C., la produzione del vino in Sicilia (forse in relazione alla colonizzazione dei Greci o degli Egiziani), più tardi presso i Sabini e gli Etruschi. Per tornare ai Romani, Orazio e Virgilio (negli anni dal 43 al 30 A.C.) scrissero lodi di varia natura al vino. È del 42 A.C., ad opera di Columella, la stesura di un ottimo manuale di viticoltura, il "De re rustica".

Verso l'anno 200 D.C. inizia la crisi dell'agricoltura, crisi che si aggrava con le invasioni barbariche e determina il progressivo abbandono delle pratiche vinicole. Con l'avvento dei Maomettani (per i quali l'alcool è proibito) si produce il totale abbandono della vite. Nel Medioevo l'enologia viene sporadicamente praticata in castelli e nei conventi. Bisogna arrivare al Rinascimento (dal XVI secolo in poi) per vedere rinnovarsi la coltivazione della vite e assistere a una grande espansione nella produzione del vino.

L'invenzione delle botti per la conservazione del vino e il suo trasporto si fa risalire ai Galli, mentre presso i Greci e i Romani venivano utilizzate le anfore. La utilizzazione delle botti condusse rapidamente alla percezione che il legno poteva trasferire aromi e profumi, arricchendo in qualità il contenuto. Come è noto, oggi l'invecchiamento del vino in botti è pratica comune e parte essenziale nella produzione di un vino di qualità.

La nascita della moderna enologia si può far risalire all'inizio del secolo XX, paradossalmente a seguito di eventi funesti: l'attacco di parassiti, legato a spostamenti di massa. Compagno l'oidio (*Uncinucula necator*), la peronospora (*Plasmopora viticola*) e la fillossera (*Phylloxera vastatrix*) della famiglia degli Afidi. I primi due sono combattuti con zolfo e rame. La fillossera richiede un brillante esempio di controllo genetico. La vite americana ha radici che sono progressivamente mutate sino a divenire resistenti alla fillossera in ragione di uno strato di sughero alla superficie radicale. La vite europea, di migliori qualità organolettiche (la *Vitis vinifera* di provenienza orientale) viene così innestata su ibridi delle viti americane (le *Vitis riparia*, *Vitis rupestris* e *Vitis berlandieri*). L'ibrido di radici di viti americane diviene così il porta-innesto delle gemme delle viti europee. Infatti la *Vitis vinifera* non-innestata e sana esiste solamente in pochissime zone d'Europa: nello Jerez, nel Colares (Lisbona) ove la sabbia impedisce lo sviluppo degli Afidi, nelle regioni della Mosella e del Douro, ove analoga protezione esercitano l'ardesia e lo scisto. Ricordiamo che l'ibrido non è un innesto. Gli ibridi di vite americana e di viti europee non possano essere coltivati in Europa (la presenza di prodotti da ibridi viene rilevata con la fotometria di assorbimento atomico). Gli incroci tra varietà europee non sono considerati ibridi.

Le moderne tecniche di vinificazione beneficiano largamente di effetti o di strumenti genuinamente fisici. Non dimentichiamo che lo stato finale del succo dell'uva doveva essere l'aceto. A bloccare questa evoluzione nefasta è l'intervento dell'uomo e il ricorso alla scienza. Per questo è preferibile, in

genere, confidare in una vinificazione sostenuta dagli ausili della tecnologia piuttosto che operare con il “fai da te” artigianale.

La buccia dell’acino contiene in sé i lieviti che determinano la fermentazione che conduce dal mosto all’alcool (*Saccharomyces ellipsoideus*). Per evitare che fermentazione si protragga sino all’aceto occorre separare il mosto dai componenti che risulterebbero alla fine dannosi o almeno ridurne la funzione. E’ questa essenzialmente la ragione della filtrazione del mosto, dopo la quale il ruolo delle particelle di lievito è notevolmente ridotto. La successiva fermentazione avviene oggi, in genere, in contenitori di acciaio. La fermentazione è un processo esotermico e la temperatura del mosto può salire a $40 - 42^{\circ}\text{C}$. In una produzione “artigianale” tali temperature causano la evaporazione di prodotti volatili, di nobili aromi di frutti e di fiori che più tardi potrebbero influire positivamente sulla qualità del vino, specie dei bianchi. Allo scopo di conservare questi tesori della natura i moderni produttori realizzano la fermentazione a freddo (circa 18°C), processo che richiede un tempo assai più lungo (tipicamente tre settimane invece di 7-8 giorni della fermentazione “naturale”). Anche la filtrazione può essere condotta a freddo, per esempio alla temperatura di -5°C , per liberarsi di infiltrati organici dannosi o sgradevoli (che solidificando precipitano mentre il vino, grazie al contenuto alcolico che ne abbassa la temperatura di solidificazione di circa 0.5°C ogni percento di alcool, può fluire). Tecniche di raffreddamento intervengono anche per incrementare il contenuto di zuccheri. Se necessario si raffreddano i grappoli sino al congelamento per ottenere la solidificazione dell’acqua e procedere quindi alla spremitura. Il mosto ottenuto in tal modo contiene più zuccheri.

Queste poche considerazioni danno una idea del perché una moderna “cantina” per la produzione del vino assomiglia piuttosto a un laboratorio di ricerca.

Bisogna peraltro aggiungere che in alcuni casi i parassiti, in genere combattuti, svolgono un ruolo particolare e anche “gradito”. E’ il caso della “*Botrytis cinerea*” o “*muffa nobile*”, frequente nelle vigne dell’area del Sauternes, in conseguenza anche del clima piuttosto umido. Forando l’acino, il parassita provoca come conseguenza una riduzione di acqua (fino 20%) e un arricchimento in zuccheri. Nel contempo si produce nel prodotto il caratteristico sapore del Sauternes da molti apprezzato, particolarmente in combinazione con il Roquefort e il “paté de foie gras”.

La leggenda vuole che il vino Sauternes sia stato per così dire scoperto in un anno di prematura primavera e di caldo estivo: la maturazione avvenne

prima del previsto, ma a seguito di improvvise cospicue piogge la raccolta venne rinviata. Con preoccupazione i viticoltori videro i grappoli coprirsi di una muffa inattesa. Decisero di tentare ugualmente la produzione del vino con la raccolta di quella “strana” uva. Il risultato finale superò ogni aspettativa: il vino risultava avere acquisito un gusto straordinario. Come si vede, molte scoperte (e non solo in Fisica), beneficiano di eventi inaspettati e non prevedibili né voluti o se preferite di un aiuto da parte della fortuna. A proposito del Sauternes possiamo aggiungere che la fama del “Chateau Ikem”, una delle più prestigiose aree di produzione, è dovuta al Principe Costantino di Russia. Infatti egli fu il primo fanatico estimatore del vino proveniente da quella area e accettò di versare una somma spropositata (20.000 franchi) per una botte, con il conseguente risultato di propagandare il produttore. In qualche modo è rimasto anche oggi il beneficio propagandistico del Principe Costantino, se è vero come è vero che il costo di una bottiglia di “Chateau Ikem” supera spesso 1000 Euro.

Molto si potrebbe aggiungere al ruolo della Fisica e della Chimica nei processi per la migliore vinificazione ma preferiamo passare alla descrizione di qualche tipico fenomeno fisico che coinvolge direttamente il vino o alcuni suoi “derivati”.

13.2 Le “lacrime” del vino

Come sanno gli estimatori (o che si atteggiavano tali e parlano a proposito di “corpo del vino” o di “glicerina”) sulla superficie interna di un bicchiere riempito a metà o meno di vino, è possibile produrre, con una leggera rotazione, alcuni rivoletti viscosi che lentamente scendono aderendo alla parete del vetro. Sono noti come “lacrime del vino” (Fig. 13.1). Come si producono? Quali proprietà le controllano? Quale elemento di giudizio sul prodotto possiamo trarne?

In idrodinamica il fenomeno è conosciuto come effetto **Marangoni**, si può osservare facilmente in ogni miscela di alcool e acqua (con il primo almeno al 20%) ed è legato essenzialmente al gradiente di tensione superficiale. A seguito della più rapida evaporazione di alcool e della più elevata tensione superficiale dell’acqua, si può produrre un gradiente di concentrazione che induce a sua volta un gradiente di tensione superficiale. Quest’ultimo produce una sollecitazione superficiale che può far salire un sottile strato lungo la parete di vetro del bicchiere, contrastando in tal modo



Fig. 13.1: Le lacrime del vino sulle pareti del bicchiere.

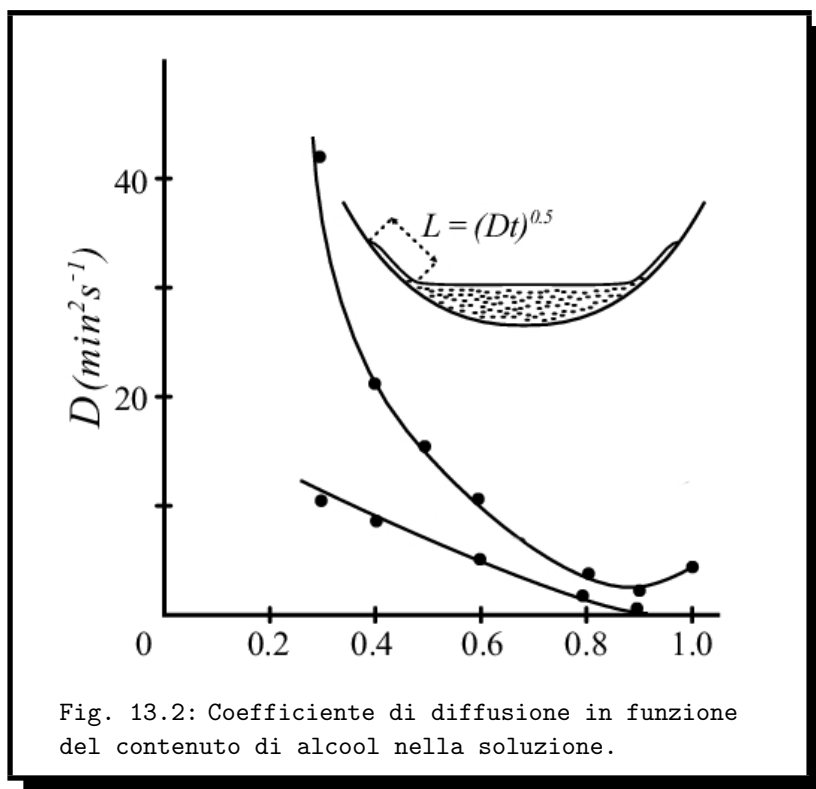
l'effetto della gravità.

Per meglio spiegare il fenomeno, due ricercatori francesi (**Fournier** e **Cazabat**) hanno recentemente studiato la cosiddetta “dinamica di bagnamento”, in termini della legge di diffusione alla Einstein del fronte dello strato $L(t)$ nella forma

$$L(t) = \sqrt{D(\varphi)t}$$

con D coefficiente di diffusione, noto anche come coefficiente di **Einstein-Smolukhovskij**. Tale coefficiente nelle formule originali per la diffusione ricavate da questi scienziati dipende dalla velocità della particelle diffondenti (per esempio piccolissime sferette di materiali in un liquido) e dal loro libero cammino medio. La scoperta dei ricercatori francesi è stata non solo che tale legge è in effetti verificata nei riguardi dello strato di liquido che bagna una superficie ma anche che il coefficiente di diffusione $D(\varphi)$ è una funzione fortemente dipendente dal contenuto alcolico φ (Fig. 13.2).

Mentre è stato confermato da questi studi che il contenuto minimo di alcool per osservare il moto spontaneo di risalita dello strato di liquido deve essere all'incirca del venti per cento ($\varphi = 0,2$), sorprendentemente e'



risultato che i valori più elevati di D si ottengono per basse concentrazioni alcoliche, benché in tali casi l'evaporazione risulti ovviamente più lenta.

La formazione delle lacrime è associata a due effetti. Una volta formati lo strato di liquido che bagna la superficie, la gravità tende ovviamente a farlo “cadere” verso il fondo del bicchiere. Dall'altra il fenomeno cosiddetto del “fingering” o formazione di grosse gocce terminali da associare alla instabilità della superficie (quale quella che provoca il caratteristico sventolio di una bandiera al vento) tipica di sistemi instabili anche rispetto a piccole deviazioni dall'equilibrio. Maggiore il contenuto di alcool, maggiore il numero di rivoletti. Tuttavia tale dipendenza non è risultata molto pronunciata: la percentuale alcolica necessaria per osservare visivamente una marcata differenza è risultata così alta che in pratica potrebbe essere un indicatore di elevato contenuto di alcool solo se si confrontasse un vino ordinario con un vino artificialmente “fortificato” da una marcata aggiunta

di alcool stesso.

Per quanto riguarda invece la glicerina (C_3H_8OH) essa contribuisce al gusto di dolce ma non al cosiddetto “corpo” del vino (ricchezza in alcool e in aromi) né, in larga misura, alla sua viscosità.

Pertanto la conclusione degli studi ai quali abbiamo fatto cenno è che dallo spettacolare fenomeno delle “lacrime” difficilmente si può trarre qualche serio elemento di valutazione sul “corpo del vino” o sulla concentrazione di glicerina.

13.3 Spumante, bolle e “fizz”

Nella regione francese nota come Champagne, nei pressi di Reims, si produceva da tempo vino bianco, di qualità piuttosto povera. Fu osservato che se immediatamente imbottigliato, anziché conservato in botti, la fermentazione non aveva termine. A seguito delle alte pressioni associate alla produzione di gas nel corso della fermentazione alcune bottiglie esplodevano (una parziale protezione era allora ottenuta con cappucci metallici e più tardi mediante bottiglie più resistenti, a pareti più spesse). Il monaco Benedettino Dom Perignon, un ragguardevole e stimato viticoltore, è considerato in particolare il creatore dello Champagne e il fondatore di tutte le tecniche per la produzione di spumanti. A soli 28 anni egli era già il responsabile delle cantine della Abbazia di Hautvillers e aveva sviluppato particolari tecniche di conservazione e fermentazione del vino in bottiglia. In realtà egli potrebbe aver appreso parte delle tecniche in una altra proprietà dei Benedettini, la Abbazia Saint Hilare a Limoux, nel sud della Francia.

Il caratteristico aspetto degli spumanti, è dovuto a una elevata quantità di anidride carbonica. Essa può essere ottenuta prolungando la fermentazione, aggiungendo zuccheri e lieviti prima a livello di mosto e successivamente nel processo di invecchiamento in bottiglia. Le bottiglie sono inclinate e frequentemente ruotate e i residui di lieviti e altri prodotti si lasciano precipitare nel collo delle stesse. La fisica interviene a questo punto: si provoca il congelamento del tratto terminale della bottiglia (in genere con azoto liquido, che a pressione ordinaria bolle a una temperatura di 77 K, pari a circa $-196^\circ C$, e tale temperatura mantiene sino alla totale evaporazione) e la rimozione della porzione non gradita.

Interessiamoci di qualche divertente fenomeno fisico che possiamo far

avvenire con lo spumante. Per esempio il copioso “spruzzo” di vino, gas e schiuma che accompagna talvolta la celebrazione di un evento di successo. Il marcato scuotimento, prima della rimozione del tappo, determina l’ingresso forzato nella soluzione, già satura di anidride carbonica, anche del gas esistente nella regione alta della bottiglia. All’apertura si ha una brusca caduta della pressione, la solubilità del gas nel liquido si riduce notevolmente e bruscamente. Pertanto vengono emesse le “bolle” accompagnate da schiuma e in parte da spumante stesso.

Gli studenti di Fisica talvolta si divertono a fare impressione sulle fanciulle della area letteraria mostrando come siano in grado di far magicamente oscillare dal basso verso l’alto e poi ancora in basso un pezzo di cioccolato in un bicchiere di spumante. Ricordando la spinta di Archimede e la adesione di bolle sulla superficie, bolle che evaporano quando il cioccolato giunge alla superficie libera del liquido, non è difficile rendersi conto del fenomeno.

Vogliamo piuttosto analizzare la fisica del caratteristico “perlage” e il particolare “frizzare” dello spumante. Il piacevole suono delle bollicine che accompagnano lo spruzzo e il mescolare dello spumante è il risultato di un processo “a valanga”. Infatti, tale “suono” è in realtà la somma di molti lievi scoppiettii di bolle individuali, come è stato recentemente dettagliato da ricercatori della Università di Liegi.

Se l’esplosione delle bolle avvenisse a tasso costante si sentirebbe un fruscio uniforme, come il disturbo sottostante nelle trasmissioni radio. E’ tipico di diversi processi acustici questa specie di “rumore bianco”, nel quale la frequenza del suono è distribuita a caso su un ampio intervallo, mentre l’ampiezza è praticamente costante. Invece, nel caso dello spumante il “fizz” è ben diverso da questo “rumore bianco” e attraverso opportuni microfoni si è potuto provare che il “suono” ha dei picchi caratteristici: infatti le bollicine non scoppiano indipendentemente ma a seguito di processi cooperativi. Ciascuna piccola esplosione individuale dura all’incirca 1/1000 di secondo mentre una numerosa serie di esse si producono in rapida successione, combinandosi a generare il segnale acustico. Il tempo tra esplosioni successive è variabile ed esse non hanno una durata preferita, comportandosi nella loro sequenza come in un processo a valanga, simile a quello che si osserva nei terremoti, nelle esplosioni solari e negli smottamenti di terra. In breve, il comportamento delle bolle durante la generazione del “fizz” è quello caratteristico di sistemi a molti componenti nei quali l’effetto globale è dovuto a una grande varietà di interazioni individuali.

Al termine di questa sezione su vini e spumanti, non è inopportuno

richiamare al lettore una terminologia poco conosciuta, benché parte di essa sia entrata nell'uso comune. Molti dei lettori sanno cosa significhi "Magnum" se non altro per aver visto come vengano celebrate vittorie di "Formula 1" con il copioso spruzzo di spumante che segue alla apertura rapida di una grossa bottiglia, dopo averla scossa. Esiste solo il Magnum nella terminologia "ortodossa" delle bottiglie di diverse dimensioni? No e nella Tabella sottostante vi elenchiamo i nomi e le proprietà di inusuali bottiglie, nella terminologia usata dai grandi specialisti di vini e di spumanti. Tale terminologia si basa su nomi biblici e non si conosce con sicurezza l'origine delle attribuzioni che riportiamo in Tabella :

Tipo	corrispondente a litri	cioé bottiglie della dimensione canonica Europea
Magnum	1.5	2
Jeroboam (in italiano Geroboamo, primo sovrano del regno di Israele)	3	4
Rehoboam (Reoboamo, figlio di Salomone)	4.5	6
Mathusalem (Matusalemme, patriarca biblico sinonimo di persona che vive molto a lungo)	6	8
Salmanazar (nome corrispondente a cinque re assiri)	9	12
Balthazar (Baldassarre, sovrano babilonese ma anche uno dei Re Magi)	12	16
Nabucodonosor (re di Babilonia)	15	20

Ora sapete che alla prossima vostra celebrazione di un felice evento potrete richiedere si proceda alla apertura di un Salmanazar di spumante per rendere entusiasti gli amici che con voi celebrano l'evento stesso.

13.4 Pastis

Nella regioni meridionali della Francia, dove il clima è spesso piuttosto caldo, è diffuso il pastis, ottenuto nel modo seguente. Prima si versa in un bicchiere una certa quantità di pastis stesso, poi si aggiunge acqua fredda. Benchè entrambi gli ingredienti siano trasparenti la loro miscela diviene opaca, quasi bianca. Perchè? Il pastis consiste in realtà di una soluzione di alcool etilico (45%) e acqua (55%) con una piccola quantità (0.2%) di anetolo, estratto dai semi d'anice (componente aromatico presente talvolta in medicinali o in profumi). L'anetolo, di colore giallognolo e fortemente ricco di aromi, si scioglie assai bene in alcool e poco o nulla in acqua (le molecole di acqua interagiscono fortemente tra di loro ma assai debolmente con quelle di anetolo. Al contrario le molecole di alcool attraggono bene sia le loro vicine dello stesso che quelle dell'anetolo). Nella bottiglia di pastis la concentrazione alcoolica è elevata e l'anetolo viene quindi facilmente disciolto nella soluzione. Quando si aggiunge l'acqua, che non accetta volentieri in soluzione le molecole di anetolo, queste ultime tendono a raggrupparsi tra di loro in "clusters" cioè in piccole gocce. La presenza di questi agglomerati rende la bevanda opaca, che costituisce essenzialmente una emulsione (come lo sono il latte, la mostarda e la maionese). La elevata diffusione di luce è legata alla dimensione degli agglomerati, confrontabile con la lunghezza d'onda della luce visibile (rileggete la sezione sul colore del cielo, al capitolo 3). In realtà, a tempi lunghi (come hanno mostrato studi di diffusione di neutroni al centro scientifico di Grenoble) gli aggregati di anetolo si attirerebbero a formare una grossa goccia, con completa separazione di fase. Tuttavia di solito non si aspetta a lungo dopo la preparazione del pastis nel bicchiere, specie se il clima è caldo...

13.5 La vodka o il "vino di pane"

Salendo verso i paesi nordici diviene difficile la coltivazione della vite e il posto del vino viene spesso occupato da bevande derivate dalla distillazione di mele (Calvados - Normandia, Nord della Francia), prugne e albicocche

(Slivovitsa - Croazia- Serbia), ginepro (Gin - Gran Bretagna) o di diversi tipi di cereali o di grano (Whisky- Gran Bretagna, Vodka - Russia). In tempi antichi, in Russia, la vodka era chiamata proprio così: il "vino di pane" (Smirnov). La vodka viene in genere consumata in Russia anche oggi allo stesso modo di come italiani e francesi fruiscono a pranzo o a cena del derivato dall'uva. Vi siete mai chiesti perché le bevande su elencate, da un punto di vista chimico, essendo essenzialmente diluizioni dell'alcool etilico ($\text{CH}_3 \text{CH}_2 \text{OH}$) in acqua, hanno comunque una gradazione intorno a 40% ?

Come prima ragione citiamo una leggenda, che ha una base fisica. Si racconta che quando Pietro il Grande introdusse il monopolio per la distillazione della vodka, divenendo così lo stato responsabile della sua qualità, gli osti cominciarono a diluire la vodka senza alcuno scrupolo. In questo modo aumentavano i loro guadagni e contemporaneamente scontentavano gli avventori per la scarsa qualità della bevanda preferita. Per stroncare questo malcostume, Pietro il Grande emanò un editto, che autorizzava i clienti a picchiare i gestori delle osterie sino alla morte se il vapore della vodka servita non poteva essere incendiato. Lestamente fu osservato che il contenuto del 40% di alcool risultava essere quello minimo per permettere la combustione sullo specchio del liquido. Su questa percentuale si fermarono gli osti, trovando un buon compromesso tra guadagno, rispetto delle norme e soddisfazione della clientela.

Un secondo e più serio motivo del numero magico è legato alla seguente circostanza, di nuovo basata su un fenomeno fisico, la espansione termica, che si manifesta nell'aumento di volume con l'aumento della temperatura e alla sua riduzione con il raffreddamento. Obbedisce a questa legge la maggior parte dei composti, tra i quali l'alcool. L'acqua, invece, costituisce un liquido anomalo: per temperature inferiori a 4°C , con il calo della temperatura il volume anziché diminuire comincia a crescere, fino al punto di congelamento dove varia bruscamente del 10%! Per questa ragione non si possano lasciare le bottiglie con l'acqua fuori casa a temperature inferiori a 0°C : l'acqua congelata, aumentando il suo volume, le romperebbe. Con la vodka ciò non accade: in Siberia le scatole con la bevanda preziosa vengono lasciate al gelo senza alcuna conseguenza. La spiegazione di questo fenomeno dipende da due fatti. Innanzitutto, la presenza di una così notevole percentuale di alcool contrasta la solidificazione (come si è già menzionato), evitando il salto repentino del volume specifico con il calo di temperatura.

In secondo luogo, al rapporto dei volumi 4 a 6 il coefficiente totale di espansione termica si trova vicino allo zero: “l’anomalia” dell’acqua si compensa con “la normalità” dell’alcool. E’ sufficiente controllare sulle tabelle di un manuale i rispettivi coefficienti di espansione termica : per l’acqua $\alpha_{H_2O} = -0.7 \cdot 10^{-3} \text{ gradi}^{-1}$ e per l’alcool $\alpha_{C_2H_5OH} = 0.4 \cdot 10^{-3} \text{ gradi}^{-1}$.

Esiste anche un terzo motivo per il magico 40%, motivo la cui scoperta si attribuisce al famoso chimico russo **Dimitri Mendeleev**. Il motivo è legato alla stabilità del contenuto alcolico percentuale a seguito della evaporazione, se si permette alle molecole di lasciare la superficie libera del liquido. Pertanto il biccherino di vodka lasciato sulla tavola mantiene la stessa gradazione alcolica anche per molte ore.

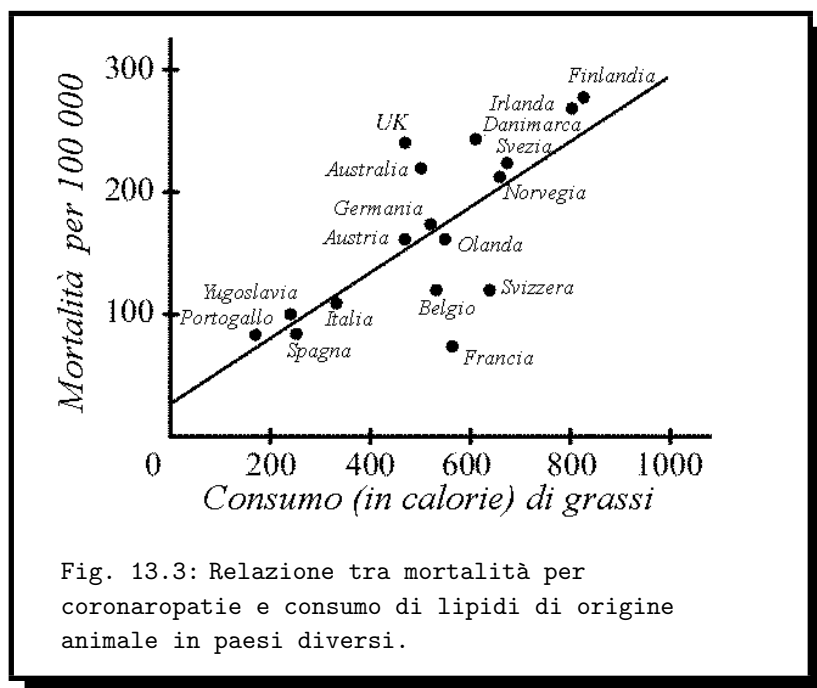
Un più importante contributo di Mendeleev è relativo un punto di minimo nell’eutettico acqua-alcool. A minore concentrazione di alcool si potrebbe produrre una separazione di fase a temperature relativamente basse. Poiché la separazione di fase è favorita dalla eventuale presenza di impurezze, questa può essere una ragione per la quale i superalcolici non esenti da impurezze varie (tipo le grappe prodotte artigianalmente) sono costretti a salire nella concentrazione di alcool per evitare la smiscelazione dei due componenti.

Alcuni chimici sostengono che un altro motivo a favore della concentrazione del 40% sia legato al salto di viscosità della soluzione dell’alcool vicino a questa concentrazione. In una soluzione debole, le molecole dell’alcool sono indipendenti e si muovono ognuna per se stessa. Quando la concentrazione raggiunge il valore del 40% ha luogo il fenomeno di polimerizzazione, le molecole dell’alcool si ordinano in lunghe catene, la viscosità ha una brusca variazione e la vodka migliora le sue proprietà organolettiche.

13.6 Vino e patologie cardio-circolatorie.

Nella figura 13.3 è riportato l’andamento del numero di morti per anno (ogni 100 000 persone) per coronaropatie in funzione del consumo (calorico, giornaliero) di lipidi di origine animale. Come si vede appare sussistere una precisa correlazione: a elevati consumi di grassi corrispondono coefficienti di mortalità elevati e viceversa. Compare nel diagramma un punto anomalo: una bassa mortalità per malattie cardiache nonostante una relativamente elevata assunzione di grassi. Si osservi, infatti, che nonostante un consumo più elevato della Gran Bretagna, nell’area francese presa in considerazione

si registra una mortalità assai minore, ridotta a quasi un quarto.



Questi dati ufficiali (risultanti dalle ricerche nell'ambito del progetto MONICA dell'Organizzazione mondiale della sanità e pubblicati in *"The Lancet"* nel 1992 da **Renaud e De Longeril**^a) dovevano già essere noti nel 1991 ad un conduttore di uno speciale televisivo dalla CBS, che per la prima volta parlò di "paradosso francese", divenuto anche noto nel seguito come "effetto Bordeaux". Si sospettò subito che il beneficio fosse legato all'assunzione giornaliera di vino rosso, comune nell'area di Bordeaux. Indagini scientifiche sistematiche, condotte in diverse aree con tipologie vinicole simili, hanno permesso di confermare in maniera incon-

^aLa correlazione riportata in Fig. 13.3 e la anomalia per il dato nell'area di Bordeaux é già stata prese in considerazione da Siliprandi, Venerando e Miotto, nel lavoro "Quando il vino è rosso" (Atti dell'Istituto Veneto di Scienze, Lettere e Arti, Tomo CLII, pag. 107 (1994)), nel quale vengono illustrate le benefiche conseguenze del consumo di vino rosso. Si veda inoltre A.Rigamonti e A.A.Varlamov, ne "Il Nuovo Saggiatore" vol.20, pag. 57 (2004).

futabile questa conclusione: “ l’assunzione di vino, rosso soprattutto e in particolare proveniente da alcune aree tipiche, determina una significativa riduzione del rischio di malattie cardiocircolatorie di vario genere”.

Per quale ragione? Il vino contiene oltre 2000 composti: acidi tartarico, citrico, malico, solforico e acetico; acetaldeide e diacetil chetoni; un’ampia varietà di esteri (in tracce) responsabili degli aromi; vanillina e acidi gallici; fenoli e antocianine responsabili della pigmentazione dei vini rossi; tracce di quasi tutti i minerali conosciuti. Fermiamo in particolare l’attenzione sui polifenoli (nella misura di circa 1 g/litro), sulle fitoalessine (presenti naturalmente nella buccia dell’uva) e sui flavonoidi. Tra le fitoalessine di particolare rilevanza appare il trans-resveratrolo, che direttamente o per effetto sinergico con altri composti manifesta una elevata attività anti-ossidante e previene anche l’invecchiamento cerebrale. Studi scientifici hanno mostrato che i polifenoli agiscono sulle lipoproteine e sulle piastrine ostacolando l’endotelina-1, molecola che provoca il restringimento dei vasi sanguigni ed è l’elemento determinante per lo sviluppo delle coronaropatie.

Senza sconfinare nel campo medico, vediamo invece se è possibile, come fisici, quantificare meglio queste osservazioni sul beneficio del vino. Tentiamo di stimare il numero di bottiglie di vino (preferibilmente rosso) che possano validamente preservarci da malattie cardiache senza provocare considerevole aumento di epatopatie. Sulla base dei dati ufficiali del progetto MONICA possiamo assumere che la probabilità di coronaropatie decresca con il numero di bottiglie di vino giornaliere b secondo la legge

$$I = I_0 \cdot e^{-b/b_i},$$

I_0 è la probabilità di malattie per una persona astemia e il valore caratteristico b_i si può assumere pari a una bottiglia di vino al dì^b.

Nei riguardi del probabile aumento del fattore di rischio di epatopatie

^bUn ipotetico consumo con $b \gg 1$ non è del tutto irrealistico. Al castello di Heidelberg, ove è conservata ancora oggi la più grande botte del mondo, che con un sistema di pompe riforniva il vino nella grande sala da pranzo, il consumo medio di vino per abitante era di due litri al dì. Il nano Percheo, buffone di corte di origini alto-atesine, consumava circa dieci bottiglie al dì. E non morì di malattia epatica ma perché, a seguito di una scommessa perduta, dovette bere due bicchieri d’acqua. Probabilmente la perdurante astinenza dal consumo di acqua, presumibilmente inquinata (a quel tempo quasi tutte le acque erano inquinate) gli procurò una gastroenterite fulminante che lo condusse alla tomba.

si può ipotizzare un aumento esponenziale:

$$C = C_0 e^{b/b_c},$$

dove C_0 è la probabilità per l'astemio e la costante caratteristica b_c si può stimare attorno a tre bottiglie di vino al dì (dai dati sulle patologie del fegato in paesi a forte consumo di super alcolici, tenendo conto del fattore di riduzione all'incirca pari a $1/4$). Sommando i due effetti si ha per la probabilità somma

$$W = I_0 \cdot e^{-b/b_i} + C_0 e^{b/b_c}.$$

Ponendo la derivata di W eguale a zero

$$b \frac{dW}{db} = -\frac{b}{b_i} I_0 \cdot e^{-b/b_i} + \frac{b}{b_c} C_0 e^{b/b_c} = 0,$$

si stima il minimo in corrispondenza al valore caratteristico b^* dato da

$$b^* = \frac{b_i b_c}{b_i + b_c} \left(\ln \frac{b_c}{b_i} + \ln \frac{I_0}{C_0} \right).$$

Pertanto

$$b^*/b_i = 0.75 (1.1 + \ln I_0/C_0) = 0.77 + 0.75 \ln I_0/C_0.$$

Supponendo, come è ragionevole, che I_0/C_0 abbia un valore all'incirca unitario si ricava una quantità ottimale di vino al dì pari a circa $0.77 b_i$, vale dire all'incirca mezzo litro di vino (rosso) al dì (meglio ai pasti).

13.7 Qualità e provenienza del vino: il metodo SNIF-NMR

Accertato pertanto il beneficio di un saggio consumo di vino, soprattutto rosso e di provenienza da particolari aree geografiche (e ovviamente non adulterato) si può porre il quesito: come stabilire la qualità e la provenienza del vino?

La fisica ha fornito un metodo scientifico: la determinazione attraverso NMR della Site – Specific Natural Isotope Fraction (SNIF). La risonanza magnetica nucleare (NMR) consiste essenzialmente nel seguente fenomeno, che descriveremo in qualche maggiore dettaglio nell'ultima parte del libro, quando tratteremo dei principi di visione all'interno del corpo umano (NMR-Imaging, noto più spesso come Risonanza Magnetica).

I nuclei, per esempio i protoni dell'atomo di idrogeno, posseggono un piccolo momento magnetico (l'analogo dell'ago della bussola). Questi momenti magnetici, se posti in un campo magnetico omogeneo, sono soggetti a un moto di precessione attorno alla direzione del campo con una frequenza Ω_L che è proporzionale al valore del campo $\mathbf{H}$. Immaginiamo ora di applicare un piccolo campo a radiofrequenza, prodotto da una bobina percorsa da corrente, in una direzione perpendicolare a quella del campo principale $\mathbf{H}$. Quando la frequenza ω di questo campo ausiliario è pari a Ω_L (ecco "la risonanza") si può produrre assorbimento di energia elettromagnetica e conseguentemente i momenti magnetici nucleari vengono rovesciati nella loro direzione rispetto al campo.

Questa descrizione vettoriale di tipo classico traduce sofisticati effetti di natura quantistica ma non è lontana dalla realtà fisica. Il punto che interessa qui è che attraverso le apparecchiature elettroniche è possibile misurare con elevatissima precisione il valore del campo magnetico di risonanza, mediante le "righe NMR" dello spettro di assorbimento. Il valore del campo magnetico al quale avviene la risonanza è quello locale al nucleo, vale a dire quello che il nucleo effettivamente sperimenta sulla base del campo esterno $\mathbf{H}$ e dei termini correttivi (assai piccoli ma osservabili) legati alle correnti elettroniche. Si comprende pertanto che le righe di risonanza cadono a diversi valori della frequenza di irraggiamento elettromagnetico, a seconda del circondario elettronico di un dato nucleo. Per esempio, per l'alcool etilico, $\text{CH}_3\text{CH}_2\text{OH}$, possiamo aspettarci righe di intensità relative 3:2:1 che cadono a frequenza diversa in corrispondenza ai protoni del gruppo molecolare CH_3 , di quelli del gruppo CH_2 e di quelli del gruppo OH . Lo spettro di risonanza diviene così una sorta di "fotografia" della configurazione molecolare.

Il metodo SNIF è stato ideato dai coniugi **Gerard e Maryvonne Martin** a Nantes negli anni 1980, allo scopo inizialmente di accertare l'arricchimento di vini mediante zuccheri. Nel 1987 è stata fondata una società, la Eurofins Scientific, che ha raccolto un "data base" *ad hoc* (che oggi contiene gli spettri NMR dei vini di Francia, Spagna, Germania e Italia). Il metodo è stato accettato ufficialmente dalla Comunità Europea (1989) e dalla Organization International de la Vigne e du Vin (OIV) e riconosciuto metodo ufficiale della Association of Official Analytical Chemistry (AOAC) in USA e Canada.

Con il metodo SNIF è oggi possibile individuare l'alcool etilico che abbia la stessa struttura chimica ma diversa origine botanica. Si può determinare

se un vino proviene realmente solo dal mosto di una particolare vigna, di una particolare regione. Il metodo si basa sul fatto che a seguito dei diversi processi di fotosintesi, di metabolismo della pianta, delle condizioni geografiche e climatiche, la frazione di deuterio rispetto all'idrogeno è diversa da zona a zona della terra e da pianta a pianta. La frazione di deuterio (D) rispetto all'idrogeno (H) è misurata usualmente in parti per milione, ppm. Vale 16.000 su Venere, 0,01 nella troposfera terrestre e sulla terra varia da 90 al polo Sud a 160 all'equatore. Il rapporto (D/H) presenta una relativamente grande variabilità. Ulteriore elemento che si presta alla caratterizzazione è costituito dalla distribuzione del deuterio sui diversi gruppi molecolari. Ancora per alcool etilico, in luogo di $\text{CH}_3\text{CH}_2\text{OH}$ possiamo avere $\text{CH}_2\text{D}-\text{CH}_2-\text{OH}$ o $\text{CH}_3-\text{CDH}-\text{OH}$ o ancora $\text{CH}_3-\text{CH}_2-\text{OD}$. Le percentuali D/H per ogni gruppo individuale possono essere ottenute dagli spettri NMR del nucleo di deuterio, i picchi corrispondenti ai diversi gruppi presentandosi a valori separati di radiofrequenza in ragione delle correzioni al campo locale prodotte dalle correnti diamagnetiche degli elettroni del gruppo stesso. Dai segnali NMR, confrontati mediante un opportuno programma computazionale con la raccolta di dati (sul "data base") si possono determinare i tipi di zuccheri aggiunti, la provenienza di vini eventualmente usati per l'arricchimento e altro.

È certamente confortante sapere che un responsabile consumo di vino, rispettoso della nobiltà di questa bevanda, è benefico e che indagini scientifiche oggettive hanno confermato e dato corpo quantitativo a un sentimento popolare di apprezzamento nei suoi confronti.

È altrettanto confortante sapere che esiste un metodo scientifico che ci consente di accertare la naturale qualità del vino, la mancanza di manipolazioni e di aggiunte, la provenienza da aree e da vitigni che garantiscono la massimizzazione dei benefici associati alla fruizione di una bottiglia di buon vino.

PARTE III

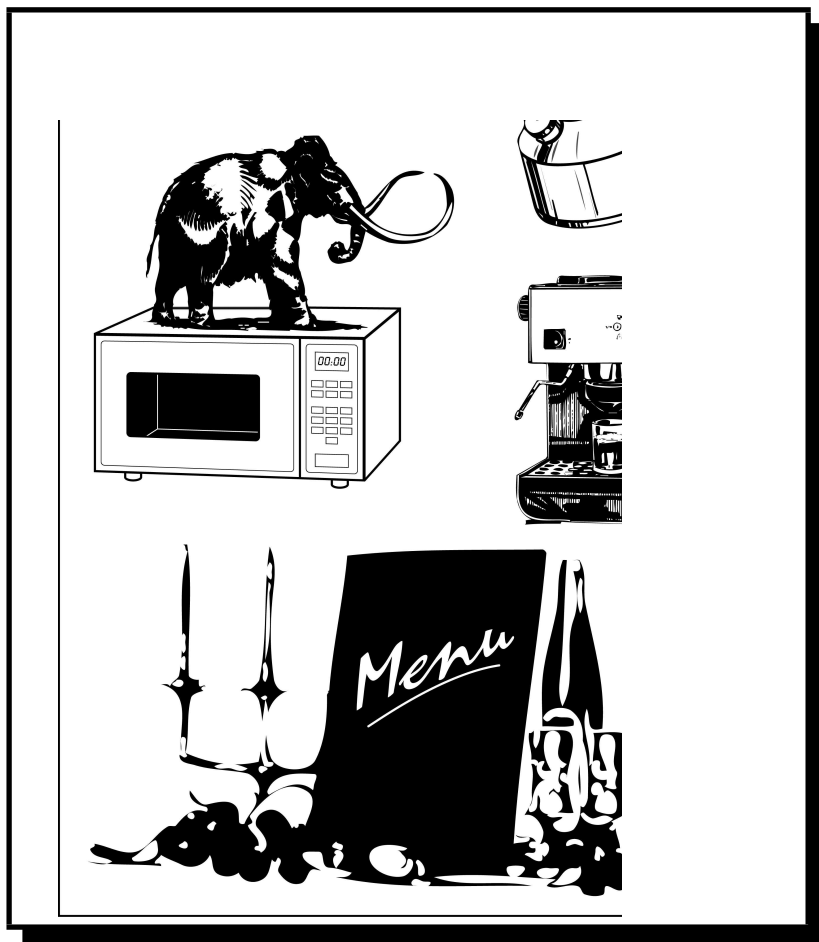
Da “Homo erectus” a “cuoco sapiens”

“Homo erectus” è comparso sulla terra all’incirca un milione di anni or sono. Nel corso di 10.000 secoli si sono prodotte grandissime trasformazioni nel suo modo di vita, nella sua conoscenza e nella sua capacità creativa. Emblematicamente, ha nel contempo avuto luogo una grande trasformazione nella manipolazione dei cibi. La “cucina” di “Homo erectus” consisteva inizialmente di un fuoco comune nel villaggio, sul quale arrostitavano carni di elefanti o bisonti, carni che erano manipolate con oggetti ricavati da pietre o ossa levigate.

L’uomo moderno manipola e fruisce dei cibi in modo fortunatamente molto diverso e ha trasformato la necessità di alimentarsi in una “arte”: frequentemente i pasti sono divenuti raffinati momenti conviviali. Alcuni ritengono che la culinaria sia quasi divenuta una scienza.

In questa parte del testo desideriamo illustrarvi come in una cucina si possano osservare molti interessanti fenomeni che sono dominio della fisica. Toccheremo diversi argomenti attinenti alla cucina, dalle modalità di cottura del forno a micro-onde, al fenomeno della cottura degli spaghetti o di un grosso tacchino, al bollire di una teiera, a quali siano le ragioni di una migliore qualità del caffè prodotto da una macchina da bar rispetto alla moka.

Si discuterà del collidere di uova, per arrivare ad accennare alla cosiddetta “cucina molecolare” che permette di ideare scientificamente nuovi metodi di preparazione e manipolazione di cibi aventi particolari proprietà.



Capitolo 14

La Signora prepara il tè.

Al rito del tè sono stati dedicati in Oriente interi trattati. Eppure, guardando ad esso da un altro punto di vista, magari mentre la Signora che ci ha invitato al tè delle cinque attende che la teiera arrivi a ebollizione, vi possiamo trovare molti fenomeni interessanti, della cui spiegazione in genere i manuali di cucina non si occupano.

Per iniziare facciamo la seguente esperienza. Mettiamo due teiere perfettamente uguali con la stessa quantità di acqua fredda su due fornelli di uguale potenza. Chiudiamo una teiera con un coperchio, lasciamo scoperta la seconda. Quale delle due bollerà per prima? Ogni esperta Signora conosce la risposta a questa domanda. Volendo portare ad ebollizione l'acqua più in fretta possibile, senza esitazione copre la pentola con un coperchio pensando che bollerà per prima la teiera con il coperchio. Ci proponiamo di motivare questo fatto dal punto di vista della fisica molecolare.

Intanto, mentre le due teiere si riscaldano, mettiamo su un terzo fornello, una teiera perfettamente uguale e con la stessa quantità d'acqua fredda contenuta nelle prime due e cerchiamo di farla bollire più in fretta (con la stessa potenza del fornello). Per far ciò è necessario in qualche modo far salire più rapidamente la temperatura dell'acqua in essa contenuta, rispetto alle prime due. Per esempio, aggiungendo nella terza teiera dell'acqua calda. È evidente che per aumentare la temperatura dell'acqua in un qualsiasi recipiente è sufficiente aggiungere dell'acqua bollente. Si potrebbe per questa via accelerare l'ebollizione nella terza teiera? Certamente no. Questo procedimento non soltanto non accelera l'ebollizione, ma anzi la rallenta.

Per convincersene, immaginiamo che l'acqua con una massa m_1 che si trova inizialmente nella teiera ad una temperatura T_1 , non sia stata

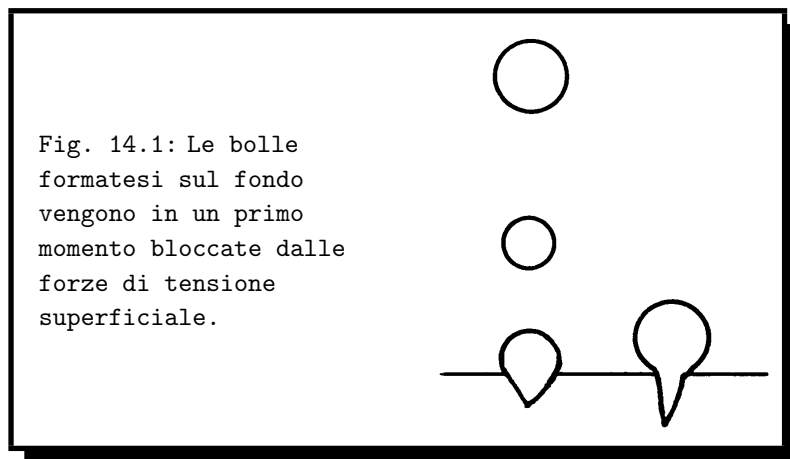
miscelata e non abbia iniziato a scambiare calore con l'acqua aggiunta, la quale abbia massa m_2 a temperatura T_2 . Il calore necessario inizialmente all'acqua m_1 per portarla ad ebollizione era $Q_1 = c m_1 (T_{eb} - T_1)$ (c calore specifico). Ora, dopo l'aggiunta della acqua a temperatura più elevata oltre al riscaldamento fino alla temperatura di ebollizione (T_{eb}) della quantità d'acqua m_1 , bisognerà riscaldare da T_2 fino a T_{eb} anche l'acqua aggiunta, con massa m_2 . Per questo la quantità totale di calore necessaria diviene

$$Q_1 = c m_1 (T_{eb} - T_1) + c m_2 (T_{eb} - T_2).$$

Anche se nella teiera si aggiunge acqua bollente, durante il travaso essa un poco si raffredda e la temperatura T_2 risulta comunque un po' più bassa di T_{eb} . La nostra irrealistica ipotesi che le porzioni d'acqua non si mescolino non ha influito in alcun modo sulla legge di conservazione dell'energia nel sistema; essa ci ha soltanto permesso di esaminare il fenomeno più direttamente e rapidamente.

Mentre abbiamo inutilmente perso tempo con la terza teiera, quella con il coperchio ha già cominciato a rumoreggiare. Cerchiamo di chiarire il motivo di questo rumore e di valutare la sua frequenza caratteristica. Possiamo ipotizzare che la prima causa del rumore siano le oscillazioni del liquido che hanno origine durante il distacco delle bolle di vapore dal fondo e dalle pareti del recipiente. Queste bollicine si formano nelle disomogeneità e nelle microfessure della superficie. Le loro dimensioni caratteristiche, fino a che non si ha l'ebollizione, sono dell'ordine di 1 mm (durante l'ebollizione esse crescono notevolmente e possono raggiungere 1 cm). Per valutare la frequenza del suono che compare nella teiera dobbiamo trovare il tempo di distacco delle bollicine dal fondo. Questo tempo caratterizza la durata della spinta che riceve il liquido nel processo di distacco delle bollicine, e, di conseguenza, il periodo delle oscillazioni che compaiono nel liquido. La frequenza del suono generato viene definita come l'inverso del tempo, $\nu \sim \tau^{-1}$. Mentre la bollicina è in quiete sul fondo su di essa agiscono due forze: la forza di Archimede che la spinge verso l'alto, $F_A = \rho_w g V_b$ (V_b è il volume della bollicina) e la forza di tensione superficiale che la trattiene sul fondo $F_S = \sigma l$ (l è la il perimetro della zona di contatto tra la bollicina e la superficie del fondo)^a. All'aumentare del volume della bollicina, la forza di Archimede cresce e ad un certo istante supera la forza di tensione superficiale. La bollicina inizia così il suo moto verso l'alto (Fig. 14.1).

^aIl contributo della forza di gravità può essere considerato trascurabile.



È intuitivo che la forza totale agente sulla bollicina quando si distacca dal fondo è dell'ordine di grandezza di F_A . La massa da associare alla bollicina durante il suo moto nel liquido non è semplicemente la sua reale massa, che ha un valore assai piccolo (massa dell'aria e del vapore in essa racchiusi), ma dalla cosiddetta “*massa aggiunta*” che, per una bollicina sferica, è

$$m^* = \frac{2}{3}\pi \rho_w r_0^3 = \frac{1}{2}\rho_w V, \quad (14.1)$$

dove ρ_w è la densità dell'acqua mentre r_0 è il raggio. Tale quantità in pratica definisce la massa di liquido intorno alla bollicina che viene coinvolta nel moto durante lo spostamento verso l'alto. Quindi l'accelerazione della bollicina all'inizio del moto è

$$a \sim \frac{F_A}{m^*} = 2g.$$

Il tempo di distacco della bollicina dal fondo può essere ora valutato considerando il moto uniformemente accelerato. In un tempo pari a

$$\tau_1 \sim \sqrt{\frac{2r_0}{a}} \sim 10^{-2} s \quad (14.2)$$

essa sale sino ad un'altezza pari al suo diametro.

La frequenza caratteristica del suono generato dal distacco delle bollicine è quindi $\nu_1 \sim \tau_1^{-1} \sim 100 Hz$. Questo valore è di un ordine di

grandezza inferiore alla frequenza percepita durante il riscaldamento della teiera (comunque molto prima dell'inizio dell'ebollizione dell'acqua)^b.

Il rumore che compare nella teiera durante il suo riscaldamento deve quindi avere anche un'altra causa. Per individuarla, seguiamo la sorte della bollicina di vapore dopo il suo distacco dal fondo. Staccandosi dal fondo caldo, dove la pressione del vapore era all'incirca uguale a quella atmosferica (altrimenti essa non potrebbe espandersi a sufficienza per l'emersione), la bollicina perviene durante la salita negli strati d'acqua superiori, non ancora sufficientemente riscaldati. Il vapore saturo che costituisce la bollicina in questo caso si raffredda, la sua pressione diminuisce e non può più compensare la pressione esterna da parte dell'acqua circostante.

A seguito di ciò la bollicina implode rapidamente o si comprime fortemente (se in essa oltre al vapore acqueo si trovava anche una certa quantità di aria) e nel liquido si propaga un impulso sonoro. L'implosione contemporanea di un gran numero di queste bollicine che si disintegrano negli strati superiori dell'acqua viene percepita come rumore. Calcoliamone la frequenza. Scriviamo l'equazione di Newton per la massa m d'acqua che si dirige all'interno della bolla nel corso dell'implosione:

$$a_r = F_p = S \Delta P,$$

ove $S = 4\pi r^2$ rappresenta l'area della superficie della bollicina, su cui agisce la forza di pressione F_p , ΔP è la differenza della pressione sul contorno, a_r è l'accelerazione della bollicina verso il suo centro. È chiaro che nel processo di implosione è coinvolta una massa uguale come ordine di grandezza al prodotto della densità dell'acqua per il volume della bollicina: $m \sim \rho_w r^3$. Quindi l'equazione di Newton può essere riscritta nella forma

$$\rho_w r^3 a_r \sim r^2 \Delta P.$$

Trascurando la pressione determinata dalla tensione della superficie incurvata della bollicina, e anche dalla piccola quantità d'aria che può essere in essa contenuta, considereremo ΔP costante (dipendente solo dalla differenza di temperatura tra gli strati d'acqua adiacenti al fondo e quelli

^bSi noti che nell'espressione non è intervenuto il coefficiente di tensione superficiale. Si può pensare che la bollicina emetta onde sonore non soltanto durante il processo di distacco dal fondo, ma durante tutto il moto di carattere accelerato, cioè fino a quando la forza di resistenza al moto da parte dell'acqua (proporzionale alla velocità) non equilibra la forza che agisce sulla bolla.

superiori). L'accelerazione $a_r = r''$ può essere scritta come r_0 / τ_2^2 , dove τ_2 è il tempo incognito di implosione della bollicina. Allora

$$\rho_w \frac{r_0^2}{\tau_2^2} \sim \Delta P,$$

da cui

$$\tau_2 \sim r_0 \sqrt{\frac{\rho_w}{\Delta P}}. \quad (14.3)$$

Vicino a $T_{eb} = 100^\circ C$ la pressione del vapore saturo cade all'incirca di $3 \cdot 10^3 Pa$ a seguito della diminuzione della temperatura di $1^\circ C$ (si veda la Tabella 14.1). Per questo si può assumere $\Delta P \sim 10^3 Pa$ e pertanto il tempo di implosione della bollicina risulta $10^{-3}s$, mentre la frequenza caratteristica del rumore che si sviluppa in conseguenza è dell'ordine di $10^3 Hz$. Questo risultato è di per se stesso più vicino alla realtà di quello valutato precedentemente.

Tabella 14.1: Dipendenza della pressione del vapore saturo dalla temperatura.

Temperatura, $^\circ C$	96.18	99.1	99.6	99.9	100	101	110.8
Pressione, kPa	88.26	98.07	100	101	101.3	105	147

Anche il fatto che a seguito dell'aumento della temperatura dell'acqua la frequenza del rumore caratteristico diminuisca gradatamente, in accordo alla equazione (14.3) (in quanto ΔP diminuisce) è un ulteriore argomento a favore di una tale origine del rumore. Immediatamente prima dell'ebollizione, le bollicine smettono di implodere anche negli strati superiori dell'acqua. Quindi come unico meccanismo d'eccitazione del suono rimane il distacco delle bollicine dal fondo: la frequenza del "canto" della teiera diminuisce notevolmente, in accordo a quanto abbiamo già visto. Dopo l'inizio dell'ebollizione la "voce" della teiera può cambiare nuovamente (soprattutto se si alza il coperchio): sono le bollicine che gorgogliano scoppiando direttamente sulla superficie dell'acqua. In questo caso è importante anche il volume d'acqua contenuto nella teiera e la forma della stessa.

Arriviamo quindi alla conclusione che il rumore della teiera, prima dell'inizio dell'ebollizione, è legato al distacco dal fondo e alla dissolvenza negli strati d'acqua superiori non ancora ben riscaldati delle bollicine di vapore.

Particolarmente interessante è lo studio di questi processi in una teiera di vetro con pareti trasparenti, durante il riscaldamento dell'acqua. Non siamo i primi ad interessarsi e risolvere i quesiti insiti in questo fenomeno. Già nel XVIII secolo, lo scienziato scozzese **Josef Black** aveva studiato il “canto” di recipienti riscaldati e stabilito come “a questo canto partecipa un duetto: le bollicine di aria riscaldata in ascesa e la vibrazione delle pareti del recipiente”.

Quindi, come ci aspettavamo, a bollire per prima è stata la teiera con il coperchio. E lo annuncia il getto di vapore che esce dal beccuccio. Qual'è la velocità di tale getto?

Non è difficile rispondere a questo quesito se osserviamo che nel processo di ebollizione praticamente tutta l'energia termica fornita alla teiera determina evaporazione. I filatelici sanno che quando si vuole staccare un francobollo da una busta, l'acqua viene versata solo sul fondo della teiera, affinché tutto il vapore che si forma esca dal beccuccio. Supporremo che nel nostro caso il beccuccio sia libero e il vapore esca all'esterno solo attraverso di esso. Ipotizziamo che avendo fornito energia per il tempo Δt , evapori una massa d'acqua ΔM . Questa può essere definita dall'equazione

$$r \Delta M = \mathcal{P} \Delta t,$$

dove r è il calore di evaporazione, mentre $\mathcal{P}$ è la potenza del riscaldatore. In questo intervallo di tempo la massa di vapore formatasi deve lasciare la teiera attraverso il beccuccio, altrimenti il vapore comincerebbe ad accumularsi sotto il coperchio. Se l'area del foro d'uscita del beccuccio è s , la densità del vapore sotto il coperchio $\rho_s(T_{eb})$ e la velocità incognita è v , si può scrivere

$$\Delta M = \rho_s(T_b) s v \Delta t.$$

La densità del vapore saturo a $T_{eb} = 373K$ si può ottenere dalla Tabella 14.1 (0.6 kg/m^3), oppure, se non si ha la tabella sotto mano, si può valutare dall'equazione di **Mendeleev-Klaapeyron**:

$$\rho_s(T_b) = \frac{m}{V} = \frac{P_s(T_{eb}) \mu_{\text{H}_2\text{O}}}{R T_{eb}}.$$

Assumendo, per la valutazione numerica, $\mathcal{P} = 500W$, $s = 2\text{ cm}^2$, $r = 2.26 \cdot 10^5 J/kg$, $P_s(T_{eb}) = 10^5 Pa$ and $R = 8.31 J/1^\circ K \cdot mole$, si ricava

$$v = \frac{\mathcal{P} R T_b}{r P_s(T_b) \mu_{H_2O} s} \sim 1 m/s.$$

Ora inizia a bollire anche la teiera senza coperchio, in notevole ritardo rispetto alla prima. E non scordiamoci che è opportuno rimuoverla con attenzione dal fornello, per evitare di scottarsi con il vapore. Che cosa “brucia” di più: il vapore o l’acqua bollente?

Prima di rispondere occorre meglio formulare questa domanda: cosa brucia di più, una determinata massa di acqua bollente o la stessa massa di vapore? Torniamo di nuovo alle valutazioni. Poniamo che il volume occupato dal vapore saturo a cento gradi sotto il coperchio della teiera sia $V_1 = 1\text{ l}$ e, per esempio, che un decimo si condensi sulla mano. Come abbiamo già visto, la densità del vapore a $T_{eb} = 100^\circ C$ è pari a 0.6 kg/m^3 . Per questo sulla mano riceveremo all’incirca una massa m_s di 0.06 g di vapore. Nella condensazione e nel successivo raffreddamento da $100^\circ C$ fino alla temperatura ambiente $T_0 \approx 20^\circ C$ verrà ceduta una quantità di calore $\Delta Q = r m_s + c m_s (T_{eb} - T_0)$.

È facile convincersi che per lo stesso effetto termico è necessaria una massa di acqua bollente quasi dieci volte maggiore. Oltre a ciò, in caso di ustione da vapore, l’area colpita sarà notevolmente più estesa. Quindi, il vapore “brucia” di più dell’acqua calda, in primo luogo a causa del notevole calore che viene ceduto durante la condensazione.

Ci siamo, però, distratti dalla discussione del nostro ipotetico esperimento con due teiere. Perché la teiera senza coperchio è arrivata all’ebollizione più tardi? La risposta è quasi evidente: durante il riscaldamento dell’acqua nella teiera senza coperchio, le molecole hanno la possibilità di abbandonare più rapidamente e senza vincoli il liquido, sottraendo in questo modo energia e allo stesso tempo raffreddando con una certa efficacia l’acqua restante nella teiera (questo processo non è altro che l’evaporazione). Per questo il riscaldatore, in questo caso, deve non soltanto portare l’acqua della teiera ad ebollizione ma anche far evaporare una parte della stessa durante il riscaldamento. È chiaro che per far ciò occorre energia maggiore (e di conseguenza più tempo) che per l’ebollizione dell’acqua nella teiera coperta, dove le molecole “veloci” che si staccano dall’acqua formano rapidamente il vapore nello spazio chiuso sotto il coperchio e ritornando all’acqua vi restituiscono la propria energia in eccesso.

Hanno in realtà luogo anche due effetti che agiscono in senso opposto a quello esaminato. In primo luogo, durante il processo di evaporazione diminuisce la massa d'acqua che dobbiamo portare ad ebollizione. In secondo luogo l'ebollizione in una teiera aperta con la normale pressione atmosferica ha luogo alla temperatura di $100^{\circ}C$. In una teiera chiusa, se essa è riempita in modo tale che i vapori non possono uscire attraverso il beccuccio, a causa dell'evaporazione prima che si pervenga all'ebollizione la pressione sulla superficie aumenta. Infatti essa è determinata dalla somma delle pressioni parziali della piccola quantità d'aria che si trova sotto il coperchio e dallo stesso vapore d'acqua. All'aumentare della pressione esterna anche la temperatura d'ebollizione dell'acqua deve essere più alta, essendo determinata dalla condizione di eguaglianza della pressione del vapore saturo nella bolla che si forma nel liquido con la pressione esterna. Quale degli effetti considerati sarà quello dominante?

Nel caso in cui nascono dubbi di questo tipo, bisogna rifarsi ad una valutazione accurata per lo meno in merito agli ordini di grandezza degli effetti in considerazione. Calcoleremo dapprima la massa d'acqua che evapora dalla teiera aperta mentre la si porta ad ebollizione.

Le molecole in un liquido interagiscono abbastanza fortemente l'un l'altra. Tuttavia, se nel solido corrispondente l'energia potenziale d'interazione supera notevolmente l'energia cinetica del moto degli atomi o delle molecole, e nel gas, al contrario, l'energia cinetica del moto caotico supera notevolmente l'energia potenziale delle molecole, in un liquido queste grandezze sono dello stesso ordine. Per questo le molecole del liquido compiono delle oscillazioni termiche attorno alle posizioni di equilibrio, raramente "saltando" in altre. "Raramente", si intende in confronto con il periodo delle oscillazioni vicino alla posizione di equilibrio. In una scala di tempo per noi naturale, perfino molto spesso: in un secondo una molecola nel liquido può cambiare la sua posizione di equilibrio milioni di volte! Tuttavia, non tutte le molecole durante i propri movimenti, anche trovandosi nei pressi della superficie, possono fuoriuscire dal liquido. Per superare le forze di interazione la molecola deve compiere un lavoro. Si può dire che l'energia potenziale della molecola nel liquido è inferiore alla sua energia potenziale nello stato di vapore di una quantità uguale al calore di evaporazione per molecola. Se r è il calore di evaporazione, il calore molare dell'evaporazione è uguale a μr e, per una molecola, $U_0 = \mu r / N_A$, (N_A numero di Avogadro). La molecola può compiere questo "lavoro d'uscita" soltanto grazie alla sua energia termica. In genere l'energia cinetica media $\bar{E}_c \approx kT$ (con $k = 1.38 \cdot 10^{-23} J / 1^{\circ}K$)

è notevolmente inferiore a U_0 . Tuttavia, secondo le leggi della fisica molecolare, nel liquido esiste comunque una certa quantità di molecole con una energia sufficientemente alta da superare la forza di attrazione ed evaporare. Il loro numero per unità di volume è dato dall'espressione

$$n_{E_k > U_0} = n_0 e^{-\frac{U_0}{kT}}, \quad (14.4)$$

dove n_0 è la concentrazione iniziale delle molecole, ed $e = 2.7182\dots$ è la base dei logaritmi naturali.

Tralasciamo ora i salti delle molecole di liquido da un posto all'altro ed immaginiamo delle molecole ad alta energia come in un gas. In un tempo Δt da una porzione di superficie di area S possono evaporare molecole ad alta energia pari a $\Delta N \sim \frac{1}{6} n S v \Delta t$, (per tale valutazione assumiamo che $\sim \frac{1}{6}$ di tutte queste molecole si avvicinano alla superficie con una velocità $v \sim \sqrt{U_0/m_0}$). Utilizzando l'espressione (14.4) ricaviamo la velocità d'evaporazione

$$\frac{\Delta N}{\Delta t} \sim \frac{n S v \Delta t}{6 \Delta t} \sim S n_0 \sqrt{\frac{U_0}{m_0}} e^{-\frac{U_0}{kT}}.$$

La massa evaporata nell'unità di tempo è pertanto

$$\frac{\Delta m}{\Delta t} \sim m_0 \frac{\Delta N}{\Delta t} \sim m_0 S n_0 \sqrt{\frac{U_0}{m_0}} e^{-\frac{U_0}{kT}}. \quad (14.5)$$

Ci sarà più agevole calcolare la quantità di acqua che evapora dalla teiera quando la temperatura aumenta di $1 K$. Per questo si utilizza la legge di conservazione dell'energia. Nel tempo Δt la teiera riceve dal fornello una quantità di calore $\Delta Q = \mathcal{P} \Delta t$ ($\mathcal{P}$ è la potenza effettiva dell'elemento riscaldante). La temperatura dell'acqua aumenta di ΔT , dato dalla relazione

$$\mathcal{P} \Delta t = c M \Delta T,$$

dove c e M sono, rispettivamente, la capacità termica e la massa dell'acqua nella teiera (trascuriamo la capacità termica della teiera). Sostituendo nella formula per la velocità di evaporazione $\Delta t = c M \Delta T / \mathcal{P}$ otteniamo

$$\frac{\Delta m}{\Delta T} = \frac{\rho c S M}{\mathcal{P}} \sqrt{\frac{r \mu_{H_2O}}{N_A m_0}} e^{-\frac{r \mu_{H_2O}}{N_A k T}} = \frac{\rho c S M \sqrt{r}}{\mathcal{P}} e^{-\frac{r \mu_{H_2O}}{N_A k T}}.$$

Durante il riscaldamento la temperatura sale da quella ambiente fino a $373 K$. Tenendo conto che la maggior fuoriuscita di massa ha luogo a

temperature abbastanza alte, assumiamo nell'esponente una temperatura media $\bar{T} = 350\text{ K}$. Per le restanti grandezze prenderemo $\Delta T = T_{eb} - T_0 = 80^\circ\text{ K}$, $S \sim 10^{-3}\text{ m}^2$, $\rho \sim 10^3\text{ kg/m}^3$, $\mu_{\text{H}_2\text{O}} = 0.018\text{ kg/mol}$ e $c = 4.19 \cdot 10^3\text{ J/kg}$, ad ottenere

$$\frac{\Delta m}{M} \approx \frac{\rho c S}{\mathcal{P}} \sqrt{r} e^{-\frac{r \mu_{\text{H}_2\text{O}}}{R \bar{T}}} (T_b - T_0) \approx 3\%.$$

Quindi durante il processo di riscaldamento senza coperchio fino alla temperatura di ebollizione alcuni percento della massa d'acqua evaporano. L'evaporazione di questa massa avviene grazie al riscaldatore e, naturalmente, favorisce l'inizio dell'ebollizione nella teiera scoperta. Per capire di quanto, è sufficiente notare che per il riscaldatore questa evaporazione è equivalente a portare ad ebollizione 1/4 di tutta la massa d'acqua nella teiera.

Ora passiamo all'esame degli effetti che rallentano l'inizio dell'ebollizione nella teiera con il coperchio. Sul primo di questi effetti, ossia dell'invariabilità della massa durante il processo di riscaldamento, non c'è nulla da aggiungere. Or ora è stato osservato che l'ebollizione del 3% della massa è equivalente al riscaldamento del 25% della massa e pertanto si può trascurare il riscaldamento di questo 3% nella teiera coperta. Il secondo effetto - l'aumento di pressione sotto il coperchio - è, analogamente, trascurabile rispetto all'evaporazione. In realtà, anche se la teiera fosse completamente riempita (il vapore non fuoriesce dal beccuccio) la pressione in eccesso rispetto a quella atmosferica non può superare il peso del coperchio diviso per la sua area (in caso contrario, il coperchio si sposterebbe, facendo uscire il vapore). Ponendo $m_{cop} = 0.3\text{ kg}$ e $S_{cop} \sim 10^2\text{ cm}^2$, otterremo

$$\Delta P \leq \frac{m_{cop} g}{S_{cop}} \approx \frac{3\text{ N}}{10^{-2}\text{ m}^2} = 3 \cdot 10^2\text{ Pa}.$$

Riferendosi nuovamente alla Tabella 14.1, osserviamo che questo aumento di pressione fa variare la temperatura di ebollizione al più di $\delta T_{eb} \sim 0.5^\circ\text{ C}$. Conseguentemente, per pervenire all'ebollizione è necessaria una quantità aggiuntiva di calore $\delta Q = c M \delta T_{eb}$.

Confrontando le grandezze $c M \delta T_{eb}$ e $r \Delta M$ vediamo che la disuguaglianza si realizza nella misura di 30:1. Quindi anche l'aumento della temperatura di ebollizione nella teiera coperta, riempita fino all'orlo, non può avvicinare sensibilmente l'effetto associato all'evaporazione dell'acqua nella teiera senza coperchio.

Sul principio dell'aumento di pressione durante il riscaldamento dell'acqua in un volume chiuso è costruita la pentola a pressione. In questo caso, al posto del beccuccio c'è un piccolo foro della valvola di sicurezza, che si apre soltanto quando venga raggiunta una certa pressione, per il resto essa è ermetica. A seguito dell'evaporazione del liquido nel volume chiuso la pressione nella pentola aumenta all'incirca fino a $1.4 \cdot 10^5 Pa$, valore al quale entra in funzione la valvola; la temperatura di ebollizione, secondo la Tabella 14.1, sale sino a $T_{eb}^* = 108^\circ C$.

Ciò permette di preparare i cibi molto più rapidamente che in una pentola normale. Dopo aver tolto la pentola dal fornello, bisogna aprirla con attenzione: la pressione infatti diminuisce rapidamente, mentre l'acqua contenuta è *surriscaldata*. Per questo una massa di liquido δm , tale che $r \delta m = c M (T_{eb}^* - T_b)$ evapora in maniera esplosiva e può causarci scottature. Il liquido in questo caso bolle immediatamente in tutto il volume della pentola.

Si può anche osservare che in alta montagna, dove la pressione atmosferica è bassa, in una pentola normale non si riesce a cuocere la carne: la temperatura di ebollizione dell'acqua può scendere fino a $70^\circ C$ (sulla cima dell'Everest la pressione è di $P_{8848} = 3.5 \cdot 10^4 Pa$). Per questo gli alpinisti nelle scalate spesso portano con se una pentola a pressione. Oltre al fatto che si può raggiungere una più alta temperatura, la pentola a pressione permette anche di risparmiare combustibile, il che compensa parzialmente il suo peso nello zaino.

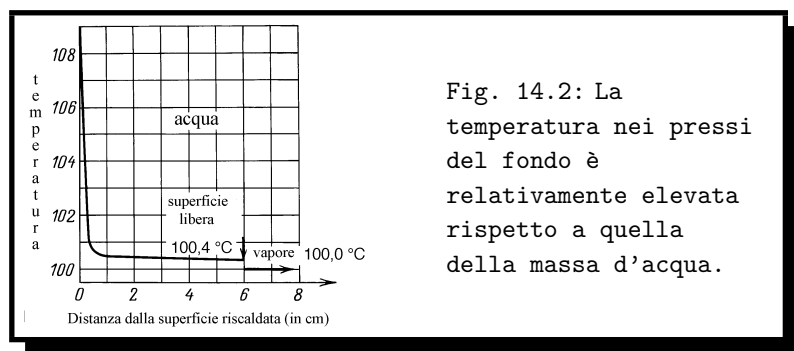
Ci siamo distratti a discutere e intanto nelle teiere l'acqua si trasforma in vapore! Togliendo dal riscaldatore la teiera con il coperchio si osserva che non smette di bollire immediatamente: per un certo tempo dal suo beccuccio continua a fuoriuscire un getto di vapore. Quale quantità d'acqua bolle dopo che la teiera è stata rimossa dal riscaldatore e per quanto tempo?

Per rispondere a queste domande, riferiamoci alla figura 14.2, nella quale è mostrata la distribuzione della temperatura dell'acqua in funzione del livello nel corso della ebollizione, in un recipiente al fondo del quale si fornisce calore. Dalla figura si vede che lo strato d'acqua adiacente al fondo, di spessore circa $\Delta H = 0.5 cm$, è fortemente surriscaldato e la sua temperatura varia da $T_f = 110^\circ C$ a $T_m = 100.5^\circ C$ (T_f è la temperatura del fondo della teiera). La temperatura della restante massa d'acqua (ipotizziamo un livello d'acqua $H = 10 cm$), in accordo al grafico, è di circa $100.5^\circ C$, mentre sulla superficie libera del liquido avviene un salto di temperatura di $\Delta T = 0.4^\circ C$.

Perciò, dopo l'interruzione dell'alimentazione di calore, nella massa d'acqua vi è una riserva di quantità di calore in eccesso rispetto all'equilibrio data da

$$\Delta Q = c \rho S \Delta H \left(\frac{T_f - T_m}{2} \right) + c \rho S H \Delta T.$$

dove S è l'area del fondo (ipotizziamo la teiera di forma cilindrica). In realtà è surriscaldato anche il fondo della teiera, ma non è difficile rendersi conto che questo fatto può essere trascurato a causa della grande capacità termica dell'acqua.



La quantità di calore ΔQ in eccesso determina pertanto l'evaporazione di uno strato di liquido δH , la cui massa δm può essere ricavata dall'equazione

$$r \delta m = \rho S \delta H r = c \rho S \left[\Delta H \left(\frac{T_f - T_m}{2} \right) + H \Delta T \right],$$

dalla quale otteniamo per δH

$$\frac{\delta H}{H} = \frac{c}{r} \left[\frac{\Delta H}{H} \left(\frac{T_f - T_m}{2} \right) + \Delta T \right] \approx 2 \cdot 10^{-3}.$$

Quindi, dopo che la teiera è stata rimossa dal fornello evapora ancora all'incirca lo 0.2% del contenuto.

Il tempo caratteristico di evaporazione di tutta la quantità d'acqua contenuta nella teiera, di massa $M = 1 \text{ kg}$, con una potenza $\mathcal{P} = 500 \text{ W}$ è

$$\tau = \frac{r M}{\mathcal{P}} \approx 5 \cdot 10^3 \text{ s}.$$

Lo 0.2% della sua massa evapora dunque in un tempo dell'ordine di 10 secondi (se si suppone che la velocità di evaporazione non cambi, come ordine di grandezza, rispetto al regime stazionario).

Così, mentre noi disquisivamo, la Signora ritiene sia giunta l'ora di versare il tè.

I popoli orientali lo berrebbero probabilmente dalle ciotole. Le ciotole pare siano state introdotte dai nomadi. La loro forma è molto comoda da trasportare, occupano poco spazio e sono resistenti. Oltre a ciò hanno un grande vantaggio rispetto ai comuni bicchieri. La forma delle ciotole è tale che lo strato superiore del liquido in esse versato, grazie alla maggiore superficie esposta, si raffredda più rapidamente che in un bicchiere, così che lo si può bere senza scottarsi. Nello stesso tempo, lo strato inferiore, che sarà bevuto successivamente, rimane caldo più a lungo. In Azerbagian esiste un recipiente ancora più particolare per bere il tè: l'armudi, a forma di pera. Anche nell'armudi la grande superficie esposta si accoppia con una più piccola superficie (in sezione) del contenitore del tè caldo (la parte inferiore del recipiente).

La vecchie tazze di porcellana venivano anch'esse realizzate in modo da presentare una grande superficie libera. I bicchieri come stoviglie per bere il tè sono entrati in uso nel XIX secolo a causa dell'aumento del prezzo delle tazze di porcellana. Questi prodotti dell'arte decorativa erano usati dalle donne, mentre gli uomini preferivano usare bicchieri di vetro, che in seguito si sono arricchiti di sotto-bicchieri d'argento con il monogramma dei padroni di casa.

Infine, chiediamoci: dal punto di vista fisico, l'argento e l'alluminio come materiali per porta-bicchieri sono appropriati? A quali requisiti deve soddisfare il materiale di cui sono fatti i porta-bicchieri?

Capitolo 15

Porchetta alle micro-onde?

Il giorno in cui l'uomo di Neandertal ha appreso come accendere il fuoco può essere considerato l'inizio dell'era che metteva fine al suo passato primordiale. Il fuoco, con il tempo avrebbe offerto la possibilità di lavorare i metalli, costruire automobili e viaggiare nello spazio e, cosa da alcuni magari ritenuta più importante, liberarsi dalla necessità di mangiare carne cruda! Una bistecca ben cotta può essere presa a simbolo della civiltà accanto, per esempio, al modello dell'atomo.

Di pari passo con lo sviluppo della civilizzazione sono cambiati anche i metodi di preparazione del cibo. Così, un fuoco di legna con un mammut a rosolare (come si è detto nell'incipit di questa Parte) è stato sostituito dapprima dal focolare domestico, poi dalla cucina a legna, a carbone e a gas, dal fornello a petrolio, quindi dai fornelli elettrici e dal grill.

Col passare dei millenni è cambiato il rapporto tra il fuoco e il cibo, anche se dal punto di vista fisico il fenomeno di fondo è rimasto invariato: il trattamento termico avviene grazie al riscaldamento del cibo mediante termo-trasmissione diretta (anche con convezione) oppure attraverso irraggiamento di radiazione elettromagnetica nell'intervallo spettrale dell'infrarosso.

Un esempio del primo metodo ci è fornito dalla cottura di ravioli o di polpettine cosiddette "dietetiche" in una pentola speciale, dove questi prodotti vengono riscaldati dal vapore dovuto all'ebollizione di acqua sottostante. Nella cottura di una minestra, alla diretta termo-trasmissione dal fondo della pentola si aggiunge anche il meccanismo di trasporto convettivo del calore dagli strati inferiori verso quelli superiori, nel corso del loro mescolamento.

Un esempio diverso è invece quello in cui il riscaldamento del cibo avviene solo attraverso l'irraggiamento a infrarossi, come nel caso di una grigliata e della cottura di uno spiedino su carboni ardenti.

Il perfezionamento del “focolare” (così chiameremo la fonte di calore per la preparazione del cibo) ha trasformato anche le stesse ricette culinarie. Grazie a ciò sono comparse nuove possibilità di preparazione di piatti molto elaborati. Per amor di verità bisogna osservare che contemporaneamente alla comparsa di nuove ricette e con il trasformarsi del “focolare”, alcuni piatti, purtroppo, scompaiono dalla nostra tavola o perché cadono nel dimenticatoio o perché cedono il posto a surrogati. Per esempio, la vera pizza napoletana può essere preparata con un minimo di ingredienti in pochi minuti, ma solo in uno speciale forno a legna ben arroventato. Per questo le vere Pizzerie sono orgogliose del proprio forno a legna, dove tutto il processo si svolge davanti ai vostri occhi. Il pizzaiolo “modella” con virtuosismo la pizza, la inforna con la sua pala: un istante ed essa si gonfia appetitosamente davanti a voi, invitandovi a mangiarla senza indugio, accompagnata da un boccale di buona birra. Se vi trovate per caso in una pizzeria moderna, con una sala lussuosa, ricca di specchi ma senza il forno a legna, non lasciatevi irretire e andatevene: è meglio rimanere digiuni che dover mangiare il surrogato che vi propinerebbero in quel luogo.

Ci siamo, però, allontanati dal tema iniziale del discorso. Facciamo un passo indietro e torniamo per il momento non in cucina... ma ad una lezione di storia della metallurgia. Qui verrete a sapere che i metodi di fusione dei metalli sono cambiati ancora più drasticamente dei metodi di preparazione dei cibi. In particolare, dopo la scoperta di **Michael Faraday** della legge di induzione elettromagnetica è stato messo a punto il metodo di elettrofusione. In sintesi, il concetto è il seguente: un pezzo di metallo viene messo in un intenso campo magnetico rapidamente variabile nel tempo. Il metallo agisce da conduttore e la forza elettromotrice (FEM) d'induzione

$$\mathcal{E} = -\frac{d\Phi}{dt}, \quad (15.1)$$

(dove Φ è il flusso magnetico che attraversa il campione) determina la formazione di flussi di induzione, o come vengono anche chiamati, flussi di Foucault. Essi si dicono anche vorticosi perché le linee di corrente, seguendo le linee del campo di induzione, sono chiuse. La dissipazione della corrente di induzione è accompagnata dall'emissione di calore. Se la FEM di induzione

è sufficientemente grande (cioè sono elevate l'ampiezza e la frequenza di variazione del campo magnetico) il calore prodotto può risultare sufficiente per la fusione del metallo. I forni di elettrofusione sono largamente utilizzati per la fusione di acciai e nella metallurgia spaziale.

E così ci siamo avvicinati all'oggetto delle nostre considerazioni. Lasciamo da parte il forno per elettrofusione e gettiamo uno sguardo nell'angolo cottura di una cucina moderna. Spesso si trova un tipo di focolare del tutto diverso, che ricorda più il forno per elettrofusione che un attrezzo di cucina. Per la cottura o per il riscaldamento di cibo precotto o surgelato viene utilizzato l'irraggiamento elettromagnetico ad iperfrequenza. Già negli anni '60, quando i cosmonauti iniziarono a trattenersi in orbita extra-terrestre sempre più a lungo, l'alimentazione da tubetti divenne inadeguata. D'altronde portare con sé nello spazio un fornello a petrolio era impossibile, per svariati motivi. In primo luogo l'ossigeno, che partecipa direttamente alla combustione, nella navicella spaziale è ovviamente preziosissimo. In secondo luogo l'assenza di gravità può rendere difficile la preparazione di molte delle ricette di Elena Molochovez, autrice di un famoso libro di cucina "Consigli alle giovani massaie".

(Pensate a come l'assenza di gravità influisca sulla possibilità di preparare una normale zuppa su un fornello a petrolio. Quali altre difficoltà vedete per la culinaria spaziale?).

La soluzione fu trovata nell'impiego per la cottura di un analogo del forno ad elettrofusione. Quasi tutti i prodotti utilizzati dall'uomo sono composti in notevole misura di acqua, che essendo un elettrolita, in una certa misura conduce corrente. Inoltre ogni molecola d'acqua è caratterizzata dal possedere un momento di dipolo elettrico, cioè un centro di cariche positive che cade a una certa distanza dal centro delle cariche negative. Per questo, sottoponendo per esempio un pezzo di carne a un campo elettromagnetico variabile, possono essere indotte correnti di dissipazione e insieme fatti oscillare i dipoli elettrici alla stessa frequenza della radiazione elettromagnetica. Questi fenomeni portano a un insieme di effetti dissipativi che comportano trasformazione di energia elettro-meccanica in calore, con aumento della temperatura e, di conseguenza, al trattamento termico di alcuni prodotti.

Un esempio di campo elettromagnetico variabile nel tempo e nello spazio è l'onda elettromagnetica. Ovviamente non tutte le onde elettromagnetiche riescono ad arrostitire una bistecca. Infatti, anche illuminassimo a lungo la

carne con la lampadina di una torcia tascabile, non riusciremo a gustarla. Per questo, evidentemente, devono essere adottate delle soluzioni particolari. Prima di tutto il campo deve essere sufficientemente intenso. Ad esempio potrebbe essere adeguato il campo di un radar. (Si dice che uccelli che capitano nella zona d'azione di radar molto potenti cadano fulminati, e non bruciati, come lessi, quasi come nel caso del barone di Munchausen). Per sicurezza tale campo andrebbe ovviamente confinato da qualche parte e limitato nel volume di localizzazione.

Come "scatole" di contenimento per le onde radar o micro-onde sono necessarie le cavità risonanti. Per il contenimento delle onde sonore è sufficiente ricorrere ad una "scatola" di legno: il corpo del violino è una tipica cassa di risonanza. In esso dopo l'eccitazione esterna possono sussistere relativamente a lungo delle onde sonore. In questo caso la fonte delle onde è rappresentata dall'archetto e dalle corde.

In modo analogo si può confinare il campo elettromagnetico in una scatola di metallo. Nella sua dimensione deve essere accomodato un numero intero di semionde di radiazione. Eccitando in questa cassa oscillazioni elettromagnetiche della necessaria frequenza, otteniamo un risuonatore con un'onda elettromagnetica i cui nodi (i punti dell'onda in cui l'ampiezza delle oscillazioni è uguale a zero) sono localizzati sulle sue pareti. Il forno a microonde o, come viene anche chiamato, il forno a iperfrequenza, costituisce proprio questa speciale cassa di risonanza, accoppiata con radiatori in miniatura tipo 'radar'. Poiché la dimensione tipica di questo strumento è di alcune decine di centimetri, si deduce immediatamente la lunghezza massima delle onde utilizzabile. Questa valutazione è facilmente verificabile, controllando la parete posteriore del forno: vedrete che la frequenza standard è $\nu = 2450 \text{ MHz}$, che corrisponde alla lunghezza dell'onda di circa $\lambda = c/\nu \approx 12 \text{ cm}$.

Continuiamo lo studio delle proprietà del forno a microonde con un esperimento che noi stessi abbiamo eseguito. Prendiamo dal freezer un grosso pezzo di carne congelato, su un piatto speciale (ne parleremo più avanti) sistemiamolo nel forno e diamo il via. Inizialmente sembra che, ad eccezione del ronzio cadenzato dei radiatori, non accada nulla. Tuttavia, col passare del tempo, attraverso il vetro del portello possiamo notare come la carne si rosola e dopo qualche decina di minuti sembri completamente cotta. Tagliandola, dopo averla estratta dal forno, possiamo osservare che all'interno del pezzo di carne vi sono parti non solo crude, ma addirittura non scongelate. Come spiegare questo fenomeno?

La più semplice ipotesi può basarsi sulla non uniforme distribuzione del campo elettromagnetico all'interno del forno. Effettivamente la dimensione della camera è dell'ordine di 30 cm, mentre la lunghezza dell'onda, come abbiamo potuto verificare prima, è di circa 12 cm. Considerando che sulle pareti devono trovarsi nodi d'onda, osserviamo che nel volume si hanno almeno tre nodi, dove l'intensità del campo è nulla. Poiché l'onda è stazionaria, la posizione di questi nodi è fissa nello spazio e il pezzo di carne cruda potrebbe trovarsi in uno di questi. Tuttavia nei moderni forni, questa difficoltà viene risolta con una lenta rotazione della piastra su cui poggia il piatto con il cibo, il che comporta un effetto di media dell'azione del campo.

Così, sembrerebbe che ponendo un intero maialino in un forno a micro-onde di misure adeguate e con il piatto ruotante, si possa dopo un tempo relativamente breve gustare "la porchetta".

In realtà, nel tentativo di preparare la porchetta alle micro-onde incontriamo un altro ostacolo, ben più importante e insuperabile. Il suo nome è *effetto pelle*. Così viene chiamata la proprietà delle correnti ad alta frequenza di scorrere solo sulla superficie di un conduttore. Poiché la frequenza dell'onda elettromagnetica nel forno è particolarmente alta, l'effetto pelle può giocare un ruolo sostanziale. Il campo elettromagnetico, attenuandosi nel profondo del grosso pezzo di carne, non riesce a portare nel centro sufficiente energia per arrostarlo.

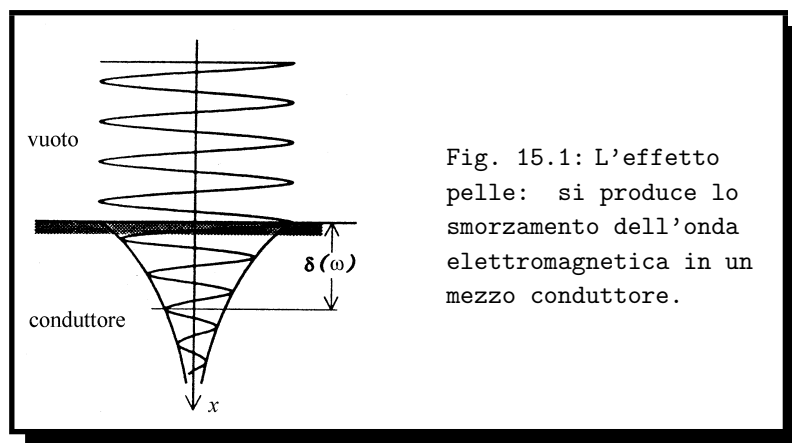
Per verificare questa ipotesi, consideriamo più dettagliatamente il fenomeno dell'effetto pelle e cerchiamo di valutare la profondità effettiva di penetrazione del campo e di spiegare come essa dipenda dalla frequenza e dalle proprietà del conduttore. Questo quesito può essere facilmente risolto con le equazioni differenziali del campo elettromagnetico (equazioni di **James Clerk Maxwell**), ma qui ci limiteremo a valutazioni qualitative.

Iniziamo con l'impostazione del problema.

Supponiamo che l'onda elettromagnetica con frequenza ω incida normalmente sul limite piatto del conduttore. All'interno l'onda inizia a indebolirsi in ampiezza, poiché l'energia si trasforma in perdite per effetto **Joule**. Si può mostrare che l'indebolimento dell'onda avviene secondo una legge, detta legge esponenziale, assai diffusa nei fenomeni naturali (si ricordi, ad esempio, la legge della disintegrazione radioattiva):

$$E(x) = E(0) e^{-x/\delta(\omega)}, \quad (15.2)$$

ove $E(0)$ è l'ampiezza del campo elettrico alla superficie, $E(x)$ è l'ampiezza alla profondità x , $\delta(\omega)$ è la profondità effettiva di penetrazione del campo (lunghezza in corrispondenza alla quale il campo si riduce al valore $E(0)/e$, con $e = 2.71 \dots$ la base dei logaritmi naturali).



Occorre ottenere un'espressione esplicita per $\delta(\omega)$. Utilizziamo il metodo delle relazioni di scala (noto ai fisici con il termine *scaling* e al quale abbiamo già fatto altre volte ricorso). È intuitivo che la profondità di penetrazione della onda deve dipendere dalla frequenza (dal momento che la corrente costante ($\omega = 0$) scorre su tutta la sezione del conduttore) e in particolare deve diminuire con l'aumento della frequenza. È naturale supporre (come di solito si fa con il metodo dello scaling) che la profondità di penetrazione dipenda dalla frequenza angolare ω in forma di legge di potenza

$$\delta(\omega) \propto \omega^\alpha,$$

e dobbiamo attenderci che α sia negativo.

È altresì evidente che la profondità di penetrazione deve dipendere dalle proprietà di conduzione della sostanza, le quali sono descritte dalla resistenza specifica ρ (o la conducibilità specifica $\sigma = 1/\rho$). L'energia dell'onda elettromagnetica nel caso dell'effetto pelle si trasforma in calore, mentre la potenza delle perdite Joule, per unità di volume, ha la forma

$$p = j E = \frac{E^2}{\rho} = \sigma E^2,$$

dove E è l'intensità del campo in un dato punto del conduttore. (Si ricordi che la legge di **Ohm** per la densità di corrente si scrive $j = \sigma E$). Quanto più efficacemente si dissipa l'energia dell'onda, tanto più rapidamente deve diminuire la sua ampiezza. Per questo la profondità di penetrazione deve dipendere dalla conducibilità specifica del mezzo:

$$\delta \propto \sigma^\beta,$$

e occorre attendersi che β , come α , risulti negativo.

Infine, ricordiamo che le grandezze elettromagnetiche nel sistema internazionale contengono una costante magnetica dimensionale $\mu_0 = 4\pi \cdot 10^{-7} H/m$. Questa costante entra nella formula per l'induzione magnetica creata dalla corrente nello spazio circostante, analogamente a come la costante ε_0 entra nella formula per il campo elettrico della carica puntiforme. Supponendo che la lunghezza di penetrazione sia la combinazione delle sole tre grandezze elencate^a, si ha

$$\delta \propto \omega^\alpha \sigma^\beta \mu_0^\gamma, \quad (15.3)$$

dove α , β e γ sono gli indici di scala. Confrontiamo ora le dimensioni della parte destra e sinistra di questa equazione. Per questo trascriviamo le dimensioni di tutte le grandezze che entrano nell'espressione (15.3):

$$[\delta] = m, \quad [\omega] = \text{sec}^{-1}, \quad [\sigma] = \Omega^{-1} \cdot m^{-1}, \quad [\mu_0] = H \cdot m^{-1}.$$

Notiamo che poichè *Henry* è l'unità di induzione, grazie alla relazione

$$\mathcal{E} = -L \frac{\Delta I}{\Delta t}$$

si può scrivere

$$H = V \cdot \text{sec}/A = \Omega \cdot \text{sec},$$

e

$$[\mu_0] = \Omega \cdot \text{sec} \cdot m^{-1}.$$

^aQuesta ipotesi equivale a trascurare la cosiddetta corrente di spostamento, con ciò assumendo che il campo magnetico sia formato solo da correnti reali e non tenendo conto del campo elettrico alternato. In caso contrario interverrebbe anche la costante dielettrica del vuoto. Alle frequenze considerate questa ipotesi si realizza con buona approssimazione.

Eguagliamo le dimensioni della parte destra e sinistra dell'equazione (15.3):

$$m = (\text{sec})^{-\alpha} (\Omega \cdot m)^{-\beta} (\Omega \cdot \text{sec} \cdot m^{-1})^{\gamma},$$

o

$$m^1 = (\Omega)^{\gamma-\beta} \cdot \text{sec}^{\gamma-\alpha} \cdot m^{-\gamma-\beta}.$$

Otteniamo così tre equazioni

$$\begin{cases} \gamma - \beta &= 0, \\ \gamma - \alpha &= 0, \\ -\gamma - \beta &= 1, \end{cases}$$

risolvendo le quali si deduce $\alpha = \beta = \gamma = -1/2$. L'equazione (15.3) prende pertanto la forma

$$\delta \propto \sqrt{\frac{1}{\mu_0 \omega \sigma}}.$$

Notiamo che la dipendenza della lunghezza di penetrazione da α e β corrisponde alle considerazioni fisiche fatte precedentemente: i coefficienti di scala sono risultati negativi. Una valutazione completa, basata sulla soluzione delle equazioni di Maxwell, conduce ad una espressione strettamente simile, con la sola differenza nel fattore numerico $\sqrt{2}$:

$$\delta = \sqrt{\frac{2}{\mu_0 \omega \sigma}}. \quad (15.4)$$

Valutiamo quindi la profondità di penetrazione alla frequenza che ci interessa ($\nu = \omega / 2\pi = 2.45 \cdot 10^9 \text{ Hz}$) per il rame (esempio di elevata conducibilità, $\sigma_c \approx 6 \cdot 10^7 \Omega^{-1} \cdot m^{-1}$) e per il muscolo ($\sigma_m \approx 2.5 \Omega^{-1} \cdot m^{-1}$). Dalla Eq. (15.4) abbiamo

$$\delta_m \approx 1 \text{ cm} \quad \text{and} \quad \delta_c \approx 10^{-3} \text{ cm}.$$

È evidente che per un buon conduttore l'effetto pelle è molto marcato: proprio per questo la resistenza dei cavi ad alte frequenze è notevolmente superiore che a basse frequenze (la corrente si muove nel sottile strato superficiale, vale a dire è come se diminuisse la sezione del cavo). Tuttavia anche per un cattivo conduttore, come è il caso del tessuto muscolare, l'effetto rimane ugualmente significativo. Anche un pezzo di carne con dimensioni

caratteristiche di 10 cm è già da considerarsi di dimensioni troppo elevate per una buona cottura: il valore del campo, e con esso la potenza, si riducono di molte volte all'interno del pezzo, così che la termo-conduzione rimane praticamente l'unica fonte di riscaldamento. Per ciò che riguarda il nostro ipotetico arrosto della porchetta, l'effetto pelle produrrebbe così un risultato assai deludente: potrà arrostarsi soltanto lo strato esterno, mentre la gran parte della carne rimarrà piuttosto mal cotta!

Non è questa la peggiore sciagura! La cuoca avveduta potrebbe porre nel forno pezzi non troppo grossi e la carne cuocerebbe uniformemente. In particolare per una fettina in forma di ciambella l'effetto pelle non è drammatico. È possibile trovare speciali impieghi di questo effetto. Per esempio, nel tempo libero pensate ad una ricetta per preparare in un forno a microonde un piatto impensabile come... un gelato avvolto da una sfoglia rosolata.

Il forno a microonde ha i suoi pregi e i suoi difetti. Da un lato, esso offre prospettive di preparazione di cibi inusuali, contribuisce a non deteriorare le vitamine, permette di creare cibi dietetici; dall'altro non permette di cuocere un semplice uovo *à la coque*. Pensiamo di mettere un uovo nel forno. Dopo l'accensione del regolatore di potenza il contenuto liquido all'interno dell'uovo si riscalda molto rapidamente e inizia l'emissione di gas. Il guscio non permette ai gas di filtrare. La pressione aumenta e... determina l'esplosione! La cella del forno riceve gli spruzzi della vostra colazione e voi, maledicendo la vostra sconsideratezza, dovrete pulire a lungo il vostro elettrodomestico.

Fig. 15.2:
L'esplosione dell'uovo
nel forno a microonde.



Ancora più importante nei riguardi dei limiti del forno a microonde è il fatto che non si possono “creare” gli arrosti. Un buon arrosto, infatti, si basa su un processo di accorpamento delle proteine che è successivo alla semplice cottura: occorre produrre una crosta superficiale che corrisponde

alla ristrutturazione delle lunghe catene proteiche attraverso legami che le interconnettono (si veda il capitolo 17). Per questo servono temperature più elevate dei 100 gradi che si possono raggiungere mediante il riscaldamento dell'acqua, elemento dominante nella maggior parte delle sostanze alimentari (se si desidera il caratteristico arrosto occorre provvedere alla creazione della crosta attraverso, per esempio, l'azione di irraggiamento con il grill).

Occorre ricordare che la potenza massima del forno è abbastanza elevata: di solito dell'ordine di un chilowatt. Mediante il regolatore di potenza potete limitarla, ma deve esserci comunque un oggetto sul quale essa possa scaricarsi. Per questo è categoricamente raccomandato di non accendere il forno privo di alimenti all'interno. Infatti, in tal caso il campo elettromagnetico riversa la sua energia sui radiatori-induttori, rischiando di distruggerli. Una situazione analoga, dal punto di vista elettrotecnico, al cortocircuito di una batteria. L'assenza di resistenza del carico (nel nostro caso questo ruolo è rappresentato per esempio dal pollo da cuocere) determina un'elevata corrente di cortocircuito, vale a dire la distruzione della batteria.

Un altro pericolo si nasconde nella scelta del recipiente utilizzato per la preparazione. Naturalmente si può utilizzare un recipiente *ad hoc* del tipo "Rutex". Si può anche ragionare sui requisiti richiesti a questo recipiente, valutare gli aspetti fisici coinvolti e allo stesso tempo risparmiare, utilizzando una vecchia pentola di creta. Il requisito principale per il recipiente è la sua "trasparenza" alla radiazione a microonde. Anche per una frequenza così elevata il materiale deve rimanere dielettrico (non molti tipi di ceramica o di vetro soddisfano questo requisito; si noti che le proprietà ottiche ed elettriche delle sostanze possono dipendere notevolmente dalla frequenza del campo elettromagnetico). Quindi in nessun caso si può mettere nella cella del forno un recipiente metallico, avvolgere un cibo pronto in un foglio di stagnola o anche utilizzare un piatto con rifiniture in oro. In un batter d'occhio il pacifico forno da cucina si trasformerà nel suo parente che vomita fuoco in un reparto metallurgico. Molte ciotole di argilla e piatti di ceramica sono invece adatti all'uso nel forno a microonde.

Per verificare se un dato recipiente è adatto, è sufficiente metterlo insieme ad un bicchiere d'acqua (pensate perché questo sia convenientemente incluso) nella cella e accendere per un paio di minuti il forno. Se la ciotola in esame è rimasta fredda, ha superato la prova.

Capitolo 16

Ab ovo.

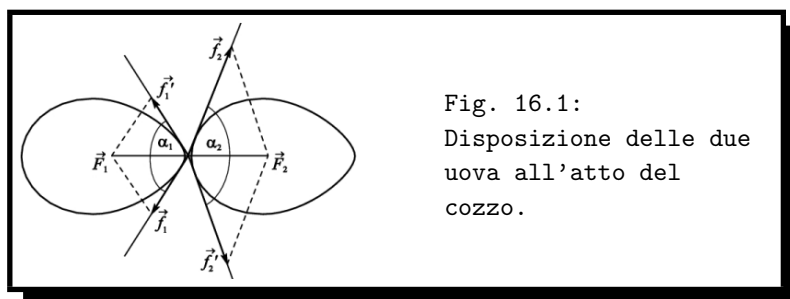
Questa espressione latina, come è noto significa letteralmente “proprio dall’inizio”. In realtà noi la useremo piuttosto nel significato di “discorso sull’uovo”, nel senso che discuteremo un poco di fisica con riferimento a questo prodotto, così comune e apprezzato in ogni cucina.

Come ricorderete dalle letture giovanili dei “Viaggi di Gulliver” una guerra senza fine tra gli imperi di Lilliput e di Blefuscu aveva preso l’avvio dopo l’editto dell’imperatore del primo, che ordinava come in ogni circostanza si dovessero rompere le uova dalla parte che terminava con la sezione più piccola. Gulliver riteneva (in accordo alla dottrina del profeta Lustrog) che il decidere da quale parte rompere le uova si trattasse in realtà di un affare privato. Benché sostanzialmente gli autori ritengano questa convinzione corretta, non è loro intenzione prendere parte alla disputa. Piuttosto, interessa loro stabilire da quale parte dell’uovo sia più “facile”, ovvero occorra minore energia, ottenere la rottura delle uova. Questo obiettivo ci permetterà qualche osservazione di fisica e darà al lettore dei suggerimenti sulla migliore strategia per riuscire vincitore da eventuali dispute che possono sorgere alla colazione mattutina o al termine del pranzo Pasquale.

Quale è la tattica più appropriata? Quale uovo usare come elemento dirompente, il grande o il piccolo? Da quale parte procedere al cozzo, dalla parte più piccola o dalla più grande? Nei riguardi di queste strategie “di attacco” è convinzione piuttosto diffusa che sia vantaggioso il lato dell’uovo “attaccante”. In realtà, se entrambe le uova muovono uniformemente non sussiste alcuna differenza nei riguardi dell’elemento che colpisce rispetto a quello che subisce il cozzo. Senza procedere a tentativi sperimentali di controllo e rompere così diverse uova per nulla, si osservi che questa prima

affermazione discende dal principio di relatività Galileiana. Nel sistema di riferimento dell'uovo "aggredito" le cose sono esattamente le stesse del sistema di riferimento dell'uovo che viene usato per l'attacco e quindi possiamo risparmiarci dispendiosi esperimenti di verifica.

Piuttosto, consideriamo più da vicino il processo di collisione, riferendoci a due uova di eguali dimensioni e forma e con la stessa resistenza dei due gusci (coefficiente di rottura σ_b). La collisione avvenga lungo l'asse comune delle due uova, quello di sinistra colpendo con la parte finale piccola, quello di destra venga invece colpito nella sua parte terminale più ampia (Fig. 16.1).



In accordo alla terza legge di Newton le forze $\mathbf{F}_1$ e $\mathbf{F}_2$ sono eguali e opposte. Sostituiamo le due con la somma di forze $\mathbf{f}_1$, $\mathbf{f}_1'$ e $\mathbf{f}_2$, $\mathbf{f}_2'$ dirette lungo le tangenti alla superfici collidenti. In conseguenza della elasticità dei gusci la collisione avviene in realtà su una superficie piccola ma finita. Lo stress f dipende dall'angolo tra le tangenti: $f = F / (2 \cos \alpha/2)$. Di conseguenza le forze di rottura agenti sui gusci sono tanto più grandi quanto maggiore è l'angolo. L'angolo, a sua volta, è definito dal raggio di curvatura della superficie: la figura mostra chiaramente come sia in realtà vantaggioso colpire l'altro uovo con la parte più piccola del vostro. Una seconda ragione in favore di questa scelta di strategia è che la bolla d'aria contenuta nella parte più ampia dell'uovo indebolisce ulteriormente la superficie corrispondente.

Si osservi che la analisi che abbiamo condotto suggerisce comunque un piccolo trucco pratico. Se il vostro contendente nella gara di rottura delle uova rivolge la parte più piccola del suo uovo verso il vostro, ebbene voi non avete che da colpire su una parte leggermente laterale. Infatti in tal caso il raggio di curvatura del guscio è minore e quindi lo sforzo associato

alla deformazione prodotta dall'urto aumenta, a vostro vantaggio^a.

In negozi di souvenir e di giocattoli si può trovare un gioco chiamato “*Flip-over top*”. La sua forma è quella di una palla tronca, con un cilindro appuntito che sale dal centro del lato piatto e termina in alto un poco arrotondato. Quando questa tozza trottola è fatta ruotare sufficientemente veloce attraverso la parte cilindrica, la boccia inferiore si comporta in modo inaspettato: dopo un certo tempo l'oggetto si capovolge e continua la sua rotazione sulla coda cilindrica. Evidentemente l'energia potenziale della parte superiore aumenta a seguito di questo capovolgimento, cioè il sistema “guadagna” energia gravitazionale anziché cederla come avviene usualmente nell'abbassarsi del baricentro di un oggetto, a spese della energia di rotazione. Questo comportamento peculiare è stato per la prima volta spiegato da **William Thomson** (anche noto altrimenti come **Lord Kelvin**) e da allora questo giocattolo è anche conosciuto come la trottola di Thomson.

Ebbene, l'uovo si può comportare in maniera simile alla trottola di Thomson! Prendete un uovo ben sodo e tenendolo per il suo asse maggiore fatelo ruotare il più velocemente possibile su una superficie piatta e liscia, per esempio una lasta di marmo. Dopo alcune rivoluzioni vedrete che l'uovo può “rizzarsi” ad assumere una posizione verticale, con il suo asse maggiore perpendicolare alla superficie sulla quale poggia. Solo dopo che gli attriti avranno quasi arrestato la rotazione, l'uovo comincerà ad abbassare il suo baricentro ruotando sempre più ampiamente e lentamente, sino ad arrestarsi poggiando su un lato, con l'asse maggiore orizzontale. Se l'uovo non è bollito abbastanza a lungo e quindi il suo interno non è totalmente solido questo giochetto non vi riuscirà in quanto gli attriti viscosi tra strati di “chiara” e tra la parte liquida e le pareti del guscio rallenteranno la rotazione e l'uovo perderà il momento angolare che gli avete fatto assumere

^aDal punto di vista della fisica le uova collidono allo stesso modo delle superfici metalliche, per esempio di due sottomarini. D'altro canto gli sforzi sulle superfici sono distribuiti allo stesso modo delle tensioni sulle pareti di una bolla di sapone. Questa analogia rende possibile collegare la pressione critica P_c al raggio di curvatura R del guscio. Si può, infatti, fare diretto ricorso alla formula di Laplace che abbiamo incontrato alla sezione bolle e gocce, vale a dire

$$P_c = \frac{2\sigma_b}{R},$$

ove lo stress critico σ_b caratterizza la resistenza del guscio. Pertanto la pressione necessaria per rompere l'uovo è inversamente proporzionale al raggio della parte di guscio che urta.

con la primitiva rotazione. Potete usare questa proprietà per accertarvi se l'uovo che desiderate ben cotto sia stato sufficientemente a lungo in acqua bollente.

Passiamo ora a discutere cosa accade all'uovo durante il trattamento termico. Cosa dobbiamo fare se vogliamo che un uovo non si rompa o non esploda (come farebbe nel forno a microonde)? Quando dovremo arrestare la bollitura se desideriamo un uovo sodo semi-indurito, noto come “à la coque”? Per quale ragione una esperta cuoca fa bollire le uova in acqua salata?

Non è usuale trovare nei manuali di cucina risposta a queste domande.

Come vedremo a proposito della cosiddetta “fisica culinaria” (capitolo 20) la bollitura di un uovo (così come altri trattamenti termici per la cottura) consiste essenzialmente nel provocare il processo di modificazione delle proteine. Ad elevata temperatura queste complesse molecole organiche si rompono in frammenti, cambiano forma e disposizione strutturale nello spazio. Enzimi o altre sostanze chimiche possono ugualmente produrre o favorire questo processo di modificazione delle proteine. Poiché le proteine del bianco d'uovo e del tuorlo sono diverse, esse subiscono il processo di denaturazione a temperature diverse. È questa apparentemente insignificante differenza che ci consente di preparare uova a “bassa bollitura”. Il bianco d'uovo inizia a modificare la sua natura a 60 gradi centigradi mentre per il tuorlo ciò comincia ad aver luogo a 63 – 65 gradi. Non è il caso di precisare meglio queste temperature in quanto si tratta di un processo graduale, che non ha luogo bruscamente a una temperatura del tutto definita. Inoltre si hanno piccole differenze nella temperatura alla quale iniziano i processi di denaturazione delle proteine, a seconda del contenuto in sale e dell' “età” dell'uovo o per altri fattori.

Dal punto di vista fisico, quando l'uovo viene immerso in acqua calda inizia un flusso di calore dal guscio verso il centro. Le equazioni per la propagazione del calore sono state da tempo sviluppate nella fisica-matematica. Pertanto allo scopo di stabilire quanto tempo impiegherà l'uovo a “bollire” noi abbiamo solo bisogno di conoscere le conducibilità termiche del bianco e del tuorlo e la forma geometrica dell'uovo stesso.

La valutazione numerica della diffusione del calore per una forma ellissoidale è stata fatta per la prima volta dal fisico inglese **Peter Barham** e nel suo libro “The Science of Cooking” si trova la formula con la quale il tempo di cottura t (in minuti) viene legato alla temperatura di ebollizione dell'acqua T_{eb} , all'asse minore dell'uovo d (in cm), alla temperatura iniziale

del tuorlo (T_0) e a quella desiderata (T_f) dello stesso.

Tale formula è

$$t = 0.15d^2 \log 2 \frac{(T_{eb} - T_0)}{(T_{eb} - T_f)}.$$

(log sta per logaritmo naturale).

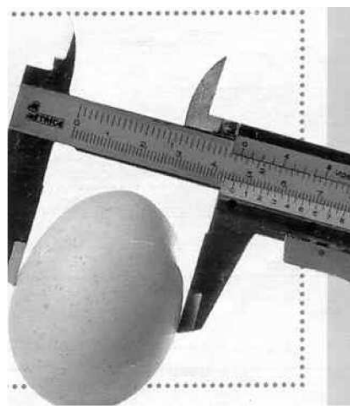
In condizioni standard, vale a dire pressione atmosferica a livello del mare, T_{eb} è ovviamente 100 gradi centigradi e sulla base di questa formula per un uovo “à la coque” non del tutto sodo (diciamo $T_f = 63^\circ C$), per un valore tipico di $d = 4\text{ cm}$ e uovo appena rimosso dal frigorifero ($T_0 = 5^\circ C$) valutiamo $t = 3\text{ min } 56\text{ sec}$.

Se l'uovo ha un diametro di 6 cm il tempo richiesto è più che doppio: $t = 8\text{ min } 50\text{ sec}$.

La formula vi insegna anche di quanto dovete aumentare il tempo di cottura se salite in montagna. Per esempio, all'altitudine di 5000 metri poiché l'acqua bolle a $88^\circ C$ vi saranno necessari, per l'uovo di 4 cm , ben $4\text{ min } 32\text{ sec}$.

La prossima volta che vorrete gustare un uovo “à la coque” di perfetto gradimento munitevi di un cronometro e di un calibro (come mostrato nella figura 16.2), fate semplici conti sulla base della formula di Barham, sarete del tutto soddisfatti e.. ringrazierete la fisica!

Fig. 16.2: La
misurazione dell'asse
minore dell'uovo.



Passiamo ora a discutere il ruolo del sale nella bollitura dell'uovo in acqua. Occorre tenere presente che la densità di un uovo fresco è leggermente

superiore a quella dell'acqua. Quindi in acqua "dolce" l'uovo si porta al fondo della pentola di bollitura. All'ebollizione forti correnti ascensionali scuotono su e giù l'uovo, che può picchiare violentemente contro il fondo o contro le pareti della pentola e il guscio si può rompere. Il bianco d'uovo finirà per fuoriuscire dalle fessure del guscio rotto e la vostra colazione è compromessa^b.

Se invece viene aggiunto sale all'acqua di bollitura (basta mezzo cucchiaino) la densità della soluzione aumenta e l'uovo non si adagia più sul fondo. Inoltre, in caso di rottura del guscio, il sale favorisce la denaturazione delle proteine del bianco, che coagula e blocca la fuoriuscita!

Ora l'uovo per la vostra colazione è ben sodo. Provate a tenerlo nelle vostre mani: dapprima non lo sentirete troppo caldo. Tuttavia, mentre il guscio si va asciugando la temperatura che voi percepirete aumenterà e dopo poco non sarete in grado di tenerlo ancora in una mano. Perché? Rispondete voi stessi.

Dopo aver dato risposta al quesito, vorrete infine staccare il guscio: vi accorgerete che il guscio è quasi attaccato all'uovo e nell'asportazione rimuove frammenti di bianco. Per evitare questo, immergete l'uovo in acqua fredda e beneficiando dei diversi coefficienti di espansione termica del bianco e del guscio, otterrete facilmente il distacco di quest'ultimo.

Abbiamo così potuto vedere come le leggi della meccanica, dell'idrodinamica e della fisica molecolare intervengano nel processo di bollitura di un uovo. Possiamo anche far intervenire i fenomeni elettrici, in realtà. Abbiamo già detto cosa accade all'uovo in un forno a micro-onde a seguito del rapido riscaldamento indotto dalle onde elettromagnetiche.

Senza dover far esplodere l'uovo possiamo comunque tenere conto che il guscio dell'uovo è un dielettrico e ripetendo le osservazioni di **Faraday** dare dimostrazione della elettricità statica.

Prendiamo un uovo fresco e pratichiamo due piccoli fori alle estremità, uno un poco più grande dell'altro. Quindi con un lungo ago sbattete il contenuto all'interno dell'uovo sino a che avrete un liquido omogeneo. Soffiate attraverso il foro più piccolo e facendo fuoriuscire il liquido, svuotate il guscio, lavatelo e lasciatelo asciugare. A questo punto sarete pronti all'esperimento. Con un sbarra di ebanite che avrete elettrizzata strofinan-

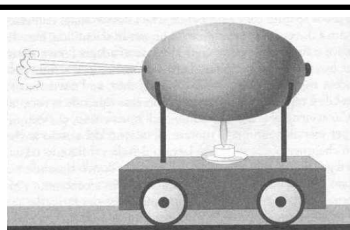
^bSe volete mostrarvi veramente cuochi raffinati, ad evitare anche il rischio di rottura del guscio a seguito della espansione termica dell'aria contenuta all'interno, potete praticare un minuscolo forellino con un spillo, meglio nella parte più larga dell'uovo.

dola con un certo vigore su lana (o anche con il pettine se lo avete appena usato su capelli secchi) tentate di avvicinare il guscio: vedrete il guscio dell'uovo seguire il pettine e potrete portarlo "a spasso" sulla tavola come un cane al guinzaglio.

Infine potrete utilizzare il guscio svuotato per realizzare un piccolo esperimento per vostro figlio o per il vostro nipotino, illustrandogli come funzionano i motori a reazione. Chiudete il foro più piccolo con del *chewing gum* e riempite il guscio a metà con acqua (troverete forse più facile invertire l'ordine di queste due operazioni e chiedetevi il perché). Sistemate il guscio su un piccolo carrettino con ruote, come illustrato in Figura 16.3. Al di sotto del guscio incendiate un batuffolo di cotone impregnato di alcool e attendete. Giunta l'acqua all'ebollizione il vapore inizierà a fuoriuscire dal forellino. In accordo al principio di azione e reazione utilizzato nei *jet*, il veicolo si metterà in moto in direzione opposta al getto di vapore.

Se non avete il carrettino potete creare una piccola nave a reazione sistemando opportunamente guscio e batuffolo sopra una lastra di sughero che lascerete galleggiare su una pozza d'acqua o utilizzando un lavandino quasi riempito di acqua.

Fig. 16.3: Il
dispositivo a reazione
mediante l'uovo.



Ancora, utilizzando il principio di **Archimede**, ecco come riconoscere un uovo fresco da un uovo ormai deteriorato. Come illustrato nella figura Fig. 16.4, introducete l'uovo in esame in un bicchiere d'acqua: se galleggia non è più il caso di utilizzarlo. Solo le uova che discendono decisamente sul fondo sono fresche, mentre un uovo di una settimana circa lo vedrete disporsi quasi verticalmente. La ragione di questo diverso comportamento sta nel fatto che mentre la spinta di Archimede è ovviamente indipendente dall'invecchiamento, con il trascorrere del tempo il peso dell'uovo diminuisce. Infatti l'invecchiamento dell'uovo determina sviluppo di idrogeno solforato che può traforare il guscio, passando attraverso minuscoli pori.

Pertanto la massa dell'uovo cambia mentre il volume rimane lo stesso e così la spinta di Archimede, legata al volume di acqua spostata, può superare la forza di gravità legata alla massa.

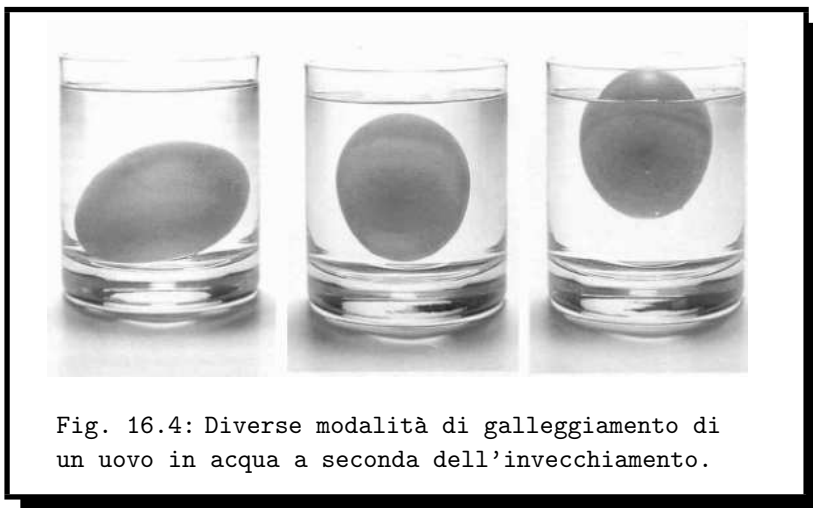


Fig. 16.4: Diverse modalità di galleggiamento di un uovo in acqua a seconda dell'invecchiamento.

Se ritenete di non potere agevolmente selezionare in questo modo le uova fresche da quelle invecchiate quando fate spesa al supermercato, potete provare un altro metodo, legato alle proprietà ottiche. Osservate attraverso l'uovo la luce di una intensa lampada: se la luce passa liberamente potete ritenere l'uovo fresco. Infatti la presenza di idrogeno solforato e di altri gas lo renderebbe opaco, in conseguenza di un assorbimento abbastanza marcato della luce visibile. Questo principio era alla base dei rudimentali ovoscopi in uso tempo fa in molte rivendite!

Capitolo 17

Il tacchino nel giorno del ringraziamento.

Come è noto negli Stati Uniti, il giorno del ringraziamento (Thanksgiving day) cade nel quarto Giovedì del mese di Novembre, alla fine della stagione del raccolto ed è inteso come celebrazione religiosa, sino dai tempi dei pellegrini e dei pionieri. Risale a **Lincoln**, sin dal 1863, l'ufficializzazione della festa del ringraziamento e la tradizione vuole che il piatto principale del pranzo sia un tacchino arrosto. Infatti colloquialmente si considera il giorno del ringraziamento anche come giorno del tacchino (Turkey Day). La scelta del tacchino in luogo, per esempio, del pollo presumibilmente fu originariamente legata alla più elevata quantità di carne fornita dal primo. Nella maggior parte dei casi i tacchini sono cotti con ripieno, composto in genere di frammenti di pane, frutti, spezie e aromi. Secondo attendibili statistiche del Dipartimento dell'Agricoltura, nel 2003 ben 45 Milioni di tacchini sono stati sacrificati negli Stati Uniti per allestire il pranzo nel giorno del ringraziamento. Non è dato sapere se anche i tacchini considerano quel Giovedì come un giorno di ringraziamento!

La massa di un tacchino può essere anche pari a 13 Kilogrammi e di conseguenza il cucinarlo adeguatamente non è affatto banale, essendo necessario ottenere un buon arrosto di tutto l'animale senza che peraltro le parti esterne vengano bruciacchiate. Il tempo di permanenza nel forno viene spesso riportato nell'avvolgimento del tacchino congelato che si acquista al supermercato. In questo capitolo noi desideriamo discutere scientificamente la stima del tempo necessario per una ottima cottura del tacchino. Il problema non è semplice, in quanto il tempo richiesto non è una funzione lineare della massa dell'animale. Ne consegue che anche un ottimo cuoco, abituato a cucinare adeguatamente un tacchino di pochi chili per una famiglia non

numerosa, senza particolari istruzioni rischia di trovarsi in serie difficoltà se deve cucinare un grosso tacchino (come quello mostrato in Fig. 17.1), del peso di dodici Kilogrammi, per una compagnia più ampia di persone. La nostra stima del tempo di cottura ottimale τ_{rst} in dipendenza dalla massa M potrà aiutarlo.



Fig. 17.1: Il tacchino del giorno di ringraziamento.

Prima di procedere alla stima precisiamo cosa si intenda per cottura della carne in termini scientifici. Il complesso muscolare consiste essenzialmente di diversi tipi di lunghe molecole organiche note come proteine. Quando la carne è cruda tali lunghe catene sono avvolte e arrotolate (Fig.17.2).



Fig. 17.2: Cottura della carne.

Nel corso del riscaldamento le catene si distendono e dopo un certo tempo di cottura esse compattano a formare una sorta di tappeto. Questo processo è noto come denaturazione proteica, alla quale abbiamo già fatto

cenno al capitolo precedente e che riprenderemo a proposito della cucina molecolare. La denaturazione proteica ha luogo a temperature relativamente basse, nell'intervallo di temperatura $55 - 80^\circ C$, questa variabilità essendo legata alla parziale sovrapposizione degli effetti termici relativi alla denaturazione delle varie proteine presenti nel muscolo, le quali hanno differente stabilità termica. Del resto, è a tutti noto che se si vuole gustare il lesso con la mostarda di Cremona per ottenere un totale compattamento delle proteine si deve operare con una piuttosto lunga bollitura a $100^\circ C$.

Tuttavia l'arrosto differisce in maniera sostanziale dalla carne lessa e per ottenere un arrosto gradevole occorre un trattamento termico a temperature considerevolmente più alte. Si deve infatti ottenere la reazione tra proteine e carboidrati nota come *reazione di Maillard*. Si tratta di una specie di "caramellizzazione" delle proteine che sono adagiate sul tappeto di cui si diceva e che determina lo speciale "*flavour*" dell'arrosto. Come appunto fu scoperto da **Camille Maillard** nel 1910 quel processo fa acquisire all'arrosto quell'aroma e quel gusto così speciale. Nella carne la reazione di Maillard avviene a $140^\circ C$ e la temperatura del forno alla quale si raccomanda di cucinare il tacchino è all'incirca $160 - 170^\circ C$. Ora la domanda è: quale tempo è richiesto affinché tale temperatura si distribuisca uniformemente su tutto il volume del tacchino?

Cerchiamo di formulare un modello del processo di cottura. Assumiamo un corpo sferico di raggio R che inizialmente alla temperatura T_0 viene posto in un forno alla temperatura T_f ; la conducibilità termica del materiale sia κ . In quale tempo la temperatura al centro della sfera raggiungerà il valore T_M ? Il processo di trasferimento del calore è regolato dalla equazione

$$\frac{\partial T(\mathbf{r}, t)}{\partial t} = -\kappa \left(\frac{\partial^2 T(\mathbf{r}, t)}{\partial x^2} + \frac{\partial^2 T(\mathbf{r}, t)}{\partial y^2} + \frac{\partial^2 T(\mathbf{r}, t)}{\partial z^2} \right), \quad (17.1)$$

ove $T(\mathbf{r}, t)$ è la temperatura al tempo t al punto $\mathbf{r} = (x, y, z)$. Evidentemente avremo che ad ogni valore del tempo la temperatura per $r = R$ è pari a T_f , mentre vorremo che per $t = \tau_{rst}$ sia $T = T_M$ per $r = 0$. Questa equazione può essere risolta esattamente, in maniera analoga a come menzionato per la cottura dell'uovo "*à la coque*". Tuttavia possiamo ottenere dei risultati significativi senza risolvere l'equazione differenziale ma ancora una volta ricorrendo a considerazioni di scala e di dimensionalità. A destra e a sinistra

della equazione (17.1) abbiamo le dimensioni

$$\frac{[T]}{[t]} = [\kappa] \frac{[T]}{[r]^2}, \quad (17.2)$$

ovvero

$$[t] = [\kappa]^{-1} [r]^2. \quad (17.3)$$

La sola grandezza fisica avente la dimensione di una lunghezza è il raggio della sfera R quindi la sola combinazione che corrisponde alla eguaglianza (17.3) è

$$\tau_{\text{rst}} = \alpha' R^2 + \tau_0. \quad (17.4)$$

Il coefficiente τ_0 dipende dalla differenza $T_f - T_M$ e dal tempo richiesto per la reazione di Maillard una volta che la temperatura al punto abbia raggiunto il valore T_M . La massa M della sfera di carne che nel nostro modello simula il tacchino è proporzionale a R^3 , quindi R dipende da $M^{1/3}$ cosicchè la formula (17.4) può essere riscritta

$$\tau_{\text{rst}} = \alpha M^{2/3} + \tau_0. \quad (17.5)$$

Il coefficiente α e il suo valore numerico dipendono dalla forma reale del tacchino, che noi invece abbiamo assunto sferico. Ricorriamo ora alla Tabella che fornisce il tempo di cottura raccomandato empiricamente sull'involuppo del tacchino acquistato al supermercato:

Massa (libbre)	Tempo di cottura (ore)
6	3
8	3.5
12	4.5
16	5.5
20	6.5
24	7.35

Riconsideriamo tali tempi di cottura riportando la dipendenza di τ^3 in funzione di M^2 (Fig. 17.3). Si nota in tale rappresentazione una legge abbastanza ben descritta da una retta. Si osservi che il fatto che τ_{rst} vada a zero se M tende a zero significa semplicemente che il tempo richiesto dalla reazione di Maillard è trascurabile rispetto a quello richiesto affinché

la temperatura al centro della sfera diventi pari a T_f , in altre parole affinché il calore possa fluire al centro del tacchino.

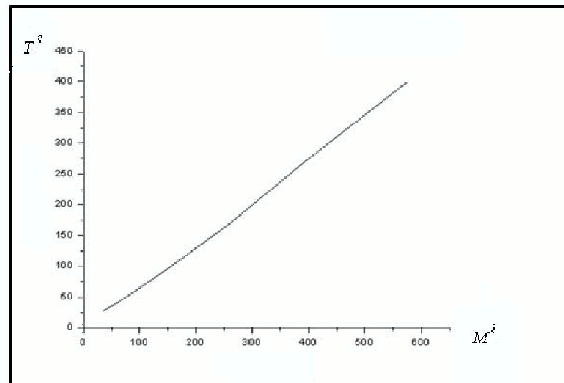


Fig. 17.3: Il cubo del tempo di cottura in funzione del quadrato della massa del tacchino.

Capitolo 18

Pasta, spaghetti e fisica.

Tutti conoscono gli spaghetti e molti hanno personalmente preparato la classica pentola d'acqua per la loro cottura. Non molti, tuttavia, si sono soffermati a meditare sui molteplici fenomeni coinvolti nel processo di preparazione di un piatto di spaghetti ben cotti. Per esempio, vi siete mai chiesti: che cosa avviene realmente nel corso del tempo in cui gli spaghetti rimangono in acqua bollente? Perché occorre un preciso intervallo di tempo per ottenere un adeguato livello di cottura? Perché questo tempo è diverso da un tipo di pasta ad un altro? Come interviene la “forma” del prodotto (spaghetto classico ma di diametro diverso oppure rigatoni o bucatini etc.) nel condizionare il tempo di cottura? Perché gli spaghetti non si rompono mai in due soli pezzi quando vengono piegati oltre un certo raggio di curvatura ma si frantumano invece in almeno tre frammenti? Perché gli spaghetti non si annodano? Quale tipo di pasta deve essere scelto per ottenere una buona adesione del sugo?

Questi interrogativi hanno stimolato l'interesse di molti ricercatori e ora cercheremo di descrivervi la fisica del processo di cottura e insieme di dare risposta ad alcune di queste o di altre domande che possono sorgere in voi mentre attendete che la cottura abbia termine.

18.1 Un cenno di storia della pasta e sulla tecnica della sua produzione

Contrariamente a quanto viene spesso sostenuto, la pasta non fu introdotta in Occidente da **Marco Polo** al ritorno dal suo viaggio in Cina (1295). In realtà, da tempi assai precedenti la storia della pasta si è snodata

lungo le coste del Mediterraneo, a partire probabilmente sin dal momento in cui l'uomo primitivo, abbandonata la vita nomade, aveva appreso la coltivazione del grano e le prime manipolazioni del prodotto. Le prime focacce, cotte su pietre roventi, sono menzionate nella Bibbia e in altri testi come il Libro della Genesi e il Libro dei Re. Nel primo millennio avanti Cristo i Greci indicavano con il termine "laganon" una pasta sotto forma di sfoglia. Il termine fu ripreso dai Romani nella forma "laganum", forse precursore delle moderne lasagne.

Anche gli Etruschi conoscevano la pasta in forma di sfoglia. Con l'impero Romano la pasta si era diffusa nell'Europa occidentale. L'uso di essiccare la pasta per prolungarne la conservazione nacque da necessità logistiche, in relazione agli spostamenti che intere tribù erano costrette a compiere. In Sicilia l'essiccazione fu introdotta dagli Arabi, attorno al X secolo. Gli antenati degli spaghetti sono presumibilmente le "trie", in forma di strisce sottili comparse nei pressi di Palermo poco dopo il primo millennio, derivate forse dalle "itryah" arabe (focaccia tagliata a strisce). Attorno al 1280, è certa la presenza in Liguria dei "maccheroni", attraverso l'inventario di una eredità redatto dal notaio genovese **Ugolino Scarpa**. Dalla storia della letteratura italiana sappiamo come la pasta abbia sollecitato la attenzione di **Jacopone da Todi**, di **Cecco Angiolieri**, per essere sublimata nella forma di maccheroni come emblema di elevata prelibatezza nel Decamerone del **Boccaccio**.

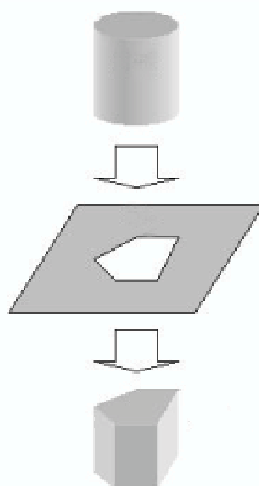
Le prime corporazioni dei "Pastai", con riconoscimento sociale e politico e proprie norme statutarie, si registrano in Italia tra il 1500 e il 1600. I maccheroni sono a quel tempo cibo delle classi agiate, soprattutto nelle regioni ove manca la produzione locale di grano duro (per esempio Napoli). L'invenzione del torchio meccanico determina una produzione di pasta a prezzi più contenuti e nel secolo XVII la pasta fa parte della alimentazione di tutte le classi sociali e praticamente diffusa in tutta l'area mediterranea. L'industria napoletana, in particolare, conquista una sorta di primato nella produzione e nella esportazione dei maccheroni e a Napoli la pasta viene venduta agli angoli delle strade, condita con formaggio grattugiato o con salsa di pomodoro e basilico. Nell'Italia settentrionale la diffusione massiccia della pasta ha luogo attorno alla fine dell'ottocento, soprattutto ad opera del pioniere **Pietro Barilla** che apre in Parma una bottega che avrebbe acquisito un ruolo di grande rilevanza nella gastronomia italiana.

I metodi di produzione industriale di pasta nelle forme oggi diffuse cos-

tituiscano un aspetto di rilievo nelle moderne tecnologie industriali relative ai processi di estrusione e di trafilatura.

L'estrusione è stata ideata e applicata allo scopo di ottenere manufatti, spesso metallici, piuttosto lunghi e con una determinata struttura della loro sezione (Si veda la Figura 18.1). Il processo di estrusione può avvenire a freddo o a caldo ed è essenzialmente basato sul principio di forzare con la compressione, mediante utensili e stampi, il passaggio di un determinato materiale attraverso opportune sagome. La trafilatura è simile alla estrusione, l'unica sostanziale differenza essendo che il materiale viene tirato all'uscita del contenitore di sagomatura e quindi viene soggetto a trazione piuttosto che a compressione. Si producono in genere cilindri, fili e tubi. Per fili metallici è possibile raggiungere sezione di diametro pari a 0.025 millimetri. Altri materiali che sono comunemente estrusi sono i polimeri, le ceramiche e ... gli alimenti.

Fig. 18.1: Il processo di estrusione è basato sul principio di forzare con la compressione il passaggio di un determinato materiale attraverso opportune sagome.



Nella Figura 18.2 sono riportati esempi di utensili di filiera per la trafilatura della pasta.



Fig. 18.2: La filiera per gli spaghetti.

18.2 Il processo di cottura degli spaghetti

Nel precedente capitolo abbiamo discusso il processo di cottura della carne e anche ottenuta un'espressione matematica per il tempo di cottura nel caso dei tacchini immolati nel giorno del ringraziamento negli Stati Uniti. Le procedure utilizzate in quel caso, sebbene con diversi coefficienti, possono essere utilizzate anche per valutare il tempo di cottura di diversi tipi di pasta.

Occorre innanzitutto stabilire chiaramente cosa significhi “cucinare la pasta”. Nella farina di grano le molecole di amido si trovano raggruppate a formare granuli di 10-30 micron di diametro, mescolati con diverse proteine. Due di queste proteine, la gliadina e la glutenina, quando entrano a contatto con l'acqua durante la lavorazione della pasta, si combinano insieme e formano una maglia continua, elastica e impermeabile che racchiude nelle sue trame i granuli di amido: il glutine. La cottura dipende proprio dalla capacità dei granuli di amido, intrappolati nel glutine al momento dell'essiccamento, di assorbire l'acqua che, nella pentola, inizia a farsi strada nella maglie di glutine e a migrare verso l'interno degli spaghetti. Intanto le molecole di amido, all'incirca alla temperatura di 70°C , si combinano a formare un composto simile ad un gel, che inizia a trattenere l'acqua. Il traguardo della cottura al dente si raggiunge nel momento in cui l'amido trasformato a gel ha assorbito la quantità d'acqua minima necessaria per renderlo tenero. Pertanto, il primo componente necessario per la cottura della pasta è lo stesso che interviene nel processo di cottura della carne: è necessario aumentare la temperatura all'interno della pasta, portandovi calore. In secondo luogo è necessario permeare d'acqua la pasta stessa. En-

trambe queste condizioni debbono essere soddisfatte contemporaneamente: riscaldare gli spaghetti al forno oppure bagnarli semplicemente nell'acqua fredda non permetterebbe di ottenere un piatto di pasta appetitoso.

Le equazioni (17.4) e (17.5) sono state ottenute come risultato dell'analisi dimensionale dell'equazione di trasferimento del calore e applicate precedentemente al caso di oggetti con simmetria sferica. Con poche modifiche possono essere applicate anche a sistemi a simmetria cilindrica, come appunto gli spaghetti. Inoltre, fortunatamente, il processo di diffusione dell'acqua dalla superficie al centro degli spaghetti è descritto dallo stesso tipo di equazione (l'equazione di diffusione) scritta per la concentrazione d'acqua invece che per la temperatura. In definitiva il tempo di cottura è legato al diametro degli spaghetti da una relazione simile alla (17.4):

$$\tau_{\text{sp}} = aD^2 + b. \quad (18.1)$$

Il coefficiente a è determinato dalle proprietà fisiche della pasta (la sua conducibilità termica, il coefficiente di diffusione dell'acqua nella pasta) mentre si può dire che il coefficiente b caratterizza la nazionalità del consumatore. Infatti, mentre il primo termine nella equazione (18.1) determina il tempo necessario affinché l'acqua e una certa quantità di calore raggiungano il centro degli spaghetti, il secondo termine misura quanto a lungo il centro degli spaghetti rimane esposto a questi due agenti. Questo spiega perché per gli Italiani, i quali preferiscono mangiare gli spaghetti "al dente", il processo di gelificazione dell'amido non debba essere completato fino al centro degli spaghetti: come vedremo in seguito, il coefficiente b debba essere in tal caso negativo. D'altra parte i consumatori di spaghetti di nazionalità diversa dall'italiana spesso ritengono che la pasta debba essere "ben cotta" e sono quindi portati a prolungare il tempo di cottura superando, anche in modo considerevole, quello indicato dal produttore sulla confezione.

Per procedere, procuriamoci al supermercato l'intera gamma di pasta di forma cilindrica: capellini, spaghetti, vermicelli, bucatini. Raccogliamo in una Tabella il tempo di cottura raccomandato (colonna "tempo di cottura"). Quindi con un calibro misuriamo il diametro di ciascuno dei differenti tipi di pasta e nella stessa Tabella, compiliamo la colonna "Diametro esterno/interno".

Tipo di pasta	Diametro esterno (mm)	Diametro interno (mm)	Tempo di cottura sperimentale (min)
capellini n.1	1.15	-	3
spaghettoni n. 3	1.45	-	5
spaghetti n. 5	1.75	-	8
vermicelli n.7	1.90	-	11
vermicelli n.8	2.10	-	13
bucatini	2.70	1	8

Tabella 18.1

Allo scopo di ottenere i valori numerici dei coefficienti a e b è sufficiente riscrivere l'equazione (18.1) utilizzando i dati di due particolari righe della nostra Tabella 18.2:

$$\begin{aligned} t_1 &= aD_1^2 + b \\ t_2 &= aD_2^2 + b \end{aligned} \quad (18.2)$$

e risolvere il sistema costituito da queste due semplici equazioni. Come esempio possiamo scegliere i dati che si riferiscono agli spaghettoni n.3 e ai vermicelli n.8. Otteniamo i risultati seguenti:

$$a = \frac{t_2 - t_1}{D_2^2 - D_1^2} = 3.4 \text{ min/mm}^2$$

$$b = \frac{D_2^2 t_1 - D_1^2 t_2}{D_2^2 - D_1^2} = -2.3 \text{ min.}$$

Poiché la pasta che abbiamo utilizzato era destinata al consumatore italiano, il tempo di cottura raccomandato sulla confezione si riferisce ad una cottura “al dente”. Questo viene evidenziato dalla circostanza che il coefficiente b è negativo. Avendo calcolato i valori numerici dei coefficienti a e b , possiamo ora verificare quali tempi di cottura si ottengono per gli altri tipi di pasta. I risultati del calcolo, riportati nell'ultima colonna, sono in buon accordo con i valori raccomandati dal produttore su base empirica, eccetto che per i due casi estremi: capellini e bucatini.

Tipo di pasta	Tempo di cottura sperimentale (min)	Tempo di cottura teorico (min)
capellini n.1	3	2.2
spaghettoni n. 3	5	5.0
spaghetti n. 5	8	8.1
vermicelli n.7	11	10.0
vermicelli n.8	13	13.0
bucatini	8	22.5 ?!

Tabella 18.2

Questa è una situazione tipica in fisica: le predizioni di un modello teorico hanno una regione di validità che dipende dalle assunzioni adottate in partenza nel modello stesso. Ad esempio, se guardiamo l'ultima riga della colonna, nel caso dei bucatini, la differenza tra il valore teorico e quello empirico- sperimentale è molto grande : 22.5 minuti contro 8 minuti. Questa contraddizione riflette una circostanza rilevante: l'intervallo di variabilità dello spessore per tutti i tipi di pasta cilindrica è molto piccolo, soltanto un millimetro. Infatti, i nostri calcoli del tempo di cottura “al dente” dei bucatini, hanno condotto a 22.5 minuti, ma nella approssimazione di cilindro pieno. In realtà, un tempo di cottura così lungo, sovracuocerebbe completamente la superficie esterna del bucatino (condizione di pasta “scotta”). Ecco perché, per un tipo di pasta così spessa si è empiricamente trovato il “trucco” di avere un “buco nel cilindro”. Durante il processo di cottura l'acqua penetra nel buco e non è più necessario trasferire calore e liquido dalla superficie esterna del cilindro sino al suo centro: la diffusione avviene da due parti!

Si può cercare di modificare la nostra formula teorica sottraendo il diametro interno da quello esterno, in questo modo avvicinando il risultato previsto a quello fenomenologico:

$$t = a (D - d)^2 + b \approx 7.5 \text{ min} .$$

In ogni caso è necessario che il buco non abbia un diametro inferiore al ~ 1 mm. In caso contrario, a causa della pressione capillare l'acqua non sarebbe in grado di penetrare al suo interno:

$$P_{\text{cap}} = \frac{\sigma (T = 100^\circ\text{C})}{d} \gtrsim 50 \text{ Pa},$$

che corrisponde alla pressione esercitata da una colonna di acqua alta un centimetro sopra il livello della pasta.

Un'altra deviazione dalla realtà sperimentale della teoria basata su un certo modello si ha nel caso di pasta molto sottile. L'origine di questo "fallimento" è evidente: per la nostra definizione di pasta "al dente" il parametro b è negativo e assume il valore: $b = -2.3 \text{ min}$. Da un punto di vista puramente formale questo significa che dovrebbe esistere una pasta così sottile che essa si troverebbe nella condizione di "cottura al dente" anche senza essere stata cotta affatto. Il suo diametro può essere determinato tramite la relazione

$$0 = aD_{\text{cr}}^2 + b,$$

che comporta

$$D_{\text{cr}} = \sqrt{|b|/a} \approx 0.82 \text{ mm}.$$

Si può osservare che il reale diametro dei capellini (1.15 mm) non è molto lontano da questo valore critico, pertanto la sottostima del tempo di cottura dei capellini è il risultato di questo limite nel nostro modello.

18.3 Nodi negli spaghetti?

Gli spaghetti, durante il processo di cottura in acqua formano intrecci complessi tra di loro ma non si annodano mai. La ragione di questo fatto può essere giustificata con il ricorso alla statistica delle catene polimeriche in soluzione, un campo relativamente nuovo nella fisica statistica^a (al capitolo della cucina molecolare ne faremo riferimento).

La probabilità che una lunga catena polimerica formi un nodo è fissata dalla equazione

$$w = 1 - \exp\left(-\frac{L}{\gamma\xi}\right),$$

dove L rappresenta la lunghezza della catena stessa, ξ è una lunghezza caratteristica alla quale il polimero può cambiare la sua direzione di $\pi/2$ e $\gamma \approx 300$ un fattore numerico piuttosto grande. Applicando questa formula agli spaghetti, dove $\xi \sim 2 \div 3 \text{ cm}$, si può stimare la lunghezza che dovrebbero avere affinché il processo di annodamento diventi significativamente

^aGli autori sono grati ad A.Y.Grosberg per l'introduzione alla teoria dei nodi.

probabile ($w \sim 0.1$) :

$$\exp\left(-\frac{L_{\min}}{\gamma\xi}\right) \approx 0.9$$

vale a dire

$$L_{\min} \approx \gamma\xi \ln 1.1 \approx 30\xi \approx 60 \div 90 \text{ cm.}$$

Con la lunghezza standard degli spaghetti, pari a 23cm, è pertanto piuttosto improbabile che l'annodamento si realizzi!

18.4 Viscosità del sugo e forma della pasta

Discutiamo il problema della migliore adesione di un sugo alla pasta nel quadro di un semplice modello. Il flusso stazionario ($Q = \Delta M / \Delta t$) di un fluido all'interno di un tubo di diametro D in presenza di una differenza di pressione ΔP è dato dalla formula di **Poiseuille**

$$Q = \frac{\pi \rho \Delta P}{2^7 \eta l} D^4,$$

dove η è la viscosità del liquido, ρ la sua densità ed l la lunghezza del tubo. D'altra parte

$$Q = \Delta M / \Delta t = \frac{1}{4} \pi \rho D^2 \Delta l / \Delta t$$

Confrontando queste due formule e attribuendo la differenza di pressione alla forza di gravità $\Delta P = \rho g l$ otteniamo

$$\frac{\Delta t}{\Delta l} = \frac{32}{\rho g} \left(\frac{\eta}{D^2} \right).$$

Pertanto la velocità con la quale il tubo viene "bagnato" dal fluido viscoso che fluisce attraverso di esso è determinata dalla combinazione η/D^2 . Per ottenere questa formula abbiamo usato il modello del flusso viscoso in un tubo verticale, in presenza di un campo di gravità. Comunque si deve osservare che l'origine dell'accelerazione è irrilevante: essa può essere dovuta alla gravità oppure impressa agli spaghetti mescolandoli vigorosamente nella zuppiera. Pertanto il tempo di mescolamento viene scritto

$$\tau_{\text{mix}} \sim \left(\frac{\eta}{D^2} \right).$$

In definitiva, nel caso di sughi più viscosi è opportuno scegliere pasta di dimensioni maggiori.

18.5 La rottura degli spaghetti

Abbiamo anche accennato all'inizio ad un altro interessante aspetto della fisica degli spaghetti, relativo al processo di frammentazione. Prendete uno spaghetti e piegatelo dolcemente sempre di più tenendolo per le due estremità. Potreste pensare che a un certo momento si spezzerà in due frammenti, all'incirca nel punto di mezzo. Invece no, vedrete che i frammenti sono sempre in numero maggiore (tre o più, sino anche a dieci): lo spaghetti si rifiuta di spezzarsi in sole due parti.

Questo bizzarro comportamento ha destato l'interesse di molti scienziati, tra i quali anche il grande **Richard Feynman**. Solo recentemente ad opera di due fisici francesi, **Auduly** e **Neukirch**, si è potuto avere una giustificazione quantitativa, che ha destato anche interesse per le possibili applicazioni nel campo delle costruzioni e della ingegneria civile. Questi scienziati hanno applicato allo spaghetti le equazioni di **Gustav Robert Kirchhoff** per modellizzare le modalità secondo le quali una sottile barra elastica reagisce alla sollecitazione che tende a farla incurvare. Il loro lavoro è stato pubblicato sulla prestigiosa rivista scientifica *Physical Review Letters* nell'Agosto 2005 ed è abbastanza difficile da riassumere in quanto quelle equazioni sono state risolte, con le appropriate condizioni al contorno, solo per via numerica.

Possiamo tuttavia spiegarvi qualitativamente l'elemento chiave che determina la frattura in più frammenti. Il fatto è il seguente. Ci si potrebbe aspettare che dopo il primo crack iniziale le due metà dello spaghetti si rilassino verso la loro posizione di equilibrio. Invece, subito dopo la prima rottura nascono delle onde elastiche che si propagano lungo lo spaghetti. Queste onde comportano un incremento locale della curvatura al di sopra del valore di rottura e accendono quindi una specie di "valanga" di rotture successive lungo la restante parte dello spaghetti nella quale si sono propagate. Queste onde "aggiuntive" stimulate dalla prima rottura vanno spegnendosi progressivamente ma sono in grado di produrre un certo numero successivo di frammentazioni. Questa brillante trattazione è stata confermata da una ripresa cinematografica ad alta velocità (Fig. 18.3).

Mentre i vostri spaghetti vanno a cottura nella pentola potrete divertirvi

Fig. 18.3: Una
istantanea della
rottura dello
spaghetto.



con qualcuno di quelli secchi a confermare il lavoro di Auduly e Neukirch!

Capitolo 19

La fisica di un buon caffè.

Il viaggiatore abituato a muoversi in paesi diversi, può notare che nel nostro periodo storico di globalizzazione e di appiattimento su standard multinazionali, mentre quasi tutte le bevande vengono offerte assai simili sia a New York sia a Katmandu, il caffè rimane invece una bevanda regionale. Bere un caffè è un'esperienza molto diversa a seconda che ci si trovi in Turchia o in Egitto, in Italia o in Francia, in Finlandia o in USA.

Se chiedete un caffè in un bar di Napoli vi verrà servito in una piccola tazzina elegante sul fondo della quale lentamente oscilla un liquido viscoso di colore quasi nero, coperto da una crema invitante. Invece, ordinando un caffè a Chicago, vi viene portato un barattolo di polistirolo di mezzo litro di volume riempito di acqua calda di colore marrone scuro.

Noi non vogliamo qui giudicare quale delle due bevande sia la più gustosa e la più salutare, ma ci limiteremo a discutere i vari metodi usati per la preparazione del caffè e i processi fisici ad essa collegati.

19.1 Caffè bollito

Un antico metodo di preparazione del caffè, arrivato fino ai giorni nostri in Finlandia e nelle zone settentrionali dei paesi scandinavi, è quello che chiameremo del bollitore. Il caffè tostato viene macinato in grani grossi, versato in acqua (circa 10 grammi ogni 150–190 *ml* di acqua) e messo in un bollitore per circa 10 min. Quindi la bevanda, senza essere filtrata, viene servita nelle tazze e lasciata riposare per qualche minuto. Non c'è nessun processo fisico interessante in questo metodo e gli autori preferiscono non discutere il sapore della bevanda ottenuta.

19.2 Caffè americano (caffettiera con filtro di carta)

Un altro metodo si basa sulla caffettiera con filtro: è molto comune negli Stati Uniti, nel nord Europa, in Germania e in Francia. Il suo principio di funzionamento è estremamente semplice e il processo richiede 6 – 8 minuti. Il caffè, di macinatura grossa, viene versato in un filtro conico costituito da carta da filtro. Quindi, l'acqua portata all'ebollizione scende bagnando i granuli di caffè macinato, penetra attraverso il filtro e si raccoglie in un contenitore di vetro. Come risultato avremo un caffè leggero: solo pochi oli riescono a passare attraverso la spessa carta da filtro. Inoltre, la macinatura a grani grossi del caffè e la mancanza di pressione non favoriscono la completa estrazione di tutti gli aromi presenti. La dose americana consiste di 5–6 grammi per 150–190 *ml* di acqua, mentre in Europa è di 10 grammi per tazza.

19.3 Caffè alla turca

Nel caffè alla “*turca*” i chicchi vengono macinati fino a ridurli in polvere e questa polvere, spesso mescolata con zucchero, viene versata in un contenitore metallico (di solito di rame od ottone) a forma di cono, che prende il nome di “*ibrik*” (Fig. 19.1).

Successivamente si versa nella *ibrik* dell'acqua e il tutto viene immerso nella sabbia rovente. In un'altra ricetta il caffè macinato si mette nel contenitore già riempito di acqua bollente (se non avete a disposizione della sabbia rovente si può utilizzare la fiamma di un fornello a gas, oppure un fornello elettrico). La trasmissione del calore dalla sabbia attraverso il fondo e le pareti della *ibrik* fa sì che il liquido si riscaldi generando correnti convettive: il liquido caldo porta con sé sulla superficie parte della polvere di caffè.

Sulla superficie, grazie alla forza di tensione superficiale, si forma “una crosta di caffè”. Piano piano il contenuto della *ibrik* arriva all'ebollizione, le bolle cominciano a rompere la crosta formando così una specie di schiuma. A questo punto si interrompe il processo togliendo l'*ibrik* dalla sabbia, in modo da non rovinare il gusto del caffè. Questa procedura si deve ripetere ancora due volte, fino a formare uno spesso strato di schiuma. Si versa quindi il contenuto in tazze e si aspetta che la polvere di caffè si depositi sul fondo. Si ottiene infine una bevanda densa e gustosa (soprattutto se la

Fig. 19.1: La Ibrik
per la preparazione
del caffè alla turca.



quantità di acqua non è elevata).

19.4 La moka italiana

La moka è la caffettiera più usata in Italia. Essa è composta da tre parti: la base dove si versa l'acqua da riscaldare, un filtro metallico cilindrico, dove si mette la polvere di caffè ottenuta da una macinazione abbastanza fine e la parte superiore, approssimativamente a forma di cono tronco rovesciato, dove esce la bevanda pronta (Fig. 19.2). Il filtro è il cuore della caffettiera moka: inferiormente al filtro è fissato un imbuto metallico che pesca molto vicino al fondo del contenitore di base.

La caffettiera moka non offre molto spazio all'inventiva del preparatore. A differenza di altri metodi di preparazione del caffè, egli deve infatti seguire la ricetta fissata dai parametri progettuali della caffettiera: versare acqua nel contenitore di base fino a che essa non raggiunga in altezza la valvola di sicurezza e il filtro deve essere riempito (circa 6 grammi di caffè per 50 ml di acqua).

Dal punto di vista fisico il processo di preparazione del caffè con la caffettiera moka è molto interessante. Il caffè macinato posto nel filtro viene pressato solo moderatamente. La moka si chiude avvitando la base

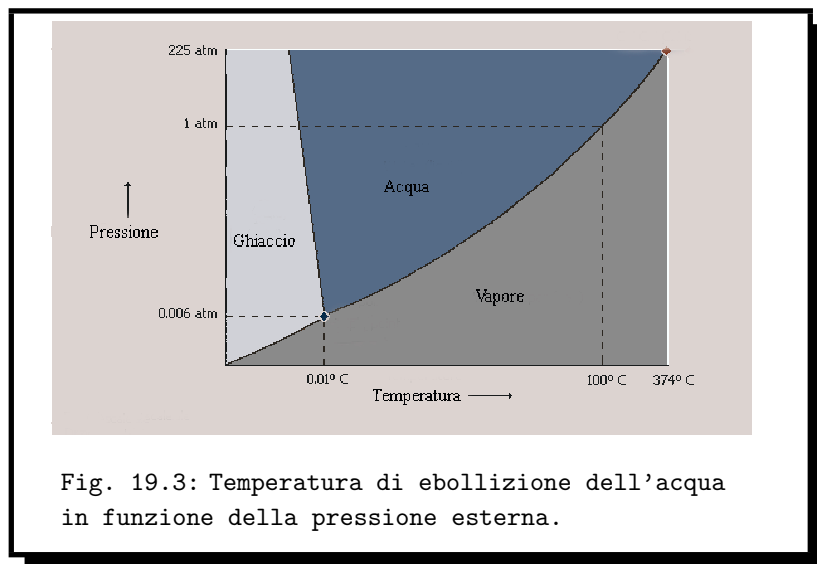


Fig. 19.2: La tipica
macchina da caffè
Moka.

e la parte superiore intorno al filtro, che copre quindi l'acqua contenuta nella base. La tenuta tra le due metà è assicurata da una guarnizione di gomma. La caffettiera viene posta su di una fiamma debole. L'acqua contenuta nella base viene riscaldata. Di conseguenza, la pressione del vapore sopra la superficie dell'acqua e sotto l'imbuto metallico aumenta rapidamente forzando l'acqua a risalire attraverso il gambo dell'imbuto e quindi attraverso la polvere di caffè contenuta nel filtro. Il liquido che a questo punto costituisce il nostro caffè, sale infine attraverso un sottile tubo fino a ricadere nella metà superiore della caffettiera. A questo punto può essere servito.

Questo metodo sembra facile da giustificare in termini fisici. Cosa muove il processo descritto? Certamente il fuoco. All'inizio si riscalda l'acqua in un volume chiuso, dove l'acqua stessa occupa la maggior parte dello spazio a disposizione. L'acqua nella caffettiera raggiunge ben presto la temperatura di $100^\circ C$ (temperatura di ebollizione al livello del mare) e la pressione del vapore saturo sopra l'acqua raggiunge il valore di 1 atm. Continuando a somministrare calore alla caffettiera, crescono sia la pressione del vapore saturo sia la temperatura dell'acqua: acqua e vapore saturo rimangono in equilibrio a pressioni e temperature sempre crescenti. La Fig. 19.3 mostra l'andamento della pressione del vapore saturo in funzione della temperatura. D'altra parte, la pressione esterna (al di sopra del filtro) rimane uguale alla pressione atmosferica. Il vapore saturo, che si trova ad una temperatura maggiore di $100^\circ C$, agisce quindi come una molla spingendo l'acqua bollente, un poco surriscaldata, attraverso la polvere di caffè contenuta nel filtro. In questo modo vengono estratti tutti gli aromi, l'essenza del caffè e

altre componenti che trasformano l'acqua in una gustosa bevanda.



Ovviamente il gusto di questa bevanda dipende dalla qualità della polvere di caffè, dalla temperatura dell'acqua e dal tempo impiegato dall'acqua per attraversare il filtro. La preparazione della miscela, la tostatura dei chicchi di caffè e la macinatura sono segreti del produttore, basati sul talento, sul lavoro e su una lunga esperienza. Da che cosa dipenda il tempo di transito dell'acqua attraverso il filtro, noi lo possiamo capire senza ricorrere allo spionaggio industriale, basandoci soltanto sulle leggi fisiche.

Nella metà del diciannovesimo secolo due ingegneri francesi **A. Darcy** e **G. Dupuis** condussero le prime osservazioni sperimentali sul fluire dell'acqua nei tubi riempiti di sabbia. Queste ricerche rappresentarono l'inizio dello sviluppo della teoria della filtrazione, che oggi si applica con successo allo studio del moto dei liquidi attraverso i solidi porosi o con crepe interconnesse. Fu quindi Darcy a formulare la cosiddetta legge di filtrazione lineare, che oggi porta il suo nome. Questa legge mette in relazione la massa di fluido che passa in un secondo (Q) attraverso un filtro di spessore L e sezione S , con la differenza di pressione ΔP alle due estremità del filtro stesso:

$$Q = \kappa \frac{\rho S}{\eta L} \Delta P. \quad (19.1)$$

In questa formula ρ e η rappresentano, rispettivamente, la densità e la viscosità del fluido. La costante κ (detta coefficiente di filtrazione) non dipende dalle caratteristiche del liquido ma solo dal mezzo poroso. Nei lavori originali di Darcy si supposeva che la differenza di pressione ΔP originasse esclusivamente dalla forza peso. In tali condizioni $\Delta P = \rho g \Delta H$, ove g è l'accelerazione di gravità e ΔH rappresenta la differenza di quota tra le due estremità del filtro che si suppone disposto verticalmente. Come si può vedere dall'analisi della formula (19.1) la costante κ ha le dimensioni di una superficie, si misura quindi, nel Sistema Internazionale, in m^2 . I valori numerici della costante sono, di solito, molto piccoli. Per esempio in un terreno sabbioso costituito da grani grossi il coefficiente di penetrazione κ assume valori dell'ordine $10^{-12} - 10^{-13} m^2$, mentre in un terreno di sabbia pressata $\kappa \sim 10^{-14} m^2$.

Cerchiamo di applicare ora la legge di Darcy allo studio della moka. Per esempio è interessante sapere fino a quale temperatura si surriscalda l'acqua bollente nella base della caffettiera. Tale temperatura si può ricavare analizzando la legge di dipendenza della temperatura di ebollizione dalla pressione esterna (Tabella 14.1) e utilizzando la legge di Darcy:

$$\Delta P = \frac{m}{\rho S t} \frac{\eta L}{\kappa}. \quad (19.2)$$

L'usuale spessore del filtro, per una moka da tre persone, è $L = 1 \text{ cm}$, la sezione vale $S = 30 \text{ cm}^2$ e la massa del caffè, all'incirca $m = 100 \text{ g}$ viene preparata in un tempo t di circa un minuto. Come coefficiente di filtrazione possiamo utilizzare quello di un terreno di sabbia grossa $\kappa \approx 10^{-13} m^2$. Con la viscosità è necessario fare attenzione poiché essa dipende dalla temperatura. Però nelle tavole si può trovare che $\eta(100^\circ \text{ C}) = 10^{-3} \text{ Pa} \cdot \text{s}$. Da questo segue che $\Delta P \approx 4 \cdot 10^4 \text{ Pa}$. La temperatura di ebollizione dell'acqua, come si vede dal grafico nella figura 19.3, è all'incirca $T^* \sim 110^\circ \text{ C}$.

Abbiamo quindi capito la base fisica del processo di preparazione del caffè in una moka italiana.

Vi sono peraltro lati oscuri in questa "macchina" che si può trasformare in una vera e propria bomba, minacciando il soffitto e le pareti della cucina, per non parlare dell'incolumità di chi si trova nei pressi. Come può accadere? È evidente che tale esito disastroso si può avere soltanto nel caso in cui la valvola di sicurezza, posta nella base della caffettiera, sia ostruita (le caffettiere vecchie sono pericolose!). Allo stesso tempo è necessario che

sia precluso anche il naturale sfogo attraverso il filtro contenente il caffè. Ad esempio, ciò può accadere quando la polvere è macinata troppo fine e molto pressata, si da divenire quasi impenetrabile per l'acqua. Se entrambe le condizioni sono verificate, la pressione nella base può crescere a valori così elevati da rompere la filettatura sulla quale è avvitata la parte superiore della caffettiera.

Per spiegare brevemente i fenomeni che si possono produrre con una polvere eccessivamente pressata si può richiamare quanto segue. A livello semi- microscopico il filtro può essere immaginato come un sistema di capillari di varie sezioni e lunghezze connessi tra di loro. L'applicabilità della legge di Darcy implica che il flusso del liquido attraverso il filtro avvenga in modo laminare. Nella realtà, eventuali brusche irregolarità nella struttura dei capillari implicano la formazione di vortici e il conseguente passaggio a moto turbolento. Questi fenomeni dissipativi richiedono, al fine di mantenere il flusso costante attraverso il filtro, un aumento della pressione applicata. L'altro motivo che stabilisce il limite inferiore per la trasparenza del filtro è legato al fenomeno della tensione superficiale. Infatti è noto che in un capillare ideale di raggio r è necessario applicare una differenza di pressione $\Delta P > 2\sigma/r$. Stimando la sezione dei capillari dell'ordine di $\sim 10^{-4} m$, e ricordando che σ è circa $0.05 N/m$, si ottiene che la soglia per la trasparenza del filtro è $10^3 Pa$. Questo valore è circa un ordine di grandezza inferiore rispetto a quello stimato per la differenza di pressione applicata al filtro alla temperatura di $110^\circ C$. In queste condizioni è chiaro che se polveri sottili di caffè vengono eccessivamente pressate nel filtro, il raggio efficace dei capillari può divenire molto minore del valore di $10^{-4} m$ stimato precedentemente e il filtro può diventare in pratica quasi impenetrabile.

Le caffettiere che si trasformano in bombe possono essere realmente molto pericolose. Partiamo dalla situazione peggiore: il filtro e la valvola sono ostruiti e 100 grammi di acqua si riscaldano in un volume chiuso appena un po' più grande di quello a disposizione dell'acqua. Per temperature vicine alle condizioni critiche dell'acqua (dove la densità del vapore è paragonabile quella dell'acqua), ad una temperatura di $T = 374^\circ C = 647 K$ tutta l'acqua diventa vapore. In linea di principio è possibile scaldare ulteriormente la moka. Tuttavia, a temperature ancora più alte, la moka dovrebbe diventare luminosa e ciò non avviene comunemente! Quindi, per un calcolo ragionevole, si può assumere che la moka si sia riscaldata sino ad una temperatura di circa $T = 600 K$. La pressione del vapore nella base della caffettiera può essere stimata a partire dall'equazione dei gas perfetti:

$$PV = \frac{m}{\mu} RT. \quad (19.3)$$

Poiché $m = 100 \text{ g}$, $V = 120 \text{ cm}^3$, $\mu = 18 \text{ g/mole}$, $R = 8.31 \text{ J/(mole} \cdot \text{K)}$, possiamo ricavare $P \approx 10^8 \text{ Pa} = 10^3 \text{ atmosfere}$. Questa è circa la pressione che si ha nella “*Fossa delle Marianne*”. Anche l’energia in gioco è impressionante e a tali temperature raggiunge il valore

$$E = \frac{5}{2} PV \sim 50 \text{ kJ}. \quad (19.4)$$

Pertanto l’esplosione può proiettare alcune parti della moka a velocità che possono raggiungere le centinaia di metri al secondo. Da queste stime si comprende che l’esplosione deve in realtà avvenire molto prima che il vapore raggiunga temperature vicine ai 600 K . Comunque si può intuire la energia che si sviluppa dentro la caffettiera grazie al riscaldamento eccessivo: è sufficiente non solo per sporcare tutta la cucina con il caffè ma anche per procurare altri guai.

Un altro curioso effetto legato all’esplosione della moka è il seguente: se siete vicino alla caffettiera quando avviene l’esplosione e siete abbastanza fortunati da non essere investiti dalle componenti metalliche scagliate intorno ad alta velocità, sarete comunque, con ogni probabilità, investiti dal getto di vapore e polvere di caffè surriscaldati. La vostra prima reazione in questo caso sarà quella di togliervi al più presto gli indumenti bagnati per evitare di rimanere ustionati. Eppure, un istante dopo vi accorgete, con sorpresa, di avvertire un senso di freddo piuttosto che di caldo. La spiegazione di questo effetto sorprendente è piuttosto semplice. Infatti l’espansione del vapore d’acqua a seguito dell’esplosione è molto rapido e il vapore non ha tempo di scambiare calore con l’ambiente: in altri termini l’espansione del vapore è adiabatica e obbedisce, con buona approssimazione, alla legge delle adiabatiche dei gas perfetti

$$TP^{\frac{1-\gamma}{\gamma}} = \text{cost}, \quad (19.5)$$

dove per l’acqua $\gamma_{H_2O} = 4/3$. In questo caso, assumendo che la temperatura e la pressione iniziale del vapore siano, rispettivamente, pari a 500 K e 10 atm , si vede che il vapore, quando raggiunge la pressione di un’atmosfera, può trovarsi ad una temperatura vicina a 200 K (-73° C).

In conclusione, il caffè preparato nella moka è forte e aromatico. Tuttavia non raggiunge la qualità di un espresso servito in un buon bar. La ragione principale è l'alta temperatura (superiore, sia pure di poco, ai $100^\circ C$) dell'acqua che viene forzata attraverso il filtro dal vapore.

Un accorgimento per migliorare la qualità del caffè preparato con la moka è quello di riscaldare la caffettiera a fuoco basso. In questo modo l'acqua passerà attraverso il filtro lentamente senza che, allo stesso tempo, il vapore nel recipiente inferiore si scaldi eccessivamente. Infine, un buon caffè potrà essere preparato con una moka quando ci si trovi in montagna in un rifugio alpino: la pressione esterna è inferiore ad 1 atm. (per esempio ad una altezza comparabile con quella dell' Everest l'acqua bolle a circa $70^\circ C$) e l'acqua surriscaldata raggiunge una temperatura limitata a $85 - 90^\circ C$ (ottimale per la preparazione di un caffè ricco di aromi).

19.5 La antica caffettiera Napoletana

La caffettiera Napoletana ricorda la moka, con la differenza che il “motore” che muove il liquido attraverso il filtro è la forza di gravità piuttosto che la forza di pressione dovuta al vapore. La caffettiera consiste di due recipienti posti l'uno sopra l'altro e separati da un filtro riempito di caffè (Fig. 19.4). L'acqua nel cilindro inferiore arriva all'ebollizione. A questo punto la caffettiera viene tolta dal fuoco e capovolta.

La filtrazione avviene quindi sotto l'azione della pressione di una colonna d'acqua alta alcuni centimetri (la sovrappressione ΔP sul filtro non raggiunge $10^3 Pa$). Il processo di preparazione del caffè è in questo caso più lento che non nella moka. Si può fare un piccolo esperimento e preparare il caffè nelle due caffettiere. Usando come base la legge di Darcy, secondo la quale il tempo di preparazione del caffè è inversamente proporzionale alla pressione, si può controllare la veridicità di quanto detto. Gli intenditori sostengono che il caffè “alla napoletana” sia migliore di quello fatto con la moka: qui il processo di filtrazione avviene più lentamente e l'aroma del caffè non viene compromesso dal contatto con l'acqua surriscaldata. I ritmi della vita moderna difficilmente lasciano tempo per tranquille discussioni nell'attesa del caffè alla napoletana. La possibilità di lunghi tempi di attesa è rimasta solo nei quadretti di vita napoletana di un tempo e nelle opere di Eduardo De Filippo.



Fig. 19.4: La
Napoletana.

19.6 Espresso

Nel passato non tutti i Napoletani erano così pazienti. Pare che, nel secolo scorso, uno degli abitanti della capitale del Regno delle Due Sicilie, che non poteva aspettare pazientemente il caffè dalla “Napoletana”, abbia convinto un suo amico, ingegnere milanese, a costruire un nuovo tipo di caffettiera, che preparasse un buon caffè aromatico e denso, in meno di mezzo minuto. Oggi una tazza di buon caffè è uno scrigno di segreti riguardanti la crescita e la raccolta dei chicchi di caffè, il modo di preparare la miscela e ovviamente la tostatura e la macinatura, come si è già fatto cenno. Dietro l’arte del caffè sta anche l’uso di una tecnologia sofisticata.

Le macchine per il caffè espresso sono di gran lunga più complesse delle semplici caffettiere illustrate precedentemente. Le macchine per l’espresso di norma si trovano unicamente in bar e ristoranti. Per gli intenditori e gli amanti del caffè esiste una versione casalinga di questa macchina (Fig. 19.5).

Nella caffettiera espresso di tipo professionale l’acqua, ad una temperatura di $90-94^{\circ}\text{C}$, viene spinta sotto una pressione di $9-16\text{ atm}$ attraverso il filtro dove è contenuto un caffè di macinatura fine. Tutto il processo dura $15-30$ secondi e consente di ottenere una o due tazzine di “espresso” di $20-30\text{ ml}$ l’una. Il meccanismo che regola il passaggio del liquido attra-

Fig. 19.5: Versione familiare di una macchina per caffè espresso.



verso il filtro si può spiegare con la legge di Darcy, come nel caso della moka. Però, a differenza della moka, la pressione applicata è dieci volte maggiore, mentre la temperatura è circa $90^{\circ}C$. Il valore della pressione più elevato consente un transito più veloce attraverso la polvere di caffè, inoltre la temperatura più bassa fa in modo che non vengano decomposte alcune componenti instabili che contribuiscono a formare il gusto del caffè. L'espresso, anche se potrebbe sembrare strano, contiene meno caffeina del caffè preparato con la moka, poiché il contatto tra l'acqua e la polvere di caffè avviene per un tempo piuttosto breve (20 – 30 secondi contro i 4 – 5 minuti della moka).

Il primo esemplare di caffettiera espresso fu presentata a Parigi nel 1855. Nelle macchine moderne, che si usano nei bar e ristoranti, l'acqua arriva alla pressione necessaria tramite una pompa. Nella macchina classica, prima si solleva la leva per immettere l'acqua, contenuta nel cilindro di riscaldamento, nel recipiente sopra il filtro. Successivamente la stessa leva viene

abbassata per spingere l'acqua attraverso il filtro. La pressione applicata al liquido è generata quindi dalla forza del braccio, moltiplicata dalla leva.

È interessante discutere su come cade il caffè nella tazza quando esce dal beccuccio della caffettiera espresso. Prima il getto di liquido esce in forma di getto continuo, successivamente diviene sempre più debole fino a gocciolare lentamente. Tentiamo di capire questo fenomeno con l'analogia di quanto si può osservare in montagna quando il sole riscalda la neve sopra un tetto. Anche in questo caso l'acqua può cadere da stalattiti di ghiaccio a gocce oppure a getto continuo. Cerchiamo di valutare quale è il flusso critico perché avvenga questa transizione tra i due tipi di gocciolamento. Supponiamo che l'acqua cada lentamente. È evidente che quando il flusso dell'acqua è troppo piccolo non si forma un getto. In fondo alla stalattite l'acqua forma una goccia, la goccia lentamente cresce, raggiunge una dimensione critica, cade e il processo si ripete. Finché il flusso è basso il processo è quasi statico. Nelle condizioni di equilibrio la goccia si stacca quando il suo peso mg supera la forza della tensione superficiale

$$F_{\sigma} = 2\pi\sigma r, \quad (19.6)$$

dove r è il raggio del collo della goccia.

Il tempo per riempire una goccia evidentemente è

$$t_d = \frac{m}{\rho Q_d}, \quad (19.7)$$

dove Q_d è la massa di fluido che passa in un secondo e ρ è la densità del fluido.

La goccia, sotto l'azione combinata della forza superficiale e del suo peso, si trova in equilibrio. Quando la sua massa arriva al valore critico e la forza dovuta alla tensione superficiale non riesce a compensare la forza di gravità il legame che tiene la goccia aderente alla stalattite si rompe.

Il tempo caratteristico τ in cui ciò avviene si può calcolare ancora utilizzando considerazioni di *scaling*, metodo del quale abbiamo già detto. In questo caso è necessario scrivere una relazione dimensionale che colleghi il tempo di distacco τ con le grandezze fisiche rilevanti per il nostro problema e cioè le dimensioni r della goccia, la tensione superficiale σ e la viscosità η del fluido. Poiché abbiamo già analizzato lo stesso problema al capitolo 6 possiamo subito scrivere

$$\tau = \frac{r\eta}{\sigma}. \quad (19.8)$$

Quindi il cambiamento di regime dal gocciolamento al getto continuo avviene quando ad una goccia che si sta formando si aggiunge altra goccia prima che la precedente sia caduta, cioè

$$t_d \sim \tau \quad \text{e} \quad \frac{m}{\rho Q_d} = \frac{r\eta}{\sigma}. \quad (19.9)$$

Ricavando la massa della goccia dalla condizione di equilibrio tramite la forza di superficie (19.9) ricaviamo infine una espressione significativa:

$$Q_d \sim \frac{2\pi\sigma^2}{\eta\rho g}. \quad (19.10)$$

Operando non con il flusso del volume dell'acqua ma con il flusso della massa Q_d avremmo potuto ottenere immediatamente questa formula dall'analisi dimensionale. Inoltre avremmo potuto ottenere che Q_d non deve dipendere dalle dimensioni della punta della stalattite (infatti la stalattite si scioglie e questo varia le dimensioni della punta della stalattite). Nel caso della caffettiera espresso, invece, la sezione del beccuccio può influire sulla valore critico del flusso. Secondo formula (19.10), però, il flusso critico non può dipendere in maniera marcata dalla sezione del beccuccio.

Capitolo 20

Gelato all'azoto liquido e cucina “molecolare”.

Il lettore avrà a questo punto potuto osservare come in una moderna cucina giochino un ruolo determinante molti fenomeni fisici, senza i quali poco si potrebbe fare per migliorare la qualità dei piatti che adornano le nostre tavole. In effetti, oggi esistono testi, come quello di **Barham** espressamente rivolti alla “*Science of Cooking*”, come si è già detto al capitolo 16.

In questo capitolo vi illustreremo come anche la tecnica stessa e i principi di base del cucinare possono essere “rivoluzionati” per realizzare particolari piatti con finalità specifiche, con nuove proprietà organolettiche che vengono fatte discendere da parametri fisici di carattere globale. Si tratta della gastronomia molecolare, talvolta nota come “fisica culinaria”. In Italia questa applicazione della fisica, forse un po’ eterodossa ma non meno interessante di altre applicazioni che interessano campi interdisciplinari, deve in larga misura la sua notorietà all’impegno di un fisico teorico della Università di Parma, **Davide Cassi**.

Da dove si può far discendere l’origine della cucina molecolare?

In genere i fisici amano lavorare su sistemi “semplici” ove è più facile e diretto lo studio dei fenomeni. Per esempio, l’atomo di idrogeno, il più semplice atomo, con un solo protone nel nucleo e un solo elettrone, è stato studiato per oltre cento anni. Dalla analisi teorica e sulla base di sempre più raffinati esperimenti, quello studio ha prodotto una fantastica rivoluzione nella fisica (con la meccanica quantistica, che descriveremo nella Parte IV), nella chimica, nella filosofia (per le implicazioni associate al carattere non-deterministico delle leggi che regolano le particelle a livello microscopico) e anche nella guerra.

Tuttavia, all'incirca verso la fine degli anni 1980, ad opera soprattutto di **Pierre-Gilles de Gennes** (fisico francese poi insignito del premio Nobel nel 1991) hanno preso l'avvio studi che si proponevano di derivare le proprietà di sistemi complessi sulla base di teorie relativamente "semplici". Un esempio può essere fornito dagli studi sulle soluzioni polimeriche (nelle quali lunghe catene di molte molecole sono diluite in liquidi), oppure nei gel (quale la gelatina usata in cucina o l'umor vitreo del nostro occhio) o ancora più semplicemente un mucchio di sabbia che forma uno stabile, o non stabile, "cono" di granelli. Da quei primi studi siamo oggi pervenuti a una nuova ed interessante disciplina, cosiddetta "fisica della complessità" e in tale ambito alcuni scienziati si dedicano per esempio alla evoluzione delle popolazioni o dei sistemi economici o alle quotazioni borsistiche (la *econofisica*). Altri fisici, appunto hanno iniziato ad applicare l'approccio ai sistemi complessi nella "gastronomia molecolare".

Il principio di base può essere individuato come segue. La cottura dei cibi non è altro che la modificazione della natura delle proteine, cioè di lunghe molecole che si rompono, si ri-aggregano in fasi diverse, come abbiamo visto per l'uovo o come accade nella crosta dell'arrosto. Allora, sulla base della fisica dei sistemi macromolecolari disordinati, possiamo tentare di individuare i parametri fisici e chimici che controllano il sistema e sviluppare nuovi metodi di manipolazione di quelle "macromolecole" che sono le proteine. In altre parole, si attaccano scientificamente i fenomeni che avvengono nei processi di preparazione dei cibi, quei meccanismi che portano a una pietanza prelibata.

Un po' esagerando, si potrebbe dire che il nuovo strumento concettuale di questa particolare applicazione della fisica è la termodinamica statistica di sistemi molecolari complessi e disordinati.

I principi base della cucina molecolare sono la valorizzazione di ingredienti e di materie prime, l'ampliamento (non la eliminazione!) della tradizione gastronomica, la attenzione ai valori nutrizionali per la creazione di piatti buoni e insieme dietetici; la realizzazione di nuove "architetture" alimentari attraverso lo studio delle proprietà chimico-fisiche degli alimenti.

Gli aspetti innovativi della gastronomia molecolare, per indicarne solo alcuni, sono le nuove frontiere per le intolleranze alimentari o le sindromi allergiche (il glutine per esempio, è sostituito con amido nei gnocchi senza la farina), la lecitina di soia impiegata in luogo dell'uovo per la limitazione del colesterolo e anche piatti più saporiti per le persone costrette a diete particolari. Si è scoperto, per esempio, che l'alcool coagula le proteine

dell'uovo senza alterarne il sapore e quindi che si può ottenere un uovo con i sapori del crudo "cuocendolo" in alcool. Il pesce può essere fritto in una miscela di zuccheri fusi anziché in olio, con dimezzamento del tempo di cottura e preservandone densità e viscosità (per evitare il sapore di dolce si avvolga il pesce in una foglia di porro).

Si possono creare meringhe avvolte in cioccolato vetroso ghiacciato, o *mousse* al cioccolato senza colesterolo o maionese senza uova e altre ricette che potrete trovare nei molti testi, di diversi autori, ormai in vendita nelle migliori librerie.

Qui noi ci limiteremo a descrivervi la preparazione di un ottimo gelato, con un metodo innovativo che ormai viene pubblicizzato anche nei quotidiani e che utilizza un "oggetto" fisico al quale abbiamo fatto spesso, e ancora faremo, riferimento: l'azoto liquido.

Innanzitutto: che cosa è il gelato e perché occorre usualmente la gelatiera, della quale invece vi insegneremo a fare a meno? A parte la crema o il sorbetto il gelato è composto da microcristalli di ghiaccio prodotti dall'abbassamento di temperatura nella gelatiera, che sono miscelati a bolle d'aria per evitarne il compattamento e rendere il gelato stesso mantecato. Lo zucchero aggiunto serve anche a ridurre le dimensioni dei cristalli di ghiaccio (nel contempo aumenta il contenuto calorico!). La rotazione dell'asse della gelatiera, con la pala, è necessaria da una parte per frantumare il ghiaccio dall'altra per evitare appunto il compattamento: la procedura per arrivare a gustare il prodotto è piuttosto lunga e le gelatiere per famiglia spesso non danno risultati soddisfacenti.

Tuttavia, ecco che una volta capiti i fenomeni fisici possiamo produrre per altra via un gelato migliore, a più basso contenuto calorico, istantaneo, più cremoso e soffice con l'ausilio della gastronomia molecolare, sfruttando la evaporazione di azoto alla temperatura di ebollizione (circa 77 K , a pressione ordinaria). La rapidissima evaporazione, una transizione di fase del primo ordine, del tutto simile alla evaporazione che si produrrebbe in acqua gettata in una fornace (e che un po' più lentamente avviene anche alla ordinaria ebollizione nella pentola della massaia), mentre determina istantaneamente microcristalli della crema medesima, provoca anche la formazione di bolle. Se ritornate a rileggere le sezioni su bolle e gocce, sulla evaporazione nella teiera, sullo spumante e il "fizz", non avrete difficoltà a comprendere cosa avviene all'atto della ebollizione dell'azoto quando dalla temperatura di 77 K viene a contatto con la vostra crema alla temperatura ambiente.

Basta disporre di un recipiente metallico (che per sua natura può sopportare lo shock termico), preparare la crema di proprio gradimento, aggiungere circa la metà del volume del preparato di azoto liquido e successivamente una altra eguale quantità, mentre con una certa energia si mescola o si sbatte il gelato con un cucchiaino di legno. E tutto è pronto per essere servito! (Fig.20.1).



Fig. 20.1: Il gelato all'azoto liquido

Il gelato non conterrà acqua o ghiaccio, sarà “naturalmente” soffice e mantecato, meno “freddo” e meno ricco delle calorie dello zucchero. Provatelo e vi ricorrete molto spesso!

Le bolle di azoto (N_2) contenute nella crema o nel sorbetto sono innocue, questo gas essendo presente nell'atmosfera che respiriamo (per il 78% in volume) e diverse industrie forniscono 10 o 15 litri di azoto a un prezzo molto contenuto, in dewar nei quali si conserva liquido per alcuni giorni.

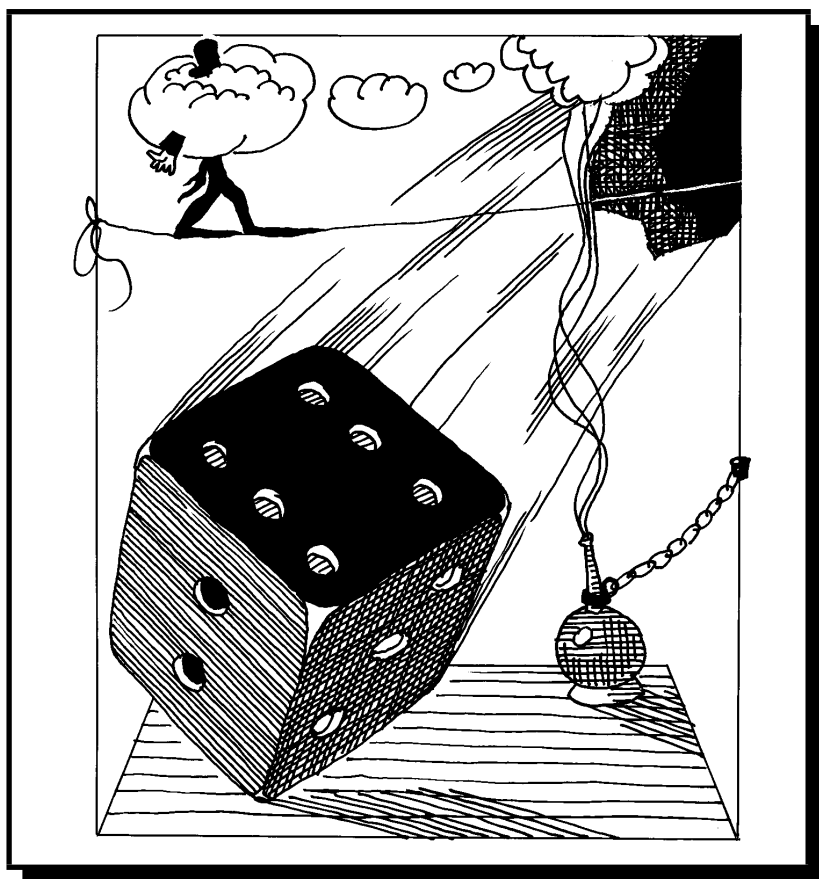
Le precauzioni da osservare sono solo quelle tipiche per la manipolazione di liquidi criogenici: evitare di scottarsi introducendo le mani all'interno del liquido (meglio usare guanti) e evitare che spruzzi possano raggiungere gli occhi (meglio munirsi di comuni occhiali).

PARTE IV

Finestre sul mondo dei quanti

In questa parte del testo descriveremo le leggi fondamentali che regolano il comportamento delle particelle costituenti la materia a livello microscopico. In particolare, discuteremo alcuni insoliti “super-fenomeni” in cui tali particelle sono coinvolte a temperature molto basse. Non è semplice descrivere questi aspetti della fisica, in quanto un discorso compiutamente soddisfacente si può fare solo con la meccanica quantistica. In realtà tale meccanica regola anche gli eventi intorno a noi, benché sia difficile accorgersene in ragione delle piccolissime modificazioni che gli effetti quantici producono nel macroscopico. Invece tale approccio è indispensabile nel microscopico e vedrete quanto insolito e meraviglioso è il mondo dei quanti!

Dal principio di esclusione dedurremo le proprietà fondamentali degli atomi, le oscillazioni di punto zero dei nuclei attorno alle posizioni di equilibrio in un solido, alcune incredibili proprietà dell’elio liquido. Il quanto più “elusivo” e di difficile materializzazione (il quanto di flusso, o “flussone”) caratterizza uno degli aspetti più affascinanti della materia: lo stato superconduttivo, del quale tratteremo lo sviluppo storico sino ai superconduttori ad alta temperatura. Quindi si descriverà lo SQUID, o magnetometro quantistico, capace di registrare campi magnetici incommensurabilmente piccoli, sino a quelli prodotti dal nostro cervello sottoposto a stimoli sonori o visivi e che apre quindi la strada alla possibile “visione” dei nostri pensieri. Con i segnali di risonanza magnetica nucleare (NMR) è già possibile “vedere” realmente all’interno del nostro corpo e breve cenno verrà fatto a questa tecnica. Infine vi illustreremo i principi della nano-scienza e dei nano-dispositivi.



Capitolo 21

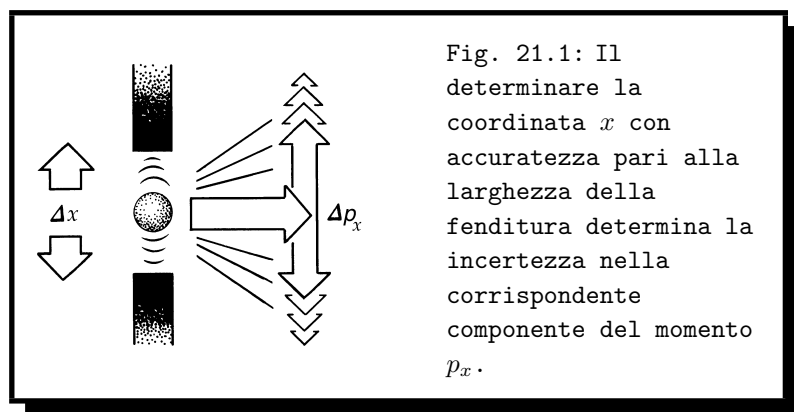
Il principio di indeterminazione.

Il fisico tedesco **Werner Heisenberg**, nel 1927, formulò il principio di indeterminazione: se nello studio di un oggetto in movimento riusciamo a definire il suo impulso con un errore Δp_x , potremo contemporaneamente definire la relativa coordinata x solo con una indeterminazione almeno pari a $\Delta x \approx \hbar / \Delta p_x$, dove $\hbar = 1.054 \cdot 10^{-34} J \cdot s$ è la costante di **Planck**.

Inizialmente questa relazione può suscitare dei dubbi. Le leggi di Newton, che si studiano alla scuola media e alle quali abbiamo più volte fatto ricorso, permettono di trovare la legge del moto di un oggetto, ossia di calcolare la dipendenza delle coordinate dal tempo. Conoscendo la legge del moto si possono trovare le componenti della velocità $\vec{v}$ (derivate delle coordinate rispetto al tempo) e dell'impulso $\vec{p} = m\vec{v}$. Ne consegue che nella fisica classica si ritiene di poter definire contemporaneamente sia la coordinata sia l'impulso e non vi è alcuna indeterminazione. Effettivamente nel mondo macroscopico in pratica tutto avviene così, mentre nel micromondo le cose stanno diversamente (Si veda la illustrazione nella figura 21.1).

21.1 Impulso e coordinata

Immaginiamo di analizzare il moto di un elettrone. Come possiamo fare? L'occhio non è adatto a questo scopo poiché la sua capacità di risoluzione è insufficiente. Allora guardiamo l'elettrone al microscopio. La risoluzione del microscopio è determinata dalla lunghezza dell'onda della luce con cui si fa la rilevazione. Per una normale luce visibile, questa grandezza è dell'ordine di 100 nm (10^{-7} m) e le particelle di dimensioni inferiori non sono visibili. Gli atomi hanno dimensioni dell'ordine di 10^{-10} m , così che è inutile sperare



di rilevarli: ancora meno è possibile l'osservazione dei singoli elettroni.

Tuttavia, immaginiamo di poterlo fare. Supponiamo di essere riusciti a costruire un microscopio che utilizzi onde elettromagnetiche a lunghezza d'onda più piccola: per esempio, raggi X o onde γ . Quanto più elevata in energia è la radiazione che utilizzeremo, ossia quanto più corta è la lunghezza d'onda corrispondente, tanto più piccoli potranno essere gli oggetti da esaminare. Sembrerebbe che questo γ -microscopio da noi immaginato sia lo strumento ideale. Grazie ad esso si potrebbe misurare la coordinata dell'elettrone con la precisione che si desidera. Come interviene allora il principio di indeterminazione?

Consideriamo più in dettaglio questo esperimento immaginario. Per ottenere informazioni sulla posizione dell'elettrone, il quanto di radiazione elettromagnetica deve dunque avere un'energia pari almeno a un γ -quanto, cioè l'energia minima $\hbar\omega$, ove ω è la frequenza angolare. Quanto minore è la lunghezza dell'onda, tanto maggiore è la quantità d'energia trasportata dal quanto. L'impulso del quanto è proporzionale alla sua energia. Nell'urto con l'elettrone, il γ -quanto trasmette parte del suo impulso. Quindi, misurando la coordinata produciamo comunque una indeterminazione dell'impulso dell'elettrone: quanto più precisa vogliamo rendere la misura della coordinata, tanto maggiore sarà l'indeterminazione nell'impulso.

Un'analisi approfondita di questo processo mostra che il prodotto delle due indeterminazioni non può mai essere reso inferiore alla costante di Planck.

Si può avere l'impressione di aver esaminato soltanto un caso particolare, di avere cioè utilizzato un "cattivo" strumento per la misura della

coordinata e che si possano fare misure molto più precise, senza “toccare” l’elettrone e senza cambiare il suo impulso. Non è così. I più grandi scienziati (tra i quali anche Einstein) hanno cercato di realizzare, almeno idealmente, una apparecchiatura che potesse misurare la coordinata di una particella e il suo impulso contemporaneamente, con una precisione più elevata di quella permessa dal principio di indeterminazione. Nessuno, però, è riuscito a farlo. Semplicemente, *non si può fare*. È una legge della natura.

Tutto ciò può sembrare alquanto nebuloso ed è difficile crearsi immediatamente un modello mentale soddisfacente. Una reale comprensione si ha soltanto dallo studio approfondito della meccanica quantistica. Per un primo approccio, tuttavia, quanto detto si può ritenere che illustri in maniera adeguata la complessa realtà fisica del microscopico.

Per capire dove si trovi il limite tra macromondo e micromondo, tentiamo una piccola stima. Per esempio, nelle esperienze sul moto di Brown si utilizzano particelle molto piccole, della dimensione di circa $1\ \mu m$ ($10^{-6}\ m$) e una massa di $10^{-10}\ g$. Benché si tratti di piccole quantità, tuttavia esse contengono un’enorme quantità di atomi. Dal principio di indeterminazione in questo caso abbiamo $\Delta v_x \Delta x \sim \hbar / m \sim 10^{-21}\ m^2 / s$.

Se, per esempio, si vuole definire la posizione di una particella con una precisione del centesimo della sua dimensione ($\Delta x \sim 10^{-8}\ m$) avremo $\Delta v_x \sim 10^{-13}\ m / s$. Si ottiene cioè una indeterminazione sulla velocità molto piccola e ciò in ragione del piccolo valore della costante di Planck. La velocità del moto di **Brown** di questa particella è all’incirca uguale a $10^{-6}\ m / s$. Come si vede l’errore sulla velocità, legato al principio di indeterminazione, è trascurabilmente piccolo (dell’ordine del decimo di milione!) anche per una così piccola frazione di materia.

È pertanto evidente che tale principio è trascurabile per corpi macroscopici. Se invece diminuiamo la massa della particella (consideriamo, per esempio, un elettrone) e aumentiamo la precisione nella definizione della coordinata ($\Delta x \sim 10^{-10}\ m$ è la tipica dimensione atomica) la indeterminazione sulla velocità diventa paragonabile alla velocità stessa.

Nello studio degli elettroni nell’atomo il principio di indeterminazione comporta effetti così marcati che non si può non tenerne conto. E le conseguenze sono sorprendenti.

21.2 Onde di probabilità di presenza

Nel modello più semplice di atomo – il modello di **Ernest Rutherford** – gli elettroni ruotano su orbite intorno al nucleo, analogamente a come i pianeti si muovono attorno al Sole. Tuttavia gli elettroni sono particelle cariche e durante la loro rotazione creano campi elettromagnetici variabili, cioè emettono radiazione: dunque perdono energia.

Ecco perché nel modello planetario dell'atomo agli elettroni spetterebbe la sorte di cadere sul nucleo e l'atomo si disintegrerebbe. Invece la stabilità degli atomi è un fatto sperimentale assolutamente certo. Era necessario, quindi, “correggere” il modello di Rutherford: ci ha pensato per primo **Niels Bohr** nel 1913.

Gli elettroni, nel modello di Bohr, possono ruotare soltanto su orbite definite, muovendosi sulle quali conservano energie rigorosamente invariate. Gli elettroni possono cambiare la loro energia solo saltando da un'orbita ad un'altra, emettendo e assorbendo quanti di radiazione in questa transizione. Tale comportamento “quantistico” degli elettroni permette di spiegare molte cose, in particolare la stabilità dell'atomo e gli spettri atomici. Questo modello ancora oggi è utilizzato per una spiegazione approssimata dei fenomeni quantistici. Tuttavia esso contraddice il principio di indeterminazione! Infatti su una tale orbita l'impulso e la coordinata possono essere definiti contemporaneamente, mentre nel micromondo, come ora sappiamo, ciò non può mai avvenire. Si è dovuto così ulteriormente correggere anche questo modello cosiddetto quasi-classico. Il quadro reale che descrive il comportamento dell'elettrone nell'atomo è molto più complesso.

Immaginate di essere riusciti a definire la posizione dell'elettrone ad un certo istante di tempo. Possiamo dire con precisione dove sarà in un momento successivo (diciamo, dopo un secondo)? No, la misura delle coordinate, come sappiamo, comporta sempre una indeterminazione nell'impulso dell'elettrone e anche utilizzando i migliori strumenti non si può predire con precisione dove l'elettrone si troverà.

Che cosa ci resta da fare in questo caso? Evidenziamo con un punto la posizione nello spazio dove abbiamo per così dire “rilevato” l'elettrone. Di nuovo indichiamo con un punto il risultato di un'altra misura delle coordinate del medesimo elettrone atomico. Ancora una nuova misura, di nuovo un punto e così via. Si vede che, anche se non si può predire in alcun modo dove si troverà con precisione il punto che individua la posizione dell'elettrone ad un istante successivo, il modo in cui si dispongono questi

punti nello spazio presenta una certa regolarità. In alcune regioni i punti sono più fitti, in altri più radi, indicando rispettivamente la maggiore o minore probabilità di trovare l'elettrone in quelle regioni.

Siamo stati costretti a rinunciare a una descrizione deterministica del movimento dell'elettrone, ma siamo in grado di predire la probabilità di trovarlo in diversi punti dello spazio. Il comportamento dell'elettrone nel micromondo viene descritto probabilisticamente! Questa descrizione sul comportamento delle particelle nel micromondo potrà non piacere, è insolita e contraddice fortemente la nostra intuizione e la nostra esperienza. Tuttavia non si può fare nulla di diverso, essendo la natura costruita in questo modo. Nel micromondo sussistono altre leggi, diverse da quelle alle quali siamo abituati nella vita di ogni giorno. Secondo l'espressione metaforica di Einstein, è come se si dovesse giocare a dadi per predire il comportamento degli elettroni. Non si può fare altrimenti^a. Nel micromondo lo stato dell'elettrone è descrivibile solo in termini della probabilità di trovarlo in diversi punti dello spazio. Nella nostra visualizzazione la misura di tale probabilità sarà la densità dei punti e si può immaginare che essi formino qualcosa di simile ad una nube, che definisce l'immagine statistico-probabilistica del comportamento dell'elettrone.

Come sono strutturate le nuvole di probabilità? Analogamente a come nella meccanica classica le leggi di Newton definiscono il moto dei corpi, così nella meccanica quantistica si scrive un'equazione, dalla cui soluzione si può trovare la "disposizione" dell'elettrone nello spazio. Questa equazione è stata concepita nel 1925 dal fisico austriaco **Ervin Schrödinger** (notate che ciò è avvenuto prima che fosse scoperto il principio di indeterminazione, ossia prima che fosse spiegata la causa che impedisce la descrizione deterministica; simili cose accadono relativamente spesso nella fisica!).

L'equazione di Schroedinger descrive quantitativamente e dettagliatamente i fenomeni atomici. Non è possibile approfondire tale aspetto senza un'analisi matematica complessa. Qui riporteremo semplicemente i risultati di tale studio e mostreremo soltanto dei casi particolari che chiariscono il comportamento probabilistico dell'elettrone nello spazio.

Nella figura 21.2 è mostrato lo schema di un esperimento sulla diffrazione degli elettroni e la fotografia dello spettro di righe che compare sullo schermo. Il quadro è completamente analogo a quello che si ha in

^aEinstein stesso non credeva alla necessità assoluta di fare ricorso ad aspetti probabilistici. Sino alla fine dei suoi giorni non accettò mai del tutto la meccanica quantistica.

caso di diffrazione della luce. Se si considera che gli elettroni si muovono lungo traiettorie rettilinee, così come devono fare secondo le leggi della fisica classica, questo esperimento non può essere spiegato.

Se invece essi si considerano “diffusi” nello spazio, il risultato dell’ esperimento diviene comprensibile. Oltre a ciò, dall’esperimento si può prevedere che la nube di probabilità di presenza possiede proprietà ondulatorie. Nel posto in cui l’ampiezza dell’onda è massima, si ha la massima probabilità che il fenomeno si verifichi.

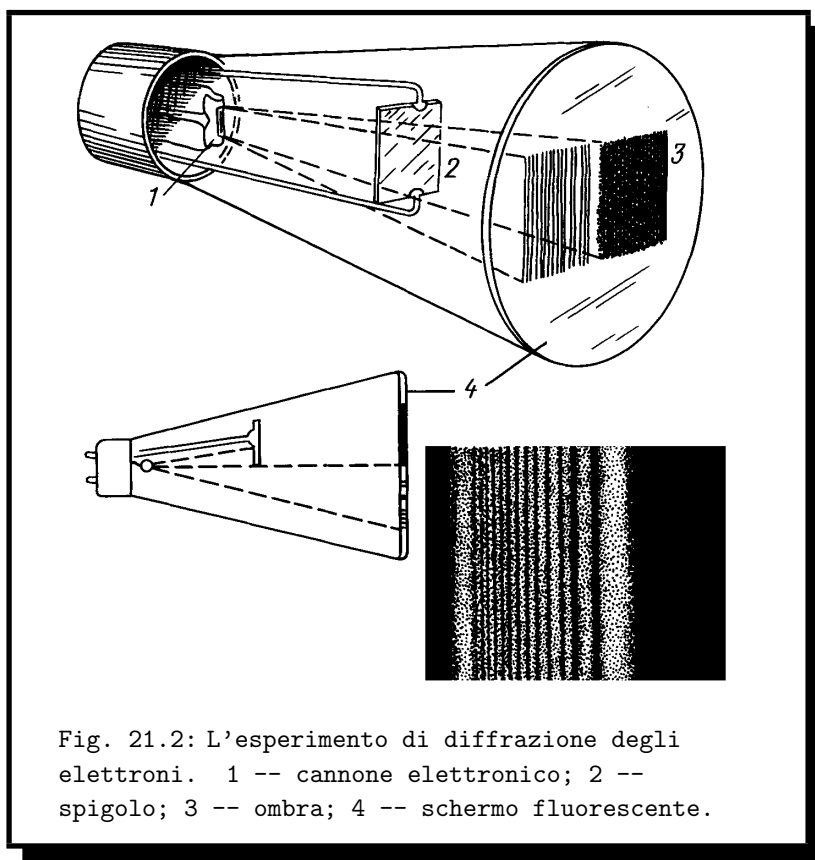


Fig. 21.2: L’esperimento di diffrazione degli elettroni. 1 -- cannone elettronico; 2 -- spigolo; 3 -- ombra; 4 -- schermo fluorescente.

La figura 21.3 evidenzia la distribuzione di probabilità di presenza dell’elettrone nell’atomo di idrogeno, per alcuni degli stati quantici nei quali esso si può trovare. Naturalmente non si tratta di fotografie di elet-

troni reali, ma di risultati di calcoli. La simmetria che caratterizza i diversi stati controlla la simmetria delle molecole e dei cristalli. Si può anche dire che queste distribuzioni costituiscono la chiave per la comprensione della bellezza delle forme regolari della natura.

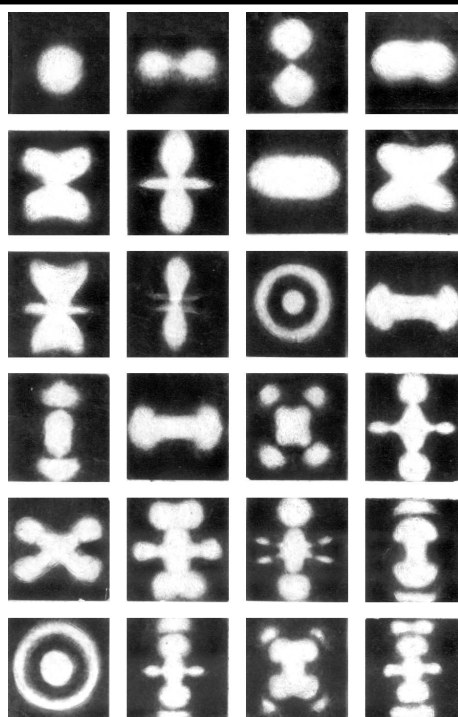


Fig. 21.3: La distribuzione probabilistica dell'elettrone in forma di nube di probabilità di presenza nei diversi stati dell'atomo di idrogeno.

Nel nostro caso quindi, le zone luminose sono quelle in cui si può trovare un elettrone con maggior probabilità. Questa distribuzione di probabilità di presenza è l'analogo delle orbite quantiche, lungo le quali si doveva muovere l'elettrone nel modello dell'atomo di Bohr. Quanto più chiara è la regione, tanto più è probabile incontrarvi un elettrone. Le figure ci ricordano le onde stazionarie, che compaiono quando un processo ondulatorio

avviene in una regione limitata dello spazio.

Quali forme sorprendenti possono assumere le nubi di probabilità! E queste figure astratte effettivamente rappresentano il comportamento degli elettroni nell'atomo e spiegano, per esempio, i livelli discreti d'energia e tutto ciò che attiene al legame chimico.

Senza entrare nel dettaglio della struttura delle nubi di probabilità, attraverso il principio di indeterminazione si può valutare la loro dimensione caratteristica. Se la nube di probabilità di presenza ha una dimensione dell'ordine di Δx , l'indeterminazione nella coordinata della particella risulta dell'ordine di Δx . Di conseguenza, l'indeterminazione sull'impulso della particella non può essere inferiore a $\Delta p_x \sim \hbar / \Delta x$. Come ordine di grandezza questa espressione definisce anche il valore minimo dell'impulso della particella.

Quanto minore è la dimensione della nube di probabilità, tanto maggiore diventa l'impulso e di conseguenza tanto più rapidamente si muove la particella all'interno della regione in cui è localizzata. Queste considerazioni generali sono sufficienti per valutare correttamente la dimensione dell'atomo. L'elettrone nell'atomo possiede energia cinetica e potenziale. L'energia cinetica dell'elettrone è l'energia del suo moto. Essa è legata all'impulso in accordo alla espressione $E_c = m v^2 / 2 = p^2 / 2m$.

L'energia potenziale dell'elettrone è l'energia dell'interazione coulombiana con il nucleo. Anche per essa sussiste una formula semplice: $E_p = -e^2/r$, dove il segno “-” corrisponde al fatto che l'energia è attrattiva (e è la carica dell'elettrone, r la distanza dall'elettrone dal nucleo). In ogni stato l'elettrone ha un determinato valore di energia totale $E = E_c + E_p$. Lo stato che comporta energia minima si chiama stato fondamentale (non eccitato). Calcoliamo la dimensione dell'atomo nello stato fondamentale. Supponiamo che l'elettrone sia distribuito su una regione di dimensione dell'ordine di r_0 . L'attrazione dell'elettrone verso il nucleo tende a far diminuire r_0 , a far “contrarre” la nube di probabilità. A ciò corrisponde la riduzione dell'energia potenziale del sistema, che è dell'ordine di $-e^2/r_0$ (diminuendo r_0 aumenta, in modulo, tale grandezza).

Se l'elettrone non possedesse energia cinetica cadrebbe sul nucleo. Tuttavia, come già sapete, si ha inevitabilmente energia cinetica in una particella confinata, proprio per il principio di indeterminazione. Esso, infatti, “impedisce” all'elettrone di cadere sul nucleo! Diminuendo r_0 aumenta l'impulso minimo della particella $p_0 \sim \hbar / r_0$ di conseguenza cresce l'energia cinetica: $E_c \sim \hbar^2 / 2m r_0^2$.

Dalla condizione di zero sulla derivata dell'energia totale otteniamo che il minimo corrisponde al valore che definisce la dimensione caratteristica della regione di localizzazione dell'elettrone nello stato fondamentale, ossia la dimensione dell'atomo:

$$r_0 \sim \frac{4\pi\epsilon_0 \hbar^2}{m e^2}. \quad (21.1)$$

La grandezza r_0 è uguale all'incirca a 0.05 nm ($5 \cdot 10^{-11} \text{ m}$). Effettivamente le dimensioni degli atomi sono di questo ordine di grandezza. È ora evidente che il principio di indeterminazione, che permette di valutare correttamente la dimensione degli atomi, è una delle leggi fondamentali del mondo atomico.

Per atomi a più elettroni esiste una particolare regolarità e anch'essa deriva direttamente dal principio di indeterminazione. Dal punto di vista sperimentale il lavoro che bisogna compiere per allontanare l'elettrone dall'atomo è definito con precisione (esso è uguale all'energia di ionizzazione E_i). Per i diversi atomi il prodotto $\sqrt{E_i}$ per la dimensione d dell'atomo è lo stesso, all'interno di un errore del $10 - 20\%$.

Il lettore, probabilmente, ha intuito di cosa si tratta: l'impulso dell'elettrone è $p \sim \sqrt{2mE}$ e la costanza del prodotto $p \cdot d \sim \hbar$ deriva dal principio di indeterminazione.

21.3 Oscillazioni di punto zero

Risultati molto interessanti si ottengono considerando le oscillazioni degli atomi nei solidi con l'aiuto del principio di indeterminazione. Gli atomi (o gli ioni) compiono delle oscillazioni intorno alle loro posizioni di equilibrio nella struttura cristallina. In genere si può ritenere che queste oscillazioni siano legate al moto termico degli atomi: quanto più alta è la temperatura, tanto più pronunciate sono le oscillazioni. Cosa accade se si abbassa la temperatura? Dal punto di vista classico, l'ampiezza delle oscillazioni deve diminuire e allo zero assoluto gli atomi si dovrebbero fermare del tutto. Secondo le leggi quantistiche ciò è possibile? Diminuzione dell'ampiezza delle oscillazioni significa, nel linguaggio quantistico, la riduzione della dimensione della nube di probabilità di presenza (cioè della regione di localizzazione della particella). Peraltro, come sappiamo, in ragione del principio di indeterminazione questo comporta un'aumento dell'impulso della particella e pertanto essa non può fermarsi.

Infatti, anche allo zero assoluto, gli atomi in qualsiasi corpo compiono delle oscillazioni. Si chiamano *oscillazioni di punto zero* e si manifestano in tutta una serie di affascinanti effetti fisici.

Tentiamo, prima di tutto, di calcolare l'energia delle oscillazioni di punto zero. In un sistema oscillante, per piccole oscillazioni, con deviazione x dalla posizione di equilibrio la forza di richiamo è $F = -kx$ (nel caso della molla k è la sua costante elastica; per un atomo facente parte di un solido, la costante è determinata dalle forze di interazione interatomica). Corrispondentemente, al sistema va associata l'energia potenziale

$$E_p = \frac{kx^2}{2} = \frac{m\omega^2 x^2}{2},$$

dove $\omega = \sqrt{k/m}$ è la pulsazione delle oscillazioni. Di conseguenza l'ampiezza delle oscillazioni $x_{\max}$ è legata all'energia immagazzinata dal sistema dalla relazione

$$E = \frac{m\omega^2 x_{\max}^2}{2}; \quad x_{\max} = \sqrt{\frac{2E}{m\omega^2}}.$$

L'ampiezza delle oscillazioni, in meccanica quantistica, definisce appunto la dimensione caratteristica della regione di localizzazione della particella, dimensione che per il principio di indeterminazione è legata all'impulso minimo. Quanto minore è l'energia delle oscillazioni, tanto minore deve essere l'ampiezza. Da altra parte, la diminuzione dell'ampiezza implica aumento dell'impulso e di conseguenza anche dell'energia della particella. L'energia minima che può possedere una particella è quindi definita dalla relazione

$$E_0 \sim \frac{p_0^2}{2m} \sim \frac{\hbar^2}{m x_0^2} \sim \frac{\hbar^2}{m} \cdot \frac{m\omega^2}{E_0}.$$

Dal confronto tra primo e ultimo membro ricaviamo $E_0 \sim \hbar\omega$. L'energia di punto zero è più grande per atomi leggeri, i quali oscillano con frequenza più alta.

L'evidenza più manifesta delle oscillazioni di punto zero è, probabilmente, l'esistenza di un liquido che non solidifica mai, anche a temperatura vicinissima allo zero assoluto (sappiate che con la tecnica di smagnetizzazione adiabatica è stato già possibile scendere alla temperatura di $10^{-7} - 10^{-8}$ gradi Kelvin!).

È intuitivo che il liquido non solidifica se l'energia cinetica delle oscillazione degli atomi è sufficiente a vincere le forze di legame associate alla struttura cristallina. In questo caso non è importante la causa originaria di tale energia cinetica, vale a dire se sia essa legata al moto termico o alle oscillazioni quantistiche di punto zero. I candidati più probabili a mantenersi liquidi a tutte le temperature sono l'idrogeno e l'elio. L'elio è anche un gas inerte. I suoi atomi interagiscono l'un l'altro molto debolmente, così che è più facile "liquefare" il solido. L'energia delle oscillazioni di punto zero nell'elio è sufficiente per ottenere questo risultato anche alle temperature prossime allo zero assoluto. Mentre l'idrogeno, sebbene i suoi atomi individuali potrebbero possedere un'energia delle oscillazioni di punto zero più alta che per l'elio, può ugualmente solidificare poiché la formazione di idrogeno molecolare aumenta le forze di legame nel solido. Tutte le altre sostanze solidificano vicino allo zero assoluto: l'elio è l'unica sostanza che a pressione normale rimane liquido. Si può anche dire che proprio il principio di indeterminazione non ne permette la solidificazione. I fisici chiamano l'elio *liquido quantistico*.

Esso possiede una altra proprietà fonte di meraviglie (che vedremo più oltre): la superfluidità. Anche questa proprietà rappresenta un fenomeno di carattere quantistico in un sistema macroscopico! **Lev Landau** diceva che l'elio liquido è la finestra che l'uomo può aprire sul mondo quantistico.

Solo ad una pressione di circa 25 atmosfere l'elio solidifica. L'elio solido non è propriamente un cristallo. In esso le oscillazioni di punto zero determinano, per esempio, un valore di energia cinetica degli atomi vicino a quello che comporterebbe il passaggio allo stato liquido. In conseguenza di ciò, a una pressione di circa 29 atmosfere a temperatura all'incirca di $2K$ la superficie del cristallo può compiere una specie di oscillazioni molto ampie, come si verifica alla superficie di separazione tra due liquidi che non miscelano (Fig.21.4). I fisici hanno chiamato l'elio solido *cristallo quantistico* e le sue proprietà sono oggetto di studi approfonditi. Torneremo più oltre sull'elio e su alcune sue stupefacenti proprietà.

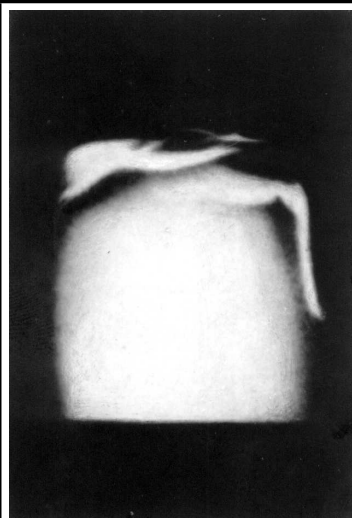


Fig. 21.4:
Oscillazioni di
interfaccia tra la
superficie dell'elio
solido (regione
brillante) e il bagno
di elio liquido.

Capitolo 22

Fenomeni meraviglianti.

Come abbiamo visto al capitolo precedente, gli elettroni si “dispongono” negli atomi obbedendo a delle leggi di carattere probabilistico totalmente diverse da quelle deterministiche che regolano gli eventi del mondo attorno a noi, quali siamo abituati a percepire quotidianamente.

Richiamate alla vostra attenzione la “fotografia” che rappresenta la ricostruzione della “nube” di probabilità di presenza di un elettrone nei suoi diversi stati quantici attorno al nucleo protonico positivo, nell’atomo di idrogeno (Fig. 21.3). Sono evidenti l’armonia e la bellezza di queste figure, nelle quali si possono cogliere la poesia del mondo quantico, delle forme regolari della natura. La simmetria che caratterizza il moto dell’elettrone negli atomi controlla anche la simmetria della materia condensata nella forma di molecole o di cristalli.

A questo proposito richiamiamo due considerazioni, una di **Einstein** e l’altra di **Heisenberg**, che bene racchiudono l’essenza della “meraviglia” anche nel “microscopico”.

“Parlando di semplicità e di bellezza (della natura) faccio della verità una questione estetica... ammiro moltissimo la semplicità e la bellezza dei modelli matematici che la natura ci offre”.

“I rapporti interni (della teoria quantistica negli atomi) mostrano nella loro astrazione matematica un grado incredibile di semplicità e bellezza, un dono che possiamo solo accettare con umiltà. Neppure Platone avrebbe potuto credere che essi fossero così belli. Essi non possono essere inventati, ma esistono dalla creazione del mondo”.

Sollecitati da queste considerazioni, nel quadro della fisica classica vogliamo ora accennare ad un esperimento che evidenzia l’evoluzione per così dire

“naturale” verso la armonia e verso il bello. Successivamente descriveremo una molecola di recente scoperta che racchiude in sé particolari elementi di simmetria e di bellezza architettonica.

22.1 Granelli vibranti

È stata condotta pochi anni or sono la seguente esperienza^a. Su una lastra rigida sono depositi molti granelli di sabbia, a caso, senza particolari disposizioni e quindi senza alcun apparente disegno. La lastra è stata quindi posta in oscillazione rigida lungo un tratto dell'asse verticale.

Per il generico granello di sabbia, che si assume a contatto nel piano della lastra con quattro granelli adiacenti, viene scritta la seguente equazione (che rappresenta, pur nella sua forma piuttosto complessa, una già notevole semplificazione):

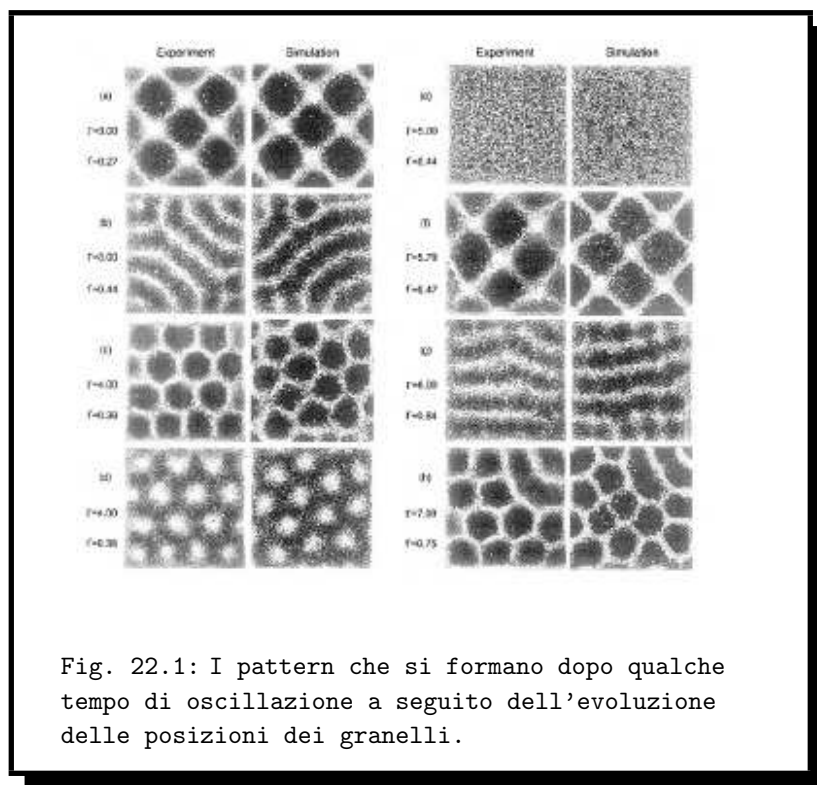
$$\frac{d^2 \vec{r}_P}{dt^2} + \gamma \frac{d \vec{r}_P}{dt} + \omega_0^2 \vec{r}_P = \vec{r}_{P+1} - 2\vec{r}_P + \vec{r}_{P-1} + \lambda \vec{r}_P r_P^2 + \vec{\epsilon} \cos \omega t$$

ove $\vec{r}_P$ è il vettore che individua istantaneamente la posizione del granello nel piano della lastra. A primo membro si riconoscono il *termine inerziale* che determina la accelerazione, il termine di *frenamento* (con coefficiente γ) e il termine di *richiamo* a frequenza propria ω_0 . A secondo membro compare una espressione approssimata dell'*interazione* tra granelli, un termine di *non-linearità* (con parametro λ) e il termine di oscillazione impressa alla piastra (rigida) a frequenza angolare ω .

La soluzione di regime (per tempi lunghi) di tale equazione fornirà la disposizione sulla lastra degli $\vec{r}_P$, cioè descriverà la formazione dell'immagine (il *pattern*) che a seguito di un fenomeno collettivo si insedia sulla piastra. L'esperimento permetterà di confermare o meno le assunzioni di base nello scrivere la suddetta equazione e la validità della nostra soluzione. Per ottenere la soluzione occorre una computazione numerica (con trasformazione di coordinate, ricorso alla equazione di Mathieu e alle risonanze parametriche e a vari tecnicismi che qui non interessano).

Nella figura 22.1 sono riportate i *pattern* che si formano dopo qualche tempo di oscillazione a seguito dell'evoluzione delle posizioni dei granelli, così come i risultati della simulazione computazionale.

^aC.Bizon et al., Physical Review Letters **80**, 57 (1998).



L'accorpamento dei granelli in forme di simmetria, armonia e "bellezza" è stupefacente. Appare che la natura, obbedendo a equazioni che l'uomo sa scrivere solo in forma approssimata e non sa quasi risolvere, evolve verso l'ordine e l'armonia.

E l'umanità? Si può ipotizzare qualcosa di analogo? Anche per noi, ancorché le nostre azioni, le nostre emozioni e i nostri sentimenti siano controllati da "equazioni" ancora più complicate di quella ipotizzata sopra, si potrebbe (in linea puramente di principio) "scrivere" equazioni con opportuni termini di sollecitazione e di interazione simili a quelli per i granelli e domandarsi verso quale "struttura" sociale e interpersonale vada evolvendo l'insieme dell'umanità? È "naturale", anche per noi come per i granelli di sabbia, l'evoluzione verso l'ordine e l'armonia?

È molto difficile dare risposte conclusive. Qualche osservazione appare possibile e forse conduce anche a qualche elemento di speranza. In fondo, in

natura operano diversi meccanismi neuro-biologici che sembrano lavorare verso una forma di selezione del bello. Per esempio, presumibilmente un meccanismo di selezione sessuale ha sfavorito i Neandertal, in ragione del non gradevole aspetto, a favore invece dell' Homo sapiens, simultaneamente presenti entrambe le specie per un tempo piuttosto lungo.

Anche oggi, in un quadro di mutata evoluzione culturale e sociale, si potrebbe ipotizzare che gli individui obbediscano a complesse equazioni (che descrivono il movimento degli elettroni e stabiliscono connessioni neurali nel cervello) che potrebbero non essere molto diverse, nei loro termini, da quella per i granelli di sabbia. Potrebbero tali equazioni favorire l'evoluzione del mondo verso ordine e armonia universale ?

Ecco il "granello di speranza", il risultato di un "pattern" mentale al quale si potrebbe pervenire, in modo analogo ai granellini di sabbia sulla lastra vibrante. Che infine possa essere questa tendenza una forma di riscatto dai limiti e dalle sofferenze della condizione umana, l'affrancamento dell'uomo dall' "Inferno" e la sua promozione verso il "Paradiso". Molti hanno già osservato come su questa navicella terrestre della quale nella prima parte del libro abbiamo descritto alcuni fenomeni ed eventi, l'umanità si trovi in un "Inferno" quando si abbandonano regole di ordine e di armonia.

Ecco alcuni versi di **Ungaretti**, nell'ode "Mio fiume anche tu, Tevere fatale".

"Vedo ora nella notte triste, imparo,
So che l'inferno s'apre sulla Terra
Su misura di quanto
L'uomo si sottrae, folle,
alla purezza della Tua passione".

22.2 Fullerene

Sino a tempi relativamente recenti (per intendersi, attorno ai primi anni 1990) si conoscevano due forme allotropiche del carbonio allo stato solido: il diamante e la grafite.

Un cristallo di diamante si può ritenere costituito da una iper-molecola di atomi di carbonio, ciascuno dei quali si lega attraverso quattro legami con i quattro vicini prossimi, in una struttura locale a tetraedro. Queste direzioni di legame tetraedrico risultano da una particolare disposizione dei

quattro elettroni più esterni dell'atomo di carbonio (disposizione detta sp^3 , o anche orbitale atomico ibrido; questi orbitali sono il risultato di proprietà quantistiche che conseguono dalle leggi fisiche alle quali abbiamo già fatto cenno a proposito degli elettroni negli atomi). Questa configurazione elettronica è presente, per esempio, nella molecola del metano CH_4 , ove i legami del carbonio sono, in questo caso, con atomi di idrogeno.

Ritornando al diamante, quella struttura spaziale di iper-molecola rigida e molto coesa conferisce al cristallo le proprietà che lo hanno reso ben noto, soprattutto nei riguardi degli effetti sulla riflessione, rifrazione e insieme diffusione della luce, ma anche per la sua durezza e resistenza.

L'atomo di carbonio, tuttavia, può anche disporre i propri elettroni in un orbitale ibrido (detto sp^2 o trigonale) che comporta massimi di probabilità di presenza degli elettroni lungo direzioni che giacciono su un piano, a 120 gradi l'uno dall'altro. In questo caso si possono produrre tre legami, per esempio come avviene nella molecola $H_2C - CH_2$, nota come etilene. Il quarto elettrone del carbonio nel caso della ibridizzazione trigonale può formare un legame, molto più debole, con atomi sovrastanti o sottostanti al piano dei legami forti, oppure rafforzare debolmente il legame $C - C$.

Nello stato solido, se gli atomi di carbonio si dispongono nella configurazione ibrida sp^2 si produrrà una struttura planare con atomi C ai vertici di un reticolo cristallino esagonale e con deboli legami tra piani adiacenti. È questa la struttura della grafite, con proprietà elettriche e ottiche diversissime dal diamante (quest'ultimo isolante e trasparentissimo, la grafite opaca al visibile e conduttrice, anche se debolmente, proprio in virtù del quarto elettrone rimasto indenne dalla ibridizzazione planare).

Come si è detto le due forme del carbonio allo stato solido erano le uniche note sino agli inizi degli anni 1990. A seguito di una serie di osservazioni e analisi speculative, con un certo contributo di eventi occasionali e di previsioni teoriche, da alcuni gruppi di ricercatori chimici e fisici in diverse parti del mondo, divenne manifesto che nelle polveri di nerofumo, nella fiamma di una candela (ogni qualvolta si ha condensazione di vapori di carbonio e quindi, con un po' di forzatura, anche sotto la pentola in cui si preparano gli spaghetti) si formavano piccole quantità di un nuovo materiale, costituito da un insieme di molecole di carbonio a struttura di un pallone da calcio. Queste molecole (C_{60}) risultano dall'assemblamento di 60 atomi di carbonio, a formare 20 esagoni e 12 pentagoni (vedi Fig. 22.2), rispettando le regole del matematico svizzero **Eulero** che nel Sette-

cento aveva calcolato essere necessari esattamente 12 pentagoni per formare uno sferoide chiuso, mentre il numero di esagoni poteva variare (infatti fu successivamente sintetizzata anche la molecola C_{70} nella quale gli esagoni sono 25, cioè una struttura più simile al pallone da rugby (vedi Fig. 22.3), oppure la molecola “gigante” C_{540} (vedi Fig. 22.4)). Oggi questo nuovo materiale si produce con relativa facilità mediante la vaporizzazione della grafite sotto irraggiamento di impulsi laser o con opportuni archi elettrici ed è stato scoperto anche nella polvere interstellare.

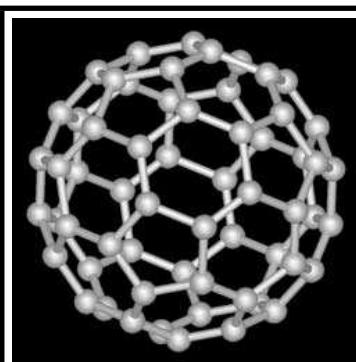


Fig. 22.2:
Architettura naturale
della molecola di
fullerene C_{60} .

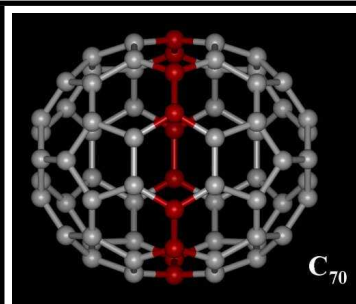


Fig. 22.3: Struttura a
pallone da rugby della
molecola C_{70} .

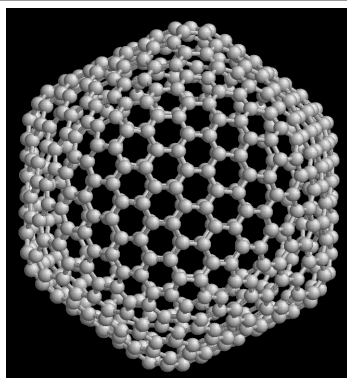
La molecola C_{60} è perfettamente sferica ed è la molecola più simmetrica che possa esistere nello spazio tridimensionale euclideo. Inoltre è caratterizzata da singolari proprietà chimiche (in particolare verso i radicali liberi) nonché da proprietà meccaniche e di stabilità (la facilità con cui può ruotare ne consente, per esempio, l'uso come elemento lubrificante per la riduzione

degli attriti). Infine, quando le molecole si raggruppano nello stato solido a formare un reticolo di tipo cubico, possono anche diventare sede della superconduzione elettrica (mediante opportune iniezioni di cariche itineranti, per quali rinviamo alla descrizione al capitolo 24 sui superconduttori).

La terza forma allotropica del carbonio, basata sulla molecola di C_{60} , destò subito enorme interesse nella chimica e nella fisica dello stato solido, in ragione dei nuovi aspetti di fondamento e delle molteplici prospettive applicative. Innumerevoli ricerche si svilupparono e diversi gruppi hanno studiato le sue proprietà e discusso teoricamente tutti gli aspetti (qualcuno la ha definita la molecola più discussa dell'universo).

Qui ci interessa rilevare l'aspetto di armoniosa architettura naturale della molecola di C_{60} (si veda la figura 22.2) e raccontarvi di un possibile rammarico per noi italiani. La molecola costituita dall'insieme di esagoni e di pentagoni è caratterizzata da un principio costruttivo strettamente simile a quello utilizzato da un architetto e inventore americano, **Buckminster Fuller**, per realizzare delle cupole geodetiche. Cosicché la molecola, e per estensione anche il solido, fu battezzato dai ricercatori americani *Fullerene*, attribuendo implicitamente a quell'architetto l'idea originale di generare queste bellissime e strutturalmente armoniche figure geometriche.

Fig. 22.4: Gigante
quasi-icosaedrale
 C_{540} .



In realtà secoli prima, un grande architetto e geometra italiano, **Piero della Francesca** aveva già individuato nell'icosaedro tronco una configurazione strutturale di particolare armonia e bellezza e ipotizzato la costruzione di cupole e strutture che ne sfruttassero le proprietà.

La priorità di Piero della Francesca nello sviluppo di problemi di costruzione e di calcolo geometrico relativi ai poliedri, mai prima disegnati in forma stereometrica (con l'eccezione, forse, di **Leonardo da Vinci**, secondo testimonianze non-dirette), certificata da documenti inoppugnabili conservati nei Musei Vaticani e raccolti nel "Libellus de Quinque Corporibus Regularibus" (parti riprodotte nella "Divina Proporzione" di **Fra Luca Pacioli**) avrebbe forse più opportunamente suggerito di attribuire al C_{60} il termine *Franceschene*. Suggestimenti avanzati in tal senso da ricercatori italiani (in particolare da **Taliani** e **Rigamonti**, a un convegno a Kirchberg, nel 1991) non ebbero successo, a seguito della ormai consolidata terminologia di origine statunitense.

Rimane la gratificante osservazione su come la potenza del genio di un italiano abbia potuto, secoli prima, ipotizzare delle geometrie di straordinaria armonia e bellezza architettonica, strutture che la natura gelosamente ha nascosto nei tempi e che solo molto più tardi sono state portate alla luce dagli studiosi.

Capitolo 23

Palle di neve e bollicine nell'elio liquido.

L'elio si trova quasi all'inizio della Tavola di **Mendeleev**, eppure sin dal momento della sua scoperta ha creato molti problemi ai fisici per la singolarità delle sue proprietà. Per la verità, questi problemi sono compensati dalla bellezza e unicità dei fenomeni che si verificano nell'elio (e anche dalle possibilità che esso offre ai ricercatori e agli ingegneri di ottenere temperature molto basse). Tra le meraviglie di questo liquido quantistico, oltre alla superfluidità (alla quale abbiamo fatto cenno al capitolo 21) vi è anche un meccanismo particolare, diverso da tutti gli altri: quello della conduzione elettrica.

Quando alla fine degli anni '50, i fisici si accingevano allo studio di questo problema, sembrò che i candidati più idonei per costituire i portatori di carica fossero gli elettroni e gli ioni positivi della ionizzazione dell'elio. Si supposeva che la carica positiva fosse trasportata non direttamente dagli ioni dell'elio (essi sono relativamente pesanti e il campo elettrico li accelera con difficoltà) ma dalle lacune.

Per capire che cosa sia una lacuna, immaginate che l'elettrone - che si trova nell'atomo di elio - ad un certo istante salti dall'atomo di appartenenza all'atomo di elio che si trova nelle vicinanze. Al posto lasciato libero da questo elettrone salta un elettrone dall'atomo attiguo, al posto di questo un terzo dall'atomo successivo e così via. Quest'altalena degli elettroni può essere raffigurata come il movimento della carica positiva nel senso opposto.

23.1 Palle di neve

In realtà questa carica positiva non esiste, si verifica semplicemente l'assenza dell'elettrone di turno dal suo "posto di lavoro": questa situazione si descrive attraverso una pseudo-particella, *la lacuna*. Tale meccanismo di trasporto di carica, di solito, opera nei semiconduttori e sembrava del tutto improbabile che potesse applicarsi anche all'elio liquido.

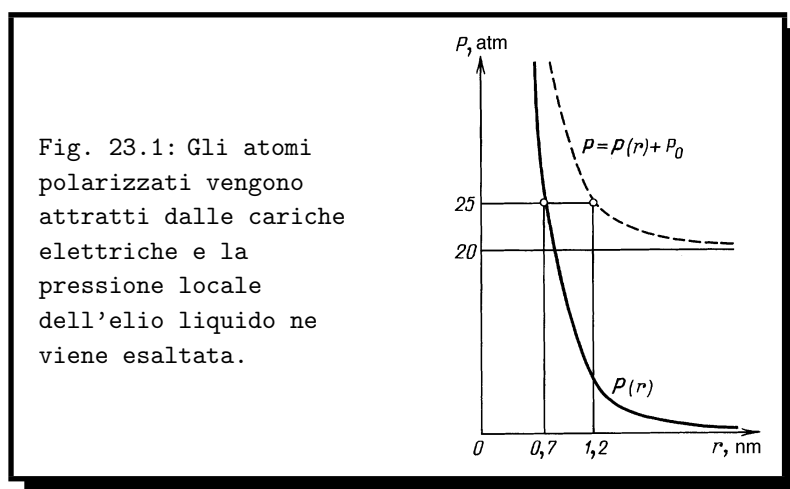
Per misurare le masse dei portatori di carica positiva e negativa nell'elio liquido sono state studiate le loro traiettorie in presenza di un campo magnetico omogeneo. Come è noto, in caso di ingresso nel campo magnetico di una particella carica, con una certa velocità iniziale, la sua traiettoria si avvita in forma di spirale. Conoscendo la velocità iniziale, l'entità del campo magnetico e avendo misurato il raggio del moto circolare della particella, si può facilmente determinare la sua massa. I risultati di questi esperimenti sono stati particolarmente inaspettati: le masse dei trasportatori di cariche positiva e negativa superavano di decine di migliaia di volte la massa dell'elettrone libero!

Naturalmente, il moto degli elettroni e delle lacune nell'elio liquido avviene attraverso gli atomi con cui essi interagiscono; per questo, le masse dei portatori possono differire dalla massa dell'elettrone libero. Tuttavia il differire di cinque ordini di grandezza è veramente singolare! Questa sorprendente differenza tra le aspettative teoriche e i dati sperimentali non era facilmente spiegabile. Per questo è stato necessario ipotizzare qualcosa di nuovo, un meccanismo sino ad allora sconosciuto.

La corretta spiegazione della struttura dei portatori di carica positiva nell'elio liquido è stata proposta dal fisico americano **Kennet Robert Atkins**. È noto che per trasformare una sostanza dallo stato liquido a quello solido non è necessario abbassare la temperatura, si può aumentare la pressione a temperatura costante e forzare in tal modo la sostanza a solidificare. Il valore della pressione alla quale la sostanza solidifica si chiama *pressione di solidificazione* (P_{sol}). È naturale che la grandezza P_{sol} dipenda dalla temperatura: quanto più la temperatura è alta, tanto più difficile è obbligare il liquido a passare alla fase solida attraverso l'aumento della pressione. All'aumentare della temperatura P_{sol} pertanto aumenta. Si è compreso che la scaltrezza del portatore di carica positiva nell'elio liquido si basa proprio sulla pressione di solidificazione relativamente piccola: a temperature abbastanza basse si ha $P_{\text{sol}} = 25$ atm., come abbiamo già avuto modo di commentare nella sezione sulle oscillazioni di punto zero

(capitolo 21).

Si è già detto come nell'elio liquido possano esistere ioni positivi He^+ . L'interazione di un tale ione con l'atomo neutro He determina un'attrazione dei suoi elettroni da parte dello ione positivo, mentre il nucleo dell'atomo, carico positivamente è respinto. Di conseguenza i centri delle cariche positiva e negativa nell'atomo neutro non sono più coincidenti. In altre parole, la presenza di ioni positivi nell'elio liquido porta alla *polarizzazione* dei suoi atomi. Questi atomi polarizzati sono attratti dallo ione positivo e ciò determina un aumento della loro concentrazione (ossia la comparsa di una densità in eccesso). Di conseguenza, si verifica anche un aumento della pressione man mano che ci si avvicina allo ione. La dipendenza della pressione in eccesso dalla distanza r dal centro dello ione positivo è mostrata graficamente nella figura 23.1.



Alla pressione di 25 atmosfere (e a temperature basse) l'elio solidifica. Pertanto, non appena la pressione vicino allo ione positivo raggiunge questo valore, un certo volume di elio intorno ad esso solidifica. Dalla figura 23.1 si può osservare come la solidificazione avvenga ad una distanza $r_0 = 0,7 \text{ nm}$ dallo ione.

Lo ione positivo risulta pertanto come inglobato nel ghiaccio della “palla di neve”. Se si applica un campo elettrico esterno, questa palla di neve inizierà a muoversi. Tuttavia, poiché essa costituisce il centro della regione della densità in eccesso, la palla non si muoverà da sola nel campo elettrico,

ma sarà accompagnata da una sorta di scorta d'onore: dietro di lei si muove tutta la coda della densità in eccesso. La massa totale del portatore della carica positiva è così determinata dalla somma di tre fattori.

Il primo è la massa della stessa "palla di neve", uguale al prodotto della densità dell'elio solido per il volume della palla. Con la pressione esterna di una atmosfera tale massa vale $32m_0$ ($m_0 = 6.7 \cdot 10^{-27} \text{ kg}$ è la massa dell'atomo di elio). Di poco meno massivo è anche il seguito che accompagna la "palla". La massa della coda di densità in eccesso che lo ione trascina con se è di $28m_0$.

Oltre a questi, nel movimento di un corpo in un liquido si ha uno scivolamento di uno strato su un altro. Questo processo dissipa energia per cui, per comunicare al corpo durante il suo movimento nel liquido una certa accelerazione è necessario applicare una forza maggiore di quella che sarebbe necessaria nel vuoto. Questa "massa", legata all'attrito tra strati, è chiamata massa aggiunta. Per la "palla" che si muove nell'elio liquido, la massa aggiunta è uguale a $15m_0$. Quindi la massa effettiva del portatore di carica positiva che si muove nell'elio liquido risulta $75m_0$.

Per illustrare il meccanismo di trasporto della carica positiva abbiamo potuto ricorrere alla fisica classica. Diversamente vanno le cose per quanto riguarda il portatore della carica negativa. Da una parte si hanno alcuni ioni carichi negativamente (possono formarsi degli ioni molecolari negativi He_2^- , tuttavia di questi ioni se ne formano pochi e nel trasferimento della carica non giocano il ruolo principale). In linea di principio il portatore di carica negativa rimane l'elettrone, la cui massa è enormemente più piccola di quanto appare dai dati sperimentali.

Bisogna considerare i prodigi che si verificano nel mondo dei quanti. L'esperimento mostra che l'elettrone, al quale abbiamo assegnato il ruolo di portatore di carica negativa nell'elio liquido... non può in realtà neppure penetrare in esso. Per capirci qualcosa, dovremo fare una piccola digressione e parlare della struttura degli atomi a molti elettroni. Il *principio di Pauli* illustra l'avversione degli atomi alla cattura di ulteriori elettroni, il che determina appunto le difficoltà che l'elettrone deve superare per penetrare nell'elio liquido.

L'energia dell'elettrone nell'atomo, come si è detto al capitolo 21, può assumere solo determinati valori discreti. A un dato valore dell'energia possono corrispondere solo alcuni limitati stati dell'elettrone, che differiscono tra loro per il carattere del "moto" (come abbiamo visto, per la forma dell'orbita, cioè, nel linguaggio quantistico, per la struttura della nube di

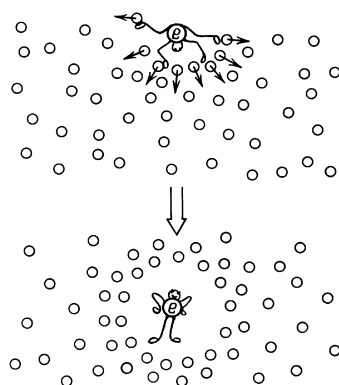
probabilità che definisce la distribuzione dell'elettrone nello spazio). Secondo il principio di **Pauli**, all'aumentare del numero di elettroni nell'atomo (del numero d'ordine dell'elemento) essi non si possono raggruppare in stati equivalenti ma devono gradatamente riempire i livelli successivi.

Inizialmente si riempie il primo stato, che corrisponde all'energia minore possibile. In tale stato, vicino al nucleo, si collocano soltanto due elettroni. Per questo nell'elio, che nella Tavola di Mendeleev ha un numero d'ordine 2, il primo stato risulta occupato dal numero massimo di elettroni accomodabili in esso. Per il terzo elettrone non c'è più posto ed esso può sistemarsi soltanto ad una certa "distanza" dal nucleo. Cosicché nell'avvicinamento dell'ulteriore elettrone ad una distanza dell'ordine della dimensione dell'atomo compaiono forze di repulsione, che ostacolano un ulteriore avvicinamento. Quindi per spingere l'elettrone all'interno dell'elio, è necessario compiere un certo lavoro di penetrazione.

23.2 Bollicine e...noccioline

I fisici italiani **Careri**, **Fasoli** e **Gaeta** hanno ipotizzato che l'elettrone scuota gli atomi ai quali non ha "diritto" di avvicinarsi troppo formando intorno a se una cavità sfericamente simmetrica, una particolare "bollicina" (Fig. 23.2). Proprio questa bollicina, insieme all'elettrone che fluttua qua e là al suo interno, è il portatore della carica negativa nell'elio liquido.

Fig. 23.2: A causa degli effetti quantistici l'elettrone non può avvicinarsi troppo agli atomi di elio e li spinge via.



Possiamo valutare le dimensioni di questa bollicina. All'aumentare della distanza, le forze di repulsione decrescono rapidamente. Allo stesso tempo sugli atomi relativamente lontani l'elettrone agisce allo stesso modo dello ione positivo, cioè li polarizza. Pertanto a grandi distanze l'interazione dell'elettrone con gli atomi di elio sarà la stessa che nel caso della "palla di neve". Di conseguenza, all'interno della "bollicina", in cui è insediato l'elettrone, la pressione in eccesso cresce con la stessa legge (si veda la figura 23.1).

Tuttavia, nel caso di una piccola pressione esterna sul limite della "bollicina" essa rimane comunque inferiore a 25 atm a causa della sua dimensione relativamente grande. Oltre a questa pressione, relativa alla densità in eccesso dell'elio polarizzato, sul contorno della "bollicina" agiscono forze di tensione superficiale dirette, come le forze di sovrappressione, verso il centro. Cosa equilibra queste forze dall'interno? Lo stesso elettrone, che non potendo penetrare negli atomi in virtù dei principi quantistici provvede al riequilibrio.

In accordo al principio di indeterminazione, del quale abbiamo parlato dettagliatamente nel capitolo 21, l'indeterminazione nella misura dell'impulso dell'elettrone è legata alla indeterminazione nella sua posizione dalla relazione $\Delta p \sim \hbar / \Delta x$. Nel caso in esame, la posizione dell'elettrone può essere definita con la precisione delle dimensioni della stessa "bollicina", ossia $\Delta x \sim R$. Di conseguenza, trovandosi all'interno della "bollicina", l'elettrone non può restare a riposo, ma deve, per così dire, dimenarsi in esso per conseguire un impulso dell'ordine di $\hbar / R$ e di conseguenza un'energia cinetica $E_c = p^2 / 2m_e \sim \hbar^2 / 2m_e R^2$.

A causa della collisione con le pareti si crea una certa pressione (ricordate l'equazione della teoria cinetico-molecolare che collega la pressione del gas con l'energia cinetica media del moto caotico delle sue particelle e la loro concentrazione: $P = \frac{2}{3}nE_c$). Tale pressione compensa la pressione totale sulla "bollicina". In altre parole, l'elettrone all'interno della "bollicina" si comporta come un gas in un recipiente chiuso, un gas composto da una sola particella!

La concentrazione di questo gas, evidentemente, è $n = 1 / V = 3 / 4\pi R^3$. Ponendo questa grandezza e l'energia cinetica $E_c \approx \hbar^2 / 2m_e R^2$ nell'espressione per la pressione, ricaviamo $P_e^{(v)} \approx \hbar^2 / 4\pi m_e R^5$.

Il calcolo quantistico accurato fornisce un valore di circa un ordine di grandezza maggiore, precisamente $P_e^{(q)} = \pi^2 \hbar^2 / 4m_e R^5$. Se la pressione esterna è bassa, la pressione sulla bollicina è determinata dalle forze di

tensione superficiale, $P_L = 2\sigma / R$. Scrivendo $P_e^{(q)} = P_L$, ricaviamo che il raggio della bollicina all'equilibrio vale

$$R_0 = \left(\frac{\pi^2 \hbar^2}{8 m_e \sigma} \right)^{\frac{1}{4}} \approx 2 \text{ nm} = 2 \cdot 10^{-9} \text{ m}.$$

Così abbiamo chiarito che i portatori di carica negativa nell'elio liquido sono le bollicine con all'interno un elettrone. La massa da attribuire a questi portatori è la massa della "palla". Tuttavia, in pratica non si ha una reale massa propria della "bollicina", essendo essa uguale alla massa dell'elettrone e quindi trascurabilmente piccola rispetto alla massa del liquido attratto, "il seguito" e la massa aggiunta. Per questo la massa totale del portatore nell'elio liquido è determinata dalla massa aggiunta più la massa della "coda" della densità in eccesso che accompagna la "bollicina" nel suo moto. A causa della grande dimensione della "bollicina" essa risulta notevolmente maggiore della massa della "palla di neve" ed è pari a $245 m_0$.

Esaminiamo ora come influisce sulle proprietà dei portatori di carica l'aumento della pressione esterna P_0 . Nella Fig. 23.1 è riportata la dipendenza dalla distanza della pressione totale vicino allo ione nell'elio liquido. Questa dipendenza per pressione arbitraria esterna $P_0^* < 20 \text{ atm}$ si ottiene traslando la dipendenza per $P_0^* = 0$ lungo l'asse delle ordinate. Come si vede dalla figura 23.1, quanto maggiore è la pressione esterna, a più grande distanza dallo ione la pressione totale diviene pari a 25 atm . Per questo, con l'aumento della pressione esterna, la "palla di neve" provoca un fenomeno di valanga, simile a quello reale: essa si copre impetuosamente di altra "neve", ossia elio solido, diventando sempre più grande. La dipendenza della dimensione della palla $r(P_0)$ dalla pressione esterna è mostrata nella figura 23.3.

Come si comporta la "bollicina" quando cresce la pressione? Fino a un certo punto, analogamente a qualsiasi bollicina in un liquido, crescendo la pressione esterna essa si comprime. Il suo raggio $R(P_0)$ diminuisce, come è mostrato nella Fig. 23.3. Tuttavia, a una pressione $P_0^* = 20 \text{ atm}$ i grafici delle dipendenze $r(P_0)$ e $R(P_0)$ intersecano; con questa pressione le dimensioni della "bollicina" e della "palla di neve" diventano entrambe uguali a 1.2 nm . Il destino della "palla di neve" con un ulteriore aumento della pressione, ci è noto: essa aumenterà marcatamente le proprie dimensioni a spese dell'elio che solidifica sulla sua superficie.

Come deve comportarsi la "bollicina", deve ancora comprimersi (linea

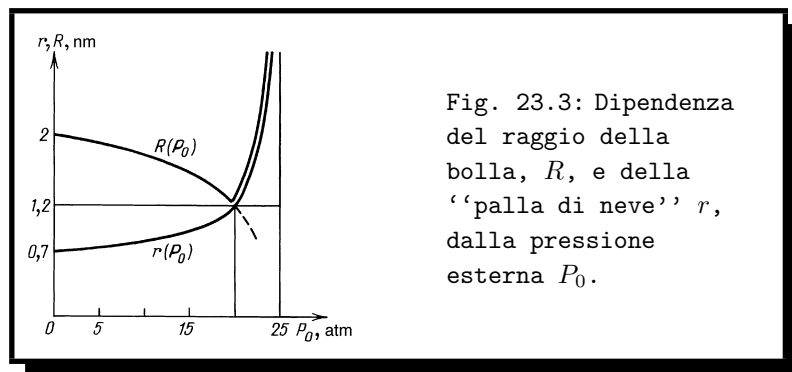


Fig. 23.3: Dipendenza del raggio della bolla, R , e della ‘palla di neve’ r , dalla pressione esterna P_0 .

tratteggiata nella figura 23.3)? Qui la “bollicina” mostra la sua “cocciutagine”: aumentando ulteriormente la pressione essa, come la “palla di neve”, inizia a coprirsi di una scorza di “ghiaccio” di elio liquido. Effettivamente con una pressione esterna $P_0^* = 20 \text{ atm}$, la pressione totale sul contorno della bollicina (ad una distanza di 1.2 nm dal suo centro) diventa uguale a 25 atm , ossia raggiunge la pressione di solidificazione dell’elio liquido. La “bollicina”, in questo caso, si “veste” di un guscio di ghiaccio, il cui raggio interno in caso di ulteriore aumento della pressione rimane all’incirca invariato, mentre quello esterno è precisamente uguale alla dimensione della palla, con la corrispondente pressione.

Così, con pressioni esterne superiori a tale valore le “bollicine” si coprono con un “guscio di ghiaccio” e si trasformano in questo modo in qualcosa di simile alle noccioline. Tuttavia il piccolo nucleo di queste “noccioline” è davvero particolare, si tratta di un elettrone, che si muove quantisticamente all’interno del guscio.

Rimane da aggiungere che con $P_0 \rightarrow 25 \text{ atm}$ il raggio esterno della “nocciolina”, così come della “palla di neve”, tende all’infinito, il che corrisponde alla solidificazione totale dell’elio liquido in tutto il volume del recipiente. Nell’elio solido, portatori di carica negativa sono le “bollicine” ghiacciate, con dimensioni uguali al raggio interno della ex “nocciolina”, all’incirca 1.2 nm . La carica positiva è trasportata dagli ioni, i residui delle ex “palle di neve”. Tuttavia in queste condizioni il trasporto delle cariche non è così semplice: la mobilità rimane di molti ordini inferiore a quelle delle “palle di neve” e delle “bollicine” nell’elio liquido.

Avreste immaginato che in un “bicchiere” di elio liquido accadessero fenomeni così peculiari? Ecco la ragione di considerare l’elio, come ha

Bollicine e...noccioline

245

detto Landau, una finestra sul mondo quantistico!

Capitolo 24

L'affascinante mondo dei superconduttori.

Il fenomeno della superconduttività consiste nel fatto che a una certa temperatura, in genere piuttosto bassa, metalli e leghe transiscono a un particolare “stato” che comporta la totale scomparsa della resistenza elettrica, si ch  la corrente circola senza perdite. Prima del 1987 le temperature alle quali si osservava questo fenomeno non erano lontane dallo zero assoluto e per ottenere la transizione allo stato superconduttivo era necessario immergere il composto in elio liquido.

Nei laboratori di tutto il mondo in realt  molti fisici e chimici erano al lavoro per scoprire composti superconduttori con temperature critiche pi  elevate, in modo che in luogo del refrigerante elio liquido si potesse, per esempio, utilizzare l’azoto liquido, che come abbiamo visto bolle alla temperatura di 77 K ,   poco costoso e facilmente reperibile.

Un certo progresso era stato ottenuto: attorno agli anni 1950 era largamente usata per i cavi superconduttori una lega di Niobio e Germanio, con temperatura critica di circa 23 K .

Nel 1986 il mondo scientifico fu messo a subbuglio dalla notizia che era stato trovata una classe di materiali, genericamente definibili come ossidi di rame, con temperature critiche attorno ai 40 K . Da questa scoperta prese l’avvio una intensa attivit  scientifica e sono stati successivamente scoperti e studiati molti altri materiali della stessa famiglia o di famiglie analoghe, con temperature critiche via via pi  alte. Si tratta dei composti oggi noti come *superconduttori ad alta temperatura critica*, assai diversi dai metalli e dalle leghe tradizionali.

Cominceremo a raccontarvi la storia di questa scoperta, per poi passare alla descrizione di affascinanti fenomeni ed effetti che si realizzano nello

stato superconduttivo in genere.

24.1 Superconduzione e la “corsa” verso alte temperature

L'articolo di **J.Georg Bednorz** e **Karl Alexander Müller** nel 1986 aveva un titolo abbastanza prudente: “About the possibility of the superconductivity at high temperature in the La-Ba-Cu-O system”. Nonostante la prudenza degli autori, quest'articolo non venne accettato dalla principale rivista americana di fisica (la “*Physical Review Letters*”) senza che alcuni dei loro Referee potessero esaminarlo criticamente. Da parte loro gli autori comprensibilmente non desideravano che tale esame critico venisse condotto per timore di “sottrazione di priorità”. D'altra parte, durante i venti anni precedenti la comunità scientifica aveva registrato numerose volte clamorosi annunci di pseudo-scoperte di un qualche mitico superconduttore ad alta temperatura critica. Probabilmente per questi precedenti l'articolo di Bednorz e Müller non fu acriticamente accolto dalla rivista. Gli autori si rivolsero allora alla rivista tedesca “*Zeitschrift fur Physik*” che lo pubblicò^a.

Alla pubblicazione del lavoro la comunità scientifica rimase piuttosto fredda. Soltanto in un laboratorio giapponese alcuni ricercatori riprodussero il fenomeno, confermando la scoperta. Dopo questa riconferma, all'inizio del 1987, il mondo scientifico fu invaso da una sorta di febbre per la ricerca di nuovi materiali superconduttori ad alta temperatura e per lo studio approfondito di quelli già individuati.

Sotto la spinta di una intensa attività di ricerca, si sintetizzarono rapidamente nuovi materiali con temperatura critica via via più alta: 45 *K* per il composto La – Sr – Cu – O, 52 *K* per il La – Ba – Cu – O sotto pressione. L'osservazione dell'effetto della pressione sulla temperatura di transizione fu determinante per gli orientamenti successivi. Lo statunitense **Paul Chu** propose infatti di utilizzare una specie di pressione interna, sostituendo gli atomi di La con atomi di Y, di dimensione più piccola e dello stesso gruppo nella tabella di Mendeleev, ritenendo che l'introduzione di atomi di Y avrebbe creato una “pressione chimica”.

In questo modo nacque il più famoso dei superconduttori ad alta tem-

^aIl premio Nobel per la Fisica del 1987 fu conferito a Bednorz e Müller. La importanza della loro scoperta è anche indicata dal fatto che tra la pubblicazione del loro articolo e il conferimento del Nobel trascorse solo un anno.

peratura critica, l' $\text{YBa}_2\text{Cu}_3\text{O}_7$, con temperatura critica $T_c = 93\text{ K}$, oggi noto con l'acronimo YBCO (Fig.24.1). Con questo composto, per la prima volta si è ottenuta la superconduzione a temperatura superiore a quella di ebollizione dell' azoto liquido. Questo refrigerante (che abbiamo già incontrato per la preparazione di un nuovo tipo di gelato, nel capitolo 20) è assai più comune, facile da maneggiare e meno costoso dell'elio liquido usato come elemento refrigerante per i superconduttori convenzionali di carattere metallico. Si aprivano così prospettive impensabili sino ad allora!

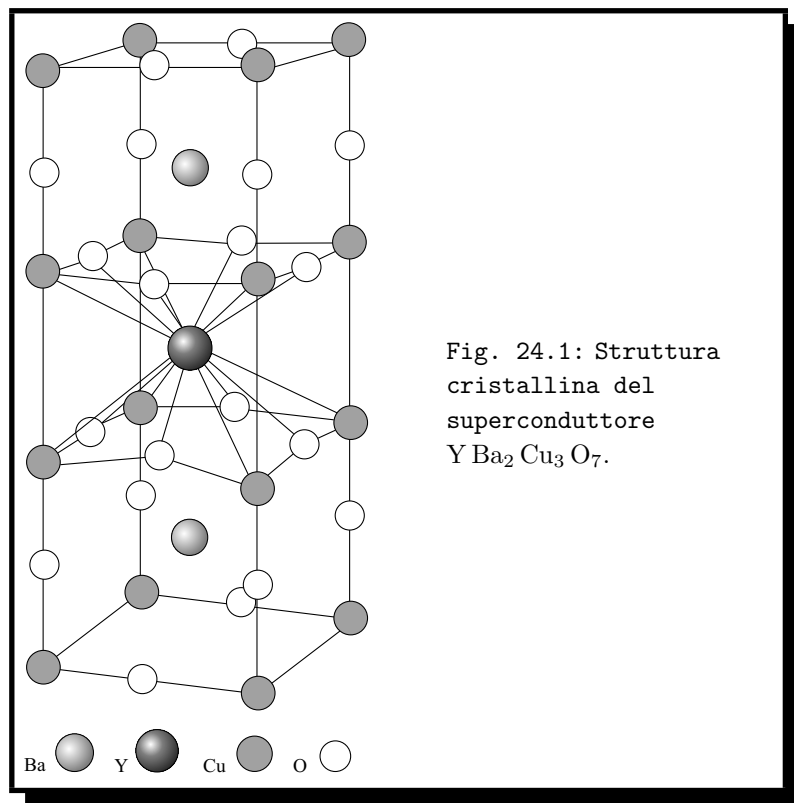
Nel 1988 fu sintetizzato un superconduttore composto da cinque elementi del tipo $\text{Ba} - \text{Ca} - \text{Sr} - \text{Cu} - \text{O}$, con temperatura critica pari a 110 K . Poco più tardi si riuscì a sintetizzare composti analoghi, basati su Hg e Tl in sostituzione di Sr, caratterizzati da temperatura critica attorno a 125 K . Il superconduttore Ba-Ca-Hg-Cu-O , sotto pressione di 300.000 atm , raggiunge la temperatura critica di 165 K , la più alta sinora osservata, corrispondente a $-108\text{ }^\circ\text{C}$ e non molto lontana dalle temperature “ambiente” minime raggiunte in Antartide!

La preparazione chimica di un composto della famiglia dei superconduttori ad alta temperatura è per la verità assai semplice. L'YBCO può essere ottenuto in poche ore in un buon laboratorio di chimica di una scuola media superiore. Questi composti superconduttori ad alta temperatura forse avrebbero potuto essere sintetizzati già nel Medio Evo!

Invece lo sviluppo della superconduttività, uno dei più interessanti e inusuali fenomeni della fisica dei solidi, ha preso altre strade.

Nel 1911, ad una riunione dell'Accademia Reale delle Scienze in Amsterdam, il fisico olandese **Kamerlingh Onnes** riferì di aver osservato la totale sparizione della resistenza elettrica nel mercurio raffreddato in elio liquido, a $4,15\text{ K}$ assoluti. In anni successivi tale fenomeno fu osservato in altri metalli, ancora fruendo delle tecniche di refrigerazione ideate da Onnes che avevano permesso di liquefare l'elio.

Le prospettive di applicazioni pratiche sembravano illimitate: trasmissione di energia elettrica senza dissipazione, magneti superpotenti, nuovi motori elettrici, trasformatori e altri dispositivi elettronici innovativi. Vi erano tuttavia due ostacoli. L'uno, le temperature critiche così basse richiedevano la liquefazione dell'elio (costoso da produrre e difficile da conservare, come si è fatto cenno) e quindi un complesso di apparecchiature complementari. Il secondo ostacolo era costituito dal fatto che lo stato di superconduzione veniva “distrutto” da un campo magnetico relativamente debole e quindi anche dalle correnti che venivano immesse in una bobina costruita



con un filo metallico superconduttore.

Altra proprietà di carattere fondamentale dello stato di superconduzione fu scoperta nel 1933: se si applica un campo magnetico, a temperatura ambiente, ad un blocco di metallo potenzialmente superconduttore e quindi lo si raffredda, quando la temperatura scende al di sotto di quella critica si osserva la completa espulsione del campo magnetico dal volume del superconduttore. È questo il cosiddetto effetto **Meißner-Ochsenfeld**.

Una seconda era nel campo della superconduttività si può ritenere quella che va da circa il 1950 sino alla scoperta di Bednorz e Muller. In quegli anni furono osservate ulteriori proprietà dello stato superconduttivo e insieme fu formulata una teoria di carattere microscopico che descriveva il fenomeno e rendeva ragione di tali proprietà. Si può sintetizzare l'insieme di queste particolarissime proprietà dicendo che lo stato di superconduzione è uno

stato quantistico su scala macroscopica. In altre parole, in un pezzo di materiale superconduttore avvengono e sono osservabili, fenomeni tipici del mondo dei quanti, che abbiamo descritto come si realizzano a livello microscopico, in atomi o molecole. Questi effetti quantici comportano, per esempio, la discretezza dei valori di un campo magnetico che attraversa la superficie circoscritta da un anello e la possibilità di costruire magnetometri che vedono entrare e uscire dal superconduttore i quanti elementari di flusso del campo magnetico, benché incommensurabilmente piccoli.

Lo sviluppo della descrizione teorica della superconduttività nei metalli, ha visto dapprima la teoria fenomenologica di **Vitaly Ginzburg** e **Lev Landau** (1950) seguita da una completa teoria di carattere microscopico ad opera di **John Bardeen**, **Leon Cooper** e **John Schrieffer** (1957). Questi scienziati mostrarono che la superconduttività è connessa alla comparsa di una particolare forza attrattiva che accoppia gli elettroni (che essendo cariche dello stesso segno tenderebbero a respingersi). Tale forza attrattiva, di carattere quantistico, si realizza in virtù della lieve deformazione che un elettrone induce nel reticolo delle cariche positive (gli ioni) del metallo. In questa deformazione “si infila” un secondo elettrone e si forma così la cosiddetta coppia di Cooper.

Le dimensioni spaziali di queste coppie in genere sono in realtà molto più elevate delle distanze tra atomi nel metallo. Seguendo la visualizzazione di Schrieffer, non si deve pensare che i due elettroni che costituiscono la coppia siano simili a una stella doppia. Piuttosto, essi devono essere immaginati come due danzatori in una grande sala, che danzano ritmicamente accoppiati ma mantenendosi separati da decine o centinaia di altri danzatori. L'energia di legame della coppia, al limite di temperatura vicina allo zero assoluto, determina il valore della temperatura critica T_c . In virtù delle leggi che regolano il mondo dei quanti, il comportamento delle coppie di Cooper differisce radicalmente da quello degli elettroni^b. In particolare,

^bLa coppia di Cooper è una “struttura” dinamica, cioè essa esiste solo in moto. Un modello illustrativo di questa attrazione tra elettroni può essere fatto considerando un elettrone, che in un metallo si muove alla velocità di Fermi $v_F \approx 10^8$ cm/sec, lungo un canale del reticolo cristallino. Ad un certo istante, quando l'elettrone si trova tra due ioni vicini, questi subiscono un breve impulso dovuto all'attrazione coulombiana con l'elettrone stesso. La durata τ di questo impulso è all'incirca il tempo di transito dell'elettrone tra due ioni vicini e quindi pari a $\tau \sim a/v_F \cong 10^{-16}$ s, essendo “a” la distanza interatomica pari a circa 10^{-8} cm. Durante questo intervallo di tempo gli ioni che hanno una frequenza di vibrazione propria tipicamente pari all'incirca $\omega \sim 10^{13}$ Hz, non cambiano praticamente posizione. Dopo mezzo periodo di vibrazione gli ioni si

esse possono occupare lo stato di più bassa energia senza alcuna limitazione nel loro numero, in apparente violazione del principio di Pauli che limita il numero di elettroni negli stati discreti di energia, come abbiamo visto per gli atomi. Pertanto, al di sotto di T_c , si hanno nel superconduttore due tipi di portatori di corrente: gli elettroni non-accoppiati e le coppie di Cooper. I primi subiscono ancora tutti i fenomeni di urto contro impurità o contro gli ioni positivi che determinano l'usuale resistenza elettrica dei metalli ordinari. Le coppie, invece, rimangono nello stato fondamentale e possono trasportare la carica $2e$ senza interagire con le impurità o con il reticolo di ioni positivi. Il trasporto di cariche da parte delle coppie non provoca pertanto dissipazione di energia e sviluppo di calore. Tutte le coppie di Cooper si muovono così in un singolo stato quantico, che si estende su tutto il materiale.

Lo sviluppo della teoria della superconduttività ha prodotto come effetto un forte impulso alla ricerca applicata. A questo proposito grande importanza ha avuto la scoperta, dovuta allo scienziato russo **Alexei A. Abrikosov**, di una classe di superconduttori (detti del secondo tipo) che si differenziano marcatamente da quelli precedentemente noti negli anni 1950 per il comportamento in campo magnetico. Prima del lavoro di Abrikosov (1957) si riteneva che il campo magnetico non potesse penetrare nella fase superconduttrice senza distruggerla e ciò è effettivamente vero in pratica per quasi tutti i metalli puri. Abrikosov ha mostrato teoricamente che esiste anche un'altra possibilità, che non comporta la completa distruzione della superconducibilità. In questo secondo caso, ad un certo valore del campo magnetico (campo critico H_{c1}) nel superconduttore si creano vortici di corrente nei quali la parte centrale è normale mentre la parte perifer-

portano ad una distanza minima e in questa configurazione si crea nel reticolo cristallino un aumento locale di densità di carica positiva. A questo istante l'elettrone che ha creato tale aumento di carica, si trova già lontano, avendo percorso una distanza pari a $v_F/\omega \cong 10^{-4} \div 10^{-5} cm$ molto più grande del parametro reticolare " a ". La nuvola di carica positiva in eccesso segue quindi come una scia l'elettrone, che passa tra gli ioni con velocità v_F . Questa "nuvola" positiva può attrarre un secondo elettrone, che peraltro si trova più vicino alla nuvola stessa che non al primo elettrone.

Questa situazione dinamica, per un intervallo di tempo abbastanza breve, è in definitiva equivalente all'attrazione reciproca di due elettroni. Da questo semplice modello di coppia di Cooper si può anche dedurre che, se il reticolo è "morbido", l'ampiezza dello spostamento degli ioni al passaggio dell'elettrone diviene elevata e quindi aumenta anche la carica in eccesso che collega e tiene legata la coppia dei due elettroni. D'altra parte, quanto più piccola è la massa M dell'atomo del reticolo, tanto più grande è la frequenza caratteristica ω e tanto maggiore è l'energia di legame della coppia.

ica è superconduttiva! Descriveremo più oltre, in un qualche dettaglio, i vortici di Abrikosov. Per ora diremo solo che un superconduttore del primo tipo può essere trasformato in un superconduttore del secondo tipo tramite l'introduzione di impurezze o di difetti. Il fatto importante ai fini applicativi è che i superconduttori di secondo tipo possono ospitare nel loro volume supercorrenti di elevata densità e sopportare quindi anche intensi campi magnetici, come descriveremo nel seguito.

Come abbiamo detto, i superconduttori ad alta temperatura che hanno fatto seguito alla scoperta di Bednorz e Muller, caratterizzano quella che si può ritenere la terza era della superconduttività. Per essi, così radicalmente diversi dai metalli convenzionali, non si dispone ancora di una teoria completa di carattere microscopico analoga a quella di Bardeen, Cooper e Schrieffer. In particolare, non è ancora stato chiarito il meccanismo di accoppiamento tra elettroni, se cioè ancora operi una forma di deformazione del reticolo cristallino oppure se intervengono meccanismi di attrazione completamente nuovi, per esempio di carattere magnetico. Pur tuttavia, si sa per certo che anche nei superconduttori ceramici ad alta T_c a base rame (tutti i superconduttori del nuovo tipo sinora scoperti hanno come elemento caratterizzante gli atomi di rame) si formano coppie di Cooper e che in essi hanno luogo tutti quei fenomeni quantistici tipici dello stato superconduttivo, che descriveremo in maggiore dettaglio nelle successive sezioni di questo capitolo.

Sulla superconduttività si basa il funzionamento di voltmetri sensibilissimi, di rivelatori e di acceleratori di particelle. Si sta ipotizzando una nuova generazione di calcolatori elettronici e prende piede una nuova branca dell'elettronica: la crio-elettronica, basata in larga misura sulla superconduttività. Con i nuovi superconduttori ad alta temperatura critica si aprono in linea di principio altre fantastiche prospettive. Il loro uso come cavi in linee elettriche farebbe risparmiare dal 20 al 30 per cento dell'energia elettrica oggi prodotta. I progetti di ottenere energia dalla fusione termonucleare, con gli enormi magneti superconduttori per il contenimento del plasma, subiranno cambiamenti significativi con l'abbandono dell'elio liquido. Gigantesche bobine superconduttrici verranno costruite per l'accumulo di energia elettrica. Apparecchiature super-sensibili per magneto-cardiogrammi, magneto-encefalogrammi, per lo studio della magneto-tellurica e altro diventeranno di uso comune. Anche se rimangono da risolvere alcuni seri problemi di carattere tecnologico, è indubbio che oggi noi sappiamo che il quasi impossibile sta divenendo accessibile.

Dopo questa introduzione di carattere generale, ora passiamo a descrivervi alcune delle manifestazioni quantistiche della superconduttività, che sono alla base di effetti veramente peculiari, così come di molteplici dispositivi.

24.2 Quantizzazione del flusso magnetico

Nel micromondo, il mondo delle molecole, degli atomi e delle particelle elementari, molte grandezze fisiche possono assumere soltanto determinati valori discreti, come si suole dire sono “quantizzate”. Per esempio, come si è già descritto, in accordo al principio di indeterminazione e alle regole di **Niels Bohr**, le energie atomiche assumono solo valori discreti. In un insieme di molte particelle - sistemi di carattere macroscopico - gli effetti quantistici di solito diventano difficili da evidenziare. Infatti, a causa del moto termico caotico si produce una sorta di media della grandezza considerata su un grande numero di suoi possibili valori e i salti quantici vengono mascherati.

Al contrario, quando un numero molto elevato di microparticelle può muoversi in maniera coordinata in un unico stato, la quantizzazione viene resa manifesta anche a livello macroscopico. Un esempio di questo tipo è fornito dalla quantizzazione del flusso magnetico in un superconduttore.

Il flusso magnetico Φ è ben noto a coloro che hanno studiato il fenomeno dell'induzione elettromagnetica. Si può definire flusso magnetico la quantità

$$\Phi = B S,$$

dove B è il modulo del vettore d'induzione magnetica e S è l'area della superficie delineata dal circuito (per semplicità consideremo l'induzione diretta lungo la normale alla superficie). Per molti sarà una scoperta sapere che il flusso magnetico prodotto dalla corrente superconduttrice che corre lungo un anello può prendere soltanto determinati valori discreti. Cerchiamo di comprendere questo fenomeno, anche se in modo semplificato. Per far ciò sarà sufficiente utilizzare la descrizione del moto di elettroni sulle orbite quantistiche che corrispondono, in questa semplificazione, alle nubi di probabilità di presenza discusse negli atomi.

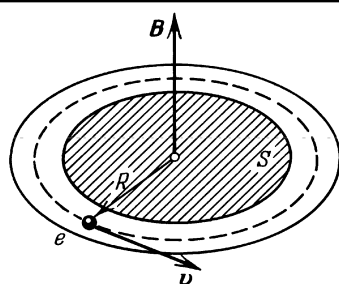
Il moto degli elettroni superconduttori nell'anello (Fig. 24.2) è simile al moto degli elettroni nell'atomo: è come se gli elettroni si muovessero su orbite gigantesche di raggio R senza alcuna collisione. Per questo è natu-

rale supporre che il loro moto obbedisca alle stesse leggi di quantizzazione del movimento degli elettroni nell'atomo. Secondo il postulato di Bohr tali orbite si selezionano mediante la seguente regola di quantizzazione: il prodotto del modulo dell'impulso dell'elettrone mv per il raggio dell'orbita R (la grandezza $m\mathbf{v} \times \mathbf{R}$ si chiama momento angolare dell'elettrone) può prendere soltanto valori discreti, ossia

$$mvR = n\hbar. \quad (24.1)$$

Qui n è un numero positivo intero, mentre la grandezza $\hbar$ che definisce il più piccolo valore (*quanto*) del momento angolare, è la costante di **Planck**, che abbiamo conosciuto nel descrivere il principio di indeterminazione. La quantizzazione di tutte le grandezze fisiche è determinata appunto da questa costante universale.

Fig. 24.2: L'elettrone in un anello conduttore.



Valutiamo ora il valore del quanto elementare di flusso. A tale scopo consideriamo un singolo elettrone sull'anello e immaginiamo di incrementare progressivamente il campo magnetico e quindi il flusso. Come noto si produce la forza elettromotrice indotta:

$$\mathcal{E}_i = -\frac{\Delta\Phi}{\Delta t}.$$

Il valore del campo elettrico è

$$E = \frac{\mathcal{E}_i}{2\pi R} = -\frac{\Delta\Phi}{2\pi R \Delta t}.$$

In virtù della seconda legge di Newton l'accelerazione della particella di

carica e risulta

$$ma = m \frac{\Delta v}{\Delta t} = - \frac{e \Delta \Phi}{2\pi R \Delta t}.$$

Eliminando Δt otteniamo

$$\Delta \Phi = - \frac{2\pi m \Delta v R}{e} = - \frac{2\pi}{e} \Delta (m v R).$$

Cosicchè il flusso magnetico concatenato con l'anello risulta proporzionale al momento angolare e in accordo alla quantizzazione di Bohr potrà assumere solo valori discreti, secondo la relazione (si veda la formula (24.1))

$$\Phi = (-) \frac{2\pi}{e} n \hbar. \quad (24.2)$$

Il valore corretto del quanto di flusso Φ_0 viene ottenuto sostituendo nella formula (24.2) la carica elettrica di una coppia:

$$\Phi_0 = \frac{2\pi \hbar}{2e} = 2.07 \times 10^{-15} \text{ Wb}.$$

Il quanto di flusso magnetico, definito nella formula precedente e anche chiamato in gergo “flussone”, è una grandezza molto piccola. Tuttavia gli strumenti moderni permettono di osservare sperimentalmente il fenomeno della quantizzazione. Un esperimento in tale senso è stato condotto da **Deaver** e **Fayrbank** (1961), con la differenza che al posto dell'anello essi hanno usato un tubo cavo superconduttore, su cui circolavano correnti. In questo esperimento è stato rilevato che il flusso magnetico attraverso l'area della sezione trasversale del tubo cambiava con un andamento a gradini.

Tuttavia in quell'esperimento la grandezza del quanto di flusso risultava la metà di quella da noi dedotta. La spiegazione di ciò è fornita dalla moderna teoria della superconduttività. Come si è già detto, nello stato superconduttore gli elettroni si uniscono in coppie ed è proprio il moto delle coppie che costituisce la corrente superconduttrice. Per questo si deduce il valore corretto del quanto del flusso magnetico solo se nella formula della quantizzazione si sostituisce la carica dell'elettrone con una carica doppia.

Lo scienziato inglese **Fritz London**, che prevede teoricamente la quantizzazione del flusso già nel 1950 (molto tempo prima che fosse compresa la reale natura dello stato superconduttore) aveva anch'egli implicitamente assunto che fossero gli elettroni, e non le coppie, a condurre la corrente.

È necessario sottolineare che la deduzione della quantizzazione del flusso magnetico da noi ottenuta, sebbene rifletta correttamente l'essenza fisica

del fenomeno, è peraltro assai semplificata. In particolare, noi l'abbiamo dedotta nei riguardi del campo magnetico associato a una corrente superconduttiva che percorre l'anello. In realtà essa vale anche per un qualunque campo magnetico esterno che venga applicato nella regione racchiusa dall'anello. Se immaginiamo di raffreddare al di sotto di T_c , dopo aver applicato il campo stesso, ebbene la quantizzazione significa che all'atto dell'entrata nello stato superconduttore vengono lanciate delle correnti che si aggiustano in modo che il flusso concatenato dell'anello sia comunque un multiplo intero di Φ_0 .

Come già abbiamo menzionato nella introduzione storico-descrittiva, nel 1933 **Meißner** e **Ochsenfeld** avevano scoperto che il campo magnetico applicato al superconduttore a temperature elevate veniva completamente espulso dal materiale quando si diminuiva la temperatura al di sotto di quella critica. Con riferimento a questo effetto è possibile comprendere la ragione del fatto che un campo magnetico può distruggere la superconduttività se portato al di sopra di un valore critico (che dipende dalla temperatura). Poiché per eccitare le correnti superficiali che schermano il campo occorre spendere una certa energia, è evidente che lo stato superconduttivo è energeticamente meno favorevole di quello normale, nel quale il campo penetra nel materiale. Quanto maggiore è il campo esterno tanto più intensa deve essere la corrente sulla superficie per garantire lo schermaggio. A un certo valore del campo, la superconduttività si autodistrugge e il metallo passa ad uno stato normale. Il campo al quale questo si verifica si chiama campo critico del superconduttore.

È importante peraltro osservare che per distruggere la superconduttività non è necessario un campo magnetico esterno. La corrente che circola lungo il conduttore crea essa stessa un campo magnetico. Quando ad un valore determinato della corrente, l'induzione di questo campo raggiunge il valore pari al campo critico, la superconduttività si annulla.

Il valore del campo critico cresce con la diminuzione della temperatura. Tuttavia anche vicino allo zero assoluto, il campo critico dei superconduttori costituiti da metalli ordinari non è alto. Nei casi più favorevoli esso è soltanto dell'ordine del migliaio di Gauss. Sembrerebbe così impossibile creare mediante i superconduttori elevati campi magnetici.

24.3 I vortici di Abrikosov

Come si è già fatto cenno alla sezione 24.1 di questo capitolo, nel 1957 il fisico teorico A.A. Abrikosov ha mostrato che in alcune leghe si poteva aumentare il valore del campo magnetico senza distruggere la superconduttività. Anche in tali leghe, analogamente al caso dei superconduttori costituiti dai metalli puri, superato un valore minimo il campo magnetico inizia a penetrare all'interno. Tuttavia, nelle leghe il campo non penetra immediatamente in tutto il volume del superconduttore. Dapprima, all'interno si formano soltanto singoli "tubicini" di linee di flusso magnetico (Fig. 24.3). In ognuno di questi tubicini è contenuta una porzione assolutamente determinata, pari al quanto di flusso magnetico $\Phi_0 = 2 \cdot 10^{-15} \text{ Wb}$.

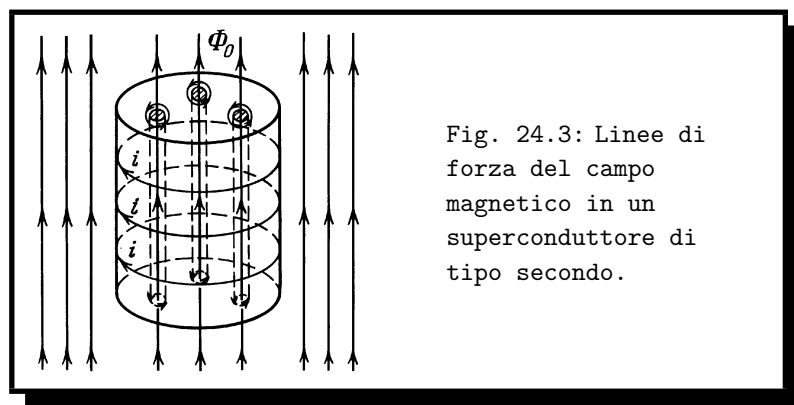


Fig. 24.3: Linee di forza del campo magnetico in un superconduttore di tipo secondo.

Quanto maggiore è il campo magnetico esterno tanti più tubicini e, di conseguenza, tanti più quanti di flusso magnetico penetrano nel superconduttore. Per questo, il flusso magnetico nel superconduttore cambia non con continuità ma a scatti, con discretezza. Qui ci imbattiamo nuovamente in un fenomeno nel quale le leggi della meccanica quantistica operano su scala macroscopica. Ogni tubicino delle linee del campo di induzione magnetica in un superconduttore è avvolto da correnti circolari che non si smorzano, che ricordano i vortici in un liquido. Ecco perché questi coaguli di linee, circondati da correnti superconduttrici, si chiamano anche *vortici di Abrikosov*. All'interno di ogni vortice la superconduttività, naturalmente, non sussiste. Tuttavia, nello spazio tra i vortici essa si conserva! Soltanto in presenza di campi molto elevati, quando i vortici diventano

così numerosi che iniziano a sovrapporsi, avviene la totale distruzione della superconduttività.

Questo insolito aspetto del fenomeno della superconduttività nelle leghe, in presenza di un campo magnetico, è stato scoperto per la prima volta lavorando solo con carta e matita, in modo squisitamente predittivo.

Oggi la tecnica sperimentale permette di osservare i vortici di Abrikosov direttamente. Per far ciò, si pone un sottile strato di polvere magnetica sulla superficie di un superconduttore (per esempio, la sezione trasversale di un cilindro). Le particelle della polvere si addensano in quelle regioni dove è penetrato il campo magnetico. La dimensione di ogni regione non è grande, di solito una frazione di micrometro. Se si guarda la superficie con un microscopio elettronico si vedranno delle macchie scure.

La figura 24.4 è una fotografia della struttura dei vortici di Abrikosov, ottenuta con questo metodo. Si vede che i vortici sono disposti periodicamente e formano una struttura reticolare analoga a quella in un reticolo cristallino. Il reticolo dei vortici è triangolare (cioè può essere ottenuto mediante una trasposizione di triangoli regolari).

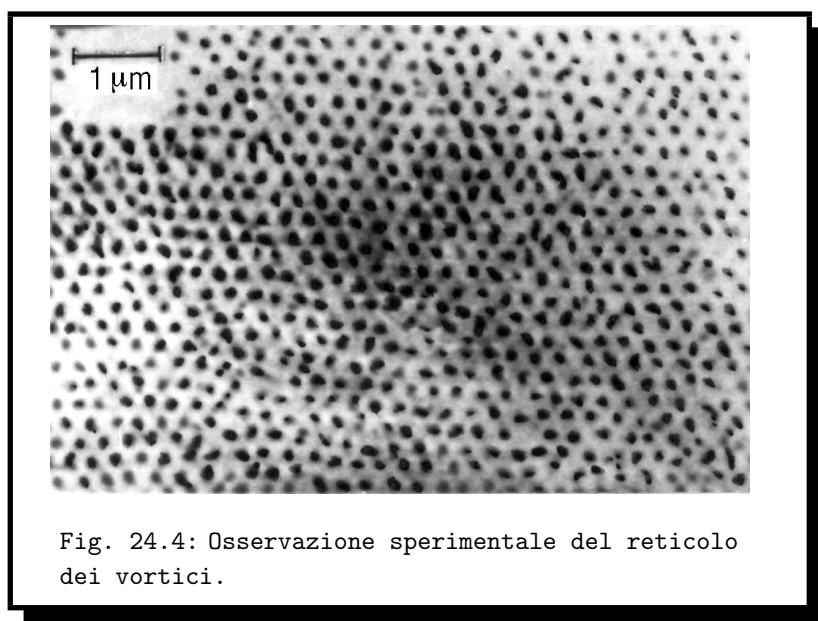


Fig. 24.4: Osservazione sperimentale del reticolo dei vortici.

Ancora più raffinata è la visualizzazione oggi ottenibile con una specie di

microscopio magnetico, nel quale elettroni vengono fatti deviare dalle linee di flusso dei vortici e attraverso successivi effetti di interferenza creano una immagine diretta degli stessi!

Quindi, ritornando alle leghe, a differenza dei metalli puri esse sono caratterizzate da due campi critici: il campo critico inferiore, al valore del quale il primo vortice penetra nel superconduttore e il campo critico superiore, che provoca una distruzione totale della superconduttività. Nell'intervallo tra questi due valori il superconduttore è penetrato da linee di vortici e si trova in un particolare *stato misto*. Come si è detto, un superconduttore con queste proprietà oggi viene chiamato superconduttore di secondo tipo, a differenza dei superconduttori di primo tipo, in cui la distruzione della superconduttività nel campo magnetico avviene totalmente, a un valore definito del campo.

Negli anni 1950 aveva preso l'avvio una vera caccia ai materiali superconduttori con alte temperature critiche e elevati valori dei campi critici. Quali metodi non sono stati utilizzati per ottenerli! Saldatura ad arco, raffreddamento rapido, polverizzazione su supporto bollente! Le leghe Nb₃Sn, Nb₃Se e Nb₃Al sono caratterizzati da campi critici superiori a 20 Tesla. Un ulteriore incremento si è avuto con i composti tripli. Infine, come si è già detto, la situazione è radicalmente mutata con i superconduttori di nuova generazione, dopo il 1987. Essi sono tutti, e marcatamente, superconduttori di tipo secondo, o di Abrikosov, con campi critici superiori elevatissimi.

Sembrerebbe così che il problema dei magneti superconduttori sia risolto. Invece la natura ha posto ancora un ostacolo sulla strada dei ricercatori. Infatti, per il solenoide superconduttore è necessario un filo che sopporti non solo un forte campo magnetico, ma anche tale che non vi si produca alcuna dissipazione. Invece i vortici, come vedremo, possono dissipare energia perché possono muoversi quando la corrente elettrica passa loro vicino.

24.4 Il “pinning”, i solenoidi superconduttori e le loro funzioni

In linea di principio si possono ottenere intensi campi magnetici facendo passare attraverso una bobina una forte corrente. Quanto maggiore è la corrente, tanto maggiore è il campo magnetico che essa crea. Tuttavia se la bobina possiede una resistenza elettrica, in essa si dissipa calore. Occorre consumare un'enorme energia per mantenere la corrente e sorgono seri

problemi per lo smaltimento del calore (altrimenti la bobina può fondere). Nel 1937 per la prima volta era stato ottenuto un campo magnetico con induzione di 10 T (quindi 100.000 Gauss). In realtà si riusciva a mantenere questo campo solamente durante la notte, quando gli altri utenti della centrale elettrica che alimentava la bobina erano disinseriti. Il calore liberatosi veniva eliminato con corrente d'acqua e ogni secondo 5 litri di acqua venivano portati ad ebollizione. Queste considerazioni mostrano come le possibilità di ottenere intensi campi magnetici in solenoidi costituiti da metalli normali siano limitate.

L'idea di impiegare la superconduttività per creare intensi campi magnetici è nata subito dopo la sua scoperta. Tutto quello che sembrava necessario era avvolgere una bobina di cavo superconduttore, chiudere le sue estremità e lanciare in questo circuito una corrente sufficientemente forte. Poiché la resistenza elettrica della bobina è uguale a zero, non si ha emissione di calore. E sebbene il raffreddamento del solenoide occorrente per instaurare lo stato di superconduzione comporti qualche difficoltà, i vantaggi ammortizzerebbero gli svantaggi. Peccato che il campo magnetico distruggesse esso stesso la superconduzione!

Vediamo invece cosa accade in presenza dei vortici di Abrikosov che abbiamo appena descritto. Dalla fisica elementare sappiamo che su un conduttore in cui circola corrente, se posto in un campo magnetico, agisce una forza. Dove è la forza di reazione, che deve comparire in base alla terza legge di Newton? Se, per esempio, il campo magnetico è creato da un altro conduttore in cui circola corrente, su di esso agisce una forza di uguale modulo, ma di direzione opposta (le forze tra conduttori in cui circola corrente vengono definite dalla legge di **Ampere**). Nel nostro caso la situazione è molto più complessa. Quando un superconduttore si trova in uno stato misto ed è attraversato da corrente, nelle zone dove si ha campo magnetico (all'interno dei vortici) compaiono forze di interazione tra la corrente e il campo. Conseguentemente, la distribuzione della corrente cambia ma anche le zone in cui è concentrato il campo magnetico non possono rimanere immobili. I vortici di Abrikosov, sotto l'azione della corrente si muovono. Come è noto, la forza agente sulla corrente è perpendicolare al campo magnetico stesso e alla direzione del filo conduttore. La forza agente sui vortici di Abrikosov è anch'essa perpendicolare al campo magnetico e alla direzione della corrente. Se, per esempio, in un superconduttore allo stato misto si crea una corrente che va da destra a sinistra, i vortici di Abrikosov nella

Fig. 24.4 inizieranno a muoversi dal basso verso l'alto o dall'alto verso il basso (a seconda della direzione del campo magnetico).

Osserviamo ora che il movimento del vortice di Abrikosov attraverso il superconduttore corrisponde allo spostamento della parte normale, non superconduttiva! Questo movimento comporta una variazione di flusso, che genera un campo elettrico, il quale provoca una dissipazione delle correnti superconduttive. Quindi, quando la corrente circola in un superconduttore che si trova nello stato misto, compare comunque una resistenza. Pertanto non si potrebbero usare questi materiali per realizzare il solenoide superconduttore.

Che fare? Bisogna impedire ai vortici di muoversi, ancorandoli laddove essi si producono. Questo si può ottenere danneggiando microscopicamente il superconduttore, creando cioè dei difetti. I difetti di solito compaiono da soli, a seguito della lavorazione meccanica o termica del materiale. La figura 24.5, per esempio, mostra una fotografia al microscopio elettronico di una pellicola di nitrato di niobio (la cui temperatura critica è di 15 K), ottenuta con la polverizzazione del metallo su una piastrina di vetro. Si vede chiaramente la struttura granulare (a colonna) del materiale. Per il vortice è molto difficile saltare attraverso il limite del grano. Ecco perché, fino ad un determinato valore della corrente (che si chiama corrente critica), i vortici rimangono immobili. La resistenza elettrica in questo caso è uguale a zero e lo stato è realmente di superconduzione. Il fenomeno di ancoraggio dei vortici si chiama *pinning*, che in inglese corrisponde a “fissare al muro o ad una tavola con puntine”.

Grazie al pinning si possono ottenere materiali superconduttori con elevati valori sia del campo critico sia della corrente critica. In questo caso, se il valore del campo critico è determinato dalle proprietà dello stesso materiale, il valore della corrente critica (più precisamente, la sua densità, ossia la corrente per sezione unitaria) dipende fortemente dal metodo di preparazione e di lavorazione del materiale. Sono state messe a punto tecnologie che permettono di ottenere materiali superconduttori caratterizzati da elevati valori di tutti i parametri critici. Per esempio, con una lega di niobio e stagno si può ottenere un materiale con densità di corrente critica di centinaia di migliaia di ampere per centimetro quadrato, campo critico superiore a 25 T e temperatura critica 18 K . Analoghi risultati sono già stati realizzati con i superconduttori ad alta temperatura.

Sono importanti anche le proprietà meccaniche del materiale. La lega di niobio e stagno di per sé è fragile e quindi diventa difficile realizzare un cavo

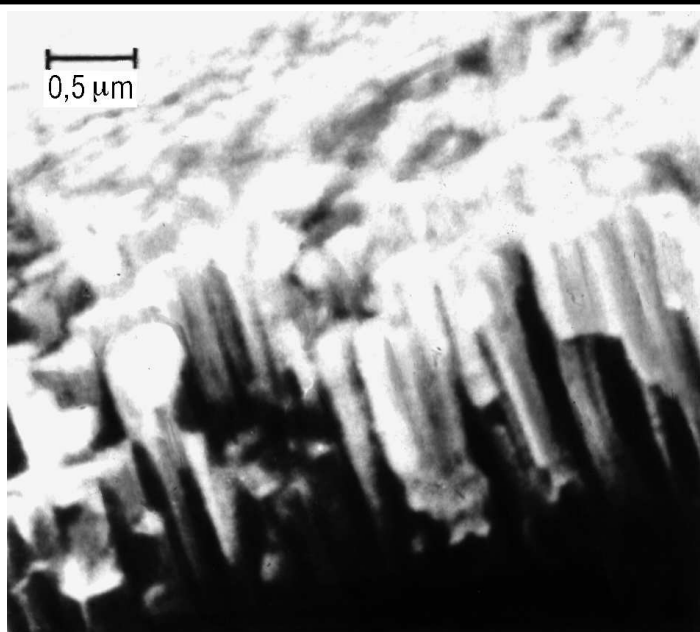


Fig. 24.5: Struttura microscopica di un film di nitrato di Niobio.

o una bobina con questo materiale. Per questo i solenoidi superconduttori sono usualmente fatti nel seguente modo: la polvere di niobio e stagno viene pigiata in un tubo di niobio. Quindi il tubo viene teso fino a diventare un filo, si avvolge la bobina e si ottiene un solenoide della lega Nb_3Sn .

Nell'industria si usano materiali più pratici dal punto di vista tecnologico, per esempio leghe di niobio titanio ($NbTi$), che possiedono una soddisfacente plasticità. Mediante questa lega si realizzano i così detti superconduttori compositi. In una barretta di rame vengono praticati molti fori in cui vengono introdotti pezzi di superconduttore. Quindi la barretta viene stirata fino ad ottenere un lungo filo. Il filo viene tagliato a pezzi e di nuovo viene introdotto in una barretta di rame, di nuovo stirato, tagliato a pezzi e così via. In questo modo si ottiene un cavo contenente fino a un milione di anime superconduttrici, che si arrotola, ottenendo una bobina. Il vantaggio consiste nel fatto che la corrente superconduttrice si distribuisce

nei cavi percorrendo tutte le anime. Rispetto ad un superconduttore anche il rame è un isolante, e in caso di unione in parallelo tra un conduttore di rame e un superconduttore tutta la corrente circola nel superconduttore, che ha resistenza nulla. Si ha anche un altro vantaggio. Ipotizziamo che in una qualche anima si sia distrutta la superconduttività. Quindi si ha un'emissione di calore e per evitare il passaggio di tutto il cavo allo stato normale è necessario che il calore prodotto localmente venga rapidamente rimosso. Il rame, che è un buon conduttore di calore, risolve brillantemente questo problema realizzando una stabilizzazione termica. Oltre a ciò, il rame garantisce buone qualità generali dei cavi.

Gli impieghi dei magneti superconduttori sono molto vari. Essi giocano un ruolo importante nella fisica delle alte energie, nello studio delle proprietà fisiche dello stato solido, sono impiegati in elettronica e anche nei trasporti. Abbiamo già fatto cenno ai progetti di treni su cuscinetti magnetici. I solenoidi superconduttori, sistemati in un vagone, creano un potente campo magnetico che col movimento del treno inducono correnti su binari speciali. In accordo alla legge di **Lentz**, il campo magnetico di queste correnti è diretto in maniera tale da ostacolare l'avvicinamento del solenoide al binario e il treno ... rimane sospeso sul piano di posa. Su una linea sperimentale di 40 chilometri, in Giappone, è stato raggiunto il record mondiale di velocità: 540 Km l'ora!

Un treno a levitazione diamagnetica mediante solenoidi superconduttori già opera a Shangai, collegando l'aeroporto al centro cittadino. Un prototipo di serie di un treno a tre vagoni che farà servizio sulla linea Tokio - Osaka è già stato sperimentato. In ogni vagone si trova un impianto refrigerante, accoppiato a un solenoide superconduttore in un modulo strutturale.

In elettrotecnica l'impiego dei magneti superconduttori diventa conveniente per la realizzazione di motori e generatori elettrici di elevatissima potenza, centinaia e più di megawatt. Gli avvolgimenti superconduttori di uno statore creano un forte campo magnetico costante, nel quale ruota un rotore di metallo normale. In questo modo si raggiunge una notevole riduzione delle dimensioni e della massa dell'impianto. Motori simili, con una potenza di alcuni megawatt, sono stati già realizzati. Sono allo studio progetti di macchine ancora più potenti. Vantaggi ancora maggiori offre l'impiego di un avvolgimento superconduttore; tuttavia per la realizzazione di questa idea sorgono alcuni problemi tecnici. Sappiamo che il campo magnetico agisce su particelle cariche in movimento con la forza di Lorentz. La direzione di questa forza è perpendicolare alla velocità della

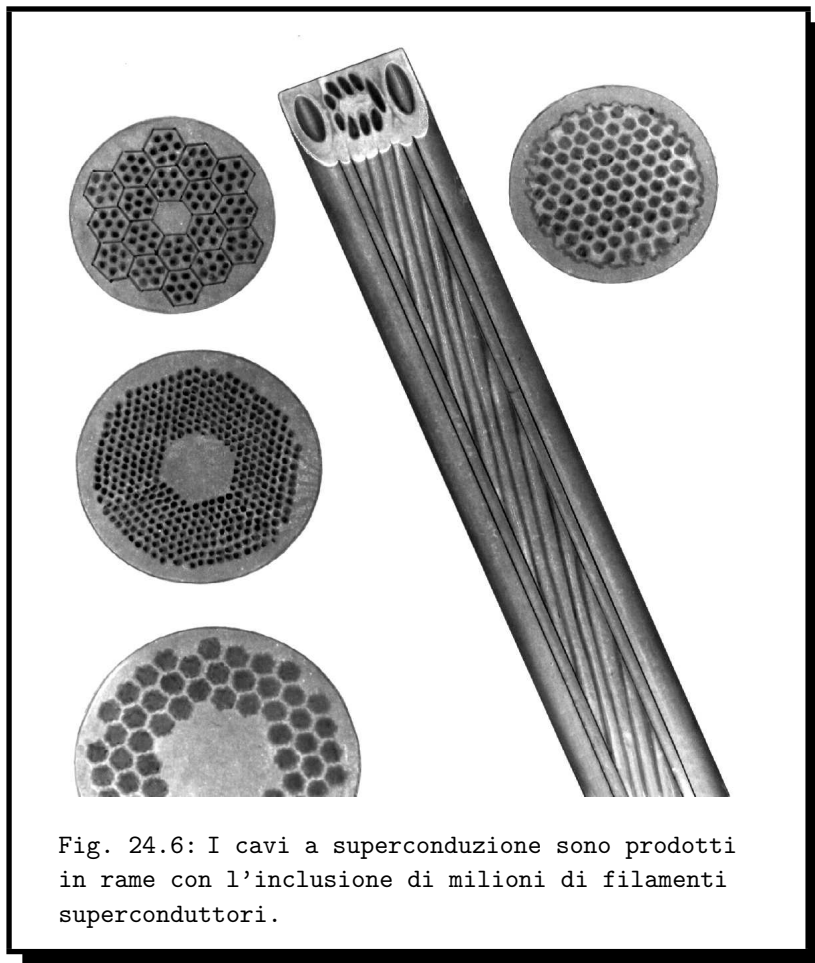
particella e ne incurva la sua traiettoria. Quanto maggiore è il campo, tanto minore è il raggio della circonferenza lungo cui si muove la particella. Proprio questo principio del contenimento delle particelle viene utilizzato negli acceleratori, nelle camere a bolle e negli impianti di fusione termonucleare controllata. I vantaggi dell'impiego di magneti superconduttori a tali finalità, senza elevato consumo di energia, sono evidenti. In Russia è già in funzione un sistema di fusione termonucleare "Tokamak-7" ed è stato progettato l'impianto "Tokamak-15", in cui si accumulerà un'energia magnetica di 600 milioni di Joule. La realizzazione di successive generazioni di sistemi analoghi, previste con energie più alte, senza l'impiego della superconduttività sarebbe impossibile.

Spesso nello studio dei solidi, delle molecole o di atomi e di nuclei è necessario creare elevati campi magnetici molto omogenei. I magneti superconduttori, in questi casi, sono insostituibili e oggi vengono largamente utilizzati nei laboratori. Piccoli solenoidi ultrasensibili accoppiati ad un sistema di raffreddamento sono già produzione industriale, in continuo aumento. Magneti superconduttori sono necessari, come vedremo, per produrre i campi magnetici impiegati nella visione all'interno del nostro corpo (*NMR imaging*).

L'energetica del futuro non è rappresentata solo da nuove fonti di energia. È necessario altresì mettere a punto efficaci metodi di conservazione e trasmissione dell'energia elettrica. Anche a questi scopi i superconduttori propongono i propri servizi. È stato messo a punto un progetto di un sistema di conservazione di energia elettrica, nel quale una gigantesca bobina superconduttrice, con diametro superiore a 100 metri, verrà sistemata in un tunnel *ad hoc*, scavato nelle montagne. In esso, mediante impianti refrigeranti ad elio liquido sarà mantenuta una temperatura vicina allo zero assoluto. In questa bobina una corrente superconduttrice, quindi senza dissipazione e non-smorzata, accumulerà un'energia di 100 Megawatt-ora.

E per ciò che riguarda il trasporto di energia senza dissipazione lungo cavi superconduttori? Le valutazioni mostrano che questo metodo diventa redditizio quando si trasmettono potenze superiori a migliaia di megawatt su una distanza di più di 20 chilometri. Molti prototipi di questi cavi sono già stato realizzati (si veda la figura 24.6). Per il momento si possono ipotizzare linee di trasmissione elettrica che trasportano l'elettricità senza perdite. Con i superconduttori a elevata temperatura critica già si producono cavi o film che, sfruttando l'azoto liquido, possono trasportare correnti all'incirca pari a un milione di Ampere per cm^2 .

Alle Olimpiadi di Pechino del 2008 è previsto che tutta la energia elettrica necessaria al funzionamento delle stesse e del villaggio olimpico sarà distribuita attraverso una rete a superconduzione.



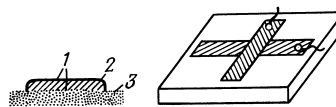
24.5 Effetto Josephson

Esaminiamo ora un altro fenomeno quantistico che si realizza in un superconduttore e che è servito per realizzare una serie di strumenti di misura veramente particolari. Si tratta dell'effetto scoperto nel 1962 da **Brian David Josephson**.

Immaginate che su una piastra di vetro (come si suole dire, su un substrato) sia stato deposto un film superconduttore (che di solito si ottiene evaporando un metallo nel vuoto). Esso viene quindi ossidato, creando sulla superficie uno strato di dielettrico (ossido) con uno spessore di alcune decine di Angstrom (vale a dire 10^{-8} cm), quindi di nuovo coperto con un film superconduttore. Si è ottenuto il cosiddetto sandwich.

Proprio nei sandwich si realizza l'effetto Josephson (per comodità di misura, esso di solito è realizzato in forma di croce, come si vede nella figura 24.7). Cominciamo ad esaminare il caso in cui i films siano allo stato normale (non superconduttore).

Fig. 24.7: La giunzione Josephson:
1 -- film metallico; 2
-- strato di ossido; 3
-- substrato.



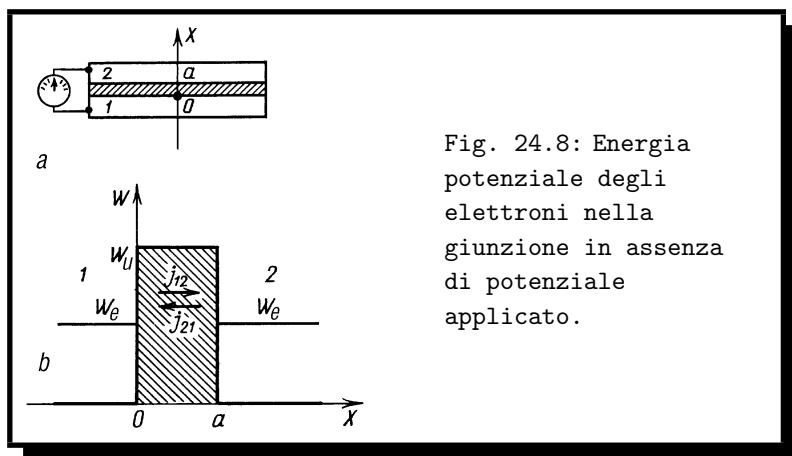
Gli elettroni di un film metallico (Fig. 24.8, *a*) possono passare nell'altro film? Sembrerebbe di no, in quanto ne sono impediti dallo strato di dielettrico posto tra di essi. Per uscire dal metallo l'elettrone deve possedere un'energia maggiore del lavoro d'estrazione, ma a basse temperature praticamente non esistono elettroni con simile energia.

Nella figura 24.8, *b* è mostrato il grafico della dipendenza dell'energia potenziale dell'elettrone dalla coordinata x (l'asse x è perpendicolare al piano del sandwich). Nel metallo l'elettrone si muove liberamente e la sua energia potenziale è uguale a zero. Per uscire nel dielettrico occorre compiere un lavoro d'estrazione W_u che è maggiore dell'energia cinetica e quindi anche dell'energia totale dell'elettrone W_e^c . Per questo si dice che

^cAnalogamente, per estrarre una molecola da un liquido si deve compiere il lavoro di evaporazione.

gli elettroni, nei film metallici, sono divisi da una barriera potenziale la cui altezza è uguale a $W_u - W_e$.

Se gli elettroni seguissero le leggi della meccanica classica, questa barriera sarebbe per loro insuperabile. Tuttavia gli elettroni sono particelle e, come abbiamo più volte ricordato, nel micromondo agiscono leggi particolari che permettono cose impossibili ai corpi macroscopici. Per esempio, un uomo con una tale energia non potrebbe certo scavalcare una montagna, ma l'elettrone può attraversarla! È come se l'elettrone scavasse sotto la montagna un tunnel e vi penetrasse, anche se la sua energia non è sufficiente per arrampicarsi sulla cima. L'*effetto tunnel* (così si chiama questo fenomeno) si spiega con le proprietà ondulatorie delle particelle (il loro "distribuirsi" nello spazio), e può essere compreso compiutamente soltanto studiando la meccanica quantistica. Rimane il fatto: esiste una certa probabilità che gli elettroni possano attraversare il dielettrico. Questa probabilità è maggiore quanto minore è l'altezza $W_u - W_e$ della barriera e quanto minore è la sua larghezza a .

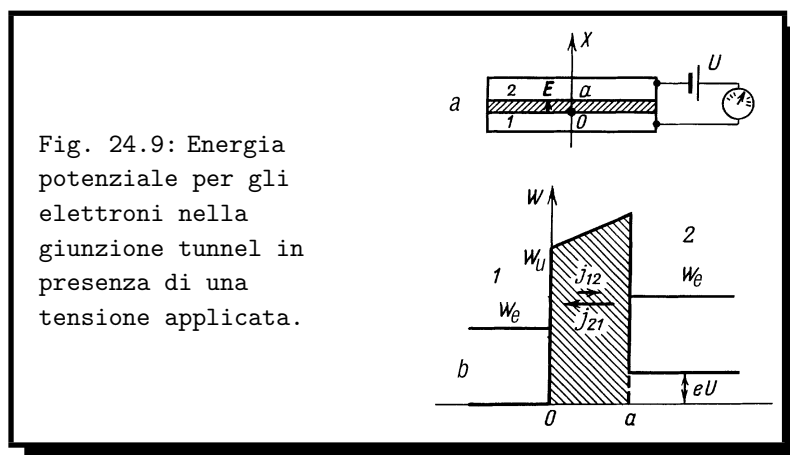


Pertanto, la sottile pellicola del dielettrico è come trasparente agli elettroni e attraverso di essa può passare la cosiddetta corrente tunnel. Tuttavia la corrente totale tunnel è uguale a zero: il numero di elettroni che passano dallo strato metallico inferiore a quello superiore, è uguale, in media, al numero di quelli che passano, al contrario, da quello superiore a quello inferiore.

Come rendere la corrente tunnel diversa da zero? Per far ciò bisogna

rompere la simmetria, per esempio collegare i film metallici ad un alimentatore, di modo che si stabilisca tra loro una tensione U (Fig. 24.9, *a*).

In tale situazione i film giocheranno il ruolo delle facce di un condensatore e nello strato del dielettrico comparirà un campo elettrico $E = U/a$. Il lavoro necessario allo spostamento della carica su una distanza x lungo la direzione del campo è dato da $A = Fx = eEx = eUx/a$. Il grafico dell'energia potenziale dell'elettrone diventa quello mostrato nella figura 24.9, *b*. Come si vede, per gli elettroni dello strato superiore ($x > a$) è più facile superare la barriera che non per gli elettroni dello strato inferiore ($x < 0$), che devono compiere un salto energetico maggiore. Di conseguenza, anche in presenza di piccole tensioni di alimentazione si innesca una corrente tunnel.



Effetti tunnel in metalli ordinari vengono utilizzati in alcuni strumenti. Ora vogliamo illustrarvi il loro impiego pratico nella superconduttività. Per questo faremo ancora un passo e immagineremo che gli strati metallici, separati dal sottile strato di dielettrico, si trovino allo stato di superconduzione. Come si comporterà il contatto tunnel superconduttore? Vedremo che la superconduttività determina effetti del tutto inaspettati.

Come si è già detto, gli elettroni del film superiore hanno un'energia eU in eccesso rispetto a quelli dello strato inferiore. Quando passano nello strato inferiore devono cedere energia per essere allo stato di equilibrio. Se il film si trovasse allo stato normale sarebbero sufficienti poche collisioni con il reticolo cristallino e l'energia in eccesso si trasformerebbe in calore.

Ora supponiamo che il film si trovi nello stato di superconduzione e quindi non è in grado di trasferire l'energia al reticolo (come abbiamo detto, in tale stato non sono efficaci le ordinarie "collisioni" degli elettroni con gli ioni del reticolo cristallino). Quindi agli elettroni non resta che emettere questa energia sotto forma di quanto di radiazione elettromagnetica.

La frequenza è legata alla tensione applicata U dalla relazione

$$\hbar\omega = 2eU.$$

Avrete notato che a destra si è scritta la carica dell'elettrone moltiplicata per due. Infatti non i singoli elettroni, ma le coppie superconduttrici si muovono nell'effetto tunnel.

Ecco il meraviglioso effetto predetto da Josephson: la tensione costante, applicata al contatto tunnel superconduttore (che si chiama anche elemento Josephson) porta alla generazione di radiazione elettromagnetica. Sperimentalmente questo effetto è stato osservato per la prima volta nell'Istituto delle basse temperature di Char'kov, dagli scienziati **I.M. Dmitrienko**, **V.M. Svistunov** e **L.K. Janson** nel 1965.

Pertanto, la prima cosa che viene in mente, se si pensa ad un' utilizzazione pratica dell'effetto Josephson, è la realizzazione di un generatore di radiazione elettromagnetica. In realtà, non è così semplice: estrarre la radiazione da una stretta fessura tra pellicole superconduttrici, dove essa si genera, è molto difficoltoso (proprio per questo la rilevazione sperimentale dell'effetto Josephson non era un compito facile). Inoltre l'intensità della radiazione è molto piccola. Per questo oggi gli elementi tipo Josephson vengono in realtà utilizzati soprattutto come rivelatori di radiazione elettromagnetica.

Questo impiego è basato sul fenomeno della risonanza tra le oscillazioni elettromagnetiche esterne dell'onda e le oscillazioni proprie che si innescano nel circuito Josephson quando si applichi una tensione costante. Per la verità, la risonanza è alla base del funzionamento di molti apparecchi: si riesce a "catturare" l'onda quando la sua frequenza coincide con la frequenza del circuito di rilevazione. Come circuito di ricezione è comodo usare un elemento tipo Josephson: è facile infatti regolare la frequenza delle sue oscillazioni proprie (cambiando la tensione); l'acutezza della risonanza, che determina la sensibilità dell'apparecchio, è molto alta. In base a questo principio sono stati realizzati i più sensibili rivelatori di radiazione elettromagnetica, che vengono utilizzati per lo studio della radiazione

dell'universo.

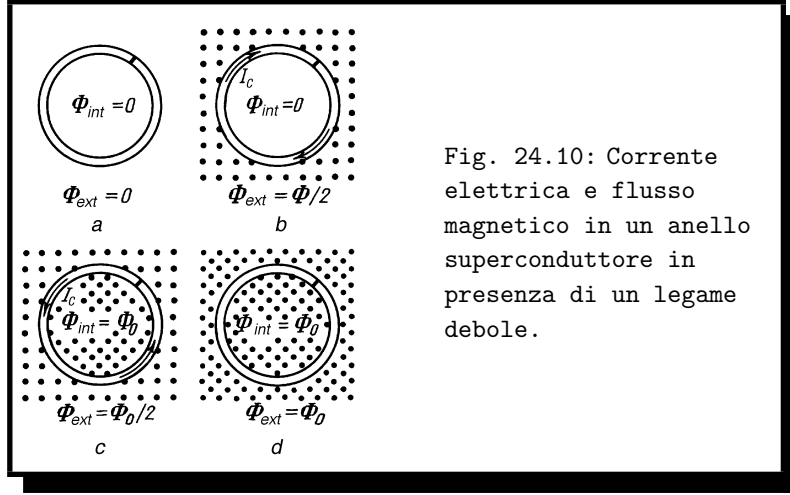
24.6 Lo SQUID, magnetometro quantistico e i suoi impieghi

L'effetto Josephson e il fenomeno di quantizzazione del flusso magnetico si utilizzano per la realizzazione di una famiglia di strumenti di misura ad alta sensibilità. Essi si chiamano “*dispositivi interferenziali quantici superconduttori*” o “*SQUID*” (dalle iniziali delle parole inglesi *Superconducting Quantum Interference Device*). Noi discuteremo il funzionamento di uno di questi strumenti: il magnetometro quantistico (uno strumento per la misura di campi magnetici deboli).

Il magnetometro quantistico più semplice consiste in un anello superconduttore, in cui è incluso un elemento Josephson (Fig. 24.10). Come già sapete, perché si crei una corrente in un normale contatto tunnel, bisogna applicarvi una tensione, anche molto piccola. Invece nel contatto superconduttore ciò non è necessario. Se nell'anello si lancia una corrente superconduttrice, essa potrà circolare anche attraverso l'elemento Josephson: le coppie superconduttrici salteranno per effetto tunnel attraverso il sottile strato del dielettrico. Questo fenomeno si chiama effetto Josephson stazionario (costante nel tempo), in contrapposizione a quello non stazionario, accompagnato dall'emissione di radiazione. Esiste, tuttavia, un valore massimo della corrente superconduttrice (viene definita *corrente critica* del contatto I_c). Con una corrente superiore a quella critica, la superconduttività nel contatto si annulla e ai suoi capi compare una tensione (l'effetto Josephson diventa non stazionario).

Così, includendo nel circuito superconduttore un elemento Josephson, non ha luogo una totale distruzione della superconduttività. Tuttavia nel circuito compare una zona, in cui la superconduttività è, per così dire, indebolita (come si dice, compare un legame debole). Proprio su questo effetto è basato l'impiego dell'elemento Josephson per la misura precisa di campi magnetici.

Cerchiamo di capire come funziona. Se il circuito fosse completamente superconduttore (non contenesse il legame debole), la corrente che attraversa la sua area sarebbe rigorosamente costante. In effetti, in accordo alla legge dell'induzione elettromagnetica, qualsiasi variazione nel campo magnetico esterno induce una forza elettromotrice (f.e.m.) d'induzione $\mathcal{E}_i = -\Delta\Phi_{\text{ext}}/\Delta t$, che comporta una modificazione della corrente nel cir-



cuito. La corrente indotta a sua volta dà origine a f.e.m. di autoinduzione $\mathcal{E}_{si} = -L \Delta I / \Delta t$. Poiché la caduta di tensione nel circuito superconduttore è uguale a zero (la resistenza è nulla) la somma algebrica di queste f.e.m. sarà uguale a zero:

$$\mathcal{E}_i + \mathcal{E}_{si} = 0,$$

ovvero

$$\frac{\Delta \Phi_{\text{ext}}}{\Delta t} + L \frac{\Delta I}{\Delta t} = 0.$$

Si deduce pertanto che, variando il flusso associato al campo magnetico esterno Φ_{ext} , la corrente superconduttiva nel circuito cambia in maniera tale che la variazione del flusso magnetico creato dalla corrente $\Phi_I = L I$ compensa la variazione legata al flusso esterno (regola di Lenz). In questo caso la corrente nel circuito rimane costante e $\Phi_{\text{int}} = \Phi_{\text{ext}} + \Phi_I = \text{cost.}$ Tale corrente non può essere modificata, a meno di portare il circuito allo stato normale (in un circuito superconduttore, ogni variazione di flusso è impedita).

Cosa succede se il circuito superconduttore contiene un legame debole? Ne consegue che la corrente attraverso questo circuito può cambiare.

Vediamo di spiegare come il flusso magnetico concatenato con l'anello superconduttore con legame debole e il valore della corrente possono cambiare quando si ha una variazione di campo magnetico. Supponiamo che

all'inizio il campo esterno e la corrente nel circuito siano uguali a zero (Fig. 24.10, *a*). Allora anche il flusso è uguale a zero. Aumentando il campo esterno, nel circuito compare una corrente superconduttiva e il flusso magnetico ad essa associato compensa esattamente il flusso esterno. Così continuerà fino a quando nel circuito non si raggiungerà il valore critico I_c (Fig. 24.10, *b*). Supponiamo che in questo momento il campo esterno comporti un flusso uguale alla metà del quanto: $\Phi_0/2$ ^d.

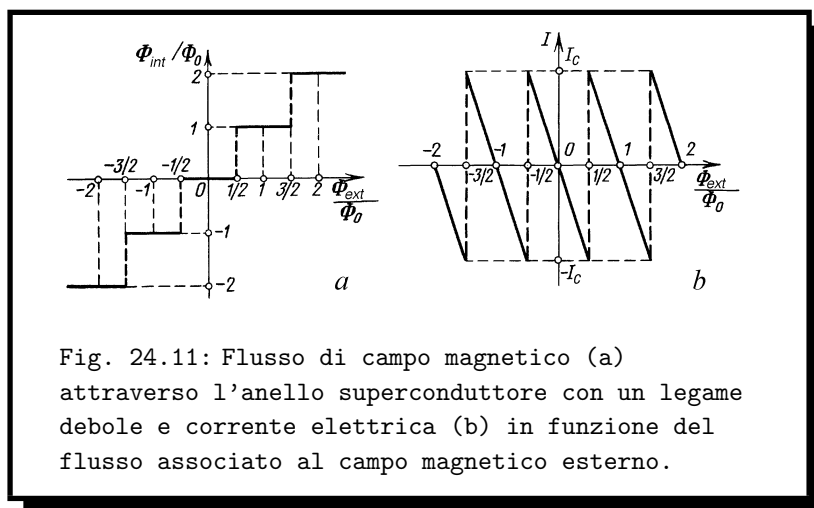
Non appena la corrente diventa uguale a I_c la superconduttività nel punto di legame debole si rompe e nel circuito entra un quanto di flusso Φ_0 (Fig. 24.10, *c*). Il flusso interno Φ_{int} aumenta, in maniera discontinua, di un'unità (il circuito superconduttore passa ad un nuovo stato quantico). Cosa succede alla corrente? Il suo valore rimane invariato, ma essa inverte la propria direzione. In effetti, se fino all'entrata del quanto del flusso Φ_0 la corrente I_c schermava completamente il flusso esterno $\Phi_0/2$, dopo la variazione di flusso essa deve portare il flusso esterno al valore Φ_0 . Per questo al momento dell'entrata del quanto di flusso la direzione della corrente cambia di scatto nel senso opposto.

Aumentando ulteriormente il campo esterno, la corrente nell'anello inizia a diminuire, la superconduttività si ripristina e il flusso all'interno dell'anello rimane uguale a Φ_0 . La corrente nel circuito diventa nulla quando il flusso esterno diventa uguale a Φ_0 (Fig. 24.10, *d*), e quindi essa inizia a fluire nella direzione opposta (nuovamente si produce lo schermaggio). Infine, con un valore del flusso dovuto al campo esterno pari a $3\Phi_0/2$, la corrente diventa di nuovo uguale a I_c la superconduttività si distrugge, entra un nuovo quanto del flusso e così via.

Nella Fig. 24.11 sono mostrati gli andamenti del flusso magnetico all'interno dell'anello e della corrente I al variare del flusso esterno Φ_{ext} (i flussi sono misurati in unità proprie, vale a dire il quanto di flusso Φ_0). Il carattere a gradini della dipendenza permette di "sentire" i singoli quanti del flusso, sebbene la loro entità sia molto piccola. Non è difficile capire perché. Il flusso magnetico all'interno del circuito superconduttore cambia di una piccola quantità $\Delta\Phi = \Phi_0$ ma repentinamente, in un piccolissimo intervallo di tempo Δt . La rapidità di variazione del flusso magnetico $\Delta\Phi/\Delta t$ è molto grande. Si può misurare, per esempio, in base alla f.e.m. d'induzione nella

^dLa corrente critica dipende da molte cause, in particolare dallo spessore dello strato del dielettrico. Cambiando lo spessore si può ottenere che il flusso magnetico creato da questa corrente, e quindi anche il flusso magnetico esterno, sia uguale a $\Phi_0/2$. Ciò semplifica l'analisi senza cambiare la sostanza del fenomeno.

speciale bobina di misura dello strumento. In questo consiste il principio di funzionamento del magnetometro quantistico.



La costruzione di un magnetometro quantistico in realtà è molto più complessa. Di solito non si usa un solo legame debole, ma diversi legami in parallelo: un'interferenza particolare delle correnti superconduttrici (più precisamente, delle onde ad esse relative e che determinano la distribuzione degli elettroni superconduttori nello spazio) porta all'aumento della precisione di misura (questi strumenti sono in realtà interferometri). L'elemento sensibile dello strumento viene collegato induttivamente alla bobina del circuito oscillante dove i salti di flusso si trasformano in impulsi di tensione, che in seguito si rafforzano. Approfondire questi dettagli di carattere più strettamente tecnico ci porterebbe oltre l'ambito del nostro libro.

Facciamo rilevare che magnetometri che misurano campi magnetici con una precisione di 10^{-15} Tesla sono già di produzione industriale. Essi trovano impiego in misure di diverso carattere alle quali faremo ora cenno.

Nel campo della geofisica lo SQUID è utilizzato per la misura dei campi magnetici terrestri, per la ricerca di riserve di petrolio o idrotermiche, in genere per la mappatura di strati di roccia (è evidente come in tali utilizzazioni sia importante disporre di SQUID che utilizzano superconduttori ad alta temperatura, cosicché si debba trasportare solo azoto liquido e non l'elio liquido).

Lo SQUID trova largo impiego nella ricerca scientifica, quando si necessita di accurate misure della suscettività paramagnetica o diamagnetica.

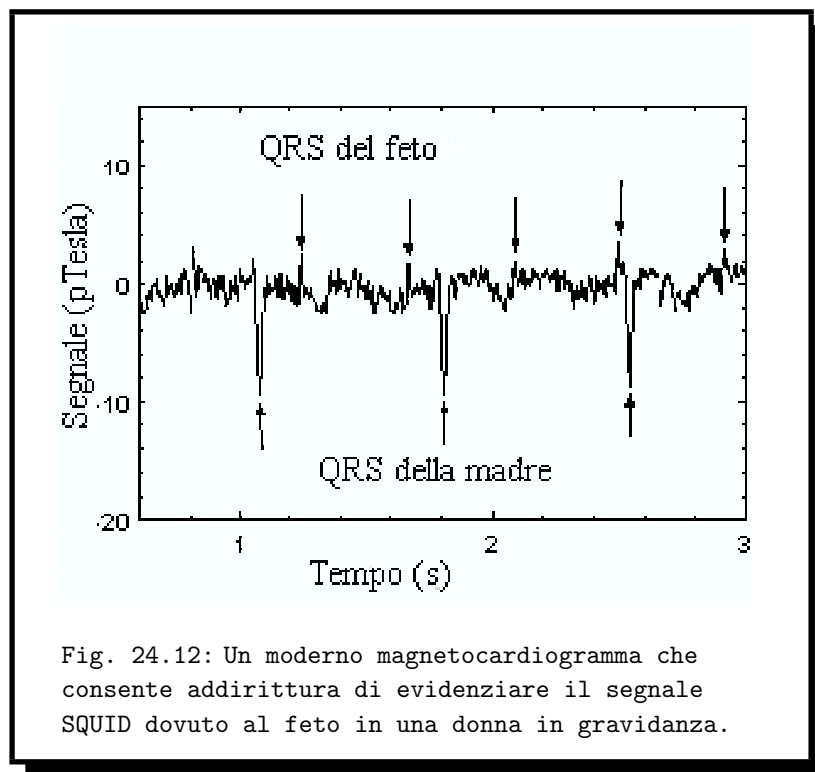
All'incirca la metà del mercato, tuttavia, interessa le applicazioni nel campo medico, per la magnetocardiografia, per la determinazione della quantità di ferro immagazzinata nel fegato, per la magneto-pneumografia o nello studio dei disordini muscolari, delle patologie cerebrali a livello prenatale, nel mapping delle funzioni della corteccia, nel morbo di Alzheimer, nell'epilessia etc.

La sensibilità di uno SQUID è all'incirca $10 - 20 \text{ femtoTesla}/Hz^{1/2}$, oggi raggiunta anche con superconduttori ad alta T_c . Per dare un'idea l'onda R di un cardiogramma o le correnti addominali comportano la generazione di oltre 1000 volte questo valore di sensibilità, su un intervallo di frequenze tipicamente di $0.3 - 80 \text{ Hz}$.

L'attività cerebrale evocata comporta un segnale di circa 50 femtoTesla . Si osservi che lo schermaggio non è indispensabile (in ragione delle tecniche gradiometriche). La mappatura della attività neuronica consente la risoluzione di un volumetto all'incirca del mm^3 . Inoltre mentre l'immagine NMR (della quale parleremo nel prossimo capitolo) fornisce in genere informazioni di carattere strutturale, lo SQUID dà accesso a informazioni di carattere funzionale. (Anche la tomografia a emissione di positroni, PET, dà informazioni funzionali, con una risposta nel tempo, ma molto lenta).

Con lo SQUID si sono già ottenuti significativi successi per la diagnosi di diverse patologie, con localizzazioni dei focolai irritativi in pazienti epilettici o la individuazione delle aree cerebrali responsabili della prosopagnosia. Presto le misure SQUID sostituiranno gran parte della elettrofisiologia ordinaria, in ragione della semplicità delle misure basate sul biomagnetismo, in quanto gli elettrodi non toccano il corpo e non si hanno artefatti legati ai contatti.

Con l'abbinamento di bobine a cavi superconduttori ad alta temperatura con SQUID convenzionali e computer sarà presumibilmente possibile il mapping simultaneo su diverse decine di aree cerebrali (su 40 punti è già stato attuato). Poiché sono già stati osservati i segnali evocati da stimoli sonori, possiamo ipotizzare non lontano il giorno in cui si potremo, per così dire, “vedere i nostri pensieri” !



Capitolo 25

Visione NMR del nostro interno.

25.1 Il magnetismo nucleare

Per comprendere come i nuclei possano avere un *momento angolare* simile a quello già descritto per gli elettroni e quindi anche un associato *momento magnetico*, possiamo immaginare che siano animati da un rapido moto di rotazione rispetto a un asse che passa per il loro baricentro. Essendo i nuclei dotati di carica elettrica ecco svilupparsi un momento magnetico, analogamente a come abbiamo visto al capitolo 24. Anche se questa visualizzazione dei nuclei come piccole trottole magnetizzate è per molti versi troppo semplicistica, ci consentirà di rendere ragione del fenomeno della *risonanza magnetica nucleare* (NMR) al quale abbiamo già fatto cenno a proposito del metodo SNIF per la certificazione della qualità e provenienza dei vini (capitolo 13).

Successivamente discuteremo i principi più generali della tecnica di *imaging*, cioè di “visione” all’interno del nostro corpo, basata appunto sulla risonanza magnetica dei nuclei.

Il momento angolare dei nuclei, più spesso chiamato *spin nucleare*, risulta un multiplo della costante fondamentale di Planck, i.e. $L = \hbar I$ ove I è un numero intero o semi-intero. Il momento magnetico nucleare $\vec{\mu}$ è per la maggior parte dei nuclei parallelo al momento angolare e si può esprimere in funzione di I attraverso una costante γ nota come rapporto giromagnetico: $\vec{\mu} = \gamma \vec{L}(I)$. Il numero quantico I rappresenta, in unità $\hbar$, la massima componente del momento angolare lungo una direzione fissata. Quasi tutti gli elementi hanno almeno un isotopo che possiede un momento magnetico nucleare, con una abbondanza relativa molto variabile. In prat-

ica noi ci riferiremo principalmente al caso del nucleo di idrogeno, cioè del protone, che è assai diffuso in natura. Il suo isotopo è il nucleo di deuterio. Anche il deuterio possiede un momento magnetico, utilizzato come abbiamo accennato per i segnali di risonanza magnetica nel metodo SNIF.

I primi segnali di risonanza magnetica dei nuclei sono stati ottenuti circa 60 anni or sono dai gruppi di **Felix Bloch** a Stanford e di **Edward Purcell** a Harvard. Subito dopo e per la prima volta nell'Europa del dopoguerra, i segnali NMR in acqua furono registrati in Pavia dal gruppo guidato da **Luigi Giulotto**.

A quei tempi le difficoltà sperimentali erano enormi. In Pavia, per ottenere le onde radio necessarie a sollecitare la *inversione* dei momenti nucleari in un campo magnetico (fenomeno, come vedremo, che costituisce la “preparazione” per registrare i segnali di risonanza e “vedere” volumetti del nostro interno) si utilizzava materiale bellico residuo dell'esercito americano e acquistato su mercatini rionali. Anche negli Stati Uniti, benché le condizioni di lavoro fossero ovviamente molto migliori che non nell'Italia del dopoguerra, pur tuttavia si lavorava in maniera largamente artigianale, inimmaginabile rispetto a come oggi si vedono gli apparati NMR in quasi tutti gli ospedali, basati sui solenoidi superconduttori, come li abbiamo descritti nella sezione 24.4.

Per dare un'idea, il magnete usato da Purcell a Harvard, costruito con materiale di ricupero della compagnia tranviaria di Boston, era così malamente calibrato che il campo magnetico si trovava in realtà ad un valore più alto di quello al quale si poteva ottenere la inversione dei momenti nucleari sotto irraggiamento con le onde radio $H_{\text{ris}} = 2\pi\nu/\gamma$, ove ν era la frequenza del generatore radio di 30 MHz. Pertanto Purcell e i suoi giovani ricercatori cercavano invano il fenomeno di risonanza. Dopo giorni e giorni di infruttuosi tentativi, disilluso e triste Purcell decise che il fenomeno non era osservabile, diede l'ordine di rinunciare e spengere la corrente che alimentava l'elettromagnete. Mentre il campo magnetico scendeva dopo l'arresto della corrente e volti tristi guardavano l'oscilloscopio dove avevano sperato di vedere i sospirati segnali, il campo attraversò il valore H_{ris} e il primo segnale di risonanza magnetica nucleare fu improvvisamente visto “passare” sullo schermo. Senza questo aiuto della buona sorte forse sarebbero trascorsi ancora molti anni prima di potere utilizzare il magnetismo nucleare.

Invece a quel tempo prese l'avvio una tecnica di ricerca scientifica che ha oggi una vastissima serie di applicazioni nella fisica dei solidi, nella chim-

ica, nella biologia, nella metrologia e nella medicina, la più appariscente e nota delle quali è appunto la “risonanza” per l’ “imaging”. In questo particolare campo di applicazione della risonanza magnetica nucleare la prima mappa bidimensionale di due provette ripiene d’acqua fu ottenuta da **Paul Lauterbur** nel 1973 mentre nel 1976 **Raymond Damadian** presentò la prima immagine NMR di un tumore in un animale.

25.2 Il principio della risonanza magnetica nucleare

Quando un nucleo a spin $I = 1/2$ come il protone, con il suo momento magnetico μ viene posto in un campo magnetico H si generano due e *solo due* livelli di energia magnetica, in conseguenza della *quantizzazione*, che fa sì che sussistano solo due possibili orientazioni del dipolo nucleare nel campo: quella parallela e quella antiparallela. Questo fenomeno è il trasposto delle quantizzazioni e dei livelli discreti che abbiamo già discusso nei riguardi di elettroni in atomi e nei superconduttori. A tali orientazioni corrispondono due livelli di energia magnetica, dal momento che l’energia di un dipolo magnetico in un campo è data dal prodotto $\mu H \cos \theta$, θ essendo l’angolo tra il dipolo e il campo. La separazione tra i due livelli di energia è quindi pari a

$$\Delta E = 2\mu H.$$

Per far mutare la orientazione del momento nucleare da parallela ad antiparallela occorre fornire una energia pari a ΔE . Noi abbiamo già imparato come la radiazione elettromagnetica, di qualunque lunghezza d’onda e quindi di qualunque frequenza angolare ω , è a sua volta quantizzata e viaggia nello spazio in forma di pacchetti di energia $\hbar\omega$ (i *fotoni*). Pertanto, se disponiamo di radiazione elettromagnetica tale che $\hbar\omega = \Delta E$ potremo far assorbire un fotone e portare il momento magnetico nucleare da parallelo ad antiparallelo. In termini di lunghezza d’onda della radiazione necessaria la condizione per ottenere questo capovolgimento è quindi

$$\lambda = \frac{2\pi c}{\omega} = \pi c \hbar / (\mu H).$$

(c è qui la velocità della luce).

Per i nuclei dell’ atomo di idrogeno, sulla base del valore del rapporto giromagnetico, in un campo magnetico di un Tesla, cioè 10.000 Oersted, la

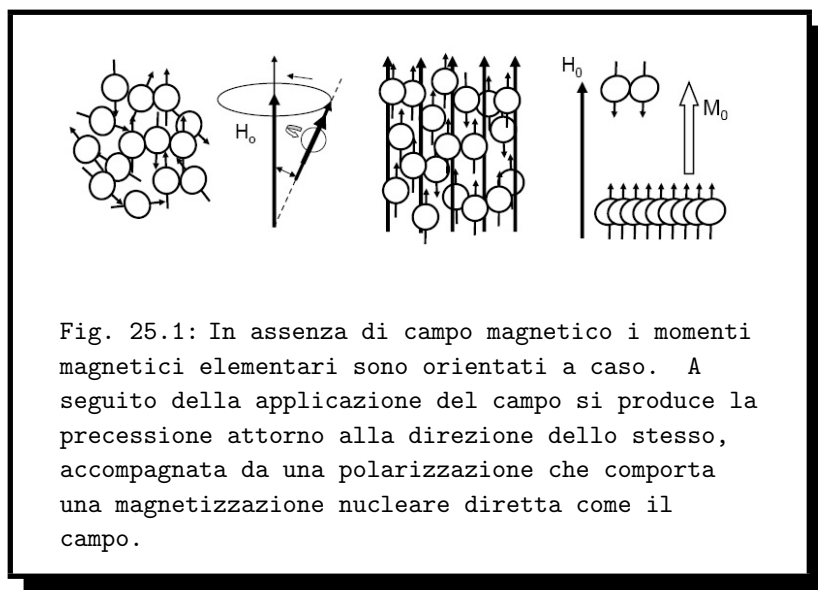
lunghezza d'onda richiesta risulta così pari $3 \cdot 10^{10} \text{ cm/sec} / 42,2 \text{ MHz}$ cioè all'incirca 7 metri, intervallo tipico delle onde delle radio-trasmissioni.

I fotoni di questa particolare lunghezza d'onda possono essere assorbiti a produrre il rovesciamento dei momenti nucleari rispetto al campo, portandoli dal livello di energia magnetica più basso a quello più alto. Tuttavia tali quanti di radiazione possono anche operare in senso opposto facendo capovolgere momenti magnetici nucleari che sono inizialmente antiparalleli al campo e portandoli a paralleli, in maniera simile a come una antenna può assorbire energia da un campo di radiazione oppure, a seconda delle condizioni di fase, cedere ulteriore energia al campo stesso. Infine, se l'irraggiamento viene applicato in forma *pulsata*, per un tempo che dipende dalla potenza del generatore di radiofrequenza ma che deve essere dell'ordine del microsecondo, è possibile anche “fermare” la inversione dei momenti magnetici a metà, cioè portarli ad essere perpendicolari al campo magnetico, uno stato di non-equilibrio quantico!

Il lettore potrà a questo punto chiedersi il perché del termine risonanza con il quale è comunemente chiamato questo fenomeno di particolare assorbimento di energia elettromagnetica. La ragione sta nella descrizione vettoriale del fenomeno stesso. Per prima cosa (per le tipiche frequenze corrispondenti alle onde medie della radio) nell'esperimento che abbiamo appena descritto non si utilizza realmente un'onda elettromagnetica nella forma che abbiamo visto avvenire nel caso di onde luminose o delle microonde del forno per la cottura di alimenti. Si opera in realtà con una bobina percorsa da una corrente prodotta da generatore a radio-frequenza (RF). Il campione di nuclei che desideriamo irraggiare viene posto sull'asse della bobina, che è disposta perpendicolarmente al campo magnetico statico H_0 prodotto dall'elettromagnete o dal solenoide superconduttore. Quando la bobina è percorsa da corrente crea sul suo asse un campo magnetico H_1 , di ampiezza molto più piccola di H_0 (tipicamente qualche unità o qualche decina di Oersted confrontato ai 10.000 o 100.000 Oersted di H_0) e che oscilla alla stessa frequenza del generatore RF. In questo modo si sfrutta totalmente l'intensità magnetica della radiofrequenza, anche se sul campione non perviene esattamente la stessa struttura di un'onda elettromagnetica che invece si riceverebbe se ci ponessimo lontano dalla sorgente che la crea.

Ora dobbiamo fare riferimento alla descrizione vettoriale dei momenti magnetici in campo magnetico. Queste minuscole trottole, oltre a ruotare molto velocemente attorno al loro asse baricentrale, ruotano anche il loro

asse attorno alla direzione del campo H_0 (si veda la Fig.25.1). Questo moto si chiama *precessione* e può essere visualizzato come quello di una trottola nel campo gravitazionale della terra quando il suo asse non è esattamente verticale. La frequenza di questa precessione è uguale a quella che abbiamo stimata per il fotone nella descrizione quantica.



Vediamo ora l'effetto di H_1 (che in luogo di campo oscillante si può equivalentemente vedere come la somma di due campi ruotanti nel piano perpendicolare alla direzione di H_0). Se la radiofrequenza è diversa da ω_{prec} l'effetto di H_1 si può ritenere trascurabile sia per il suo piccolo valore sia perché il suo moto di rotazione e il moto di precessione dei momenti magnetici sono *sfasati* l'uno rispetto all'altro e quindi gli effetti si annullano nel corso del tempo. Quando la frequenza RF di H_1 e quella di precessione ω_{prec} coincidono, allora la precessione dei momenti magnetici e la rotazione di H_1 procedono *in fase*, si ha cioè *risonanza* tra i due moti. Questo avverrà a un preciso valore del campo magnetico, se abbiamo fissato la frequenza RF. Durante il moto risonante di H_1 gli effetti di questo debole campo su $\vec{\mu}$ si sommano nel corso del tempo e tendono a spostare l'asse del nucleo dalla direzione originale. A questo punto scatta il comportamento quantistico: solo due sono le orientazioni possibili per $\vec{L}$ o per $\vec{\mu}$ e quindi il nucleo,

forzato ad abbandonare la sua iniziale quantizzazione lungo il campo deve saltare alla direzione opposta!

Ecco la *risonanza* della descrizione vettoriale e la *meccanica quantistica dei momenti angolari* al lavoro nel produrre il rovesciamento dei momenti magnetici.

Come possiamo “accorgerci” che i nuclei si sono “rovesciati” rispetto al campo statico? Non è difficile con le moderne tecniche di risonanza nucleare ad impulsi: possiamo spegnere rapidissimamente il generatore RF che ha creato H_1 , “accendere” istantaneamente un ricevitore RF che ha la bobina stessa come elemento di ricezione all’ingresso (come la maggior parte delle radio) e registrare i segnali di precessione dei nuclei mentre essi tendono a ritornare verso la situazione di equilibrio, cioè di parallelismo con il campo, dalla quale li avevamo sollecitati a rovesciarsi. Questo segnale, prodotto nella stessa bobina che precedentemente era servita per sollecitare i momenti magnetici e che ora raccoglie le forze elettromotrici indotte dal moto di precessione “libera” degli stessi, si chiama FID, acronimo da *Free Induction Decay*. Il segnale raccolto nel regime del tempo, a partire dall’istante in cui si è spento il trasmettitore RF, può essere processato da un computer e trasformato come spettro nel campo delle frequenze.

Da questa descrizione potete immaginare che uno apparato NMR sarà molto diverso dai più comuni spettrometri nel campo del visibile o dagli strumenti per trasmettere e ricevere gli “spettri di segnali”.

Ora dobbiamo fare cenno a una interessante e utilissima “complicazione”. Nella materia condensata i nuclei non sono affatto isolati. Possono interagire tra di loro e con tutti gli altri gradi di libertà microscopici che ne controllano la distribuzione statistica sui due livelli di energia, al variare della temperatura. L’effetto della temperatura nel causare promozione di particelle su stati di energia più alta l’abbiamo in diverse occasioni già richiamato, per esempio per la evaporazione di molecole da liquidi. Le molteplici interazioni, la loro natura elettrica o magnetica e la loro dinamica microscopica (così come la loro dipendenza dalla temperatura o da campi esterni) sono oggi oggetto delle ricerche in fisica della materia condensata.

Una particolare interazione, che ci ha già interessato a proposito dello SNIF e della risonanza dei nuclei di deuterio dell’alcool etilico nei vini (capitolo 13), è l’interazione dei nuclei con le correnti elettroniche. Tali correnti, diverse da atomo a atomo a seconda del “circondario” di elettroni, creano dei deboli campi magnetici che correggono il campo effettivo agente sui nuclei. Pertanto si producono degli spostamenti (*shift*) nelle frequenze

di risonanza, le quali non cadono più al valore corrispondente ad H_0 ma a frequenze leggermente diverse a seconda delle correzioni dovute alle *correnti elettroniche*. Da spettri NMR ottenuti in alta risoluzione si può risalire, attraverso il cosiddetto *chemical shift*, alla particolare configurazione di atomi o gruppi di atomi nella molecola attorno a un determinato nucleo. Si attua così una specie di fotografia della posizione locale dei diversi nuclei.

Abbiamo fatto cenno come il deuterio potesse sostituire l'idrogeno nelle tre posizioni dei gruppi molecolari CH_3 oppure CH_2 o OH e come nel metodo SNIF i segnali di risonanza corrispondenti ai tre gruppi di atomi si presentassero a diverse frequenze e con diversa intensità, in modo da consentire una corrispondenza con spettri campione, da utilizzare per la classificazione del vino. Anche nell'*imaging* NMR in medicina si realizza una corrispondenza tra i segnali di risonanza e la posizione nello spazio di un volumetto di atomi, per poi ricostruire una immagine che dipende dalla densità locale di nuclei, come discuteremo nella prossima sezione.

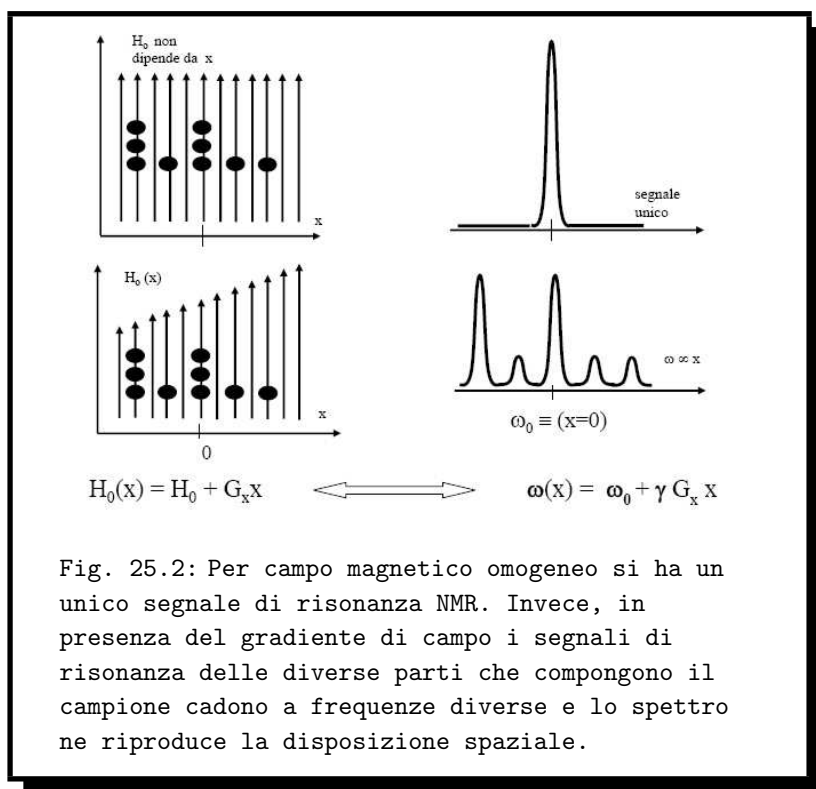
25.3 Come si attua l' "Imaging"

Abbiamo sinora implicitamente assunto che a prescindere dalle deboli correzioni dovute alle correnti elettroniche (che ci sono servite per una prima osservazione sulla dipendenza "dal posto" all'interno di una molecola delle frequenze di risonanza) il campo magnetico fosse *omogeneo*, cioè con lo stesso valore in ogni punto di un campione di molti atomi che vogliamo studiare.

Come suggerito la prima volta da Lauterbur, se si realizza una opportuna disomogeneità nel campo (si introduce cioè il cosiddetto *gradiente di campo magnetico*) e si hanno in tal modo diverse frequenze di risonanza, dagli spettri si può risalire alla posizione spaziale del campione. Il principio della tecnica è molto semplice, anche se poi nella pratica per ottenere la reale visione di un organo interno occorre raffinare molto la metodologia e disporre di potenti computer per pilotare gli impulsi RF, i gradienti di campo magnetico e processare i segnali raccolti dalle bobine.

Immaginiamo di disporre lungo l'asse x delle piccole sferette contenenti acqua, come in Fig.25.2 e mediante delle bobine aggiuntive (rispetto a quella che produce il campo H_0 diretto secondo z) supponiamo di aver creato un gradiente di campo lungo x , con la intensità del campo che aumenta con continuità muovendosi dalla sinistra verso destra. Il segnale

di risonanza magnetica si avrà quindi a valori diversi della frequenza RF e lo spettro sarà costituito da cinque distinti picchi, che avranno altezze proporzionali alle quantità di acqua nei diversi valori della coordinata x e quindi di intensità 3,1,3,1,1. Conoscendo il valore del gradiente di campo possiamo legare lo *spettro* alla *posizione*, per così dire “leggendo” i tre picchi non in termini della loro frequenza ma in termini della coordinata, con una densità di atomi di idrogeno che vale 3,1,3,1,1 alle diverse posizioni.



Immaginiamo ora di introdurre nella regione del campo magnetico una disposizione di sferette più complicata e nel contempo di raccogliere gli spettri NMR in corrispondenza a gradienti di campo applicati lungo direzioni diverse. Per estensione del caso più semplice potete intuire che attraverso operazioni di codifica frequenza di risonanza-posto sia possibile ricostruire una “immagine” della disposizione di sferette. Una tecnica di ricostruzione

delle immagini simile a quella che abbiamo schematizzato si realizza nelle tecniche tomografiche.

Nell' *imaging NMR*, invece dei segnali di FID possiamo operare sui segnali detti delle "eco di spin". Queste eco si presentano spontaneamente al tempo $2t$, se dopo un primo impulso al tempo zero che porta i momenti nucleari perpendicolari ad H_0 ne applichiamo al tempo t un secondo di lunghezza temporale doppia^a. Dalla Fourier trasformazione dei segnali delle eco si può ottenere l'immagine in una varietà di modi, che hanno in comune il principio di fondo ma sono tecnicamente molto diversi.

Una prima serie di impulsi RF "prepara" i nuclei con una certa disposizione angolare rispetto ad H_0 . Quindi nel corso del tempo mentre i nuclei tendono a "rilassare" verso la condizione di equilibrio (che vuole i momenti magnetici in maggioranza allineati parallelamente, come abbiamo detto)^b, vengono applicati sia impulsi di corrente sulle bobine aggiuntive che creano i gradienti di campo magnetico lungo direzioni diverse, sia ulteriormente applicati altri impulsi RF che determinano la formazione delle eco di spin. Una complessa varietà di segnali viene raccolta dal ricevitore, viene digitalizzata e inviata al computer che li acquisisce (generalmente occorrono molte acquisizioni per una stessa serie di impulsi RF e dei gradienti, ecco la ragione per la quale spesso il paziente deve restare alcune decine di minuti all'interno del solenoide che crea H_0).

Terminata una serie di acquisizioni il computer, mediante algoritmi molto veloci, inizia la "elaborazione" dei segnali e mediante la analisi armonica stabilisce la corrispondenza tra la intensità del segnale a una data frequenza e la densità di nuclei risonanti a un dato posto. Al termine, il computer presenta sullo schermo una "visione" bidimensionale di una determinata sezione della parte del corpo del paziente (oppure anche una intera raffigurazione tridimensionale).

^aSi noti che nelle tecniche di risonanza pulsata la durata dell'impulso RF che sollecita i nuclei è dell'ordine del microsecondo mentre tipicamente, per i nuclei dell'idrogeno in condizioni di elevata mobilità come in acqua e nei tessuti molli, la durata del segnale di FID o quello di comparsa della eco di spin, è dell'ordine delle centinaia di millisecondi.

^bIl tempo caratteristico che descrive tale ritorno all'equilibrio è chiamato T_1 mentre l'analogo tempo nel corso del quale le componenti dei momenti nucleari perpendicolari al campo tendono a mediarsi a zero è noto come T_2 . Per non complicare la presentazione degli aspetti di fondo del fenomeno che conduce all'imaging ometteremo i riferimenti al ruolo che possono avere i tempi di rilassamento nel determinare un certo tipo di "visualizzazione" di organi o di parte di essi o nel differenziare segnali provenienti da nuclei in diversa condizione di mobilità (si veda peraltro la Fig. 25.3).

25.4 Stupefacenti “visioni”

Per apprezzare appieno la immagine del vostro interno (per esempio una sezione del vostro cervello che il fisico medico ha realizzato senza aprirvi la testa!) occorre considerare per prima cosa che si tratta realmente di una “visione” e non di “ombra” che il minore o maggiore assorbimento di raggi X produce all’atto di una “radiografia”.

L’occhio umano è un rivelatore diretto per raccogliere quei segnali che gli elettroni, emettendo onde elettromagnetiche alle frequenze dello spettro visibile, hanno inviato verso di noi: tutto ciò che noi vediamo direttamente con i nostri occhi è legato alle accelerazioni di elettroni. Sfortunatamente, o fortunatamente, le onde elettroniche emesse dalle parti interne del nostro corpo non possono giungere direttamente al nostro occhio e noi vediamo solo l’ “esterno” dei corpi.

D’altra parte, come abbiamo mostrato, i nuclei possono essere forzati ad emettere onde elettromagnetiche, a frequenze molto minori di quelle del visibile, quelle delle onde radio, con una frequenza lievemente diversa a seconda di dove essi si trovano. A tali frequenze l’occhio umano è ovviamente insensibile, ecco perché dobbiamo raccogliere le immagini dovute alla emissione di onde da parte dei nuclei con una complessa apparecchiatura elettronica e una procedura computazionale. Tuttavia si tratta di una vera visione dell’ interno di un oggetto o di un corpo umano.

A tale straordinario successo dell’uomo hanno concorso una varietà di avanzamenti nel pensiero scientifico: la comprensione quantistica dei momenti angolari, la teoria della interazione radiazione-materia, l’ elettronica digitale, la matematica e gli algoritmi di trasformazione, le tecniche computazionali.

Molti e rilevanti sono i vantaggi dell’imaging NMR rispetto ad altre tecniche diagnostiche. L’operatore può facilmente scegliere quale sezione del corpo del paziente evidenziare o anche registrare simultaneamente i segnali da più sezioni. In particolare, pilotando opportunamente i gradienti di campo magnetico, può ottenere l’immagine *sagittale* o *coronale*, cioè in piani verticali, di una sezione centrale del nostro cranio o spostata a destra o a sinistra (cose praticamente impossibili con le tecniche di radiografia X). L’operatore può anche “restringere” il campo di osservazione, arrivando cioè a raccogliere i segnali di risonanza, e quindi a visualizzare, solo uno specifico organo o una parte anche piccola di esso, con un grande risoluzione. Da non trascurare infine la possibilità di misurare direttamente, in situ, i

coefficienti di *autodiffusione* (in quanto il segnale della eco di spin dipende dalla rapidità con la quale una certa frazione di nuclei si sposta nel gradiente di campo), dipendenti dalla viscosità locale, nonché i vettori di flusso di liquidi corporei, come i *flussi di sangue*. Per una determinata patologia, giocando su diversi parametri quali ad esempio lunghezza e frequenza degli impulsi e tempi di acquisizione dei segnali, l'operatore può arrivare a una particolare combinazione di sollecitazioni sui momenti nucleari così da ottenere, per esempio, un aumento del contrasto (si veda la illustrazione in Fig.25.3).

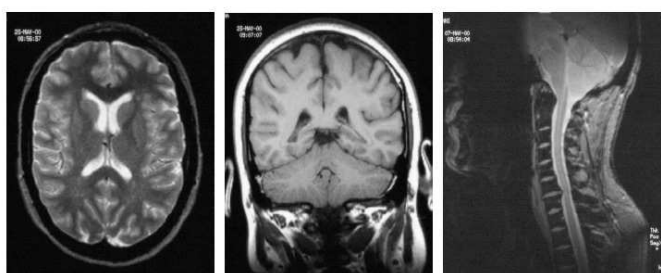


Fig. 25.3: Immagini del cranio e del rachide che con diversi contrasti (ottenuti sfruttando i diversi tempi di rilassamento) evidenziano la materia bianca, la materia grigia, il rachide e il fluido cerebro-spinale, con ottimo dettaglio anatomico.

In forma molto concisa, si può dire che per ogni elemento di immagine (*pixel*) corrispondente a un piccolo volumetto^c, si possono ricavare tutte le informazioni utili, comprese in alcuni casi le mappe di concentrazione di particolari specie chimiche.

Abbiamo detto che per questioni legate alla sensibilità, cioè al rapporto

^cLa risoluzione che si vuole ottenere dipende da molte variabili e un suo incremento comporta la necessità di rinunciare o di ridurre qualche altro parametro utile. Con bobine RF dedicate, si può arrivare a vedere volumetti di 2 micron per 2 micron, su uno spessore di 200 micron.

segnale-rumore, occorre in genere raccogliere e sommare molti segnali di FID o di eco per potere successivamente ottenere una adeguata immagine. Quindi i tempi necessari sono in genere piuttosto lunghi, dell'ordine delle decine di minuti.

Nel 1977 il fisico inglese **Peter Mansfield**^d ha ideato una sequenza di gradienti di campo che non è particolarmente vantaggiosa per la qualità delle immagini, ma è straordinariamente veloce: fornisce una immagine da un solo segnale di FID (in pratica in circa 50 millisecondi). Con tale tecnica (detta dell'eco planare) è oggi possibile vedere realmente il cuore mentre esso pulsa, in una sorta di film in cui si susseguono sullo schermo contrazioni ed espansioni atriali e ventricolari.^e.

Avreste mai immaginato che la fisica potesse produrre una simile stupefacente “visione” ?

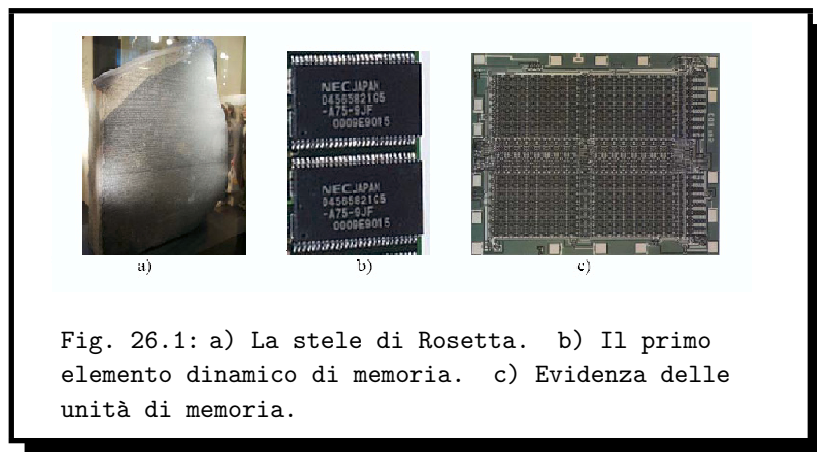
^dMansfield, assieme a Lauterbur, è stato insignito nel 2003 del premio Nobel per la medicina

^eGli autori ringraziano il Dr. Fabio Tedoldi per le informazioni fornite sulle moderne tecniche di imaging NMR.

Capitolo 26

Nanoscienza e futuro del computer.

Nella Fig. 26.1(a) potete vedere la famosa Stele di Rosetta, una pietra in granito scuro di $114\text{ cm} \times 72\text{ cm}$, realizzata in onore del faraone Tolomeo V Epifanio (196 a.C.) in occasione del primo anniversario della sua incoronazione. Nel 1822 questo storico monumento permise a **Jean-François Champollion** di spiegare il significato dei geroglifici Egizi ^a.



Nel 1970 gli ingegneri della Nippon Electric Company (NEC) realizzarono un elemento di memoria a carattere dinamico (Fig. 26.1(b)) cos-

^aLa Stele di Rosetta riporta un'iscrizione con tre differenti forme di scrittura: Geroglifico, Egiziano demotico e Greco antico. Poiché il Greco antico era conosciuto, la stele di Rosetta offrì una chiave decisiva per poter procedere alla comprensione dei geroglifici.

tituito da 1.024 unità di memoria (i piccoli rettangoli della figura 26.1c) disposte in quattro griglie di 32 colonne e di un ugual numero di righe. Le dimensioni di questo primo elemento erano di $0.28\text{ cm} \times 0.89\text{ cm}$ e questo “*chip*” poteva contenere tutto il testo della Stele di Rosetta.

26.1 Pietre miliari nell’era dei computer

L’inizio di una nuova era nel campo della conservazione, trasmissione ed elaborazione della informazione si può far risalire al 1936, con la pubblicazione da parte di **Alan Turing** dell’articolo “*On Computable Numbers with an Application to the Entscheidungsproblem*”. Tale lavoro conteneva i principali concetti di base per la struttura logica di un computer universale. Dieci anni più tardi apparve la prima generazione di computer, che operava utilizzando le valvole elettroniche (tubi termoionici del tipo di quelli usati negli apparecchi radio o televisivi di lontana memoria). Quelli si possono considerare tempi preistorici dell’era dei computer. La maggior parte di quelle prime macchine erano in realtà destinate alla verifica di modalità teoriche di funzionamento. Le dimensioni e il peso di quei “dinosauri” di computer erano ragguardevoli, spesso richiedendo interi edifici per ospitarli.

Poco dopo si realizzò un considerevole progresso nel campo dell’elettronica. Nel 1947 **Walter Brattain**, **William Shockley** e **John Bardeen** (premi Nobel per la fisica nel 1956) inventarono il primo transistor a semiconduttore (realizzato con punte di contatto e che aveva dimensioni dell’ordine del centimetro). I transistor sono stati progressivamente impiegati in ambito elettronico come amplificatori di segnali elettrici o come interruttori elettronici comandati da segnali elettrici e nel seguito hanno sostituito praticamente quasi del tutto i tubi termoionici.

Nel 1958 **Jack Kilby** realizzò il primo “*chip*” a semiconduttore costituito da coppie di transistor. Poco dopo comparvero i primi micro-circuiti integrati che compattavano decine o centinaia di transistor su uno stesso cristallo semiconduttore come substrato.

Queste scoperte segnarono l’inizio di quello che si può ritenere il secondo periodo nella era dei computer (1955-1964). All’incirca nello stesso periodo si iniziò ad utilizzare come memoria di raccolta dati dei dispositivi basati su materiali magnetici, lontani progenitori dei moderni dischi rigidi.

Parallelamente allo sviluppo del materiale *hardware* fece balzi in avanti la teoria riguardante la architettura e la struttura logica dei computer. Il

geniale matematico e fisico **Johan von Neuman** nel Giugno del 1954 nel lavoro “*The first draft of a report on the EDVAC*”^b formulò una meravigliosa descrizione non solo del computer a elettronica digitale ma altresì delle sue proprietà logiche. Il computer divenne da allora un oggetto di interesse scientifico in sé stesso. Questa è la ragione per cui anche oggi spesso gli scienziati chiamano il computer “la macchina di von Neuman”.

Negli anni 1965-1974 furono prodotti i computer di terza generazione basati sui circuiti integrati. Memorie di più alta capacità basate su semiconduttori sostituirono quelle basate sui materiali magnetici, del tipo di quelle tuttora in uso nei personal computer (PC) come *random access memory* (RAM).

Rilevanti passi in avanti della cibernetica teorica, della fisica e della tecnologia determinarono un veloce sviluppo delle tecniche computeristiche nell'ultimo terzo del secolo ventesimo. Divenne parimenti possibile ridurre drasticamente le dimensioni e i prezzi dei computer, sino a farli divenire di uso comune, a partire all'incirca dagli anni 80.

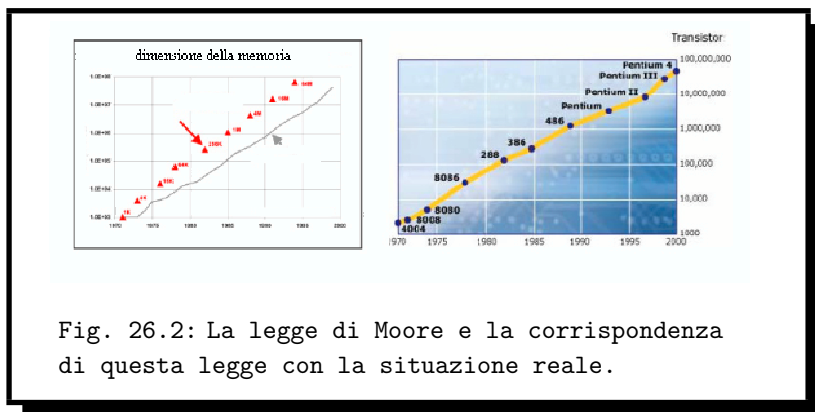
Incredibili progressi nella tecnologia hanno permesso di ottenere in un unico dispositivo da migliaia a decine di milioni di transistor, in una superficie di 3 cm^2 , la stessa area utilizzata inizialmente da Kilby. Gli attuali dispositivi di memoria possono oggi immagazzinare tutte le immagini e i testi della American Congress Library (150 milioni di libri e documenti) in un computer da scrivania.

La legge euristica di **Moore** (in una delle sue varie formulazioni) fissa la crescita esponenziale negli anni della capacità di memoria e del numero dei componenti sul circuito integrato di silicio più economico. Nella Fig. 26.2 si può vedere la corrispondenza di questa legge con la reale crescita, nell'arco di 30 anni.

Una specie di ossessione dei produttori di computer è stato l'incrementare il numero di transistor per unità di superficie e la frequenza dell'orologio di controllo che sincronizza le varie operazioni delle unità (il *clock*) e quindi la velocità di operazione.

Nel 2004 il processore Prescott della Intel, progettato sulla base delle regole della tecnologia nota come dei 90 nm (che in pratica significa disporre in un millimetro quadrato diversi milioni di transistor) conteneva 150 milioni di transistor, operando con il clock alla frequenza di 3.4 GHz .

^bElectronic Discrete Variable Automatic Computer è stato il primo elaboratore dotato di programmi memorizzati.



All'inizio del 2007 i produttori hanno iniziato a muoversi verso il traguardo dei 65 nm e sono in vista le tecnologie per quello dei 45 nm .

Tuttavia, questo modo diretto di incrementare la potenza di calcolo è quasi al capolinea. Una evidente restrizione a ulteriori aumenti con le tecniche convenzionali è posta dalla velocità di propagazione dei segnali elettromagnetici nei circuiti, che ovviamente non può superare la velocità della luce, quale che sia la tecnica di realizzazione del processore. Benché la velocità della luce, $c = 300.000\text{ Km/s}$, appaia enormemente grande, si consideri che la frequenza di 3.4 GHz comporta che l'intervallo tra due operazioni è dell'ordine di $\delta t = 1/\nu \approx 3 \cdot 10^{-10}\text{ sec}$. Ciò significa che la separazione spaziale tra due unità coinvolte in due successive operazioni non può essere maggiore di $L = c\delta t \approx 10\text{ cm}$! Quindi se si scegliesse la strada di incrementare la frequenza dell'orologio di controllo occorrerebbe aumentare il numero di unità di transistor, vale a dire incrementare ulteriormente la miniaturizzazione, lungo la direttrice che è stata percorsa sinora.

Tuttavia ulteriore miniaturizzazione è conflittuale con le proprietà fisiche degli elementi di base dei moderni processori. L'aspetto negativo della miniaturizzazione è il surriscaldamento legato all'aumento della dissipazione per effetto Joule quando una corrente, anche se di debole intensità, percorre centinaia di milioni di transistor. Un'altra restrizione alla miniaturizzazione è legata all'aumento del campo elettrico che si ha riducendo lo spessore dello strato di ossido isolante tra gli elettrodi del transistor: si finirebbe per arrivare alla rottura dielettrica (processo di scarica). Una terza limitazione è da associarsi alle scale caratteristiche che descrivono le fluttuazioni nella

distribuzione delle impurezze, che possono divenire dello stesso ordine delle dimensioni del dispositivo.

Di conseguenza le CPU (*central processing unit*) di minori dimensioni devono usare alimentazione più bassa in quanto la corrente elettrica tende a traforare il materiale dielettrico. Ma riducendo la tensione si aumenta la probabilità di un fenomeno denominato “spargimento”, dove la corrente diventa rumorosa e difficile da controllare. Pertanto una riduzione della tensione non è cosa che i progettisti perseguono, mentre al contrario diventerebbe cosa della quale preoccuparsi nel caso si scendesse a dimensioni ancora minori. In conclusione l’idea di aumentare a forza bruta la frequenza del *clock* si può considerare sostanzialmente fallace.

26.2 XXI secolo: alla ricerca di un nuovo paradigma

Poiché appariva sempre più difficile l’aumento della frequenza del *clock* nei tradizionali processori, nel tentativo di aumentare le prestazioni dei computer i due principali realizzatori di *chip*, Intel e AMD, nel corso del 2005 sono stati praticamente costretti a iniziare il passaggio al nuovo tipo di architettura, quella detta del *dual core*, che significa che due nuclei di elaborazione vengono insediati sullo stesso *chip*. Nel corso del 2007 è prevista la produzione dei primi *chip multi core*, a 4 *core*. Le CPU *dual core* e *multi core* uniscono 2 o più processori indipendenti, le rispettive *cache* (deposito temporaneo) e i *cache controller* in un singolo blocco. Questo tipo di architettura consente di aumentare la potenza di calcolo senza aumentare la frequenza di lavoro, a tutto vantaggio del calore dissipato.

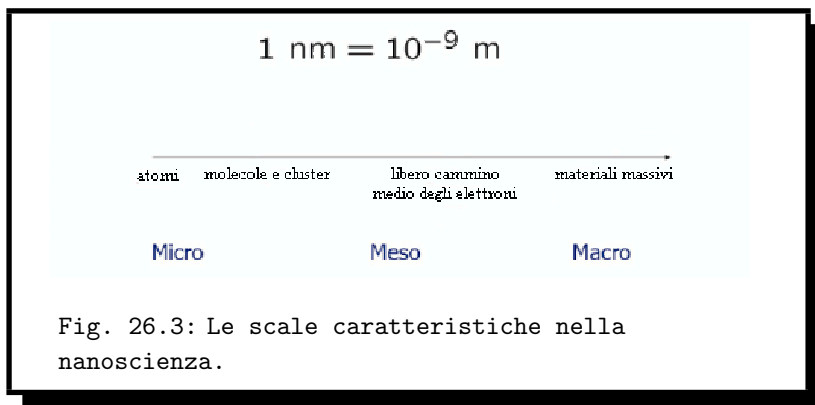
Tuttavia la tecnologia *multicore* non è una panacea, in quanto può favorire un aumento della velocità di calcolo solo per alcune specifiche applicazioni. In primo luogo per fruire dei vantaggi di questi nuovi processori è necessario sviluppare un software completamente nuovo. In un sistema *multicore* si possono comunque produrre effetti collettivi di disturbo, che comportano un calo nella affidabilità di tali architetture. Benché lo sviluppo della tecnologia *multicore* sia certamente una soluzione perseguita oggi, è comunque certo che il potenziamento delle capacità di calcolo dei computer dovrà trovare vie innovative. È necessario prendere atto che sono ormai stati praticamente raggiunti i limiti del mondo macroscopico, regime nel quale sono sostanzialmente le leggi della fisica classica che nei circuiti del processore controllano il comportamento elettrico di trasferimento dei

segnali mediante gli elettroni.

Questa è essenzialmente la ragione per la quale nell'ultima decade gli scienziati hanno tentato di elaborare nuove vie nello sviluppo dei dispositivi di calcolo, basate su logiche diverse da quelle di **Turing** e di **von Neuman**, con elementi base per il processo computazionale diversi dai transistor. Gli elementi di base relativi a questi approcci innovativi comportano lo sviluppo di nuovi algoritmi (il computer quantistico) e di dispositivi a dimensioni nanometriche. Ora tenteremo di far sì che il lettore abbia un'idea di che cosa si tratta.

26.3 Dove sono i confini tra il mondo microscopico e quello macroscopico?

Come si è detto poco sopra, la prospettiva più vicina nella tecnologia dei computer è quella associata alla litografia dei 45 *nm*. Come si confronta questa scala di distanza alle dimensioni caratteristiche del mondo dei quanti? Le scale caratteristiche della nanoscienza sono riportate nella Fig. 26.3.



La prima scala di lunghezza relativamente microscopica, ma per la quale ancora si può fare ricorso alla fisica classica o quasi e che dobbiamo tenere presente nel prefigurare la riduzione delle dimensioni delle unità elementari in un processore, è il libero cammino medio degli elettroni in uno strato metallico. Ricordiamoci che quando si deriva la legge di Ohm si assume che il moto elettronico abbia carattere diffusivo. Questo significa che

nell'ordinario trasporto di carica elettronica (il flusso di corrente) giocano un ruolo determinante i processi di diffusione (*scattering*) degli elettroni da parte di impurezze o di difetti naturalmente presenti nel metallo, il quale non può essere interamente costituito dalla ideale disposizione di atomi su un reticolo regolare per dimensioni relativamente grandi. Poiché la concentrazione di impurezze o di difetti che comportano il moto diffusivo è relativamente elevata, il trasporto è costituito da una specie di media sulle posizioni dei difetti stessi. Pertanto i fili metallici che collegano l'un l'altro i transistor possono essere considerati dei resistori classici, che peraltro obbediscono alla legge di Ohm solo sino a quando le loro dimensioni L rimangono molto maggiori del libero cammino medio l_e : $L \gg l_e$, (nei film metallici prodotti con lo *sputtering* l_e è dell'ordine di diversi nanometri). Nel caso che si realizzi la condizione opposta, il moto degli elettroni avrà carattere balistico, con gli elettroni stessi che in luogo di muoversi a velocità (media) costante si muoveranno di moto accelerato. Le proprietà di questo filo metallico si discosteranno marcatamente dal comportamento ohmico.

Una altra scala di lunghezza che in pratica caratterizza i limiti superiori del mondo quantico è la lunghezza d'onda di **de Broglie** dell'elettrone: $\lambda_F = h/p_e$. Abbiamo già visto che caratteristiche del mondo quantico sono la quantizzazione e l'interferenza. Pertanto quando la minore delle dimensioni di un resistore diventasse confrontabile con la lunghezza d'onda di de Broglie ($L_z \sim \lambda_F$) l'elettrone si dovrà comportare secondo le leggi della meccanica quantistica. Si può dire che il moto dell'elettrone nella direzione corrispondente è quantizzato e l'intero resistore diviene un pozzo quantico. Il valore di λ_F dipende marcatamente dalla concentrazione di elettroni nell'elettrodo e nei metalli ordinari risulta dell'ordine delle distanze di scala atomica. Peraltro in un semiconduttore λ_F può essere molto più grande e il confinamento quantistico diviene cruciale. In particolare, il confinamento quantico in una direzione permette la realizzazione di un gas bidimensionale di elettroni, che è uno dei blocchi costitutivi dei moderni dispositivi elettronici.

Una terza lunghezza di scala che caratterizza le particelle nel mondo dei quanti è la lunghezza di coerenza di fase l_ϕ . Il processo di *scattering* di un elettrone contro una impurezza cambia la direzione del momento p_e dell'elettrone stesso ma non ne cambia il modulo e pertanto non ne cambia la energia. Il processo è detto di scattering elastico. Ora, nella meccanica quantistica esiste una importante caratteristica - la fase della funzione

d'onda degli elettroni- strettamente legata alla energia. Di conseguenza, a seguito di un anche elevato numero di processi di *scattering* l'elettrone continua a conservare la propria fase e può quindi prendere parte a processi di interferenza. In alcuni processi, invece, l'elettrone può cambiare la direzione del suo momento magnetico elementare o urtare, anziché contro una impurezza, contro uno ione vibrante. Come conseguenza può cambiare la propria energia e quindi anche la fase quantistica. La distanza caratteristica percorrendo la quale l'elettrone, a seguito di processi casuali di *scattering* vede la propria fase variare di 2π (il che corrisponde a una totale perdita di memoria) è detta appunto lunghezza di coerenza di fase l_ϕ . Tale grandezza determina una nuova scala quantica: quando la dimensione spaziale dell'elemento del processore diviene minore della lunghezza di coerenza di fase ($L < l_\phi$) possono prodursi fenomeni di interferenza quantistica.

Infine, ricordiamoci che nella interpretazione classica il valore della carica elettronica è considerato come infinitamente piccolo e questo consente di ritenere pressoché continuo il flusso di corrente che fluisce, per esempio, nel processo di carica o di scarica di un condensatore. Se la capacità elettrica di un elemento si riduce a valori così piccoli che la energia elettrostatica di singolo elettrone $e^2/2C$ (con C capacità del dispositivo) diventa confrontabile con altre energie caratteristiche del sistema, per esempio con la energia elettronica nel potenziale di gate eV_g o con la energia termica $k_B T$, possono aver luogo i fenomeni quantistici dei quali abbiamo parlato a proposito del tunneling di elettroni (Capitolo 24).

Dispositivo convenzionale	Dispositivo mesoscopico
diffusivo, $L \gg l_e$	balistico, $L \lesssim l_e$
incoerente, $L \gg l_\phi$	coerenza di fase, $L \lesssim l_\phi$
nessuna quantizzazione, $L \gg \lambda_F$	quantizzazione da confinamento, $L \lesssim \lambda_F$
nessun effetto di carica elettronica finita, $e^2/2C \ll k_B T$	effetti di carica da singolo elettrone, $e^2/2C \gtrsim k_B T$

Le interrelazioni tra i diversi effetti e le diverse lunghezze di scala alle quali abbiamo fatto cenno, comportano una articolata varietà nei meccanismi di trasporto della carica in sistemi nanoscopici. Le specifiche proprietà dei diversi regimi possono essere sfruttate per una varietà di applicazioni. In questa sede non è possibile scendere a una dettagliata discussione delle

varie proprietà fisiche dei possibili “candidati” ad essere gli elementi dei nuovi computer e ci limiteremo ora a discutere solo alcune di tali proprietà.

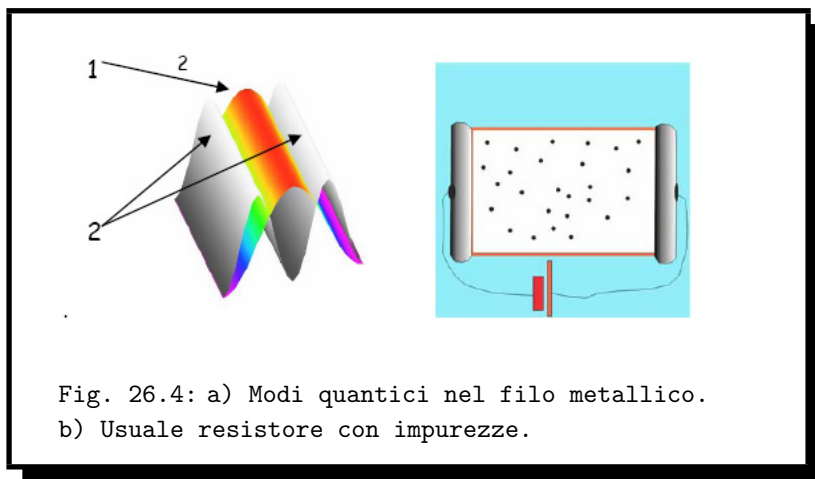
26.4 Fili quantici e punte di contatto quantiche

Consideriamo come avviene il flusso di carica attraverso il cosiddetto “filo quantico”, vale a dire un sottile canale che ospita portatori di carica elementari ma che a differenza dei fili metallici convenzionali si può considerare privo di impurezze o di altri difetti del reticolo cristallino.

Come si è già detto, quando la dimensione trasversale w di un conduttore diventa confrontabile con la lunghezza d'onda elettronica λ_F l'elettrone si trova confinato all'interno di una specie di pozzo quantico e il suo moto nella direzione corrispondente è quantizzato. Ciò comporta che la parte di energia legata al moto trasversale possa assumere solo valori discreti E_n , mentre peraltro il moto lungo l'asse del filo rimane sostanzialmente quello libero, privo di vincoli quantici. D'altra parte la energia totale dell'elettrone rimane costante e quindi al crescere della energia E_n deve corrispondere la diminuzione della energia associata al moto longitudinale. Energia minore significa maggiore lunghezza d'onda corrispondente. Quindi ad ogni livello di energia E_n corrisponde una specifica onda piana con caratteristica lunghezza d'onda λ_n . Pertanto, mentre si propaga lungo un filo quantico l'elettrone ricorda piuttosto un insieme di modi ondulatori in una guida d'onda (si veda la Fig. 26.4 a) piuttosto che una particella che diffonda in un mezzo convenzionale con i relativi processi di *scattering* (Fig. 26.4 b) ^c.

Ciascun modo contribuisce al processo di trasferimento della carica e la conducibilità complessiva è data dalla somma di tutti i contributi. Il contributo al trasferimento di carica da un singolo canale quantico può essere stimato, come abbiamo in altre occasioni già fatto, ricorrendo a consider-

^c Al Capitolo 2 è stata discussa la propagazione del suono nella guida d'onda del canale sonoro subacqueo, senza che si tenesse in conto alcun processo di quantizzazione. Per la verità anche nel mondo macroscopico fenomeni di quantizzazione sussistono. Quando la lunghezza d'onda del suono è confrontabile con la profondità, la guida d'onda diviene “confinante” per tali onde e si possono propagare solo definiti modi sonori. A quel punto la guida d'onda inizia a “deteriorarsi”, in maniera simile a quanto accade nel microfono ad acqua di Bell (capitolo 7). La lunghezza d'onda caratteristica perché ciò avvenga deve essere dell'ordine del kilometro e quindi la frequenza del suono dell'ordine di $\nu = c/\lambda \sim 1\text{Hz}$. Se la frequenza aumenta il numero di nodi nelle onde stazionarie trasverse cresce e si può ritornare alla descrizione del continuo, come si è fatto appunto al capitolo 2.



azioni di *scaling*, cioè ad analisi dimensionali. Nel caso ideale esso non può dipendere dalle proprietà del filo e deve essere la diretta combinazione di costanti universali. La sola combinazione compatibile con la dimensione della resistenza elettrica è h/e^2 , grandezza che non potrebbe comparire nella teoria classica. Si dimostra che in effetti il valore e^2/h determina la conduttanza massima che si può associare a un singolo modo che si propaga lungo il cavo quantico.

Nella Fig.26.5 si può vedere la realizzazione sperimentale di una punta di contatto a carattere quantico. Si vede anche che la conduttanza alla temperatura $T = 1.7\text{ K}$ invece di essere costante (in accordo al comportamento classico) è una funzione che varia con discretezza al variare della tensione del *gate* (la tensione applicata esternamente che controlla la larghezza effettiva del canale di conduzione).

La variazione per ciascuno dei gradini corrisponde al valore $2e^2/h$, in accordo alla nostra stima speculativa basata sulla analisi dimensionale.

Il carattere quantico della conduttanza scompare usualmente per $L > 2\text{ }\mu\text{m}$, oppure viene smussato a seguito di un innalzamento della temperatura (si veda la curva continua, che corrisponde alla temperatura di 20 K).

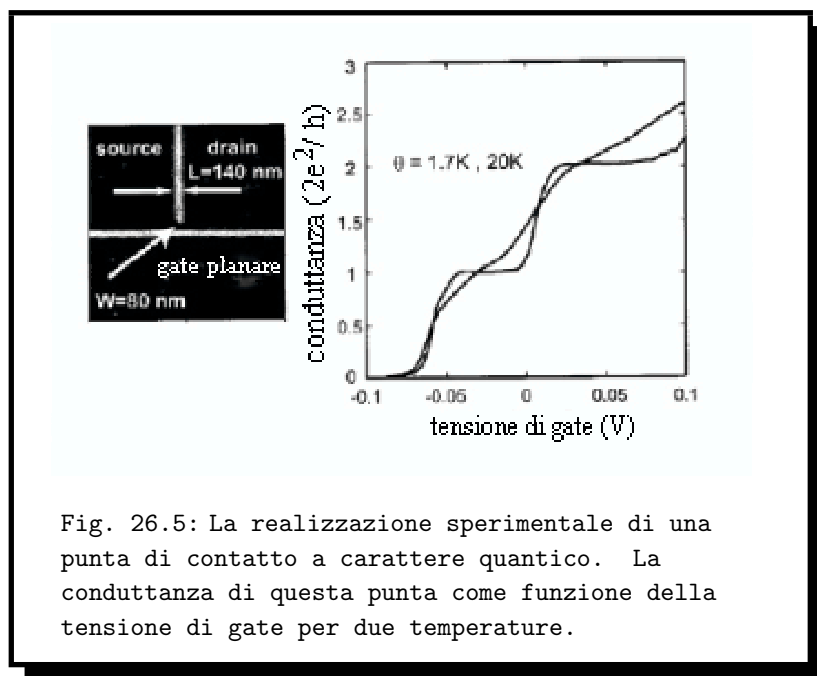
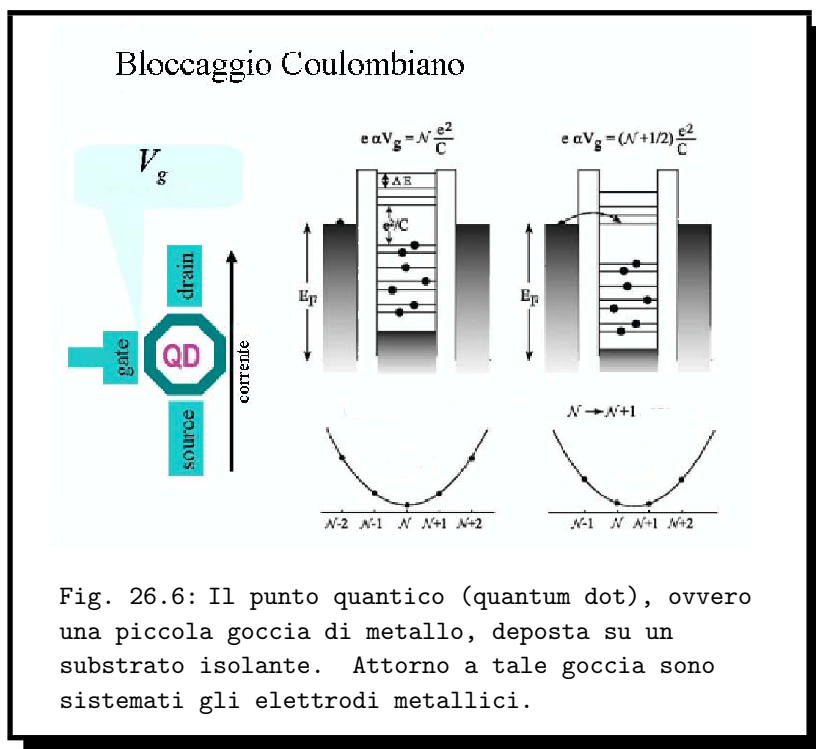


Fig. 26.5: La realizzazione sperimentale di una punta di contatto a carattere quantico. La conduttanza di questa punta come funzione della tensione di gate per due temperature.

26.5 Il "Coulomb blockade" e il transistor a singolo elettrone

Consideriamo ora le proprietà del cosiddetto *punto quantico* (*quantum dot*), ovvero una piccola "goccia" di metallo di dimensione dell'ordine della decina di nanometri, deposta su un substrato isolante. Attorno a tale "goccia" sono sistemati gli elettrodi metallici, vale a dire il *gate* che può modificare il potenziale elettrostatico del *dot*, gli elettrodi "*source*" e "*drain*", che possono fornire o ricevere gli elettroni dal *dot* stesso, come mostrato nella Fig. 26.6.

Assumiamo che il *dot* contenga N elettroni in eccesso, quindi una carica totale $Q = Ne$. La energia elettrostatica del *dot* in assenza di potenziale esterno è $Q^2/2C = N^2e^2/2C$. La capacità elettrica C del *dot* è usualmente molto piccola e questa energia è legata alla repulsione delle cariche elettroniche interne al *dot* stesso. Mediante il gate possiamo applicare un potenziale V_g al *dot* e quindi la energia comprende anche il lavoro del campo elettrico



sulla carica che viene trasferita al *dot*:

$$E(N, V) = -V_g N e + \frac{N^2 e^2}{2C}$$

La derivata di questa funzione parabolica rispetto ad N è nulla in corrispondenza a $N = CV_g/e$, laddove la funzione ha formalmente il suo minimo. Ricordiamoci peraltro che N deve essere un numero intero. Pertanto, al variare della tensione del *gate* si possono produrre diverse situazioni (si veda la fig. 26.7). In un primo caso, quando $V_g = eN/C$, il minimo nella curva corrisponde a uno stato reale con un numero intero di elettroni. Nel caso in cui, invece, $V_g = e(N + 1/2)/C$, il minimo cade formalmente a un numero semi-intero di elettroni, il che non è possibile. I valori per lo stato reale del sistema più vicini a questo risultato formale si hanno con i numeri interi N e $N + 1$. Si osservi l'importante fatto che entrambi i casi

comportano lo stesso valore dell'energia:

$$E(N = CV_g/e - 1/2, V_g) = E(N = CV_g/e + 1/2, V_g) = CV_g^2/2 + e^2/8C.$$

Questo significa che al valore $V_g = eN/C$ della tensione del *gate* lo stato con $(N + 1)$ elettroni sul *dot* è separato dallo stato con N elettroni dalla energia $e^2/2C$: la legge della conservazione della energia “blocca” il trasporto tra la sorgente e il *drain*. Vice versa, quando $V_g = e(N + 1/2)/C$, l'elettrone addizionale può trasferirsi senza costo, in altre parole il *dot* è “aperto”. Osserviamo così che questo dispositivo quantico può operare come un particolare transistor interruttore. A basse temperature, laddove $k_B T \ll e^2/2C$, la corrente che può fluire attraverso questo “transistor a singolo elettrone” è pressoché nulla per ogni tensione del *gate*, con esclusione dei valori $V_g(N) = e(N + 1/2)/C$ quando si produce il salto dell'elettrone in eccesso: la conduttanza mostra quindi degli stretti picchi a definiti valori della tensione del *gate*.

Si può ipotizzare che con l'uso di dispositivi similari a questi transistor a singolo elettrone si possano in futuro sviluppare circuiti logici che operano sulle più piccole correnti possibili, con una minima dissipazione di calore.

26.6 Osservazioni conclusive

In questo capitolo abbiamo presentato in modo in realtà lontano dalla profondità e completezza dell'argomento, alcuni semplici aspetti di base di una nuova e affascinante area di ricerca scientifica, che coinvolge la fisica e la chimica e che è oggi chiamata “nanoscienza”.

Possiamo solo aggiungere che vi sono altri possibili candidati che potrebbero operare come elementi di dispositivi quantici e che sono attualmente oggetto di studio. Essi sono, ad esempio, il gas bidimensionale di elettroni, i nano-tubi di carbonio e altri dispositivi molecolari o sistemi nano-elettromeccanici. Altri possibili dispositivi per computazione quantica sono i sistemi coinvolti nella spintronica (elettronica che manipola individualmente gli elettroni anche in base al loro momento angolare di spin) o nella combinazione di superconduttività e di magnetismo, a livello “nano”.

I sistemi nanoscopici per loro natura implicano la comprensione dei peculiari fenomeni coinvolti nel trasporto quantico, del ruolo delle interazioni elettrone-elettrone e del disordine, del ruolo dei contatti e di tutto l'insieme elettromagnetico che è coinvolto nel dispositivo.

Gli scienziati sono al lavoro, ma bisogna dire che si è oggi ancora piuttosto lontani dalla comprensione profonda di tutti questi aspetti.

