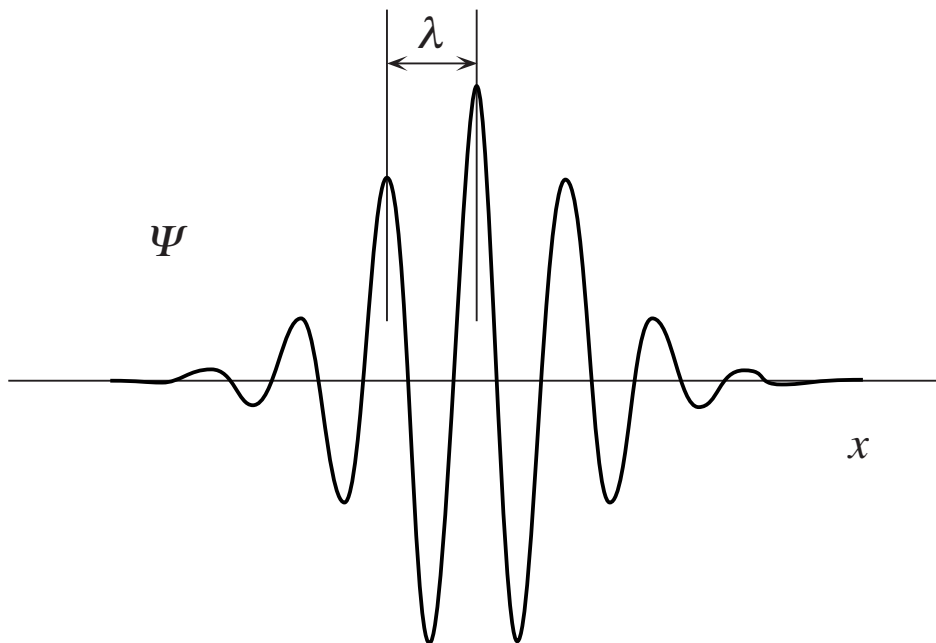


INTRODUCCIÓN A LA MECÁNICA CUÁNTICA

JULIO GRATTON



PRÓLOGO

Las presentes notas se basan en los apuntes que preparé en 2000 para el Curso de Física 4, y hacen pareja con *Termodinámica e Introducción a la Mecánica Estadística*. El estudiante debería leer ambas pues son complementarias. Esta edición ha sido ampliada considerablemente respecto de la versión primitiva. Usamos siempre las unidades gaussianas en el desarrollo de la teoría. Sin embargo, al considerar ejemplos y al dar valores numéricos se emplean a veces unidades prácticas o que pertenecen a otros sistemas. Por lo tanto el lector debe tener el debido cuidado en el empleo de las fórmulas.

Existe una extensa bibliografía que el estudiante puede consultar con provecho. Si bien todos los temas del programa de Física 4 se tratan en estas notas y en *Termodinámica e Introducción a la Mecánica Estadística*, se recomienda a los estudiantes que consulten y lean otros textos, para familiarizarse con la literatura y dado que algunos temas se tratan en ellos con mayores detalles o con enfoques diferentes. Asimismo, es fascinante conocer la historia de la Mecánica Cuántica, para apreciar cómo se fueron desarrollando los conceptos que se introducen en el Curso. El alumno no debe desdeñar obras que se han escrito hace ya muchos años, pues muchas de ellas son excelentes, y a veces mejores que otras más recientes. Entre los innumerables libros que se han escrito sobre la Mecánica Cuántica puedo indicar los siguientes:

(a) de carácter introductorio:

1. R. Eisberg y R. Resnik, *Física Cuántica*, Limusa
2. L. R. Argüello, *Física Moderna*, Answer Just in Time.
3. J. C. Wilmott, *Física Moderna*, Limusa.
4. R. Eisberg, *Fundamentos de Física Moderna*, Limusa.
5. S. Borowitz, *Fundamentals of Quantum Mechanics*, Benjamin.
6. R. H. Dicke y J. P. Wittke, *Introduction to Quantum Mechanics*, Addison-Wesley.
7. R. P. Feynman, R. B. Leighton y M. Sands, *The Feynman Lectures on Physics, Vol 3, Quantum Mechanics*, Addison-Wesley.
8. F. Mandl, *Quantum Mechanics*, Butterworths.
9. D. Park, *Introduction to Quantum Theory*, Mc Graw-Hill.
10. S. Gasiorowicz, *Quantum Physics*, Wiley.

(b) más avanzados:

11. G. Baym, *Lectures in Quantum Mechanics*, Benjamin.
12. D. Bohm, *Quantum Theory*, Prentice-Hall.
13. A. S. Davydov, *Quantum Mechanics*, Addison-Wesley.
14. P. A. M. Dirac, *The Principles of Quantum Mechanics*, Oxford.
15. L. D. Landau y E. M. Lifschitz, *Quantum Mechanics (Nonrelativistic Theory)*, Addison-Wesley.
16. A. Messiah, *Quantum Mechanics*, Wiley.
17. E. Merzbacher, *Quantum Mechanics*, Wiley.
18. L. I. Schiff, *Quantum Mechanics*, Mc. Graw-Hill.

(c) de carácter histórico, excelentes aunque requieren un buen conocimiento de la Mecánica Cuántica para ser apreciados en todo su valor:

19. A. Pais, *"Subtle is the Lord ..."*, Oxford.
20. A. Pais, *Inward bound*, Oxford.

21. A. Pais, *Niels Bohr's times*, Oxford.

También puede resultar provechoso consultar las diferentes voces en la *Encyclopaedia Britannica*, dado que han sido escritas por distinguidos especialistas, así como realizar búsquedas en la Web.

Pido disculpas por las erratas que pueden haber quedado en en estas notas pese a la revisión y agradeceré que se me ponga al corriente de las que fueran detectadas.

Julio Gratton

Buenos Aires, enero de 2003.

INDICE

1. Introducción	1
2. Naturaleza atómica de la materia y la electricidad	3
La hipótesis atómica	3
Evidencias de la naturaleza atómica de la materia	3
Pesos atómicos y la Tabla Periódica de los elementos	5
La Teoría Cinética	5
Tamaño de los átomos	6
La atomicidad de la carga eléctrica	7
Los rayos catódicos	7
El electrón	8
El experimento de Millikan y la cuantificación de la carga	8
3. Estructura atómica	10
Cargas atómicas positivas	10
La dispersión de rayos X y la cantidad de electrones de cada átomo	10
El modelo atómico de Thomson	12
Radioactividad	13
La dispersión de partículas α por los átomos y el fracaso del modelo de Thomson	13
El modelo de Rutherford y el núcleo atómico	15
La constante que está faltando	18
4. Radiación, fotones y la constante de Planck	20
Introducción	20
La teoría de Planck de la radiación de cuerpo negro	20
El postulado de Planck	23
El efecto fotoeléctrico	23
Teoría cuántica de Einstein del efecto fotoeléctrico	25
El efecto Compton	27
La emisión de rayos X	31
Creación y aniquilación de pares	33
La naturaleza dual de la radiación electromagnética	35
5. La Teoría Cuántica Antigua	37
Introducción	37
El espectro atómico	37
Los postulados de Bohr	38
Teoría de Bohr del átomo con un electrón	39
El espectro de líneas de rayos X	43
Refinamientos del modelo de Bohr	44
El principio de correspondencia	48
El experimento de Franck y Hertz	49
Constantes fundamentales, consideraciones dimensionales y escalas	51
Crítica de la Teoría Cuántica Antigua	52

6. Propiedades ondulatorias de la materia	53
El postulado de Broglie	53
Algunas propiedades de las ondas piloto	54
El experimento de Davisson y Germer	56
Interpretación de la regla de cuantificación de Bohr	57
El principio de incerteza	58
Interpretación física de Heisenberg del principio de incerteza	62
La relación de incerteza entre la energía y el tiempo	63
La dispersión de un paquete de ondas	64
El principio de complementaridad	65
7. La teoría de Schrödinger	68
Introducción	68
La ecuación de Schrödinger	68
Interpretación de la función de onda	70
La ecuación de Schrödinger independiente del tiempo	72
Cuantificación de la energía en la teoría de Schrödinger	73
Valores esperados y operadores diferenciales	76
Propiedades matemáticas de operadores lineales en espacios funcionales	79
8. El formalismo de la Mecánica Cuántica	87
Introducción	87
Propiedades de las funciones de onda y de las autofunciones de la energía	88
Normalización en una caja	90
Relaciones de conmutación	90
Autoestados de una variable dinámica	91
Mediciones simultáneas y operadores que conmutan	92
Las relaciones de incerteza de Heisenberg	94
Constantes del movimiento y ecuaciones del movimiento para operadores	95
El límite clásico	97
Representación coordenadas y representación impulsos	98
Transformaciones unitarias	100
9. Soluciones de la ecuación de Schrödinger	103
Introducción	103
La partícula libre	103
El potencial escalón	106
Penetración de una barrera de potencial	110
El oscilador armónico simple	114
10. Fuerzas centrales, momento angular y átomo de hidrógeno	118
Introducción	118
Propiedades del momento angular	118
Magnitud del momento angular	119
Traslaciones y rotaciones infinitesimales y sus generadores	120
Fuerzas centrales y conservación del momento angular	121

Energía cinética y momento angular	121
Reducción del problema de fuerzas centrales	122
Dependencia angular de las autofunciones	123
El problema de autovalores para L_z y la cuantificación espacial	125
Teoría elemental del efecto Zeeman	125
Autovalores y autofunciones de L^2	126
La ecuación radial	131
Estados ligados de átomos con un solo electrón	133
11. El Spin	138
El momento angular intrínseco	138
La evidencia espectroscópica	138
La hipótesis de Uhlenbeck y Goudsmit	139
El experimento de Stern y Gerlach	140
El spin como una variable dinámica	142
Los spinores y la teoría del spin en forma matricial	144
Spin y rotaciones	146
Las matrices de Pauli	150
Operadores de momento magnético y momento angular intrínseco	152
12. Átomos con varios electrones, el Principio de Exclusión y la Tabla Periódica	156
Descripción de un átomo con varios electrones	156
El método del campo autoconsistente	159
Propiedades de los elementos	161
El Principio de Exclusión	164
El Principio de Exclusión y la estructura atómica	165
La unión química y otras interacciones entre átomos	170
13. Partículas idénticas	183
La indistinguibilidad y la función de onda de un sistema de varias partículas idénticas	183
Funciones de onda simétricas y antisimétricas	185
Bosones y Fermiones	186
Sistemas de partículas independientes	187
El Principio de Exclusión de Pauli	188
Las interacciones de intercambio	189
El átomo de helio	196
14. Segunda cuantificación	200
Segunda cuantificación de sistemas de Bosones	200
Segunda cuantificación de sistemas de Fermiones	213
Ecuaciones de movimiento para operadores de campos de Fermiones y Bosones	219
Conexión de la segunda cuantificación con la Teoría Cuántica de Campos	221
El método de Hartree-Fock	223
15. Las Estadísticas Cuánticas	229
El límite clásico	229
La función de partición de un sistema de partículas idénticas sin interacción	230

Las distribuciones de Bose-Einstein y de Fermi-Dirac	232
El gas perfecto de Bosones	233
El gas de fotones y la radiación de cuerpo negro	242
La emisión y absorción de fotones	246
El gas perfecto de Fermiones	257
El modelo de electrón libre de los metales	263
El equilibrio de las enanas blancas	265

1. INTRODUCCIÓN

La Mecánica Cuántica se ocupa del comportamiento de la materia y la radiación en las escalas atómica y subatómica. De esta forma procura describir y explicar las propiedades de las moléculas, los átomos y sus constituyentes: electrones, protones, neutrones, y otras partículas más esotéricas como los quarks y los gluones. Esas propiedades incluyen las interacciones de las partículas entre sí y con la radiación electromagnética.

El comportamiento de la materia y la radiación en la escala atómica presenta aspectos peculiares; de acuerdo con ello las consecuencias de la Mecánica Cuántica no siempre son intuitivas ni fáciles de entender. Sus conceptos chocan con las nociones que nos resultan familiares porque derivan de las observaciones cotidianas de la naturaleza en la escala macroscópica. Sin embargo, no hay razones en virtud de las cuales el comportamiento del mundo atómico y subatómico deba seguir las mismas pautas que los objetos de nuestra experiencia diaria.

El desarrollo de las ideas básicas de la Mecánica Cuántica comenzó a principios del siglo pasado, como consecuencia de una serie de descubrimientos y observaciones que pusieron en evidencia las graves dificultades de la Física Clásica para interpretar las propiedades del átomo y sus partes constituyentes así como las propiedades de la radiación electromagnética y su interacción con la materia. Esos descubrimientos revolucionaron las nociones hasta entonces sustentadas por los físicos y plantearon una asombrosa cantidad de enigmas, cuya solución obligó a realizar un profundo replanteo de los fundamentos y conceptos básicos de la Física.

El estudio de la Mecánica Cuántica es importante por varias razones. En primer lugar porque pone de manifiesto la metodología esencial de la Física. En segundo lugar porque tuvo un éxito formidable ya que permitió dar respuestas válidas a casi todos los problemas en los cuales se la ha aplicado. En tercer lugar porque es la herramienta teórica básica para numerosas disciplinas de gran importancia, como la Química Física, la Física Molecular, Atómica y Nuclear, la Física de la Materia Condensada y la Física de Partículas.

Subsiste, sin embargo, una curiosa paradoja alrededor de la Mecánica Cuántica. A pesar de su notable éxito en todas las cuestiones de interés práctico en las que se la ha aplicado, sus fundamentos contienen aspectos aún no aclarados en forma completamente satisfactoria. En particular, cuestiones relacionadas con el proceso de medición.

Una característica esencial de la Mecánica Cuántica, que la diferencia de la Mecánica Clásica, es que en general es imposible por razones de principio, efectuar una medición sobre un sistema sin perturbarlo. Pero los detalles de la naturaleza de esta perturbación, y el punto exacto en que ella ocurre son asuntos aún oscuros y controvertidos. Por estos motivos la Mecánica Cuántica atrajo algunos de los más brillantes científicos del siglo XX, que han erigido con ella un majestuoso y elegante edificio intelectual.

Este es un curso introductorio. Por lo tanto pondremos el énfasis sobre el desarrollo de los conceptos básicos de la Mecánica Cuántica, sin entrar en los detalles de algunas técnicas de cálculo y formalismos, dado que estos temas se estudian en otros cursos.

En los Capítulos 2 a 4 de estas notas pasaremos revista a estos temas desde una perspectiva histórica, y mostraremos que el comportamiento de las partículas atómicas y de la radiación no se puede describir adecuadamente mediante las nociones clásicas de partícula y onda. Estos conceptos, que derivan de la experiencia a nivel macroscópico, no son adecuados en la escala atómica y por lo tanto deben ser abandonados y reemplazados por una nueva teoría, que es precisamente la Mecánica Cuántica. Por razones de espacio no entraremos en los detalles prácticos y

técnicos de los experimentos que contribuyeron a echar las bases de la Mecánica Cuántica y en cambio sugerimos al lector que recurra a la bibliografía para satisfacer su natural curiosidad. Recomendamos enfáticamente que realice estas lecturas complementarias para adquirir una adecuada cultura científica.

En el Capítulo 5 presentamos a la Teoría Cuántica Antigua, por su interés histórico y porque constituyó, a pesar de sus falencias, el primer intento exitoso en resolver algunos de los problemas y paradojas surgidas del estudio del átomo.

En los Capítulos 6 y 7 introducimos las ideas fundamentales de la Mecánica Cuántica moderna, y en los siguientes Capítulos desarrollamos el formalismo de la teoría y mostramos su aplicación por medio de algunos ejemplos.

Estas notas dejan de lado gran parte de las extensiones y aplicaciones de la Mecánica Cuántica. En particular, no tratamos ni la Mecánica Cuántica Relativística, ni las Teorías de Campos. Tampoco incursionamos en las aplicaciones al núcleo atómico, a las partículas subnucleares y a la materia condensada.

2. NATURALEZA ATÓMICA DE LA MATERIA Y LA ELECTRICIDAD

La hipótesis atómica

El concepto del átomo, en la forma que fuera aceptado por los científicos desde 1600 hasta 1900, se basó en las ideas de filósofos griegos del siglo V AC. Fueron Leucippo de Mileto y su discípulo Demócrito de Abdera quienes originaron la filosofía atómica, introduciendo la noción de un constituyente último de la materia, que denominaron *átomo* (es decir, *indivisible* en la lengua griega). Demócrito creía que los átomos eran uniformes, sólidos, duros, incompresibles e indestructibles y que se movían en número infinito por el espacio vacío; según sus ideas, las diferencias de forma y tamaño de los átomos determinaban las propiedades de la materia. Estas especulaciones fueron luego continuadas por Epicuro de Samos.

Si bien la teoría atómica griega es significativa del punto de vista histórico y filosófico, carece de valor científico, pues no se funda en observaciones de la naturaleza, ni en mediciones, pruebas y experimentos. Para los griegos, la ciencia constituía tan sólo *un* aspecto de su sistema filosófico, mediante el cual buscaban una teoría general que explicara el Universo. Con este fin ellos usaban casi exclusivamente la matemática y el razonamiento, cuando hablaban de la Física. Fue así que Platón y Aristóteles atacaron la teoría atómica sobre bases filosóficas y no científicas. En efecto, mientras Demócrito creía que la materia no se podía mover en el espacio sin el vacío, y que la luz consistía del rápido movimiento de partículas a través del vacío, Platón rechazaba la idea que atributos como “bondad” o “belleza” fueran simplemente “manifestaciones mecánicas de átomos materiales”. Del mismo modo, Aristóteles no aceptaba la existencia del vacío, pues no podía concebir que los cuerpos cayeran con igual rapidez en un vacío. El punto de vista Aristotélico prevaleció en la Europa medioeval, y la ciencia de los teólogos Cristianos se basó en la revelación y la razón, motivo por el cual las ideas de Demócrito fueron repudiadas por considerárselas materialistas y ateas.

Evidencias de la naturaleza atómica de la materia

Con el Renacimiento dio comienzo la nueva ciencia experimental, y se pusieron en duda los puntos de vista Aristotélicos hasta entonces dominantes. Tan pronto como Galileo expresó su creencia de la existencia del vacío (en 1638), los científicos comenzaron a estudiar las propiedades del aire y del vacío (parcial), para poner a prueba los méritos relativos de la ortodoxia Aristotélica y de la teoría atómica. Así fue que Robert Boyle en 1658 comenzó sus estudios sistemáticos sobre la elasticidad del aire que lo llevaron a establecer en 1662 la Ley que lleva su nombre¹. Como conclusión de sus experimentos, Boyle escribió que toda materia está constituida por partículas sólidas de una única clase, dispuestas en *moléculas* de modo de dar a los materiales sus diferentes propiedades. Cuarenta años después, en 1704, Isaac Newton, en su libro *Optiks*, expuso su visión del átomo, semejante a las de Demócrito y de Boyle. Fue así como las antiguas especulaciones acerca de una partícula dura e indivisible fueron lentamente reemplazadas por una teoría científica basada en resultados experimentales y en deducciones matemáticas. Pero fueron necesarios más de 2000 años antes que los físicos modernos comprendieran que el átomo es *divisible*, y que no es ni duro, ni sólido, ni inmutable.

¹ Redescubierta en 1672 en forma independiente por el físico francés Edme Mariotte.

En el curso del siglo XIX se acumuló gran parte de la evidencia de que la materia está compuesta por átomos. Recapitulamos brevemente los hitos más relevantes.

En primer lugar debemos citar algunas leyes de la química. Mencionamos en primer término la Ley de las proporciones constantes, descubierta por Joseph Proust en 1794:

Ley de las proporciones constantes:

cuando se unen elementos químicos para formar un determinado compuesto, las proporciones en peso de los elementos que se combinan son siempre las mismas.

Dicha Ley fue extendida en 1808 por John Dalton, quien propuso la

Ley de las proporciones múltiples:

cuando dos elementos se combinan de distintas formas para dar lugar a diferentes compuestos, los pesos de uno de los dos elementos que se combinan con un peso definido del otro, guardan una relación simple entre sí.

La teoría atómica permite explicar estas leyes. Toda cantidad macroscópica de algún elemento químico consta de gran número de *átomos* de dicho elemento. Todos esos átomos tienen el mismo peso (o masa), que es característico del elemento. Cuando dos elementos se combinan para formar un compuesto, los átomos de los elementos se combinan en una proporción simple, para dar lugar a una *molécula* del compuesto.

Por ejemplo, si se forma óxido cúprico a partir de cobre y oxígeno, se encuentra siempre que 63.5 g de cobre se combinan con 16 g de oxígeno. A partir de los mismos elementos también se puede formar óxido cuproso, pero en este caso 63.5 g de cobre siempre se combinan con 8 g de oxígeno. La hipótesis atómica explica estos hechos diciendo que los pesos atómicos del cobre y el oxígeno están en la relación 63.5:16, y que el óxido cúprico es CuO mientras que el óxido cuproso es Cu₂O. Gracias a esta hipótesis tan simple se pudieron explicar cuantitativamente los pesos de combinación que se observaron en química.

Casi simultáneamente, Joseph-Louis Gay Lussac (1808) encontró que en el estado gaseoso, no sólo los pesos sino también los volúmenes que participan en las reacciones químicas siguen leyes sencillas, siempre y cuando los gases se comporten según las leyes de los gases ideales:

Ley de Gay-Lussac:

en cada gas que se forma o se descompone, los volúmenes de los gases componentes y compuestos guardan relaciones simples entre sí.

Si comparamos esta ley con las anteriores se llega a la conclusión que el volumen de un gas está relacionado con el número de partículas del mismo, y como consecuencia de ello Amedeo Avogadro formuló en 1811² la Ley que lleva su nombre:

Ley de Avogadro:

volúmenes iguales de distintos gases, en las mismas condiciones de temperatura y presión, contienen el mismo número de moléculas.

Esta ley implica que, a una misma temperatura y presión, una cantidad de gas cuyo peso es igual al peso molecular³ ocupa siempre el mismo volumen específico sin importar de que gas se trate.

² El trabajo de Avogadro fue ignorado durante casi 50 años, y su Ley fue aceptada por la comunidad científica recién en 1858.

A temperatura y presión normales, o sea 0 °C y 1 Atm, este volumen es de 22.4 litros. El número de moléculas en un mol⁴ se denomina *número de Avogadro*, y su valor es

$$N_0 = 6.023 \times 10^{23} \quad (2.1)$$

El número de Avogadro se puede determinar de distintas maneras; la más precisa se basa en medir las distancias atómicas por difracción de rayos X.

Pesos atómicos y la Tabla Periódica de los elementos

A medida que se descubrieron más y más elementos a lo largo del siglo XIX, los científicos se comenzaron a preguntar qué relación existe entre las propiedades físicas de los elementos y su peso atómico. De esta forma, durante la década de 1860 se propusieron varios esquemas. En 1869, el químico Dmitry Ivanovich Mendeleev introdujo la Tabla Periódica, basada sobre los pesos atómicos determinados a partir de la teoría de Avogadro de las moléculas diatómicas. Encontró que si se ordenaban a los elementos según su peso atómico, se ponía en evidencia una característica *periodicidad* de sus propiedades. Los elementos que tienen propiedades químicas semejantes, o bien tienen pesos atómicos aproximadamente iguales (como ocurre con el grupo Pt, Ir y Os), o bien tienen pesos atómicos que aumentan de manera uniforme (como K, Rb y Cs). Dejando de lado el Hidrógeno pues es anómalo, Mendeleev ordenó⁵ los 63 elementos entonces conocidos en seis grupos, de acuerdo con su valencia. Al observar que las propiedades químicas cambian gradualmente a medida que aumenta el peso atómico, Mendeleev predijo la existencia de nuevos elementos en todos los casos en que había un “hueco” en la secuencia de pesos atómicos de elementos consecutivos dentro del ordenamiento propuesto. Por lo tanto su sistema, además de ser una forma de clasificación, sirvió también de herramienta para la investigación. Al mismo tiempo, la Tabla Periódica dejó planteados interrogantes muy importantes para cualquier futura teoría del átomo: ¿de dónde proviene el patrón de los pesos atómicos? ¿cuál es el origen de la periodicidad de las propiedades químicas de los elementos?

La Teoría Cinética

La hipótesis atómica se fortaleció aún más debido al éxito de la *Teoría Cinética*, la cual trata los gases como compuestos por un número muy grande de moléculas que se desplazan en el vacío con velocidades distribuidas al azar, cuya magnitud promedio se relaciona con la temperatura. De esta forma se pueden calcular las propiedades mecánicas y térmicas de los gases (ecuación de estado, viscosidad, conductividad térmica, etc.) en términos de la masa, el tamaño y la velocidad de las moléculas. El primero en desarrollar esta teoría fue Daniel Bernoulli (1738), quien la empleó para deducir la Ley de Boyle, basado en la idea que la presión se debe al choque de las moléculas del gas con las paredes del recipiente que lo contiene. Sin embargo su trabajo fue rechazado por más de un siglo⁶. La teoría fue vuelta a formular en forma independiente por John

³ Expresado (por ej.) en g.

⁴ Un mol es la cantidad de sustancia cuyo peso es igual al peso molecular expresado en g.

⁵ El número de orden de cada elemento dentro de la Tabla Periódica se denomina hoy número atómico, y como veremos en el Capítulo 3, es igual al número de electrones que posee el átomo.

⁶ Fundamentalmente porque se daba más crédito a las ideas de Newton, según las cuales la presión se originaba debido a una supuesta repulsión entre las moléculas. Incluso los genios se equivocan!

Herapath (1820) y por John James Waterston (1845) quien fue el primero en deducir la equipartición de la energía. Sin embargo, estos trabajos corrieron la misma suerte que el de Bernoulli y fueron ignorados⁷. Recién después de los trabajos de James Prescott Joule (1840), que desacreditaron la Teoría del Calórico al mostrar que el calor es una forma de energía, el camino quedó despejado para la aceptación de la Teoría Cinética. Fue así que Rudolf Clausius desarrolló en 1857 la matemática correspondiente, y luego James Clerk Maxwell y Ludwig Eduard Boltzmann completaron su desarrollo alrededor de 1860.

En este contexto corresponde mencionar un fenómeno que constituye una de las comprobaciones más evidentes de la hipótesis atómica. Si se suspende un objeto diminuto dentro de un gas, también es bombardeado por las moléculas. Como el número de las moléculas es finito, no se establece un equilibrio exacto en cualquier instante y en consecuencia el objeto se mueve en forma aleatoria. Un botánico, Robert Brown, fue el primero (1828) en observar este fenómeno (que en su honor se denomina *movimiento Browniano*) al observar bajo el microscopio una suspensión de granos de polen en agua. Mucho tiempo después, en 1908, Jean Perrin usó el movimiento Browniano para determinar el número de Avogadro, basado en la analogía entre las partículas suspendidas en el líquido y las moléculas en la atmósfera⁸. La teoría correspondiente había sido publicada por Albert Einstein y Marian Ritter von Smoluchowski en 1905. Luego del trabajo de Perrin ya nadie cuestionó la existencia de los átomos.

Tamaño de los átomos

Las primeras estimaciones modernas del tamaño de los átomos fueron realizadas por Joseph Lotschmidt en 1865, y se basaron en los resultados de la Teoría Cinética. No las comentaremos aquí. Será suficiente decir que conociendo el número de Avogadro, podemos calcular el tamaño de los átomos de un sólido (por ejemplo un metal) si suponemos que están ubicados uno junto al otro de modo que los átomos vecinos se tocan entre sí. Si $\mathcal{A}$ es la masa atómica⁹ de la sustancia y ρ su densidad, el radio r de un átomo está dado por

$$r = \frac{1}{2} \left(\frac{\mathcal{A}}{N_0 \rho} \right)^{1/3} \quad (2.2)$$

Si se hace este cálculo para varios elementos, desde el litio ($A = 7$) al plomo ($A = 207$) se encuentra que r varía entre 1.3×10^{-8} cm y 1.55×10^{-8} cm. De acuerdo con estas estimaciones todos los átomos tienen un tamaño del orden de 10^{-8} cm, es decir 1 Å. Otros métodos, como la teoría cinética, dan resultados semejantes. Se debe notar que no hemos definido con precisión lo que significa el “radio de un átomo”, y por lo tanto tenemos que tener cuidado con el uso de este término; en particular no sabemos todavía cómo varía la interacción entre dos átomos como función de la distancia.

⁷ En esos tiempos estaba en boga la Teoría del Calórico, motivo por el cual no se aceptaba la idea que el calor estuviera relacionado con el movimiento de las moléculas.

⁸ La variación de la densidad del aire con la altura depende del balance entre la gravedad (que tiende a hacer descender las moléculas) y la agitación térmica (que tiende a expandir el aire). La relación entre densidad y altura para partículas Brownianas en suspensión obedece a un balance semejante.

⁹ La masa atómica $\mathcal{A}$ es la masa de N_0 átomos. El valor numérico de $\mathcal{A}$, cuando se expresa en múltiplos enteros de la masa del átomo de hidrógeno, se denomina *número de masa* y se indica con A .

La atomicidad de la carga eléctrica

Los experimentos sobre la electrólisis realizados por Michael Faraday a partir de 1832 demostraron que la cantidad de sustancia liberada en un electrodo de una cuba electrolítica por el paso de una carga eléctrica Q es proporcional a la masa equivalente de la sustancia (la masa atómica dividida por la valencia). La constante de proporcionalidad F se denomina *Faraday*, y se tiene:

$$M = \frac{QA}{vF} \quad (2.3)$$

donde M es la masa liberada de la sustancia, Q es la carga transferida y v es la valencia. El valor de un Faraday es

$$F = 96500 \text{ C/mol equivalente} \quad (2.4)$$

El hecho que la masa liberada es estrictamente proporcional a la carga total transferida sugiere que la carga es transportada por los iones mismos. La hipótesis más simple es que cada ion lleva una carga qv , es decir, un múltiplo de una *carga elemental* q . Entonces la carga necesaria para liberar un mol es N_0qv y por lo tanto, puesto que para un mol $M = A$, tendremos que

$$q = F / N_0 \approx 4.8 \times 10^{-10} \text{ u.e.s.} \approx 1.60 \times 10^{-19} \text{ C} \quad (2.5)$$

Si combinamos la hipótesis atómica con los resultados de la electrólisis se concluye que cada ion está asociado con una carga determinada qv , que es un múltiplo de la carga elemental q . Por lo tanto la carga eléctrica tiene también una naturaleza atómica y en el electrolito cada ion lleva un número de “átomos de carga” igual a su valencia.

Los rayos catódicos

A temperatura y presión normales los gases no conducen la electricidad, hasta que la intensidad del campo eléctrico es tal que se produce una chispa. Sin embargo, si se tiene un par de electrodos en un recipiente cerrado y se reduce la presión a menos de 10 mm Hg, al aplicar algunos kV entre los electrodos se observa una descarga brillante, con colores y patrones llamativos. Si se reduce aún más la presión, la región oscura que está delante del cátodo se extiende paulatinamente hasta que a una presión de unos 10^{-3} mm Hg llena todo el recipiente. No obstante, sigue pasando corriente eléctrica. Si se hace un orificio en el ánodo, se observa un resplandor verdoso en la pared del tubo de vidrio detrás del orificio. Los agentes que producen este resplandor viajan en línea recta desde el orificio del ánodo, cosa que se puede verificar por la sombra que produce cualquier objeto que se interponga entre el ánodo y la pared de vidrio. Si se coloca una rueda de paletas en la trayectoria, comienza a girar, lo que muestra que los agentes llevan cantidad de movimiento. Dichos agentes se denominaron *rayos catódicos*¹⁰.

¹⁰ Los rayos catódicos fueron descubiertos por Julius Plücker en 1858 e investigados por William Crookes en 1879, quien encontró que se desvían en un campo magnético, y que la dirección de la desviación sugiere que se trata de partículas de carga negativa. Sin embargo la verdadera naturaleza de los rayos catódicos fue tema de controversia. Una prueba crucial consistió en estudiar el efecto sobre los mismos de un campo eléctrico. En 1892 Heinrich Hertz llevó a cabo un experimento que tuvo resultados negativos. J. J. Thomson consideró que ello se debía a que el vacío no había sido suficientemente bueno en el experimento de Hertz, y decidió repetirlo con un vacío mejor.

El electrón

Joseph John Thomson realizó varios experimentos en 1896-7 sobre un haz fino de rayos catódicos producido colimando los rayos que salen del orificio del ánodo. Comprobó que todos son desviados por igual por un campo eléctrico transversal a su trayectoria, y por el sentido de la desviación dedujo que todos tienen la misma carga negativa que indicamos con $-e$, de modo que se trata de algún tipo de partícula. Estudiando la desviación concluyó que, si estas partículas tienen una masa m , resulta

$$\frac{mv^2}{e} = \text{cte.} \quad (2.6)$$

Aplicando un campo magnético al haz de rayos catódicos observó una desviación, a partir de la cual determinó la relación carga/masa de las partículas. El valor actualizado de esa relación es:

$$\frac{e}{m} = (1.7598 \pm 0.0004) \times 10^8 \text{ C/g} \quad (2.7)$$

Los experimentos de electrólisis permiten también calcular una relación carga/masa. Si consideramos esta relación para el elemento más liviano, o sea el hidrógeno, resulta

$$\frac{Q}{M} = 9.57 \times 10^4 \text{ C/g} \quad (2.8)$$

Si comparamos la relación carga/masa (2.7) con la (2.8) obtenemos

$$\frac{(e/m)}{(Q/M)} = 1.84 \times 10^3 \quad (2.9)$$

Por lo tanto, o las partículas de los rayos catódicos son mucho más livianas que el átomo de hidrógeno, o bien llevan una carga casi dos mil veces mayor a la del ion hidrógeno. Esta última hipótesis parece tan poco lógica que Thomson propuso que tanto las partículas de los rayos catódicos como el ion de hidrógeno llevan cargas de igual valor absoluto y que las partículas de los rayos catódicos, que denominó *electrones*, son mucho más livianas que los átomos.

El experimento de Millikan y la cuantificación de la carga

En realidad hasta ahora sólo podemos afirmar que la carga *promedio* en la electrólisis está cuantificada, pues si bien supusimos que cada átomo lleva una carga qv , en realidad solo es necesario suponer que la carga promedio vale qv . Del mismo modo, los experimentos de Thomson se podrían explicar suponiendo que dentro de los átomos hay algún material especial con carga negativa y con una relación carga/masa sumamente elevada, y que en las descargas eléctricas algunos fragmentos de ese material son emitidos por los átomos del gas o por el cátodo. Todavía no hemos presentado pruebas concluyentes de que ese material especial sólo puede existir en cantidades discretas, a parte las deducciones que se pueden hacer a partir de la electrólisis.

Las pruebas cruciales fueron aportadas por Robert A. Millikan, quien en 1909 estudió la caída en el aire por efecto de la gravedad de minúsculas gotas cargadas eléctricamente. Mediante un par de electrodos podía introducir un campo eléctrico vertical y así estudió el movimiento de las gotas, con y sin el campo eléctrico. No vamos a reproducir aquí las fórmulas (ver J. C. Wilmott,

Física Atómica), pero se puede mostrar que de esta forma se puede calcular la carga de las gotas individuales. Lo que encontró Millikan es que dichas cargas son siempre múltiplos enteros de la carga más pequeña que puede llevar una gota. De este modo quedó demostrado que la carga está cuantificada. El valor del cuanto de carga es

$$e = 4.80273 \times 10^{-10} \text{ u.e.s.} \approx 1.603 \times 10^{-19} \text{ C} \quad (2.10)$$

valor que coincide con el que se deduce de la electrólisis (ec. (2.5)). Combinando este valor con la relación carga/masa (2.7) podemos obtener la masa del electrón:

$$m_e = (9.108 \pm 0.012) \times 10^{-28} \text{ g} \quad (2.11)$$

Recogiendo los resultados que hemos presentado hasta aquí, podemos concluir que a comienzos del siglo XX había quedado demostrada la naturaleza atómica de la materia. Sin embargo, en contradicción con las ideas primitivas acerca del átomo, éste resultaba ser un objeto *compuesto*, y uno de sus componentes, el electrón, había sido identificado.

3. ESTRUCTURA ATÓMICA

Cargas atómicas positivas

De acuerdo con los resultados reseñados en el Capítulo 2 sabemos que la materia está formada por átomos cuyo radio es de unos 10^{-8} cm. Estos átomos están constituidos, al menos en parte, por electrones. Puesto que son eléctricamente neutros es obvio que la carga debida a los electrones que contienen debe estar equilibrada por una carga positiva igual.

Además de los rayos catódicos, en los experimentos con tubos de descarga se observaron partículas con carga positiva que emanan del ánodo, que fueron denominadas *rayos positivos*¹. En 1898 Wilhelm Wien investigó estos rayos y encontró que tienen una relación masa/carga más de 1000 veces mayor que la de los electrones. Puesto que esta relación es comparable con la relación masa/(carga del electrón) de los átomos residuales del tubo de descarga, se sospechó que los rayos positivos son *iones* (átomos cargados positivamente porque les faltan uno o más electrones) provenientes del gas presente en el tubo. En 1913 Thomson refinó el dispositivo de Wien para separar los diferentes iones y medir sus relaciones masa/carga. Así determinó la presencia de iones de varios estados de carga (es decir, átomos que han perdido uno, dos, tres, ..., etc. electrones), que aparecían como diferentes trazas en una placa fotográfica. Al realizar sus experimentos con neón, observó que los haces de iones del mismo estado de carga producían *dos* trazas en vez de una. Los químicos habían atribuido al Ne un número de masa de 20.2, pero las trazas observadas por Thomson sugerían números atómicos de 20.0 (la traza más intensa) y 22.0 (la más débil). Concluyó entonces que el Ne consiste de una mezcla de dos variedades, que denominó *isótopos*²: la más abundante, ^{20}Ne (con número de masa 20.0) y la más escasa ^{22}Ne (de número de masa 22.0). Más tarde se descubrió un tercer isótopo, el ^{21}Ne , que también está presente en cantidades diminutas³. Estos resultados demostraron que la hipótesis de Dalton, que todos los átomos de un dado elemento tienen la misma masa, está en error. La técnica de Thomson fue perfeccionada por Francis W. Aston, quien desarrolló el espectrógrafo de masa en 1919 y lo utilizó para analizar cerca de 50 elementos en los años siguientes, lo que le permitió descubrir que la mayoría tienen isótopos.

En conclusión, no hay ninguna evidencia de que en el átomo exista una partícula positiva equivalente al electrón. Por lo tanto la carga positiva está de alguna manera (no trivial) asociada con la masa del átomo.

La dispersión de rayos X y la cantidad de electrones de cada átomo

Las pruebas que los átomos de una dada especie contienen un número definido de electrones provienen de los experimentos sobre dispersión de rayos X.

¹ A veces también llamados rayos canales. Fueron descubiertos por Eugen Goldstein en 1886.

² El término isótopo proviene del griego y significa “igual lugar”. Se refiere al hecho que las variedades del mismo elemento ocupan igual lugar en la Tabla Periódica, pues tienen (casi exactamente) las mismas propiedades químicas. Ya en 1886 Crookes avanzó la idea que todos los átomos tienen pesos atómicos enteros y que los elementos cuyos números de masa tienen valores no enteros son en realidad mezclas. Asimismo, la existencia de isótopos fue sospechada por Frederick Soddy en 1910, al estudiar los productos del decaimiento radioactivo del torio.

³ Hoy sabemos que de cada 1000 átomos de Ne, 909 son de ^{20}Ne , 88 son ^{22}Ne y 3 son ^{21}Ne .

Los rayos X fueron descubiertos por Wilhelm Conrad Roentgen en 1895, al realizar experimentos de descargas en gases. Cuando la diferencia de potencial entre el cátodo y el ánodo es de algunos kV, al ser bombardeado por los electrones, el ánodo emite una radiación penetrante que se denominó radiación X. En poco tiempo se pudo mostrar que los rayos X son radiación electromagnética de longitud de onda muy corta, debido a que

- Los rayos X se producen cuando electrones energéticos impactan sobre un objeto sólido. En estas circunstancias los electrones sufren una violenta desaceleración, y de acuerdo con la teoría electromagnética un electrón acelerado o desacelerado emite radiación.
- Haga y Wind encontraron en 1899 que los rayos X se difractan al pasar por una rendija muy fina, lo que muestra que son un fenómeno ondulatorio. El tamaño de la figura de difracción indica que la longitud de onda es del orden de 10^{-8} cm.
- En 1906 Charles Glover Barkla mostró que los rayos X se pueden polarizar, lo que indica que son ondas transversales.
- En 1912 Max von Laue desarrollo una técnica para medir la longitud de onda de los rayos X, basada en la difracción por una red cristalina.

Cuando la radiación electromagnética incide sobre una partícula cargada, ésta oscila por efecto del campo eléctrico de la onda, y al ser acelerada emite radiación. La intensidad total de la radiación emitida por una carga acelerada está dada por la fórmula

$$R = \frac{2}{3} \frac{e^2 a^2}{c^3} \quad (3.1)$$

donde e es la carga, a es la aceleración, c la velocidad de la luz y estamos usando unidades Gaussianas. Aquí lo que nos interesa es que la energía emitida es proporcional al cuadrado de la aceleración. En un campo eléctrico E , la aceleración de una carga es eE/m , de modo que la cantidad de energía dispersada por la carga es proporcional a e^4/m^2 . Puesto que la masa de los electrones es mucho menor que la de los demás constituyentes del átomo, es obvio que la dispersión de los rayos X por los átomos se debe esencialmente a los electrones.

Es posible entonces usar la dispersión de rayos X para estimar el número de electrones de un átomo, siempre y cuando la longitud de onda de los rayos X que se emplean sea *menor* que la distancia entre los electrones, pues en este caso las oscilaciones de los diferentes electrones casi nunca estarán en fase (suponemos que las distancias que separan a los electrones no siguen un patrón regular). En este caso las radiaciones emitidas por los diferentes electrones son incoherentes y podemos sumar sus intensidades. La intensidad total de la radiación dispersada es entonces proporcional al número de electrones. Hay un segundo requisito que se debe cumplir, esto es, que la frecuencia de los rayos X sea mucho mayor que cualquiera de las frecuencias naturales de oscilación de los electrones dentro del átomo, de modo que podamos tratar los electrones como si fueran libres. Esto, en la práctica, significa que la longitud de onda de los rayos X debe ser bastante menor que 10^{-8} cm. Con estas hipótesis (omitimos los detalles del cálculo⁴) J. J. Thomson mostró que la intensidad total de la radiación difundida por un electrón bajo la influencia de una onda electromagnética de intensidad I es

⁴ Ver por ejemplo W. Panofsky y M. Phillips, Classical Electricity and Magnetism, Addison-Wesley.

$$R = \sigma_T I \text{ J/s} \quad , \quad \sigma_T = \frac{8\pi}{3} \left(\frac{e^2}{mc^2} \right)^2 = 6.66 \times 10^{-25} \text{ cm}^2 \quad (3.2)$$

donde σ_T se denomina *sección eficaz de Thomson* del electrón. La (3.2) también se puede escribir en términos del *radio clásico del electrón*:

$$r_0 = \frac{e^2}{mc^2} = 2.817938 \times 10^{-13} \text{ cm} \quad (3.3)$$

en la forma

$$\sigma_T = \frac{8\pi}{3} r_0^2 \quad (3.4)$$

Si hay n electrones por unidad de volumen, la intensidad difundida por un trozo de materia de sección unidad y espesor dx será

$$ndx\sigma_T I \quad (3.5)$$

Esta es la energía perdida por el rayo, que por lo tanto sufre una variación de intensidad dada por

$$dI = -n\sigma_T I dx \quad (3.6)$$

La (3.6) implica que la intensidad del haz se atenúa exponencialmente al atravesar un espesor del material, es decir:

$$I = I_0 e^{-n\sigma_T x} \quad (3.7)$$

Luego midiendo la atenuación de los rayos X que atraviesan un determinado espesor de materia se puede determinar n y de allí el número Z de electrones de cada átomo. Los resultados experimentales muestran que Z es *aproximadamente* igual a la mitad del número de masa A .

El modelo atómico de Thomson

Para avanzar más fue necesario realizar otros experimentos, sugeridos por un modelo propuesto por J. J. Thomson para explicar los datos conocidos, y que hacía nuevas predicciones aún no verificadas. De acuerdo con este modelo el átomo es una esfera de carga positiva de unos 10^{-8} cm de radio, con electrones en su interior como las pasas de uva dentro de un pan dulce. Suponiendo que la carga positiva está distribuida uniformemente, cabe esperar que también los electrones estén distribuidos uniformemente, pues así la carga neta en toda esfera con centro en el centro del átomo es nula en promedio y la distribución de cargas es estable. Este modelo está de acuerdo con todas las propiedades del átomo conocidas en su momento. Además, si se perturban las posiciones de los electrones, éstos oscilarán y por lo tanto emitirán radiación de una determinada frecuencia. Luego se explica, al menos cualitativamente, la emisión de luz por los átomos (se puede ver, sin embargo, que muy difícilmente pueda haber un acuerdo cuantitativo con el espectro de la radiación emitida que se observa).

El experimento crucial para probar el modelo fue realizado por Ernest Rutherford y sus colaboradores en 1911 y consistió en el estudio de la dispersión de partículas α por átomos. Veremos

que el modelo de Thomson predice que el número de partículas α dispersadas en ángulos grandes es despreciable. Pero el experimento desmintió esta predicción, y la explicación de las observaciones llevó a Rutherford a proponer el *modelo nuclear*, según el cual el átomo está constituido por un pequeño núcleo central donde está concentrada la carga positiva y casi toda la masa, rodeado por electrones en movimiento, de forma que el conjunto es totalmente neutro.

Radioactividad

Entre 1896 y 1898 Antoine-Henri Becquerel, Pierre Curie y Maria Curie (Maria Sklodowska) descubrieron que algunos elementos pesados como el uranio y el torio emiten espontáneamente radiaciones penetrantes, capaces de velar una placa fotográfica. Haciendo pasar un haz colimado de estas radiaciones a través de un campo magnético, se encontraron tres componentes que fueron denominados radiación α , β y γ . Los rayos γ no son desviados por el campo magnético, lo que indica que no tienen carga eléctrica; en cambio, los rayos α y β se desvían, mostrando que los primeros tienen carga positiva y los segundos, negativa. Estos experimentos se realizaron bajo vacío. Introduciendo aire en el dispositivo, se observó que bastan pocos centímetros de aire para detener la radiación α , pero no las otras dos componentes. Interponiendo láminas de distintos espesores se encontró que pocos mm de un material denso son suficientes para detener la radiación β ; en cambio, la radiación γ sólo disminuye apreciablemente si se interpone un bloque de plomo de varios cm de espesor. Actualmente sabemos que la radiación β está compuesta por electrones de gran energía⁵, y que los rayos γ son radiación electromagnética de longitud de onda extremadamente corta. La naturaleza de las partículas α fue descubierta por Rutherford, quien encontró que se trata de átomos de Helio doblemente ionizados⁶.

Al estudiar la radioactividad del torio, Rutherford y Frederick Soddy descubrieron en 1902 que la radioactividad está asociada con profundos cambios dentro del átomo, que lo transforman en un elemento distinto. Encontraron que el torio produce continuamente una sustancia químicamente diferente, que es intensamente radioactiva. Si el elemento así producido se separa del torio, desaparece con el correr del tiempo, dado que a su vez se *transmuta* en otro elemento. Observando este proceso, Rutherford y Soddy formularon la ley del decaimiento exponencial, que establece que en cada unidad de tiempo, decae una fracción fija del elemento radioactivo.

El descubrimiento de la radioactividad y la transmutación de los elementos obligó a los científicos a modificar radicalmente sus ideas sobre la estructura atómica, pues demostró que *el átomo no es ni indivisible ni inmutable*. En vez de ser simplemente un receptáculo inerte que contiene electrones, se vio que el átomo puede cambiar de forma y emitir cantidades prodigiosas de energía. Además, las radiaciones mismas sirvieron de instrumento para investigar el interior del átomo.

La dispersión de partículas α por los átomos y el fracaso del modelo de Thomson

Si se hace incidir un haz colimado de partículas α sobre una hoja delgada (por ej. una lámina de oro de 10^{-4} cm de espesor) se observa que casi todas la atraviesan con una leve pérdida de energía, y que la mayoría son desviadas menos de 1° desde su dirección original. Sin embargo, una pequeña fracción sufre desviaciones mayores y alrededor de una en cada 10^4 se desvía en 90° o

⁵ No siempre los rayos β llevan carga negativa; algunas sustancias radioactivas emiten positrones, de modo que en ese caso los rayos β llevan carga positiva.

⁶ Se recomienda al alumno leer en la bibliografía citada la descripción de estos experimentos.

más. Estos fueron los resultados de los experimentos realizados por Rutherford y sus colaboradores Hans Geiger y Ernest Marsden en 1911. Vamos a ver ahora el significado de dicho resultado.

Al atravesar la lámina, las partículas α de hecho atraviesan los átomos. En su trayectoria son desviadas por los campos eléctricos debidos a las cargas internas de los átomos. Corresponde aclarar que en este caso el efecto de los electrones es despreciable, debido a su masa muy pequeña (aproximadamente 10^{-4} veces la masa de las partículas α): es fácil estimar que el orden de magnitud del ángulo de máxima desviación en una colisión entre una partícula α y un electrón atómico es de apenas 10^{-4} radianes. En consecuencia se pueden ignorar las colisiones con los electrones y basta considerar los efectos de las cargas atómicas positivas.

De acuerdo con el modelo de Thomson, los átomos constan de cargas positivas esféricas de unos 10^{-8} cm de radio, con los electrones distribuidos en su interior. Estas esferas están densamente empaquetadas en la lámina, por lo tanto si ésta tiene 10^{-4} cm de espesor, la partícula α atravesará aproximadamente 10^4 átomos. Se trata entonces de un problema de dispersión múltiple, y la desviación final de la partícula α es la suma de las desviaciones producidas por cada átomo que atravesó. Estas desviaciones tienen sentidos distribuidos al azar, de modo que podemos estimar la probabilidad de que ocurra una determinada desviación final si conocemos la desviación promedio debida a cada átomo. Se trata de un problema análogo al del “paseo al azar” y para nuestro propósito es suficiente una estimación grosera de las magnitudes de interés.

Consideremos el choque de una partícula α cuya carga es ze ($z = 2$) y cuya cantidad de movimiento es $p = mv$ con una esfera de carga positiva Ze y radio R . Suponiendo que el ángulo ϕ de desviación es pequeño, podemos escribir

$$\phi \approx \frac{\Delta p}{p} = \frac{F \Delta t}{p} \quad (3.8)$$

Podemos estimar la fuerza F como zZe^2 / R^2 y el tiempo que dura la colisión como $\Delta t \approx R/v$. Obtenemos entonces

$$\phi \approx \frac{Ze^2}{ER} \quad (3.9)$$

donde E es la energía cinética de la partícula α . Sustituyendo en (3.9) los valores típicos en estos experimentos ($Z = 80$, $e \approx 4.8 \times 10^{-10}$ u.e.s., $R \approx 10^{-10}$ m y $E \approx 5$ MeV) se obtiene

$$\phi \approx 2 \times 10^{-4} \text{ radianes} \quad (3.10)$$

Combinando un número muy grande n de estas colisiones, y suponiendo que las desviaciones de las mismas están distribuidas al azar, se obtiene la siguiente expresión para la probabilidad $P(\Phi)d\Phi$ de que la desviación total esté comprendida entre Φ y $\Phi + d\Phi$:

$$P(\Phi)d\Phi = \frac{1}{n\phi^2} e^{-\frac{\Phi^2}{2n\phi^2}} \Phi d\Phi \quad (3.11)$$

El valor medio cuadrático del ángulo de desviación total es

$$\Phi_{rms} = \left(\overline{\Phi^2} \right)^{1/2} = \phi \sqrt{n} \quad (3.12)$$

Los resultados experimentales de Rutherford y sus colaboradores mostraron que entre 0° y 3° la (3.11) describe correctamente la dispersión de las partículas α , con un valor de $\Phi_{rms} \approx 1^\circ$, lo que implica una desviación promedio por átomo de $0.01^\circ \approx 1.5 \times 10^{-4}$ radianes. De modo que en la región de desviaciones pequeñas, la concordancia con la predicción del modelo es excelente.

Pero se presenta una grave discrepancia para las desviaciones grandes, pues el experimento muestra que alrededor de 1 de cada 10^4 partículas se desvía en 90° o más. En cambio, la (3.11) predice, por ejemplo, que la probabilidad que la desviación sea mayor que 10° es de 2×10^{-22} . En consecuencia, el modelo de Thomson *no* describe correctamente las desviaciones en ángulos grandes.

El problema no tiene arreglo posible. Si, por ejemplo, disminuimos el radio R de la carga positiva para así aumentar ϕ , (pues $\phi \propto 1/R$), al disminuir el tamaño de los átomos disminuye el número n de colisiones (pues $n \propto R^2$) y en consecuencia Φ_{rms} se mantiene constante, independientemente de R . La conclusión es que la dispersión múltiple *nunca* puede producir las desviaciones a grandes ángulos que se observan.

Claramente, la condición $\Phi_{rms} = \text{cte.}$ vale solo para radios tales que $n \gg 1$. Para radios muy pequeños puede haber *una sola colisión*, y entonces el ángulo de desviación es el que corresponde a un único evento. La ec. (3.9) sugiere que esto podría ocurrir para $R \approx 10^{-12}$ cm. Pero por otra parte *todos* los métodos para medir radios atómicos indican $R \approx 10^{-8}$ cm.

Está claro entonces que el modelo de Thomson se debe descartar.

El modelo de Rutherford y el núcleo atómico

De resultados de la evidencia que acabamos de comentar, Rutherford propuso en 1911 que el átomo tiene un *núcleo* central diminuto donde reside toda la carga positiva y la mayor parte de la masa, y que los electrones giran alrededor de él⁷. Veremos ahora que este modelo está de acuerdo con los resultados de los experimentos de dispersión de partículas α , pero para eso primero tenemos que formular una teoría diferente para dicha dispersión.

La fórmula de dispersión de Rutherford

En este caso, al analizar la dispersión de partículas α , el núcleo central se puede considerar como una carga puntual. A partir de este modelo se puede obtener una fórmula para la dispersión de partículas α que concuerda muy bien con los resultados experimentales.

Las hipótesis básicas que permiten deducir dicha fórmula son:

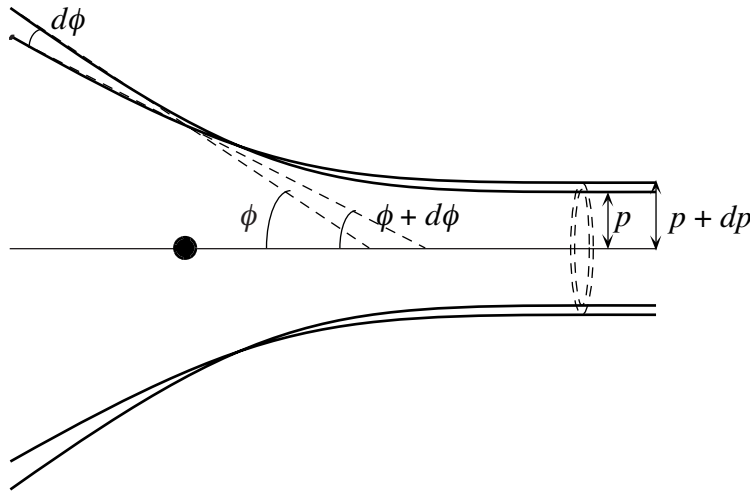
- La dispersión se debe a la interacción entre la partícula α y el núcleo, y sólo es significativa si la trayectoria pasa cerca del núcleo. Esto implica que los choques son raros y por lo tanto el problema es de una única colisión.
- La fuerza entre la partícula α y el núcleo sigue la ley de Coulomb hasta distancias muy pequeñas.

⁷ Es interesante recordar que el físico japonés Hantaro Nagaoka había propuesto en 1904 un modelo atómico planetario semejante al de Rutherford, pero no fue tomado en cuenta por sus contemporáneos debido al problema de estabilidad que se presenta con esa clase de modelos. En efecto, el electromagnetismo predice que una carga acelerada irradia; por lo tanto un electrón que gira alrededor de un núcleo debería perder energía y caer sobre él.

- Se puede ignorar el efecto de los electrones.

Nosotros para simplificar la discusión vamos a suponer además que el núcleo está fijo en el referencial del laboratorio y por brevedad citaremos los resultados sin dar la demostración, ya que la misma se puede encontrar desarrollada en la bibliografía citada y en la mayoría de los textos de Mecánica. Suponemos que la lámina tiene un espesor d , y que hay n átomos por unidad de volumen. La geometría de la colisión se indica en la Fig. 3.1. El núcleo está fijo en el origen y la partícula α se aproxima desde la derecha con velocidad inicial v . Si no hubiera desviación, la partícula pasaría a una distancia p del núcleo. La distancia p se denomina *parámetro de impacto*. Después de la dispersión la partícula se ha desviado en un ángulo ϕ y tiene la velocidad final v_f . Dado que el núcleo se supone inmóvil, tendremos que $|v| = |v_f| = v$, y como se conserva el momento angular, es sencillo verificar que

$$\cot(\phi/2) = \frac{2Ep}{zZe^2} \quad (3.13)$$



donde $E = mv^2/2$ es la energía de la partícula incidente. Como en un choque frontal la máxima distancia de acercamiento es

$$a = \frac{zZe^2}{E} \quad (3.14)$$

podemos escribir la (3.13) como

$$\cot(\phi/2) = 2p/a \quad (3.15)$$

A partir de estos resultados es fácil mostrar que la probabilidad por unidad de ángulo sólido de que una partícula α de energía E sea

Fig. 3.1. Dispersión de una partícula α por un núcleo.

dispersada en un ángulo ϕ es:

$$\frac{dP(\phi)}{d\Omega} = \frac{nd}{16} \left(\frac{zZe^2}{E} \right)^2 \csc^4(\phi/2) = \frac{nda^2}{16} \csc^4(\phi/2) \quad (3.16)$$

donde $d\Omega = 2\pi \sin\phi d\phi$. La (3.16) es la famosa fórmula de *dispersión de Rutherford*. Se puede observar que el número de partículas dispersadas es proporcional a Z^2 , de modo que el estudio de la dispersión de partículas α permite determinar Z , un dato que no se conocía bien en su tiempo, pues sólo se sabía que $Z \approx A/2$ gracias a los resultados de la dispersión de rayos X.

Las predicciones de la (3.16) fueron verificadas en 1913 por Geiger y Marsden y confirmaron de modo sorprendente dicha fórmula, y así convalidaron el modelo atómico de Rutherford. Corresponde también mencionar que la (3.16) vale también en el marco de la Mecánica Cuántica.

Dispersión hacia adelante

Para analizar la aplicación de la fórmula de Rutherford a los experimentos de dispersión de partículas α es preciso aclarar primero algunas cuestiones. Se puede ver de la (3.16) que cuando ϕ

es muy pequeño, $dP(\phi)/d\Omega$ es mayor que la unidad, y tiende al infinito para $\phi \rightarrow 0$. Claramente los experimentos indican que esto no ocurre. Pero debemos tener en cuenta que no tenemos derecho de extender la validez de la (3.16) a valores muy pequeños de ϕ , porque en ese caso no se cumplen nuestras hipótesis básicas: primero porque no se puede ignorar el efecto de los electrones, segundo porque no es cierto que la desviación total es el resultado de una única colisión. El principal efecto de los electrones (además de ocasionar ellos mismos dispersiones múltiples) es el de “apantallar” la carga nuclear, de forma que una partícula α que pasa lejos del núcleo no experimenta *toda* la repulsión Coulombiana del mismo. Si bien por ahora no sabemos cómo es la distribución de carga debida a los electrones, lo que sí sabemos es que a una distancia del núcleo del orden del radio atómico R , el apantallamiento es completo y por lo tanto las partículas α que inciden con parámetros de impacto mayores que R no experimentan desviación. Si el número de átomos por unidad de área del blanco es bajo, de manera que $nd\pi R^2 < 1$, entonces la probabilidad de dispersión es menor que la unidad. Esto es lo que ocurre, por ejemplo, con blancos gaseosos de poco espesor (un blanco gaseoso de 1 cm de espesor a 10^{-2} mm Hg contiene aproximadamente 4×10^{14} átomos/cm², luego $nd\pi R^2 \approx 0.1$). En cambio cuando $nd\pi b^2 > 1$ ocurre dispersión múltiple. En este caso, la probabilidad total de dispersión *no* es igual a la probabilidad de dispersión por átomo multiplicada por el número de átomos en el blanco, pues una vez que una partícula α ha chocado una vez, los choques subsiguientes *no* aumentan la probabilidad de dispersión sino que solamente modifican la desviación, como ya se vio anteriormente. Por lo tanto, para láminas metálicas, la distribución angular de las partículas α dispersadas varía con continuidad desde la distribución de Rutherford para ángulos grandes hasta la distribución de dispersión múltiple para ángulos pequeños, de modo tal que la probabilidad total de dispersión es siempre igual o menor que la unidad.

El tamaño del núcleo

La fórmula de Rutherford (3.16) se obtuvo suponiendo que el núcleo y las partículas α son cargas puntuales. Esta aproximación vale siempre y cuando la distancia de máximo acercamiento sea mayor que el radio del núcleo. En consecuencia se puede decir que el radio del núcleo debe ser *seguramente menor* que el valor más pequeño del acercamiento máximo para el cual la (3.16) da todavía el valor correcto del número de partículas α dispersadas.

En la dispersión de las partículas α emitidas por el radio ($E = 5.3$ MeV) por un blanco de cobre (para el cual $Z = 29$, como se deduce a partir del número de partículas dispersadas) se observó que la ley (3.16) vale hasta 180° . En este caso de la (3.14) se deduce que

$$a = 1.58 \times 10^{-12} \text{ cm} \quad (3.17)$$

y por lo tanto se concluye que el radio del núcleo del cobre es menor o igual que dicha cantidad.

Comentarios sobre el modelo de Rutherford

En conclusión, los experimentos de Rutherford demostraron que el átomo contiene un núcleo de muy pequeñas dimensiones (radio menor o igual a 10^{-12} cm), donde reside toda la carga positiva y casi toda la masa del átomo. También mostraron que la carga positiva, medida en unidades de

la carga electrónica, es decir el número Z , es igual (dentro del error experimental) al número atómico del elemento⁸.

Por todo ello Rutherford propuso un modelo en el que el átomo está formado por un núcleo con una carga positiva igual a Ze , alrededor del cual giran Z electrones, de manera que su movimiento equilibra dinámicamente la atracción Coulombiana que sobre ellos ejerce el núcleo. La extensión del movimiento electrónico determina que el tamaño del átomo sea de aproximadamente 10^{-8} cm. Este concepto sugiere que las propiedades químicas de un elemento están determinadas por el número de electrones de sus átomos (y no por la masa atómica, como se pensaba hasta entonces), y también sugiere que deben existir elementos para todos los valores posibles de Z , hasta que Z se hace demasiado grande y el núcleo se vuelve inestable pues probablemente ya no puede mantener más carga positiva en su interior.

Sin embargo, es evidente que el modelo todavía es incompleto y presenta serias fallas. En primer lugar no se da todavía ninguna idea acerca de cómo los electrones determinan las propiedades químicas. En segundo lugar tampoco se explica porqué todos los átomos de una misma especie tienen aparentemente el mismo tamaño, y además, porqué los átomos de especies muy distintas tienen casi exactamente el mismo tamaño. Comentaremos este asunto en la siguiente Sección. Por último, el modelo enfrenta una objeción aún más grave: al girar alrededor del núcleo, los electrones sufren una aceleración continua y por lo tanto deberían irradiar ondas electromagnéticas y perder energía, y en consecuencia caer en espiral hacia el núcleo. Veremos más adelante que la solución de esta dificultad requiere abandonar la física clásica.

La constante que está faltando

Consideremos el caso más sencillo, el del átomo de Hidrógeno, pues es suficiente para entender la esencia de la dificultad de explicar los tamaños atómicos. Supongamos que el electrón gira alrededor del núcleo en una órbita circular. La fuerza centrípeta es la fuerza de Coulomb y por lo tanto tendremos

$$\frac{mv^2}{r} = \frac{e^2}{r^2} \quad (3.18)$$

que nos da la relación

$$rv^2 = \frac{e^2}{m} \quad (3.19)$$

En esta ecuación r puede tener cualquier valor, pues la (3.19) sólo nos dice que si escalamos r por un factor k^2 es necesario escalar v por el factor k^{-1} . Se puede verificar que la (3.19) equivale a la Tercera Ley de Kepler.

El origen del problema reside en que las únicas constantes que intervienen en el modelo clásico del átomo de Rutherford son la carga y la masa del electrón, y con ellas *no podemos* formar una longitud característica.

Existe una manera de combinar e y m con una constante universal clásica, de modo de formar una longitud, y es usar la velocidad de la luz, c . Se construye así la constante

⁸ Es decir el número por el cual se ordenan los elementos en la Tabla Periódica.

$$r_e \equiv \frac{e^2}{mc^2} = 2.8 \times 10^{-13} \text{ cm} \quad (3.20)$$

Esta longitud es el *radio clásico del electrón*, que ya mencionamos en conexión con la dispersión Thomson, pero no tiene nada que ver con las dimensiones atómicas, ya que surge de comparar el equivalente energético de la masa del electrón con la energía del campo eléctrico del mismo electrón. Es obvio que la velocidad de la luz puede intervenir sólo cuando se consideran efectos relativísticos, y este *no* es el caso del átomo pues se ve de inmediato a partir de la (3.19) que si $r \approx 10^{-8}$ cm resulta que $v \ll c$.

Adelantándonos a lo que veremos más adelante en detalle, es fácil ver que el dilema se resuelve si introducimos la *constante de Planck* h ($h = 6.63 \times 10^{-27}$ erg s). Es posible entonces formar la longitud

$$a_0 \equiv \frac{(h/2\pi)^2}{me^2} = 0.529 \times 10^{-8} \text{ cm} \quad (3.21)$$

Esta longitud, que se denomina *radio de Bohr*, tiene efectivamente el orden de magnitud *correcto*. Esto sugiere que la constante de Planck desempeña algún papel con respecto del tamaño de los átomos, aunque por el momento no sabemos cuál, ni porqué.

En 1913 Niels Bohr reformó el modelo de Rutherford del átomo de hidrógeno, postulando que sólo estaban permitidas las órbitas circulares cuyo momento angular fuera un múltiplo entero de $\hbar = h/2\pi$, es decir, aquellas órbitas que cumplen

$$mvr = n\hbar \quad , \quad n = 1, 2, 3, \dots \quad (3.22)$$

Sustituyendo (3.22) en (3.19) obtenemos

$$r = \frac{n^2 \hbar^2}{me^2} = n^2 a_0 \quad (3.23)$$

Por lo tanto, de acuerdo con la condición de Bohr (3.22), a_0 es el radio de la primera órbita permitida ($n = 1$) y representa efectivamente el radio del átomo de hidrógeno.

La hipótesis de Bohr se encuadra dentro de lo que hoy se denomina Teoría Cuántica Antigua, que tuvo considerable éxito, pues permitió interpretar varias propiedades atómicas y también resolver la paradoja de los calores específicos que resulta como consecuencia del Teorema de Equipartición. Sin embargo esta teoría presenta varios inconvenientes y finalmente fue abandonada cuando se introdujo la Mecánica Cuántica moderna. En el Capítulo 5 la comentaremos más en detalle. Pero primero conviene discutir algunas características de la radiación electromagnética y de sus interacciones con la materia, que fueron históricamente las que llevaron a introducir la constante de Planck y además mostraron que las propiedades de la radiación no se pueden explicar satisfactoriamente en términos del concepto clásico de onda.

4. RADIACIÓN, FOTONES Y LA CONSTANTE DE PLANCK

Introducción

En el Capítulo 3 vimos que el modelo atómico de Rutherford está de acuerdo con los resultados de los experimentos de dispersión de partículas α , pero que no hay ningún concepto de la física clásica que permite explicar el tamaño de los átomos. Aparentemente hay una constante de la naturaleza, que aún no sabemos cómo interviene en la teoría, que determina ésta y otras propiedades no clásicas de la materia y la radiación. Cuando Rutherford formuló su modelo ya se conocía esa constante: se trata de la *constante de acción de Planck*, que fuera introducida por Max Planck cuando presentó un artículo sobre la radiación de cuerpo negro en la Sociedad Física Alemana a fines de 1900. En este Capítulo presentaremos algunas evidencias de la naturaleza universal de dicha constante, en lo referente a fenómenos en los que interviene la radiación electromagnética.

La teoría de Planck de la radiación de cuerpo negro

Ya estudiamos varios aspectos de la radiación de cuerpo negro¹ y por lo tanto no volveremos sobre ello, pero queremos recordar que al estudiar la Mecánica Estadística mostramos que la cantidad de modos de radiación electromagnética de frecuencia comprendida entre ν y $\nu + d\nu$ presentes en una cavidad de volumen V está dada por

$$N(\nu)d\nu = \frac{8\pi V}{c^3} \nu^2 d\nu \quad (4.1)$$

y que la aplicación del teorema de equipartición, según el cual a cada modo de oscilación del campo electromagnético le corresponde en el equilibrio una energía media $\bar{\epsilon} = kT$ (k es la constante de Boltzmann), lleva a la distribución espectral de Rayleigh-Jeans²:

$$u(\nu, T)d\nu = \frac{8\pi\nu^2}{c^3} \bar{\epsilon} d\nu = \frac{8\pi\nu^2 kT}{c^3} d\nu \quad (4.2)$$

donde $u(\nu, T)$ es la densidad de energía (energía por unidad de volumen) en el intervalo de frecuencia $(\nu, \nu + d\nu)$. La (4.2) contradice la experiencia y es a todas luces absurda pues al integrarla sobre todas las frecuencias predice una densidad de energía divergente (la “catástrofe ultravioleta”). En esa época se conocía la *Ley de Wien*, que se deduce a partir de argumentos puramente termodinámicos, y que establece que

$$u(\nu, T) = c_1 \nu^3 g(c_2 \nu / T) \quad , \quad c_1, c_2 = \text{cte.} \quad (4.3)$$

donde g es una función desconocida. La (4.3) implica que el máximo de $u(\nu, T)$ ocurre para una frecuencia ν_m que es proporcional a la temperatura, esto es $\nu_m \sim T$. Por otra parte las medi-

¹ Ver *Termodinámica e introducción a la Mecánica Estadística*, Capítulos 15 y 18.

² Deducida por Lord Rayleigh (John William Strutt) en 1900. En su trabajo, publicado en *Nature* en 1905, estimó mal el número de modos, por un factor 8 en exceso, error que fue corregido ese mismo año por James Jeans. También corresponde mencionar que en 1905 Einstein obtuvo la fórmula (4.2) en su artículo sobre el efecto fotoeléctrico que se comenta más adelante.

ciones de $u(\nu, T)$ que se conocían antes de 1900 cubrían solamente el rango de frecuencias altas, más allá del máximo de $u(\nu, T)$. En base a esos datos Wien propuso en 1896 la fórmula empírica

$$u(\nu, T) = A\nu^3 e^{-B\nu/T} \quad , \quad A, B = \text{cte.} \quad (4.4)$$

que tiene la forma (4.3) y por lo tanto cumple la Ley de Wien, y con elecciones adecuadas de las constantes permite un buen ajuste de las mediciones disponibles hasta ese momento.

En 1900 Otto Lummer y Ernst Pringsheim, y también Heinrich Rubens y Ferdinand Kurlbaum llevaron a cabo mediciones muy precisas en un rango de frecuencias bajas que hasta entonces no de había estudiado y encontraron que la fórmula (4.4) está en desacuerdo con los valores medidos, según los cuales para frecuencias muy bajas se tiene

$$u(\nu, T) \sim \nu^2 T \quad (4.5)$$

Al conocer estos resultados³ Planck propuso una fórmula que interpolara entre (4.4) y (4.5). Dicha fórmula es la siguiente:

$$u(\nu, T)d\nu = \frac{8\pi\nu^2}{c^3} \frac{h\nu}{e^{h\nu/kT} - 1} d\nu \quad (4.6)$$

y es la célebre *distribución de Planck*. Ajustando la (4.6) a la distribución espectral observada, Planck determinó el valor de h (denominada *constante de Planck*) como 6.55×10^{-27} erg s, un valor muy cercano al actual⁴, con el cual obtuvo un perfecto acuerdo con las mediciones. No conforme con esto procuró comprender porqué la (4.6) concuerda tan bien con los resultados experimentales y con ese fin trató de darle una base teórica. No discutiremos aquí los argumentos de Planck, que se explican en detalle en el Capítulo 18 de *Termodinámica e Introducción a la Mecánica Estadística*, sólo diremos que llegó a la conclusión que la energía de los osciladores materiales de la pared de la cavidad, que están en equilibrio con la radiación, se forma a partir de un número *finito* de *cuantos* o parcelas de energía, cada uno de energía $\varepsilon = h\nu$, con lo cual introdujo la *cuantificación* de la energía. Pero en realidad fue Einstein quien aclaró⁵ el significado de la fórmula de Planck y de las hipótesis sobre la cuales se basa, al expresar que la energía de los osciladores de Planck puede tomar sólo valores que son *múltiplos enteros* del cuanto $h\nu$, y que en la absorción y emisión de radiación, la energía del oscilador cambia en forma discontinua.

A continuación vamos a mostrar cómo se debe modificar la teoría clásica para obtener el resultado correcto (4.6). Recordemos el origen de la Ley de Equipartición. Básicamente es una consecuencia de la distribución de Boltzmann, que en este caso expresa que la probabilidad $P(\varepsilon)d\varepsilon$ que un modo de oscilación tenga una energía entre ε y $\varepsilon + d\varepsilon$ vale

$$P(\varepsilon)d\varepsilon = \frac{e^{-\varepsilon/kT}}{Z} \quad , \quad Z = \int_0^{\infty} e^{-\varepsilon/kT} d\varepsilon \quad (4.7)$$

³ El lector puede notar que (4.5) es el resultado de la teoría de Rayleigh-Jeans, que fue publicada recién en 1905, después de los hechos que estamos relatando. Por lo tanto Planck no conocía el resultado clásico (4.2).

⁴ El valor actual de h es $6.6260755 \times 10^{-27}$ erg s.

⁵ En su artículo de 1906 sobre la teoría del calor específico de un sólido.

En términos de la distribución de Boltzmann, $\bar{\varepsilon}$ se expresa como

$$\bar{\varepsilon} = \frac{\int_0^{\infty} \varepsilon P(\varepsilon) d\varepsilon}{\int_0^{\infty} e^{-\varepsilon/kT} d\varepsilon} = \frac{\int_0^{\infty} \varepsilon e^{-\varepsilon/kT} d\varepsilon}{\int_0^{\infty} e^{-\varepsilon/kT} d\varepsilon} \quad (4.8)$$

y al calcular las integrales en la (4.8) se obtiene precisamente $\bar{\varepsilon} = kT$.

La gran contribución de Planck consistió en que advirtió que podía obtener el resultado que buscaba si trataba ε como una variable *discreta*, en lugar de la variable *continua* que es en la Física Clásica. A partir de esto Einstein llegó a la conclusión que para llegar a la (4.6) era preciso suponer que la energía de cada *modo de oscilación del campo de radiación* toma sólo ciertos valores discretos, múltiplos enteros de $h\nu$, en lugar de cualquier valor. Supongamos entonces que los valores *permitidos* de la energía son

$$\varepsilon_{\nu} = 0, h\nu, 2h\nu, 3h\nu, \dots \quad (4.9)$$

Siendo así las integrales en la (*4.7) se deben reemplazar por *sumas*:

$$\bar{\varepsilon}_{\nu} = \frac{\sum_{n=0}^{\infty} nh\nu P(nh\nu)}{\sum_{n=0}^{\infty} P(nh\nu)} = \frac{\sum_{n=0}^{\infty} \frac{nh\nu}{kT} e^{-nh\nu/kT}}{\sum_{n=0}^{\infty} \frac{1}{kT} e^{-nh\nu/kT}} = h\nu \frac{\sum_{n=0}^{\infty} n e^{-n\alpha}}{\sum_{n=0}^{\infty} e^{-n\alpha}} = -h\nu \frac{d}{d\alpha} \left(\ln \sum_{n=0}^{\infty} e^{-n\alpha} \right) \quad (4.10)$$

donde $\alpha = h\nu/kT$. Pero

$$\sum_{n=0}^{\infty} e^{-n\alpha} = \frac{1}{1 - e^{-\alpha}} \quad (4.11)$$

y por lo tanto

$$\bar{\varepsilon}_{\nu} = -h\nu \frac{d}{d\alpha} \left[\ln \left(\frac{1}{1 - e^{-\alpha}} \right) \right] = h\nu \frac{d}{d\alpha} \ln(1 - e^{-\alpha}) = \frac{h\nu e^{-\alpha}}{1 - e^{-\alpha}} = \frac{h\nu}{e^{\alpha} - 1} \quad (4.12)$$

Resulta entonces⁶

$$\bar{\varepsilon}_{\nu} = \frac{h\nu}{e^{h\nu/kT} - 1} \quad (4.13)$$

Usando este valor de $\bar{\varepsilon}_{\nu}$ en la (4.2) en lugar de kT se obtiene la (4.6), con lo cual queda aclarado el origen de dicha fórmula.

Recapitulando, no es preciso alterar la distribución de Boltzmann, sino que hay que postular que la energía de los osciladores de la radiación misma puede tomar solamente los valores discretos (4.9) y no cualquier valor, como predice la teoría clásica.

⁶ Esta deducción de la distribución de Planck basada en suponer que la energía de cada modo de oscilación del campo de radiación está cuantificada fue dada por Debye en 1910.

El postulado de Planck

La hipótesis de Planck se puede generalizar y enunciar como un postulado del modo siguiente:

Postulado de Planck:

cualquier ente físico con un grado de libertad mecánico, cuya coordenada generalizada realiza oscilaciones armónicas simples sólo puede poseer valores discretos de la energía, dados por

$$\varepsilon = nh\nu \quad , \quad n = 0, 1, 2, \dots \quad (4.14)$$

donde ν es la frecuencia de la oscilación y h es una constante universal.

Si la energía del ente obedece el postulado de Planck se dice que está *cuantificada*, los niveles de energía permitidos se llaman *estados cuánticos* y el número entero n se denomina *número cuántico*.

Se podría tal vez pensar que hay sistemas físicos cuyo comportamiento no está de acuerdo con el postulado de Planck, por ejemplo sistemas macroscópicos como el péndulo. Pero es fácil constatar que debido a la pequeñez de h , en esos casos la diferencia de energía $h\nu$ entre los niveles es insignificante: no se puede medir la energía de un sistema macroscópico con suficiente precisión como para verificar si el postulado de Planck se cumple o no. Las consecuencias del postulado de Planck se ponen de manifiesto solamente cuando ν es muy grande y/o ε es tan pequeña que $\Delta\varepsilon = h\nu$ es del mismo orden que ε . Un caso es el que acabamos de estudiar: la radiación de alta frecuencia del cuerpo negro.

Cabe subrayar que Planck restringió la cuantificación de la energía a las oscilaciones de los *electrones* radiantes de la pared de la cavidad del cuerpo negro, pues pensaba que una vez emitida, la energía electromagnética se esparcía por el espacio en forma de ondas. La extensión de la cuantificación de la energía a la radiación misma se debe a Einstein, quien (como veremos al tratar el efecto fotoeléctrico) propuso que la energía radiante está cuantificada en paquetes concentrados que hoy llamamos *fotones* o *cuantos de luz*.

También debemos mencionar que el postulado de Planck resuelve el problema de los calores específicos de los sólidos, cuya dependencia con la temperatura no se describe correctamente mediante la teoría clásica basada en la equipartición de la energía. El problema es enteramente análogo al de la radiación de cuerpo negro y el lector lo puede encontrar tratado en el Capítulo 18 de *Termodinámica e introducción a la Mecánica Estadística*.

En lo que sigue examinaremos los procesos de interacción entre la radiación y la materia. Consideraremos el efecto fotoeléctrico, el efecto Compton y la creación de pares, que implican la absorción o dispersión de radiación, y procesos de emisión de radiación como el *bremssstrahlung* (o radiación de frenamiento) y la aniquilación de pares. En todos estos casos la evidencia experimental indica que en sus interacciones con la materia la radiación se comporta como corpúsculo, a diferencia de su naturaleza ondulatoria cuando se propaga.

El efecto fotoeléctrico

El efecto fotoeléctrico fue descubierto por Hertz en 1887, cuando observó que una descarga eléctrica entre dos electrodos se produce más fácilmente si sobre uno de ellos incide luz ultravioleta. Poco después, los trabajos de Wilhelm Hallwachs (1888), J. J. Thomson (1899) y Philipp L. A. Lenard (1900) demostraron que la luz ultravioleta facilita la descarga porque provoca la emisión de electrones desde la superficie del cátodo y determinaron las características

de dicha emisión, un fenómeno que se denominó *efecto fotoeléctrico*. La Fig. 4.1 muestra el esquema de un experimento para estudiar el efecto fotoeléctrico

La Fig. 4.2 muestra la corriente fotoeléctrica como función de la diferencia de potencial entre cátodo y ánodo. Se observa que para V suficientemente grande, i alcanza un valor límite, o de *saturación*, para el cual todos los electrones emitidos por el cátodo son colectados por el ánodo. La corriente de saturación es proporcional a la intensidad del haz de luz que incide sobre el cátodo. Si V se hace negativo, la corriente no cae de inmediato a cero, lo que sugiere que los electrones son emitidos con cierta energía cinética, de modo que algunos alcanzan el otro electrodo a pesar que el campo eléctrico se opone a su movimiento. Sin embargo, para cierto valor negativo V_0 , llamado *potencial de frenamiento*, la corriente fotoeléctrica se anula.

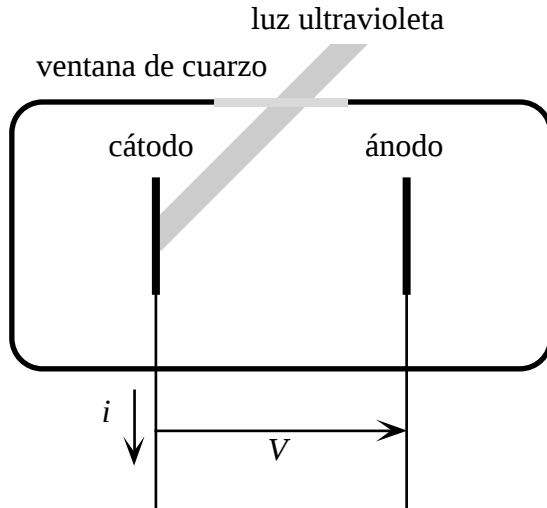


Fig. 4.1. Experimento para estudiar el efecto fotoeléctrico. El dispositivo está bajo vacío. El voltaje V entre el cátodo y el ánodo se puede variar de forma continua, y se mide la corriente i .

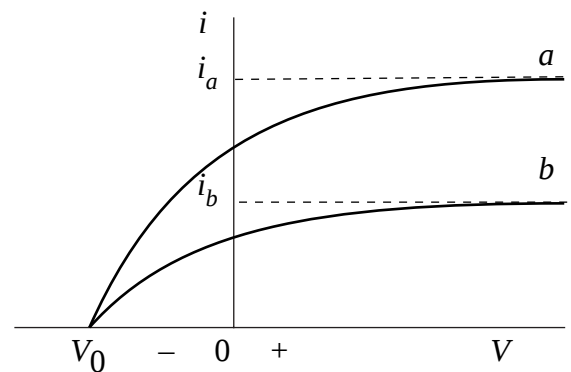


Fig. 4.2. Corriente vs. voltaje en el aparato de la fig. 4.1. La curva b corresponde a luz cuya intensidad es la mitad de la de la curva a .

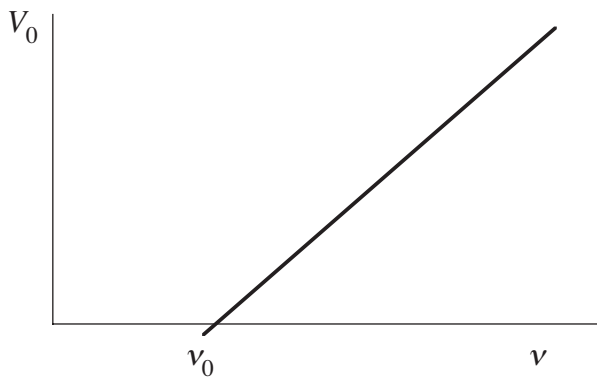


Fig. 4.3. Potencial de frenamiento en función de la frecuencia de la luz.

debajo de la cual no hay efecto fotoeléctrico.

Hay tres aspectos fundamentales del efecto fotoeléctrico que no se pueden explicar en términos de la teoría ondulatoria clásica de la luz:

- Según la teoría ondulatoria, la intensidad del haz luminoso es proporcional al cuadrado de la amplitud E del campo eléctrico oscilante de la onda. Como la fuerza sobre el electrón es eE , la energía cinética de los fotoelectrones debería aumentar con la intensidad del haz, pero el

La energía cinética máxima de los fotoelectrones es entonces

$$K_{\max} = eV_0 \quad (4.15)$$

y es *independiente de la intensidad de la luz*.

El comportamiento del potencial de frenamiento V_0 como función de la frecuencia ν de la luz fue estudiado por Millikan en 1914 y se obtuvo un gráfico lineal como el de la Fig. 4.3. Se observa una frecuencia de corte ν_0 bien definida, que depende del material del cátodo, por

experimento muestra que $K_{\max}$ es independiente de la intensidad del haz; esto fue probado sobre un rango de intensidades de 10^7 .

- Según la teoría ondulatoria el efecto fotoeléctrico debería ocurrir para cualquier frecuencia, con tal que el haz tenga suficiente intensidad como para suministrar la energía necesaria para emitir los fotoelectrones. Pero el experimento muestra que para cada superficie hay una frecuencia de corte ν_0 , por debajo de la cual no hay efecto fotoeléctrico, sin que importe cuán intenso sea el haz de luz.
- Si la energía que adquiere el fotoelectrón es absorbida de la onda, debe tenerse presente que la sección eficaz de absorción para un electrón en un metal difícilmente sea mucho mayor que la sección transversal de un átomo. Por otra parte, en la teoría clásica, la energía luminosa está uniformemente distribuida sobre el frente de onda. Por lo tanto, si la intensidad de la luz es baja, cabría esperar que exista un tiempo de retraso fácilmente medible, entre el instante en que la luz comienza a incidir sobre el cátodo y el momento en que es emitido el fotoelectrón, pues durante ese intervalo el electrón irá absorbiendo la energía del haz hasta acumular la que necesita para escapar⁷. Sin embargo nunca se observó tal retraso.

Teoría cuántica de Einstein del efecto fotoeléctrico

En 1905 Einstein, influenciado por los trabajos de Lenard, puso en duda la teoría clásica de la luz y varios años antes de los experimentos de Millikan propuso que la energía luminosa está cuantificada en paquetes concentrados, a los que hoy llamamos fotones. El argumento de Einstein se apoyaba en que los experimentos de interferencia y difracción de la luz, sobre los cuales se basa la teoría ondulatoria, se efectúan en situaciones en que el número de fotones es muy grande. Por lo tanto sus resultados representan el promedio de los comportamientos de los fotones individuales, lo que explica porqué en esos experimentos no manifiestan los fotones.

Por cierto, los experimentos clásicos de interferencia y difracción muestran de manera definitiva que los fotones no viajan desde donde son emitidos hasta donde son finalmente absorbidos del mismo modo que lo haría una partícula clásica: viajan como ondas, en el sentido que los cálculos basados en la propagación de ondas explican correctamente los patrones de interferencia y de difracción, que dependen, como dijimos, del modo en que se desplazan en promedio los fotones. Pero Einstein no se preocupó de la propagación de la radiación, sino de como es emitida y absorbida. Pensó que si la energía contenida en las ondas electromagnéticas de frecuencia ν sólo puede ser un múltiplo entero de $h\nu$, entonces en el proceso de emisión se producen *cuantos* de energía electromagnética, cada uno de los cuales lleva una energía $h\nu$. Einstein supuso que esos cuantos están localizados inicialmente en una pequeña región del espacio, y que se mantienen localizados mientras se alejan de la fuente con la velocidad c . Supuso además que la energía de cada paquete o fotón está relacionada con su frecuencia de acuerdo con la ecuación

$$E = h\nu \quad (4.16)$$

También supuso que en el efecto fotoeléctrico cada fotón es *completamente absorbido* por un electrón.

⁷ Una sencilla estimación muestra que si la energía necesaria para extraer un fotoelectrón es de 2 eV y la sección eficaz de absorción es igual al área transversal de un átomo, con una fuente luminosa de 1 W a 1 m de distancia se precisarían 200 s para que el electrón adquiriera la energía necesaria para escapar.

Cuando un electrón es emitido por el fotocátodo, su energía cinética es entonces

$$K = h\nu - w \quad (4.17)$$

donde $h\nu$ es la energía del fotón absorbido y w es el trabajo necesario para extraer el electrón del metal. Este trabajo toma en cuenta el efecto de los campos eléctricos atractivos debidos a los átomos de la superficie y las pérdidas de energía cinética del electrón causadas por las colisiones que sufre hasta que sale de la superficie. Algunos electrones están ligados más fuertemente que otros, o soportan más colisiones en el trayecto. Por lo tanto es lógico suponer que hay una energía cinética máxima con la cual un fotoelectrón puede ser emitido, que se tiene cuando la energía de unión del electrón es la mínima posible y cuando éste no pierde energía cinética por colisiones. Esta energía cinética máxima es

$$K_{\max} = K_{\max}(\nu) = h\nu - w_0 \quad (4.18)$$

donde w_0 , que se denomina la *función trabajo*, es la energía mínima necesaria para extraer un electrón del metal y es una propiedad del metal del cátodo.

Veremos ahora que las hipótesis de Einstein explican satisfactoriamente el efecto fotoeléctrico, lo que no se logra en cambio con la teoría ondulatoria clásica.

- El primer hecho, que $K_{\max}$ es independiente de la intensidad de la luz, es consecuencia de la (4.18). De acuerdo con la teoría de Einstein, la intensidad de la luz es proporcional al *número* de fotones que llegan a la superficie (por unidad de tiempo y de área), por lo tanto la corriente fotoeléctrica es proporcional a la intensidad, pero cada proceso individual de emisión es *independiente* de la intensidad y sólo depende de la *frecuencia* de la radiación.
- La existencia de una *frecuencia de corte*, es también una consecuencia de la (4.18). En efecto, para cada metal, existe una frecuencia ν_0 tal que

$$h\nu_0 = w_0 \quad (4.19)$$

y para esa frecuencia $K_{\max}(\nu_0) = 0$. En otras palabras, un fotón de la frecuencia ν_0 tiene justamente la energía suficiente para extraer un fotoelectrón, sin que le sobre nada que pueda aparecer como energía cinética del electrón. Si $\nu < \nu_0$, los fotones no tienen energía suficiente para extraer fotoelectrones y no hay efecto fotoeléctrico, por grande que sea la intensidad (o sea el flujo de fotones) del haz de luz.

- La *ausencia del tiempo de retraso* también se explica, pues la energía necesaria para que el electrón sea emitido se suministra en paquetes concentrados. Tan pronto la iluminación del fotocátodo deja de ser nula, hay por lo menos *un* fotón, que puede ser absorbido provocando la emisión de un fotoelectrón.

Si sustituimos $K_{\max} = eV_0$ (ec. (4.15)) en la ecuación de Einstein (4.18) obtenemos

$$V_0 = \frac{h\nu}{e} - \frac{w_0}{e} \quad (4.30)$$

Por lo tanto, la teoría de Einstein predice que el potencial de frenamiento V_0 es una función *lineal* de ν , en perfecto acuerdo con los resultados experimentales (Fig. 4.3). La pendiente de la curva experimental permite determinar

$$\frac{h}{e} = 3.9 \times 10^{-15} \text{ V s} \quad (4.31)$$

Multiplicando esta relación por la carga electrónica se puede determinar h . Examinando cuidadosamente estos datos Millikan encontró $h = 6.57 \times 10^{-34} \text{ J s}$ con una precisión del 0.5%, valor que concuerda con el que obtuvo Planck. Es ciertamente notable el acuerdo entre estas dos determinaciones de h usando fenómenos y teorías completamente diferentes, y es una prueba de la validez de la suposición fundamental del cuanto de luz, o fotón.

Actualmente el concepto de fotón se usa en todo el espectro electromagnético y no sólo en el visible o cerca de él. La energía de los fotones varía en 18 órdenes de magnitud desde las ondas de VLF ($\nu \approx 10^4 \text{ Hz}$) a los fotones más energéticos de los rayos cósmicos ($\nu \approx 10^{22} \text{ Hz}$).

Cabe mencionar que para que un fotón pueda ser absorbido, como ocurre en el efecto fotoeléctrico, es preciso que el electrón esté *ligado* a un átomo o a un sólido. Un electrón libre no puede absorber un fotón pues, como veremos luego, en tal proceso no se pueden conservar simultáneamente la energía y la cantidad de movimiento (por la misma razón un electrón libre tampoco puede emitir un fotón). Para que el electrón absorba (o emita) un fotón debe intervenir en el proceso una tercera partícula, pues sólo así se pueden conservar tanto la energía como la cantidad de movimiento. Debido a la gran masa de un átomo o de un sólido, el sistema puede absorber la cantidad de movimiento necesaria para el balance sin adquirir una cantidad apreciable de energía. Por lo tanto la ecuación de la energía (4.18) sigue siendo válida, y el efecto es posible porque además del electrón emitido existe una partícula pesada que absorbe la cantidad de movimiento necesaria para conservar el impulso. Para fotones de energía comparable a las de los rayos X o superiores, el efecto fotoeléctrico es un mecanismo importante de absorción. A energías todavía mayores se vuelven importantes otros procesos que estudiaremos más adelante. Por último, es preciso recalcar que en el modelo de Einstein un fotón de frecuencia ν tiene *exactamente* la energía $h\nu$, y no múltiplos enteros de $h\nu$. Según el modelo de Einstein, la radiación de la cavidad de un cuerpo negro constituye un “gas de fotones”. Aplicando este concepto, años después de la deducción de Planck, Satyendra Nath Bose volvió a obtener la fórmula⁸ (4.6).

El efecto Compton

En 1923 los experimentos de Arthur Holly Compton dieron una nueva confirmación de la naturaleza corpuscular de la radiación. Compton hizo incidir un haz colimado de rayos X de longitud de onda λ bien definida sobre un blanco de grafito y midió la intensidad y la longitud de onda de los rayos dispersados en varias direcciones. Se observó que aunque el haz incidente consiste esencialmente de *una única* longitud de onda, en los rayos X dispersados en direcciones que forman un ángulo θ no nulo con la dirección del haz, aparecen *dos* longitudes de onda: una es la misma λ de la radiación incidente y la otra, λ' , es mayor, esto es $\lambda' = \lambda + \Delta\lambda$. Este corrimiento $\Delta\lambda$, denominado *corrimiento Compton*, varía con el ángulo en que se observan los rayos X dispersados.

⁸ Es interesante mencionar que el artículo de Bose fue rechazado por el árbitro de la revista Philosophical Magazine, a la cual lo había presentado; de resultados de ello Bose lo envió a Einstein, quien lo tradujo personalmente al alemán y lo hizo publicar el Zeitschrift für Physik. Posteriormente Einstein generalizó el método de Bose y lo aplicó a partículas con masa. Las teorías de Bose y de Einstein se tratan en el Capítulo 15.

La presencia de una longitud de onda *diferente* de la de la radiación incidente en la radiación dispersada *no se puede explicar* si se considera la radiación X como una onda electromagnética clásica. De acuerdo con el Electromagnetismo clásico el campo eléctrico de la onda incidente, que oscila con la frecuencia ν , actúa sobre los electrones libres del blanco y los fuerza a oscilar con esa misma frecuencia. Los electrones oscilantes irradian en todas las direcciones ondas electromagnéticas, de frecuencia igual a la de la oscilación. Por lo tanto, según la teoría clásica, la radiación dispersada tiene la misma frecuencia y longitud de onda que la radiación incidente.

Ante estos hechos, Compton (y también independientemente Peter Debye) interpretó los resultados del experimento suponiendo que la radiación X incidente está compuesta por fotones,

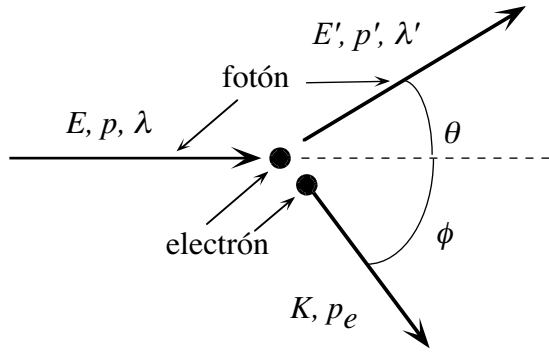


Fig. 4.4. Efecto Compton.

cada uno de los cuales lleva una energía $E = h\nu$ y que esos fotones chocan con los electrones libres del blanco. Según este punto de vista, los fotones que han chocado constituyen la radiación dispersada. Como en la colisión el fotón cede parte de su energía al electrón, el fotón dispersado debe tener una energía menor, $E' < E$, luego una frecuencia menor $\nu' = E'/h$ y entonces una longitud de onda mayor $\lambda' = c/\nu'$. Así se explica cualitativamente el corrimiento $\Delta\lambda = \lambda' - \lambda$.

Los experimentos indican que la frecuencia de la radiación dispersada es *independiente* del material

del blanco, lo que significa que la colisión no involucra átomos completos. Por lo tanto Compton supuso que la dispersión se debe a la colisión entre un fotón y un particular electrón del blanco, y además supuso que los electrones que participan en estos procesos se pueden considerar *libres* e inicialmente en reposo. Se pueden justificar estas suposiciones teniendo en cuenta que la energía de un fotón X es varios órdenes de magnitud mayor que la de un fotón ultravioleta, y al estudiar el efecto fotoeléctrico vimos que esta última es comparable a la energía mínima con la cual un electrón está ligado a un metal.

Vamos a analizar ahora cuantitativamente la colisión entre un fotón y un electrón (Fig. 4.4). Si suponemos que el fotón es una partícula de energía $E = h\nu$ que se mueve con la velocidad de la luz, es evidente que su masa en reposo debe ser nula, por lo tanto su energía total relativística es puramente cinética. La cantidad de movimiento del fotón se puede obtener de la relación general entre la energía relativística E , la cantidad de movimiento p y la masa en reposo m :

$$E^2 = c^2 p^2 + m^2 c^4 \quad (4.32)$$

Puesto que para un fotón $m = 0$, resulta

$$p = E/c = h\nu/c = h/\lambda \quad (4.33)$$

Cabe observar que el Electromagnetismo clásico conduce también a la relación $p = E/c$ donde p y E son respectivamente la cantidad de movimiento y la energía de la radiación, por unidad de volumen.

La conservación de la cantidad de movimiento (ver Fig. 4.4) implica

$$p = p' \cos \theta + p_e \cos \phi \quad (4.34)$$

y

$$p' \sin \theta = p_e \sin \phi \quad (4.35)$$

donde p' y p_e son la cantidad de movimiento del fotón y del electrón luego de la colisión y θ, ϕ son los ángulos entre $\mathbf{p}$ y $\mathbf{p}'$, $\mathbf{p}_e$, respectivamente. Eliminando ϕ entre estas ecuaciones resulta

$$p^2 + p'^2 - 2pp' \cos \theta = p_e^2 \quad (4.36)$$

La conservación de la energía relativística total requiere

$$E + m_e c^2 = E' + m_e c^2 + K \quad (4.37)$$

donde K es la energía cinética del electrón después de la colisión. Usando la (4.33), la condición de conservación de la energía se escribe como

$$c(p - p') = K \quad (4.38)$$

Usando la (4.32) podemos escribir K en función de p_e :

$$(m_e c^2 + K)^2 = c^2 p_e^2 + m_e^2 c^4 \quad (4.39)$$

que se reduce a

$$2Km_e + K^2 / c^2 = p_e^2 \quad (4.40)$$

Usando ahora la (4.38), la condición de conservación de la energía se escribe como

$$2m_e c(p - p') + (p - p')^2 = p_e^2 \quad (4.41)$$

Finalmente podemos eliminar p_e entre la (4.36) y la (4.41) para obtener la relación entre p y p' :

$$m_e c(p - p') = pp'(1 - \cos \theta) \quad (4.42)$$

que podemos escribir en la forma

$$\frac{1}{p'} - \frac{1}{p} = \frac{1}{m_e c} (1 - \cos \theta) \quad (4.43)$$

Recordando la (4.33), la (4.43) se puede llevar a la forma

$$\lambda' - \lambda = \frac{h}{m_e c} (1 - \cos \theta) \quad (4.44)$$

La (4.44), que expresa entonces la conservación de la energía y la cantidad de movimiento en la colisión es la *ecuación de Compton*. La cantidad

$$\lambda_C = \frac{h}{m_e c} = 2.42631058 \times 10^{-10} \text{ cm} \quad (4.45)$$

se denomina *longitud de onda Compton* del electrón. En términos de λ_C el corrimiento Compton se expresa como

$$\Delta\lambda = \lambda_C(1 - \cos\theta) \quad (4.46)$$

Se puede observar que el corrimiento Compton depende solamente del ángulo de dispersión θ , y es independiente de la longitud de onda del fotón incidente.

La ecuación de Compton describe correctamente los corrimientos observados. De acuerdo con la (4.46) $\Delta\lambda$ varía desde 0 para $\theta = 0$ (que corresponde a una colisión rasante en la que el fotón incidente no sufre desviación) hasta $2\lambda_C$ para $\theta = 180^\circ$ (colisión frontal, en la cual el fotón regresa en la misma dirección desde la cual vino). En experimentos posteriores Compton y A. W. Simon (1925) detectaron también el electrón de retroceso que resulta de la colisión y comprobaron que aparece simultáneamente con la radiación X dispersada. También verificaron la predicción sobre la energía y dirección de movimiento del electrón dispersado. Con esto quedaron reivindicadas las ideas de Einstein y ya nadie dudó de la realidad de los fotones.

Falta todavía explicar porqué para ángulos $\theta \neq 0$ se observan fotones dispersados con la misma longitud de onda de la radiación incidente. Hasta aquí hemos supuesto que el electrón que sufre la colisión está libre. Esta suposición es razonable, incluso si el electrón está inicialmente ligado, siempre y cuando la energía cinética que adquiere en la colisión sea mucho mayor que su energía de ligadura. Sin embargo, si el electrón está *fuertemente* ligado, o si la energía del fotón incidente es muy *pequeña*, el electrón no es expulsado del átomo. Entonces la colisión ocurre, en la práctica, con el átomo como un todo y podemos repetir las consideraciones anteriores acerca de la conservación de la energía y la cantidad de movimiento, pero en lugar de la masa m_e del electrón, tenemos que poner la masa m_a del átomo, pues éste retrocede como un todo durante la colisión. Puesto que m_a es varios miles de veces mayor que m_e , el corrimiento Compton para colisiones con electrones fuertemente ligados es insignificante e imposible de detectar.

En conclusión: en el experimento se observan los fotones dispersados por los electrones débilmente ligados, que son expulsados del átomo por la colisión, y esos fotones cambian su longitud de onda, pero también se observan fotones dispersados por otros electrones fuertemente ligados, que siguen ligados después de la colisión; estos fotones no cambian su longitud de onda.

El proceso de dispersión sin cambio de la longitud de onda no es otra cosa que la *dispersión de Thomson*, de la cual hemos ya hablamos en el Capítulo 3. La correspondiente teoría fue desarrollada por Thomson en 1900 en a partir del Electromagnetismo clásico. Pese a que la deducción de la dispersión Thomson es diferente de la presente interpretación cuántica de la dispersión Compton, ambas explican los mismos hechos, pues la dispersión Thomson es el límite de la dispersión Compton, cuando la longitud de onda de la radiación incidente es muy larga, esto es, cuando $\lambda = \lambda_C / \lambda \rightarrow 0$. En ese límite los resultados clásico y cuántico coinciden.

La física clásica falla al explicar la dispersión de radiación en el dominio de las frecuencias muy altas, es decir longitudes de onda muy cortas. Ocurre aquí algo semejante a la “catástrofe ultravioleta”, donde las predicciones clásicas fallan groseramente al intentar describir la radiación de cuerpo negro. Esto se relaciona con el valor de la constante de Planck. Para longitudes de onda muy largas, ν es muy pequeña y como h también es muy pequeña, la “granulosidad” $h\nu$ de la energía electromagnética es imperceptible y se confunde con el continuo de la física clásica. Pero para longitudes de onda suficientemente cortas $h\nu$ ya no es pequeño y los efectos cuánticos son importantes.

Para completitud citamos la *fórmula de Klein-Nishina* que expresa la sección eficaz de la dispersión Compton⁹:

$$\sigma_C = \sigma_T f(\Lambda) \quad , \quad \Lambda = \frac{h\nu}{m_e c^2} = \frac{\lambda_C}{\lambda} \quad (4.47)$$

donde $\sigma_T = 8\pi r_0^2 / 3$ es la sección eficaz clásica de Thomson; los efectos cuánticos están contenidos en la función $f(\Lambda)$, que tiene la expresión complicada que damos a continuación:

$$f(\Lambda) = \frac{3(1+\Lambda)}{4\Lambda^2} \left(\frac{2+2\Lambda}{1+2\Lambda} - \frac{\ln(1+2\Lambda)}{\Lambda} \right) - \frac{3}{4} \left(\frac{1+3\Lambda}{(1+2\Lambda)^2} - \frac{\ln(1+2\Lambda)}{2\Lambda} \right) \quad (4.48)$$

Obsérvese que Λ es igual al cociente entre la energía del fotón y la energía en reposo del electrón. Para Λ pequeños, se ve fácilmente que

$$f(\Lambda) = 1 - 2\Lambda + \frac{26}{5}\Lambda^2 + \dots \quad (4.49)$$

En el límite $\Lambda \rightarrow 0$ se tiene $f(\Lambda) = 1$ y la (4.47) se reduce a la fórmula clásica de Thomson. Por último, podemos ver que el balance de la cantidad de movimiento (4.36)

$$p^2 + p'^2 - 2pp' \cos \theta = p_e^2 \quad (4.50)$$

y de la energía

$$2m_e c(p - p') + (p - p')^2 = p_e^2 \quad (4.51)$$

no se pueden satisfacer simultáneamente si el fotón es absorbido por el electrón libre. En efecto, si el fotón fuera absorbido se tendría $p' = 0$, y la (4.50) y (4.51) se reducen, respectivamente a

$$p^2 = p_e^2 \quad \text{y} \quad 2m_e c p + p^2 = p_e^2 \quad (4.52)$$

que no se pueden satisfacer simultáneamente. Por lo tanto *un electrón libre no puede absorber un fotón*. Del mismo modo se puede ver que *un electrón libre no puede emitir un fotón*.

La emisión de rayos X

Los rayos X se producen en un tubo de descarga cuando electrones de alta energía acelerados por una diferencia de potencial de $\approx 10^4$ V, se frenan al penetrar en el material del ánodo¹⁰. Según la física clásica, el frenamiento de los electrones resulta en la emisión de un *espectro continuo* de radiación electromagnética, que se denomina *radiación de frenamiento* o también *bremsstrahlung* (que significa lo mismo en idioma alemán). La observación muestra que, además del espectro continuo de bremsstrahlung, también se emiten *líneas características* del material del ánodo, pero por el momento no nos ocuparemos de ellas.

La característica más notable de la emisión de rayos X por estos dispositivos es que para una dada energía del electrón hay una longitud de onda mínima λ_m bien definida. Por ejemplo para

⁹ Obtenida por Oskar Klein y Yoshio Nishina en 1928 a partir de la teoría relativística de Dirac.

¹⁰ Este fenómeno se llama *Efecto Volta*.

electrones de 40 keV, $\lambda_m = 0.311 \text{ \AA}$. Mientras la forma de la distribución espectral depende tanto del material del blanco como del potencial V de aceleración, λ_m depende sólo de V y tiene el mismo valor para todos los materiales. El Electromagnetismo clásico no explica este hecho, porque no hay motivo para que no se puedan emitir radiaciones de cualquier longitud de onda.

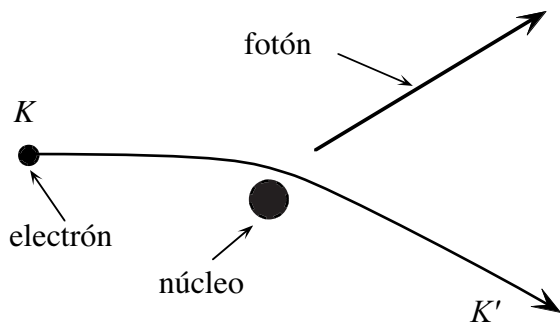


Fig. 4.5. Emisión de rayos X por bremsstrahlung.

Sin embargo si se considera que la radiación X consiste de fotones la explicación es muy simple. La Fig. 4.5 muestra un esquema del proceso elemental de emisión de un fotón por bremsstrahlung. Un electrón de energía cinética inicial K se frena al chocar con un núcleo y la energía que pierde aparece como un fotón de rayos X. La interacción Coulombiana del electrón con el núcleo hace que éste absorba cierta cantidad de movimiento, y el frenamiento del electrón produce la emisión del fotón. Como la masa del núcleo es mucho mayor que la del electrón, se puede ignorar la energía que adquiere en la colisión.

Si K' es la energía cinética del electrón después de la colisión, la energía del fotón emitido es

$$h\nu = \frac{hc}{\lambda} = K - K' \quad (4.53)$$

Los electrones del haz incidente pueden perder diferentes cantidades de energía en las colisiones, y un dado electrón queda en reposo después de cierto número de ellas. Es así que muchos electrones producen en conjunto un espectro continuo que se extiende desde λ_m hasta $\lambda \rightarrow \infty$, de acuerdo con las pérdidas de energía de las colisiones individuales. Un fotón de longitud de onda λ_m se emite cuando el electrón pierde toda su energía en un único choque. Cuando esto ocurre, $K' = 0$, y entonces $K = hc / \lambda_m$. Puesto que K es igual a la energía eV que adquirió el electrón al acelerarse en la diferencia de potencial V del tubo de rayos X, se tiene

$$\lambda_m = \frac{hc}{eV} \quad (4.54)$$

Por lo tanto λ_m corresponde a la conversión de toda la energía cinética del electrón en un único fotón. La (4.54) muestra que la existencia de λ_m es un fenómeno puramente cuántico¹¹.

El bremsstrahlung es, como vemos, una suerte de “efecto fotoeléctrico inverso”: en el efecto fotoeléctrico se absorbe un fotón cuya energía y cantidad de movimiento van al electrón y al núcleo, en el bremsstrahlung se crea un fotón cuya energía y cantidad de movimiento provienen de la colisión de un electrón con un núcleo¹². El bremsstrahlung no sólo ocurre en los tubos de rayos X, sino que tiene lugar también en muchas situaciones en la naturaleza y en el laboratorio.

¹¹ La existencia de λ_m fue predicha por Einstein en 1906, y la evidencia experimental fue obtenida por William Duane y Franklin Hunt en 1915. Por ese motivo λ_m se denomina *límite de Duane-Hunt*.

¹² En 1909 Johannes Stark aplicó la hipótesis del cuanto de luz al bremsstrahlung y concluyó que la conservación de la cantidad de movimiento requiere que el cuanto de luz posea un impulso igual a $h\nu/c$. Esta fue la primera vez que se incluyó explícitamente el fotón en el balance de la cantidad de movimiento de un proceso elemental.

Creación y aniquilación de pares

Además del efecto fotoeléctrico y el efecto Compton, hay otro proceso mediante el cual un fotón de alta energía puede perder su energía al interactuar con la materia. Se trata de la *creación de pares*, que es un ejemplo de conversión de la energía radiante en energía de masa en reposo y energía cinética de partículas. Aquí vamos a comentar brevemente este proceso y su proceso inverso, la *aniquilación de pares*, pero no daremos sus tratamientos completos ya que requieren la teoría cuántica relativística de campos.

Creación de pares

Hay hoy una amplísima evidencia experimental acerca de la creación de pares, pero no existe explicación de este fenómeno en la física clásica. En la Fig. 4.6 se muestra un esquema de un proceso de este tipo, que consiste en que un fotón pierde toda su energía $h\nu$ en una colisión con un núcleo, creando un *par* de partículas integrado por un *electrón* y un *positrón* y dándoles además cierta energía cinética. El positrón es una partícula idéntica en todas sus propiedades al

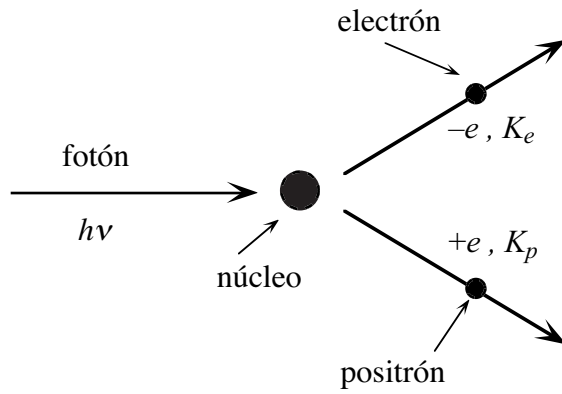


Fig. 4.6. Producción de un par electrón-positrón.

electrón, excepto que el signo de su carga (y de su momento magnético) es opuesto al del electrón: esto es, el positrón es un electrón con carga positiva en lugar de negativa. El positrón es lo que se denomina la *antipartícula* del electrón. La teoría cuántica de campos predice que *cada partícula tiene una correspondiente antipartícula*, y esta predicción se ha verificado siempre. En general, cuando se habla de “par” se entiende un par formado por una partícula y su correspondiente antipartícula. En la creación de pares electrón-positrón la energía involucrada en el retroceso del núcleo es despreciable debido a

su masa, y por lo tanto el balance de la energía (relativística) es

$$h\nu = E_e + E_p = (m_e c^2 + K_e) + (m_p c^2 + K_p) = 2m_e c^2 + K_e + K_p \quad (4.55)$$

puesto que $m_p = m_e$. En esta fórmula los subíndices e y p se refieren al electrón y al positrón, respectivamente. La energía cinética del positrón es ligeramente mayor que la del electrón debido a que la carga positiva del núcleo acelera al positrón y frena al electrón. La carga eléctrica total se conserva, pues el fotón no tiene carga y la carga neta del par es nula.

Haciendo los balances de energía y cantidad de movimiento, es fácil verificar que un fotón no puede desaparecer en el vacío creando un par: es necesaria la presencia del núcleo (que puede absorber cantidad de movimiento sin afectar sensiblemente el balance de energía) para permitir la conservación simultánea de la cantidad de movimiento y la energía.

De la (4.55) se ve que hay una energía mínima, o *umbral*, que debe tener el fotón para crear un par:

$$h\nu_{\min} = 2m_e c^2 \approx 1.02 \text{ MeV} \quad (4.56)$$

que corresponde a una longitud de onda $\lambda = \lambda_C / 2 = 0.012 \text{ \AA}$. Si la longitud de onda del fotón es menor que este valor, el fotón crea el par con energía cinética no nula.

Como se ve, la creación de pares es un fenómeno de alta energía, que ocurre con rayos X cortos o rayos γ . Se observa en la naturaleza debido a los fotones de alta energía de los rayos cósmicos y a los rayos γ emitidos por sustancias radioactivas, y en el laboratorio debido a los fotones de bremsstrahlung producidos en los aceleradores de partículas.

Si el fotón tiene suficiente energía es posible crear otros pares partícula-antipartícula, por ejemplo pares protón-antiprotón, y otros más. Como la masa en reposo del electrón es la menor entre las partículas posibles, el umbral para la creación de pares electrón-positrón es el más bajo.

Aniquilación de pares

El proceso inverso a la creación de pares es la aniquilación de pares. Si una partícula y su antipartícula, por ejemplo un electrón y un positrón se encuentran inicialmente en reposo y próximos entre sí, se unen y se aniquilan mutuamente. En este proceso *desaparece* la materia y en su lugar aparece energía radiante. Puesto que la cantidad de movimiento inicial del conjunto es nula, y el proceso debe conservar la cantidad de movimiento, no es posible crear un único fotón. El proceso más probable es la creación de dos fotones, con cantidades de movimiento iguales pero de signo opuesto. Pero también es posible, aunque mucho menos probable, la aniquilación con creación de tres (o más) fotones.

Consideremos la aniquilación con creación de dos fotones. La conservación de la cantidad de movimiento requiere $p_1 = p_2$. Como para un fotón $p = h\nu/c$, lo anterior implica $\nu_1 = \nu_2 = \nu$. Puesto que las energías cinéticas de las partículas del par son nulas, la conservación de la energía requiere

$$h\nu = m_e c^2 = 0.511 \text{ MeV} \quad (4.57)$$

Esto corresponde a una longitud de onda igual a

$$\lambda = \frac{h}{m_e c} = \lambda_C \quad (4.58)$$

es decir, igual a la longitud de onda Compton.

Los positrones que se crean en el proceso de producción de pares pierden energía al atravesar la materia en sucesivas colisiones, hasta que se combinan con un electrón para formar un sistema ligado denominado *positronio*, en el cual el electrón y el positrón se mueven alrededor de su centro de masa común. El positronio tiene una vida muy corta ya que en un tiempo de alrededor de 10^{-10} s desde su formación, el electrón y el positrón se aniquilan.

Comentarios

La primera evidencia experimental de la existencia del positrón fue obtenida por Carl David Anderson en 1932 al investigar la radiación cósmica¹³. En 1933 Patrick M. Blackett y Giuseppe Occhialini observaron por la primera vez la creación de pares electrón-positrón, utilizando una cámara de niebla. Este descubrimiento permitió explicar el origen de una discrepancia entre las

¹³ La radiación cósmica consiste en un flujo de partículas cargadas y fotones de gran energía que inciden sobre la Tierra provenientes de fuentes extraterrestres.

mediciones de la atenuación de rayos X de gran energía por diversos materiales y la teoría entonces existente. La atenuación predicha por la teoría era muy pequeña en comparación con la observada, pues la creación de pares es el mecanismo de absorción que domina para fotones de alta energía. Pero la mayor importancia de estos descubrimientos fue que confirmó la teoría cuántica relativística de Dirac, que en 1928 predijo la existencia del positrón y los procesos de creación y aniquilación de pares.

La naturaleza dual de la radiación electromagnética

Acabamos de estudiar varios fenómenos que muestran que la radiación electromagnética se comporta como un conjunto de partículas, los fotones, que intervienen individualmente en los procesos elementales de emisión, absorción y dispersión. Sin embargo, los fenómenos de interferencia y difracción muestran que la radiación es un fenómeno ondulatorio. Por lo tanto la radiación electromagnética se comporta como *onda* en ciertas circunstancias y como *corpúsculo* en otras. Veremos pronto que las partículas atómicas, como el electrón y el protón, también exhiben la misma clase de dualidad, pues además de comportarse como partículas, en ciertas circunstancias se comportan como ondas, pudiendo dar lugar a fenómenos de interferencia y difracción. Esta *dualidad onda-partícula* es característica de todos los entes de escala atómica o menor, y no es compatible con nuestra experiencia a nivel macroscópico ni con la descripción dada por la física clásica. Sin embargo, veremos oportunamente que la Mecánica Cuántica permite reconciliar la coexistencia de los aspectos corpusculares y ondulatorios de estos entes.

Es oportuno subrayar en este contexto que los conceptos de “onda” y de “partícula” son *extrapolaciones* de experiencias a nivel macroscópico. Así, en la Mecánica Newtoniana se trata un planeta como la Tierra como una partícula. Creemos en ese concepto porque la Mecánica Newtoniana (o la relativística, si la velocidad es muy elevada) permite calcular correctamente el movimiento. Aplicando los mismos métodos podemos describir fenómenos de escala más pequeña, como el movimiento de partículas de polvo o de las gotas de un aerosol. Pero este no es un motivo suficiente para dar por sentado que el concepto *macroscópico* de partícula se puede extrapolar indefinidamente, hasta la escala del átomo y más allá de ella. Por supuesto, se pueden hacer extrapolaciones, pero *sólo hasta que se encuentra que ya no funcionan*.

Con el concepto de onda sucede lo mismo. Podemos ver con nuestros ojos las ondas en la superficie del agua de un estanque. Pero aún del punto de vista clásico, sabemos que no podemos extrapolar el concepto indefinidamente, pues al llegar a las dimensiones moleculares la noción misma de “superficie del agua” pierde sentido. En esa escala encontraremos un gran número de moléculas que se mueven aparentemente al azar, y sólo después de promediar el comportamiento de grandes grupos de moléculas podemos recuperar los conceptos de superficie y de onda. Estos comentarios muestran que el problema surge cuando intentamos extrapolar conceptos macroscópicos como “superficie” y “onda”, derivados de nuestra experiencia cotidiana, hasta dominios donde carecen de validez.

Lo que indican claramente los fenómenos que hemos estudiado en este Capítulo y otros que veremos más adelante, es que *los conceptos clásicos de partícula y onda no se pueden extrapolar a la escala atómica*. En esa escala no es lícito establecer una distinción entre “partícula” y “onda” en el sentido clásico. Por otra parte, esa distinción está implícita en el planteo tradicional de la Mecánica, por la forma misma con la cual se definen las variables dinámicas del sistema. Por eso, como veremos, en la Mecánica Cuántica se parte de un planteo radicalmente diferente.

Yendo ahora específicamente al caso de las ondas electromagnéticas, podemos decir lo siguiente. No podemos ver las ondas electromagnéticas del mismo modo que las olas en la superficie del agua. Pero tenemos varias razones para creer que son también un fenómeno ondulatorio. En efecto:

- podemos observar los fenómenos de interferencia y difracción, que se asocian con los fenómenos ondulatorios;
- la distribución de energía en el espacio y el tiempo se predice correctamente por medio de la teoría de Maxwell para todas las longitudes de onda, desde prácticamente infinito hasta alrededor de unos $0.02 \text{ \AA}$;
- en el caso de las ondas de radio podemos medir, además de la longitud de onda, también la amplitud y la fase.

Veamos estos puntos con más detalle, comenzando por el último. Hay un límite a la intensidad más pequeña que se puede medir, dado por la agitación térmica de las moléculas en la antena de detección, que corresponde a un flujo de alrededor de 10^{10} fotones/s. Por lo tanto, el último punto nos dice que *muchos fotones* que actúan en conjunto sobre los electrones de una antena tienen la apariencia de un campo electromagnético clásico.

Por otra parte, es fácil demostrar que en un interferómetro óptico ordinario se puede trabajar con intensidades que corresponden a tener en un dado instante *un único fotón* dentro del instrumento. Y en esas condiciones *se observa interferencia*. Un caso extremo fue estudiado por G. I. Taylor en 1909, quien demostró que se obtiene el patrón de interferencia habitual, usando una fuente luminosa tan débil que la fotografía demoró *tres meses* en registrarse.

Podemos analizar el significado de lo anterior si imaginamos sustituir la placa fotográfica por varios fotoelectrodos muy pequeños, de modo que detectando el electrón emitido podemos determinar en cuál de ellos incidió el fotón. La energía luminosa total captada por cualquier detector es proporcional al número de fotones que llegaron al mismo. Si se efectúa ese experimento, se encuentra que a medida que transcurre el tiempo, el número de fotones que se registra en cada detector tiende al valor de la intensidad que predice la teoría ondulatoria. Pero esto es cierto *sólo en promedio*, y si el número total de fotones es pequeño pueden ocurrir grandes fluctuaciones. Por otra parte, cada fotón individual llega a un sólo detector.

La figura de interferencia dada por dos rendijas depende de que la luz pase por *ambas* (esto se comprueba tapando una de ellas). El hecho que cuando la luz atraviesa ambas rendijas se produce interferencia demostró que es un fenómeno ondulatorio y por ese motivo se descartó la teoría corpuscular de Newton. Significa esto que los fotones se pueden dividir en dos? Recordemos que por nuestro aparato pasa un sólo fotón por vez. Podemos salir de dudas, poniendo un detector detrás de cada rendija. Si hacemos eso no veremos más la figura de interferencia, pero podremos determinar si el fotón pasa por una sola rendija o por ambas. Lo que resulta es que el fotón pasa o por una rendija, o por la otra, y no por ambas a la vez. Por lo tanto, el fotón no se divide. Sin embargo, si el fotón no pasa por ambas rendijas, ¿cómo hace para interferir consigo mismo?

Una pista para resolver este dilema consiste en pensar que la onda electromagnética nos dice algo, no acerca de *dónde está exactamente* el fotón, sino acerca de *la probabilidad de encontrarlo en determinado sitio*. Si suponemos que esa probabilidad está relacionada con la intensidad de la onda, se resuelve el problema. En el Capítulo 15 se muestra como el comportamiento dual onda-corpúsculo de la radiación queda incorporado al formalismo de la Mecánica Cuántica.

5. LA TEORÍA CUÁNTICA ANTIGUA

Introducción

El intento de resolver el problema de la inestabilidad del átomo de Rutherford llevó a Niels Bohr a formular en 1913 una teoría simple de la estructura atómica, uno de cuyos mayores méritos fue que permitió explicar el espectro de la radiación electromagnética emitida por ciertos átomos. Dicha teoría fue luego perfeccionada por William Wilson, Jun Ishiwara, Planck, Arnold Sommerfeld y otros, y dio lugar a lo que hoy llamamos la Teoría Cuántica Antigua. Si bien esta teoría fue luego abandonada, cumplió en su momento un rol importante para el desarrollo de la Mecánica Cuántica moderna. Por este motivo daremos aquí una reseña de la misma. Pero antes es necesario mencionar brevemente algunos aspectos sencillos del espectro atómico.

El espectro atómico

En contraste con el espectro continuo de la radiación térmica, la radiación electromagnética emitida por un átomo libre consiste de un conjunto discreto de longitudes de onda. Cada una de estas longitudes de onda recibe el nombre de *línea*, pues así es como aparece en las placas fotográficas obtenidas con los espectrógrafos. Cada especie atómica tiene su propio espectro, integrado por un conjunto de líneas características. Este hecho tiene gran importancia práctica, pues permite identificar los elementos presentes en una fuente de luz. Por esta razón durante el siglo XIX se dedicó mucho esfuerzo a medir con precisión los espectros de los diferentes átomos. Tales espectros son muy complicados y generalmente constan de centenares de líneas.

El más sencillo de todos los espectros es el del átomo de hidrógeno, lo que no es sorprendente puesto que se trata del átomo más simple ya que tiene un solo electrón. Por este motivo, y también por razones históricas y teóricas, presenta mucho interés.

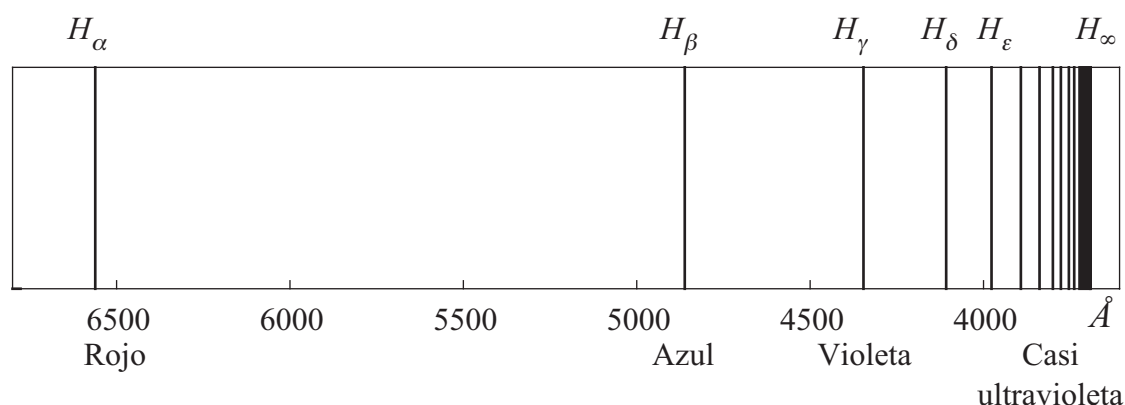


Fig. 5.1. Parte visible del espectro del hidrógeno.

En el visible (Fig. 5.1), el espectro del hidrógeno presenta una serie regular de líneas que comienza en el rojo y termina en el violeta; el espaciado entre líneas contiguas decrece paulatinamente hasta que se llega al límite de la serie, que se encuentra para 3645.6 Å. La regularidad de este espectro llevó a muchos a buscar fórmulas empíricas que representasen las longitudes de onda de las líneas. La fórmula correcta fue hallada en 1885 por Johann Jakob Balmer (un profesor de escuela secundaria) y es, expresada en términos del número de onda $k \equiv 2\pi/\lambda$,

$$\frac{1}{\lambda_{2,n}} = \frac{k_{2,n}}{2\pi} = R_H \left(\frac{1}{2^2} - \frac{1}{n^2} \right) \quad , \quad n = 3, 4, 5, \dots \quad (5.1)$$

donde

$$R_H = 109677.576 \pm 0.012 \text{ cm}^{-1} \quad (5.2)$$

se denomina *constante de Rydberg* para el átomo de hidrógeno. Posteriormente, gracias al trabajo de varios espectroscopistas entre quienes tuvo un rol preponderante Johannes Rydberg (1890), se encontraron fórmulas semejantes para diferentes series, del tipo

$$\frac{k_{m,n}}{2\pi} = R_H \left(\frac{1}{m^2} - \frac{1}{n^2} \right) \quad , \quad m = 1, 2, 3, \dots \quad , \quad n = n+1, n+2, n+3, \dots \quad (5.3)$$

La *serie de Balmer* (5.1) corresponde a $m = 2$; la serie correspondiente a $m = 1$ se encuentra en el ultravioleta y se denomina *serie de Lyman*. Las correspondientes a $m = 3, 4$ y 5 se encuentran en el infrarrojo y se llaman *serie de Paschen*, *Brackett* y *Pfund*, respectivamente.

Las fórmulas para las series de los átomos alcalinos (sodio, potasio, rubidio, cesio, ...) tienen la misma estructura general:

$$\frac{k_{m,n}}{2\pi} = R \left(\frac{1}{(m - a_m)^2} - \frac{1}{(n - b_m)^2} \right) \quad (5.4)$$

donde R es la constante de Rydberg para el elemento, m es el entero que identifica la serie, a_m y b_m son constantes para la serie, y n es el entero variable. El valor de R es el mismo para todos los elementos dentro de un 0.05% y aumenta levemente al aumentar el peso atómico.

Ligado con el espectro de emisión que estuvimos comentando, está el espectro de absorción, que se obtiene cuando se emplea una fuente luminosa que emite un espectro continuo y se interpone entre la fuente y el espectrógrafo un recipiente que contiene un vapor atómico. En este caso se encuentra que la placa está velada, salvo en algunas líneas que corresponden a las longitudes de onda que son absorbidas por los átomos del vapor. A cada línea del espectro de absorción le corresponde una línea del espectro de emisión. Pero la inversa no es cierta: no todas las líneas de emisión se observan en el espectro de absorción. Por ejemplo, en el espectro de absorción del hidrógeno sólo aparecen normalmente las líneas de la serie de Lyman. Pero si el gas está a una temperatura muy elevada, como ocurre en algunas estrellas, en el espectro de absorción también se ven las líneas de la serie de Balmer.

Los postulados de Bohr

Toda teoría de la estructura atómica debe explicar estas características del espectro, así como muchas otras que no hemos comentado. La gran precisión de las medidas espectroscópicas impone además severos requerimientos sobre las predicciones teóricas.

Por otra parte, el espectro del hidrógeno es completamente inexplicable en términos clásicos, del mismo modo que son inexplicables otros aspectos del átomo tales como su tamaño, la existencia del núcleo y su estabilidad, como hemos visto en el Capítulo 3. El gran mérito de Bohr reside en que reconoció la necesidad de abandonar la Física Clásica, y en consecuencia tuvo la audacia de proponer que varias leyes de la Mecánica y del Electromagnetismo no se cumplen en la escala atómica. De esta forma consiguió dar un paso de fundamental importancia, que indicó la dirección en que había que avanzar para superar el punto muerto al cual había llegado la Física Teórica.

La teoría de Bohr tiene gran sencillez matemática y concuerda cuantitativamente con los datos espectroscópicos del hidrógeno. Sin embargo, los postulados sobre los cuales se basa parecen artificiosos. Son los siguientes:

Postulados de Bohr:

- El electrón se mueve en una órbita circular alrededor del núcleo bajo la influencia de la atracción Coulombiana de éste, obedeciendo las leyes de la mecánica clásica.
- Dentro de las infinitas órbitas clásicas, el electrón se mueve sólo en aquellas en las que el momento angular orbital L tiene los valores $L = n\hbar = nh/2\pi$, donde $n = 1, 2, 3, \dots$
- Cuando el electrón se mueve en una órbita permitida, no irradia energía electromagnética a pesar de ser acelerado constantemente y por lo tanto su energía total E permanece constante.
- Un electrón que se mueve inicialmente en una órbita de energía E_i puede cambiar discontinuamente su movimiento y pasar a moverse en otra órbita de energía E_f ; cuando esto ocurre se emite un fotón cuya frecuencia es $\nu = (E_i - E_f)/h$.

Se debe notar la diferencia entre la cuantificación de Bohr del momento angular y la cuantificación de Planck, que se refiere a la energía total de una partícula (por ejemplo un electrón que realiza oscilaciones armónicas simples). Veremos que la cuantificación del momento angular implica también la cuantificación de la energía total, pero de una forma diferente a la de Planck. El tercer postulado resuelve *manu militari* el problema de la estabilidad debido a la radiación del electrón acelerado, mediante el simple expediente de postular que esta característica de la teoría clásica *no vale* para el electrón cuando se mueve en una órbita permitida. Este postulado se basa en la observación experimental de que los átomos son estables, a pesar que la teoría clásica predice lo contrario. El cuarto postulado está ligado al postulado de Einstein sobre la energía de un fotón (ec. (4.41)).

Se puede observar que los postulados de Bohr mezclan de manera arbitraria la física clásica con la no clásica, y en tal sentido son intelectualmente insatisfactorios. Por ejemplo, se supone que el electrón se mueve según las leyes de la mecánica clásica, y al mismo tiempo se afirma que su momento angular está cuantificado; se supone que el electrón obedece la ley de Coulomb, pero acto seguido se lo exime de cumplir la regla que toda carga acelerada irradia.

Sin embargo, se puede argumentar que no nos debemos sorprender si las leyes de la física clásica (basadas en nuestra experiencia con sistemas macroscópicos) no son completamente válidas cuando tratamos con un sistema microscópico como el átomo. En última instancia, la justificación de los postulados de Bohr (y de cualquier postulado, en realidad) reside en si describen correctamente los resultados experimentales, o no.

Teoría de Bohr del átomo con un electrón

Consideremos un átomo formado por un núcleo de masa M y carga $+Ze$, y un electrón de masa m_e y carga $-e$ (por ej. un átomo de hidrógeno, un átomo de helio ionizado una vez, uno de litio doblemente ionizado, etc.) que gira alrededor del núcleo en una órbita circular de radio r con la velocidad v . Por ahora, supongamos que el núcleo se puede considerar fijo (o sea, $M = \infty$). Se cumple entonces que la fuerza de Coulomb debe ser igual a la fuerza centrípeta:

$$\frac{Ze^2}{r^2} = \frac{m_e v^2}{r} \quad (5.5)$$

Puesto que la fuerza de Coulomb es central, se conserva el momento angular

$$L = m_e v r = \text{cte.} \quad (5.6)$$

Reemplazando en (5.5) tenemos

$$Ze^2 = \frac{L^2}{m_e r} \quad (5.7)$$

Despejando r obtenemos

$$r = \frac{L^2}{m_e Ze^2} \quad (5.8)$$

Aplicando la regla de cuantificación

$$L = n\hbar \quad , \quad n = 1, 2, 3, \dots \quad (5.9)$$

donde hemos definido

$$\hbar \equiv h / 2\pi \quad (5.10)$$

encontramos que las órbitas permitidas tienen radios dados por

$$r_n = \frac{n^2 \hbar^2}{m_e Ze^2} \quad (5.11)$$

que son proporcionales al cuadrado del *número cuántico* n . La órbita más pequeña ($n = 1$) del átomo de hidrógeno ($Z = 1$) tiene un radio igual a

$$r_1 = a_0 \equiv \frac{\hbar^2}{m_e e^2} = 0.529 \times 10^{-8} \text{ cm} \quad (5.12)$$

que se denomina *radio de Bohr*. Veremos en seguida que esta órbita es la de menor energía.

La velocidad del electrón también está cuantificada, y su valor es

$$v_n = \frac{n\hbar}{m_e r_n} = \frac{Ze^2}{n\hbar} = \frac{Z}{n} \alpha c \quad (5.13)$$

donde hemos introducido la constante adimensional

$$\alpha \equiv \frac{e^2}{\hbar c} = 7.297 \times 10^{-3} \approx 1/137 \quad (5.14)$$

que se llama *constante de la estructura fina* por motivos que se verán en breve. Dicha constante es una medida de la fuerza de la interacción electromagnética y juega un rol muy importante en la electrodinámica cuántica.

Para la órbita más pequeña ($n = 1$) del átomo de hidrógeno tenemos

$$v_1 = \frac{e^2}{\hbar} = \alpha c = 2.2 \times 10^8 \text{ cm/s} \quad (5.15)$$

que es menos del 1% de la velocidad de la luz, por lo tanto está bien haber usado en nuestros cálculos la mecánica clásica no relativística.

Calculemos la energía total del electrón:

$$E_n = T_n + V_n = \frac{1}{2} m_e v_n^2 - \frac{Ze^2}{r_n} = -\frac{1}{2} m_e v_n^2 \quad (5.16)$$

Sustituyendo (5.11) y (5.13) en (5.16) obtenemos

$$E_n = T_n + V_n = -\frac{m_e}{2n^2} \left(\frac{Ze^2}{\hbar} \right)^2 = \frac{Z^2}{n^2} E_1, \quad n = 1, 2, 3, \dots \quad (5.17)$$

donde

$$E_1 = -\frac{1}{2} \frac{m_e e^4}{\hbar^2} = -\frac{1}{2} \frac{e^2}{a_0} = -\frac{1}{2} m_e c^2 \alpha^2 = -13.6 \text{ eV} \quad (5.18)$$

Por lo tanto la cuantificación del momento angular implica la cuantificación de la energía total. La información contenida en la (5.17) se suele presentar en un diagrama de niveles de energía (ver Fig. (5.2)). El estado estable, o sea el de energía mínima, corresponde a $n = 1$ y su energía es $E_1 = -13.6 \text{ eV}$ para el átomo de hidrógeno.

Podemos calcular la frecuencia $\nu_{m,n}$ del fotón emitido por el electrón al pasar del estado n al estado m ($m < n$). De acuerdo con el cuarto postulado de Bohr

$$\nu_{m,n} = \frac{E_n - E_m}{h} = \frac{m_e}{4\pi\hbar} \left(\frac{Ze^2}{\hbar} \right)^2 \left(\frac{1}{m^2} - \frac{1}{n^2} \right) \quad (5.19)$$

El correspondiente número de onda $k_{m,n} = 2\pi\nu_{m,n}/c$ está dado por

$$k_{m,n} = \frac{m_e Z^2 e^4}{2c\hbar^3} \left(\frac{1}{m^2} - \frac{1}{n^2} \right) = 2\pi R_\infty \left(\frac{1}{m^2} - \frac{1}{n^2} \right), \quad m < n \quad (5.20)$$

donde:

$$R_\infty = \frac{m_e Z^2 e^4}{4\pi c \hbar^3} = \frac{Z^2}{4\pi \tilde{\lambda}_C} \alpha^2 \quad (5.21)$$

y $\tilde{\lambda}_C \equiv \lambda_C / 2\pi$ (λ_C es la longitud de onda Compton del electrón, ec. (4.58)). Las ecuaciones (5.17) y (5.21) contienen las predicciones principales de la teoría de Bohr.

Veamos primero la emisión de radiación por un átomo de Bohr. Las ecuaciones mencionadas nos dicen lo siguiente:

- El estado normal del átomo es el de mínima energía o sea el estado con $n = 1$, que se suele denominar *estado fundamental* o *estado base*. Puesto que no hay otro estado con energía

menor, este estado es estable. El radio de la órbita correspondiente (ec. (5.12)) determina el tamaño del átomo (con un solo electrón), que resulta ser del orden de magnitud correcto.

- En determinadas circunstancias el átomo puede absorber energía (por causa de las colisiones, como ocurre en una descarga eléctrica, o por otro mecanismo), pasando a un estado de energía mayor, con $n > 1$.
- El átomo emite ese exceso de energía, obedeciendo la tendencia común de todos los sistemas físicos, y regresa al estado fundamental. Esto se consigue mediante una serie de transiciones en las que el electrón pasa sucesivamente a estados de energía más baja, hasta que finalmente llega al estado fundamental.
- En la gran variedad de procesos de excitación y desexcitación que ocurren en la fuente de luz cuyo espectro se está registrando se producen todas las transiciones posibles y por lo tanto se emite el espectro completo. A partir de la ec. (5.17) podemos calcular los números de onda de todas las líneas del espectro. Es fácil verificar que de esa manera se obtienen las fórmulas de las series de Lyman, Balmer, Paschen, etc. (Fig. 5.3). El valor de R_∞ concuerda con el valor experimental de la constante de Rydberg.

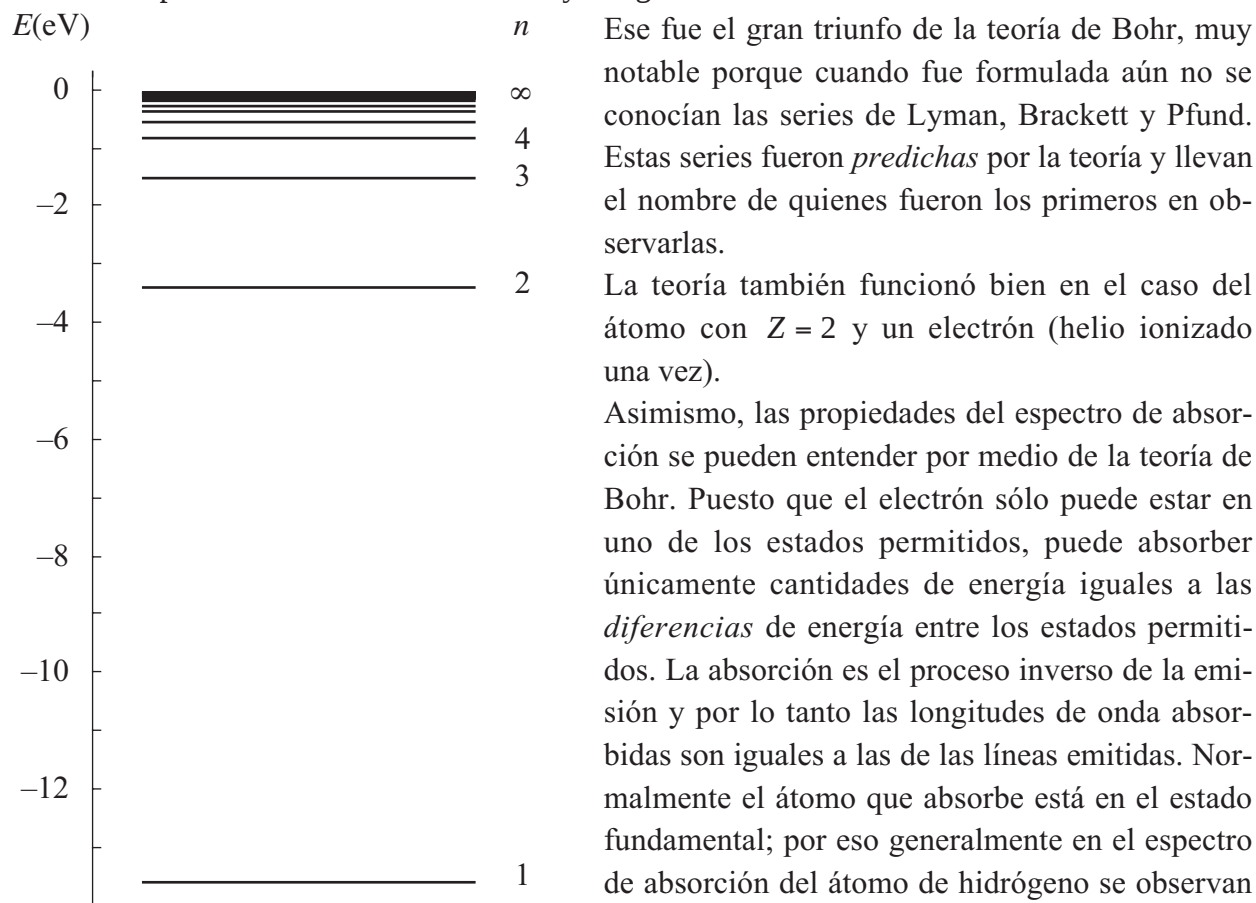


Fig. 5.2. Niveles de energía del átomo de hidrógeno de acuerdo con la teoría de Bohr.

tura necesaria se puede estimar a partir de la distribución de Boltzmann, y es fácil ver que debe ser del orden de los 10^5 °K.

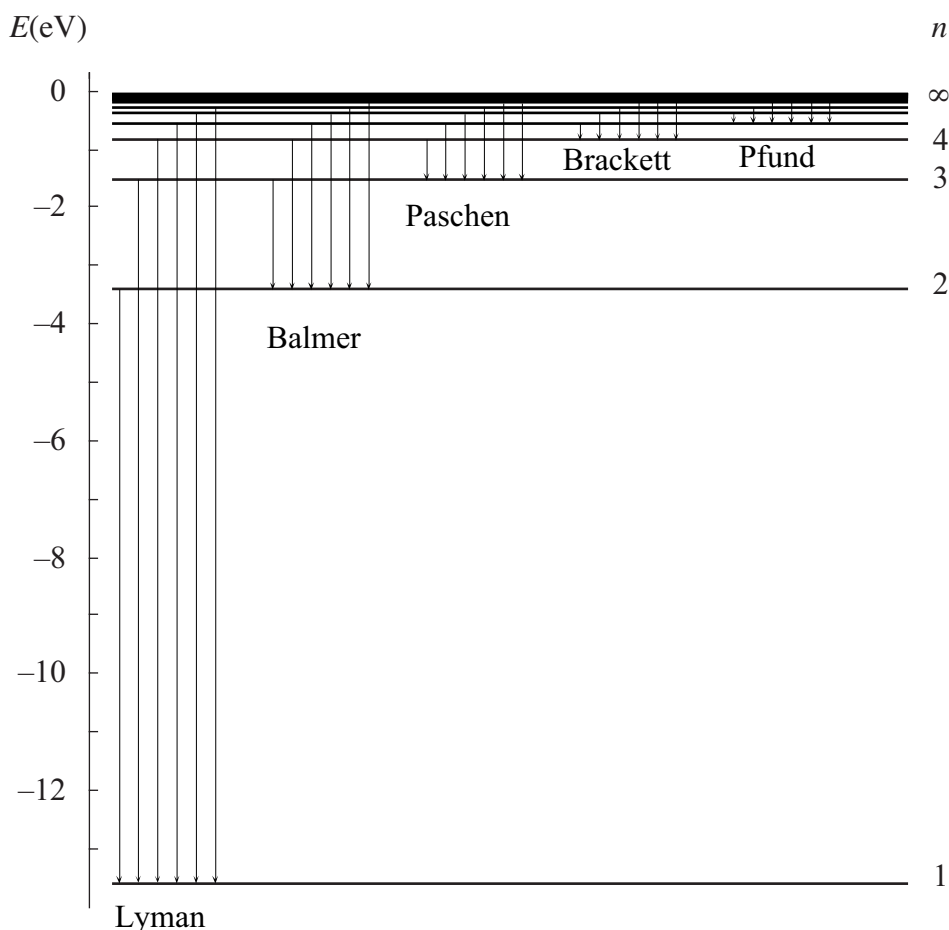


Fig. 5.3. Origen de las cinco series espectrales conocidas del átomo de hidrógeno de acuerdo con la teoría de Bohr.

El espectro de líneas de rayos X

Es oportuno mencionar aquí las investigaciones de Henry G. J. Moseley sobre las líneas espectrales de rayos X de elementos pesados, que influyeron sobre la teoría que estamos presentando ya que Bohr se mantuvo al corriente de las mismas. Moseley se basó en los trabajos de Barkla, y para medir las longitudes de onda de los rayos X utilizó las técnicas de difracción sobre cristales desarrolladas por Sir Lawrence Bragg y su hijo William. Comparando los rayos X emitidos por diferentes elementos, encontró que tienen frecuencias características que varían de acuerdo con un patrón regular. Sin embargo, la diferencia de frecuencia no está gobernada por la diferencia de masa entre los elementos, sino por la diferencia entre las *cargas eléctricas* de sus núcleos, es decir, las diferencias entre los respectivos números atómicos¹.

Sin entrar en mayores detalles, podemos decir que el espectro de línea de rayos X se debe a transiciones de los electrones más internos de un átomo. Para ellos, la carga nuclear no está apantallada por los demás electrones (que se mueven en órbitas de radio mayor), por lo tanto se puede aplicar la fórmula (5.17) que obtuvimos en la Sección precedente. En la terminología de rayos X, los niveles de energía correspondientes a $n = 1, 2, 3, \dots$ se denominan $K, L, M, \dots$

En su primer experimento de 1913, Moseley estudió los rayos X de la serie K (asociada con las transiciones al nivel de energía K) para los elementos hasta el Zn y el año siguiente extendió sus investigaciones hasta el Au, usando las líneas de la serie L (transiciones al nivel L). Las fórmulas

¹ Fue precisamente Moseley quien introdujo el término “número atómico” para designar a Z .

que obtuvo para las frecuencias se relacionan muy estrechamente con las fórmulas de Bohr para átomos con un solo electrón, y muestran que las frecuencias son proporcionales a Z^2 . Así, las series K y L son análogas a las series de Lyman y de Balmer del átomo de hidrógeno.

La regularidad de las diferencias en las frecuencias del espectro de rayos X permitió a Moseley ordenar por número atómico todos los elementos desde el Al hasta el Au. Pudo observar así que en algunos casos el orden dado por los pesos atómicos es incorrecto. Por ejemplo, el peso atómico del Co es mayor que el del Ni, pero sus números atómicos son 27 y 28, respectivamente. Cuando Mendeleyev construyó su Tabla Periódica se tuvo que basar en los pesos atómicos, pero colocó al Cu y el Ni en el orden inverso (es decir, el orden correcto de acuerdo con Z) para que sus propiedades químicas encajaran mejor.

En algunos lugares de la Tabla Periódica, Moseley encontró diferencias en Z mayores que la unidad, y predijo correctamente la existencia de elementos aún no descubiertos. Puesto que hay un único elemento para cada Z , los científicos pudieron finalmente confiar en la completitud de la Tabla Periódica.

Refinamientos del modelo de Bohr

Mencionaremos brevemente varias correcciones y perfeccionamientos del modelo de Bohr.

Corrección por masa nuclear finita

El hecho que el núcleo tiene una masa finita se puede tener en cuenta fácilmente si en todas las ecuaciones de la Sección anterior en lugar de la masa del electrón se emplea la masa reducida del sistema electrón-núcleo, dada por

$$\mu = \frac{m_e M}{m_e + M} \quad (5.22)$$

La fórmula para los números de onda se escribe ahora

$$k_{m,n} = 2\pi R_M Z^2 \left(\frac{1}{m^2} - \frac{1}{n^2} \right) \quad , \quad m < n \quad (5.23)$$

donde:

$$R_M = \frac{\mu Z^2 e^4}{4\pi c \hbar^3} \quad (5.24)$$

Con la corrección por el efecto de masa finita, la teoría de Bohr concuerda con los datos espectroscópicos dentro de un 0.003%.

La teoría de Wilson-Sommerfeld

El acierto de la teoría de Bohr acentuó el carácter misterioso de sus postulados básicos, uno de los cuales es la relación entre la cuantificación de Bohr del momento angular y la cuantificación de Planck de la energía total de un oscilador armónico simple. Este asunto se aclaró en 1916 cuando Wilson y Sommerfeld enunciaron una regla de cuantificación para *cualquier sistema cuyas coordenadas varían periódicamente con el tiempo*. Esta regla permitió ampliar el dominio de aplicación de la teoría cuántica e incluye como casos particulares las cuantificaciones de Planck y Bohr. La podemos enunciar así:

Regla de cuantificación de Wilson-Sommerfeld:

en un sistema cuántico, toda coordenada q que varía periódicamente en el tiempo satisface la condición de cuantificación

$$\oint p_q dq = n_q h \quad , \quad n_q = 1, 2, 3, \dots \quad (5.25)$$

donde p_q es el impulso asociado a q , y la integración se efectúa sobre un período.

Vamos a verificar que las reglas de Bohr y de Planck son casos particulares de la (5.28).

Si una partícula se mueve en una órbita alrededor de un centro de fuerzas podemos describir el movimiento en el plano de la órbita usando las coordenadas polares r y θ , donde r es la distancia al centro y θ es el ángulo medido desde una dirección fija. El impulso conjugado a r es $p_r = m\dot{r}$ y el impulso conjugado a θ es el momento angular $L = mr^2\dot{\theta}$, que es una constante del movimiento.

Cuando la partícula realiza un movimiento circular uniforme con radio r_0 no es necesario aplicar la condición (5.25) a la coordenada radial, pues $p_r = 0$. La aplicación de (5.22) a la coordenada angular nos da

$$\oint p_q dq = \oint L d\theta = 2\pi L = n_\theta h \quad (5.26)$$

Podemos escribir la (5.26) en la forma

$$L = n_\theta \hbar \quad (5.27)$$

que es precisamente la condición de cuantificación de Bohr (5.9) si identificamos n_θ con n .

Consideremos ahora el caso de una partícula que realiza oscilaciones armónicas con la frecuencia ν . La posición de la partícula se especifica dando la coordenada x , que cumple

$$x = x_0 \sin(2\pi\nu t + \varphi_0) = x_0 \sin(\omega t + \varphi_0) \quad (5.28)$$

donde x_0 es la amplitud de la oscilación y φ_0 es la fase inicial. El impulso de la partícula es

$$p = m\dot{x} = 2\pi\nu m x_0 \cos(2\pi\nu t + \varphi_0) = \omega m x_0 \cos(\omega t + \varphi_0) \quad (5.29)$$

La condición de Wilson-Sommerfeld nos dice entonces

$$\oint p_q dq = \oint p dx = nh \quad (5.30)$$

El cálculo de la integral nos da

$$\begin{aligned} \oint p dx &= \frac{1}{m} \int_0^{T=2\pi/\omega} p^2 dt = \omega^2 m x_0^2 \int_0^T \cos^2(\omega t + \varphi_0) dt \\ &= \omega m x_0^2 \int_0^{2\pi} \cos^2(\alpha) d\alpha = \pi \omega m x_0^2 = \frac{E}{\nu} \end{aligned} \quad (5.31)$$

donde

$$E = \frac{1}{2} \omega^2 m x_0^2 \quad (5.32)$$

es la energía total del oscilador. Usando la (5.31) en la (5.30) obtenemos finalmente

$$E = nh\nu \quad (5.33)$$

que es precisamente la condición de cuantificación de Planck.

Órbitas elípticas y la teoría relativística de Sommerfeld

Si quitamos la restricción que el movimiento del electrón sea circular, tenemos que aplicar las condiciones (5.25) no sólo al movimiento azimutal, sino también al movimiento radial. Se obtienen entonces dos condiciones de cuantificación:

$$\oint L d\theta = 2\pi L = n_\theta h \quad , \quad \oint p_r dr = n_r h \quad (5.34)$$

La primera de ellas nos da la condición de cuantificación del momento angular que ya conocemos para las órbitas circulares:

$$L = n_\theta \hbar \quad , \quad n_\theta = 1, 2, 3, \dots \quad (5.35)$$

La segunda condición (que no se aplica a las órbitas circulares) nos lleva después de algunas cuentas a una relación entre los semiejes a y b de la elipse:

$$L \left(\frac{a}{b} - 1 \right) = n_r \hbar \quad , \quad n_r = 0, 1, 2, \dots \quad (5.36)$$

Conviene ahora definir el *número cuántico principal* n como

$$n \equiv n_\theta + n_r \quad (5.37)$$

que coincide con el único número cuántico que usamos para tratar las órbitas circulares. De acuerdo con (5.35) y (5.36) n puede tomar los valores

$$n = 1, 2, 3, \dots \quad (5.38)$$

Para un valor fijo de n , el *número cuántico azimutal* n_θ puede tomar los valores

$$n_\theta = 1, 2, \dots, n \quad (5.39)$$

y el *número cuántico radial* n_r vale entonces

$$n_r = n - n_\theta \quad (5.40)$$

Usando las fórmulas de la Mecánica para el movimiento en un campo de fuerzas centrales, se puede mostrar entonces (omitimos los detalles por brevedad) que

$$a = \frac{n^2 \hbar^2}{\mu Z e^2} \quad , \quad b = a \frac{n_\theta}{n} \quad , \quad E = -\frac{\mu Z^2 e^4}{2n^2 \hbar^2} \quad (5.41)$$

donde μ es la masa reducida del electrón.

Se puede ver que (salvo por la sustitución $\mu \rightarrow m_e$) la primera de estas ecuaciones es idéntica a la (5.11), que da el radio de las órbitas circulares de Bohr; la segunda muestra que la forma de las órbitas depende del cociente n_θ / n y la tercera equivale a la (5.16). Por lo tanto para cada valor del número cuántico principal hay n diferentes órbitas permitidas, que difieren por el valor del número cuántico azimutal. Una de ellas, la que corresponde a $n_\theta = n$, es la órbita circular de Bohr y las otras son elípticas. Todas esas n órbitas tienen el mismo valor de la energía, pues E depende sólo del número cuántico principal. Esta circunstancia se expresa diciendo que el correspondiente nivel de energía está *degenerado*.

La degeneración de los niveles de energía correspondientes a las órbitas con el mismo n y diferente n_θ es consecuencia de que la fuerza de Coulomb depende de la inversa del cuadrado de la distancia, y de que tratamos el problema como no relativístico. Por eso esta degeneración se suele denominar “accidental”. Si se calculan las correcciones relativísticas, cosa que hizo Sommerfeld, la degeneración se rompe y los niveles de energía del mismo n pero diferentes n_θ se separan, dando lugar a lo que se llama la *estructura fina* del espectro del hidrógeno.

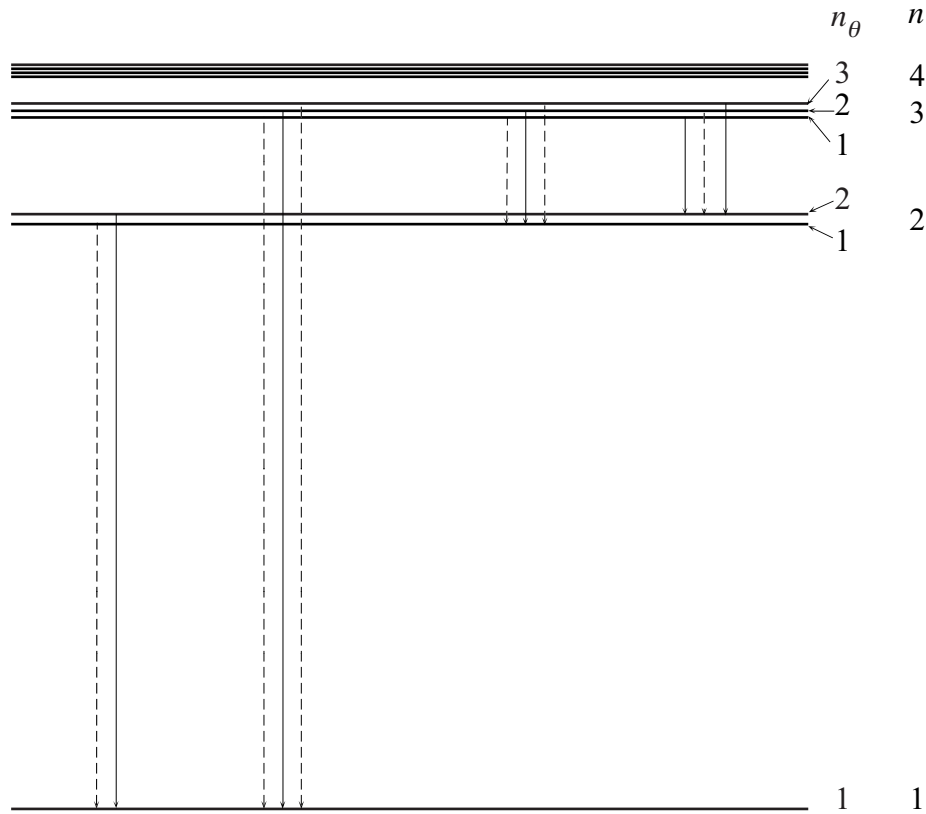


Fig. 5.4. Estructura fina de algunos niveles del átomo de hidrógeno. Se ha exagerado la separación para hacerla visible. Las flechas de trazos indican transiciones que no se observan.

Puesto que $v/c \approx \alpha \approx 10^{-2}$, como ya vimos, las correcciones relativísticas son del orden de $(v/c)^2 \approx \alpha^2 \approx 10^{-4}$. El cálculo es largo y tedioso, pero se puede mostrar que

$$E = -\frac{m_e Z^2 e^4}{2n^2 \hbar^2} \left[1 + \frac{\alpha^2 Z^2}{n} \left(\frac{1}{n_\theta} - \frac{3}{4n} \right) \right] \quad (5.42)$$

Vemos entonces que la separación de los niveles degenerados es proporcional a α^2 . Ese es el motivo por el cual α recibió el nombre de *constante de la estructura fina*, aunque como veremos más adelante su significado es mucho más general. En la Fig. 5.4 se ve la estructura fina de los primeros niveles del átomo de hidrógeno. Las flechas indican las transiciones que producen las líneas espectrales. Las que se observan en el espectro están indicadas con líneas llenas, y con líneas punteadas se indican transiciones que no se observan experimentalmente. Los números de onda de las líneas, calculadas a partir de la (5.42), concuerdan muy bien con los observados. Examinando la figura se ve que sólo se producen las transiciones para las cuales

$$\Delta n_{\theta} \equiv n_{\theta_f} - n_{\theta_i} = \pm 1 \quad (5.43)$$

La ec. (5.43) se denomina *regla de selección*.

El principio de correspondencia

La existencia de reglas de selección no se explica por medio de la teoría que hemos desarrollado hasta ahora y por ese motivo, buscando darles una justificación teórica, Bohr introdujo en 1923 un postulado adicional:

Principio de correspondencia:

- Para todo sistema físico, las predicciones de la teoría cuántica deben corresponder a las predicciones clásicas para valores grandes de los números cuánticos que especifican al sistema.
- Una regla de selección vale para cualquier valor del número cuántico n correspondiente. Por lo tanto toda regla de selección que se aplica en el límite clásico (n grande) se aplica también en el dominio cuántico (n pequeño).

Es obvio que las predicciones de la teoría cuántica deben corresponder a las predicciones clásicas en aquél límite en que el sistema se comporta clásicamente. La primera parte del principio de correspondencia expresa que ese límite se encuentra en el dominio de los números cuánticos grandes, y se apoya en hechos conocidos, como por ejemplo que la teoría de Rayleigh-Jeans del espectro del cuerpo negro concuerda con la teoría de Planck para ν pequeño. A partir de la ec. (4.11) se ve que

$$\lim_{\nu \rightarrow 0} \bar{\epsilon} = \lim_{\nu \rightarrow 0} \frac{h\nu}{e^{\frac{h\nu}{kT}} - 1} = kT \quad (5.44)$$

y por lo tanto

$$\bar{\epsilon} = \bar{n}h\nu \rightarrow kT \quad (5.45)$$

de modo que el valor promedio del número cuántico que especifica la energía de las ondas electromagnéticas de frecuencia ν debe aumentar a medida que disminuye ν .

La segunda parte del principio de correspondencia es una hipótesis razonable, pues no parece lógico que una regla de selección valga sólo para un dominio limitado del número cuántico involucrado.

A modo de ejemplo, podemos aplicar el principio de correspondencia a un oscilador armónico simple de frecuencia ν , cargado eléctricamente. De acuerdo con la teoría cuántica los estados de

energía de este sistema están dados por la ecuación $E_n = nh\nu$, y las teorías cuántica y clásica coinciden para $n \rightarrow \infty$. Puesto que el oscilador está cargado puede emitir o absorber radiación electromagnética. De acuerdo con la teoría clásica, el sistema emite radiación de frecuencia ν debido al movimiento acelerado de la carga. De acuerdo con la teoría cuántica, el oscilador emite un fotón de frecuencia $\nu' = (E_i - E_f)/h = (n_i - n_f)\nu$ cuando efectúa una transición desde el estado n_i al estado n_f . Por lo tanto, la primera parte del principio de correspondencia requiere que $\nu' = \nu$ y en consecuencia en el límite clásico se debe cumplir la regla de selección

$$\Delta n_{\text{emisión}} \equiv n_f - n_i = -1 \quad (5.46)$$

Aplicando un razonamiento semejante a la absorción de radiación se llega a la conclusión que en el límite clásico se cumple la regla de selección

$$\Delta n_{\text{absorción}} \equiv n_f - n_i = +1 \quad (5.47)$$

La segunda parte del principio de correspondencia nos dice que las reglas de selección (5.46) y (5.47) valen en todo el dominio cuántico. El estudio del espectro vibracional de moléculas diatómicas muestra que en efecto esto es cierto.

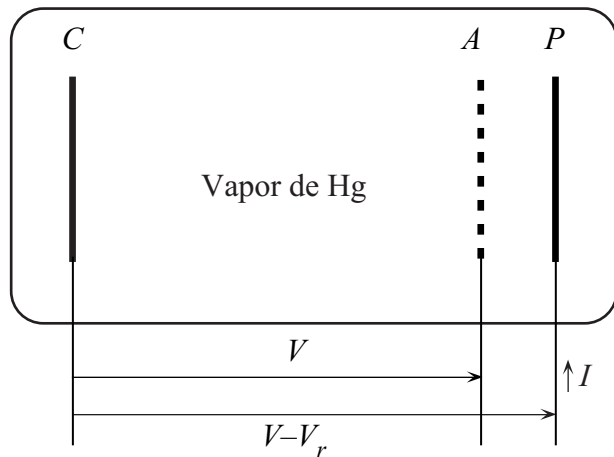
Mediante el estudio de los espectros atómicos y moleculares se encontraron empíricamente numerosas reglas de selección, gran parte de las cuales se pudieron entender aplicando el principio de correspondencia, aunque a veces surgieron ambigüedades.

Se debe notar, sin embargo, que la teoría cuántica moderna *no precisa* invocar el principio de correspondencia para explicar las reglas de selección, pues éstas surgen como consecuencia de leyes generales de conservación, sin necesidad de postulados adicionales.

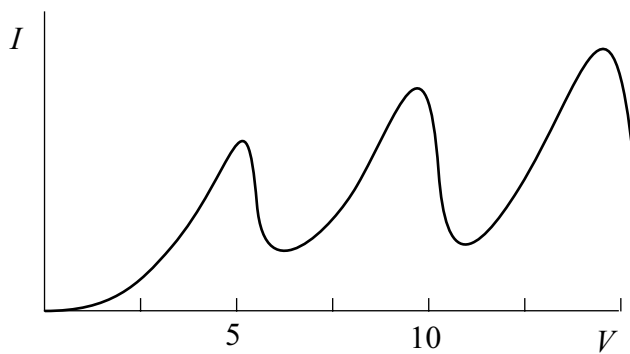
El experimento de Franck y Hertz

La confirmación directa que los estados de energía del átomo están cuantificados provino de un sencillo experimento realizado en 1914 por James Franck y Gustav Hertz. En este experimento (Fig. 5.5a) un cátodo caliente C emite electrones que son acelerados hacia un ánodo A en forma de grilla por una diferencia de potencial V , y pasan a través de él para ser recogidos por una placa colectora P , que está a un potencial $V_p = V - V_r$. El dispositivo contiene el gas o vapor de los átomos que se quiere investigar, a baja presión. El experimento consiste en determinar la corriente I debida a los electrones recogidos por la placa como función de V .

El primer experimento se realizó con vapor de Hg, y los resultados se muestran en la Fig. 5.5b. Se observa que para V pequeño, I aumenta con V , pero cuando se llega a 4.9 V la corriente cae abruptamente. Esto indica que cuando alcanzan una energía cinética de 4.9 eV, los electrones comienzan bruscamente a interactuar con los átomos de Hg, y una fracción importante de ellos pierde toda su energía cinética al excitar los átomos. Si V es apenas mayor que 4.9 V, este proceso de interacción ocurre justo delante de A y los electrones que han perdido su energía cinética ya no la pueden recuperar en el resto de su viaje hacia el ánodo; por lo tanto son rechazados por el potencial de frenamiento V_r y no llegan a la placa. La caída abrupta de I cuando $V = 4.9$ V indica que los electrones de menos de 4.9 eV no pueden transferir su energía a los átomos de Hg. De esto se concluye que los niveles de energía del átomo de Hg están cuantificados, y que 4.9 eV es la diferencia de energía entre el estado fundamental y el primer estado excitado.



(a)



(b)

Fig. 5.5. Experimento de Franck y Hertz: (a) esquema del dispositivo, (b) resultados.

Si esta conclusión es correcta, debe existir en el espectro de emisión del Hg una línea que corresponde a la transición del primer estado excitado al estado fundamental, con una frecuencia dada por $\nu = (4.9 \text{ eV})/h$, lo que corresponde a una longitud de onda de $2536 \text{ \AA}$. Eso fue lo que observaron Franck y Hertz, quienes comprobaron que mientras $V < 4.9 \text{ eV}$ el vapor de Hg no emite ninguna línea espectral, pero cuando el potencial es ligeramente mayor el espectro muestra una única línea de emisión de $2536 \text{ \AA}$.

El experimento de Franck y Hertz es una clara prueba que los estados de energía del átomo de Hg están cuantificados, y permite medir directamente las diferencias de energía entre los estados cuánticos. Si se extiende el estudio a diferencias de potencial mayores aparecen, en efecto, otras bruscas caídas de la corriente. Algunas de éstas se deben a que si V es suficientemente grande, en su recorrido de C a A los electrones pueden excitar dos o más veces el nivel de 4.9 eV , pero otras caídas se deben a la excitación de estados diferentes. A partir de los valores de V correspondientes a estas caídas

se pueden determinar las diferencias de energía entre esos estados y el estado fundamental.

Otra manera de determinar experimentalmente el esquema de niveles de un átomo es medir su espectro, para construir un conjunto de niveles compatible con el mismo. Esto no es fácil en la práctica, puesto que el espectro, y por lo tanto el esquema de niveles, suele ser muy complicado. Por otra parte, el método espectroscópico tiene la virtud de ser muy preciso.

Toda vez que las diferencias de energía entre estados atómicos se determinaron por ambos métodos, el espectroscópico y el de Franck y Hertz, los resultados fueron coincidentes.

Por encima del estado discreto más alto, es decir para $E \geq 0$, están los estados de energía que consisten de un electrón no ligado y un átomo ionizado. La energía del electrón libre no está cuantificada, por lo tanto los estados del electrón con $E \geq 0$ forman un continuo. Es posible excitar el átomo desde el estado fundamental hasta un estado del continuo, si se suministra a un electrón una energía mayor que la *energía de ionización*, que para el átomo de Hg es de 10.4 eV (13.6 eV para el átomo de hidrógeno). También es posible el proceso inverso, esto es que un átomo ionizado capture un electrón libre en uno de los estados ligados del átomo neutro. En este proceso se emite radiación de una frecuencia mayor que la que corresponde al límite de la serie correspondiente al nivel en cuestión. El valor exacto de la frecuencia depende de la energía inicial $E \geq 0$ del electrón libre, y puesto que E puede tener cualquier valor, el espectro tiene una parte continua más allá del límite de la serie. Esto se puede observar en el caso del Hg, aunque con dificultad.

Constantes fundamentales y escalas de la Física Atómica

Lo que hasta ahora hemos visto acerca de la física del átomo y de la interacción entre la radiación electromagnética y la materia nos muestra que en ellas intervienen cuatro constantes físicas universales:

Constantes fundamentales de la Física Atómica:

- e^2 , el cuadrado de la carga eléctrica fundamental ($e = 4.80321 \times 10^{-10}$ u.e.s.)
- $c = 2.99792458 \times 10^{10}$ cm/s, la velocidad de la luz
- $m_e = 9.1093897 \times 10^{-28}$ g, la masa del electrón
- $\hbar = 1.05457 \times 10^{-27}$ erg s, la constante de Planck

En el sistema Gaussiano, y tomando como dimensiones fundamentales longitud (L) tiempo (T) y energía (E), la dimensionalidad de esas constantes es:

$$[e^2] = EL \quad , \quad [c] = LT^{-1} \quad , \quad [m_e] = EL^{-2}T^{-2} \quad , \quad [\hbar] = ET \quad (5.48)$$

En virtud del Teorema Pi, hay una única combinación adimensional independiente de estas cuatro constantes, que es la *constante de la estructura fina*

$$\alpha = \frac{e^2}{\hbar c} = 1/137.0359895 \quad (5.49)$$

Con las tres constantes clásicas e^2 , c , m_e podemos formar una longitud característica que es el *radio clásico del electrón*

$$r_0 = \frac{e^2}{m_e c^2} = 2.817938 \times 10^{-13} \text{ cm} \quad (5.50)$$

un tiempo característico

$$\tau = \frac{e^2}{m_e c^3} = \frac{r_0}{c} = 0.939963 \times 10^{-23} \text{ s} \quad (5.51)$$

y una energía característica, que es el equivalente de la masa en reposo del electrón:

$$\varepsilon = m_e c^2 \quad (5.52)$$

Todas las longitudes características que aparecen en la teoría se pueden entonces expresar en términos de r_0 y α :

- la longitud de onda Compton del electrón

$$\lambda_C \equiv \frac{\lambda_C}{2\pi} = \frac{\hbar}{m_e c} = \frac{r_0}{\alpha} = 0.386159 \times 10^{-10} \text{ cm} \quad (5.53)$$

- el radio de Bohr

$$a_0 = \frac{\hbar^2}{m_e e^2} = \frac{r_0}{\alpha^2} = 0.529177 \times 10^{-8} \text{ cm} \quad (5.54)$$

- la longitud de onda característica del espectro atómico

$$\tilde{\lambda} \approx \frac{2n^2 c \hbar^3}{Z^2 m_e e^4} = \frac{2n^2}{Z^2} \tilde{\lambda}_0, \quad \tilde{\lambda}_0 = \frac{c \hbar^3}{m_e e^4} = \frac{r_0}{\alpha^3} = 7.25163 \times 10^{-7} \text{ cm} \quad (5.55)$$

Vemos en consecuencia que las longitudes características de los fenómenos de escala atómica que involucran electrones y radiación electromagnética guardan entre sí relaciones de escala determinadas por la constante de la estructura fina α . Concretamente:

$$r_0 : \tilde{\lambda}_C : a_0 : \tilde{\lambda}_0 = 1 : \alpha^{-1} : \alpha^{-2} : \alpha^{-3} \quad (5.56)$$

Relaciones análogas existen entre los tiempos y energías características. Estas relaciones tienen un carácter *fundamental* pues provienen de las propiedades de la interacción electromagnética y la naturaleza cuántica de los fenómenos atómicos, y son *independientes* de la particular teoría que los describe.

Crítica de la Teoría Cuántica Antigua

Hemos visto en este Capítulo que la teoría de Bohr y su extensión por Wilson y Sommerfeld, que constituyen la *Teoría Cuántica Antigua*, tuvieron importantes éxitos. Sin embargo debemos señalar las siguientes limitaciones y defectos:

- La teoría se aplica solamente a sistemas periódicos en el tiempo, lo que excluye muchos sistemas físicos.
- Permite calcular las energías de los estados permitidos y las frecuencias de la radiación emitida o absorbida en las transiciones entre esos estados, pero no predice el tiempo característico involucrado en una transición.
- Sólo se aplica a los átomos con un electrón, y aquellos que tienen muchos aspectos en común con los átomos de un electrón (como los metales alcalinos), pero falla si se la intenta aplicar al átomo de helio, que tiene dos electrones.
- Por último, la teoría no es intelectualmente satisfactoria, pues se mezclan en ella de forma arbitraria aspectos clásicos con aspectos cuánticos.

Puesto que algunas de estas objeciones son de carácter fundamental, los físicos de la época se esforzaron por desarrollar una nueva teoría cuántica que no padeciera estas limitaciones ni fuera pasible de objeciones. Este esfuerzo logró el objetivo cuando Werner Heisenberg en 1924 (y luego Max Born y Pascual Jordan) propuso su *dinámica de matrices* y Erwin Schrödinger en 1925, apoyándose en una idea propuesta en 1924 por Louis-Victor de Broglie, desarrolló la *mecánica ondulatoria*. Pese a que su forma es muy distinta, las teorías de Heisenberg y de Schrödinger son completamente equivalentes y su contenido es idéntico, como fue demostrado por Schrödinger. El planteo axiomático de la Mecánica Cuántica se completó poco después por medio de la *teoría de las transformaciones* de Paul A. M. Dirac y Pascual Jordan. Pero de eso nos ocuparemos en los próximos capítulos.

6. PROPIEDADES ONDULATORIAS DE LA MATERIA

El postulado de Broglie

El desarrollo de la Mecánica Cuántica comenzó con una idea muy simple pero revolucionaria que fue expuesta en 1924 por Louis-Victor de Broglie en su Tesis Doctoral. Inspirado por el comportamiento dual onda-corpúsculo de la radiación, de Broglie especuló sobre la posibilidad que también la materia tuviera un comportamiento dual, esto es que las entidades físicas que consideramos como partículas (electrones, átomos, bolas de billar, etc.) pudieran en determinadas circunstancias manifestar propiedades ondulatorias.

Hemos visto que la naturaleza corpuscular de la radiación electromagnética se pone en evidencia cuando se estudia su interacción con la materia (emisión y absorción, efecto fotoeléctrico, efecto Compton, creación y aniquilación de pares, etc.). Por otra parte, su naturaleza ondulatoria se manifiesta por la forma con que se propaga, dando lugar a los fenómenos de interferencia y difracción. Esta situación se puede describir diciendo que la radiación electromagnética es una onda que al interactuar con la materia manifiesta un comportamiento corpuscular. Con igual derecho podemos también decir que consta de partículas (los fotones) cuyo movimiento está determinado por las propiedades de propagación de ciertas ondas que les están asociadas. En realidad ambos puntos de vista son aceptables. Pensando en términos de la segunda alternativa y razonando por analogía, de Broglie exploró la idea que el movimiento de una *partícula* está gobernado por la propagación de ciertas *ondas piloto* asociadas con ella. Ciertamente, es muy sugestivo el hecho que la constante de Planck juegue un rol crucial, tanto para el comportamiento de los electrones del átomo (como lo muestra claramente el éxito de la teoría de Bohr) como para la interacción de la radiación con la materia. En el caso de la radiación, h está vinculado con los aspectos corpusculares de la misma. No es absurdo entonces especular sobre la posibilidad que en el caso de una partícula como el electrón, h esté relacionado con alguna clase de comportamiento ondulatorio.

Cuando de Broglie publicó su trabajo aún no se habían observado comportamientos ondulatorios asociados con el movimiento de una partícula, aunque el tema había sido investigado en varias ocasiones. Pero esta falta de evidencia no es excluyente, pues si en esas ocasiones la longitud de onda de las ondas piloto hubiese sido muy corta no hubiera sido posible observar los aspectos ondulatorios. Esta situación se da también en la Óptica, donde para observar interferencia o difracción es preciso que las diferencias de caminos ópticos sean del orden de la longitud de onda de la luz empleada. Cuando esto no ocurre, la propagación de la luz se puede describir adecuadamente mediante la óptica geométrica, que es en esencia una teoría corpuscular.

Para determinar la longitud de onda de las ondas piloto, de Broglie procedió por analogía a lo que se hace con la radiación electromagnética, considerada como un conjunto de fotones. De acuerdo con la ecuación de Einstein, la frecuencia de un fotón de energía E es

$$\nu = E / h \quad (6.1)$$

La longitud de onda se calcula mediante la relación usual

$$\lambda = v_f / \nu \quad (6.2)$$

donde v_f es la velocidad de fase de la onda. Para el caso del fotón, $v_f = c$ de modo que

$$\lambda = c / \nu = hc / E \quad (6.3)$$

Recordando que la cantidad de movimiento del fotón es $p = E / c$, tenemos que

$$\lambda = h / p \quad (6.4)$$

En consecuencia, por analogía con las ecs. (6.1) y (6.4), se puede formular el

Postulado de Broglie:

La longitud de onda y la frecuencia de la onda piloto asociada a una partícula de impulso p y energía relativística total E están dadas por

$$\lambda = h / p \quad , \quad \nu = E / h \quad (6.5)$$

y el movimiento de la partícula está regido por la propagación de las ondas piloto.

La longitud de onda de la onda piloto se llama *longitud de onda de Broglie* de la partícula.

Algunas propiedades de las ondas piloto

En la descripción del movimiento de la partícula por medio de la onda piloto está implícita la hipótesis que la posición de la partícula está determinada por la onda, en el sentido que la probabilidad de encontrar la partícula en un dado lugar está relacionada con la amplitud de la onda en ese lugar. Si bien aún no conocemos la ecuación que rige la propagación de la onda piloto, podemos hacer algunas afirmaciones sobre su comportamiento, basadas en las propiedades generales de las ondas. Para simplificar trataremos una sola dimensión espacial, pues la generalización a tres dimensiones es obvia. En primer lugar, para que se puedan presentar interferencia y difracción, las ondas piloto deben cumplir el principio de superposición. Por lo tanto podemos construir paquetes de onda de la forma:

$$\Psi(x, t) = \int_{-\infty}^{+\infty} A(k') e^{-i\varphi(k')} e^{i(k'x - \omega't)} dk' \quad (6.6)$$

que es una superposición de ondas planas del tipo¹

$$\Psi_{k'} = e^{i(k'x - \omega't)} \quad (6.7)$$

con una distribución espectral dada por la función real $A(k')$ y fases $\varphi(k')$, donde

$$k' = \frac{2\pi}{\lambda'} = \frac{p'}{\hbar} \quad , \quad \omega' = \omega(k') = 2\pi\nu' = \frac{E'}{\hbar} \quad (6.8)$$

¹ El lector no se debe sentir incómodo porque $\Psi(x, t)$ y Ψ_k sean complejas. Es más sencillo para el cálculo usar exponenciales imaginarias en lugar de senos o cosenos y de últimas, si se desea, se puede siempre tomar la parte real de $\Psi(x, t)$ y Ψ_k . Además, veremos en el Capítulo 7 que las ondas asociadas con las partículas *son* complejas.

Como $E' = m'c^2$ y $p' = m'v'$ (donde $m' = m_0/[1 + (v'/c)^2]^{1/2}$, m_0 es la masa en reposo de la partícula y v' su velocidad), cada onda monocromática (6.7) se propaga con la velocidad de fase

$$v'_f = \frac{\omega'}{k'} = \frac{E'}{p'} = \frac{c^2}{v'} \quad (6.9)$$

La (6.9) muestra que v_f no coincide con la velocidad de la partícula y es mayor que c , pero éste no es un inconveniente, pues el movimiento de una partícula localizada está descrito por el grupo (6.6) y no por las ondas monocromáticas (6.7), que no están localizadas sino que se

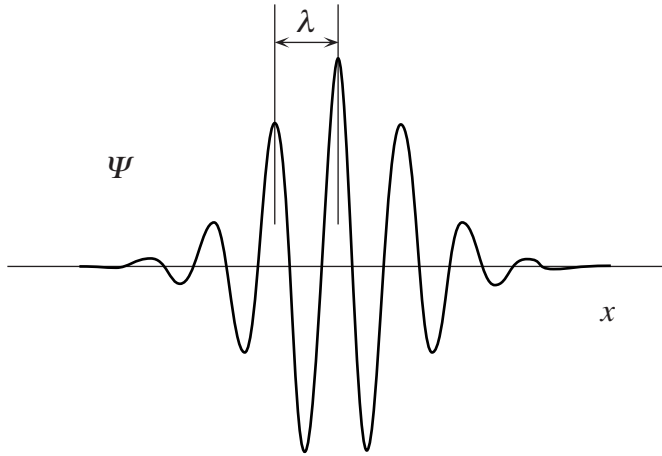


Fig. 6.1. Paquete de ondas.

extienden a todo el espacio. Es importante observar que v_f depende de k , lo que implica que hay *dispersión* y por lo tanto el grupo de ondas se distorsiona y se ensancha a medida que se propaga. En esto las ondas piloto difieren de las ondas electromagnéticas en el vacío, para las cuales la velocidad de fase c no depende de k , y por lo tanto un paquete de ondas se propaga sin cambio de forma.

Supongamos que la distribución espectral $A(k')$ tiene su máximo en $k' = k$ y difiere de cero sólo para k'

próximo a k . En tal caso, las únicas contribuciones significativas en la superposición (6.6) provienen de los valores de k' próximos a k . El paquete tendrá la forma representada (cualitativamente) en la Fig. 6.1 y estará localizado en el lugar donde las diferentes ondas planas con $k' = k + k''$ (donde $k'' \ll k$) que lo componen están en fase. Calculamos entonces la fase, conservando términos lineales en k'' y despreciado las potencias superiores; resulta:

$$\begin{aligned} k'x - \omega(k')t - \varphi(k') &= (k + k'')x - \omega(k + k'')t - \varphi(k + k'') \\ &= kx - \omega(k)t - \varphi(k) + \left(x - \frac{\partial \omega}{\partial k'}t - \frac{\partial \varphi}{\partial k'} \right)_k k'' \end{aligned} \quad (6.10)$$

La posición del paquete está dada por el valor de x para el cual se anula el coeficiente de k'' en la (6.10), esto es:

$$x = \left(\frac{\partial \omega}{\partial k'} \right)_k t + \left(\frac{\partial \varphi}{\partial k'} \right)_k = v_g t + x_0, \quad x_0 \equiv \left(\frac{\partial \varphi}{\partial k'} \right)_k \quad (6.11)$$

La (6.11) muestra que el paquete se desplaza con la *velocidad de grupo*, dada por

$$v_g \equiv \left(\frac{\partial \omega}{\partial k'} \right)_k = \frac{\partial E}{\partial p} \quad (6.12)$$

Ahora bien, de $E^2 = c^2 p^2 + m_0^2 c^4$ resulta

$$\frac{\partial E}{\partial p} = c^2 \frac{p}{E} \quad (6.13)$$

y puesto que $E = mc^2$ y $p = mv$, donde v es la velocidad de la partícula, obtenemos que

$$v_g = v \quad (6.14)$$

esto es, el grupo se desplaza con la velocidad de la partícula, como debe ser para que la descripción del movimiento dada por la onda piloto sea consistente. Notar que de (6.9) y (6.14) se obtiene que $v_g v_f = c^2$

El experimento de Davisson y Germer

Claramente, los aspectos ondulatorios del movimiento de una partícula sólo se manifiestan si la longitud de onda de Broglie (6.5) es del orden de magnitud de alguna dimensión característica del experimento y dada la pequeñez de h , esto no es fácil de conseguir. Por ejemplo, una partícula de polvo cuya masa es de 10^{-11} g y que se desplaza con una velocidad de 1 cm/s tiene una longitud de onda de Broglie del orden de 10^{-15} cm, que es despreciable en comparación con el tamaño de cualquier sistema físico (recordemos que el núcleo atómico tiene un radio del orden de 10^{-12} cm). Por consiguiente no se puede verificar el postulado de Broglie estudiando el movimiento de partículas macroscópicas.

Consideremos ahora un electrón de 10 eV de energía, que es del orden de magnitud de la energía cinética del electrón en un átomo de hidrógeno. En este caso resulta

$$\lambda \approx 3.9 \times 10^{-8} \text{ cm} \quad (6.15)$$

que si bien es pequeña, es del orden del tamaño de un átomo y por lo tanto de la distancia interatómica en un cristal. Esto sugiere que cuando un haz de electrones se refleja sobre un cristal, o lo atraviesa, se pueda observar la difracción de la onda piloto.

Los primeros en observar este efecto fueron Clinton Davisson y Lester Germer, en 1927. En su experimento hicieron incidir un haz de electrones de 54 eV (cuya longitud de onda de Broglie es de 1.67 Å) sobre la superficie de un cristal de níquel (distancia interatómica $d = 2.15$ Å), y midieron la cantidad $N(\theta)$ de electrones dispersados a distintos ángulos θ . Encontraron que $N(\theta)$ tiene un pico para $\theta \approx 50^\circ$ (ver Fig. 6.2). Este resultado *prueba cualitativamente el postulado de Broglie*. En efecto, el pico sólo se puede explicar como el efecto de la interferencia constructiva de las ondas dispersadas por los átomos regularmente espaciados sobre la superficie del cristal. Se debe observar que se trata de la interferencia de las ondas asociadas a *un único electrón*, y que provienen de varias partes del cristal. Esto se puede demostrar empleando un haz de intensidad tan pequeña que en todo instante un único electrón está viajando en el aparato; en este caso se observa que la distribución angular de los electrones dispersados no cambia.

Los resultados de Davisson y Germer también confirman *cuantitativamente* el postulado de Broglie. Recordemos la conocida fórmula de la red de difracción (ley de Bragg):

$$d \sin \theta = n \lambda \quad (6.16)$$

Si suponemos que el pico a 50° corresponde a difracción del primer orden ($n = 1$) la (6.16) nos da $\lambda = 1.65 \text{ \AA}$, que dentro de la precisión del experimento concuerda con el valor de la longitud de onda calculada mediante la (6.5). Para voltajes de aceleración mayores se puede observar también un segundo pico (correspondiente a $n = 2$).

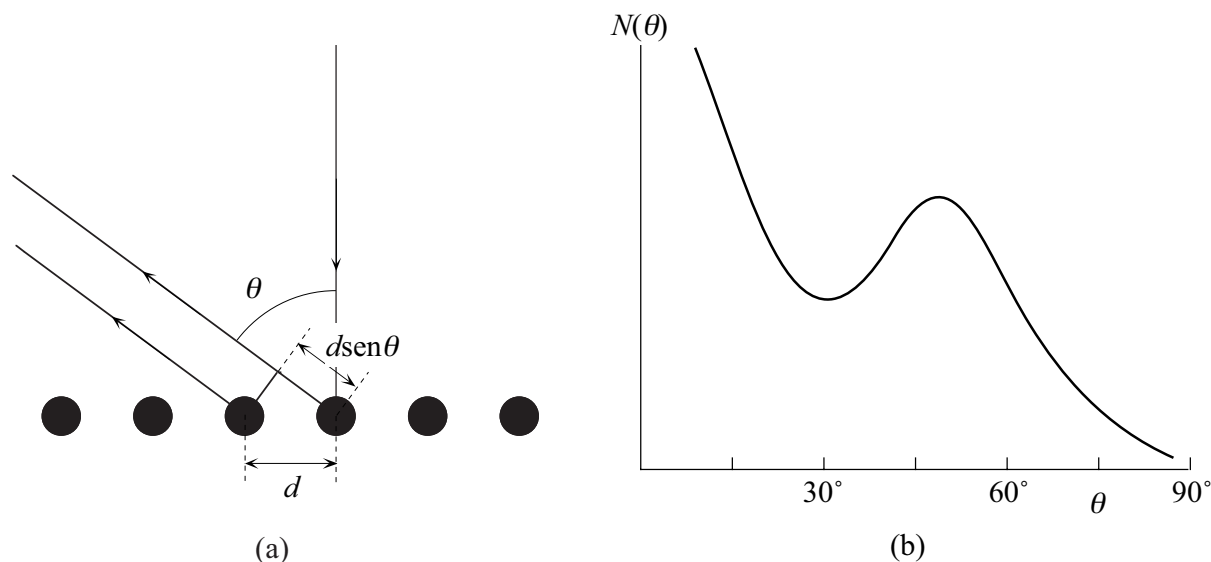


Fig. 6.2. Difracción de electrones: (a) esquema geométrico, (b) resultados.

En 1928 George P. Thomson (hijo de J. J. Thomson) observó la difracción de electrones en la transmisión a través de cristales. Poco después, Immanuel Estermann, Otto Frisch y Otto Stern encontraron efectos de difracción al dispersar átomos de He en la superficie de un cristal de LiF. Desde entonces se observaron muchos otros ejemplos de estos efectos, y la validez del postulado de Broglie quedó confirmada más allá de toda duda.

Vamos a ver ahora que las propiedades ondulatorias del electrón permiten identificar las razones físicas detrás de los hasta entonces misteriosos postulados de la teoría de Bohr y Sommerfeld (Capítulo 5).

Interpretación de la regla de cuantificación de Bohr

La longitud de onda de Broglie de un electrón cuya energía cinética es del orden de la energía cinética del electrón en el átomo de hidrógeno, es del mismo orden de magnitud que el tamaño del átomo. Por ese motivo es sensato esperar que las propiedades de las ondas piloto sean de fundamental importancia para el movimiento del electrón dentro del átomo. Por otra parte hay una importante diferencia entre el movimiento de una partícula libre que hemos considerado hasta ahora y el movimiento del electrón en un estado ligado. En el caso de una partícula libre la onda piloto es una onda viajera. En el caso de un electrón que recorre repetidamente una órbita, cabe esperar que la onda asociada sea estacionaria.

En 1924 de Broglie mostró que las propiedades de las ondas estacionarias permiten interpretar la regla de cuantificación de Bohr del momento angular

$$L = mvr = pr = nh/2\pi \quad , \quad n = 1, 2, 3, \dots \quad (6.17)$$

En efecto, si sustituimos en esta ecuación la expresión del impulso $p = h / \lambda$, obtenemos

$$2\pi r = n\lambda \quad (6.18)$$

es decir: *la circunferencia de las órbitas permitidas contiene un número entero de longitudes de onda de Broglie*. Este es el significado de la condición de cuantificación de Bohr.

Si pensamos que el electrón recorre reiteradamente su órbita, la (6.18) es precisamente la condición necesaria para que la onda piloto asociada al movimiento del electrón se combine coherentemente consigo misma en sucesivos recorridos, de manera que se forme una onda estacionaria. Si se violara esa condición, al cabo de cierto número de vueltas, la onda piloto interferiría destructivamente consigo misma y su intensidad total se anularía. Puesto que la intensidad de la onda piloto se relaciona con la probabilidad de encontrar la partícula, esto implica que el electrón no se puede encontrar en una órbita que no cumpla la (6.18).

Más en general, en el caso de una partícula que efectúa un movimiento periódico se puede demostrar lo mismo. Por consiguiente llegamos a la

Interpretación de Broglie de las reglas de cuantificación de la Teoría Cuántica Antigua:

el requerimiento que la onda piloto asociada sea estacionaria equivale a pedir que el movimiento de la partícula cumpla las condiciones de cuantificación de Wilson-Sommerfeld.

Corresponde subrayar la enorme importancia conceptual de la interpretación de Broglie, que por fin aclara el origen físico de las reglas de cuantificación que hasta entonces era misterioso. Veremos más adelante que las propiedades de las ondas estacionarias tienen una importancia fundamental en la teoría de Schrödinger. Se mostrará que todo estado de energía definida del electrón está descrito por una onda estacionaria, y de resultados de ello todas las características observables de esos estados son *independientes del tiempo*, entre ellas la distribución de la carga eléctrica. Ese es el motivo porqué un electrón no emite ondas electromagnéticas cuando se encuentra en uno de los estados permitidos del átomo.

El principio de incerteza

La descripción del movimiento en términos de la onda piloto trae como consecuencia inevitable que no existe ningún estado de una partícula en el cual se puedan conocer con exactitud y simultáneamente su posición y su cantidad de movimiento. Esto es una consecuencia de propiedades generales de las ondas de cualquier naturaleza (y por lo tanto también de las ondas piloto), y del hecho que, de acuerdo con la interpretación de Broglie, la partícula está localizada donde la onda piloto tiene una amplitud no nula.

Para simplificar consideremos el movimiento de una partícula libre en una dimensión espacial x (la generalización a tres dimensiones es trivial). Supongamos que conocemos con exactitud la cantidad de movimiento p_x de la partícula; la relación (6.4) nos dice entonces que la longitud de onda de la onda piloto debe ser *exactamente* $\lambda_x = h / p_x$. La onda piloto es pues una onda monocromática que se extiende desde $x = -\infty$ a $x = +\infty$, del tipo

$$\Psi = Ae^{i(k_x x - \omega t)} \quad , \quad k_x = 2\pi / \lambda_x = p_x / \hbar \quad , \quad \omega = 2\pi\nu = E / \hbar \quad (6.19)$$

Luego una partícula cuyo impulso se conoce con exactitud puede tener cualquier posición.

La onda piloto que describe una partícula localizada en el entorno de algún x no puede ser monocromática sino que debe ser un paquete del tipo (6.6) (Fig. 6.1), esto es

$$\Psi = \int_{-\infty}^{+\infty} A(k'_x) e^{-i\varphi(k'_x)} e^{i(k'_x x - \omega' t)} dk'_x \quad (6.20)$$

En un instante dado, por ejemplo $t = 0$, tendremos (recordando las (6.10) y (6.11)) que

$$\Psi(x, 0) = \int_{-\infty}^{+\infty} A(k'_x) e^{-i\varphi(k'_x)} e^{ik'_x x} dk'_x = e^{+i[k_x x - \varphi(k_x)]} \int_{-\infty}^{+\infty} A(k_x + k''_x) e^{ik''_x (x - x_0)} dk''_x \quad (6.21)$$

Las ondas monocromáticas se superponen en fase en $x = x_0$ y por lo tanto $\Psi(x, 0)$ es máxima allí. Para $x \neq x_0$, las diferentes ondas con $k'_x = k_x + k''_x$ tienen desfases dados por $k''_x (x - x_0)$ y habrá interferencia destructiva cuando $|k''_x (x - x_0)| \approx \pi$. Si el ancho de la distribución espectral del paquete es Δk_x (esto es, si $A(k_x + k''_x)$ difiere apreciablemente de cero sólo si $|k''_x| < \Delta k_x$), su extensión espacial Δx es entonces (aproximadamente) $\Delta x \approx 2\pi / \Delta k_x$, y en consecuencia, recordando que $p_x = \hbar k_x$, resulta

$$\Delta x \Delta p_x \approx h \quad (6.22)$$

Este argumento muestra que *hay una relación* entre la incerteza de la posición de la partícula (dada por la extensión Δx del paquete) y la incerteza de su impulso (dada por el ancho $\Delta k_x = \Delta p_x / \hbar$ de la distribución espectral del mismo).

Se debe notar que la relación (6.22) es aproximada, porque no dimos aún una definición precisa de Δx y Δp_x . Esta definición depende de la relación entre Ψ y la probabilidad de encontrar la partícula en un determinado lugar, que aún no hemos especificado. Veremos en el próximo Capítulo que dicha probabilidad es proporcional a $|\Psi|^2$. Por lo tanto es natural definir Δx como la desviación standard desde la media, calculada con una distribución de probabilidad proporcional a $|\Psi|^2$. De modo análogo, Δp_x se puede definir en términos de $|A|^2$.

La herramienta matemática apropiada para manejar expresiones del tipo (6.20) es la *transformación* (o integral) *de Fourier*. La *transformada de Fourier* de la función $f(r)$ se indica con $F(s)$ y está definida por

$$F(s) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} f(r) e^{irs} dr \quad (6.23)$$

La transformación de Fourier se puede invertir por medio de la *integral de inversión*

$$f(r) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} F(s) e^{-irs} ds \quad (6.24)$$

y se dice que $F(s)$ es la *antitransformada de Fourier* de $f(r)$. Las transformadas de Fourier figuran en tablas (ver por ejemplo Gradshteyn y Ryzhik, *Tables of Integrals, Series and Pro-*

ducts, Academic Press, 1980) o se calculan numéricamente. De interés para nosotros son la transformada de Fourier de una constante, que es proporcional a la *función delta de Dirac*:

$$f(r) = 1 \quad , \quad F(s) = (2\pi)^{1/2} \delta(s) \quad (6.25)$$

y la transformada de Fourier de una Gaussiana, que es también una Gaussiana:

$$f(r) = e^{-r^2/2a^2} \quad , \quad F(s) = ae^{-a^2s^2/2} \quad (6.26)$$

De la (6.25) vemos que el paquete de ondas que describe una partícula perfectamente localizada en $x = 0$, es decir, tal que $\Psi(x) \propto \delta(x)$ se obtiene como una superposición del tipo (6.20) con $A(k'_x) = \text{cte.}$, es decir, una superposición de todos los posibles valores de k'_x , y por lo tanto de p_x . Por consiguiente, *si una partícula está exactamente localizada en una posición, su cantidad de movimiento es completamente indeterminada.*

El caso de una partícula cuya cantidad de movimiento se conoce con exactitud pero su posición está completamente indeterminada, y el de una partícula cuya posición se conoce con exactitud pero su cantidad de movimiento está completamente indeterminada son casos extremos. En general se conoce la cantidad de movimiento con una incerteza Δp_x y la posición con una incerteza Δx . En ese caso, las propiedades de la integral de Fourier permiten encontrar la relación entre Δx y Δp_x . Por ejemplo, supongamos tener un paquete Gaussiano de ancho Δk_x , esto es

$$A(k'_x) \propto e^{-(k'_x/2\Delta k_x)^2} \quad (6.27)$$

Entonces la (6.26) muestra que en $t = 0$, $\Psi(x)$ es una Gaussiana de ancho Δx :

$$\Psi(x) \propto e^{-(x/2\Delta x)^2}, \quad (6.28)$$

donde los anchos Δk_x y Δx de las respectivas distribuciones de probabilidad (proporcionales a $|\Psi|^2$ y $|A|^2$) cumplen

$$\Delta x \Delta k_x = 1/2 \quad (6.29)$$

y recordando que $k_x = p_x / \hbar$, obtenemos que en $t = 0$ un paquete Gaussiano cumple²

$$\Delta x \Delta p_x = \hbar/2 \quad (6.30)$$

La teoría de la transformación de Fourier permite demostrar que en general los anchos de $f(r)$ y de su transformada $F(s)$, definidos como las *desviaciones standard* desde las medias (calculadas en términos de $|f|^2$ y $|F|^2$), cumplen la relación

$$\Delta r \Delta s \geq 1/2 \quad (6.31)$$

² Se puede mostrar que para tiempos diferentes del inicial, el ancho Δx es mayor.

donde el signo = se cumple sólo cuando $f(r)$ y su transformada $F(s)$ son Gaussianas. En general la relación entre las incertezas de la posición y el impulso de una partícula es entonces:

$$\Delta x \Delta p_x \geq \hbar / 2 \quad (6.32)$$

Generalizando lo anterior a tres dimensiones, llegamos al:

Principio de incerteza de Heisenberg:

si x, y, z son las coordenadas de una partícula y p_x, p_y, p_z son los respectivos impulsos conjugados, se cumple que

$$\begin{aligned} \Delta x \Delta p_x &\geq \hbar / 2 \\ \Delta y \Delta p_y &\geq \hbar / 2 \\ \Delta z \Delta p_z &\geq \hbar / 2 \end{aligned} \quad (6.33)$$

en el caso de coordenadas angulares, entre cada coordenada θ y el correspondiente momento angular L_θ se cumple la relación de incerteza

$$\Delta \theta \Delta L_\theta \geq \hbar / 2 \quad (6.34)$$

Observemos que el principio de incerteza no establece restricciones sobre productos del tipo $\Delta x \Delta p_y, \Delta x \Delta y, \Delta p_x \Delta p_y$, etc.

La (6.32) establece solamente un *límite inferior* al producto $\Delta x \Delta p_x$. Es perfectamente posible tener situaciones en que $\Delta x \Delta p_x \gg \hbar$; esto ocurre cuando nuestras mediciones de la posición y el impulso no alcanzan la máxima precisión posible, compatible con el principio de incerteza. Puesto que $\hbar$ es muy pequeño, es muy difícil que en una medición en escala macroscópica $\Delta x \Delta p_x$ sea comparable con $\hbar$, y por ese motivo el principio de incerteza es *irrelevante* en los experimentos de la Mecánica Clásica. Sin embargo, sus consecuencias son muy importantes cuando se consideran las distancias y los impulsos de los sistemas atómicos y nucleares. Como prueba de ello vamos a mostrar que *el tamaño del átomo está determinado por el principio de incerteza*.

Consideremos, para simplificar, un átomo de hidrógeno. Podemos expresar la energía del electrón como la suma de la energía cinética $\mathbf{p}^2 / 2m_e$ más la energía potencial $-e^2 / r$:

$$E = \frac{\mathbf{p}^2}{2m_e} - \frac{e^2}{r} \quad (6.35)$$

La incerteza de la posición del electrón es $\Delta r \approx r$, y por consiguiente la incerteza de su impulso es $\Delta \mathbf{p} \approx \hbar / \Delta r$, por lo tanto $\mathbf{p}^2 \approx (\Delta \mathbf{p})^2 \approx \hbar^2 / (\Delta r)^2$. Sustituyendo en (6.35) obtenemos

$$E = \frac{\hbar^2}{2m_e(\Delta r)^2} - \frac{e^2}{\Delta r} \quad (6.36)$$

El estado de menor energía se obtiene requiriendo que

$$\frac{dE}{d\Delta r} = -\frac{1}{(\Delta r)^2} \left(\frac{\hbar^2}{m_e \Delta r} - e^2 \right) = 0 \quad (6.37)$$

lo que nos da

$$\Delta r = \frac{\hbar^2}{m_e e^2} = a_0 \quad (6.38)$$

es decir $r = a_0$. Por lo tanto:

El radio de Bohr, que nos da el tamaño del átomo de hidrógeno, se obtiene pidiendo que la energía del átomo tenga el mínimo valor compatible con el principio de incerteza.

Interpretación física de Heisenberg del principio de incerteza

En nuestra presentación, el principio de incerteza surge como una consecuencia matemática de la hipótesis de Broglie, que asocia a cada partícula una onda piloto que describe su movimiento. Pero debemos recordar que cuando Werner Heisenberg introdujo en 1927 el principio de incerteza, sus argumentos no se basaron en la hipótesis de Broglie, sino en las propiedades corpusculares de la radiación electromagnética y sus consecuencias sobre el proceso de medición. De esta forma, Heisenberg puso de manifiesto que existe un límite natural insuperable a la precisión con la que se pueden medir simultáneamente la posición y el impulso de una partícula, debido a que *la medición de una de estas cantidades inevitablemente perturba a la partícula de modo tal que deja incierto el valor de la otra cantidad*. Este es el origen físico del principio de incerteza.

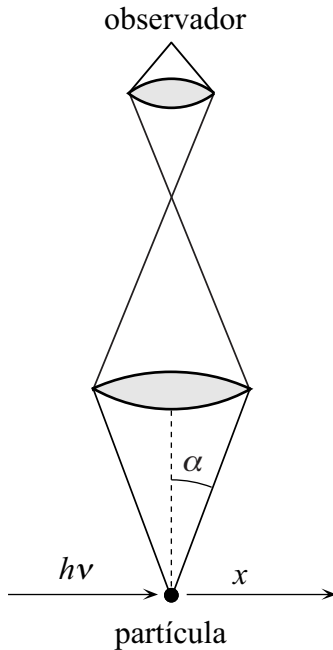


Fig. 6.3. Microscopio de Heisenberg.

La naturaleza de los argumentos de Heisenberg se entiende si examinamos un “experimento de pizarrón” ideado por Bohr en 1928. Consideremos el dispositivo de la Fig. 6.3. Con el microscopio se desea determinar la posición instantánea de la partícula, que podemos ver por medio de los fotones que dispersa cuando se la ilumina. El poder de resolución del microscopio es $\lambda/\text{sen}\alpha$ (λ es la longitud de onda y α es el semiángulo subtendido por el objetivo del microscopio), por lo tanto la indeterminación de la medida es

$$\Delta x \cong \frac{\lambda}{\text{sen}\alpha} \quad (6.39)$$

Supongamos que basta ver un fotón para llevar a cabo la medida. Claramente, el microscopio capta el fotón cuando éste es dispersado en un ángulo comprendido entre $-\alpha$ y $+\alpha$. Luego la incerteza de la componente x del impulso del fotón después de la dispersión es

$$(\Delta p_x)_{\text{fotón}} = 2p \text{sen}\alpha = \frac{2h \text{sen}\alpha}{\lambda} \quad (6.40)$$

pues $p_{\text{fotón}} = h/\lambda$. Ahora bien, como la componente x del impulso del fotón se puede conocer exactamente antes de la dispersión (no hace falta conocer la coordenada x del cuanto), la conservación de la cantidad de movimiento implica que la partícula adquiere un impulso cuya magnitud es incierta en una cantidad igual a la incerteza del impulso del fotón, es decir

$$\Delta p_x = (\Delta p_x)_{\text{fotón}} = \frac{2h \sin \alpha}{\lambda} \quad (6.41)$$

Por lo tanto, en el instante de la medición tenemos que

$$\Delta x \Delta p_x \approx 2h > h \quad (6.42)$$

Si se usa luz de longitud de onda más corta, la medida de la posición es más precisa, pero al mismo tiempo aumenta la incerteza del impulso de la partícula.

Esta discusión muestra que el principio de incerteza es consecuencia de la cuantificación de la radiación electromagnética, y se origina porque para observar la partícula es preciso dispersar por lo menos *un* fotón. En otras palabras, es imposible observar la partícula por medio de una iluminación que le imparta un impulso arbitrariamente pequeño. Debido a la cuantificación, el fotón es el intermediario indispensable entre la partícula y el instrumento de medida, y perturba la partícula de una manera incontrolable e impredecible. Por lo tanto es imposible, después de la medida, conocer con exactitud la posición y la cantidad de movimiento de la partícula. Las relaciones de incerteza (6.32) y (6.33) expresan que la constante de Planck es la medida de la *magnitud mínima* de esa perturbación incontrolable³.

La relación de incerteza entre la energía y el tiempo

Consideremos un grupo de ondas del tipo

$$\Psi = \int_{-\infty}^{+\infty} A(k'_x) e^{-i\varphi(k'_x)} e^{i(k'_x x - \omega' t)} dk'_x, \quad \omega' = \omega(k'_x) \quad (6.43)$$

cuya longitud es Δx . El tiempo que necesita el grupo para recorrer la distancia Δx es

$$\Delta t = \frac{\Delta x}{v_g} = \frac{\Delta x}{v_x} \quad (6.44)$$

donde v_x es la componente x de la velocidad de la partícula. Por lo tanto Δt es la incerteza con la cual se conoce el instante en el cual el grupo pasa por un determinado lugar. Pero igualmente, podemos interpretar que Δt es el intervalo de tiempo durante en cual un observador ubicado en una posición fija x puede llevar a cabo mediciones sobre la partícula.

Por otra parte el grupo es una superposición de ondas de diferentes k'_x , por lo tanto de diferentes frecuencias, y por ende de diferentes energías $E' = \hbar\omega' = \hbar\omega(k'_x)$. De la relación

³ En el libro de Heisenberg *The Physical Principles of Quantum Mechanics* (Dover, 1930) el lector puede encontrar una extensa discusión del principio de incerteza y de numerosos ‘experimentos de pizarrón’ que ilustran su origen físico.

$$\frac{\partial \omega}{\partial k_x} = \frac{\partial E}{\partial p_x} = (v_g)_x = v_x \quad (6.45)$$

resulta que la incerteza en p_x implica una incerteza en la energía de la partícula, dada por

$$\Delta E = v_x \Delta p_x \quad (6.46)$$

Tomando el producto de (6.44) por (6.46) y usando la relación de incerteza $\Delta x \Delta p_x \geq \hbar/2$ obtenemos entonces una relación entre el tiempo Δt durante el cual se observa la partícula y la incerteza ΔE de su energía:

$$\Delta t \Delta E \geq \hbar/2 \quad (6.47)$$

La interpretación de Heisenberg de la relación de incerteza (6.47) es más amplia y se enuncia de la manera siguiente:

Relación de incerteza entre la energía y el tiempo:
una medida de la energía de un sistema efectuada durante el tiempo Δt tiene una incerteza ΔE , y se cumple que $\Delta t \Delta E \geq \hbar/2$.

Veremos más adelante ejemplos de situaciones donde se aplica esta relación de incerteza.

La dispersión de un paquete de ondas

Consideremos un paquete de ondas que describe una partícula libre, de la forma (6.6):

$$\Psi(x, t) = \int_{-\infty}^{+\infty} A(k') e^{i(k'x - \omega' t)} dk' \quad (6.48)$$

y sea Δk_x es el ancho de la distribución espectral $A(k')$. En la (6.48) pusimos $\varphi(k') = 0$ de modo que el grupo está localizado en $x = 0$ cuando $t = 0$.

En $t = 0$ la (6.48) se reduce a

$$\Psi(x, 0) = e^{ikx} \int_{-\infty}^{+\infty} A(k'') e^{ik''x} dk'' \quad (6.49)$$

donde $k' = k + k''$. Para estudiar la evolución temporal del paquete desarrollamos el exponente imaginario de la (6.48) en potencias de k'' , pero a diferencia de lo que hicimos antes (ec. (6.10)) vamos a conservar los términos cuadráticos. Tenemos entonces:

$$\begin{aligned} k'x - \omega(k')t &= (k + k'')x - \omega(k + k'')t = kx - \omega(k)t + \left(x - \frac{\partial \omega}{\partial k} t\right) k'' - \frac{1}{2} \frac{\partial^2 \omega}{\partial k^2} t k''^2 \\ &= kx - \omega(k)t + (x - vt) k'' - \frac{\hbar}{2m} t k''^2 \end{aligned} \quad (6.50)$$

donde para escribir el último renglón usamos que $\partial \omega / \partial k = v_g = v = p/m = \hbar k/m$, y por lo tanto que $\partial^2 \omega / \partial k^2 = \hbar/m$. Sustituyendo (6.50) en (6.48) obtenemos

$$\Psi(x, t) = e^{i[kx - \omega(k)t]} \int_{-\infty}^{+\infty} A(k') e^{i\left[(x-vt)k'' - \frac{\hbar}{2m} t k''^2\right]} dk'' \quad (6.51)$$

Comparando la (6.51) con la (6.49) vemos que *si ignoráramos* el término cuadrático en k'' de la exponencial en el integrando de la (6.51), se tendría que $\Psi(x, t) = e^{-i\omega(k)t} \Psi(x - vt, 0)$, esto es, el grupo se movería sin cambiar de forma (la fase $e^{-i\omega(k)t}$ es irrelevante). Es precisamente el término $\hbar t k''^2 / 2m$ el que da lugar a la dispersión del grupo, al introducir un desfase creciente con el tiempo entre las diferentes ondas monocromáticas que lo componen. La condición de interferencia destructiva es ahora

$$\left| \Delta k \left(\delta x_t - \frac{\hbar}{2m} t \Delta k \right) \right| \approx \pi \quad (6.52)$$

donde hemos escrito $\delta x_t = x - vt$ para indicar el apartamiento desde el centro del grupo. De la (6.52) resulta que el ancho del grupo es entonces

$$\Delta x_t = 2\delta x_t \approx \frac{2\pi}{\Delta k} + \frac{\hbar}{m} t \Delta k \quad (6.53)$$

que podemos escribir en la forma

$$\Delta x_t \approx \frac{h}{\Delta p} + \frac{\Delta p}{m} t = \Delta x_0 + \Delta v t \quad (6.54)$$

La (6.54) muestra que el ancho del paquete crece linealmente con el tiempo desde su valor mínimo $\Delta x_0 \approx h / \Delta p$ para $t = 0$, dado por el principio de incerteza.

Este resultado es coherente con lo que se obtiene clásicamente. En efecto, en la Mecánica Clásica, si en $t = 0$ determinamos la posición de una partícula con una incerteza Δx_0 y su velocidad con una incerteza Δv (debido a las limitaciones de los instrumentos de medida), después de transcurrir un tiempo t la incerteza de la posición es precisamente la que resulta de la ec. (6.54). En este sentido, la fórmula (6.54) no nos dice nada nuevo. La novedad es que ahora, a diferencia de lo que ocurre en la Mecánica Clásica, Δx_0 y Δv *no son independientes* pues están relacionados por el principio de incerteza $\Delta x_0 \Delta v \approx h / m$. De resultados de eso la (6.54) se puede escribir

$$\Delta x_t = \Delta x_0 + \frac{h}{m \Delta x_0} t \quad (6.55)$$

De la (6.55) vemos que cuanto menor es la incerteza Δx_0 , tanto más rápidamente crece Δx_t .

El principio de complementariedad

El principio de incerteza permite resolver las aparentes paradojas que se originan en la dualidad onda-corpúsculo de la radiación y la materia. Si se intenta determinar si la radiación es una onda o un corpúsculo, resulta que todo experimento que fuerza a radiación a exhibir su carácter ondulatorio, al mismo tiempo suprime las manifestaciones de su carácter corpuscular,

y viceversa. Es decir, en una misma situación experimental no se pueden observar a la vez los aspectos ondulatorio y corpuscular. Lo mismo ocurre con la materia. De resultados de ello las evidencias obtenidas bajo distintas condiciones experimentales no se pueden captar en una única imagen, sino que son complementarias, en el sentido que sólo la totalidad de los fenómenos agota la posible información sobre el objeto del estudio. Esto es consecuencia de que a nivel atómico es imposible separar netamente el comportamiento de los objetos (fotones, electrones, etc.) de la interacción con el instrumento de medida, que define las condiciones bajo las cuales aparece el fenómeno. Esta es la esencia del *principio de complementariedad* de Bohr: los conceptos de onda y partícula no se contradicen sino que se complementan.

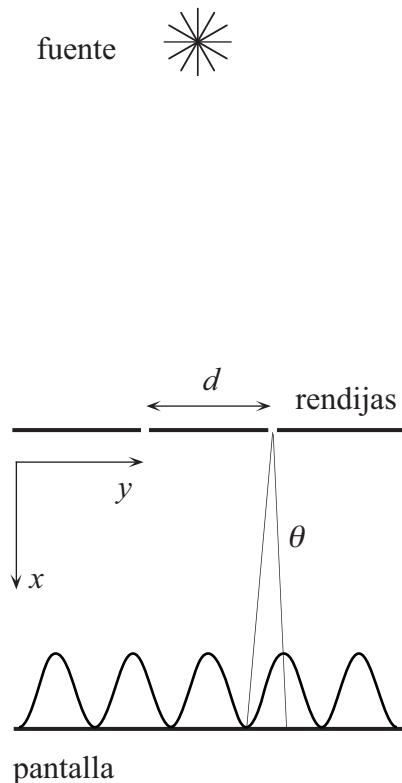


Fig. 6.4. Interferencia por dos rendijas.

Consideremos, por ejemplo, el experimento de Young de interferencia de luz por dos rendijas (o en forma equivalente, la interferencia de electrones) cuyo esquema se da en la Fig. 6.4. La distancia entre las rendijas es d .

Del punto de vista ondulatorio, la onda original se divide en dos ondas coherentes al pasar por las rendijas, y la superposición de ambas produce las características franjas de interferencia en la pantalla.

Supongamos ahora que reemplazamos la pantalla por un mosaico de minúsculos fotocátodos, de manera que midiendo la corriente producida por la emisión del correspondiente fotoelectrón podemos determinar en qué lugar de la pantalla ha llegado cada fotón. Igualmente, la distribución de los fotoelectrones sigue el mismo patrón de interferencia. Sin embargo, cada fotón individual llega a un lugar bien definido de la pantalla: el del fotocátodo donde fue absorbido.

Si se piensa en el fotón como un corpúsculo, parece lógico pensar que tiene que pasar o por una, o por la otra de las rendijas, pero en este caso se plantea una paradoja, pues a primera vista parece absurdo que el movimiento del fotón más allá de las rendijas esté influenciado por la presencia de la rendija por la cual no pasó.

La paradoja proviene de suponer que el fotón debe pasar por una rendija o por la otra. Tal afirmación no tiene sentido a menos que se determine experimentalmente por cuál de las dos rendijas ha pasado. Sin embargo, resulta que esa determinación es imposible de lograr, sin destruir el patrón de interferencia.

En efecto, para determinar por cuál rendija pasa el fotón habría que poner un detector en cada rendija. Pero el detector inevitablemente interactúa con el fotón y altera la trayectoria que de otra forma seguiría. El principio de incerteza permite demostrar que si el detector tiene suficiente resolución espacial como para poder determinar por cuál rendija pasó el fotón, entonces la perturbación que produce en el impulso del fotón causa una desviación que destruye el patrón de interferencia de dos rendijas.

Supongamos, para concretar, que con nuestro detector determinamos la coordenada y del fotón con una incerteza

$$\Delta y \ll d. \quad (6.56)$$

con lo cual podemos asegurar por cuál rendija pasó.

En este proceso de detección hay una interacción que cambia el impulso del fotón; de resultas de ello la componente y del impulso tendrá una incerteza Δp_y . Para no destruir el patrón de interferencia, Δp_y debe ser tal que

$$\Delta p_y / p_x \ll \theta \quad (6.57)$$

donde p_x es la componente x del impulso del fotón y θ es el ángulo subtendido desde la rendija por la posición de un máximo y la del mínimo adyacente, que vale

$$\theta = \lambda / 2d \quad (6.58)$$

Sustituyendo (6.58) en (6.57) obtenemos

$$\Delta p_y \ll p_x \lambda / 2d = h / 2d \quad (6.59)$$

pues $\lambda \equiv h / p_x$. Multiplicando ahora las desigualdades (6.56) y (6.59) obtenemos que la condición que se debe satisfacer para determinar por cuál rendija pasa el fotón, *sin destruir el patrón de interferencia*, es

$$\Delta y \Delta p_y \ll h / 2 \quad (6.60)$$

Pero esta condición viola el principio de incerteza. Por lo tanto *no podemos saber por cuál rendija pasó el fotón y al mismo tiempo ver la figura de interferencia*. Esto significa que el problema que nos preocupa es ilusorio.

Veremos luego otros ejemplos en que el principio de incerteza ayuda a resolver conflictos aparentes entre los aspectos ondulatorios y corpusculares de un ente.

7. LA TEORÍA DE SCHRÖDINGER

Introducción

El trabajo de Broglie llamó la atención de Einstein, quien lo consideró muy importante y lo difundió entre los físicos. Inspirado en las ideas allí expuestas, Erwin Schrödinger desarrolló entre 1925 y 1926 su teoría de la *mecánica ondulatoria*, que es una de las maneras en que se presenta la Mecánica Cuántica. Corresponde mencionar que casi simultáneamente, Werner Heisenberg desarrolló un enfoque alternativo: la *mecánica matricial*. En la teoría de Heisenberg no se consideran ondas piloto; en su lugar se manejan las variables dinámicas como x , p_x , etc., que se representan mediante matrices. Los aspectos cuánticos se introducen en dicha teoría por medio del principio de incerteza, que se expresa por medio de las propiedades de conmutación de las matrices. El principio de incerteza es en realidad *equivalente* al postulado de Broglie, y las teorías de Heisenberg y de Schrödinger son idénticas en contenido aunque de forma aparentemente muy distinta. Pero esto no fue comprendido en seguida, y en un primer momento hubo ácidas polémicas entre los sostenedores de una y otra, hasta que Schrödinger en 1928 demostró la equivalencia de ambas. Debido a que la teoría de Schrödinger se presta mejor para un tratamiento introductorio no entraremos en los detalles de la teoría de Heisenberg¹.

La ecuación de Schrödinger

Aunque el postulado de Broglie es correcto, no es todavía una teoría completa del comportamiento de una partícula, pues no se conoce la ecuación que rige la propagación de la onda piloto. Por eso pudimos analizar la propagación de la onda piloto únicamente en el caso de una partícula libre y no sabemos aún como tratar una partícula sometida a fuerzas. Falta, además, una relación cuantitativa entre la onda y la partícula, que nos diga de qué forma la onda determina la probabilidad de observar la partícula en un determinado lugar.

Pese a que el postulado de Broglie es consistente con la relatividad (restringida), Schrödinger se limitó a desarrollar una teoría no relativística². Además abandonó el término “onda piloto” y llamó *función de onda* a la función $\Psi(x,t)$ y a la onda en sí; nosotros usaremos esa terminología de ahora en más. En consecuencia Schrödinger adoptó como punto de partida las ecuaciones

$$\lambda = h / p \quad (7.1)$$

y

$$v = E / h \quad (7.2)$$

pero supuso que la energía está dada por la expresión no relativística

¹ El lector interesado los puede encontrar en el libro de Heisenberg ya citado en el Capítulo 6.

² Tenía buenas razones para ello. El desarrollo de la Mecánica Cuántica Relativística es mucho más complejo y difícil y sólo se pudo completar muchos años después. La razón fundamental es que en una teoría relativística consistente no se pueden introducir fuerzas (que implican acción a distancia) y además hay que tomar en cuenta los procesos de creación y aniquilación de materia (ver Capítulo 4). Esto implica que se tiene que cuantificar el campo electromagnético (y el que describe cualquier otro tipo de fuerzas que se quiera considerar) y los campos que describen las partículas.

$$E = p^2 / 2m + V \quad (7.3)$$

donde m es la masa en reposo. Esta diferencia en la definición de E cambia el valor de v . Pero cabe señalar que los experimentos de difracción que demuestran la validez de la (7.1), en realidad no dicen nada acerca de la (7.2). Además veremos que el valor de v no tiene importancia en la teoría de Schrödinger. Asimismo, la (7.3) da el valor correcto de la velocidad de grupo para una partícula libre ($V = V_0 = \text{cte.}$), pues usando (7.2) resulta

$$\hbar\omega = \hbar^2 k^2 / 2m + V_0 \quad (7.4)$$

donde $k = 2\pi / \lambda$ y $\omega = 2\pi\nu$, y por lo tanto

$$v_g = \frac{\partial\omega}{\partial k} = \hbar k / m = p / m = v \quad (7.5)$$

La ecuación que rige la propagación de las ondas de materia debe satisfacer entonces los siguientes requisitos:

- Debe ser consistente con las ecs. (7.1), (7.2) y (7.3), esto es con $\hbar\omega = \hbar^2 k^2 / 2m + V$.
- Debe ser lineal en $\Psi(x,t)$, para que podamos superponer funciones de onda y reproducir efectos de interferencia y difracción como los que observaron Davisson y Germer.
- El impulso de una partícula libre es constante; por lo tanto cuando $V = V_0 = \text{cte.}$ la ecuación debe tener soluciones de onda viajera como vimos en el Capítulo 6.

Consideremos entonces una partícula libre. Para satisfacer el tercer requisito, la describiremos mediante una onda viajera Ψ_f que depende del argumento $(kx - \omega t)$. Puesto que para hacer aparecer los factores ω y k^2 de la ec. (7.4) hay que derivar la función de onda una vez respecto del tiempo y dos veces respecto de la coordenada, la ecuación buscada podría ser de la forma

$$\alpha\hbar \frac{\partial\Psi_f}{\partial t} = \frac{\beta\hbar^2}{2m} \frac{\partial^2\Psi_f}{\partial x^2} + V_0\Psi_f \quad (7.6)$$

donde α y β son constantes a determinar. Como la (7.6) es lineal, satisface el segundo requisito. Es fácil ver que ondas viajeras del tipo $\Psi_f = \cos(kx - \omega t)$ ó $\Psi_f = \sin(kx - \omega t)$ *no cumplen* con primer requisito, pues no se puede recoger en un factor común la dependencia espacial y temporal, como haría falta para obtener la (7.4). Pero con un poco de álgebra se ve que una combinación lineal del tipo $\Psi_f = \cos(kx - \omega t) + \gamma\sin(kx - \omega t)$ permite obtener la (7.4) si $\gamma = \pm i$, $\alpha = \gamma$ y $\beta = -1$. Lo usual es tomar $\gamma = +i$, luego $\alpha = i$ y entonces obtenemos de la (7.6) la ecuación

$$i\hbar \frac{\partial\Psi_f}{\partial t} = -\frac{\hbar^2}{2m} \frac{\partial^2\Psi_f}{\partial x^2} + V_0\Psi_f \quad (7.7)$$

que cumple los tres requisitos que estipulamos. La solución de onda viajera es pues de la forma

$$\Psi_f = \cos(kx - \omega t) + i\sin(kx - \omega t) = e^{i(kx - \omega t)} \quad (7.8)$$

y como adelantamos en el Capítulo 6 es *necesariamente* compleja.

La ec. (7.8) se obtuvo para el caso especial $V = V_0 = \text{cte.}$. Vamos a *postular* que cuando $V = V(x,t)$, la ecuación diferencial que describe la función de onda tiene la misma forma. Se obtiene así la

Ecuación de Schrödinger:

$$i\hbar \frac{\partial \Psi(x,t)}{\partial t} = -\frac{\hbar^2}{2m} \frac{\partial^2 \Psi(x,t)}{\partial x^2} + V(x,t)\Psi(x,t) \quad (7.9)$$

Como se ve, es una ecuación en derivadas parciales lineal y del segundo orden, y a diferencia de otras ecuaciones diferenciales de la física contiene el imaginario i . De resultados de eso *sus soluciones son necesariamente funciones complejas*.

La (7.9) vale para una dimensión espacial. En tres dimensiones la ecuación de Schrödinger es:

$$i\hbar \frac{\partial \Psi(\mathbf{r},t)}{\partial t} = -\frac{\hbar^2}{2m} \nabla^2 \Psi(\mathbf{r},t) + V(\mathbf{r},t)\Psi(\mathbf{r},t) \quad (7.10)$$

donde ∇^2 es el operador Laplaciano.

Interpretación de la función de onda

La función de onda $\Psi(x,t)$ es inherentemente compleja y por lo tanto *no se puede medir* con un instrumento real. Pero ésta es una característica deseable pues nos impide atribuir a la función de onda una existencia física como (por ejemplo) la de las olas de la superficie del agua. En realidad $\Psi(x,t)$ no es más que un instrumento de cálculo que sólo tiene significado en el contexto de la teoría de Schrödinger de la que forma parte. Esto queda de manifiesto claramente si se considera que la función de onda *no aparece* en la teoría de Heisenberg, y sin embargo los resultados físicos de ambas teorías son *equivalentes*. Veremos más adelante otras consecuencias de estos hechos. Pero lo dicho no implica que la función de onda carece de interés físico. Veremos, en efecto, que *la función de onda contiene toda la información sobre la partícula asociada, compatible con el principio de incerteza*.

Para obtener esa información hay que relacionar $\Psi(x,t)$ con las variables dinámicas de la partícula asociada. Como dijimos en el Capítulo 6, hay una relación entre la intensidad de $\Psi(x,t)$ en un punto (x, t) y la densidad de probabilidad $P(x,t)$ de encontrar a la partícula en el entorno de ese punto. Sin embargo, es obvio que no podemos igualar una cantidad compleja como $\Psi(x,t)$ con $P(x,t)$ que es una magnitud real. La relación correcta entre $\Psi(x,t)$ y $P(x,t)$ fue propuesta en 1926 por Max Born:

Postulado de Born:

si en el instante t se lleva a cabo una medida para ubicar la partícula descrita por la función de onda $\Psi(x,t)$, entonces la probabilidad $P(x,t)dx$ de que el valor de x se encuentre entre x y $x + dx$ es

$$P(x,t)dx = \frac{\Psi^*(x,t)\Psi(x,t)dx}{\int_{\text{todo } x} \Psi^*(x,t)\Psi(x,t)dx} = \frac{|\Psi(x,t)|^2 dx}{\int_{\text{todo } x} |\Psi(x,t)|^2 dx} \quad (7.11)$$

donde $\Psi^*(x,t)$ indica el complejo conjugado de $\Psi(x,t)$.

Para que el postulado de Born tenga sentido, la integral

$$N_{\Psi} \equiv \int_{\text{todo } x} \Psi^*(x,t)\Psi(x,t)dx \quad (7.12)$$

(que se denomina la *norma* de Ψ) debe existir. En otras palabras, toda función de onda aceptable debe ser de cuadrado integrable, lo que implica que $\Psi(x,t)$ debe tender a cero con suficiente rapidez para $x \rightarrow \pm\infty$. Se debe notar que las soluciones de onda viajera del tipo (7.8) *no* son de cuadrado integrable y por lo tanto *no* son funciones de onda aceptables. Sin embargo no las vamos a descartar pues a partir de ellas, mediante la integral de Fourier, se pueden formar paquetes que sí son de cuadrado integrable y por lo tanto funciones de onda aceptables. Volveremos más adelante sobre este tema.

En adelante vamos a suponer que la función de onda está *normalizada* de modo que la integral del cuadrado de su módulo (extendida a todo el dominio de la coordenada) es igual a la unidad:

$$N_{\Psi} = \int_{\text{todo } x} \Psi^*(x,t)\Psi(x,t)dx = \int_{\text{todo } x} |\Psi(x,t)|^2 dx = 1 \quad (7.13)$$

cosa que podemos lograr siempre, pues la ecuación de Schrödinger es lineal y si Ψ' es una solución no normalizada, entonces

$$\Psi = \frac{\Psi'}{\sqrt{\int_{\text{todo } x} \Psi'^* \Psi' dx}} \quad (7.14)$$

es también solución y está normalizada. Además mostraremos en breve que si la función de onda está normalizada en un dado instante t , permanece normalizada en todo otro instante t' . En consecuencia podemos escribir la (7.11) como

$$P(x,t)dx = \Psi^*(x,t)\Psi(x,t)dx = |\Psi(x,t)|^2 dx \quad (7.15)$$

La cantidad $\Psi^*\Psi = |\Psi|^2$ es siempre *real* y por lo tanto el postulado de Born no es inconsistente. Pero esto no prueba todavía la validez de dicho postulado, pues $\Psi^*\Psi$ no es la única función real que se puede obtener a partir de Ψ . Sin embargo se puede dar un argumento de *plausibilidad* del modo siguiente. La Ψ satisface la ecuación de Schrödinger

$$i\hbar \frac{\partial \Psi}{\partial t} = -\frac{\hbar^2}{2m} \frac{\partial^2 \Psi}{\partial x^2} + V\Psi \quad (7.16)$$

Tomando el complejo conjugado de la (7.16) obtenemos

$$-i\hbar \frac{\partial \Psi^*}{\partial t} = -\frac{\hbar^2}{2m} \frac{\partial^2 \Psi^*}{\partial x^2} + V\Psi^* \quad (7.17)$$

Ahora multiplicamos la (7.16) por Ψ^* y la (7.17) por Ψ y restamos para obtener

$$i\hbar \left(\Psi^* \frac{\partial \Psi}{\partial t} + \Psi \frac{\partial \Psi^*}{\partial t} \right) = -\frac{\hbar^2}{2m} \left(\Psi^* \frac{\partial^2 \Psi}{\partial x^2} - \Psi \frac{\partial^2 \Psi^*}{\partial x^2} \right) \quad (7.18)$$

que se puede poner en la forma

$$\frac{\partial}{\partial t}(\Psi^*\Psi) = -\frac{\partial}{\partial x} \left[\frac{i\hbar}{2m} \left(\Psi \frac{\partial \Psi^*}{\partial x} - \Psi^* \frac{\partial \Psi}{\partial x} \right) \right] \quad (7.19)$$

Si ahora definimos la *corriente de probabilidad* como

$$J(x,t) = \frac{i\hbar}{2m} \left(\Psi \frac{\partial \Psi^*}{\partial x} - \Psi^* \frac{\partial \Psi}{\partial x} \right) \quad (7.20)$$

la (7.19) se escribe como

$$\frac{\partial P}{\partial t} = -\frac{\partial J}{\partial x} \quad (7.21)$$

Esta ecuación expresa la *conservación de la probabilidad*. En efecto, integrando la (7.21) entre x_1 y x_2 obtenemos

$$\frac{\partial}{\partial t} \int_{x_1}^{x_2} P(x,t) dx = J(x_1,t) - J(x_2,t) \quad (7.22)$$

que nos dice que la variación de la probabilidad de encontrar la partícula en el intervalo (x_1, x_2) está dada por el flujo neto de la corriente de probabilidad que entra en dicho intervalo.

Si extendemos el intervalo de integración a todo x , resulta

$$\frac{\partial}{\partial t} \int_{-\infty}^{+\infty} P(x,t) dx = \frac{\partial}{\partial t} \int_{-\infty}^{+\infty} \Psi^*(x,t) \Psi(x,t) dx = \frac{dN_\Psi}{dt} = 0 \quad (7.23)$$

puesto que Ψ y $\partial \Psi / \partial x$, y entonces J , se anulan en $x = \pm\infty$. Por lo tanto la (7.23) nos dice que si la función de onda está normalizada en un dado instante t , permanece normalizada en todo otro instante t' .

La ecuación de Schrödinger independiente del tiempo

Cuando $V = V(\mathbf{r})$ no depende del tiempo, la ecuación de Schrödinger (7.10) admite soluciones separables de la forma

$$\Psi(\mathbf{r},t) = \psi(\mathbf{r})\phi(t) \quad (7.24)$$

Las soluciones de la forma (7.24) son una familia de soluciones particulares de la (7.10) que reviste particular interés, como veremos ahora.

Sustituyendo (7.24) en (7.10) y dividiendo por $\psi(\mathbf{r})\phi(t)$ obtenemos

$$i\hbar \frac{1}{\phi(t)} \frac{d\phi(t)}{dt} = \frac{1}{\psi(\mathbf{r})} \left(-\frac{\hbar^2}{2m} \nabla^2 \psi(\mathbf{r}) + V(\mathbf{r})\psi(\mathbf{r}) \right) \quad (7.25)$$

El miembro izquierdo de la (7.25) depende solo de t y el miembro derecho depende solo de las variables espaciales $\mathbf{r}$. Por lo tanto ambos miembros deben ser iguales a una constante C , que se

denomina *constante de separación*, y entonces la ecuación de Schrödinger se separa en dos ecuaciones, una para $\phi(t)$ y la otra para $\psi(\mathbf{r})$:

$$i\hbar \frac{1}{\phi(t)} \frac{d\phi(t)}{dt} = C \quad , \quad \frac{1}{\psi(\mathbf{r})} \left(-\frac{\hbar^2}{2m} \nabla^2 \psi(\mathbf{r}) + V(\mathbf{r})\psi(\mathbf{r}) \right) = C \quad (7.26)$$

La ecuación para $\phi(t)$ se integra de inmediato. A menos de un factor constante, la solución es:

$$\phi(t) = e^{-iCt/\hbar} \quad (7.27)$$

que es una función compleja puramente oscilatoria con frecuencia $\nu = C/h$. Recordando que la (7.2) nos dice que $\nu = E/h$, donde E es la energía de la partícula, debemos identificar la constante de separación C con la energía total E . Luego la segunda de las (7.26) se escribe como la

Ecuación de Schrödinger independiente del tiempo:

$$-\frac{\hbar^2}{2m} \nabla^2 \psi(\mathbf{r}) + V(\mathbf{r})\psi(\mathbf{r}) = E\psi(\mathbf{r}) \quad (7.28)$$

Esta ecuación determina la *dependencia espacial* $\psi(\mathbf{r})$ de las funciones de onda separables del tipo (7.24). Como se ve es una ecuación del segundo orden, que no contiene factores imaginarios y por lo tanto sus soluciones *no* son necesariamente complejas. Sus soluciones $\psi_E(\mathbf{r})$, que dependen del valor de la constante de separación E , se denominan *autofunciones* (o *funciones propias*) de la energía, y los correspondientes valores de E se denominan *autovalores* (o *valores propios*) de la energía.

Finalmente, la función de onda se escribe como

$$\Psi_E(\mathbf{r}, t) = \psi_E(\mathbf{r}) e^{-iEt/\hbar} \quad (7.29)$$

Veremos en seguida que al resolver la ecuación de Schrödinger independiente del tiempo (7.28) sólo se obtienen soluciones aceptables para ciertos valores particulares de E . De esta forma aparece la cuantificación de la energía.

Cuantificación de la energía en la teoría de Schrödinger

Para ver de manera sencilla como aparece la cuantificación de la energía, consideremos las soluciones $\psi_E(x)$ de la ecuación de Schrödinger independiente del tiempo en el caso de una sola coordenada x . La (7.28) se escribe entonces

$$\frac{\hbar^2}{2m} \psi_{xx} = [V(x) - E]\psi \quad (7.30)$$

Para que una solución sea aceptable ψ y ψ_x deben ser uniformes, finitas y continuas para todo x , pues en caso contrario la función de onda Ψ violaría alguno de estos requisitos, y entonces la densidad de probabilidad o la corriente de probabilidad no estarían bien definidas en algún lugar. Si además E y V son finitos, entonces la (7.30) nos dice que ψ_{xx} es también finita.

Supongamos ahora que resolvemos la (7.30) para un potencial $V(x)$ como se muestra en la Fig. 7.1, por ejemplo por integración numérica. Para eso comenzamos por asignar arbitrariamente un

valor de E (como el que se indica en la Fig. 7.1 con la recta horizontal). Vemos así que los puntos de retorno clásicos x' y x'' (determinados por la condición $E = V(x)$) dividen el eje x en tres regiones: $x < x'$ (región I) y $x > x''$ (región III), donde $V(x) - E > 0$ y por lo tanto $\psi_{xx} > 0$, y $x' < x < x''$ (región II), donde $V(x) - E < 0$ y entonces $\psi_{xx} < 0$.

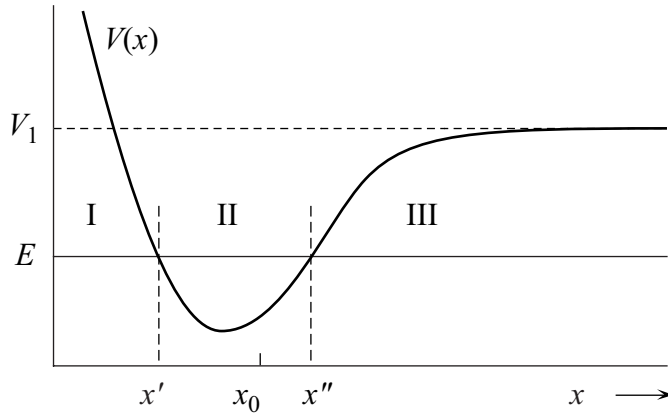


Fig. 7.1. Diagrama de la energía que muestra la relación entre la energía potencial y la energía total en un caso típico.

Comencemos por integrar la (7.30) en la dirección *positiva* de x a partir del punto x_0 , asignando los valores iniciales de $\psi(x_0)$ y $\psi_x(x_0)$. Dada la linealidad de la ecuación (7.30) podemos siempre fijar $\psi(x_0) = 1$, de modo que el único parámetro que podemos variar para obtener una solución aceptable es $\psi_x(x_0)$.

De acuerdo a lo dicho, en la región II ψ es cóncava hacia abajo donde $\psi > 0$ y cóncava hacia arriba donde $\psi < 0$. Luego ψ oscila y se mantiene acotada.

En la región III, en cambio, ψ es cóncava hacia arriba donde $\psi > 0$ y cóncava hacia

abajo donde $\psi < 0$. Entonces ψ se dispara hacia $+\infty$ (si es positiva) como en la curva (a) de la Fig. 7.2, o hacia $-\infty$ (si es negativa) como en la curva (b), y en ambos casos, tanto más rápidamente cuanto mayor es $|\psi_x(x_0)|$. Por supuesto ninguno de estos comportamientos es aceptable. Solamente para un *único* valor $\psi_x(x_0)_{III}$ de $\psi_x(x_0)$ se encontrará una ψ que tenga un comporta-

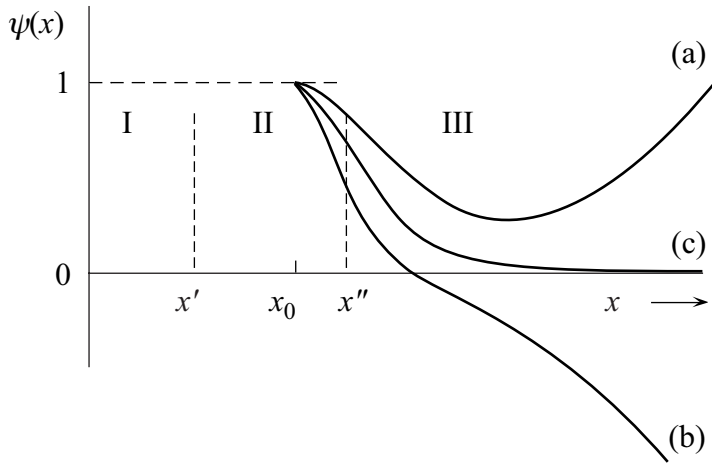


Fig. 7.2. Intentos de integración numérica de la ecuación de Schrödinger independiente del tiempo.

miento aceptable como el de la curva (c). En este caso ψ tiende asintóticamente al eje x , y cuanto más se aproxima a él, tanto menor es su concavidad, de modo que no se vuelve hacia arriba. Cuando se resuelve numéricamente el problema, $\psi_x(x_0)_{III}$ se determina por prueba y error.

Al integrar la (7.30) desde x_0 en la dirección *negativa* de x , en la región I sucede lo mismo que en la región III: en general ψ se dispara hacia $+\infty$ o hacia $-\infty$, salvo para un *único* valor $\psi_x(x_0)_I$ de $\psi_x(x_0)$, que da un comportamiento aceptable para $x \rightarrow -\infty$.

Usualmente tendremos $\psi_x(x_0)_I \neq \psi_x(x_0)_{III}$. Por lo tanto la (7.30) *no tiene soluciones aceptables* para el valor de E que hemos fijado, pues cualquiera sea el valor de $\psi_x(x_0)$ que se elija, ψ diverge para $x \rightarrow +\infty$, o para $x \rightarrow -\infty$, o en ambos casos.

Sin embargo, al repetir el procedimiento con diferentes valores de E , se encuentra en general que existen algunos valores *especiales* $E_1, E_2, E_3, \dots$ (ordenados de menor a mayor) para los cuales se cumple que $\psi_x(x_0)_I = \psi_x(x_0)_{III}$, y entonces la ecuación de Schrödinger independiente del tiempo *tiene soluciones aceptables* $\psi_1(x), \psi_2(x), \psi_3(x), \dots$. La Fig. 7.3 muestra (cualitativamente) el aspecto de las tres primeras soluciones aceptables; se ve que en las regiones I y III el

comportamiento de todas ellas es semejante, mientras que en la región II se observa que $\psi_1(x)$ no se anula nunca, $\psi_2(x)$ tiene un nodo, $\psi_3(x)$ tiene dos nodos, y así sucesivamente.

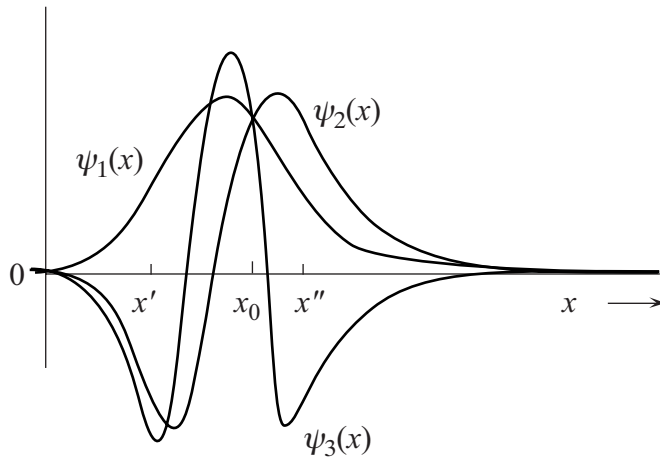


Fig. 7.3. Las tres primeras soluciones aceptables de la ecuación de Schrödinger independiente del tiempo en un caso típico.

no son infinitesimales, porque las diferencias entre las derivadas segundas de ψ en x_0 son *finitas*. Por lo tanto los valores permitidos de la energía están bien separados y forman un *conjunto discreto*.

Resulta entonces que *la energía de una partícula cuya energía potencial es independiente del tiempo está cuantificada*, pues sólo para un conjunto discreto de valores $E_1, E_2, E_3, \dots$ de la energía (que se denominan *niveles de energía*) hay soluciones aceptables de la ecuación de Schrödinger independiente del tiempo.

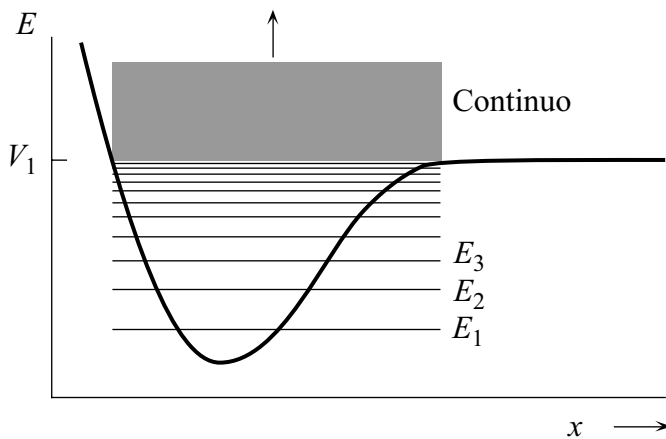


Fig. 7.4. Niveles permitidos de energía en un caso típico.

$V(x) - E < 0$, la ψ oscila y se mantiene finita para $x \rightarrow +\infty$. Pero nuestra discusión anterior muestra entonces que para *cualquier* valor de E ($E > V_1$) se puede *siempre* elegir un valor

De la Fig. 7.3 y de la ec. (7.30) se puede ver que $E_2 > E_1$. Consideremos el punto x_0 donde $\psi_1(x) = \psi_2(x)$. De la Fig. se reconoce que

$$|(\psi_2)_{xx}|_{x_0} > |(\psi_1)_{xx}|_{x_0} \quad (7.31)$$

de modo que

$$|V(x_0) - E_2| > |V(x_0) - E_1| \quad (7.32)$$

y por lo tanto $E_2 > E_1$. De manera semejante se puede ver que $E_3 > E_2$, etc.

Debe quedar claro que las diferencias de energía $(E_2 - E_1)$, $(E_3 - E_2)$, ... etc.

Lo dicho vale siempre y cuando el valor de la energía total sea tal que haya *dos* puntos de retorno. Si $V(x)$ es tal como se muestra en la Fig. 7.1, esto es cierto si $E < V_1$, donde V_1 es el valor límite que alcanza V cuando $x \rightarrow +\infty$; en este caso hay en general un número *finito*³ de niveles discretos de la energía para los cuales $E < V_1$ (Fig. 7.4).

Cuando $E > V_1$ la situación es diferente, pues hay *un único punto de retorno* (el punto x'), que divide el eje x en las regiones I ($x < x'$) y II ($x > x'$). Puesto que en la región II se cumple que

³ Puede haber excepciones cuando la energía potencial se aproxima muy lentamente al valor límite, porque en ese caso la separación entre los niveles de energía se hace muy pequeña a medida que nos aproximamos a V_1 , y si eso ocurre puede haber *infinitos* niveles discretos de energía por debajo de V_1 . Un ejemplo importante es el de la energía potencial Coulombiana, que tiende muy lentamente al valor límite $V_1 = 0$.

$\psi_x(x_0)_I$ de $\psi_x(x_0)$, tal que ψ tenga un comportamiento aceptable para $x \rightarrow -\infty$. Por lo tanto todos los valores $E > V_l$ de la energía están permitidos. Se dice entonces que los niveles de energía forman un *continuo*. Debe quedar claro que si $V(x)$ alcanza valores límite tanto a la izquierda como a la derecha, los valores permitidos de la energía forman un continuo para todas las energías mayores que el menor de esos valores límite.

Resumiendo, hemos llegado al siguiente resultado:

Cuantificación de la energía en la teoría de Schrödinger:

- si la relación entre $V(x)$ y la energía total E es tal que una partícula clásica está confinada (ligada), los niveles de energía permitidos son discretos.
- si esa relación es tal que una partícula clásica no está ligada, la energía total puede tomar cualquier valor, de modo que los niveles de energía forman un continuo.

El conjunto de todos los niveles permitidos de la energía se denomina *espectro* de energía. En general, entonces, el espectro de energía de una partícula cuya energía potencial es $V(x)$ tiene una parte discreta y una parte continua. Sólo cuando $V(x)$ crece sin límite para $x \rightarrow \pm\infty$ el espectro es completamente discreto. Veremos que las conclusiones a que hemos llegado a partir de nuestra discusión cualitativa se verifican en todos los casos específicos que estudiaremos.

Valores esperados y operadores diferenciales

Consideremos una partícula y su función de onda $\Psi(x,t)$. Conforme al postulado de Born, la probabilidad de encontrar la partícula en el intervalo entre x y $x + dx$ es:

$$P(x,t)dx = \Psi^*(x,t)\Psi(x,t)dx \quad (7.33)$$

Por consiguiente el valor medio de una serie de mediciones de x (realizadas al mismo tiempo t), que designaremos con $\bar{x}$ y llamaremos el *valor esperado* de x , está dado por

$$\bar{x} = \int_{-\infty}^{+\infty} xP(x,t)dx = \int_{-\infty}^{+\infty} \Psi^*(x,t)x\Psi(x,t)dx \quad (7.34)$$

En la (7.34) hemos escrito los factores del integrando en un orden particular, por razones que resultarán claras en breve. Es fácil ver que el valor esperado de cualquier función $f(x,t)$ de x y t (por ejemplo $V(x,t)$) está dado por

$$\overline{f(x,t)} = \int_{-\infty}^{+\infty} \Psi^*(x,t)f(x,t)\Psi(x,t)dx \quad (7.35)$$

puesto que las mediciones que se realizan para obtener el valor medio se hacen en el mismo tiempo t .

La coordenada x y funciones como $V(x,t)$ son ejemplos de *variables dinámicas* que caracterizan el estado de la partícula. Otras variables dinámicas de interés son el impulso p y la energía total E (y en el caso de movimiento en tres dimensiones, las tres componentes del impulso $\mathbf{p}$ y del momento angular $\mathbf{L}$). Consideremos el impulso. Por analogía con (7.34) y (7.35) podríamos pensar que el valor esperado de p se calcule como

$$\bar{p} = \int_{-\infty}^{+\infty} \Psi^*(x,t) p \Psi(x,t) dx \quad (7.36)$$

El problema de la (7.36) es que para calcular la integral se precisaría expresar p en términos de x y t . En un problema clásico, una vez resueltas las ecuaciones del movimiento se puede siempre expresar p como función de x y t , pero *no* podemos hacer lo mismo en el caso cuántico, pues el principio de incerteza establece que p y x no se pueden conocer simultáneamente con exactitud. Luego hay que encontrar otra forma de expresar el integrando en términos de x y t .

La forma que buscamos se puede inferir considerando el caso de una partícula libre, cuya función de onda (ec. 7.8) es⁴

$$\Psi_f = e^{i(kx - \omega t)} \quad (7.37)$$

Si derivamos respecto de x obtenemos $\partial \Psi_f / \partial x = ik \Psi_f$, que recordando la relación $k = p / \hbar$ se puede escribir como

$$-i\hbar \frac{\partial}{\partial x} \Psi_f = p \Psi_f \quad (7.38)$$

La (7.38) indica que existe una relación entre la variable dinámica p y el *operador diferencial* $-i\hbar \partial / \partial x$:

$$p \leftrightarrow -i\hbar \frac{\partial}{\partial x} \quad (7.39)$$

Una relación semejante existe entre la energía E y el operador $i\hbar \partial / \partial t$

$$E \leftrightarrow i\hbar \frac{\partial}{\partial t} \quad (7.40)$$

pues la (7.37) nos dice que

$$i\hbar \frac{\partial}{\partial t} \Psi_f = E \Psi_f \quad (7.41)$$

Estas asociaciones no están limitadas al caso de la partícula libre. Consideremos por ejemplo la relación (7.3) que define la energía E :

$$E = p^2 / 2m + V \quad (7.42)$$

Si en esta relación reemplazamos las variables dinámicas por los operadores asociados resulta

⁴ El lector podría objetar que usemos la función (7.37), que no es de cuadrado integrable. La objeción es correcta: el procedimiento riguroso es considerar un paquete de ondas de cuadrado integrable, pero casi monocromático. El cálculo es más laborioso, pero se llega al mismo resultado. Veremos en breve una forma más cómoda de proceder, basada en normalizar funciones como la (7.37) con la delta de Dirac.

$$i\hbar \frac{\partial}{\partial t} = -\frac{\hbar^2}{2m} \frac{\partial^2}{\partial x^2} + V(x,t) \quad (7.43)$$

Esta ecuación entre operadores implica que para toda función de onda $\Psi(x,t)$ se cumple que

$$i\hbar \frac{\partial \Psi(x,t)}{\partial t} = -\frac{\hbar^2}{2m} \frac{\partial^2 \Psi(x,t)}{\partial x^2} + V(x,t)\Psi(x,t) \quad (7.44)$$

que es la ecuación de Schrödinger (7.9). Esto implica que *postular las relaciones (7.39) y (7.40) entre las variables dinámicas y los correspondientes operadores es equivalente a postular la ecuación de Schrödinger*. Esta fue en realidad la vía que siguió Schrödinger para obtener su ecuación, y es un método poderoso para obtener la ecuación de Schrödinger en casos más complicados que los que discutimos hasta ahora.

Aplicando la (7.39) en la (7.36) obtenemos la expresión del valor esperado del impulso:

$$\bar{p} = -i\hbar \int_{-\infty}^{+\infty} \Psi^*(x,t) \frac{\partial \Psi(x,t)}{\partial x} dx \quad (7.45)$$

y análogamente usando la (7.40) se obtiene el valor esperado de la energía

$$\bar{E} = i\hbar \int_{-\infty}^{+\infty} \Psi^*(x,t) \frac{\partial \Psi(x,t)}{\partial t} dx \quad (7.46)$$

Usando la (7.43) obtenemos también

$$\bar{E} = i\hbar \int_{-\infty}^{+\infty} \Psi^*(x,t) \left(-\frac{\hbar^2}{2m} \frac{\partial^2}{\partial x^2} + V(x,t) \right) \Psi(x,t) dx \quad (7.47)$$

En general, podemos calcular el valor esperado de cualquier variable dinámica $f(x,p,t)$ como

$$\overline{f(x,p,t)} = \int_{-\infty}^{+\infty} \Psi^*(x,t) \hat{f}(x, -i\hbar \partial / \partial x, t) \Psi(x,t) dx \quad (7.48)$$

donde el operador $\hat{f}(x,p,t)$ se obtiene reemplazando⁵ p por el operador $-i\hbar \partial / \partial x$ en la función $f(x,p,t)$ que especifica la variable dinámica f .

Estos resultados muestran que la función de onda, además de determinar la densidad de probabilidad y la corriente de probabilidad como ya vimos, contiene mucha más información, pues por medio de la (7.48) nos permite calcular el valor esperado de *cualquier* variable dinámica. En realidad, la función de onda contiene *toda* la información que podemos llegar a obtener acerca de la partícula, habida cuenta del principio de incerteza.

⁵ Hay que tomar ciertas precauciones, que se verán oportunamente, cuando en la expresión de f aparecen productos de dos o más variables dinámicas, ya que en general el producto de operadores no es conmutativo.

Propiedades matemáticas de operadores lineales en espacios funcionales

Como acabamos de ver, en la Mecánica Cuántica las variables dinámicas se representan por medio de operadores. Por lo tanto conviene repasar algunas propiedades de los operadores, que vamos a presentar sin pretensión de rigor matemático.

Al calcular la norma de la función de onda y los valores esperados de variables dinámicas nos encontramos con expresiones del tipo

$$\int_D \xi^* \eta d\tau = (\xi, \eta) \quad (7.49)$$

que se suelen llamar *producto interno* o *producto escalar* de las funciones ξ y η , por analogía con el producto escalar ordinario de dos vectores. La integral se efectúa sobre el dominio D de las variables de que dependen las funciones. Por ejemplo, en las (7.45), (7.48) la variable de integración es x y la integral va de $-\infty$ a $+\infty$; en tres dimensiones podríamos tener $d\tau = dx dy dz$ o bien $d\tau = r^2 \sin\theta dr d\theta d\phi$ (en este caso la integral sobre r va de 0 a ∞ , la integral sobre θ va de 0 a π y la integral sobre ϕ va de 0 a 2π), u otras expresiones según el sistema de coordenadas usado. Para tener expresiones compactas usaremos la notación (ξ, η) para indicar el producto escalar. Daremos por sobreentendido que las funciones $\xi, \eta, \zeta, \varphi, \dots$ etc. son de cuadrado integrable y satisfacen oportunas condiciones en el contorno de D , y por ahora supondremos que el dominio D de integración en (7.52) es finito. Más adelante veremos como proceder cuando D es infinito.

Producto escalar, norma y ortogonalidad

El *producto escalar* o *producto interno* (ξ, η) de dos funciones ξ y η (en este orden) es un número complejo y se cumple que

$$(\eta, \xi) = (\xi, \eta)^* \quad , \quad (\eta, a\xi + b\varphi) = a(\eta, \xi) + b(\eta, \varphi) \quad (7.50)$$

donde a, b son dos números complejos cualesquiera. La *norma* N de una función se define como

$$N = (\xi, \xi) \quad (7.51)$$

La norma es siempre real y positiva, salvo para la función nula $\xi = 0$ para la cual $N = 0$. Una función se dice *normalizada* cuando su norma es 1:

$$(\xi, \xi) = 1 \quad (7.52)$$

Dos funciones ξ y η se dicen *ortogonales* cuando su producto escalar es nulo:

$$(\xi, \eta) = 0 \quad (7.53)$$

Sistema ortonormal

Una serie de Fourier es un caso especial de *expansión de una función arbitraria en serie de funciones ortogonales*. Pasaremos revista a las propiedades generales de esas expansiones.

Sea un conjunto finito (o infinito numerable) de funciones

$$\varphi_1, \varphi_2, \varphi_3, \dots \quad (7.54)$$

definidas en D . El conjunto se dice *ortonormal* (ortogonal y normalizado) si se cumple

$$(\varphi_i, \varphi_j) = \delta_{ij} \quad \text{para todo } i, j \quad (7.55)$$

Queremos representar una función f arbitraria de cuadrado integrable definida en D mediante una serie de las φ_j . Primero supongamos que la serie tiene un número finito n de términos de la forma

$$f'_n = \sum_{j=1}^n a_j \varphi_j \quad (7.56)$$

donde los a_j son coeficientes numéricos. Usaremos como criterio de la “bondad” de la aproximación el de minimizar la norma de la diferencia entre f y f'_n , (el error standard); pediremos entonces que

$$M_n = (f - f'_n, f - f'_n) \quad (7.57)$$

sea mínimo. Con algunos cálculos que omitimos por brevedad se encuentra que el mínimo se tiene cuando

$$a_j = (\varphi_j, f) \quad (7.58)$$

y en ese caso el error standard es

$$M_n = (f, f) - \sum_{j=1}^n |a_j|^2 \quad (7.59)$$

Puesto que por definición M_n es no negativo, se obtiene la *desigualdad de Bessel*

$$(f, f) \geq \sum_{j=1}^n |a_j|^2 \quad (7.60)$$

Si el conjunto (7.54) es infinito, podríamos esperar que al tomar n más y más grande se obtenga una serie que represente cada vez mejor a f . Esta expectativa intuitiva es correcta, siempre y cuando el conjunto ortonormal (7.54) sea *completo*. El sistema (7.54) se dice *completo* cuando para cada número positivo ε arbitrariamente pequeño existe un entero finito $n_0 = n_0(\varepsilon)$, tal que para todo $n > n_0$ se cumple que $M_n < \varepsilon$. En ese caso, se dice que la serie

$$f' = \sum_{j=1}^{\infty} a_j \varphi_j \quad , \quad (7.61)$$

con los coeficientes dados por la (7.58), *converge en la media* a f , y se escribe

$$f = \sum_{j=1}^{\infty} a_j \varphi_j \quad (7.62)$$

La igualdad (7.62) se debe interpretar que significa que la serie del miembro derecho es igual a *f casi en todas partes*⁶ del dominio D . En general no es trivial demostrar que un dado conjunto ortonormal es completo. Nosotros vamos a *suponer*, cuando sea necesario, que los sistemas ortonormales que nos interesan son efectivamente completos.

El método de ortogonalización de Schmidt

Hasta ahora hemos supuesto que el sistema ortonormal φ_j está asignado. En realidad para hacer expansiones alcanza con que las funciones del sistema sean linealmente independientes. Pero es más cómodo trabajar con un sistema ortonormal. Por lo tanto es útil tener un procedimiento sistemático para ortogonalizar un conjunto de funciones f_j linealmente independientes. A continuación describimos el método de Schmidt.

Sea N_1 la norma de f_1 . La primera función ortonormal es entonces

$$\varphi_1 = \frac{1}{\sqrt{N_1}} f_1 \quad (7.63)$$

La segunda función ortonormal es

$$\varphi_2 = \frac{1}{\sqrt{N_2}} [f_2 - \varphi_1(\varphi_1, f_2)] \quad (7.64)$$

donde N_2 es la norma de la función encerrada por el corchete de la (7.64).

La tercera función ortonormal es

$$\varphi_3 = \frac{1}{\sqrt{N_3}} [f_3 - \varphi_1(\varphi_1, f_3) - \varphi_2(\varphi_2, f_3)] \quad (7.65)$$

donde N_3 es la norma de la función encerrada por el corchete de la (7.65).

Continuando de esta manera se pueden ortogonalizar y normalizar sucesivamente todas las f_j .

Función delta de Dirac y relaciones de clausura

La función *delta de Dirac* $\delta(x - a) = \delta(a - x)$ es una función par definida por:

$$\bullet \quad \delta(x - a) = 0 \quad \text{para } x \neq a \quad (7.66)$$

$$\bullet \quad \int_b^c \delta(x - a) dx = \begin{cases} 1 & \text{si } a \text{ está dentro del intervalo } (b, c) \\ 0 & \text{si } a \text{ está fuera del intervalo } (b, c) \end{cases} \quad (7.67)$$

La función delta de Dirac se puede definir como la forma límite de una curva con un pico agudo que se hace más y más alto y más y más angosto, de forma tal que el área debajo de la curva se mantiene constante. Por ejemplo se puede usar una Gaussiana, y entonces

$$\delta(x - a) = \lim_{\varepsilon \rightarrow 0} \frac{1}{\varepsilon \sqrt{\pi}} e^{-\frac{(x-a)^2}{\varepsilon^2}} \quad (7.68)$$

⁶ Se dice que dos funciones f_1 y f_2 son iguales *casi en todas partes* cuando difieren en un conjunto de puntos de medida nula.

Sea $f(x)$ es una función arbitraria de x definida en el entorno de $x = a$ (el intervalo de integración incluye $x = a$); es fácil entonces demostrar las siguientes propiedades de integrales que contienen funciones delta:

- (I):

$$\int f(x)\delta(x-a)dx = f(a) \quad (7.69)$$

- (II) si $\delta^{(n)}(x-a)$ indica la derivada n -ésima de $\delta(x-a)$ respecto de su argumento, entonces

$$\int f(x)\delta^{(n)}(x-a)dx = (-1)^n f^{(n)}(a) \quad (7.70)$$

- (III) si $y = y(x)$, entonces

$$\int f(x)\delta(y(x))dx = \sum_i \frac{f(x_i)}{|y'(x_i)|} \quad (7.71)$$

donde los puntos x_i son las raíces reales de $y(x) = 0$ en el intervalo de integración. La ec. (7.71) es equivalente a decir que

$$\delta(y(x)) = \sum_i \frac{\delta(x-x_i)}{|y'(x_i)|} \quad (7.72)$$

La función delta se puede definir en más de una dimensión. Por ejemplo

$$\delta(\mathbf{r} - \mathbf{a}) = \delta(x - a_x)\delta(y - a_y)\delta(z - a_z) \quad (7.73)$$

Retornamos ahora a nuestra discusión de las expansiones ortonormales. Escribimos la (7.62) para una función arbitraria f :

$$f(x) = \sum_{j=1}^{\infty} a_j \varphi_j(x) = \sum_{j=1}^{\infty} (\varphi_j^*, f) \varphi_j(x) = \int_D \left\{ \sum_{j=1}^{\infty} \varphi_j(x) \varphi_j^*(x') f(x') dx' \right\} \quad (7.74)$$

donde hemos intercambiado el orden de suma e integral. Comparando la (7.74) con la (7.69) obtenemos que

$$\sum_j \varphi_j(x) \varphi_j^*(x') = \delta(x' - x) \quad (7.75)$$

La (7.75) se denomina *relación de clausura* o de *completitud*, y se cumple para todo sistema ortonormal completo. La (7.75) se denomina a veces “teorema de expansión” debido al siguiente teorema: si f y g son funciones arbitrarias y las φ_j forman un sistema completo, se cumple

$$(f, g) = \sum_j (f, \varphi_j)(\varphi_j, g) \quad (7.76)$$

La demostración se basa en la definición del producto escalar y en la relación de completitud (7.75).

Integrales de Fourier

Hasta ahora hemos considerado expansiones en términos de un sistema ortonormal φ_j que consiste de una infinidad numerable de funciones. Cuando el dominio D es infinito el índice discreto j se convierte en un parámetro continuo k . También puede suceder que aparezca una combinación de valores discretos e intervalos continuos del índice. En general, la transición matemática rigurosa desde un dominio finito con un índice discreto a un dominio infinito con un parámetro continuo no es trivial. Nosotros indicaremos los pasos formales para el caso de la integral de Fourier, sin pretensión de rigor⁷.

Consideremos la integral de Fourier

$$f(x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} F(k) e^{ikx} dk \quad (7.77)$$

y la transformación inversa

$$F(k) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} f(x) e^{-ikx} dx \quad (7.78)$$

En estas fórmulas las variables continuas k y x están en pie de igualdad, pero para establecer una conexión con nuestra anterior discusión podemos observar que la (7.77) es análoga a la expansión en serie (7.62) mientras que la (7.78) es análoga a la expresión del coeficiente (7.58) de la serie, donde la función ortonormal $\varphi_j(x)$ se ha reemplazado por

$$\varphi(k, x) = \frac{1}{\sqrt{2\pi}} e^{ikx} \quad (7.79)$$

Es fácil verificar que la *ortonormalidad* de las funciones (7.79) se expresa como

$$\frac{1}{2\pi} \int_{-\infty}^{+\infty} e^{i(k-k')x} dx = \delta(k - k') \quad (7.80)$$

Análogamente, combinando la (7.77) con la (7.78) obtenemos la *relación de clausura* o de *completitud*:

$$\frac{1}{2\pi} \int_{-\infty}^{+\infty} e^{ik(x-x')} dk = \delta(x - x') \quad (7.81)$$

Un resultado importante de las integrales de Fourier es el teorema de Parseval, que establece que si f y g son dos funciones cuyas transformadas de Fourier son F y G , se cumple que

$$\int_{-\infty}^{+\infty} f^*(x) g(x) dx = \int_{-\infty}^{+\infty} F^*(k) G(k) dk \quad (7.82)$$

⁷ El lector interesado en el detalle matemático puede consultar, por ejemplo, el libro de E. C. Titchmarsh *Introduction to the Theory of Fourier Integrals* Oxford Univ. Press, 1948.

Una demostración no rigurosa de la (7.82) se obtiene sustituyendo las integrales (7.77) en el primer miembro de la (7.82), intercambiando el orden de integración, y usando la relación de ortonormalidad (7.80).

Si en (7.82) ponemos $g = f$ resulta

$$\int_{-\infty}^{+\infty} f^*(x)f(x)dx = \int_{-\infty}^{+\infty} F^*(k)F(k)dk \quad (7.83)$$

que muestra que la norma de $f(x)$ en el dominio infinito de x es igual a la norma de $F(k)$ en el dominio infinito de k . Esta forma especial del teorema de Parseval es el equivalente para la integral de Fourier de la (des)igualdad de Bessel (pues la ec. (7.59) se convierte en una igualdad para un sistema completo).

Este resultado justifica emplear funciones que no se pueden normalizar a la unidad (como las ondas planas) para construir paquetes. En efecto, si las funciones del sistema ortonormal se normalizan a la delta de Dirac, la (7.83) garantiza que el paquete de ondas está normalizado si $F(k)$ es de cuadrado integrable.

Operadores lineales

Un operador A se define como un ente que actúa sobre una función ξ para producir otra función η :

$$\eta = A\xi \quad (7.84)$$

Nosotros estamos interesados únicamente en operadores lineales, que cumplen las condiciones

$$A(c\xi) = cA(\xi) \quad , \quad A(\xi + \eta) = A\xi + A\eta \quad (7.85)$$

El producto de dos operadores A y B se define como

$$AB\xi = A(B\xi) \quad (7.86)$$

Debe notarse que en general

$$AB\xi \neq BA\xi \quad (7.87)$$

a menos que A y B conmuten. El *conmutador* de dos operadores se indica con $[A, B]$ y se define como

$$[A, B] = AB - BA \quad (7.88)$$

Núcleo de un operador lineal

Un operador lineal K definido por

$$g = Kf \quad (7.88)$$

se puede escribir en forma explícita como

$$g(x) = \int_D k(x, x') f(x') dx' \quad (7.89)$$

La función $k(x, x')$ se denomina *núcleo* del operador K . Pese a que la (7.89) tiene el aspecto de una ecuación integral lineal que no incluye operadores diferenciales, es fácil ver que con la ayuda de la función delta de Dirac se obtienen operadores diferenciales. Por ejemplo si

$$K = \frac{d^2}{dx^2} + \alpha^2 \quad (7.90)$$

entonces

$$k(x, x') = \left(\frac{d^2}{dx^2} + \alpha^2 \right) \delta(x - x') \quad (7.91)$$

Operador adjunto

El operador $K^\dagger$ adjunto de K se define como el operador que cumple

$$(\xi, K\eta) = (K^\dagger \xi, \eta) \quad (7.92)$$

para todo par de funciones ξ y η arbitrarias. Usando la (7.89) encontramos que el *núcleo adjunto* $k^\dagger(x, x')$ está dado por

$$k^\dagger(x, x') = [k(x', x)]^* \quad (7.93)$$

Operadores Hermitianos

Un operador K para el cual $K^\dagger = K$ se dice *Hermitiano* o *autoadjunto*. Para un operador autoadjunto el producto escalar (f, Kf) es real para cualquier f , pues:

$$(f, Kf) = (K^\dagger f, f) = (Kf, f) = (f, Kf)^* \quad (K \text{ Hermitiano}) \quad (7.94)$$

Esta propiedad de los operadores Hermitianos es de fundamental importancia, pues una hipótesis básica de la Mecánica Cuántica es que los valores esperados de las magnitudes observables se expresan en la forma (f, Kf) y son siempre reales. Por lo tanto *todos los observables deben estar representados por operadores Hermitianos*.

Un operador que cumple $K^\dagger = -K$ se dice *anti-Hermitiano*.

Autovalores y autofunciones de un operador Hermitiano

El problema de autovalores para un operador lineal K se expresa como

$$K\varphi = \lambda\varphi \quad (7.95)$$

Las soluciones de (7.95) existen sólo para ciertos valores particulares de λ , indicados con λ_j , que se denominan *autovalores* o *valores propios* del operador K . Las correspondientes soluciones $\varphi_j(x)$ se denominan *autofunciones* o *funciones propias* de K . Un ejemplo lo constituyen las

soluciones de la ecuación de Schrödinger independiente del tiempo, que son las autofunciones del operador que representa la energía total de un sistema cuántico. El conjunto de autovalores de K se denomina el *espectro* de K .

Por lo dicho, los autovalores y autofunciones de K cumplen que

$$K\varphi_j = \lambda_j\varphi_j \quad (7.96)$$

Una propiedad muy importante de los operadores Hermitianos es que *todos sus autovalores son reales*. Para demostrarlo basta tomar el producto escalar de la (7.96) por φ_j :

$$(\varphi_j, K\varphi_j) = \lambda_j(\varphi_j, \varphi_j) \quad (7.97)$$

Puesto que $(\varphi_j, K\varphi_j)$ y (φ_j, φ_j) son ambos reales, λ_j es real. Vale también la propiedad recíproca, esto es, un operador todos cuyos autovalores son reales es Hermitiano.

Vamos a mostrar ahora que las autofunciones que corresponden a diferentes autovalores son ortogonales. Para ver esto, consideremos dos autofunciones φ_i y φ_j que satisfacen

$$K\varphi_i = \lambda_i\varphi_i \quad , \quad K\varphi_j = \lambda_j\varphi_j \quad (7.98)$$

Ahora tomamos el producto escalar de φ_j (a la izquierda) por la primera de estas ecuaciones y el producto escalar de la segunda ecuación por φ_i (a la derecha) y restamos:

$$(\varphi_j, K\varphi_i) - (K\varphi_j, \varphi_i) = (\lambda_i - \lambda_j^*)(\varphi_j, \varphi_i) \quad (7.99)$$

Puesto que K es Hermitiano el primer miembro de la (7.99) se anula y como λ_j es real obtenemos

$$0 = (\lambda_i - \lambda_j)(\varphi_j, \varphi_i) \quad (7.100)$$

Las condiciones de ortogonalidad (7.100) muestran que las autofunciones correspondientes a diferentes autovalores son ortogonales. En el caso que se presente degeneración, esto es que a un autovalor le correspondan dos o más autofunciones, siempre se las puede ortogonalizar por el método que indicamos antes.

En consecuencia de lo anterior, *a partir de las autofunciones de un operador Hermitiano es siempre posible construir un sistema ortonormal* (de ser necesario ortogonalizando las autofunciones correspondientes a los autovalores degenerados). Surge entonces el problema de la completitud de ese sistema. La completitud puede ser difícil de probar rigurosamente para los operadores de la Mecánica Cuántica, aunque por razones físicas parece plausible que se cumpla y no se conocen hasta ahora excepciones. Por lo tanto nosotros vamos a suponer que *toda variable dinámica u observable está representado por un operador Hermitiano cuyas autofunciones forman un sistema ortonormal completo*.

8. EL FORMALISMO DE LA MECÁNICA CUÁNTICA

Introducción

En el Capítulo 7 vimos que el estado de un sistema cuántico se describe por medio de la función de onda $\Psi(\mathbf{r}, t)$, y que las variables dinámicas se representan mediante operadores Hermitianos, cuyos autovalores son siempre reales y cuyas autofunciones forman un sistema ortonormal completo. La función de onda contiene toda la información de interés físico acerca del sistema, en el sentido que a partir de ella se puede calcular el valor esperado de cualquier variable dinámica. Los operadores que representan las variables dinámicas se construyen a partir de sus expresiones clásicas¹ mediante ciertas prescripciones, que se introducen en la teoría como postulados. Veremos ahora algunas consecuencias de estos hechos, así como ciertas extensiones del formalismo. Volvamos por un momento a las soluciones de la ecuación de Schrödinger independiente del tiempo. Dicha ecuación expresa el problema de autovalores del operador Hamiltoniano² $\hat{H}$ que representa la energía total del sistema en términos de las variables dinámicas. Para una partícula que se mueve a lo largo del eje x en un campo de fuerzas que derivan de una energía potencial $V(x)$ este operador es

$$\hat{H} = \frac{\hat{p}^2}{2m} + V(\hat{x}, t) = -\frac{\hbar^2}{2m} \frac{\partial^2}{\partial x^2} + V(x, t) \quad (8.1)$$

Para casos más complicados hemos enunciado reglas generales que nos permiten construir la expresión apropiada de $\hat{H}$. Es fácil verificar explícitamente que el operador (8.1) es Hermitiano (si la función de onda cumple adecuadas condiciones de contorno); por otra parte los resultados de nuestro estudio cualitativo implican que $\hat{H}$ es Hermitiano pues sus autovalores son reales. El conjunto de las autofunciones de $\hat{H}$ permite formar entonces un sistema ortonormal completo

$$\Psi_j(x, t) = e^{-iE_j t / \hbar} \psi_j(x) \quad (8.2)$$

Aquí, el índice j toma valores discretos para las autofunciones correspondientes a la parte discreta del espectro de energía, y es un parámetro continuo $j \equiv E$ para las autofunciones de la parte continua del espectro de E , que escribiremos como

$$\Psi_E(x, t) = e^{-iEt / \hbar} \psi_E(x) \quad (8.3)$$

Las autofunciones de la parte discreta del espectro cumplen la condición de ortonormalidad

$$\int_{-\infty}^{+\infty} \Psi_i^*(x, t) \Psi_j(x, t) dx = \int_{-\infty}^{+\infty} \psi_i^*(x) \psi_j(x) dx = \delta_{ij} \quad (8.4)$$

mientras que las autofunciones del continuo están normalizadas a la función delta de Dirac:

¹ Esto es cierto cuando se trata de observables que tienen un análogo clásico. Veremos más adelante que será necesario introducir en la Mecánica Cuántica nuevos observables que no corresponden a ninguna variable dinámica clásica, como el spin del electrón.

² En esta Sección indicamos con $\hat{}$ los operadores, para distinguirlos de las variables dinámicas que representan.

$$\int_{-\infty}^{+\infty} \Psi_{E'}^*(x,t) \Psi_E(x,t) dx = \int_{-\infty}^{+\infty} \psi_{E'}^*(x) \psi_E(x) dx = \delta(E' - E) \quad (8.5)$$

de modo semejante a la $\varphi(k, x) = (2\pi)^{-1/2} e^{ikx}$ en el caso de la transformada de Fourier.

Propiedades de las funciones de onda y de las autofunciones de la energía

Gracias a la completitud del sistema (8.2), toda solución $\Psi(x,t)$ de la ecuación de Schrödinger

$$i\hbar \frac{\partial \Psi(x,t)}{\partial t} = -\frac{\hbar^2}{2m} \frac{\partial^2 \Psi(x,t)}{\partial x^2} + V(x)\Psi(x,t) \quad (8.6)$$

se puede expandir en términos de las Ψ_j :

$$\begin{aligned} \Psi(x,t) &= \sum_{\text{discreto}} a_j \Psi_j(x,t) + \int_{\text{continuo}} a(E) \Psi_E(x,t) dE \\ &= \sum_{\text{discreto}} a_j e^{-iE_j t/\hbar} \psi_j(x) + \int_{\text{continuo}} a(E) e^{-iEt/\hbar} \psi_E(x) dE \end{aligned} \quad (8.7)$$

En virtud de la ortonormalidad de las autofunciones, los coeficientes de (8.7) están dados por

$$a_j = (\Psi_j, \Psi) = \int_{-\infty}^{+\infty} \Psi_j^*(x,t) \Psi(x,t) dx = e^{iE_j t/\hbar} \int_{-\infty}^{+\infty} \psi_j^*(x) \Psi(x,t) dx \quad (8.8)$$

para la parte discreta del espectro y por

$$a(E) = (\Psi_E, \Psi) = \int_{-\infty}^{+\infty} \Psi_E^*(x,t) \Psi(x,t) dx = e^{iEt/\hbar} \int_{-\infty}^{+\infty} \psi_E^*(x) \Psi(x,t) dx \quad (8.9)$$

para la parte continua. Notar que en general los coeficientes a_j y $a(E)$ dependen de t .

Usando las condiciones de ortonormalidad (8.3) y (8.4) se encuentra que la norma de $\Psi(x,t)$ es

$$(\Psi, \Psi) = \int_{-\infty}^{+\infty} \Psi^*(x,t) \Psi(x,t) dx = \sum_{\text{discreto}} a_j^* a_j + \int_{\text{continuo}} a(E)^* a(E) dE \quad (8.10)$$

Es fácil verificar que la norma de Ψ no depende de t como ya habíamos demostrado en el Capítulo anterior. En adelante vamos a suponer que Ψ está normalizado a la unidad:

$$\int_{-\infty}^{+\infty} \Psi^*(x,t) \Psi(x,t) dx = \sum_{\text{discreto}} a_j^* a_j + \int_{\text{continuo}} a(E)^* a(E) dE = 1 \quad (8.11)$$

Usando la expansión (8.7) podemos ver que la densidad de probabilidad se expresa como

$$\begin{aligned}
P(x, t) &= \Psi(x, t)^* \Psi(x, t) \\
&= \sum_{j \text{ discreto}} a_j^* a_j \psi_j^*(x) \psi_j(x) + \sum_{j, k \text{ discreto}} (1 - \delta_{kj}) a_k^* a_j e^{-i(E_j - E_k)t/\hbar} \psi_k^*(x) \psi_j(x) + \\
&\quad \sum_{j \text{ discreto}} \int_{E \text{ continuo}} \left(a_j^* a(E) e^{-i(E - E_j)t/\hbar} \psi_j^*(x) \psi_E(x) + \text{c.c.} \right) dE + \\
&\quad \int_{E \text{ continuo}} a(E)^* a(E) \psi_E^*(x) \psi_E(x) dE + \\
&\quad \iint_{E, E' \text{ continuo}} [1 - \delta(E' - E)] a(E')^* a(E) e^{i(E' - E)t/\hbar} \psi_{E'}^*(x) \psi_E(x) dE' dE
\end{aligned} \tag{8.12}$$

donde c.c. indica el complejo conjugado de la expresión precedente. En general, la densidad de probabilidad depende del tiempo. Sin embargo, en el caso particular en que la partícula tiene una energía bien definida, de modo que su función de onda es una de las autofunciones (8.2) de la energía, la densidad de probabilidad es independiente del tiempo:

$$P = P(x) = \Psi_j^*(x, t) \Psi_j(x, t) = \psi_j^*(x) \psi_j(x) \tag{8.13}$$

Por ese motivo esos estados se denominan *estados estacionarios*. Es fácil verificar que en un estado estacionario la corriente de probabilidad es nula cuando ψ es real.

Calculemos el valor esperado de la energía en un estado descrito por la función de onda $\Psi(x, t)$. Usando la expansión (8.7) de $\Psi(x, t)$ en términos de las autofunciones de $\hat{H}$ resulta

$$\bar{E} = (\Psi, \hat{H} \Psi) = i\hbar \int_{-\infty}^{+\infty} \Psi^*(x, t) \frac{\partial \Psi(x, t)}{\partial t} dx = \sum_{\text{discreto}} E_j a_j^* a_j + \int_{\text{continuo}} E a(E)^* a(E) dE \tag{8.14}$$

Observando la (8.14) podemos interpretar que $a_j^* a_j$ es la probabilidad que al efectuar una medida de la energía obtengamos un valor igual a E_j (esto es, la probabilidad que el sistema se encuentre en el estado discreto j) y que $a(E)^* a(E) dE$ es la probabilidad que al medir la energía obtengamos un valor comprendido en el intervalo entre E y $E + dE$ (esto es, la probabilidad que el sistema se encuentre un estado del continuo con energía entre E y $E + dE$).

Se debe notar que si Ψ no es un estado estacionario las a_j y $a(E)$ son funciones del tiempo; sin embargo es fácil verificar que si $V = V(x)$,

$$\frac{d\bar{E}}{dt} = \sum_{\text{discreto}} E_j \frac{d(a_j^* a_j)}{dt} + \int_{\text{continuo}} E \frac{d(a(E)^* a(E))}{dt} dE = 0 \tag{8.15}$$

de modo que el valor esperado de la energía se conserva, como ocurre clásicamente.

Observando las (8.11) y (8.14) se aclara porqué las autofunciones del continuo se normalizan con la delta de Dirac (ec. (8.4)): se asegura así que se cumpla con el postulado de Born.

Para concluir esta discusión es oportuno hacer un comentario acerca de la *completitud* del sistema de autofunciones. En realidad no hemos demostrado que el sistema (8.2) tenga esa propiedad. Sin embargo, por razones físicas parece razonable que así sea. Esto se puede ver con el siguiente argumento. Supongamos que el sistema (8.2) no fuera completo. Eso implicaría que

existe algún estado $\Psi(x,t)$ de la partícula para el cual la norma de la diferencia entre Ψ y su expansión en términos de las Ψ_j no es nula. Esto es, la función

$$\Phi(x,t) = \Psi(x,t) - \left(\sum_{\text{discreto}} a_j \Psi_j(x,t) + \int_{\text{continuo}} a(E) \Psi_E(x,t) dE \right) \quad (8.15)$$

no sería nula “casi en todas partes”. En consecuencia $\Phi(x,t)$ representaría un estado de la partícula, solución de la ecuación de Schrödinger (8.5), para el cual existe una probabilidad no nula de encontrar la partícula en alguna parte, pero para el cual *no se puede* medir la energía. Como esto es físicamente absurdo, concluimos que $\Phi(x,t)$ debe ser nulo (casi en todas partes) y por consiguiente el sistema (8.2) debe ser completo.

En adelante, para simplificar la notación, escribiremos siempre

$$\Psi(x,t) = \sum_{j=1}^{\infty} a_j \Psi_j(x,t) = \sum_{j=1}^{\infty} a_j e^{-iE_j t/\hbar} \psi_j(x) \quad (8.16)$$

dando por *sobreentendido* que la sumatoria de la (8.16) se debe interpretar como una *suma* sobre la parte discreta del espectro y una *integral* sobre la parte continua, como en la (8.7).

Normalización en una caja

Existen maneras de evitar que las autofunciones del espectro continuo se tengan que normalizar con la delta de Dirac. Se trata de artificios matemáticos mediante los cuales el espectro continuo se transforma en discreto. Mostraremos en qué consisten para una dimensión espacial, pues la generalización a más dimensiones es obvia. La idea básica es suponer que x está limitado a un intervalo finito $(x_0, x_0 + L)$ de longitud L y que la función de onda cumple oportunas condiciones de contorno en x_0 y $x_0 + L$. Se suelen emplear dos tipos de condiciones de contorno:

- paredes rígidas: en este caso se imponen las condiciones

$$\Psi(x_0, t) = \Psi(x_0 + L, t) = 0 \quad (8.17)$$

- condiciones periódicas: se pide que

$$\Psi(x_0, t) = \Psi(x_0 + L, t) \quad (8.18)$$

Si repetimos el razonamiento cualitativo del Capítulo 7 acerca de las soluciones de la ecuación de Schrödinger independiente del tiempo, teniendo en cuenta que las soluciones tienen que cumplir las condiciones $\psi(x_0) = \psi(x_0 + L) = 0$ (caso (8.17)) o bien $\psi(x_0) = \psi(x_0 + L)$ (caso (8.18)), es fácil ver que los autovalores de la energía son siempre discretos.

Relaciones de conmutación

En lo sucesivo omitiremos la $\hat{}$ que empleamos para distinguir los observables de los operadores que los representan, pues no puede surgir confusión. Sean dos observables F y G cualesquiera representados por los operadores Hermitianos F y G . Es fácil verificar que el adjunto del producto FG es:

$$(FG)^\dagger = G^\dagger F^\dagger = GF \quad (8.19)$$

de donde se desprende que FG es Hermitiano si y sólo si F y G conmutan. Luego, si bien x y p_x son observables, el producto xp_x no lo es, pues los operadores x e $-i\hbar\partial/\partial x$ no conmutan. En efecto, si Ψ una función cualquiera:

$$(xp_x - p_x x)\Psi = \left(-i\hbar x \frac{\partial}{\partial x} + i\hbar \frac{\partial}{\partial x} x\right)\Psi = i\hbar\Psi \quad (8.20)$$

de donde obtenemos la siguiente relación fundamental entre operadores:

$$[x, p_x] \equiv xp_x - p_x x = i\hbar \quad (8.21)$$

De manera semejante se puede mostrar que

$$[y, p_y] \equiv yp_y - p_y y = i\hbar \quad , \quad [z, p_z] \equiv zp_z - p_z z = i\hbar \quad (8.22)$$

Todos los demás productos de las coordenadas cartesianas y sus impulsos conjugados son conmutativos, esto es $xy = yx$, $xp_y = p_y x$, $p_x p_y = p_y p_x$, etc.

Autoestados de una variable dinámica

Como veremos en breve, el hecho que x y p_x no conmutan está vinculado con la relación de incerteza entre x y p_x , es decir con la imposibilidad de tener valores bien definidos de ambas variables dinámicas. Como primer paso para encontrar esa relación, vamos a determinar los estados de un sistema para los cuales una variable dinámica F representada por el operador Hermitiano F tiene un valor bien definido, es decir, los estados Ψ_f en los cuales al medir F obtenemos con certeza un cierto valor f . Es fácil verificar que para que eso ocurra, Ψ debe ser una autofunción de F . En efecto, Ψ_f debe satisfacer la condición:

$$\bar{F} = (\Psi_f, F\Psi_f) = f \quad (8.23)$$

y la condición

$$\begin{aligned} \Delta F^2 &= \overline{(F - \bar{F})^2} = (\Psi_f, (F - \bar{F})^2 \Psi_f) = (\Psi_f, (F - f)^2 \Psi_f) \\ &= ((F - f)\Psi_f, (F - f)\Psi_f) = 0 \end{aligned} \quad (8.24)$$

donde para llegar al último renglón usamos la (8.23) y el hecho que F es Hermitiano. De la (8.24) se desprende entonces que se debe cumplir

$$(F - f)\Psi_f = 0 \quad \text{esto es} \quad F\Psi_f = f\Psi_f \quad (8.25)$$

Luego Ψ_f es una autofunción de F . Los estados descriptos por autofunciones de F se suelen denominar *autoestados* de F . La (8.25) nos dice lo siguiente:

- Una variable dinámica representada por el operador Hermitiano F tiene un valor bien definido f cuando el sistema se encuentra en el autoestado de F correspondiente al autovector f .

La ecuación de Schrödinger independiente del tiempo $H\psi_n = E_n\psi_n$ no es más que un caso particular de la (8.25): las autofunciones Ψ_n son los únicos estados en los cuales el sistema tiene una energía bien definida E_n .

Puesto que el operador Hermitiano F tiene un sistema completo de autofunciones ortonormales Ψ_{f_j} , cualquier estado Ψ se puede representar como

$$\Psi = \sum_j a_j \Psi_{f_j} \quad \text{con} \quad a_j = (\Psi_{f_j}, \Psi) \quad (8.26)$$

y el valor esperado de F en el estado Ψ es

$$\bar{F} = (\Psi, F\Psi) = \sum_j f_j |a_j|^2 \quad (8.27)$$

Las conclusiones son entonces las siguientes:

- El resultado de medir una variable dinámica F de un sistema que se encuentra en un estado cualquiera Ψ es *siempre* uno de los autovalores de F .
- Si el sistema se encuentra en un estado cualquiera Ψ , la probabilidad que al medir F se obtenga el autovalor f está dada por $|a_j|^2 = |(\Psi_{f_j}, \Psi)|^2$.
- Inmediatamente después de efectuar una medición de F cuyo resultado fue el autovalor f , el sistema se encuentra en el autoestado correspondiente a f .

Mediciones simultáneas y operadores que conmutan

Puesto que un observable F tiene un valor bien definido f sólo cuando el sistema se encuentra en un autoestado de F , es evidente que:

- Dos observables F y G pueden ambos tener valores bien definidos si y solo si el sistema está en un autoestado de ambos operadores.

Cuando F y G tienen valores bien definidos para todos los autoestados de uno de ellos (por ejemplo F) se dice que son *compatibles*.

Para que dos observables sean compatibles es necesario que los operadores que los representan tengan un sistema completo de autofunciones en común. En tal caso podemos designar las autofunciones mediante los autovalores f_j de F y g_k de G , es decir

$$F\Psi_{f_j g_k} = f_j \Psi_{f_j g_k} \quad , \quad G\Psi_{f_j g_k} = g_k \Psi_{f_j g_k} \quad (8.28)$$

Si multiplicamos por F la segunda de estas ecuaciones, por G la primera y restamos, resulta

$$(FG - GF)\Psi_{f_j g_k} = 0 \quad (8.29)$$

Puesto que la (8.29) vale para todas las autofunciones de un sistema completo, se debe cumplir

$$[F, G] = FG - GF = 0 \quad (8.30)$$

Recíprocamente, si F y G conmutan hay un sistema completo de autofunciones comunes a ambos. En efecto, sea Ψ_{g_k} la autofunción de G que satisface $G\Psi_{g_k} = g_k\Psi_{g_k}$. Entonces

$$G(F\Psi_{g_k}) = F(G\Psi_{g_k}) = g_k(F\Psi_{g_k}) \quad (8.31)$$

de modo que $F\Psi_{g_k}$ es también autofunción de G correspondiente al mismo autovalor g_k . Hay dos posibilidades, según si el autovalor g_k es degenerado o no lo es.

Si g_k no es degenerado, le corresponde una única autofunción Ψ_{g_k} , luego la (8.31) implica que

$$F\Psi_{g_k} = \lambda\Psi_{g_k} \quad (8.32)$$

donde λ es un número, de modo que Ψ_{g_k} es también una autofunción de F correspondiente al autovalor λ . Si en cambio g_k es degenerado y le corresponden (digamos) n autofunciones $\Psi'_{g_k,i}$, la (8.31) implica que

$$F\Psi'_{g_k,i} = \sum_{r=1}^n \Psi'_{g_k,r} f_{ri} \quad \text{con} \quad f_{ri} = (\Psi'_{g_k,r}, F\Psi'_{g_k,i}) \quad (8.33)$$

donde las cantidades f_{ri} se denominan *elementos de matriz* del operador F en la base $\Psi'_{g_k,i}$, y por ser F Hermitiano cumplen $f_{ir} = f_{ri}^*$. Vemos entonces que en general las $\Psi'_{g_k,i}$ no son autofunciones de F . Pero es *siempre* posible construir un nuevo conjunto $\Psi_{g_k,s}$ de autofunciones de G correspondientes al autovalor g_k , de la forma

$$\Psi_{g_k,s} = \sum_{i=1}^n d_{si} \Psi'_{g_k,i} \quad (8.34)$$

y determinar los coeficientes d_{si} de modo que $\Psi_{g_k,s}$ sea también autofunción de F , es decir que

$$F\Psi_{g_k,s} = \lambda\Psi_{g_k,s} \quad (8.35)$$

En efecto, si sustituimos la (8.34) en (8.35) resulta

$$F\Psi_{g_k,s} = \sum_{i=1}^n d_{si} F\Psi'_{g_k,i} = \sum_{i=1}^n \sum_{r=1}^n d_{si} f_{ri} \Psi'_{g_k,r} = \lambda \sum_{i=1}^n d_{si} \Psi'_{g_k,i} \quad (8.36)$$

Por lo tanto se debe cumplir que

$$\sum_{i=1}^n \sum_{r=1}^n d_{si} (f_{ri} - \lambda \delta_{ri}) \Psi'_{g_k,r} = 0 \quad (8.37)$$

La (8.37) nos da un sistema de n ecuaciones lineales para las n incógnitas d_{si} :

$$\sum_{i=1}^n d_{si} (f_{ri} - \lambda \delta_{ri}) = 0 \quad , \quad s \text{ fijo} \quad , \quad r = 1, 2, \dots, n \quad (8.38)$$

Este sistema de ecuaciones tiene soluciones no triviales si, y solo si, es nulo el determinante de los coeficientes:

$$\det(f_{ri} - \lambda \delta_{ri}) = 0 \quad (8.39)$$

La ec. (8.39) es una ecuación de grado n en el autovalor λ , que se denomina *ecuación característica*. Se puede mostrar que en virtud de que $f_{ir} = f_{ri}^*$, la (8.36) tiene siempre n raíces reales que nos dan los autovalores $\lambda_1, \lambda_2, \dots, \lambda_n$ de F . Para cada autovalor λ_s , el sistema (8.38) permite calcular los coeficientes d_{sj} y por lo tanto la autofunción $\Psi_{g_k, s}$. Estas nuevas funciones son evidentemente autofunciones simultáneas de F y G . Queda entonces demostrado que:

- La condición necesaria y suficiente para que dos observables sean compatibles es que conmuten los operadores que los representan.
- Cuando dos observables F y G están representados por operadores que conmutan, son compatibles y existe un sistema completo de autofunciones comunes a ambos.

Las relaciones de incerteza de Heisenberg

Acabamos de ver que cuando dos observables conmutan se pueden especificar simultáneamente. En cambio,

- Cuando dos observables F y G no conmutan no pueden ambos tener valores bien definidos simultáneamente.

La medida de la inevitable falta de precisión en el conocimiento de F y G está determinada por su conmutador:

$$[F, G] = FG - GF = iC \quad (8.40)$$

Es fácil verificar que si F y G son Hermitianos, C también es Hermitiano. Ahora vamos a demostrar que en cualquier estado $\Psi(x, t)$ las incertezas

$$\Delta F = \left[\overline{(F - \bar{F})^2} \right]^{1/2} \quad \text{y} \quad \Delta G = \left[\overline{(G - \bar{G})^2} \right]^{1/2} \quad (8.41)$$

de F y G cumplen la relación

$$\Delta F \Delta G \geq \frac{1}{2} | \bar{C} | \quad (8.42)$$

Para ver esto conviene primero demostrar que si u y w son dos funciones *cualesquiera*, se cumple la *desigualdad de Schwarz*:

$$(u, u)(w, w) \geq \frac{1}{4} [(u, w) + (w, u)]^2 \quad (8.43)$$

En efecto, sea λ un parámetro real. Entonces si $w \neq -\lambda u$ se cumple que

$$0 < (\lambda u + w, \lambda u + w) = \lambda^2 (u, u) + \lambda [(u, w) + (w, u)] + (w, w) \quad (8.44)$$

Por lo tanto la ecuación

$$\lambda^2 (u, u) + \lambda [(u, w) + (w, u)] + (w, w) = 0 \quad (8.45)$$

no tiene raíces reales y por consiguiente su discriminante es negativo, esto es

$$[(u, w) + (w, u)]^2 - 4(u, u)(w, w) < 0 \quad (8.46)$$

de donde si $w \neq -\lambda u$ obtenemos la (8.43) con el signo $>$. Si en cambio $w = -\lambda u$ es fácil verificar que la (8.43) se cumple con el signo $=$.

Sea ahora $\Psi = \Psi(x, t)$ un estado cualquiera de nuestro sistema, y pongamos

$$u = (F - \bar{F})\Psi \quad y \quad w = -i(G - \bar{G})\Psi \quad (8.47)$$

Entonces

$$(u, u) = ((F - \bar{F})\Psi, (F - \bar{F})\Psi) = (\Psi, (F - \bar{F})^2 \Psi) = \Delta F^2 \quad (8.48)$$

y análogamente se encuentra que

$$(w, w) = \Delta G^2 \quad (8.49)$$

Por otra parte

$$\begin{aligned} (u, w) + (w, u) &= -i((F - \bar{F})\Psi, (G - \bar{G})\Psi) + i((G - \bar{G})\Psi, (F - \bar{F})\Psi) \\ &= -i(\Psi, (F - \bar{F})(G - \bar{G})\Psi) + i(\Psi, (G - \bar{G})(F - \bar{F})\Psi) \\ &= -i(\Psi, (FG - F\bar{G} - \bar{F}G + \bar{F}\bar{G})\Psi) + i(\Psi, (GF - G\bar{F} - \bar{G}F + \bar{G}\bar{F})\Psi) \quad (8.50) \\ &= -i(\Psi, (FG - GF)\Psi) \\ &= (\Psi, C\Psi) \end{aligned}$$

Sustituimos ahora las (8.48)-(8.50) en la desigualdad de Schwarz (8.43) y obtenemos el resultado buscado:

Relación general de incerteza para operadores que no conmutan:

cuando dos observables F y G no conmutan, las incertezas de F y G cumplen la relación

$$\Delta F \Delta G \geq \frac{1}{2} | \bar{C} | \quad (8.51)$$

donde $C = -i[F, G]$.

Como caso particular sea $F = x$ y $G = p_x$. Entonces, de la (8.21) $[x, p_x] = i\hbar$, luego $C = \hbar$ y obtenemos de (8.51) la relación de incerteza:

$$\Delta x \Delta p_x \geq \frac{1}{2} \hbar \quad (8.52)$$

como habíamos ya anticipado en el Capítulo 6.

Constantes del movimiento y ecuaciones del movimiento para operadores

Sea F un observable (que puede depender explícitamente del tiempo) y $\bar{F}$ su valor esperado en un estado cualquiera $\Psi(x, t)$. La variación de $\bar{F}$ con el tiempo está dada por

$$\frac{d}{dt}\bar{F} = \frac{d}{dt}(\Psi, F\Psi) = \left(\frac{\partial\Psi}{\partial t}, F\Psi\right) + \left(\Psi, F\frac{\partial\Psi}{\partial t}\right) + \left(\Psi, \frac{\partial F}{\partial t}\Psi\right) \quad (8.53)$$

Usando la ecuación de Schrödinger

$$H\Psi = i\hbar\frac{\partial\Psi}{\partial t} \quad (8.54)$$

donde H es el operador Hamiltoniano del sistema, podemos escribir la (8.53) en la forma

$$\begin{aligned} \frac{d}{dt}\bar{F} &= \frac{i}{\hbar}(H\Psi, F\Psi) - \frac{i}{\hbar}(\Psi, FH\Psi) + \left(\Psi, \frac{\partial F}{\partial t}\Psi\right) \\ &= \frac{i}{\hbar}(\Psi, [H, F]\Psi) + \left(\Psi, \frac{\partial F}{\partial t}\Psi\right) \end{aligned} \quad (8.55)$$

es decir

$$\frac{d}{dt}\bar{F} = \frac{i}{\hbar}\overline{[H, F]} + \overline{\frac{\partial F}{\partial t}} \quad (8.56)$$

La (8.56) muestra que si F conmuta con H y no depende explícitamente del tiempo,

$$\frac{d}{dt}\bar{F} = 0 \quad \text{para todo } \Psi(x, t) \quad (8.57)$$

Cuando se cumple la (8.57) el observable F se denomina *constante del movimiento* y se dice que se *conserva*. En vista de lo demostrado precedentemente podemos afirmar lo siguiente:

- Un observable F que no depende explícitamente del tiempo y conmuta con el Hamiltoniano del sistema es *constante del movimiento*. Existe entonces un sistema completo de autofunciones de H y F , es decir un conjunto completo de estados estacionarios en los cuales F está bien definido y que por lo tanto se pueden caracterizar por los autovalores E_j de H y F_k de F .
- El conjunto de todas las constantes de movimiento del sistema junto con el Hamiltoniano forman un *sistema completo de observables compatibles*. Por lo tanto existe un sistema completo de estados estacionarios que se pueden caracterizar por los autovalores E_j de H y de las constantes del movimiento.

Se puede *definir* un operador dF/dt que se denomina *derivada total de F con respecto del tiempo* como el operador para el cual

$$\frac{d\bar{F}}{dt} = \frac{d}{dt}\bar{F} \quad \text{para todo } \Psi(x, t) \quad (8.58)$$

Podemos escribir la (8.56) en términos de dF/dt :

$$\frac{d\bar{F}}{dt} = \frac{i}{\hbar}\overline{[H, F]} + \overline{\frac{\partial F}{\partial t}} \quad (8.59)$$

a partir de la cual podemos escribir la siguiente ecuación entre operadores:

$$\frac{dF}{dt} = \frac{i}{\hbar}[H, F] + \frac{\partial F}{\partial t} \quad (8.60)$$

Para una partícula que se mueve en un campo de fuerzas que deriva de una energía potencial $V(\mathbf{r})$ el Hamiltoniano es:

$$H = \frac{\mathbf{p}^2}{2m} + V(\mathbf{r}) \quad (8.61)$$

donde $\mathbf{p}$ es un operador vectorial cuyas componentes son los operadores p_x, p_y, p_z y $\mathbf{r}$ es un operador vectorial cuyas componentes son los operadores x, y, z . Entonces, la velocidad está representada por el operador vectorial

$$\mathbf{v} = \frac{d\mathbf{r}}{dt} = \frac{i}{\hbar}[H, \mathbf{r}] = \frac{i}{2\hbar m}[\mathbf{p}^2, \mathbf{r}] \quad (8.62)$$

Es fácil verificar que $[p_x^2, x] = -2i\hbar p_x$ y análogamente para las demás componentes. Por lo tanto resulta

$$\mathbf{v} = \frac{\mathbf{p}}{m} \quad (8.63)$$

La (8.63) es formalmente idéntica a la relación clásica entre velocidad e impulso, pero aquí se trata de una relación entre operadores.

Del mismo modo, la aceleración está representada por el operador vectorial

$$\mathbf{a} = \frac{d\mathbf{v}}{dt} = \frac{i}{\hbar}[H, \mathbf{v}] = \frac{i}{\hbar m}[H, \mathbf{p}] = \frac{i}{\hbar m}[V(\mathbf{r}), \mathbf{p}] \quad (8.64)$$

Si recordamos que $\mathbf{p} = -i\hbar\nabla$ resulta que $[V(\mathbf{r}), \mathbf{p}] = i\hbar\nabla V(\mathbf{r})$ y por lo tanto obtenemos

$$m\mathbf{a} = -\nabla V(\mathbf{r}) \quad (8.65)$$

que también coincide formalmente con la expresión clásica, esto es la Segunda Ley de Newton, pero el significado de la (8.65) es diferente pues se trata de una relación entre operadores como la (8.63).

El límite clásico

Las ecuaciones (8.63) y (8.65) implican que para cualquier estado $\Psi(x, t)$ los valores esperados de la velocidad y la aceleración de una partícula que se mueve en un campo de fuerzas que deriva de una energía potencial $V(\mathbf{r})$ cumplen las relaciones

$$\bar{\mathbf{v}} = \frac{\bar{\mathbf{p}}}{m} \quad (8.66)$$

y

$$m\bar{\mathbf{a}} = -\overline{\nabla V(\mathbf{r})} \quad (8.67)$$

Estas relaciones son idénticas a las ecuaciones clásicas, salvo que se refieren a los valores esperados. Cuando las imprecisiones de nuestras medidas de los impulsos y las posiciones de la partícula descrita por $\Psi(x,t)$ son muy grandes en comparación con las incertezas de dichas magnitudes, los valores medidos coinciden con los valores esperados. Por lo tanto en ese caso no es necesario distinguir entre valores medidos y valores esperados. En esas condiciones los valores medidos cumplen las ecuaciones de la Mecánica Clásica. Este resultado fue obtenido por Paul Ehrenfest (1927) y se conoce como

Teorema de Ehrenfest:

la Mecánica Cuántica coincide con la Mecánica Clásica en el límite en que el principio de incerteza deja de tener relevancia.

Por otra parte fuera de este límite la Mecánica Cuántica predice varios interesantes fenómenos que no tienen análogo clásico como veremos en breve.

Representación coordenadas y representación impulsos

Podemos avanzar en nuestro entendimiento de la Mecánica Cuántica si hacemos un análisis de Fourier de la función de onda, en la forma

$$\Psi(\mathbf{r}, t) = \frac{1}{(2\pi\hbar)^{3/2}} \int \Phi(\mathbf{p}, t) e^{\frac{i}{\hbar} \mathbf{p} \cdot \mathbf{r}} d\mathbf{p} \quad (8.68)$$

donde $\Phi(\mathbf{p}, t)$ es la transformada de Fourier de $\Psi(\mathbf{r}, t)$ en el instante t

$$\Phi(\mathbf{p}, t) = \frac{1}{(2\pi\hbar)^{3/2}} \int \Psi(\mathbf{r}, t) e^{-\frac{i}{\hbar} \mathbf{p} \cdot \mathbf{r}} d\mathbf{r} \quad (8.69)$$

Se pueden obtener fácilmente algunas propiedades importantes de Φ usando la relación de clausura en tres dimensiones, que se escribe por analogía con la (7.81) como

$$\frac{1}{(2\pi\hbar)^{3/2}} \int e^{\frac{i}{\hbar} \mathbf{p} \cdot \mathbf{r}} d\mathbf{p} = \frac{1}{(2\pi)^3} \int e^{i\mathbf{k} \cdot \mathbf{r}} d\mathbf{k} = \delta(\mathbf{r}) \quad (8.70)$$

Introduciendo la representación de Fourier (8.68) y usando las propiedades de la función delta de Dirac para intercambiar el orden de las integraciones, podemos mostrar que

$$\int |\Psi(\mathbf{r}, t)|^2 d\mathbf{r} = \int |\Phi(\mathbf{p}, t)|^2 d\mathbf{p} \quad (8.71)$$

y que el valor esperado de $\mathbf{p}$ se expresa como

$$\bar{\mathbf{p}} = \int \Psi(\mathbf{r}, t)^* \left(\frac{\hbar}{i} \nabla \right) \Psi(\mathbf{r}, t) d\mathbf{r} = \int \Phi(\mathbf{p}, t)^* \mathbf{p} \Phi(\mathbf{p}, t) d\mathbf{p} \quad (8.72)$$

mientras que el valor esperado de $\mathbf{r}$ está dado por

$$\bar{\mathbf{r}} = \int \Psi(\mathbf{r}, t)^* \mathbf{r} \Psi(\mathbf{r}, t) d\mathbf{r} = \int \Phi(\mathbf{p}, t)^* \left(-\frac{\hbar}{i} \nabla_{\mathbf{p}} \right) \Phi(\mathbf{p}, t) d\mathbf{p} \quad (8.73)$$

donde $\nabla_{\mathbf{p}}$ es el operador gradiente en el espacio de los impulsos. La fórmula (8.71) implica que si Ψ está normalizada a la unidad para todo t , Φ queda también automáticamente normalizada. Por otra parte, puesto que

$$\frac{1}{\sqrt{2\pi\hbar}} e^{-\frac{i}{\hbar}\mathbf{p}\cdot\mathbf{r}} \quad (8.74)$$

representa un estado de impulso $\mathbf{p}$ bien definido, Φ representa la amplitud con la cual está presente dicho estado en la función de onda Ψ . Por ejemplo, si $\Phi(\mathbf{p}, t)$ tiene un pico bien marcado para un determinado valor $\mathbf{p} = \mathbf{p}(t)$, en el instante t nuestro estado es un estado cuyo impulso está bastante bien definido y se parece (en ese instante) a una onda plana.

Esta observación, en conjunción con la (8.72), nos lleva a interpretar que $|\Phi(\mathbf{p}, t)|^2$ es la probabilidad de encontrar (en el instante t) que el valor del impulso de la partícula está comprendido entre $\mathbf{p}$ y $\mathbf{p} + d\mathbf{p}$. Diremos entonces que $|\Phi(\mathbf{p}, t)|^2$ es la *densidad de probabilidad en el espacio de los impulsos*.

Hay una evidente analogía entre la (8.72) y la (7.34), y entre la (8.73) y la (7.36). Del punto de vista matemático, esto es simplemente una consecuencia de las relaciones de reciprocidad entre la transformación y la antitransformación de Fourier, ecs. (8.68) y (8.69). De resultados de ello, tanto la función de onda Ψ (definida en el espacio de las coordenadas) como la Φ (definida en el espacio de los impulsos), son descripciones *igualmente válidas* del estado del sistema. Dada una de ellas, la otra se puede calcular a partir de la (8.69) o de la (8.68). Del punto de vista físico, esta reciprocidad es una manifestación de la complementariedad y de la naturaleza ondulatoria de la materia.

De acuerdo con los postulados de la Mecánica Cuántica, el valor esperado de una función $f(\mathbf{r})$ de las coordenadas está dado por

$$\overline{f(\mathbf{r}, t)} = \int f(\mathbf{r}, t) |\Psi(\mathbf{r}, t)|^2 d\mathbf{r} \quad (8.75)$$

Si sustituimos la expresión (8.68) de Ψ en la (8.75) podemos mostrar (si la función de onda se anula con suficiente rapidez a distancias grandes, para que se pueda intercambiar el orden de las integraciones y los términos de superficie se puedan despreciar) que se obtiene

$$\overline{f(\mathbf{r}, t)} = \int \Phi(\mathbf{p}, t)^* f\left(-\frac{\hbar}{i}\nabla_{\mathbf{p}}, t\right) \Phi(\mathbf{p}, t) d\mathbf{p} \quad (8.76)$$

Análogamente, a partir de la (8.72) se puede mostrar que para toda función analítica $g(\mathbf{p}, t)$ se tiene

$$\overline{g(\mathbf{p}, t)} = \int \Phi(\mathbf{p}, t)^* g(\mathbf{p}, t) \Phi(\mathbf{p}, t) d\mathbf{p} \quad (8.77)$$

En general, para un operador $h(\mathbf{r}, \mathbf{p}, t)$ que es una función analítica de $\mathbf{r}$ y $\mathbf{p}$ y eventualmente del tiempo, siempre y cuando su expresión no contenga productos de operadores que no conmutan (del tipo $r_i^\alpha p_i^\beta$), se puede mostrar que

$$\overline{h(\mathbf{r}, \mathbf{p}, t)} = \int \Psi(\mathbf{r}, t)^* h\left(\mathbf{r}, \frac{\hbar}{i} \nabla_{\mathbf{r}}, \mathbf{p}, t\right) \Psi(\mathbf{r}, t) d\mathbf{r} = \int \Phi(\mathbf{p}, t)^* h\left(-\frac{\hbar}{i} \nabla_{\mathbf{p}}, \mathbf{p}, t\right) \Phi(\mathbf{p}, t) d\mathbf{p} \quad (8.78)$$

A partir de las fórmulas clásicas, estas prescripciones permiten escribir las expresiones para los operadores cuánticos correspondientes en la mayoría de los casos de interés.

Todavía no sabemos como proceder cuando nuestros operadores contienen expresiones del tipo $r_i^\alpha p_i^\beta$ es decir, productos de operadores que no conmutan³ (por ejemplo, el operador $\mathbf{r} \cdot \mathbf{p}$). En estos casos postularemos que la (8.78) sigue valiendo, pero es necesario entonces agregar alguna prescripción acerca del *orden* adecuado de los factores, cosa que veremos oportunamente cuando se presente el caso.

La ec. (8.78) establece que a toda magnitud física le está asociado un operador lineal que, cuando se lo interpone entre la función de onda y su conjugada compleja, permite obtener por integración el valor esperado de la correspondiente magnitud. La forma explícita del operador depende, evidentemente, de si describimos el estado del sistema mediante la función de onda Ψ en el espacio de las coordenadas o la función de onda Φ en el espacio de los impulsos. En el primer caso se dice que se está usando la *representación coordenadas*, en el segundo caso que se usa la *representación impulsos*.

Del punto de vista práctico, es evidente que al tratar una partícula libre es preferible usar la representación impulsos. Pero si la energía potencial no es constante, generalmente es mejor usar la representación coordenadas, porque trabajar con un operador de la forma $V(i\hbar \nabla_{\mathbf{p}})$ puede ser incómodo. En el caso del oscilador armónico, cuyo Hamiltoniano clásico es $p^2/2m + \omega^2 r^2/2$, las dos representaciones son igualmente convenientes.

Claramente, ambas formas de describir al sistema son equivalentes y la elección de una u otra es una cuestión de mera conveniencia. Más adelante veremos que es ventajoso considerar ambas descripciones como dos representaciones de una formulación más general y abstracta de la Mecánica Cuántica.

Transformaciones unitarias

El hecho que exista la libertad de elegir la forma de describir un sistema se debe a que en la Mecánica Cuántica, las cantidades medibles (y que por lo tanto se comparan con los resultados experimentales) se calculan *siempre* a partir de expresiones del tipo⁴

$$(f, g) \quad (8.79)$$

y por lo tanto están sólo *indirectamente* relacionadas con f y g . Por consiguiente, cualquier transformación del tipo

$$f' = Uf \quad (8.80)$$

³ El problema aquí es que el análisis de Fourier nos permite describir nuestro estado, o en términos de las autofunciones de una componente impulso, o en términos de las autofunciones de la coordenada conjugada. Pero el principio de incerteza nos impide encontrar autofunciones de ambos operadores: tales autofunciones simultáneas no existen.

⁴ Aquí usamos la notación compacta del producto escalar que introdujimos en el Capítulo 7.

que a cada función f de nuestro espacio funcional $\mathcal{F}$ le haga corresponder una función f' de otro espacio funcional $\mathcal{F}'$, de forma tal que preserve el producto escalar (8.80), es decir, tal que se cumpla que

$$(f', g') = (f, g) \quad (8.81)$$

para todo par de funciones f, g de $\mathcal{F}$, nos da una descripción equivalente de nuestro sistema. Este es justamente el caso de la transformación de Fourier, que a la función $\Psi(\mathbf{r}, t)$ del espacio de las coordenadas le hace corresponder la función $\Phi(\mathbf{p}, t)$ del espacio de los impulsos, pero esta no es la única transformación del tipo (8.80) posible.

Veamos qué condiciones generales tiene que satisfacer el operador U que induce la transformación (8.80) para que se cumpla la (8.81). Es obvio, para empezar, que U debe ser un operador lineal y que debe tener inversa, para que las descripciones en los espacios $\mathcal{F}$ y $\mathcal{F}'$ sean equivalentes. Pero además, recordando la definición del operador adjunto, tenemos que

$$(f', g') = (Uf, Ug) = (U^\dagger Uf, g) \quad (8.82)$$

Comparando la (8.82) con la (8.81) vemos que para que el producto escalar de dos funciones cualesquiera se conserve, se debe cumplir que

$$U^\dagger = U^{-1} \quad (8.83)$$

de modo que

$$U^\dagger U = U U^\dagger = 1 \quad (8.84)$$

donde con 1 indicamos el operador identidad. Un operador que cumple la (8.83) se denomina *unitario*.

Como se puede apreciar de las fórmulas (8.72), (8.73), (8.75), (8.76) y (8.78) para el caso de las representaciones coordenadas e impulso, la transformación U induce también transformaciones de los operadores que representan variables dinámicas. De esta forma, el operador A definido en $\mathcal{F}$ se transforma en un operador A' definido en $\mathcal{F}'$, y para que las dos representaciones sean equivalentes se debe cumplir

$$(f, Ag) = (f', A'g') \quad (8.85)$$

para cualesquiera $f, g \in \mathcal{F}$. Usando la (8.80) y la condición de unitariedad (8.83), el miembro izquierdo de (8.85) se puede escribir como

$$(f, Ag) = (U^{-1}f', AU^{-1}g') = (U^\dagger f', AU^\dagger g') = (f', UAU^\dagger g') \quad (8.86)$$

Comparando entonces la (8.86) con la (8.85) vemos que la transformación del operador A está dada por

$$A' = UAU^\dagger \quad (8.87)$$

En resumidas cuentas:

- Dada una transformación unitaria U que nos lleva de $\mathcal{F}$ a $\mathcal{F}'$, la descripción del sistema en términos de las funciones f y los operadores A definidos en $\mathcal{F}$ es equivalente a la descripción en términos de las funciones f' y los operadores A' definidos en $\mathcal{F}'$.
- Las relaciones entre f y f' y entre A y A' están dadas, respectivamente, por las ecs. (8.80) y (8.87).

En realidad no hace falta que $\mathcal{F}'$ sea un espacio funcional. Basta que sea un espacio vectorial cualquiera, con tal que se pueda establecer una correspondencia biunívoca del tipo (8.80) entre los elementos de $\mathcal{F}$ y $\mathcal{F}'$, que preserve el producto escalar (esto es, que cumpla la (8.81)). Veremos más adelante ejemplos concretos de estas transformaciones.

Por lo tanto, resulta que mediante transformaciones inducidas por operadores unitarios podemos obtener diferentes descripciones del mismo sistema cuántico, todas equivalentes. Esta es la idea básica en que se funda la equivalencia entre el formalismo de matrices de Heisenberg y la mecánica ondulatoria de Schrödinger, y también la teoría de las transformaciones de Dirac y Jordan. Volveremos sobre estas cuestiones más adelante, pues primero conviene que el lector se familiarice con la descripción dada por la teoría de Schrödinger, antes de introducir formalismos más generales, pero por eso mismo más abstractos y por lo tanto menos intuitivos.

9. SOLUCIONES DE LA ECUACIÓN DE SCHRÖDINGER

Introducción

En este Capítulo aplicaremos el formalismo de Schrödinger de la Mecánica Cuántica para estudiar las soluciones de algunos problemas sencillos en una dimensión. El propósito de estos ejemplos es que el lector se familiarice con las técnicas de cálculo y que vea el origen de algunas de las curiosas propiedades de las soluciones de la ecuación de Schrödinger. Comenzaremos por el caso más sencillo, que es la partícula libre. Luego consideraremos el potencial escalón y la barrera de potencial, para poner en evidencia un importante fenómeno que es puramente cuántico: la penetración de una barrera o “efecto túnel”. Finalmente trataremos el oscilador armónico simple y mostraremos que sus niveles de energía están cuantificados, de una forma ligeramente diferente de la que resulta del Postulado de Planck de la Teoría Cuántica Antigua. En todos los casos vamos a emplear la representación coordenadas.

La partícula libre

El Hamiltoniano de una partícula libre en una dimensión espacial x es

$$H = \frac{p^2}{2m} \quad (9.1)$$

Claramente p conmuta con H , de modo que el impulso es constante del movimiento. La ecuación de Schrödinger independiente del tiempo es

$$-\frac{\hbar^2}{2m} \frac{d^2\psi}{dx^2} = E\psi \quad (9.2)$$

y como sabemos tiene soluciones para cualquier valor de $E \geq 0$. Por lo tanto el espectro de autovalores de H es continuo y las correspondientes autofunciones de la energía (normalizadas a la delta de Dirac) son las ondas planas

$$\psi_{E,k} = \frac{1}{\sqrt{2\pi}} e^{ikx}, \quad E = \frac{\hbar^2 k^2}{2m} = \frac{p^2}{2m} \quad (9.3)$$

Las correspondientes funciones de onda son

$$\Psi_{E,k} = e^{-iEt/\hbar} \psi_{E,k} = \frac{1}{\sqrt{2\pi}} e^{i(kx - \omega t)}, \quad \omega = \frac{E}{\hbar} = \frac{\hbar k^2}{2m} \quad (9.4)$$

Cada autovalor E de la energía es doblemente degenerado, pues corresponde a dos valores del impulso:

$$p = \hbar k = \pm \sqrt{2mE} \quad (9.5)$$

El signo $+$ en la (9.5) corresponde a una partícula que se mueve hacia la derecha y el signo $-$ a una partícula que se mueve hacia la izquierda. Puesto que los estados estacionarios (9.4) tienen un impulso bien definido, la posición de la partícula está totalmente indeterminada y es igualmente probable encontrarla en cualquier parte.

Las funciones de onda (9.4) forman un sistema completo, de modo que cualquier estado $\Psi(x, t)$ solución de la ecuación de Schrödinger

$$i\hbar \frac{\partial \Psi(x, t)}{\partial t} = H\Psi(x, t) \quad (9.6)$$

se puede representar como una superposición de las (9.4) de la forma

$$\Psi(x, t) = \int_{-\infty}^{+\infty} a(k) \Psi_{E,k} dk = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} a(k) e^{i(kx - \omega t)} dk \quad (9.7)$$

donde

$$a(k) = (\Psi_{E,k}, \Psi) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} e^{-i(kx - \omega t)} \Psi(x, t) dx \quad (9.8)$$

Puesto que

$$(\Psi, \Psi) = \int_{-\infty}^{+\infty} a(k)^* a(k) dk \quad (9.9)$$

la $\Psi(x, t)$ está normalizada si la distribución espectral $a(k)$ lo está. Mediante paquetes de ondas de la forma (9.7) podemos describir estados en los cuales la partícula está localizada y estudiar su movimiento, como ya hicimos anteriormente.

Si no queremos trabajar con las autofunciones del continuo, podemos normalizar las autofunciones de la energía en un intervalo de longitud L . Las soluciones permitidas normalizadas en el intervalo $(x_0, x_0 + L)$ que satisfacen $\psi(x_0) = \psi(x_0 + L) = 0$ (normalización en una caja) son ondas estacionarias de la forma

$$\psi(x - x_0) = (2/L)^{1/2} \sin[k_n(x - x_0)] \quad , \quad k_n = n\pi/L \quad , \quad n = 1, 2, 3, \dots \quad (9.10)$$

y los correspondientes autovalores discretos de la energía son

$$E_n = \hbar^2 n^2 \pi^2 / 2mL^2 \quad , \quad n = 1, 2, 3, \dots \quad (9.11)$$

Si usamos en cambio la condición de contorno periódica $\psi(x_0) = \psi(x_0 + L)$, las soluciones permitidas (normalizadas en el intervalo $(x_0, x_0 + L)$) son ondas viajeras de la forma

$$\psi(x - x_0) = (1/L)^{1/2} e^{ik_n(x - x_0)} \quad , \quad k_n = \pm 2n\pi/L \quad , \quad n = 0, 1, 2, \dots \quad (9.12)$$

y los autovalores discretos de la energía son

$$E_n = 2\hbar^2 n^2 \pi^2 / mL^2 \quad , \quad n = 0, 1, 2, \dots \quad (9.13)$$

Mientras L sea mucho mayor que el tamaño de la región de interés, estos procedimientos no tienen efectos significativos y L no aparece en los resultados de los cálculos de las cantidades de interés físico.

Esta manera de hacer las cosas es útil para las aplicaciones a la Mecánica Estadística, en las que conviene trabajar con el espectro discreto para poder contar el número de estados en un dado intervalo de energía. Con miras a esas aplicaciones vamos a calcular la densidad de estados por unidad de intervalo de la energía para una partícula libre en tres dimensiones. Normalizaremos las autofunciones de la energía en un cubo de arista L .

Las soluciones permitidas satisfacen las condiciones de contorno $\psi(0, y, z) = \psi(L, y, z) = 0$, $\psi(x, 0, z) = \psi(x, L, z) = 0$, $\psi(x, y, 0) = \psi(x, y, L) = 0$, y son ondas estacionarias de la forma

$$\psi(x, y, z) = (2/L)^{3/2} \sin(k_x x) \sin(k_y y) \sin(k_z z) \quad (9.14)$$

con

$$k_x = n_x \pi / L, \quad k_y = n_y \pi / L, \quad k_z = n_z \pi / L, \quad n_x, n_y, n_z = 1, 2, 3, \dots \quad (9.15)$$

y los correspondientes autovalores discretos de la energía son

$$E_n = \hbar^2 n^2 \pi^2 / 2mL^2, \quad n^2 = n_x^2 + n_y^2 + n_z^2 \quad (9.16)$$

Podemos dar una imagen geométrica de la (9.16) representando cada autovalor como un punto de coordenadas (n_x, n_y, n_z) en un espacio de tres dimensiones. Claramente, el número $N(E)$ de estados con energía menor o igual a cierto valor E es igual a la cantidad de puntos (n_x, n_y, n_z) que cumplen la condición $E_n \leq E$, que usando la (9.16) se puede escribir como

$$n_x^2 + n_y^2 + n_z^2 \leq \frac{2mL^2 E}{\hbar^2 \pi^2} \quad (9.17)$$

Por lo tanto, $N(E)$ es igual al volumen del primer octante de una esfera del espacio (n_x, n_y, n_z) con centro en el origen y radio $\rho(E) = (L/\hbar\pi)\sqrt{2mE}$. Entonces

$$N(E) = \frac{1}{8} \left(\frac{4}{3} \pi \rho(E)^3 \right) = \frac{1}{8} \frac{4}{3} \pi \frac{L^3}{\pi^3 \hbar^3} (2mE)^{3/2} = \frac{4}{3} \pi \frac{V}{h^3} (2mE)^{3/2} \quad (9.18)$$

donde $V = L^3$ es el volumen del cubo. Diferenciando la (9.18) podemos calcular la densidad de estados por unidad de intervalo de la energía, definida por $dN = f(E)dE$:

$$f(E) = 4\pi m \frac{V}{h^3} (2mE)^{1/2} \quad (9.19)$$

Usando la relación $p = \sqrt{2mE}$, podemos escribir N en términos del módulo de la cantidad de movimiento en la forma

$$N = \frac{4}{3} \pi \frac{V p^3}{h^3} \quad (9.20)$$

Se puede observar que la cantidad

$$\mathcal{V}_f = \frac{4}{3} \pi p^3 V \quad (9.21)$$

es el “volumen” del *espacio de las fases* (x, y, z, p_x, p_y, p_z) accesible a nuestra partícula, que está contenida dentro del cubo de arista L , y cuya cantidad de movimiento tiene un módulo menor o igual que p . Por lo tanto la (8.19) se puede escribir también en la forma

$$\mathcal{V}_f = h^3 N \quad (9.22)$$

Este resultado se suele expresar diciendo que

Cada estado ocupa un volumen h^3 del espacio de las fases.

La densidad de estados por unidad de intervalo del módulo de la cantidad de movimiento es

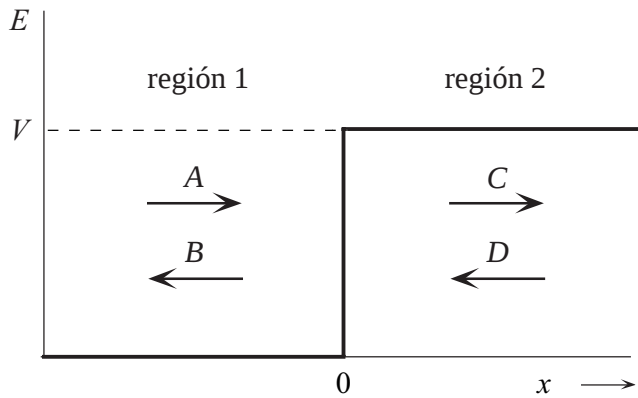
$$\frac{dN}{dp} = f(p) = \frac{4\pi p^2 V}{h^3} \quad (8.23)$$

como se puede obtener de inmediato diferenciando la (9.20).

Obtuvimos estos resultados utilizando funciones de onda normalizadas en una caja cúbica de arista L , pero es fácil verificar que se obtiene el mismo resultado si se usan funciones de onda normalizadas con condiciones periódicas de contorno en un cubo de arista L . Más en general, se puede mostrar que los resultados (9.18)-(9.20), (9.22) y (9.23) son independientes de la forma de la caja y sólo dependen de su volumen V .

El potencial escalón

Sea una partícula que se mueve en una dimensión y cuya energía potencial es (Fig. 9.1):



$$V(x) = \begin{cases} 0 & , \quad x < 0 \\ V = \text{cte.} & , \quad x > 0 \end{cases} \quad (9.24)$$

El Hamiltoniano es entonces

$$H = \begin{cases} \frac{p^2}{2m} & , \quad x < 0 \\ \frac{p^2}{2m} + V & , \quad x > 0 \end{cases} \quad (9.25)$$

Fig. 9.1. Potencial en escalón.

Para $x < 0$ (región 1) la ecuación de Schrödinger independiente del tiempo es

$$-\frac{\hbar^2}{2m} \frac{d^2\psi}{dx^2} = E\psi \quad (9.26)$$

cuya solución general (no normalizada) para un dado valor de E es

$$\psi_E^{(1)} = Ae^{ikx} + Be^{-ikx} \quad , \quad k = +\sqrt{2mE} / \hbar \quad (9.27)$$

donde A y B son constantes a determinar. La corriente de probabilidad de la solución (9.27) es

$$J_E^{(1)} = \frac{\hbar k}{m} (|A|^2 - |B|^2) = v(|A|^2 - |B|^2) \quad (9.28)$$

donde $v = \hbar k / m$ es el módulo de la velocidad en la región 1. Por lo tanto $J_E^{(1)}$ es la diferencia entre la corriente $v |A|^2$ que va hacia la derecha y la corriente $-v |B|^2$ que va hacia la izquierda. Para $x > 0$ (región 2) la ecuación de Schrödinger independiente del tiempo es

$$-\frac{\hbar^2}{2m} \frac{d^2\psi}{dx^2} = (E - V)\psi \quad (9.29)$$

Aquí tenemos que distinguir dos casos, según si E es mayor o menor que V .

Si $E > V$ la solución general de (9.29) es

$$\psi_E^{(2)} = Ce^{ik'x} + De^{-ik'x} \quad , \quad k' = +\sqrt{2m(E - V)} / \hbar \quad (9.30)$$

donde C y D son constantes a determinar. La corriente de probabilidad de la solución (9.30) es

$$J_E^{(2)} = \frac{\hbar k'}{m} (|C|^2 - |D|^2) = v' (|C|^2 - |D|^2) \quad (9.31)$$

donde $v' = \hbar k' / m$ es el módulo de la velocidad en la región 2. Por lo tanto $J_E^{(2)}$ es la diferencia entre la corriente $v' |C|^2$ que va hacia la derecha y la corriente $-v' |D|^2$ que va hacia la izquierda.

Si $E < V$ la solución general de (9.29) es una superposición de ondas *evanescentes*:

$$\psi_E^{(2)} = Ce^{-\kappa x} + De^{\kappa x} \quad , \quad \kappa = +\sqrt{2m(V - E)} / \hbar \quad (9.32)$$

Pero en nuestro caso el término $De^{\kappa x}$ no es aceptable pues diverge para $x \rightarrow \infty$. Debemos entonces poner $D = 0$ en la (9.32) y queda

$$\psi_E^{(2)} = Ce^{-\kappa x} \quad (9.33)$$

La corriente de probabilidad correspondiente a la solución (9.33) es nula.

Ahora tenemos que empalmar en $x = 0$ la solución de la región 1 con la de la región 2. Para eso observamos que $\psi(x)$ y $d\psi(x)/dx$ deben ser continuas, para que la corriente de probabilidad sea continua y por lo tanto se conserve la probabilidad. Debemos pedir entonces

$$\left[\psi_E^{(1)}(x) \right]_{x=0} = \left[\psi_E^{(2)}(x) \right]_{x=0} \quad , \quad \left[\frac{d}{dx} \psi_E^{(1)}(x) \right]_{x=0} = \left[\frac{d}{dx} \psi_E^{(2)}(x) \right]_{x=0} \quad (9.34)$$

Vamos a considerar por separado los casos $E > V$ y $E < V$.

Caso $E > V$

En este caso las condiciones (9.34) nos dan dos ecuaciones

$$A + B = C + D \quad , \quad kA - kB = k'C - k'D \quad (9.35)$$

para determinar las cuatro constantes que figuran en las (9.35). Pero como se debe cumplir la condición de normalización, en realidad sólo tres de ellas son independientes. Por lo tanto hay infinitas maneras de satisfacer las (9.35). Esto corresponde al hecho que podemos tener una solución que corresponde a que la partícula llega al escalón viniendo desde la izquierda, otra solu-

ción que corresponde a que llega viniendo desde la derecha, y también cualquier combinación lineal de ambas. Por lo tanto para $E > V$ el espectro de la energía es continuo y cada autovalor $E > V$ es doblemente degenerado.

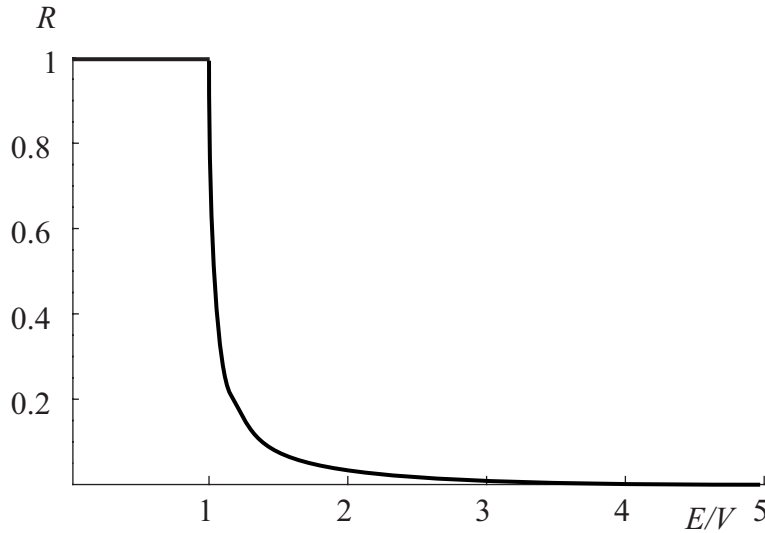
Partícula que viene desde la izquierda

Para una partícula que llega desde la izquierda debemos tener $D \equiv 0$, de modo que si $E > V$ las (9.35) se reducen a

$$A + B = C \quad , \quad kA - kB = k'C \quad (9.36)$$

Resolviendo este sistema obtenemos

$$\frac{B}{A} = \frac{k - k'}{k + k'} \quad , \quad \frac{C}{A} = \frac{2k}{k + k'} \quad (9.37)$$



El *coeficiente de reflexión* del potencial escalón se define como el cociente entre la corriente reflejada y la corriente incidente:

$$R = \frac{v |B|^2}{v |A|^2} = \frac{(k - k')^2}{(k + k')^2} \quad (9.38)$$

El valor de R (Fig. 9.2) depende solamente del cociente $\rho = E/V$:

$$R = \frac{(\sqrt{\rho} - \sqrt{\rho - 1})^2}{(\sqrt{\rho} + \sqrt{\rho - 1})^2} \quad (9.39)$$

Fig. 9.2. Coeficiente de reflexión del potencial en escalón.

Puesto que $B/A > 0$ la onda incidente y la reflejada están en fase.

El coeficiente de transmisión se define como

$$T = \frac{v' |C|^2}{v |A|^2} = \frac{4kk'}{(k + k')^2} = \frac{4\sqrt{\rho}\sqrt{\rho - 1}}{(\sqrt{\rho} + \sqrt{\rho - 1})^2} \quad (9.40)$$

Es fácil verificar que se cumple $R + T = 1$, de manera que se conserva la probabilidad.

Comparando con el caso clásico podemos notar una importante diferencia, pues una partícula clásica con $E > V$ no sufre reflexión al llegar al escalón de $V(x)$: simplemente sigue de largo con una diferente velocidad. En cambio una partícula cuántica tiene una probabilidad no nula de ser reflejada.

Partícula que viene desde la derecha

Si la partícula llega desde la derecha tenemos $A \equiv 0$, de modo que las (9.35) se reducen a

$$B = C + D \quad , \quad -kB = k'C - k'D \quad (9.41)$$

y entonces

$$\frac{B}{D} = \frac{2k'}{k+k'} \quad , \quad \frac{C}{D} = \frac{k'-k}{k+k'} \quad (9.42)$$

A diferencia de antes, la onda incidente y la reflejada están en contrafase, pues $C/D < 0$. El coeficiente de reflexión es:

$$R = \frac{v' |C|^2}{v' |D|^2} = \frac{(k-k')^2}{(k+k')^2} \quad (9.43)$$

y su valor es el mismo que para el caso de partícula incidente desde la izquierda y otro tanto ocurre con el coeficiente de transmisión.

Caso $E < V$

Si $E < V$ la solución en la región 2 es una exponencial decreciente, y las condiciones (9.34) nos dan

$$A + B = C \quad , \quad ik(A - B) = -\kappa C \quad (9.44)$$

de donde podemos obtener

$$\frac{B}{A} = \frac{ik + \kappa}{ik - \kappa} = e^{i\alpha} \quad , \quad \alpha = \pi + 2 \arctan(\kappa/k) \quad (9.45)$$

y

$$\frac{C}{A} = \frac{2ik}{ik - \kappa} = 1 + e^{i\alpha} \quad (9.46)$$

y en consecuencia la solución es (a menos de una fase irrelevante $e^{i\alpha/2}$)

$$\psi_E = \begin{cases} 2A \cos\left(kx - \frac{\alpha}{2}\right) & , \quad x < 0 \\ 2A \cos\left(\frac{\alpha}{2}\right) e^{-\kappa x} & , \quad x > 0 \end{cases} \quad (9.47)$$

que representa una onda estacionaria en la región 1 y una onda evanescente en la región 2. La onda incidente Ae^{ikx} se refleja *totalmente* (pues $|B|^2 = |A|^2$) y aunque ψ tiene un valor no nulo en la región 2, no hay penetración permanente. La corriente de probabilidad es nula en todas partes. El espectro de la energía para $E < V$ es continuo y a cada autovalor $E < V$ le corresponde una única autofunción, por lo tanto no hay degeneración.

La solución (9.47) predice que la partícula se puede encontrar en la región $x > 0$, clásicamente inaccesible. Sin embargo para observarla en esa región es preciso determinar su posición con una incerteza del orden $\Delta x \approx 1/\kappa$, y entonces $\Delta p > \hbar / \Delta x \approx \hbar \kappa = \sqrt{2m(V - E)}$ por el principio de incerteza. Por lo tanto, la incerteza de la energía es

$$\Delta E = \frac{\Delta p^2}{2m} = V - E \quad (9.48)$$

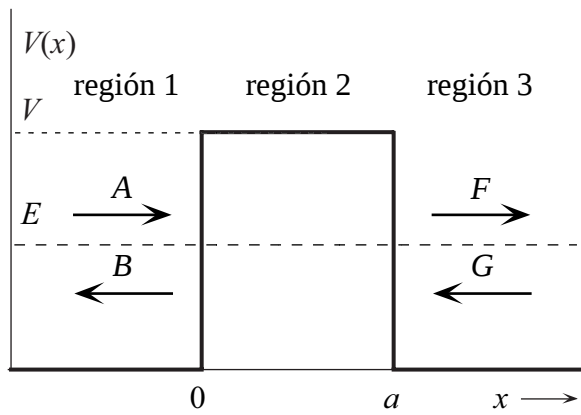
Luego si observamos a la partícula en la región prohibida por la mecánica clásica, en realidad no podemos saber si su energía total es menor que V .

El lector notará que los resultados cuánticos del caso $E > V$ son análogos a los que se obtienen en la Óptica para la reflexión y transmisión de ondas que inciden perpendicularmente sobre la interfase que separa dos medios de diferente índice de refracción. En cambio, el resultado cuántico para $E < V$ es análogo a la reflexión total interna.

Las autofunciones que hemos obtenido al estudiar este problema no han sido aún normalizadas. Si se desea se las puede normalizar como corresponde a las autofunciones del espectro continuo, esto es con la delta de Dirac.

Penetración de una barrera de potencial

Sea una partícula que se mueve en una dimensión y cuya energía potencial es (Fig. 9.3):



$$V(x) = \begin{cases} 0 & , \quad x < 0 \\ V & , \quad 0 < x < a \\ 0 & , \quad a < x \end{cases} \quad (9.49)$$

Fig. 9.3. Barrera de potencial.

Nos interesa estudiar el caso en que la partícula llega a la barrera desde la izquierda (región 1, $x < 0$) con una energía $E < V$. En ese caso la Mecánica Clásica predice que la partícula se refleja en $x = 0$ y no puede llegar a la región 3 ($x > a$). Vamos a mostrar que la Mecánica Cuántica predice en cambio que la partícula puede atravesar la barrera.

De acuerdo con lo visto anteriormente, las soluciones generales de la ecuación de Schrödinger independiente del tiempo en las tres regiones en que se divide el eje x son:

$$\begin{aligned} \text{región 1: } x < 0 & \quad , \quad \psi_E^{(1)} = Ae^{ikx} + Be^{-ikx} \quad , \quad k = +\sqrt{2mE} / \hbar \\ \text{región 2: } 0 < x < a & \quad , \quad \psi_E^{(2)} = De^{\kappa x} + Ce^{-\kappa x} \quad , \quad \kappa = +\sqrt{2m(V-E)} / \hbar \\ \text{región 3: } x > a & \quad , \quad \psi_E^{(3)} = Fe^{ikx} + Ge^{-ikx} \quad , \quad k = +\sqrt{2mE} / \hbar \end{aligned} \quad (9.50)$$

Las condiciones de empalme en $x = 0$ son

$$\left[\psi_E^{(1)}(x) \right]_{x=0} = \left[\psi_E^{(2)}(x) \right]_{x=0} \quad , \quad \left[\frac{d}{dx} \psi_E^{(1)}(x) \right]_{x=0} = \left[\frac{d}{dx} \psi_E^{(2)}(x) \right]_{x=0} \quad (9.51)$$

y nos dan

$$A + B = D + C \quad , \quad A - B = -i \frac{\kappa}{k} (D - C) \quad (9.52)$$

Si resolvemos el sistema (9.52) para A y B en términos de D y C resulta

$$A = \frac{1}{2} \left(1 - i \frac{\kappa}{k} \right) D + \frac{1}{2} \left(1 + i \frac{\kappa}{k} \right) C \quad , \quad B = \frac{1}{2} \left(1 + i \frac{\kappa}{k} \right) D + \frac{1}{2} \left(1 - i \frac{\kappa}{k} \right) C \quad (9.53)$$

Las (9.53) se pueden escribir en forma matricial:

$$\begin{pmatrix} A \\ B \end{pmatrix} = \begin{pmatrix} M_{11}^0 & M_{12}^0 \\ M_{21}^0 & M_{22}^0 \end{pmatrix} \begin{pmatrix} D \\ C \end{pmatrix} \quad (9.54)$$

donde

$$M_{11}^0 = M_{22}^0 = \frac{1}{2} \left(1 - i \frac{\kappa}{k} \right) \quad , \quad M_{12}^0 = M_{21}^0 = \frac{1}{2} \left(1 + i \frac{\kappa}{k} \right) \quad (9.55)$$

Las condiciones de empalme en $x = a$ son

$$\left[\psi_E^{(1)}(x) \right]_{x=a} = \left[\psi_E^{(2)}(x) \right]_{x=a} \quad , \quad \left[\frac{d}{dx} \psi_E^{(1)}(x) \right]_{x=a} = \left[\frac{d}{dx} \psi_E^{(2)}(x) \right]_{x=a} \quad (9.56)$$

y nos dan

$$De^{\kappa a} + Ce^{-\kappa a} = Fe^{ika} + Ge^{-ika} \quad , \quad De^{\kappa a} - Ce^{-\kappa a} = i \frac{k}{\kappa} (Fe^{ika} - Ge^{-ika}) \quad (9.57)$$

Si resolvemos el sistema (9.57) para D y C en términos de F y G resulta

$$\begin{aligned} D &= \frac{e^{-\kappa a + ika}}{2} \left(1 + i \frac{k}{\kappa} \right) F + \frac{e^{-\kappa a - ika}}{2} \left(1 - i \frac{k}{\kappa} \right) G \\ C &= \frac{e^{\kappa a + ika}}{2} \left(1 - i \frac{k}{\kappa} \right) F + \frac{e^{\kappa a - ika}}{2} \left(1 + i \frac{k}{\kappa} \right) G \end{aligned} \quad (9.58)$$

Escribimos las (9.58) en forma matricial:

$$\begin{pmatrix} D \\ C \end{pmatrix} = \begin{pmatrix} M_{11}^a & M_{12}^a \\ M_{21}^a & M_{22}^a \end{pmatrix} \begin{pmatrix} F \\ G \end{pmatrix} \quad (9.59)$$

donde

$$\begin{aligned} M_{11}^a &= \frac{e^{-\kappa a + ika}}{2} \left(1 + i \frac{k}{\kappa} \right) \quad , \quad M_{12}^a = \frac{e^{-\kappa a - ika}}{2} \left(1 - i \frac{k}{\kappa} \right) \\ M_{21}^a &= \frac{e^{\kappa a + ika}}{2} \left(1 - i \frac{k}{\kappa} \right) \quad , \quad M_{22}^a = \frac{e^{\kappa a - ika}}{2} \left(1 + i \frac{k}{\kappa} \right) \end{aligned} \quad (9.60)$$

De la (9.54) y la (9.59) obtenemos

$$\begin{pmatrix} A \\ B \end{pmatrix} = \begin{pmatrix} M_{AF} & M_{AG} \\ M_{BF} & M_{BG} \end{pmatrix} \begin{pmatrix} F \\ G \end{pmatrix} \quad (9.61)$$

donde

$$M_{AF} = M_{11}^0 M_{11}^a + M_{12}^0 M_{21}^a, \quad M_{AG} = M_{11}^0 M_{12}^a + M_{12}^0 M_{22}^a \quad (9.62)$$

$$M_{BF} = M_{21}^0 M_{11}^a + M_{22}^0 M_{21}^a, \quad M_{BG} = M_{21}^0 M_{12}^a + M_{22}^0 M_{22}^a$$

La ec. (9.61) vincula las amplitudes de las ondas a izquierda y derecha de la barrera, y por lo tanto permite resolver cualquier problema de reflexión y transmisión que se desee (para $E < V$, se entiende). En el presente caso, nos interesa estudiar la transmisión de una partícula que llega a la barrera desde la izquierda. Por lo tanto pondremos $G = 0$ en la (9.61) y entonces resulta

$$\frac{F}{A} = \frac{1}{M_{AF}} = \frac{e^{-ika}}{\cosh \kappa a + \frac{i}{2} \left(\frac{\kappa}{k} - \frac{k}{\kappa} \right) \sinh \kappa a} \quad (9.63)$$

El coeficiente de transmisión de la barrera es entonces

$$T = \left| \frac{F}{A} \right|^2 = \frac{1}{\cosh^2 \kappa a + \frac{1}{4} \left(\frac{\kappa}{k} - \frac{k}{\kappa} \right)^2 \sinh^2 \kappa a} = \frac{1}{1 + \frac{\sinh^2 \kappa a}{4 \frac{E}{V} \left(1 - \frac{E}{V} \right)}} \quad (9.64)$$

Si la barrera es alta (o sea E/V no es muy próximo a 1) y ancha ($\kappa a \gg 1$) de modo que transmite poco, entonces $\sinh \kappa a \approx e^{\kappa a} / 2$ y la expresión de T toma una forma sencilla:

$$T \approx 16 e^{-2\kappa a} \frac{E}{V} \left(1 - \frac{E}{V} \right) \quad (9.65)$$

La penetración de la barrera se suele denominar *efecto túnel* y es una manifestación del carácter ondulatorio de la partícula. El mismo fenómeno aparece para cualquier tipo de ondas. En Óptica se lo conoce con el nombre de “reflexión interna total frustrada”.

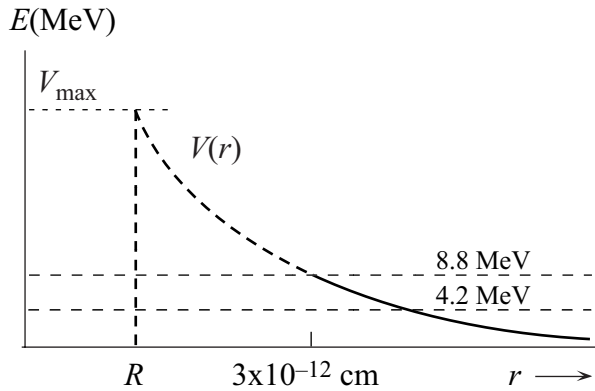


Fig. 9.4. Emisión de partículas α por un núcleo de ^{238}U .

El efecto túnel permite explicar una paradoja que se presenta en la emisión de partículas α por núcleos radioactivos. Como ejemplo, consideremos el elemento ^{238}U . Mediante la dispersión por el núcleo del ^{238}U de partículas α de 8.8 MeV emitidas por el ^{212}Po se determinó la energía potencial $V(r)$ de la partícula α y se encontró que coincide con la que proviene de la ley de Coulomb, por lo menos hasta la distancia de 3×10^{-12} cm que es hasta donde puede llegar una partícula α de 8.8 MeV. Por otra parte los experimentos de dispersión de partículas α por núcleos livianos muestran que

$V(r)$ se desvía del comportamiento $1/r$ cuando $r < R$ (R es el radio del núcleo), porque a distancias menores actúan las fuerzas nucleares que son atractivas. Cuando se desarrolló la Mecánica Cuántica no se conocía todavía el valor preciso de R para los núcleos pesados, pero era obvio que para el ^{238}U debía ser seguramente menor que 3×10^{-12} cm. Ahora bien, el núcleo del ^{238}U emite ocasionalmente partículas α . Se supuso entonces que dichas partículas están presentes

dentro del núcleo, al cual están ligadas por el potencial $V(r)$. A partir de estos argumentos se concluyó¹ que la forma de $V(r)$ es la que se indica cualitativamente en la Fig. 9.4. Por otra parte la energía cinética de las partículas α emitidas por el ^{238}U es de 4.2 MeV (medida muy lejos del núcleo, donde $V(r) = 0$). Por lo tanto se presenta una situación paradójica, pues clásicamente es inexplicable que una partícula α se emita con una energía *menor* que la que corresponde al tope de la barrera.

La paradoja fue resuelta en 1928 por George Gamow, Edward Condon y Ronald P. Gurney en términos del efecto túnel de la Mecánica Cuántica. Para la energía potencial de la Fig. 9.4 no se puede aplicar la expresión (9.64) del coeficiente de transmisión, pero se puede mostrar que

$$T = e^{-2 \int_a^b \kappa(r) dr}, \quad \kappa(r) = +\sqrt{2m[V(r) - E]}/\hbar \quad (9.66)$$

donde a y b son los puntos de retorno clásicos. La probabilidad que una partícula α que llega a la barrera la atraviese es igual a T . Por unidad de tiempo, la partícula α que va y viene dentro del núcleo, choca con la barrera $N \approx v/2R$ veces. Por lo tanto la probabilidad de emisión por unidad de tiempo es

$$\lambda \approx vT/2R \quad (9.67)$$

Tomando $v = (2E/m)^{1/2}$ y $R \approx 9 \times 10^{-13}$ cm (valor que infirieron del análisis de Rutherford de la dispersión de partículas α por núcleos livianos) Gamow, Condon y Gurney obtuvieron valores de λ que concuerdan razonablemente con los que se infieren a partir de los tiempos característicos del decaimiento radioactivo, pese a que para diferentes elementos hay enormes variaciones de λ (por ejemplo $\lambda = 5 \times 10^{-18} \text{ s}^{-1}$ para el ^{238}U y $\lambda = 2 \times 10^6 \text{ s}^{-1}$ para el ^{212}Po), que se deben a que λ depende muy fuertemente de E (la forma y la altura de la barrera son aproximadamente las mismas para todos los emisores α).

Corresponde mencionar aquí que Johannes W. Geiger y John M. Nuttall propusieron en 1912 una ley empírica de la forma $\log \lambda = a + b \log E$, para relacionar la probabilidad de emisión λ con la energía E de las partículas α emitidas por diferentes sustancias radioactivas. Dicha ley reproduce razonablemente bien los datos medidos, pero el valor de la constante b implica que λ depende de una potencia de E extraordinariamente alta, alrededor de 90. En esos tiempos, los fundamentos teóricos de la Ley de Geiger-Nuttall eran, por supuesto, desconocidos.

Por estos motivos, la aplicación exitosa de la Mecánica Cuántica a la emisión de partículas α constituyó uno de los apoyos más sólidos a la nueva teoría, además de ilustrar muy claramente la dualidad onda-partícula.

Hay muchos otros problemas sencillos en una dimensión (como el pozo cuadrado de potencial, etc.), que se pueden resolver fácilmente por medio de las técnicas que hemos presentado aquí. Por razones de brevedad no los vamos a tratar, pero el lector interesado los puede encontrar desarrollados en la bibliografía.

¹ Esta conclusión fue confirmada por experimentos posteriores con partículas α de energía suficientemente grande como para investigar el potencial para todo r .

El oscilador armónico simple

La energía potencial de un oscilador armónico simple es

$$V(x) = \frac{1}{2} m \omega^2 x^2 \quad (9.68)$$

donde ω es la frecuencia clásica del oscilador y m la masa. La forma (9.68) de $V(x)$ es de gran importancia práctica, pues *es una aproximación para cualquier energía potencial arbitraria en el entorno de un punto de equilibrio estable*. El oscilador armónico simple es también importante porque el comportamiento de sistemas tales como las vibraciones de un medio elástico y del campo electromagnético en una cavidad se pueden describir como la superposición de un número infinito de osciladores armónicos simples. Al cuantificar esos sistemas nos encontramos entonces con la mecánica cuántica de muchos osciladores armónicos lineales de diferentes frecuencias. Por tal motivo, todas las teorías de campos modernas utilizan los resultados que vamos a obtener.

El Hamiltoniano del oscilador armónico simple es

$$H = \frac{p^2}{2m} + \frac{1}{2} m \omega^2 x^2 \quad (9.69)$$

y la ecuación de Schrödinger independiente del tiempo es entonces:

$$-\frac{\hbar^2}{2m} \frac{d^2\psi}{dx^2} + \frac{1}{2} m \omega^2 x^2 \psi = E \psi \quad (9.70)$$

Para aligerar las fórmulas introducimos en lugar de x y p los operadores adimensionales ξ y η definidos por:

$$x = \xi \sqrt{\frac{\hbar}{m\omega}}, \quad p = \eta \sqrt{m\hbar\omega} \quad (9.71)$$

Es fácil verificar que ξ y $\eta = -i d/d\xi$ cumplen la relación de conmutación

$$[\xi, \eta] = \xi\eta - \eta\xi = i \quad (9.72)$$

En términos de ξ y η el Hamiltoniano se escribe

$$H = \hbar\omega\mathcal{H}, \quad \mathcal{H} = \frac{1}{2}(\eta^2 + \xi^2) \quad (9.73)$$

y la ec. (9.70) en la forma

$$\frac{d^2\psi}{d\xi^2} + (2\varepsilon - \xi^2)\psi = 0, \quad E = \hbar\omega\varepsilon \quad (9.74)$$

Veamos el comportamiento de $\psi(\xi)$ para $\xi \rightarrow \pm\infty$. Para valores finitos de ε es fácil verificar que

$$\psi(\xi \rightarrow \pm\infty) \rightarrow e^{-\xi^2/2} \quad (9.75)$$

de manera que ψ tiene el comportamiento de una Gaussiana. Es inmediato verificar por sustitución directa en la (9.61) que

$$\psi_0(\xi) = e^{-\xi^2/2} \quad (9.76)$$

es una solución de la (9.74) y corresponde al autovalor $\varepsilon = 1/2$. En efecto, si $\varepsilon = 1/2$ se cumple:

$$\frac{d^2\psi_0}{d\xi^2} + (2\varepsilon - \xi^2)\psi_0 = -\psi_0 + \xi^2\psi_0 + (2\varepsilon - \xi^2)\psi_0 = 0 \quad (9.77)$$

Para encontrar las demás autofunciones y autovalores vamos a usar una sencilla y elegante técnica de operadores, que es diferente de los métodos que se emplean habitualmente en los textos elementales de Mecánica Cuántica. Hacemos así porque esta técnica es el prototipo de otras semejantes que se aplican en una variedad de problemas.

Nuestro método se funda en las propiedades de conmutación de ciertos operadores no Hermitianos oportunamente definidos, y permite encontrar sistemáticamente mediante un procedimiento recursivo todas las autofunciones y sus correspondientes autovalores a partir de ψ_0 y de ε . Para eso definimos el operador

$$a = \frac{1}{\sqrt{2}}(\xi + i\eta) = \frac{1}{\sqrt{2}}\left(\xi + \frac{d}{d\xi}\right) \quad (9.78)$$

que por supuesto no es Hermitiano, y su adjunto

$$a^\dagger = \frac{1}{\sqrt{2}}(\xi - i\eta) = \frac{1}{\sqrt{2}}\left(\xi - \frac{d}{d\xi}\right) \quad (9.79)$$

En términos de a y $a^\dagger$ el operador $\mathcal{H}$ se expresa como

$$\mathcal{H} = a^\dagger a + \frac{1}{2} \quad (9.80)$$

El conmutador de a y $a^\dagger$ es

$$aa^\dagger - a^\dagger a = 1 \quad (9.81)$$

Puesto que $\mathcal{H}$ y $a^\dagger a$ conmutan, las autofunciones de $\mathcal{H}$ y $a^\dagger a$ son las mismas, de modo que para encontrar los estados estacionarios es suficiente resolver el problema de autovalores de $a^\dagger a$. Si llamamos λ_n ($n = 0, 1, 2, \dots$) a los autovalores y ψ_n las correspondientes autofunciones, la ecuación que queremos resolver es

$$a^\dagger a \psi_n = \lambda_n \psi_n \quad (9.82)$$

Primero vamos a demostrar que los autovalores no pueden ser negativos. De la (9.82) obtenemos

$$(\psi_n, a^\dagger a \psi_n) = (a \psi_n, a \psi_n) = \lambda_n (\psi_n, \psi_n) \quad (9.83)$$

donde usamos la definición de operador adjunto. Puesto que la norma de una función no puede ser negativa, concluimos que

$$\lambda_n \geq 0 \quad (9.84)$$

Si ψ_k es una autofunción de $a^\dagger a$, entonces $a^\dagger \psi_k$ es también una autofunción; en efecto usando la relación de conmutación (9.81) vemos que:

$$(a^\dagger a) a^\dagger \psi_k = a^\dagger (a^\dagger a + 1) \psi_k = (\lambda_k + 1) a^\dagger \psi_k \quad (9.85)$$

Por lo tanto $a^\dagger \psi_k$ es una autofunción con autovalor $\lambda_k + 1$. Del mismo modo se obtiene

$$(a^\dagger a) a \psi_k = a (a^\dagger a - 1) \psi_k = (\lambda_k - 1) a \psi_k \quad (9.86)$$

que muestra que $a \psi_k$ es una autofunción de $a^\dagger a$ con autovalor $\lambda_k - 1$. Debido a estas propiedades $a^\dagger$ y a se denominan *operador de subida* y *operador de bajada*, respectivamente. Operando reiteradamente con $a^\dagger$ y a sobre una autofunción ψ_k dada, podemos generar nuevas autofunciones correspondientes a diferentes autovalores, del mismo modo como se suben o se bajan los peldaños de una escalera. Sin embargo, la condición (9.84) limita la cantidad de veces que se puede aplicar el operador de bajada, porque cuando se llega un autovalor $0 \leq \lambda_0 < 1$, la aplicación del operador de bajada no permite ya encontrar una nueva autofunción, pues sería una autofunción correspondiente a un autovalor que viola la condición (9.84). Por lo tanto para el peldaño más bajo de la escalera ($n = 0$) se debe cumplir

$$a^\dagger a \psi_0 = \lambda_0 \psi_0 \quad , \quad 0 \leq \lambda_0 < 1 \quad (9.97)$$

y también

$$a \psi_0 = 0 \quad (9.88)$$

y por consiguiente el menor autovalor de $a^\dagger a$ es

$$\lambda_0 = 0 \quad (9.89)$$

Partiendo entonces de ψ_0 y de $\lambda_0 = 0$ podemos obtener todas las demás autofunciones y autovalores por aplicación reiterada del operador de subida $a^\dagger$. Pero nosotros ya conocemos ψ_0 , que está dado por la (9.76):

$$\psi_0(\xi) = e^{-\xi^2/2} \quad (9.90)$$

Por consiguiente la n -ésima autofunción, y su correspondiente autovalor son

$$\psi_n \propto (a^\dagger)^n \psi_0(\xi) = \left(\frac{1}{\sqrt{2}} \left(\xi - \frac{d}{d\xi} \right) \right)^n e^{-\xi^2/2} \quad , \quad \lambda_n = n \quad (9.91)$$

Usando la (9.74) y la (9.80) obtenemos que

$$H\psi_n = \hbar\omega(n + \frac{1}{2})\psi_n \quad (9.92)$$

y por lo tanto los autovalores de la energía son

$$E_n = \hbar\omega(n + \frac{1}{2}) \quad , \quad n = 0, 1, 2, \dots \quad (9.93)$$

Observamos que a diferencia del caso clásico, la energía del oscilador *no es nula* en el estado fundamental ($n = 0$) sino que todavía vale $\hbar\omega/2$. Este resultado de la teoría de Schrödinger difiere del que se obtuvo en la Teoría Cuántica Antigua a partir de los postulados de cuantificación Planck y de Wilson-Sommerfeld. La energía $\hbar\omega/2$ se denomina *energía de punto cero* del oscilador armónico y su existencia es un fenómeno cuántico que se puede entender en base al principio de incerteza.

Veamos ahora las expresiones explícitas de las autofunciones (9.91). Es fácil verificar que ψ_n tiene la forma

$$\psi_n = C_n H_n(\xi) e^{-\xi^2/2} \quad (9.94)$$

donde $H_n(\xi)$ es un polinomio de grado n y C_n es una constante de normalización todavía no especificada. Los polinomios $H_n(\xi)$ se denominan *polinomios de Hermite* y se suelen definir de modo que el coeficiente de la potencia más elevada de ξ sea 2^n . Los primeros polinomios de Hermite son:

$$\begin{aligned} H_0(\xi) &= 1 & H_3(\xi) &= 8\xi^3 - 12\xi \\ H_1(\xi) &= 2\xi & H_4(\xi) &= 16\xi^4 - 48\xi^2 + 12 \\ H_2(\xi) &= 4\xi^2 - 2 & H_5(\xi) &= 32\xi^5 - 160\xi^3 + 120\xi \end{aligned} \quad (9.95)$$

Los polinomios de Hermite satisfacen la ecuación diferencial

$$\frac{d^2 H_n}{d\xi^2} - 2\xi \frac{dH_n}{d\xi} + 2nH_n = 0 \quad (9.96)$$

Una forma simple de definir los polinomios de Hermite es

$$H_n = (-1)^n e^{\xi^2} \frac{d^n}{d\xi^n} e^{-\xi^2} \quad (9.97)$$

Se puede ver que los polinomios de Hermite tienen *paridad definida* dada por $(-1)^n$ y que sus n raíces son todas reales. Por lo tanto ψ_n tiene n nodos.

Las autofunciones normalizadas son:

$$\psi_n(\xi) = \frac{1}{\sqrt{2^n n!}} \left(\frac{m\omega}{\pi\hbar} \right)^{1/4} H_n(\xi) e^{-\xi^2/2} \quad , \quad \xi = x \sqrt{\frac{m\omega}{\hbar}} \quad (9.98)$$

En la bibliografía citada el lector puede encontrar gráficos de las autofunciones (9.98).

10. FUERZAS CENTRALES, MOMENTO ANGULAR Y ÁTOMO DE HIDRÓGENO

Introducción

Nos ocuparemos ahora del movimiento de una partícula en tres dimensiones. En la Mecánica Cuántica, del mismo modo que en la Mecánica Clásica hay una clase muy importante de problemas que surgen cuando las fuerzas son centrales, esto es, cuando la energía potencial depende sólo de la distancia r desde la partícula a un punto fijo (que tomaremos como el origen de coordenadas). El Hamiltoniano es entonces (μ indica la masa de la partícula)

$$H = \frac{\mathbf{p}^2}{2\mu} + V(r) = -\frac{\hbar^2}{2\mu} \nabla^2 + V(r) \quad (10.1)$$

Puesto que una fuerza central no produce momento respecto del origen, clásicamente se conserva el *momento angular* (orbital), que se define como

$$\mathbf{L} = \mathbf{r} \times \mathbf{p} \quad (10.2)$$

En la Mecánica Clásica, la conservación del momento angular para problemas de fuerzas centrales tiene importantes consecuencias: la órbita de la partícula está contenida en un plano (normal a la dirección del momento angular) y se cumple la Ley de las Áreas (la velocidad areolar es constante y proporcional al módulo del momento angular). De resultados de ello, si se describe el movimiento en el plano de la órbita, el problema es separable en coordenadas polares con origen en el centro de fuerzas. Gracias a esto el análisis del movimiento en tres dimensiones se simplifica enormemente, pues queda reducido esencialmente al estudio del movimiento radial, esto es, un movimiento *unidimensional* bajo la acción de la fuerza central más la fuerza centrífuga (cuyo valor está determinado por el módulo del momento angular). De acuerdo con el principio de correspondencia cabe esperar que el momento angular juegue también un rol muy importante en la Mecánica Cuántica. Veremos, en efecto, que debido a la conservación del momento angular se simplifica considerablemente la solución del problema, ya que la ecuación de Schrödinger independiente del tiempo se separa en coordenadas polares esféricas. En consecuencia las autofunciones de la energía se pueden escribir como el producto una función de los ángulos (que depende del momento angular y no depende de la forma de la energía potencial $V(r)$) y una función radial donde está comprendida la totalidad de los efectos de $V(r)$.

Propiedades del momento angular

El operador $\mathbf{L}$ que representa el momento angular se obtiene reemplazando $\mathbf{p}$ en la (10.2) por $-i\hbar\nabla$. Aquí no hay dificultad con operadores que no conmutan, pues en la (10.2) sólo aparecen productos como xp_y , yp_x , etc.. Comenzamos por obtener las relaciones de conmutación del momento angular usando las reglas de conmutación entre las componentes de $\mathbf{r}$ y $\mathbf{p}$. Por ejemplo, es fácil verificar que

$$\begin{aligned} [L_x, x] &= 0 & , & \quad [L_x, p_x] = 0 \\ [L_x, y] &= i\hbar z & , & \quad [L_x, p_y] = i\hbar p_z \\ [L_x, z] &= -i\hbar y & , & \quad [L_x, p_z] = -i\hbar p_y \end{aligned} \quad (10.3)$$

Relaciones semejantes valen para los conmutadores de L_y y L_z con las componentes de $\mathbf{r}$ y $\mathbf{p}$. A partir de ellas podemos calcular las relaciones de conmutación de las componentes de $\mathbf{L}$ entre sí:

$$[L_x, L_y] = i\hbar L_z \quad , \quad [L_y, L_z] = i\hbar L_x \quad , \quad [L_z, L_x] = i\hbar L_y \quad (10.4)$$

El hecho que las componentes de $\mathbf{L}$ no conmutan implica que *no existen en general estados en que todas las componentes del momento angular tengan valores bien definidos.*

Si aplicamos la relación general (8.51):

$$\Delta F \Delta G \geq \frac{1}{2} |\overline{[F, G]}| \quad (10.5)$$

obtenemos que

$$\Delta L_x \Delta L_y \geq \frac{\hbar}{2} |\overline{L_z}| \quad , \quad \Delta L_y \Delta L_z \geq \frac{\hbar}{2} |\overline{L_x}| \quad , \quad \Delta L_z \Delta L_x \geq \frac{\hbar}{2} |\overline{L_y}| \quad (10.6)$$

Las (10.6) muestran que las tres componentes de $\mathbf{L}$ se pueden determinar con exactitud sólo si los valores esperados de las mismas son nulos, esto es si

$$\overline{\mathbf{L}} = 0 \quad (10.7)$$

Además, si las componentes de $\mathbf{L}$ están bien definidas se debe cumplir

$$\Delta L_j^2 = \overline{(L_j - \overline{L_j})^2} = \overline{L_j^2} = 0 \quad , \quad j = x, y, z \quad (10.8)$$

y por lo tanto para tener $\Delta \mathbf{L} = 0$ es preciso que además de la (10.7) se cumpla también

$$\overline{L_x^2} = \overline{L_y^2} = \overline{L_z^2} = 0 \quad (10.9)$$

En consecuencia el único estado ψ en el cual se pueden conocer con exactitud las tres componentes de $\mathbf{L}$ es un autoestado de las tres componentes con autovalores nulos, tal que:

$$\mathbf{L}\psi = 0 \quad (10.10)$$

Magnitud del momento angular

El cuadrado de la magnitud de $\mathbf{L}$ es

$$\mathbf{L}^2 = L_x^2 + L_y^2 + L_z^2 \quad (10.11)$$

El valor medio de $\mathbf{L}^2$ en un estado cualquiera ϕ es siempre mayor o igual que el valor medio del cuadrado de una cualquiera de las componentes de $\mathbf{L}$, digamos L_z , esto es:

$$\overline{\mathbf{L}^2} = \overline{L_x^2} + \overline{L_y^2} + \overline{L_z^2} \geq \overline{L_z^2} \quad (10.12)$$

puesto que

$$\overline{L_x^2}, \overline{L_y^2} \geq 0 \quad (10.13)$$

El signo = en (10.12) vale sólo si $\mathbf{L}\psi = 0$. En efecto

$$\begin{aligned}\overline{L_x^2} &= (\phi, L_x^2 \phi) = (L_x \phi, L_x \phi) = 0 \Rightarrow L_x \phi = 0 \\ \overline{L_y^2} &= (\phi, L_y^2 \phi) = (L_y \phi, L_y \phi) = 0 \Rightarrow L_y \phi = 0\end{aligned}\quad (10.14)$$

y entonces usando

$$L_z = -\frac{i}{\hbar} [L_x, L_y] \quad (10.15)$$

encontramos que

$$L_z \phi = 0 \quad (10.16)$$

y en vista de (10.14) y (10.16) concluimos que

$$\overline{L^2} > \overline{L_z^2} \quad (10.17)$$

para todo ϕ , excepto cuando se cumple la condición (10.10). Vemos entonces que, salvo cuando todas las componentes de $\mathbf{L}$ son nulas, la magnitud del momento angular supera el máximo valor de cualquiera de sus componentes.

Es fácil verificar que L^2 conmuta con todas las componentes de $\mathbf{L}$:

$$[L^2, L] = 0 \quad (10.18)$$

Por lo tanto L^2 y una cualquiera de sus componentes, digamos L_z , poseen un sistema completo de autofunciones en común.

Traslaciones y rotaciones infinitesimales y sus generadores

Hay una estrecha relación entre el impulso y las traslaciones, y entre el momento angular y las rotaciones rígidas. Sea $\phi(\mathbf{r})$ una función *escalar* diferenciable arbitraria. Si desplazamos el valor de ϕ a un nuevo punto $\mathbf{r} + \mathbf{a}$ (Fig. 10.1) obtenemos una nueva función $\Phi(\mathbf{r})$ tal que

$$\Phi(\mathbf{r} + \mathbf{a}) = \phi(\mathbf{r}) \quad (10.19)$$

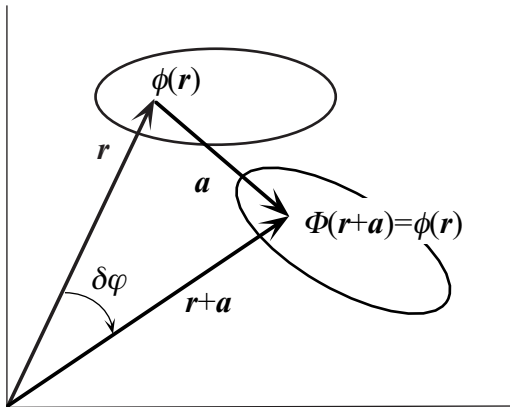


Fig. 10.1. Rotación rígida.

Para una traslación infinitesimal $\delta \mathbf{a}$:

$$\Phi(\mathbf{r}) = \phi(\mathbf{r} - \delta \mathbf{a}) = \phi(\mathbf{r}) - \delta \mathbf{a} \cdot \nabla \phi(\mathbf{r}) \quad (10.20)$$

y por lo tanto

$$\delta \phi(\mathbf{r}) = \Phi(\mathbf{r}) - \phi(\mathbf{r}) = -\frac{i}{\hbar} \delta \mathbf{a} \cdot \mathbf{p} \phi(\mathbf{r}) \quad (10.21)$$

Por este motivo el operador $\mathbf{p}$ se denomina *generador de traslaciones infinitesimales*.

Si el desplazamiento es una rotación infinitesimal de un ángulo $\delta \varphi$ alrededor de un eje que pasa por el origen,

$$\mathbf{a} = \delta\boldsymbol{\varphi} \times \mathbf{r} \quad (10.22)$$

donde $\delta\boldsymbol{\varphi}$ es un vector cuya magnitud es $\delta\varphi$ y que tiene la dirección del eje de rotación. La variación de ϕ es entonces:

$$\delta\phi(\mathbf{r}) = \Phi(\mathbf{r}) - \phi(\mathbf{r}) = -(\delta\boldsymbol{\varphi} \times \mathbf{r}) \cdot \nabla\phi(\mathbf{r}) = -\delta\boldsymbol{\varphi} \cdot (\mathbf{r} \times \nabla)\phi(\mathbf{r}) \quad (10.23)$$

es decir

$$\delta\phi(\mathbf{r}) = -\frac{i}{\hbar} \delta\boldsymbol{\varphi} \cdot \mathbf{L}\phi(\mathbf{r}) \quad (10.24)$$

Por este motivo el operador $\mathbf{L}$ se denomina *generador de rotaciones infinitesimales*.

De la (10.10) y la (10.24) se desprende que la función de onda ψ del estado en el cual las tres componentes de $\mathbf{L}$ están bien definidas es *invariante* ante una rotación infinitesimal arbitraria. Por lo tanto $\psi = \psi(r)$ depende sólo de r y no de los ángulos y describe un estado *esféricamente simétrico*. Tales estados de momento angular nulo se denominan *estados S*.

Dejamos como ejercicio para el lector comprobar que la relación fundamental (10.24) implica las relaciones de conmutación (10.3).

Fuerzas centrales y conservación del momento angular

Es fácil verificar usando las relaciones de conmutación (10.3) (lo dejamos como ejercicio) que

$$[\mathbf{L}, \mathbf{p}^2] = 0 \quad (10.25)$$

Además, poniendo $\phi = V(r)\psi(\mathbf{r})$ en la (10.24) obtenemos

$$\delta[V(r)\psi(\mathbf{r})] = V(r)\delta\psi(\mathbf{r}) \Rightarrow -\frac{i}{\hbar} \delta\boldsymbol{\varphi} \cdot \mathbf{L}V(r)\psi(\mathbf{r}) = -\frac{i}{\hbar} \delta\boldsymbol{\varphi} \cdot V(r)\mathbf{L}\psi(\mathbf{r}) \quad (10.26)$$

Puesto que $\delta\boldsymbol{\varphi}$ y $\psi(\mathbf{r})$ son arbitrarios, concluimos que

$$[\mathbf{L}, V(r)] = 0 \quad (10.27)$$

Si $V = V(r)$, en virtud de las (10.25) y (10.27) obtenemos

$$[\mathbf{L}, H] = 0 \quad (10.28)$$

y por lo tanto *en presencia de fuerzas centrales el momento angular es una constante del movimiento*, tal como ocurre en la Mecánica Clásica.

Energía cinética y momento angular

En Mecánica Clásica se usa la identidad

$$\mathbf{L}^2 = (\mathbf{r} \times \mathbf{p})^2 = r^2 p^2 - (\mathbf{r} \cdot \mathbf{p})^2 \quad (10.29)$$

para expresar la energía cinética en términos de la componente radial del impulso y de L^2 que es constante del movimiento, y así el problema de fuerzas centrales se reduce a un problema unidimensional equivalente. Aquí tenemos que proceder con cuidado, pues

$$(\mathbf{r} \times \mathbf{p})^2 \neq r^2 p^2 - (\mathbf{r} \cdot \mathbf{p})^2 \quad (10.30)$$

debido a que $\mathbf{r}$ y $\mathbf{p}$ no conmutan. La expresión correcta es

$$L^2 = (\mathbf{r} \times \mathbf{p}) \cdot (\mathbf{r} \times \mathbf{p}) = r^2 p^2 - \mathbf{r}(\mathbf{r} \cdot \mathbf{p}) \cdot \mathbf{p} + 2i\hbar \mathbf{r} \cdot \mathbf{p} \quad (10.31)$$

lo cual se verifica con un poco de paciencia desarrollando el producto $(\mathbf{r} \times \mathbf{p}) \cdot (\mathbf{r} \times \mathbf{p})$. Puesto que

$$\mathbf{r} \cdot \mathbf{p} = -i\hbar r \frac{\partial}{\partial r} \quad (10.32)$$

la (10.31) se escribe

$$L^2 = r^2 p^2 + \hbar^2 r^2 \frac{\partial^2}{\partial r^2} + 2\hbar^2 r \frac{\partial}{\partial r} = r^2 p^2 + \hbar^2 \frac{\partial}{\partial r} \left(r^2 \frac{\partial}{\partial r} \right) \quad (10.33)$$

y como L conmuta con cualquier función de r , podemos escribir la energía cinética en la forma

$$T \equiv \frac{p^2}{2\mu} = \frac{L^2}{2\mu r^2} - \frac{\hbar^2}{2\mu r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial}{\partial r} \right) \quad (10.34)$$

Debemos recordar que para llegar a esta expresión hemos dividido la (10.33) por r^2 ; por lo tanto la (10.34) no vale para $r = 0$, salvo en el caso especial de los estados de momento angular nulo, para los cuales $L\psi = 0$.

Reducción del problema de fuerzas centrales

Cuando L^2 es constante del movimiento existe un sistema de autofunciones $\psi_{E,\lambda}(\mathbf{r})$ de H y de L^2 con autovalores E y $\lambda\hbar^2$, respectivamente:

$$H\psi_{E,\lambda}(\mathbf{r}) = E\psi_{E,\lambda}(\mathbf{r}) \quad , \quad L^2\psi_{E,\lambda}(\mathbf{r}) = \lambda\hbar^2\psi_{E,\lambda}(\mathbf{r}) \quad (10.35)$$

Aquí λ es un número puro, pues $\hbar$ tiene dimensiones de momento angular.

Usando la (10.34) la ecuación de Schrödinger independiente del tiempo se escribe en la forma

$$\left[-\frac{\hbar^2}{2\mu r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial}{\partial r} \right) + \frac{\hbar^2 \lambda}{2\mu r^2} + V(r) \right] \psi_{E,\lambda}(\mathbf{r}) = E\psi_{E,\lambda}(\mathbf{r}) \quad (10.36)$$

Puesto que L^2 conmuta con cualquier función de r , todas las soluciones de la (10.36) se pueden obtener como combinaciones lineales de soluciones separables en coordenadas esféricas, de la forma

$$\psi_{E,\lambda}(\mathbf{r}) = R_{E,\lambda}(r)Y_\lambda(\theta, \varphi) \quad (10.37)$$

donde $Y_\lambda(\theta, \varphi)$ es una solución de

$$L^2 Y_\lambda(\theta, \varphi) = \lambda \hbar^2 Y_\lambda(\theta, \varphi) \quad (10.38)$$

y la autofunción radial $R_{E,\lambda}(r)$ satisface la ecuación diferencial ordinaria

$$\left[-\frac{\hbar^2}{2\mu r^2} \frac{d}{dr} \left(r^2 \frac{d}{dr} \right) + \frac{\hbar^2 \lambda}{2\mu r^2} + V(r) \right] R_{E,\lambda}(r) = E R_{E,\lambda}(r) \quad (10.39)$$

Si hacemos la sustitución

$$R_{E,\lambda}(r) = \frac{u_{E,\lambda}(r)}{r} \quad (10.40)$$

encontramos que u satisface la ecuación radial:

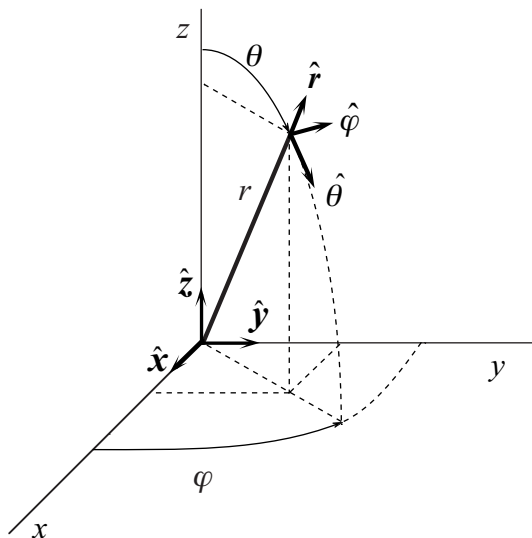
$$-\frac{\hbar^2}{2\mu} \frac{d^2 u_{E,\lambda}}{dr^2} + \left[\frac{\hbar^2 \lambda}{2\mu r^2} + V(r) \right] u_{E,\lambda} = E u_{E,\lambda} \quad (10.41)$$

Notar el término $\hbar^2 \lambda / 2\mu r^2$ que se suma a la energía potencial. Este término se denomina *potencial centrífugo* pues representa el potencial del cual deriva la fuerza centrífuga.

La (10.41) se parece a la ecuación de Schrödinger independiente del tiempo *en una dimensión* que estudiamos en el Capítulo 9, pero las condiciones de contorno para $u_{E,\lambda}$ son *diferentes* pues r no puede ser negativo. Por ejemplo, para que ψ sea finita en todas partes, u se debe anular en $r = 0$. Dado que las condiciones de contorno dependen de $V(r)$ dejaremos para más adelante su discusión y estudiaremos primero los autovalores y autofunciones del momento angular, que son aspectos comunes a todos los problemas de fuerzas centrales.

Dependencia angular de las autofunciones

Conviene expresar el momento angular en coordenadas polares esféricas (Fig. 10.2)



$$\begin{aligned} x &= r \sin \theta \cos \varphi \\ y &= r \sin \theta \sin \varphi \\ z &= r \cos \theta \end{aligned} \quad (10.42)$$

El operador ∇ se escribe entonces

$$\nabla = \hat{\mathbf{r}} \frac{\partial}{\partial r} + \hat{\boldsymbol{\theta}} \frac{1}{r} \frac{\partial}{\partial \theta} + \hat{\boldsymbol{\varphi}} \frac{1}{r \sin \theta} \frac{\partial}{\partial \varphi} \quad (10.43)$$

donde

$$\begin{aligned} \hat{\mathbf{r}} &= \hat{\mathbf{x}} \sin \theta \cos \varphi + \hat{\mathbf{y}} \sin \theta \sin \varphi + \hat{\mathbf{z}} \cos \theta \\ \hat{\boldsymbol{\theta}} &= \hat{\mathbf{x}} \cos \theta \cos \varphi + \hat{\mathbf{y}} \cos \theta \sin \varphi - \hat{\mathbf{z}} \sin \theta \\ \hat{\boldsymbol{\varphi}} &= -\hat{\mathbf{x}} \sin \varphi + \hat{\mathbf{y}} \cos \varphi \end{aligned} \quad (10.44)$$

Fig. 10.2. Coordenadas polares esféricas.

La (10.43) muestra que (a diferencia de las compo-

nentes cartesianas) las componentes polares del impulso no conmutan entre sí.

El momento angular se puede expresar como

$$\mathbf{L} = \mathbf{r} \times \mathbf{p} = -i\hbar r \hat{\mathbf{r}} \times \nabla = -i\hbar \left(\hat{\boldsymbol{\phi}} \frac{\partial}{\partial \theta} - \hat{\boldsymbol{\theta}} \frac{1}{\sin \theta} \frac{\partial}{\partial \varphi} \right) \quad (10.45)$$

De la (10.45) es evidente que $\mathbf{L}$ conmuta con cualquier función de r y con cualquier derivada con respecto de r como ya sabíamos.

Usando (10.45) y (10.44) obtenemos

$$\begin{aligned} L_x &= -i\hbar \left(-\sin \varphi \frac{\partial}{\partial \theta} - \cos \varphi \cot \theta \frac{\partial}{\partial \varphi} \right) \\ L_y &= -i\hbar \left(\cos \varphi \frac{\partial}{\partial \theta} - \sin \varphi \cot \theta \frac{\partial}{\partial \varphi} \right) \\ L_z &= -i\hbar \frac{\partial}{\partial \varphi} \end{aligned} \quad (10.46)$$

y finalmente obtenemos la expresión de $\mathbf{L}^2$ en coordenadas polares esféricas:

$$\mathbf{L}^2 = -\hbar^2 \left[\frac{1}{\sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial}{\partial \theta} \right) + \frac{1}{\sin^2 \theta} \frac{\partial^2}{\partial \varphi^2} \right] \quad (10.47)$$

El problema de autovalores de $\mathbf{L}^2$ es entonces

$$\left[\frac{1}{\sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial}{\partial \theta} \right) + \frac{1}{\sin^2 \theta} \frac{\partial^2}{\partial \varphi^2} \right] Y_\lambda(\theta, \varphi) = -\lambda Y_\lambda(\theta, \varphi) \quad (10.48)$$

Podemos resolver esta ecuación por separación de variables poniendo

$$Y(\theta, \varphi) = \Theta(\theta) \Phi(\varphi) \quad (10.49)$$

y resulta

$$\frac{\sin^2 \theta}{\Theta} \left[\frac{1}{\sin \theta} \frac{d}{d\theta} \left(\sin \theta \frac{d\Theta}{d\theta} \right) + \lambda \Theta \right] = -\frac{1}{\Phi} \frac{d^2 \Phi}{d\varphi^2} \quad (10.50)$$

Si indicamos la constante de separación como m^2 (veremos en seguida que m es real en virtud de las condiciones de contorno que imponemos) obtenemos una ecuación diferencial ordinaria para Φ :

$$\frac{d^2 \Phi}{d\varphi^2} + m^2 \Phi = 0 \quad (10.51)$$

y otra para Θ :

$$\frac{1}{\sin \theta} \frac{d}{d\theta} \left(\sin \theta \frac{d\Theta}{d\theta} \right) - \frac{m^2}{\sin^2 \theta} \Theta + \lambda \Theta = 0 \quad (10.52)$$

El problema de autovalores para L_z y la cuantificación espacial

Para resolver la (10.51) vamos a requerir que

$$\Phi(\varphi + 2\pi) = \Phi(\varphi) \quad (10.53)$$

para que la función de onda sea *univaluada*. Esto nos restringe a las soluciones

$$\Phi(\varphi) = e^{im\varphi} \quad , \quad m = 0, \pm 1, \pm 2, \dots \quad (10.54)$$

que son autofunciones del operador L_z , pues de acuerdo con la (10.46):

$$L_z e^{im\varphi} = -i\hbar \frac{\partial}{\partial \varphi} e^{im\varphi} = m\hbar e^{im\varphi} \quad , \quad m = 0, \pm 1, \pm 2, \dots \quad (10.55)$$

Puesto que L_z conmuta con L^2 ambos operadores tienen un sistema de autofunciones en común, por lo cual el resultado (10.55) no es sorprendente.

Incidentalmente, se puede observar que la (10.51) admite soluciones de la forma $\cos m\varphi$ y $\sin m\varphi$ que también cumplen la condición (10.53), pero no son autofunciones de L_z . Asimismo L_x y L_y conmutan con L^2 y por lo tanto también cualquiera de ellas tiene un sistema de autofunciones en común con L^2 , pero esas autofunciones tienen una forma mucho más complicada debido a que usamos el eje z como el eje de las coordenadas polares esféricas y por eso en las expresiones de L_x y L_y figuran θ y φ , mientras que en la de L_z figura solamente φ .

La ec. (10.55) muestra que los autovalores de L_z son $m\hbar$; por consiguiente una medición de L_z puede dar como resultado únicamente múltiplos enteros de $\hbar$. Puesto que la dirección del eje z es arbitraria, esto implica que *el momento angular respecto de cualquier eje está cuantificado* y que una medición del mismo sólo puede dar como resultado uno de los valores discretos

$$0, \pm \hbar, \pm 2\hbar, \pm 3\hbar, \dots \quad (10.56)$$

Este hecho se conoce como *cuantificación espacial*.

El número m se suele denominar *número cuántico magnético* debido al papel que juega en la descripción del efecto de un campo magnético uniforme $\mathbf{B}$ sobre una partícula cargada que se mueve en un campo de fuerzas centrales.

Teoría elemental del efecto Zeeman

Cuando un átomo está sometido a un campo magnético externo, se observa que las líneas espectrales se subdividen en varias componentes. Este fenómeno se denomina *efecto Zeeman*.

Se puede demostrar que cuando una partícula cuya carga es q se mueve en un campo magnético externo descrito por el potencial vectorial $\mathbf{A}$, el operador $\mathbf{p}$ se debe reemplazar por

$$\mathbf{p} \rightarrow \mathbf{p} - \frac{q}{c} \mathbf{A} \quad (10.57)$$

Por consiguiente el Hamiltoniano H se debe reemplazar por:

$$H = \frac{1}{2\mu} \mathbf{p}^2 + V \rightarrow H' = \frac{1}{2\mu} \left(\mathbf{p} - \frac{q}{c} \mathbf{A} \right)^2 + V \quad (10.58)$$

Para un campo magnético *uniforme* $\mathbf{A} = (\mathbf{B} \times \mathbf{r})/2$; la ecuación de Schrödinger independiente del tiempo se escribe entonces (despreciando términos cuadráticos en $\mathbf{B}$) como

$$\left[\frac{p^2}{2\mu} - \frac{q}{2\mu c} \mathbf{B} \cdot \mathbf{L} + V(r) \right] \psi = E' \psi \quad (10.59)$$

Por lo tanto en primera aproximación el efecto de un campo magnético es agregar a la energía un término de energía magnética $-\mathbf{M} \cdot \mathbf{B}$, donde $\mathbf{M}$ es el momento magnético efectivo del sistema y está relacionado con el momento angular por

$$\mathbf{M} = \frac{q}{2\mu c} \mathbf{L} \quad (10.60)$$

Esta ecuación muestra que la relación entre las magnitudes del momento magnético y del momento angular de un electrón ($q = -e$) es la constante $e/2\mu c$. Es usual expresar dicha constante como

$$\frac{e}{2\mu c} = \frac{g_L \beta_B}{\hbar} \quad , \quad g_L = 1 \quad , \quad \beta_B = \frac{e\hbar}{2\mu c} = 0.927 \times 10^{-20} \text{ erg/gauss} \quad (10.61)$$

La cantidad β_B se denomina *magnetón de Bohr* y es la unidad natural para los momentos magnéticos atómicos; el número g_L se llama *factor giromagnético* (o *factor de Landé*) orbital y lo introducimos explícitamente pese a ser igual a la unidad, pues veremos que hay otros factores giromagnéticos que tienen un valor diferente.

Si se toma el eje z en la dirección del campo magnético, la perturbación magnética $-\mathbf{M} \cdot \mathbf{B}$ contiene sólo el operador L_z y por lo tanto las autofunciones simultáneas de H y L_z son también autofunciones de H' . Entonces es inmediato comprobar que en presencia del campo magnético los autovalores de la energía de un electrón pasan a ser:

$$E \rightarrow E' = E - \frac{q}{2\mu c} m\hbar = E + \beta_B m \quad (10.62)$$

Este resultado explica el origen del término “número cuántico magnético” para designar a m . La ecuación (10.62) implica que los estados con el mismo E y diferente m , que en ausencia del campo magnético estaban degenerados, se separan debido a que la interacción con el campo depende de la orientación de $\mathbf{L}$ (y por lo tanto del momento magnético $\mathbf{M}$) con respecto de $\mathbf{B}$.

Esta es la teoría elemental del efecto Zeeman. Para la mayoría de las aplicaciones es necesario refinarla, para tomar en cuenta el spin del electrón y otras correcciones a la energía que son importantes a menos que B sea muy grande.

Autovalores y autofunciones de L^2

Volvemos ahora al estudio de la ec. (10.52):

$$\frac{1}{\sin\theta} \frac{d}{d\theta} \left(\sin\theta \frac{d\Theta}{d\theta} \right) - \frac{m^2}{\sin^2\theta} \Theta + \lambda\Theta = 0 \quad (10.63)$$

Conviene hacer el cambio de variable

10. Fuerzas centrales, momento angular y átomo de hidrógeno

$$\xi = \cos \theta \quad , \quad F(\xi) = \Theta(\theta) \quad (10.64)$$

que nos lleva a la ecuación

$$\frac{d}{d\xi} \left[(1 - \xi^2) \frac{dF}{d\xi} \right] - \frac{m^2}{1 - \xi^2} F + \lambda F = 0 \quad (10.65)$$

En el caso particular $m = 0$ la (10.65) se reduce a la conocida *ecuación de Legendre*:

$$\frac{d}{d\xi} \left[(1 - \xi^2) \frac{dF}{d\xi} \right] + \lambda F = 0 \quad (10.66)$$

Puesto que la ecuación de Legendre es invariante frente al cambio $\xi \rightarrow -\xi$ (es decir $\theta \rightarrow \pi - \theta$) basta buscar soluciones de paridad definida (esto es funciones simétricas o antisimétricas respecto del plano $z = 0$).

La solución regular de la (10.66) se puede expandir en serie

$$F(\xi) = \sum_k a_k \xi^k \quad (10.67)$$

Sustituyendo (10.67) en (10.66) obtenemos una relación de recurrencia para los coeficientes de la serie:

$$a_{k+2} = \frac{k(k+1) - \lambda}{(k+1)(k+2)} a_k \quad (10.68)$$

La (10.68) muestra que para k par se cumple $a_{-2} = 0$ y para k impar $a_{-1} = 0$, de acuerdo a la hipótesis que F es regular en $\xi = 0$. Por otra parte, si la serie no termina para algún valor finito de k , se tiene que

$$\left(\frac{a_{k+2}}{a_k} \right)_{k \rightarrow \infty} \rightarrow 1 \quad (10.69)$$

y por lo tanto la serie diverge para $\xi \rightarrow \pm 1$. Por la misma razón excluimos la segunda solución linealmente independiente de la (10.66) que diverge logarítmicamente en $\xi \rightarrow \pm 1$.

Concluimos que para tener una solución aceptable, la serie se debe cortar en un valor finito $k = l$, donde l es un entero no negativo, y en ese caso F es un polinomio de grado l . Entonces los valores permitidos de λ son:

$$\lambda = l(l+1) \quad , \quad l = 0, 1, 2, \dots \quad (10.70)$$

El número l se llama *número cuántico del momento angular orbital*. Por tradición los correspondientes autoestados del momento angular orbital se designan con los símbolos $S, P, D, F, G, H, \dots$. Si hay varias partículas en el campo de fuerzas centrales se usan las minúsculas $s, p, d, f, \dots$.

$g, h, \dots$ para designar los estados de momento angular de cada partícula y se reservan las mayúsculas $S, P, D, F, G, H, \dots$ para el momento angular orbital total¹.

Los autovalores de L^2 son por lo tanto

$$l(l+1)\hbar^2 = 0, 2\hbar^2, 6\hbar^2, 12\hbar^2, \dots \quad (10.71)$$

Las soluciones polinomiales de la (10.66) se suelen definir como

$$P_l(\xi) = \frac{1}{2^l l!} \frac{d^l}{d\xi^l} (\xi^2 - 1)^l \quad (10.72)$$

y se denominan *polinomios de Legendre*. El factor constante en (10.72) se introduce para que

$$P_l(\pm 1) = (\pm 1)^l \quad (10.73)$$

Los primeros polinomios de Legendre son

$$\begin{aligned} P_0(\xi) &= 1 & P_3(\xi) &= \frac{1}{2}(5\xi^3 - 3\xi) \\ P_1(\xi) &= \xi & P_4(\xi) &= \frac{1}{8}(35\xi^4 - 30\xi^2 + 3) \\ P_2(\xi) &= \frac{1}{2}(3\xi^2 - 1) & P_5(\xi) &= \frac{1}{8}(63\xi^5 - 70\xi^3 + 15\xi) \end{aligned} \quad (10.74)$$

Puesto que los $P_l(\xi)$ son autofunciones del operador Hermitiano L^2 , es evidente que son ortogonales:

$$\int_{-1}^{+1} P_{l'}(\xi) P_l(\xi) d\xi = \delta_{l'l} \frac{2}{2l+1} \quad (10.75)$$

Una vez conocidas las soluciones de la ecuación de Legendre (10.66) no es difícil encontrar también las soluciones de la (10.65) con $m \neq 0$. Para eso usaremos una técnica semejante a la que empleamos en el Capítulo 9 al estudiar el oscilador armónico. Consideremos los operadores no Hermitianos

$$L_+ = L_x + iL_y, \quad L_- = L_x - iL_y \quad (10.76)$$

cuya forma explícita es

$$L_+ = \hbar e^{i\varphi} \left(\frac{\partial}{\partial \theta} + i \cot \theta \frac{\partial}{\partial \varphi} \right), \quad L_- = -\hbar e^{-i\varphi} \left(\frac{\partial}{\partial \theta} - i \cot \theta \frac{\partial}{\partial \varphi} \right) \quad (10.77)$$

Es fácil comprobar que L_+ y L_- cumplen las siguientes relaciones de conmutación:

$$[L^2, L_{\pm}] = 0, \quad [L_z, L_+] = \hbar L_+, \quad [L_z, L_-] = -\hbar L_- \quad (10.78)$$

y satisfacen las siguientes identidades que nos serán de utilidad:

¹ Los primeros cuatro símbolos son las iniciales de los nombres de las series espectrales Sharp, Principal, Diffuse y Fundamental.

10. Fuerzas centrales, momento angular y átomo de hidrógeno

$$L_-L_+ = \mathbf{L}^2 - L_z^2 - \hbar L_z \quad , \quad L_+L_- = \mathbf{L}^2 - L_z^2 + \hbar L_z \quad (10.79)$$

Indicamos ahora con

$$Y_l^m(\theta, \varphi) = \Theta_l^m(\theta)e^{im\varphi} \quad (10.80)$$

las autofunciones de $\mathbf{L}^2$ y L_z . Conocemos ya algunas de ellas, las que corresponden a $m = 0$:

$$Y_l^0(\theta, \varphi) = \Theta_l^0(\theta) = P_l(\cos\theta) \quad (10.81)$$

y ahora mostraremos como se obtienen las demás.

Si aplicamos L_z a la función $L_+Y_l^m$ y usamos la (10.78) obtenemos:

$$L_z(L_+Y_l^m) = \{L_+L_z + [L_z, L_+]\}Y_l^m = \{L_+L_z + \hbar L_+\}Y_l^m = (m+1)\hbar(L_+Y_l^m) \quad (10.82)$$

Esto muestra que $L_+Y_l^m$ es una autofunción de L_z correspondiente al autovalor $(m+1)\hbar$, y por lo tanto es proporcional a Y_l^{m+1} . Escribimos entonces

$$L_+Y_l^m = C_+(l, m)\hbar Y_l^{m+1} \quad (10.83)$$

donde $C_+(l, m)$ es la constante de proporcionalidad. De igual modo si aplicamos L_z a la función $L_-Y_l^m$. resulta

$$L_z(L_-Y_l^m) = \{L_-L_z + [L_z, L_-]\}Y_l^m = \{L_-L_z - \hbar L_-\}Y_l^m = (m-1)\hbar(L_-Y_l^m) \quad (10.84)$$

y por lo tanto $L_-Y_l^m$ es una autofunción de L_z correspondiente al autovalor $(m-1)\hbar$, y debe ser proporcional a Y_l^{m-1} ; entonces

$$L_-Y_l^m = C_-(l, m)\hbar Y_l^{m-1} \quad (10.85)$$

donde $C_-(l, m)$ es una constante.

Por consiguiente, aplicando reiteradamente los operadores L_+ y L_- podemos obtener a partir de Y_l^0 todas las diferentes Y_l^m para el l dado.

Observemos que para cada l , la condición (10.17):

$$\overline{\mathbf{L}^2} > \overline{L_z^2} \quad (10.86)$$

implica

$$\lambda = l(l+1) \geq m^2 \quad (10.87)$$

En consecuencia para cada l debe haber un valor máximo de m , $m = q$, tal que la aplicación de L_+ a Y_l^q no nos da una nueva autofunción, es decir

$$L_+Y_l^q = 0 \quad (10.88)$$

Multiplicando la (10.88) por L_- y usando la (10.79) obtenemos

$$L_- L_+ Y_l^q = (\mathbf{L}^2 - L_z^2 - \hbar L_z) Y_l^q = \hbar^2 [l(l+1) - q(q+1)] Y_l^q = 0 \quad (10.89)$$

de donde resulta que

$$q(q+1) = l(l+1) \quad (10.90)$$

La (10.90) nos da dos posibles valores de q : l y $-(l+1)$. Obviamente debemos descartar el segundo y por lo tanto el máximo valor de m es $m = l$.

Análogamente debe haber un valor mínimo de m , $m = q'$ tal que

$$L_- Y_l^{q'} = 0 \quad (10.91)$$

Multiplicando la (10.91) por L_+ y usando la (10.79) obtenemos

$$L_+ L_- Y_l^{q'} = (\mathbf{L}^2 - L_z^2 + \hbar L_z) Y_l^{q'} = \hbar^2 [l(l+1) - q'(q'-1)] Y_l^{q'} = 0 \quad (10.92)$$

de donde resulta que

$$q'(q'-1) = l(l+1) \quad (10.93)$$

Las raíces de (10.93) son $-l$ y $l+1$ y puesto que la segunda de ellas se debe descartar, el mínimo valor de m es $m = -l$.

En consecuencia los posibles valores de m para un l dado son:

$$m = 0, \pm 1, \dots, \pm l \quad (10.94)$$

o sea en total $2l+1$ valores.

Usando las identidades (10.79) podemos calcular los valores de las constantes $C_+(l, m)$ y $C_-(l, m)$, los resultados son:

$$C_+(l, m) = \sqrt{(l-m)(l+m+1)} \quad , \quad C_-(l, m) = \sqrt{(l+m)(l-m+1)} \quad (10.95)$$

Las autofunciones normalizadas $Y_l^m(\theta, \varphi)$ de $\mathbf{L}^2$ y L_z se denominan *armónicos esféricos* y su forma explícita para $m \geq 0$ es

$$Y_l^m(\theta, \varphi) = \sqrt{\frac{2l+1}{4\pi} \frac{(l-m)!}{(l+m)!}} (-1)^m e^{im\varphi} P_l^m(\cos\theta) \quad (m \geq 0) \quad (10.96)$$

Los armónicos esféricos con superíndice *negativo* se definen como

$$Y_l^m = (-1)^m (Y_l^{-m})^* \quad (m < 0) \quad (10.97)$$

En la (10.96) P_l^m indica las *funciones asociadas de Legendre*, que son soluciones de la ecuación (10.65), y se definen a partir de los polinomios de Legendre como

$$P_l^m(\xi) = (1-\xi^2)^{m/2} \frac{d^m}{d\xi^m} P_l(\xi) \quad (10.98)$$

Los armónicos esféricos forman un sistema ortonormal completo, de modo que

$$\int_{\varphi=0}^{2\pi} d\varphi \int_{\theta=0}^{\pi} [Y_l^m(\theta, \varphi)]^* Y_l^{m'}(\theta, \varphi) \sin\theta d\theta = \delta_{ll'} \delta_{mm'} \quad (10.99)$$

Podemos observar que en una *inversión* de coordenadas, en la cual

$$\theta \rightarrow \pi - \theta, \quad \varphi \rightarrow \varphi + \pi \quad (10.100)$$

la función $e^{im\varphi}$ queda multiplicada por $(-1)^m$ y $P_l^m(\cos\theta)$ por $(-1)^{l+m}$. Por consiguiente la *paridad* de los armónicos esféricos está dada por $(-1)^l$.

Los primeros armónicos esféricos son

$$\begin{aligned} Y_0^0 &= \frac{1}{\sqrt{4\pi}} \\ Y_1^0 &= \sqrt{\frac{3}{4\pi}} \cos\theta = \sqrt{\frac{3}{4\pi}} \frac{z}{r} \\ Y_1^{\pm 1} &= \mp \sqrt{\frac{3}{8\pi}} e^{\pm i\varphi} \sin\theta = \mp \sqrt{\frac{3}{8\pi}} \frac{x \pm iy}{r} \\ Y_2^0 &= \sqrt{\frac{5}{16\pi}} (3\cos^2\theta - 1) = \sqrt{\frac{5}{16\pi}} \frac{2z^2 - x^2 - y^2}{r^2} \\ Y_2^{\pm 1} &= \mp \sqrt{\frac{15}{8\pi}} e^{\pm i\varphi} \cos\theta \sin\theta = \mp \sqrt{\frac{15}{8\pi}} \frac{(x \pm iy)z}{r^2} \\ Y_2^{\pm 2} &= \sqrt{\frac{15}{32\pi}} e^{\pm 2i\varphi} \sin^2\theta = \sqrt{\frac{15}{32\pi}} \frac{(x \pm iy)^2}{r^2} \end{aligned} \quad (10.101)$$

La ecuación radial

Volvemos ahora a la discusión de la ecuación radial. De lo visto antes, concluimos que en presencia de fuerzas centrales las soluciones de la ecuación de Schrödinger independiente del tiempo son de la forma

$$\psi(r, \theta, \varphi) = \frac{u(r)}{r} Y_l^m(\theta, \varphi) \quad (10.102)$$

Puesto que r no cambia por una inversión de coordenadas, esas funciones de onda tienen la misma paridad que los armónicos esféricos. Por lo tanto para l *par* tendremos estados de paridad *par* y para l *impar* estados de paridad *impar*.

La función de onda radial $u(r)$ debe ser solución de la ec. (10.41):

$$-\frac{\hbar^2}{2\mu} \frac{d^2 u_{E,\lambda}}{dr^2} + \left[\frac{\hbar^2 l(l+1)}{2\mu r^2} + V(r) \right] u_{E,\lambda} = E u_{E,\lambda} \quad (10.103)$$

Puesto que los armónicos esféricos están normalizados, la normalización de la función de onda para las autofunciones correspondientes a los autovalores discretos de la (10.103) requiere que

10. Fuerzas centrales, momento angular y átomo de hidrógeno

$$\int \int \psi^* \psi r^2 dr d\Omega = \int_0^\infty u^* u dr = 1 \quad (10.104)$$

La normalización de las autofunciones del espectro continuo requiere

$$\int_0^\infty u_E^* u_{E'} dr = \delta(E - E') \quad (10.105)$$

En la mayoría de los casos de interés podemos suponer que $V(r)$ es finito en todas partes, salvo eventualmente en el origen, y que cerca de $r = 0$

$$V(r) \approx cr^\alpha, \quad (r \rightarrow 0, \quad \alpha = \text{entero} \geq -1) \quad (10.106)$$

Además vamos a suponer que

$$V \rightarrow 0, \quad (r \rightarrow \infty) \quad (10.107)$$

Recordemos que la (10.34) no vale en el origen y por lo tanto la (10.103) vale sólo para $r \neq 0$ y la debemos complementar con una condición de contorno en $r = 0$. Sin entrar en mayores detalles podemos decir que la condición de contorno adecuada se obtiene imponiendo que H sea Hermitiano. En los casos que nos interesan aquí, esto se consigue usando la condición

$$u(0) = 0 \quad (10.108)$$

Vale la pena aclarar que esta sencilla condición es *suficiente pero no necesaria*. En ciertos casos, por ejemplo en la teoría relativística del átomo de hidrógeno, se pueden aceptar funciones de onda con una singularidad débil en el origen.

Cuando $l \neq 0$, en muchos casos podemos despreciar $V(r)$ cerca del origen en comparación con el potencial centrífugo ($\propto r^{-2}$). Si esto ocurre, la (10.103) se reduce a

$$\frac{d^2 u}{dr^2} - \frac{l(l+1)}{r^2} u = 0, \quad (l \neq 0, \quad r \rightarrow 0) \quad (10.109)$$

cuya solución general es

$$u = Ar^{l+1} + Br^{-l} \quad (10.110)$$

Puesto que $l \geq 1$ en la (10.110), la condición de contorno (10.108) exige que $B = 0$. Por lo tanto para todos los estados con $l \geq 1$ tenemos que

$$u = Ar^{l+1}, \quad (l \neq 0, \quad r \rightarrow 0) \quad (10.111)$$

Para los estados S no hay término centrífugo en la (10.103) y entonces el comportamiento de $u(r)$ cerca del origen depende de la forma de $V(r)$.

Si $V \rightarrow 0$ para $r \rightarrow \infty$ la ecuación radial se reduce para r grande a:

$$\frac{d^2 u}{dr^2} + \frac{2\mu E}{\hbar^2} u = 0, \quad (r \rightarrow \infty) \quad (10.112)$$

Cuando $E > 0$ las soluciones de esta ecuación son oscilatorias, por consiguiente todos los valores de E están permitidos y el espectro de autovalores es continuo. En cambio cuando $E < 0$ las soluciones de la (10.112) son exponenciales y para que ψ sea normalizable debemos excluir la solución exponencialmente creciente. Luego en este caso tendremos

$$u(r) \propto e^{-\kappa r} \quad , \quad \kappa = +\sqrt{\frac{-2\mu E}{\hbar^2}} \quad (r \rightarrow \infty) \quad (10.113)$$

Las autofunciones que tienen este comportamiento representan *estados ligados* y las condiciones de contorno permiten en general sólo ciertos valores discretos de la energía.

Para estudiar los estados ligados conviene transformar la ecuación radial poniendo

$$u(r) = \rho^{l+1} e^{-\rho} w(\rho) \quad , \quad \rho = \kappa r \quad (10.114)$$

pues de esta forma eliminamos de la función incógnita $w(r)$ las partes que describen el comportamiento ya conocido de $u(r)$ para $r \rightarrow 0$ y $r \rightarrow \infty$. Sustituyendo (10.114) en (10.103) obtenemos la ecuación

$$\frac{d^2 w}{d\rho^2} + 2\left(\frac{l+1}{\rho} - 1\right) \frac{dw}{d\rho} + \left(\frac{V}{E} - 2\frac{l+1}{\rho}\right) w = 0 \quad (10.115)$$

entre cuyas soluciones deberemos elegir las que satisfacen las condiciones de contorno en el origen y en el infinito.

Estados ligados de átomos con un solo electrón

En este caso la energía potencial es

$$V(r) = -\frac{Ze^2}{r} \quad (10.116)$$

donde Ze es la carga del núcleo y $-e$ la del electrón (para el átomo de hidrógeno $Z = 1$).

En virtud de la discusión precedente $u(r)$ se comporta como r^{l+1} cerca del origen, y, dada la forma (10.116) de $V(r)$, resulta que *también* para los estados S tenemos ese comportamiento.

Conviene introducir el parámetro

$$\rho_0 \equiv \frac{Ze^2 \kappa}{|E|} = \frac{Ze^2}{\hbar} \sqrt{\frac{2\mu}{|E|}} \quad (10.117)$$

de modo que

$$\frac{V}{E} = \frac{\rho_0}{\rho} \quad (10.118)$$

y entonces la (10.115) se escribe como:

$$\rho \frac{d^2 w}{d\rho^2} + 2(l+1-\rho) \frac{dw}{d\rho} + [\rho_0 - 2(l+1)]w = 0 \quad (10.119)$$

Autovalores de la energía

Para resolver la ecuación (10.119) ensayamos un desarrollo en serie de la forma

$$w(\rho) = \sum_k a_k \rho^k \quad (10.120)$$

Introduciendo esta expresión en la (10.119) obtenemos la relación de recurrencia

$$a_{k+1} = \frac{2(k+l+1) - \rho_0}{(k+1)(k+2l+1)} a_k \quad (10.121)$$

De la (10.121) se desprende que $a_{-1} = 0$ y por lo tanto la serie comienza con un término constante $a_0 \neq 0$. Por otra parte, para valores grandes de k

$$a_{k+1} \approx \frac{2}{k} a_k, \quad k \gg 1 \quad (10.122)$$

y entonces para $\rho \rightarrow \infty$ resulta

$$w(\rho) \propto e^{2\rho} \quad (10.123)$$

Este comportamiento es inaceptable, pues nos da

$$u(r) \approx \rho^{l+1} e^{-\rho} w(\rho) \propto e^{\rho} \quad (10.124)$$

Por lo tanto la serie debe terminar de modo que $w(\rho)$ sea un polinomio.

Si N es el grado de dicho polinomio, tendremos que $a_N \neq 0$ y $a_{N+1} = 0$ y la (10.121) nos da la condición

$$\rho_0 = 2(N+l+1) \quad \text{con} \quad N = 0, 1, 2, \dots, \quad l = 0, 1, 2, \dots \quad (10.125)$$

Esta condición implica la *cuantificación de la energía* y es equivalente a la fórmula (5.17) a partir de la cual se obtiene la fórmula de Rydberg para las series espectrales de átomos hidrogenoides.

Para ver esto conviene definir

$$n \equiv N + l + 1 = \frac{\rho_0}{2} \quad (10.126)$$

y reemplazar ρ_0 por su expresión (10.117). Resulta entonces

$$E = E_n = -\frac{\mu Z^2 e^4}{2\hbar^2 n^2} \quad (10.127)$$

que es idéntica a la fórmula (5.17) que obtuvimos aplicando la condición de cuantificación de Bohr de la Teoría Cuántica Antigua.

También vemos que la extensión radial de la región donde está localizado el electrón está determinada por

$$\frac{1}{\kappa} = \frac{n}{Z} a \quad , \quad a = \frac{\hbar^2}{\mu e^2} \quad (10.128)$$

donde a es el radio de Bohr.

Puesto que por definición N es un entero no negativo, es obvio que n es un entero positivo que cumple

$$n \geq 1 \quad (10.129)$$

El estado fundamental corresponde $n = 1$ y $l = 0$ y su energía es de ≈ -13.6 eV para el hidrógeno. Hay infinitos niveles discretos de energía, con punto de acumulación en $E = 0$ para $n \rightarrow \infty$.

De acuerdo con los resultados de nuestro estudio, las autofunciones del Hamiltoniano del electrón de un átomo hidrogenoide se pueden elegir para que sean además autofunciones de L^2 y de L_z y en tal caso están caracterizadas por los tres números cuánticos n, l, m . Puesto que E_n depende sólo del número cuántico principal, todos los niveles de energía, exceptuando el nivel fundamental, son degenerados. Para cada nivel E_n ($n > 1$) hay n valores posibles de l :

$$l = 0, 1, 2, \dots, n-1 \quad (10.130)$$

y para cada uno de ellos hay $2l + 1$ valores posibles del número cuántico magnético m . En consecuencia para cada nivel de energía hay

$$\sum_{l=0}^{n-1} (2l + 1) = n^2 \quad (10.131)$$

autofunciones ψ_{nlm} linealmente independientes. Se dice en este caso que el *grado de degeneración* de cada nivel es n^2 . En realidad, si se toma en cuenta el spin del electrón el grado de degeneración se duplica (también para el estado fundamental).

La aparición de degeneración está *siempre* asociada con las simetrías del sistema. La mayoría de las veces las simetrías son evidentes. Por ejemplo, la degeneración con respecto del número cuántico magnético se origina en la ausencia de una orientación privilegiada y refleja la invariancia del sistema con respecto de las rotaciones alrededor del origen. Claramente esta degeneración está presente para cualquier potencial central. En cambio, la degeneración de los estados con un mismo n y diferente l es peculiar del potencial Coulombiano. Cualquier apartamiento del potencial de la dependencia $1/r$ elimina esta degeneración, porque en ese caso los autovalores de la energía dependen no sólo de n sino también de l . Por este motivo esta degeneración se suele llamar *degeneración accidental*, pero esta denominación no debe llevar a confusión pues no ocurre por casualidad. Su origen se encuentra en una sutil simetría en el espacio de los impulsos, y se relaciona con el hecho que para el potencial de Coulomb la ecuación de Schrödinger independiente del tiempo no sólo es separable en coordenadas esféricas, sino también en coordenadas parabólicas.

Estudio de las autofunciones

Resumiremos sin demostración algunas de las propiedades más importantes de las autofunciones radiales de los estados ligados del potencial de Coulomb.

Las soluciones polinomiales de la (10.119) se relacionan con cierta clase de polinomios ortogonales $L_N^p(z) = L_{q-p}^p(z)$ denominados *polinomios asociados de Laguerre*.

Los *polinomios de Laguerre* de grado q se pueden definir como

$$L_q(z) = e^z \frac{d^q}{dz^q} (e^{-z} z^q) \quad (10.132)$$

y los polinomios asociados de Laguerre de grado $q - p$ como:

$$L_{q-p}^p(z) = (-1)^p \frac{d^p}{dz^p} L_q(z) \quad (10.133)$$

Se encuentra que

$$L_{q-p}^p(0) = \frac{(q!)^2}{(q-p)! p!} \quad (10.134)$$

El número de ceros de $L_{q-p}^p(z)$ es $N = q - p = n - l - 1$, y la integral de normalización es

$$\int_0^\infty e^{-z} z^{2l+2} [L_{n-l-1}^{2l+1}(z)]^2 dz = \frac{2n(n+1)!}{(n-l-1)!} \quad (10.135)$$

Finalmente, reuniendo todos los resultados que obtuvimos podemos escribir las autofunciones normalizadas para los estados ligados de átomos hidrogenoides en la forma:

$$\psi_{nlm}(r, \theta, \varphi) = \left\{ (2\kappa)^3 \frac{(n-l-1)!}{2n(n+l)!} \right\}^{1/2} e^{-\kappa r} (2\kappa r)^l L_{n-l-1}^{2l+1}(2\kappa r) Y_l^m(\theta, \varphi) \quad , \quad \kappa = \frac{Z}{na} \quad (10.136)$$

A veces en lugar de las (10.136) se usan las autofunciones reales:

$$\bar{\psi}_{nlm}(r, \theta, \varphi) \propto e^{-\kappa r} (2\kappa r)^l L_{n-l-1}^{2l+1}(2\kappa r) P_l^m(\theta, \varphi) \begin{cases} \cos m\varphi \\ \sin m\varphi \end{cases} \quad (10.137)$$

Se puede ver que, salvo en el caso de los estados S , las $\bar{\psi}_{nlm}(r, \theta, \varphi)$ tienen n superficies nodales (contando $r = 0$ como una superficie). En efecto $L_{n-l-1}^{2l+1}(2\kappa r)$ se anula para $n - l - 1$ valores de r , $P_l^m(\theta)$ se anula para $l - m$ valores de θ , $\cos m\varphi$ y $\sin m\varphi$ se anulan para m valores de φ , y por último r^l tiene un nodo en $r = 0$ si $l \neq 0$.

El hecho que la función de onda es nula en el origen para todos los estados menos para los estados S tiene consecuencias muy importantes, pues por este motivo sólo los electrones atómicos en los estados S tienen una probabilidad apreciable de encontrarse en el interior del núcleo atómico. Por lo tanto son los únicos que pueden interactuar con el núcleo, dando lugar a procesos como la *captura electrónica* (una forma de decaimiento β^+ alternativa a la emisión de un positrón) y la *conversión interna* (una forma de desexcitación nuclear en la cual en vez de emitir radiación γ , el núcleo transfiere su exceso de energía a un electrón atómico).

El significado cuántico del radio de Bohr se puede apreciar si se observa que la función de onda del estado fundamental es

10. Fuerzas centrales, momento angular y átomo de hidrógeno

$$\psi_{1,0,0}(r, \theta, \varphi) = \left(\frac{Z^3}{\pi a^3} \right)^{1/2} e^{-Zr/a} \quad (10.138)$$

y el valor esperado de r en este estado es

$$\bar{r} = 4 \left(\frac{Z}{a} \right)^3 \int_0^\infty e^{-2Zr/a} r^3 dr = \frac{3a}{2Z} \quad (10.139)$$

La densidad de probabilidad en el estado fundamental es máxima en $r = a/Z$.

Existen por supuesto varias correcciones a la simple fórmula (10.127) de los niveles de energía:

- El efecto de la *masa finita* del núcleo, que se puede tomar en cuenta usando en lugar de μ la masa reducida del sistema electrón-núcleo.
- Los efectos del spin del electrón y las correcciones relativísticas por la velocidad del electrón (*estructura fina*).
- Los efectos de la interacción magnética entre el electrón y el núcleo (*estructura hiperfina*).
- Los efectos debidos a la interacción entre el electrón y el campo electromagnético (*corrimiento de Lamb*).

No nos ocuparemos de los últimos dos en este curso.

11. EL SPIN

El momento angular intrínseco

Hasta aquí nos ocupamos de la descripción cuántica de una partícula, considerada como una masa puntiforme sin estructura interna. En consecuencia supusimos que el estado de la misma se puede especificar por completo mediante la función de onda Ψ , que es una función *escalar* de las variables espaciales x, y, z y del tiempo.

Es ciertamente notable que esta imagen tan simple permite explicar muchos aspectos de la física del átomo y del núcleo. Sin embargo, estos logros no nos deben hacer perder de vista que este modelo sólo explica los aspectos más gruesos del átomo y del núcleo. Muchos detalles se pudieron entender recién a partir del descubrimiento que muchas partículas, incluyendo electrones, protones y neutrones no están completamente descriptas por el modelo de masa puntiforme sin estructura¹. En otras palabras, el estado de esas partículas *no está especificado por completo* por una función escalar de las coordenadas y del tiempo. En efecto, la evidencia empírica muestra que es necesario atribuir a esas partículas un *momento angular intrínseco* o *spin*, y asociado con él, un *momento magnético intrínseco*.

La evidencia espectroscópica

Mencionaremos primero la evidencia espectroscópica acerca del spin y del momento magnético intrínseco, dado que fue a partir de ella que se introdujeron esos atributos del electrón.

La expresión que obtuvimos para los niveles de energía del átomo de hidrógeno (ec. (10.127)) *no explica por completo* el espectro observado, por cuanto muchas de las líneas muestran desdoblamientos en varias componentes, dando lugar a una *estructura fina* de los niveles de energía que no se puede explicar mediante la teoría que desarrollamos en el Capítulo 10.

En 1916, en el marco de la Teoría Cuántica Antigua, Sommerfeld dio una explicación *aparentemente* satisfactoria de la estructura fina (ver el Capítulo 5), basada sobre las correcciones relativísticas, en virtud de las cuales la energía de los niveles (que en el caso no relativístico depende sólo del número cuántico principal) pasa a depender también del número cuántico azimutal. La separación de los niveles así obtenida coincide con la que se observa en el hidrógeno y el helio ionizado y también con la que se observa en las líneas de rayos X de los átomos pesados. Por ese motivo la explicación de Sommerfeld fue aceptada durante varios años.

Sin embargo, poco antes del desarrollo de la Mecánica Cuántica, aparecieron problemas con la teoría de Sommerfeld, en relación con el espectro de los átomos alcalinos. El estado fundamental de un átomo alcalino tiene una estructura electrónica muy simple, y su espectro óptico se puede interpretar suponiendo que el átomo consiste esencialmente de un ion inerte alrededor del cual se mueve el único electrón de valencia. Por lo tanto se comporta básicamente como un átomo de hidrógeno, con la salvedad que el potencial en el que se mueve el electrón de valencia no es Coulombiano, pues la carga del núcleo está apantallada por los electrones del ion. Entonces, aún sin tener en cuenta el efecto relativístico (que en este caso se puede despreciar), la energía de los niveles depende del número cuántico azimutal.

Sin embargo, en los átomos alcalinos se observa que los niveles correspondientes a valores dados de *ambos* números cuánticos (principal y azimutal) aparecen a veces *desdoblados ulterior-*

¹ Por mucho tiempo los protones y neutrones se consideraron elementales. Hoy sabemos que son compuestos.

mente. También se encontró que las separaciones de esos dobletes se pueden representar (formalmente) mediante la fórmula relativística de Sommerfeld. Pero Millikan y Bowen, y también Alfred Landé, que hicieron este descubrimiento en 1924, señalaron que dicha fórmula no se podía aceptar en este caso, por cuanto el número cuántico azimutal es *el mismo* para ambas componentes del doblete y por lo tanto la causa del fenómeno no se podía explicar.

La hipótesis de Uhlenbeck y Goudsmit

En 1926 Samuel A. Uhlenbeck y George E. Goudsmit, dos estudiantes graduados, explicaron del misterio y mostraron que las dificultades se resolvían si se atribuía al electrón una nueva propiedad: la de poseer un *momento angular* $\mathbf{S}$ y un *momento magnético* $\mathbf{M}_s$ intrínsecos, tal como ocurriría si un cuerpo cargado eléctricamente girara alrededor de un eje que pasa por él².

La evidencia espectroscópica muestra que la *magnitud* del momento angular intrínseco del electrón está dada por

$$S^2 = s(s+1)\hbar^2 \quad (11.1)$$

donde s , el *número cuántico de spin*, vale

$$s = 1/2 \quad (11.2)$$

y que la componente S_z del momento angular intrínseco alrededor de un eje z de orientación arbitraria sólo puede tomar los valores

$$S_z = \hbar m_s, \quad m_s = \pm s = \pm 1/2 \quad (11.3)$$

Para explicar los desdoblamientos de la estructura fina y del efecto Zeeman, se encontró que el momento magnético intrínseco asociado con el spin del electrón debe valer

$$\mathbf{M}_s = -\frac{e}{\mu c} \mathbf{S} \quad (11.4)$$

de manera que ($\beta_B = e\hbar/2\mu c$ indica el magnetón de Bohr, ec. (10.61)):

$$\frac{e}{\mu c} = \frac{g_S \beta_B}{\hbar}, \quad g_S = 2 \quad (11.5)$$

y por lo tanto el factor giromagnético de spin es *el doble* del factor giromagnético orbital.

En breve tiempo Uhlenbeck y Goudsmit (y otros como Wolfgang Pauli, Heisenberg, Jordan, Sommerfeld, etc.) mostraron que la introducción del spin resuelve todas las dificultades entonces conocidas y por lo tanto este nuevo atributo del electrón se aceptó, junto a los ya conocidos de carga y masa. En efecto, se encontró que la aparición de dobletes en el espectro de los metales alcalinos se debe exclusivamente a la interacción entre el momento magnético orbital y el momento magnético intrínseco del electrón de valencia. En cuanto a la estructura fina de los niveles de los átomos hidrogenoides, se vio que resulta de una particular combinación de los

² Se puede mencionar que en 1921 Compton había ya especulado sobre la posibilidad que el electrón tuviera un momento angular y un momento magnético intrínsecos, pero no elaboró ulteriormente la idea.

la relatividad, que por una curiosa *coincidencia* dan un resultado idéntico al que obtuvo originalmente Sommerfeld a partir de la Teoría Cuántica Antigua.

El spin quedó así incorporado a la Mecánica Cuántica como un postulado adicional, que es perfectamente compatible con los demás postulados de la teoría pero *no es* una consecuencia lógica de los mismos, sino que se introduce en base a la evidencia experimental. Hay que aclarar aquí que el spin del electrón *no se puede* imaginar como el resultado de la rotación de una distribución de cargas. Es imposible formular un modelo de ese tipo sin incurrir en graves dificultades. Además, la evidencia experimental más reciente tiende a indicar que el electrón es una partícula puntual³. Por consiguiente no es lícito interpretar el spin en términos de modelos clásicos, como esferas de carga en rotación o cosas parecidas.

Posteriormente (1928) Dirac desarrolló una *teoría relativística* del electrón, en la que no es preciso postular el spin sino que éste aparece en forma natural como una consecuencia de la invariancia Lorentz de las ecuaciones que lo describen. La ecuación de Dirac también predice correctamente el valor del factor giromagnético⁴. Estos resultados muestran que el spin está íntimamente relacionado con la relatividad. Con la teoría de Dirac el spin adquirió una firme base teórica.

El experimento de Stern y Gerlach

Si bien en su momento no se tuvo conciencia de ello, en realidad el *momento magnético intrínseco del electrón* ya había sido medido en 1922 por Otto Stern y Walter Gerlach, en un experimento muy interesante que ilustra varios conceptos importantes para la interpretación de la Mecánica Cuántica. En el experimento se intentaba medir el momento magnético de átomos y moléculas haciendo pasar un haz atómico (o molecular) colimado a través de un campo magnético fuertemente inhomogéneo (Fig. 11.1).

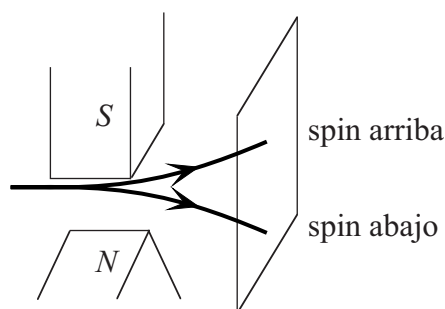


Fig. 11.1. Esquema del experimento de Stern y Gerlach.

De acuerdo con la física clásica, sobre un momento magnético $\mathbf{M}$ sometido a un campo magnético no uniforme $\mathbf{B}$ se ejerce una fuerza dada por

$$\mathbf{F} = \nabla(\mathbf{M} \cdot \mathbf{B}) \quad (11.6)$$

La fuerza (11.6) es la única que actúa sobre un átomo ya que éste es neutro. En el experimento, las partículas se hacían pasar por una región en la cual la variación de la dirección de $\mathbf{B}$ era muy pequeña, pero su magnitud variaba muy fuertemente con la posición. En ese caso la (11.6) nos da (aproximadamente)

$$\mathbf{F} = M_B \nabla B \quad (11.7)$$

donde M_B indica la proyección de $\mathbf{M}$ en la dirección de $\mathbf{B}$.

³ La cota superior del tamaño del electrón que resulta de los experimentos de dispersión a muy altas energías es menor que 10^{-16} cm.

⁴ Al orden más bajo en la constante de la estructura fina. El valor exacto de g_s se calcula por medio de la Electrodinámica Cuántica.

La deflexión se mide estudiando las trazas que el haz de partículas deja en una pantalla, y a partir de ella se puede determinar F y de ahí finalmente M_B .

Los resultados de los experimentos fueron notables. Clásicamente se esperaba obtener en la pantalla una traza continua, correspondiente a los valores de M_B entre $-M$ y $+M$. En cambio, se observaron *varias trazas equidistantes*. Este resultado es una clara demostración de la naturaleza cuántica del momento magnético del átomo. Puesto que el vector $\mathbf{M}$ puede asumir solamente ciertas direcciones discretas en el espacio, dicho fenómeno se denomina *cuantificación espacial*. Stern y Gerlach también obtuvieron resultados cuantitativos (aunque de poca precisión). Encontraron que los valores permitidos de M_B variaban en pasos iguales desde un mínimo $-M$ a un máximo $+M$. Por convención, se suele designar al valor *máximo* de la proyección de $\mathbf{M}$ como el *momento magnético* de la partícula.

De acuerdo con la ley de Ampère el movimiento de los electrones atómicos, además de producir el momento angular orbital $\mathbf{L}$, también produce un momento magnético $\mathbf{M}$ que cumple la relación *clásica*

$$\mathbf{M} = -\frac{e}{2\mu c} \mathbf{L} \quad (11.8)$$

Aplicando el principio de correspondencia, esperamos entonces que en Mecánica Cuántica valga la relación

$$\mathbf{M} = -\frac{e}{2\mu c} \mathbf{L} = -\frac{g_L \beta_B}{\hbar} \mathbf{L} \quad (11.9)$$

Puesto que cualquier componente de $\mathbf{L}$ tiene $2l + 1$ autovalores

$$m = 0, \pm 1, \pm 2, \dots, \pm l \quad (11.10)$$

esperamos que los autovalores de M_B sean

$$M_B = -\frac{g_L \beta_B}{\hbar} m = 0, \pm \frac{g_L \beta_B}{\hbar}, \pm 2 \frac{g_L \beta_B}{\hbar}, \dots, \pm l \frac{g_L \beta_B}{\hbar} \quad (11.11)$$

y como en el experimento de Stern y Gerlach l es entero, esperamos un número *impar* de trazas. Sin embargo, en un experimento clásico realizado con átomos de Ag (que tiene un único electrón de valencia) se observaron *dos* trazas, esto es, un número *par*, y el valor del momento magnético resultó ser

$$M = -\frac{e\hbar}{2\mu c} = -\beta_B \quad (11.12)$$

Las extraordinarias implicaciones de este experimento no fueron comprendidas de inmediato. En cambio se intentó interpretar los resultados en términos de las relaciones (11.9) y (11.10), suponiendo que el electrón de valencia estaba en un estado P (con $l = 1$) y que por alguna razón el estado con $m = 0$ no se presentaba. Se trataba así de explicar porqué se observan solamente dos trazas. Esta hipótesis tan artificiosa tiene la virtud que también explica el valor del momento magnético medido, pero igualmente es implausible, pues normalmente el estado de más baja energía de un átomo no es P sino S .

Antes de proseguir conviene examinar la validez de algunos de los argumentos que usamos para interpretar los resultados del experimento de Stern y Gerlach. La ec. (11.7) es puramente clásica y podría caber la duda que no sea lícito aplicarla al momento magnético cuantificado. La respuesta a esta duda es que en todo experimento hay aspectos que se describen correctamente mediante las leyes clásicas, pues éstas son las leyes que gobiernan lo que experimentan nuestros sentidos, por medio de los cuales (en última instancia y en forma indirecta) nos relacionamos con lo que sucede en los átomos y los núcleos. Puesto que las partículas del haz empleado en el experimento de Stern y Gerlach tienen masa grande, es correcto representarlas mediante paquetes que se dispersan muy lentamente y por lo tanto ese movimiento se puede describir mediante las leyes clásicas. Ese es el motivo por el cual el electrón, que es en realidad el objeto de estudio en el experimento, debe viajar junto al átomo. Si se intentara realizar el mismo experimento con electrones libres, los efectos de interferencia cuántica destruirían el patrón de trazas y no se podría obtener un resultado útil.

Volviendo a los resultados del experimento de Stern y Gerlach con átomos de Ag, las varias interpretaciones que se intentaron al principio fueron harto insatisfactorias y solamente la hipótesis de Uhlenbeck y Goudsmit permitió dar una explicación adecuada.

En efecto, si suponemos que el átomo de Ag está en un estado S (como es lógico), entonces, puesto que $l = 0$, la (12.12) mide en realidad el valor máximo de la componente (en la dirección de $\mathbf{B}$) del momento magnético *intrínseco*. Pero a diferencia del momento magnético orbital, en virtud de (11.3) y (11.4) el momento magnético intrínseco puede tener solamente *dos* proyecciones:

$$M_B = \pm \frac{e\hbar}{2\mu c} \quad (11.13)$$

Por lo tanto *se explica tanto la aparición de dos trazas como su separación*.

Esta interpretación quedó confirmada en 1927 con el experimento de Phipps y Taylor, quienes emplearon la técnica de Stern y Gerlach para medir el momento magnético de átomos de hidrógeno. Este experimento es muy significativo, pues la teoría del átomo de hidrógeno es bien conocida y predice sin lugar a dudas que el estado fundamental es un estado S y entonces el único valor posible de m es $m = 0$. Por lo tanto, en ausencia de otro momento magnético diferente del que proviene del movimiento orbital del electrón, se tendría $M_B = 0$ y el campo magnético no afectaría al haz. Sin embargo Phipps y Taylor encontraron que el haz se separa en *dos* componentes desviadas simétricamente.

Se puede descartar que el momento magnético debido al cual se produce la división del haz provenga del núcleo, puesto que un momento magnético nuclear tendría una magnitud del orden de $e\hbar/2\mu_N c$, donde μ_N es la masa del núcleo⁵ (del protón en el caso del hidrógeno). Pero el momento magnético nuclear es *tres órdenes de magnitud menor* que el que midieron Phipps y Taylor. Por lo tanto es inevitable concluir que el momento magnético responsable de la separación del haz reside en el electrón.

El spin como una variable dinámica

De acuerdo con lo anterior vamos a introducir el spin en la teoría, para lo cual agregamos a las variables dinámicas x, y, z que describen la posición del electrón una cuarta *variable de spin*, que

⁵ La cantidad $\beta_N = e\hbar/2\mu_N c$ se denomina *magnetón nuclear*.

indicaremos con σ . A σ le asignamos sentido físico asociando las dos posibles proyecciones del momento magnético intrínseco $\mathbf{M}_s$ que se miden en el experimento de Stern y Gerlach con dos diferentes valores de σ . De esta forma asociaremos (arbitrariamente):

$$\begin{aligned}\sigma = +1 & \quad \text{con} \quad M_{s,B} = -\frac{e\hbar}{2\mu c} \\ \sigma = -1 & \quad \text{con} \quad M_{s,B} = +\frac{e\hbar}{2\mu c}\end{aligned}\tag{11.14}$$

Frecuentemente se dice “spin arriba” para indicar a $\sigma = +1$ y “spin abajo” para indicar a $\sigma = -1$ (Fig. 11.1). La función de onda depende ahora también de la variable de spin y suponemos que los postulados de la Mecánica Cuántica se aplican a la nueva variable independiente del mismo modo que a las demás variables. Por ejemplo

$$|\Psi(x, y, z, +1, t)|^2 dx dy dz \tag{11.15}$$

es la probabilidad que en el instante t la partícula se encuentre cerca de x, y, z y que, *además*, tenga “spin arriba”, esto es, que la proyección de su momento magnético intrínseco en la dirección de $\mathbf{B}$ sea $M_{s,B} = -\beta_B$.

Esta pequeña generalización de la teoría es suficiente. Todas las fórmulas de los Capítulos precedentes se pueden extender sin dificultad a la Mecánica Cuántica con spin. Donde antes integrábamos sobre las variables espaciales, ahora tenemos que introducir también una suma sobre los dos valores que puede asumir la variable de spin. Por ejemplo, la normalización de la función de onda es ahora

$$\begin{aligned}\sum_{\sigma} \iiint dx dy dz |\Psi(x, y, z, \sigma, t)|^2 = \\ \iiint dx dy dz |\Psi(x, y, z, +1, t)|^2 + \iiint dx dy dz |\Psi(x, y, z, -1, t)|^2 = 1\end{aligned}\tag{11.16}$$

Igual que antes, a cada variable dinámica le corresponde un operador lineal Hermitiano.

Entre todos los estados estacionarios $\psi(x, y, z, \sigma)$ hay una clase especial que es muy fácil de tratar. Son aquellos en que ψ es separable en la forma

$$\psi(x, y, z, \sigma) = \vartheta(x, y, z)\chi(\sigma) \tag{11.17}$$

Nos interesan dos particulares funciones $\chi(\sigma)$ que llamamos $\alpha(\sigma)$ y $\beta(\sigma)$ y que se definen como

$$\begin{aligned}\chi(\sigma) = \alpha(\sigma) & \quad \text{si} \quad \chi(+1) = 1 \quad , \quad \chi(-1) = 0 \\ \chi(\sigma) = \beta(\sigma) & \quad \text{si} \quad \chi(+1) = 0 \quad , \quad \chi(-1) = 1\end{aligned}\tag{11.18}$$

Es decir, α corresponde a “spin arriba” y describe una partícula con un valor definido $\sigma = 1$. Análogamente β corresponde a “spin abajo” y describe una partícula con $\sigma = -1$.

El estado más general es una combinación lineal de estados separables, cosa que permite enormes simplificaciones. Por lo tanto, para una $\psi(x, y, z, \sigma)$ arbitraria tendremos

$$\psi(x, y, z, \sigma) = \psi(x, y, z, +1)\alpha(\sigma) + \psi(x, y, z, -1)\beta(\sigma) \tag{11.19}$$

De todo esto podemos extraer una importante conclusión: las variables espaciales por un lado, y la variable de spin por el otro, se pueden estudiar por separado y reunir en cualquier estadio del proceso.

Los spinores y la teoría del spin en forma matricial

Por dicho anteriormente podemos dejar de lado por el momento las variables de posición y desarrollar una *mecánica cuántica de spin* para una partícula cuyo estado está determinado por una función $\chi(\sigma)$ que depende solamente de la variable de spin. Esta teoría es simple porque σ toma solamente dos valores, pero muy útil pues es el paradigma para la descripción cuántica de cualquier sistema cuyos estados se pueden representar como la superposición de un número *finito* de estados independientes (dos en el caso del spin), y hay muchos problemas de la Mecánica Cuántica en los cuales un formalismo de dos estados es una buena aproximación.

Para desarrollar esta teoría es muy conveniente usar matrices. Para ello representaremos los estados especiales α y β como matrices de una columna:

$$\alpha = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \quad , \quad \beta = \begin{pmatrix} 0 \\ 1 \end{pmatrix} \quad (11.20)$$

y el estado general de spin $\chi(\sigma)$ como

$$\chi(\sigma) = \begin{pmatrix} \chi(+1) \\ \chi(-1) \end{pmatrix} = \chi(+1)\alpha + \chi(-1)\beta = \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} = c_1\alpha + c_2\beta \quad (11.21)$$

En la (11.21) hemos puesto

$$c_1 = \chi(+1) \quad , \quad c_2 = \chi(-1) \quad (11.22)$$

donde c_1, c_2 son números complejos, el cuadrado de cuyo módulo representa la probabilidad de encontrar la partícula con spin arriba y abajo, respectivamente. Por lo tanto la normalización de $\chi(\sigma)$ es

$$|c_1|^2 + |c_2|^2 = (c_1^* \ , \ c_2^*) \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} = 1 \quad (11.23)$$

La matriz (11.21) de una columna y dos filas se denomina *spinor*⁶, y es un ente matemático con peculiares y bien definidas propiedades de transformación bajo rotaciones. Se lo puede considerar como un *vector de componentes complejas en un espacio vectorial lineal de dos dimensiones*. Esta última denominación no debe llevar a confusión ya que $\chi(\sigma)$ *no es* un vector en el espacio ordinario.

Si definimos el *spinor adjunto* de χ como

$$\chi^\dagger = (c_1^* \ , \ c_2^*) \quad (11.24)$$

la (11.23) se puede escribir como

⁶ Los spinores fueron introducidos en 1912 por Élie-Joseph Cartan al estudiar las representaciones del grupo de las rotaciones.

$$\chi^\dagger \chi = 1 \quad (11.25)$$

Dados dos spinores χ , χ' , su *producto escalar* (complejo) se define como

$$\chi^\dagger \chi' = (c_1^*, c_2^*) \begin{pmatrix} c_1' \\ c_2' \end{pmatrix} = c_1^* c_1' + c_2^* c_2' \quad (11.26)$$

Dos spinores son *ortogonales* cuando su producto escalar es nulo. Por ejemplo α y β son dos spinores ortonormales. *Todo par de spinores ortonormales forma una base*, en términos de la cual se puede representar cualquier otro spinor (ver la ec. (11.21)).

Hemos postulado que las magnitudes físicas se representan mediante *operadores lineales Hermitianos*. Consideremos primero los operadores lineales en general. Si F es un operador lineal, su efecto sobre cualquier spinor se puede definir en términos de su efecto sobre los spinores básicos α y β :

$$F\alpha = F_{11}\alpha + F_{12}\beta \quad , \quad F\beta = F_{21}\alpha + F_{22}\beta \quad (11.27)$$

Es inmediato verificar que los coeficientes F_{ij} ($i, j = 1, 2$) caracterizan por completo a F . En efecto

$$\begin{aligned} \xi = F\chi &= c_1 F\alpha + c_2 F\beta = c_1(F_{11}\alpha + F_{12}\beta) + c_2(F_{21}\alpha + F_{22}\beta) \\ &= d_1\alpha + d_2\beta \end{aligned} \quad (11.28)$$

donde d_1 y d_2 son las componentes de ξ .

La ec. (11.28) se puede escribir en forma matricial como

$$\begin{pmatrix} d_1 \\ d_2 \end{pmatrix} = \begin{pmatrix} F_{11} & F_{12} \\ F_{21} & F_{22} \end{pmatrix} \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} \quad (11.29)$$

y por lo tanto $F\xi$ se representa mediante la ecuación matricial (11.29). De aquí en más usaremos el mismo símbolo para designar el estado y el spinor que lo representa y también el mismo símbolo para designar la magnitud física y el operador y la matriz que la representan.

El operador *adjunto* o *conjugado Hermitiano* $F^\dagger$ del operador F se define como el operador cuya matriz es la matriz *transpuesta y conjugada* de F :

$$F^\dagger = \begin{pmatrix} F_{11}^* & F_{21}^* \\ F_{12}^* & F_{22}^* \end{pmatrix} \quad (11.30)$$

El *valor esperado* de F se define como

$$\bar{F} = \chi^\dagger F\chi = (c_1^*, c_2^*) \begin{pmatrix} F_{11} & F_{12} \\ F_{21} & F_{22} \end{pmatrix} \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} \quad (11.31)$$

El valor esperado de una magnitud física F en cualquier estado debe ser real. Es fácil verificar que la condición necesaria y suficiente para que esto ocurra es que la matriz F sea Hermitiana.

Todos los resultados que ya conocemos para los operadores Hermitianos (Capítulos 7 y 8) se pueden trasladar las matrices Hermitianas, simplemente reemplazando la palabra “operador” por

“matriz”. En efecto, dada una base en el espacio de los spinores, existe una correspondencia biunívoca entre los operadores lineales y las matrices de 2×2 . En realidad el presente formalismo es más simple que el de los Capítulos 7 y 8 porque el número máximo de spinores linealmente independientes es 2 y por lo tanto no surgen problemas con la completitud y por supuesto tampoco con el espectro continuo.

Spin y rotaciones

La primera magnitud física que vamos a considerar es el momento magnético intrínseco⁷ $\mathbf{M}_s$, y en particular su componente en la dirección del campo magnético $\mathbf{B}$. Tomaremos el eje z en la dirección de $\mathbf{B}$ y en consecuencia la componente z de $\mathbf{M}_s$ está representada por la matriz Hermitiana

$$M_{s,z} = -\beta_B \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad (11.32)$$

puesto que sus autovalores son $\mp\beta_B$ y sus autovectores son α y β . Nos preguntamos ahora cómo se representan las demás componentes de $\mathbf{M}_s$.

Para determinar las matrices $M_{s,x}$ y $M_{s,y}$ vamos a estipular que las tres componentes del valor esperado $\overline{\mathbf{M}_s}$ del momento magnético intrínseco se deben transformar por efecto de una rotación del mismo modo que las tres componentes de un vector.

Para imponer esta condición debemos determinar primero cómo se transforma un spinor bajo una rotación. Consideremos una rotación de un ángulo θ alrededor del eje $\hat{\mathbf{n}}$ y sean

$$\chi = \begin{pmatrix} c_1 \\ c_2 \end{pmatrix}, \quad \chi' = \begin{pmatrix} c'_1 \\ c'_2 \end{pmatrix} \quad (11.33)$$

el spinor antes de la rotación y el spinor rotado, respectivamente.

Claramente la relación entre χ y χ' debe ser lineal, para preservar la superposición lineal de dos estados, y por lo tanto estará representada por una matriz U :

$$\chi' = U\chi \quad (11.34)$$

cuyos elementos dependen solamente de los parámetros de la rotación, esto es θ y $\hat{\mathbf{n}}$.

Puesto que χ' debe estar normalizado si χ lo está se debe cumplir que

$$\chi'^{\dagger}\chi' = \chi^{\dagger}U^{\dagger}U\chi = \chi^{\dagger}\chi \quad (11.35)$$

y como χ es arbitrario la (11.35) implica que

$$U^{\dagger}U = 1 \quad (11.36)$$

donde con 1 (en *cursiva*) indicamos la matriz identidad. Una matriz que cumple la (11.36) se dice *unitaria*.

⁷ Procedemos así porque $\mathbf{M}_s$ es una magnitud que se puede medir directamente, cosa que no ocurre con el momento angular intrínseco. Es mejor dejar para más adelante la discusión de cómo se introduce en la teoría el momento angular intrínseco.

De la (11.36) resulta que el determinante de una matriz unitaria es un número complejo de módulo 1:

$$|\det U|^2 = 1 \quad (11.36)$$

y que la inversa de U es

$$U^{-1} = U^\dagger \quad (11.37)$$

Se acostumbra decir que la matriz unitaria U , que corresponde a la rotación $(\theta, \hat{n})$ en la cual $\chi \rightarrow \chi'$, *representa* dicha rotación.

Si U_1 representa una rotación R_1 , y U_2 representa una rotación R_2 , entonces la matriz $U_2 U_1$ representa la rotación que resulta de realizar primero R_1 y luego R_2 . Pero el mismo efecto se puede obtener directamente efectuando una única rotación R_3 , representada por U_3 . Por lo tanto, las matrices unitarias que representan rotaciones deben tener la propiedad grupal

$$U_3 = e^{i\varphi} U_2 U_1 \quad (11.38)$$

donde hemos introducido una fase arbitraria, dado que todos los spinors $e^{i\varphi} \chi$ representan el mismo estado.

Consideremos ahora una rotación infinitesimal $(d\theta, \hat{n})$. Tal rotación debe diferir muy poco de la matriz identidad. Por lo tanto hasta el orden $d\theta$ tendremos:

$$U_R = 1 - \frac{i}{2} d\theta \hat{n} \cdot \boldsymbol{\sigma} \quad (11.39)$$

donde con $\boldsymbol{\sigma}$ indicamos tres matrices Hermitianas constantes $\sigma_x, \sigma_y, \sigma_z$, aún no determinadas (deben ser Hermitianas para que U_R sea unitaria al primer orden en $d\theta$; en cuanto al factor $1/2$, lo hemos puesto por conveniencia).

Si conocemos estas tres matrices podemos construir U para cualquier rotación finita por aplicación sucesiva de muchas rotaciones infinitesimales, es decir por integración de la (11.39). Efectivamente, gracias al Teorema de Euler (que establece que para cualquier rotación dada en tres dimensiones se puede siempre encontrar un eje fijo de modo que la rotación se puede reducir a una rotación alrededor de dicho eje) es suficiente considerar *rotaciones finitas alrededor de un eje fijo* ($\hat{n}$ constante). Entonces la (11.39) se integra fácilmente y da

$$U_R = \lim_{N \rightarrow \infty} \left(1 - \frac{i}{2} \frac{\theta}{N} \hat{n} \cdot \boldsymbol{\sigma} \right)^N = e^{-\frac{i}{2} \theta \hat{n} \cdot \boldsymbol{\sigma}} \quad (11.40)$$

Para que la matriz (u operador) (11.40) represente una rotación aún falta satisfacer la condición (11.38), que impone fuertes restricciones sobre la forma de las matrices $\sigma_x, \sigma_y, \sigma_z$ aún desconocidas. Pero en vez de proceder directamente elaborando dicha condición, es mejor seguir una vía alternativa que nos llevará al mismo resultado, y se basa en estudiar las propiedades de transformación bajo rotaciones de un *operador vectorial*, esto es, un operador $\mathcal{A}$ tal que $\overline{\mathcal{A}} = \hat{x} \overline{\mathcal{A}_x} + \hat{y} \overline{\mathcal{A}_y} + \hat{z} \overline{\mathcal{A}_z}$ se transforma como un vector ordinario.

Es importante observar que $\boldsymbol{\sigma}$ es un operador vectorial. En efecto, para una rotación infinitesimal tenemos que

$$\chi' \left[1 - \frac{i}{2} d\theta \hat{\mathbf{n}} \cdot \boldsymbol{\sigma} \right] \chi \quad (11.41)$$

Si tomamos el producto escalar de la (11.41) por χ resulta

$$\chi^\dagger \chi' \left[1 - \frac{i}{2} d\theta \hat{\mathbf{n}} \cdot \boldsymbol{\sigma} \right] \chi = \chi^\dagger \chi - \frac{i}{2} d\theta \hat{\mathbf{n}} \cdot \bar{\boldsymbol{\sigma}} \quad (11.42)$$

donde el valor esperado de $\boldsymbol{\sigma}$ se toma respecto del estado χ . Consideremos ahora el efecto de realizar una *rotación finita arbitraria* de χ y χ' . Los productos escalares $\chi^\dagger \chi'$ y $\chi^\dagger \chi$ son invariantes bajo la transformación unitaria que representa esta rotación. Por lo tanto el producto escalar $\hat{\mathbf{n}} \cdot \bar{\boldsymbol{\sigma}}$ es también *invariante ante rotaciones*. Puesto que $\hat{\mathbf{n}}$ es un vector, eso significa que también $\bar{\boldsymbol{\sigma}}$ es un vector y por lo tanto $\boldsymbol{\sigma}$ es un operador vectorial.

Un operador vectorial $\mathbf{A}$ debe satisfacer ciertas relaciones matemáticas que vamos a obtener ahora mismo. Una rotación R se representa en el espacio ordinario por medio de una matriz ortogonal

$$R = \begin{pmatrix} R_{xx} & R_{xy} & R_{xz} \\ R_{yx} & R_{yy} & R_{yz} \\ R_{zx} & R_{zy} & R_{zz} \end{pmatrix} \quad (11.43)$$

Bajo la misma rotación, los spinores se transforman mediante la correspondiente matriz unitaria U_R , esto es

$$\chi' = U_R \chi \quad (11.44)$$

Por lo tanto se debe cumplir que

$$\chi'^\dagger \mathbf{A} \chi' = \chi^\dagger U_R^\dagger \mathbf{A} U_R \chi = R \chi^\dagger \mathbf{A} \chi \quad (11.45)$$

Esto es:

$$\begin{pmatrix} \chi^\dagger U_R^\dagger A_x U_R \chi \\ \chi^\dagger U_R^\dagger A_y U_R \chi \\ \chi^\dagger U_R^\dagger A_z U_R \chi \end{pmatrix} = \begin{pmatrix} R_{xx} & R_{xy} & R_{xz} \\ R_{yx} & R_{yy} & R_{yz} \\ R_{zx} & R_{zy} & R_{zz} \end{pmatrix} \begin{pmatrix} \chi^\dagger A_x \chi \\ \chi^\dagger A_y \chi \\ \chi^\dagger A_z \chi \end{pmatrix} \quad (11.46)$$

de donde obtenemos

$$\chi^\dagger (U_R^\dagger A_x U_R) \chi = \chi^\dagger (R_{xx} A_x + R_{xy} A_y + R_{xz} A_z) \chi \quad (11.47)$$

y ecuaciones semejantes para $\chi^\dagger (U_R^\dagger A_y U_R) \chi$ y $\chi^\dagger (U_R^\dagger A_z U_R) \chi$. Puesto que χ es arbitrario, resulta finalmente

$$\begin{aligned} U_R^\dagger A_x U_R &= R_{xx} A_x + R_{xy} A_y + R_{xz} A_z \\ U_R^\dagger A_y U_R &= R_{yx} A_x + R_{yy} A_y + R_{yz} A_z \\ U_R^\dagger A_z U_R &= R_{zx} A_x + R_{zy} A_y + R_{zz} A_z \end{aligned} \quad (11.48)$$

Las (11.48) valen para cualquier rotación. Consideremos ahora una rotación infinitesimal de un ángulo ε alrededor del eje z , para la cual

$$R = \begin{pmatrix} 1 & -\varepsilon & 0 \\ \varepsilon & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}, \quad U_R = 1 - \frac{i}{2} \varepsilon \sigma_z \quad (11.49)$$

Si sustituimos (11.49) en (11.48) vemos que las componentes de un operador vectorial deben satisfacer las siguientes condiciones *necesarias*:

$$\begin{aligned} \sigma_z A_x - A_x \sigma_z &= 2iA_y \\ \sigma_z A_y - A_y \sigma_z &= -2iA_x \\ \sigma_z A_z - A_z \sigma_z &= 0 \end{aligned} \quad (11.50)$$

Del mismo modo, considerando rotaciones infinitesimales alrededor del eje x y del eje y obtenemos las condiciones:

$$\begin{aligned} \sigma_x A_y - A_y \sigma_x &= 2iA_z \\ \sigma_x A_z - A_z \sigma_x &= -2iA_y \\ \sigma_x A_x - A_x \sigma_x &= 0 \end{aligned} \quad (11.51)$$

y

$$\begin{aligned} \sigma_y A_z - A_z \sigma_y &= 2iA_x \\ \sigma_y A_x - A_x \sigma_y &= -2iA_z \\ \sigma_y A_y - A_y \sigma_y &= 0 \end{aligned} \quad (11.52)$$

En particular, como $\boldsymbol{\sigma}$ mismo es un operador vectorial, poniendo $\mathbf{A} \rightarrow \boldsymbol{\sigma}$ en las (11.50)-(11.52) obtenemos las siguientes relaciones de conmutación entre las matrices σ_x , σ_y , σ_z :

$$\begin{aligned} \sigma_x \sigma_y - \sigma_y \sigma_x &= 2i\sigma_z \\ \sigma_y \sigma_z - \sigma_z \sigma_y &= 2i\sigma_x \\ \sigma_z \sigma_x - \sigma_x \sigma_z &= 2i\sigma_y \end{aligned} \quad (11.53)$$

Estas relaciones se suelen escribir en forma compacta como

$$\boldsymbol{\sigma} \times \boldsymbol{\sigma} = 2i\boldsymbol{\sigma} \quad (11.54)$$

Las relaciones fundamentales (11.53) son las condiciones que deben cumplir las matrices σ_x , σ_y , σ_z , para que matrices del tipo

$$U_R = e^{-\frac{i}{2} \theta \hat{n} \cdot \boldsymbol{\sigma}} \quad (11.55)$$

representen rotaciones. Se puede mostrar que estas condiciones son necesarias y suficientes.

Es importante destacar que, en realidad, los argumentos que presentamos en esta discusión de las rotaciones (desde la ec. (11.34) en adelante) *no dependen* de la *dimensionalidad* de las matrices χ , χ' , σ y U_R . Es decir, valen también si χ y χ' son matrices de n filas y una columna y entonces σ y U_R son matrices de $n \times n$. Por lo tanto nuestros resultados son generales y valen para todos los operadores de momento angular.

Las matrices de Pauli

Nos restringimos ahora al caso $n = 2$ y nos ocuparemos de determinar las matrices Hermitianas σ_x , σ_y , σ_z de 2×2 que satisfacen las relaciones de conmutación (11.53).

Claramente, σ_z está determinada en gran medida por nuestra elección de los spinores básicos α y β . En efecto, α corresponde a “spin arriba” y describe una partícula con un valor definido $\sigma = 1$ (es decir $M_{s,B} = -\beta_B$) y β corresponde a “spin abajo” y describe una partícula con $\sigma = -1$ (o sea $M_{s,B} = +\beta_B$). Por lo tanto en un estado cualquiera

$$\chi = \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} \quad (11.56)$$

$|c_1|^2$ y $|c_2|^2$ son las probabilidades de encontrar “spin arriba” y “spin abajo”, respectivamente. Es obvio que una rotación alrededor del eje z (que por convención hemos tomado en la dirección del campo magnético) no tiene ningún efecto sobre estas probabilidades. Por lo tanto si realizamos una rotación infinitesimal de un ángulo ε alrededor del eje z , $|c_1|^2$ y $|c_2|^2$ no deben cambiar. Eso significa que

$$U_R = 1 - \frac{i}{2} \varepsilon \sigma_z \quad (11.57)$$

es una matriz diagonal, luego también σ_z es diagonal. Además, si tomamos la traza de la primera de las (11.53) obtenemos

$$\text{Tr}(\sigma_z) = 0 \quad (11.58)$$

puesto que $\text{Tr}(AB) = \text{Tr}(BA)$. Por consiguiente σ_z tiene la forma

$$\sigma_z = \begin{pmatrix} \lambda & 0 \\ 0 & -\lambda \end{pmatrix}, \quad (\lambda \text{ real}) \quad (11.59)$$

Para determinar σ_x , σ_y y el valor de λ conviene definir las matrices auxiliares no Hermitianas

$$\sigma_+ = \sigma_x + i\sigma_y, \quad \sigma_- = \sigma_x - i\sigma_y = \sigma_+^\dagger \quad (11.60)$$

Se verifica fácilmente que estas matrices cumplen las siguientes reglas de conmutación

$$\begin{aligned} [\sigma_z, \sigma_+] &= 2\sigma_+ \\ [\sigma_z, \sigma_-] &= -2\sigma_- \\ [\sigma_+, \sigma_-] &= 4\sigma_z \end{aligned} \quad (11.61)$$

Ahora ponemos

$$\sigma_+ = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \quad (11.62)$$

con a, b, c, d a determinar, y sustituimos (11.59) y (11.62) en la primera de las (11.61). Resulta:

$$\begin{pmatrix} 0 & 2\lambda b \\ -2\lambda c & 0 \end{pmatrix} = 2 \begin{pmatrix} a & b \\ c & d \end{pmatrix} \quad (11.63)$$

Por lo tanto debe ser $a = d = 0$, $(\lambda - 1)b = 0$ y $(\lambda + 1)c = 0$.

De estas ecuaciones inferimos que λ puede valer $+1$ ó -1 . La elección es arbitraria pues el otro valor está presente en la (11.59). Elegimos entonces $\lambda = +1$ y por lo tanto resulta $c = 0$. Obtenemos así

$$\sigma_+ = \begin{pmatrix} 0 & b \\ 0 & 0 \end{pmatrix}, \quad \sigma_- = \sigma_+^\dagger = \begin{pmatrix} 0 & 0 \\ b^* & 0 \end{pmatrix} \quad (11.64)$$

Por último, sustituyendo estas matrices en la última de las (11.61) encontramos que $|b|^2 = 4$. La fase de b no se puede inferir de las relaciones de conmutación fundamentales (11.53) y por consiguiente es arbitraria. Elegimos entonces $b = 2$ y de esta forma llegamos al resultado final:

$$\sigma_x = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad \sigma_y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, \quad \sigma_z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad (11.65)$$

Las (11.65) son las *matrices de spin de Pauli*. Están completamente determinadas (a menos del orden de filas y columnas, que corresponde a la elección entre $\lambda = +1$ y $\lambda = -1$, y a menos de la fase de b) por las relaciones de conmutación (11.53) y por nuestra elección de los autospinores α y β de σ_z como spinores básicos.

Es fácil verificar (lo dejemos como ejercicio) las siguientes propiedades de las matrices de Pauli:

$$\begin{aligned} \sigma_x^2 = \sigma_y^2 = \sigma_z^2 = 1 \\ \sigma_x \sigma_y = i\sigma_z, \quad \sigma_y \sigma_z = i\sigma_x, \quad \sigma_z \sigma_x = i\sigma_y \end{aligned} \quad (11.66)$$

Vemos entonces que las matrices de Pauli además de Hermitianas son también *unitarias* y que tienen *traza nula* (por las (11.65)). Asimismo, cualquier par de matrices de Pauli diferentes *anti-conmutan*, por ejemplo

$$\sigma_x \sigma_y + \sigma_y \sigma_x = 0 \quad (11.67)$$

Además las cuatro matrices $I, \sigma_x, \sigma_y, \sigma_z$ son linealmente independientes, y por lo tanto cualquier matriz A de 2×2 se puede expresar como

$$A = \lambda_0 I + \lambda_x \sigma_x + \lambda_y \sigma_y + \lambda_z \sigma_z = \lambda_0 I + \boldsymbol{\lambda} \cdot \boldsymbol{\sigma} \quad (11.68)$$

Si A es Hermitiana, todos los coeficientes deben ser reales. Si U es unitaria, se puede mostrar (con un poco de paciencia) a partir de la (11.68) que se la puede expresar siempre en la forma

$$U = e^{i\gamma} (1 \cos \omega + i \hat{\mathbf{n}} \cdot \boldsymbol{\sigma} \sin \omega) \quad (11.69)$$

donde γ y ω son ángulos reales y $\hat{n}$ es un versor. Se puede también mostrar que la (11.69) se puede escribir en la forma

$$U = e^{i\gamma + i\omega \hat{n} \cdot \sigma} \quad (11.70)$$

Cualquier matriz unitaria de 2×2 se puede escribir de esta forma. Comparando estas expresiones con la (11.40) vemos que cualquier matriz unitaria de la forma (11.69) u (11.70) con $\gamma = 0$ representa una rotación de un ángulo $\theta = -2\omega$ alrededor del eje $\hat{n}$. Para $\gamma = 0$ se tiene $\det U = 1$ y se dice que U es *unimodular*. El conjunto de todas las matrices unitarias y unimodulares de 2×2 constituye el grupo $SU(2)$.

Usando la (11.69), la matriz (11.40) que representa una rotación se puede escribir en la forma

$$U_R = e^{-\frac{i}{2}\theta \hat{n} \cdot \sigma} = 1 \cos\left(\frac{\theta}{2}\right) - i\hat{n} \cdot \sigma \sin\left(\frac{\theta}{2}\right) \quad (11.71)$$

Una consecuencia inmediata de la (11.71) es que para una rotación de 2π se obtiene $U = -1$, es decir, una rotación de 360° alrededor de cualquier eje cambia el signo de todas las componentes de un spinor, a diferencia de lo que ocurre con los vectores (y tensores en general) que vuelven a sus valores originales al dar una vuelta completa. Debe notarse que esta propiedad no afecta los valores esperados ni los elementos de matriz pues éstos dependen *cuadráticamente* de los spinores. Para comparar el diferente comportamiento de spinores y vectores bajo rotaciones consideremos por ejemplo lo que ocurre para una rotación de un ángulo θ alrededor del eje x : para un spinor se tiene de (11.71), (11.45) y (11.34):

$$\begin{aligned} c'_1 &= c_1 \cos\left(\frac{\theta}{2}\right) - ic_2 \sin\left(\frac{\theta}{2}\right) \\ c'_2 &= -ic_1 \sin\left(\frac{\theta}{2}\right) + c_2 \cos\left(\frac{\theta}{2}\right) \end{aligned} \quad (11.72)$$

mientras que para un vector se tiene

$$\begin{aligned} A'_x &= A_x \\ A'_y &= A_y \cos \theta - A_z \sin \theta \\ A'_z &= A_y \sin \theta + A_z \cos \theta \end{aligned} \quad (11.73)$$

Hemos visto previamente que σ es un operador vectorial. Se puede mostrar que en el espacio de los spinores σ es el *único operador vectorial*, a menos de un factor constante. En efecto, tomando la traza de las (11.50)-(11.52) se verifica que la traza de todas las componentes de un operador vectorial cualquiera A es nula. A partir de ahí se encuentra que

$$A = k\sigma, \quad k = \text{cte.} \quad (11.74)$$

Operadores de momento magnético y momento angular intrínseco

Puesto que el operador vectorial σ es esencialmente único, a partir de la (11.32) concluimos que el momento magnético intrínseco está dado por

$$\mathbf{M}_s = -\beta_B \boldsymbol{\sigma} \quad (11.75)$$

Hasta ahora no hemos hablado todavía del *momento angular intrínseco*, que también fue postulado por Uhlenbeck y Goudsmit. Esto se debe a que, a diferencia del momento magnético, no es una magnitud que se pueda medir en forma directa con facilidad, y no quisimos entrar en el detalle de los argumentos originales de Uhlenbeck y Goudsmit, pues se fundan en aspectos de la teoría de los espectros atómicos que no hemos estudiado. Pero con base a lo que ya vimos, podemos presentar dos argumentos que muestran que el electrón debe poseer un momento angular intrínseco:

- Tiene un momento magnético, y si este momento magnético proviene de alguna clase de corriente interna debida a la circulación de materia cargada, entonces cabe esperar que junto con el momento magnético exista también un momento angular.
- Si el electrón que se mueve alrededor del núcleo tiene un momento magnético, no se puede conservar el momento angular, a menos que posea un momento angular intrínseco además del momento angular orbital.

Veamos más en detalle el segundo de estos argumentos. Un momento magnético que se mueve en un campo eléctrico experimenta una fuerza que deriva de una energía potencial dada por

$$V = \mathbf{M}_s \cdot \frac{\mathbf{v}}{c} \times \mathbf{E} \quad (11.76)$$

Para un campo central, $\mathbf{E} = f(r)\mathbf{r}$, y entonces V es proporcional a $\mathbf{M}_s \cdot \mathbf{v} \times \mathbf{r}$ o, lo que es equivalente, a

$$V \propto \mathbf{M}_s \cdot \mathbf{r} \times \mathbf{p} = \mathbf{M}_s \cdot \mathbf{L} = -\beta_B \boldsymbol{\sigma} \cdot \mathbf{L} \quad (11.77)$$

en virtud de la (11.75). El factor de proporcionalidad depende solamente de la coordenada radial r . De resultados de esto, además del potencial central, aparece en el Hamiltoniano un término de interacción proporcional a $\mathbf{M}_s \cdot \mathbf{L}$, que implica que la energía del electrón depende de la orientación relativa del momento magnético y del momento angular orbital $\mathbf{L}$. Es evidente entonces que $\mathbf{L}$, cuyas componentes no conmutan entre sí, *no es una constante del movimiento*. Por lo tanto el momento angular no se conservará, a menos que el electrón participe en el balance del momento angular en virtud de un *momento angular intrínseco* asociado con $\mathbf{M}_s$.

Un término del Hamiltoniano proporcional a $\boldsymbol{\sigma} \cdot \mathbf{L}$ se suele llamar *interacción spin-órbita*, y como dijimos al comienzo del Capítulo, esta clase de interacción aparece en los átomos como una corrección magnética y relativística al potencial electrostático central. En los núcleos aparece también una interacción spin-órbita, pero su origen es distinto que en el átomo, pues proviene de la interacción fuerte entre los nucleones y tiene por consiguiente efectos muy importantes.

Veamos ahora cómo debe ser el momento angular intrínseco para que haya conservación del momento angular en presencia de la interacción spin-órbita. Puesto que sabemos que $\mathbf{L}$ no es constante del movimiento, procuraremos definir el momento angular intrínseco $\mathbf{S}$ de modo tal que el *momento angular total*

$$\mathbf{J} = \mathbf{L} + \mathbf{S} \quad (11.78)$$

sea una constante del movimiento.

Dado que $\mathbf{S}$ debe ser un operador vectorial, tiene que ser proporcional a $\boldsymbol{\sigma}$. Ponemos entonces

$$\mathbf{S} = a\boldsymbol{\sigma} \quad (11.79)$$

Para entender como opera $\mathbf{J}$, recordamos que la función de onda del electrón es un spinor de la forma

$$\psi = \begin{pmatrix} \psi_1(x, y, z) \\ \psi_2(x, y, z) \end{pmatrix} \quad (11.80)$$

y en consecuencia la (11.78) significa que

$$\mathbf{J} = \mathbf{L} + a\boldsymbol{\sigma} \quad (11.81)$$

Por lo tanto el operador $\mathbf{L}$ actúa sólo sobre las coordenadas x, y, z mientras que $\boldsymbol{\sigma}$ acopla las dos componentes del spinor. Claramente $\mathbf{L}$ y $\boldsymbol{\sigma}$ conmutan.

Determinamos ahora a requiriendo que $\mathbf{J}$ conmute con $\boldsymbol{\sigma} \cdot \mathbf{L}$, lo cual asegura que sea una constante del movimiento. Por ejemplo, para la componente z tenemos

$$\begin{aligned} [\boldsymbol{\sigma} \cdot \mathbf{L}, J_z] &= \sigma_x [L_x, L_z] + \sigma_y [L_y, L_z] + a[\sigma_x, \sigma_z] L_x + a[\sigma_y, \sigma_z] L_y \\ &= -i\hbar \sigma_x L_y + i\hbar \sigma_y L_x - 2ia\sigma_y L_x + 2ia\sigma_x L_y \\ &= i(2a - \hbar)(\sigma_x L_y - \sigma_y L_x) = 0 \end{aligned} \quad (11.82)$$

Por consiguiente debemos tener $a = \hbar/2$, lo que nos da

$$\mathbf{S} = \frac{\hbar}{2} \boldsymbol{\sigma} \quad (11.83)$$

Por lo tanto, como ya anticipamos al comienzo de este Capítulo, cualquier componente del momento angular intrínseco tiene los dos autovalores

$$m_s \hbar, \quad m_s = \pm s = \pm 1/2 \quad (11.84)$$

El máximo valor de una componente del spin $\mathbf{S}$ es $s = 1/2$ (en unidades de $\hbar$), y por eso se dice que *el electrón tiene spin 1/2*.

Usando la primera de las (11.66) obtenemos también que

$$\mathbf{S}^2 = \frac{\hbar^2}{4} \boldsymbol{\sigma} \cdot \boldsymbol{\sigma} = s(s+1)\hbar^2 \mathbf{1} = \frac{3\hbar^2}{4} \mathbf{1} \quad (11.85)$$

y por lo tanto cualquier spinor es autospinor de $\mathbf{S}^2$ con autovalor $3\hbar^2/4$

A partir de las relaciones de conmutación de las matrices de Pauli es inmediato verificar que las componentes de $\mathbf{S}$ cumplen las relaciones de conmutación

$$\begin{aligned}
S_x S_y - S_y S_x &= i\hbar S_z \\
S_y S_z - S_z S_y &= i\hbar S_x \\
S_z S_x - S_x S_z &= i\hbar S_y
\end{aligned}
\tag{11.86}$$

y que las relaciones de conmutación entre las componentes del momento angular total $\mathbf{J}$ son

$$\begin{aligned}
J_x J_y - J_y J_x &= i\hbar J_z \\
J_y J_z - J_z J_y &= i\hbar J_x \\
J_z J_x - J_x J_z &= i\hbar J_y
\end{aligned}
\tag{11.87}$$

Tanto las (11.86) como las (11.87) son idénticas a las relaciones de conmutación (10.4) del momento angular orbital $\mathbf{L}$. En general *el operador que representa el momento angular de cualquier sistema cuántico satisface reglas de conmutación de la forma* (11.87).

El operador unitario que transforma la ψ de un electrón con spin bajo una rotación infinitesimal está dado por

$$U_R = 1 - \frac{i}{2} \varepsilon \hat{\mathbf{n}} \cdot \boldsymbol{\sigma} - i\varepsilon \hat{\mathbf{n}} \cdot \mathbf{L} = 1 - i\varepsilon \hat{\mathbf{n}} \cdot \mathbf{S} - i\varepsilon \hat{\mathbf{n}} \cdot \mathbf{L} = 1 - i\varepsilon \hat{\mathbf{n}} \cdot \mathbf{J} \tag{11.88}$$

donde hemos usado la (10.24) y la (11.39). La (11.88) muestra que cuando se toma en cuenta el spin, *el generador de rotaciones infinitesimales es el momento angular total $\mathbf{J}$* , y la conservación del momento angular (total) es simplemente una consecuencia de la invariancia del Hamiltoniano bajo rotaciones.

Para una rotación finita tendremos

$$U_R = e^{-i\theta \hat{\mathbf{n}} \cdot \mathbf{J}} \tag{11.89}$$

Para cualquier sistema cuántico el operador momento angular $\mathbf{J}$ es *por definición* el generador de rotaciones infinitesimales y toda rotación finita se puede expresar de la forma (11.89).

12. ÁTOMOS CON VARIOS ELECTRONES, EL PRINCIPIO DE EXCLUSIÓN Y LA TABLA PERIÓDICA

En el Capítulo 10 estudiamos los átomos con un solo electrón, y vimos que con la introducción del spin (Capítulo 11) la Mecánica Cuántica proporciona una teoría perfectamente satisfactoria, si se toman en cuenta todas las correcciones (que nosotros no hemos desarrollado en sus detalles por razones de brevedad, pero que el lector puede encontrar tratadas en extenso en la bibliografía). Ahora queremos ver, al menos en forma cualitativa, cómo se puede aplicar la Mecánica Cuántica a los átomos con varios electrones. En particular, nos gustaría entender en términos generales los esquemas de niveles de energía de átomos más complicados y explicar, por ejemplo, la tabla periódica de los elementos, la existencia de elementos químicamente inertes como los gases nobles, etc.. Veremos que la Mecánica Cuántica, tal como la desarrollamos hasta aquí, no es todavía suficiente para esos fines, y que es necesario introducir un postulado adicional: el principio de exclusión de Pauli.

Descripción de un átomo con varios electrones

Consideremos los estados estacionarios de un átomo con N electrones. La correspondiente función de onda es solución de la ecuación de Schrödinger independiente del tiempo, que en primera aproximación, despreciando el efecto spin-órbita y otras interacciones magnéticas se puede escribir como:

$$\mathcal{H}\psi_N = \left[\sum_{i=1}^N \left(\frac{\mathbf{p}_i^2}{2\mu} - \frac{Ze^2}{r_i} \right) + \sum_{\substack{i,j=1 \\ i>j}}^N \frac{e^2}{|\mathbf{r}_i - \mathbf{r}_j|} \right] \psi_N = E\psi_N \quad (12.1)$$

Aquí $\mathbf{r}_i$ y $\mathbf{p}_i$ son la posición y la cantidad de movimiento del i -ésimo electrón y Ze es la carga del núcleo. Cabe observar que ψ_N depende de las $3N$ variables de posición y N variables de spin del conjunto de los electrones.

Este problema es el análogo cuántico del problema de muchos cuerpos de la Mecánica clásica, e igual que éste, no tiene solución en forma cerrada, ni siquiera para $N = 2$. Si bien en principio se puede resolver numéricamente la (12.1), tal proceder tendría poca utilidad pues sería difícil, o imposible, interpretar el significado de los resultados que se obtuvieran, como así calcular otras propiedades del átomo, además de los niveles de energía. En casos complicados como éste, es preferible formular un modelo aproximado, que retenga la esencia del problema exacto, pero que al mismo tiempo permita entender la naturaleza física de la solución y estudiar la influencia de los diversos parámetros. Como no existe una solución exacta en forma cerrada, la confianza que el modelo es correcto se funda sobre el éxito que tenga en predecir con razonable exactitud las magnitudes físicas de interés.

La complicación de la (12.1) proviene de la repulsión electrostática entre los electrones

$$\sum_{\substack{i,j=1 \\ i>j}}^N \frac{e^2}{|\mathbf{r}_i - \mathbf{r}_j|} \quad (12.2)$$

debido a que cada término de la (12.2) depende de las coordenadas de *dos* partículas.

Si no existieran esos términos, la (12.1) sería la ecuación de Schrödinger para N electrones, cada uno de los cuales se mueve en el campo producido por la carga nuclear Ze , sin interactuar con los demás. Este problema más simple es separable, y cada electrón se puede describir por medio de una función de onda del tipo estudiado en el Capítulo 10, con el agregado de la variable de spin. La función de onda total ψ_N se puede expresar entonces en términos del producto de las funciones de onda de los electrones individuales, y la energía total es la suma de las energías de cada uno de ellos, con lo que el problema queda resuelto.

Esta observación sugiere el programa a seguir para desarrollar nuestro modelo: tenemos que procurar aproximar la expresión (12.2) para transformarla en una suma, cada uno de cuyos términos sea función de las coordenadas de un sólo electrón, esto es, reemplazar

$$\sum_{\substack{i,j=1 \\ i>j}}^N \frac{e^2}{|\mathbf{r}_i - \mathbf{r}_j|} \quad \text{por} \quad \sum_{i=1}^N V_i'(r_i) \quad (12.3)$$

donde cada $V_i'(r_i)$ es un potencial *central*¹ que representa en forma *aproximada* el efecto sobre el i -ésimo electrón de la repulsión que ejercen sobre él los $N - 1$ electrones restantes. Con esta aproximación tenemos que

$$\mathcal{H} \cong \sum_{i=1}^N H_i \quad \text{con} \quad H_i = \frac{\mathbf{p}_i^2}{2\mu} - V_i(r_i) \quad , \quad V_i(r_i) = -\frac{Ze^2}{r_i} + V_i'(r_i) \quad (12.4)$$

y entonces la ecuación de Schrödinger (aproximada) independiente del tiempo es

$$\left(\sum_{i=1}^N H_i \right) \psi_N = \mathcal{E} \psi_N \quad (12.5)$$

La (12.5) admite soluciones separables de la forma

$$\psi_N(\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N) = \psi_{k_1}(\mathbf{r}_1) \psi_{k_2}(\mathbf{r}_2) \dots \psi_{k_i}(\mathbf{r}_i) \dots \psi_{k_N}(\mathbf{r}_N) \quad (12.6)$$

correspondientes al autovalor

$$\mathcal{E} = E_{k_1} + E_{k_2} + \dots + E_{k_i} + \dots + E_{k_N} \quad (12.7)$$

donde las $\psi_{k_i}(r_i)$ son soluciones de las ecuaciones

$$H_i \psi_{k_i}(\mathbf{r}_i) = \left[\frac{\mathbf{p}_i^2}{2\mu} - V_i(r_i) \right] \psi_{k_i}(\mathbf{r}_i) = E_{k_i} \psi_{k_i}(\mathbf{r}_i) \quad , \quad i = 1, \dots, N \quad (12.8)$$

¹ Esta suposición involucra dos aproximaciones. La primera consiste en reemplazar la interacción sobre el i -ésimo electrón debida a los demás electrones por el campo eléctrico debido a la distribución media de carga de los mismos. La segunda estriba en suponer que dicha distribución media de carga tiene simetría esférica. Es verdad que la distribución de carga de un electrón en un estado s es esféricamente simétrica, y también es cierto que lo es la que proviene de una subcapa (n, l) llena. Pero cuando hay subcapas parcialmente llenas, la correspondiente distribución media de carga no es simétrica. Por lo tanto la segunda aproximación implica reemplazar estas distribuciones de carga no simétricas por un promedio sobre todas las direcciones.

En la (12.6) y la (12.8) hemos omitido escribir en las ψ_{k_i} la variable de spin, que daremos por sobreentendida en lo que sigue, cuando corresponda. Cada una de las (12.8) es la ecuación de Schrödinger independiente del tiempo para *una única partícula* en un campo de fuerzas centrales dado por $V_i(r_i)$. Se puede notar que las ecuaciones (12.8) no cambian si se permutan entre sí las variables $\mathbf{r}_i$ ($i = 1, 2, \dots, N$). Por lo tanto podemos omitir los subíndices de las variables y escribir

$$H_i \psi_{k_i}(\mathbf{r}) = \left[\frac{\mathbf{p}^2}{2\mu} - V_i(r) \right] \psi_{k_i}(\mathbf{r}) = E_{k_i} \psi_{k_i}(\mathbf{r}) \quad , \quad i = 1, \dots, N \quad (12.9)$$

Por lo que vimos en los Capítulos 10 y 11, las soluciones de las (12.9) se pueden escribir como

$$\psi_{k_i}(\mathbf{r}, \sigma) = \psi_{nlm_l m_s}(\mathbf{r}, \sigma) = \mathcal{R}_{i,nl}(r) Y_l^{m_l}(\theta, \varphi) \chi_{m_s}(\sigma) \quad (12.10)$$

donde las $\mathcal{R}_{i,nl}(r)$ son las soluciones de las ecuaciones radiales apropiadas y σ es la variable de spin. Por lo tanto el estado de cada electrón queda especificado por los cuatro números cuánticos n, l, m_l, m_s y el conjunto de N de tales cuaternas (que indicamos con k), una para cada electrón, especifica el estado del átomo. Claramente, esta promete ser una manera sencilla de atacar el problema, y vale la pena examinar si se puede llevar adelante nuestro programa de manera razonable.

Conviene ahora detenerse un momento para analizar el significado de lo que estamos haciendo y hacer varias aclaraciones. En primer lugar es importante observar que (habiendo despreciado los efectos de spin) la energía E_{k_i} de cada electrón depende solamente de los números cuánticos n y l , y no de m_l y m_s . Cada par de valores (n, l) define lo que se denomina una *subcapa* del átomo, la cual se designa por medio de un número (que es el valor de n) seguido de una letra ($s, p, d, f, \dots$, de acuerdo con el valor 0, 1, 2, 3, ... de l). Así, por ejemplo, $3d$ designa la subcapa $n = 3$, $l = 2$. Por lo dicho, la energía de cada electrón es la misma para todos los $2(2l + 1)$ estados de cada subcapa y por consiguiente la energía total del átomo está determinada si se conoce cuántos electrones hay en cada subcapa. Esto es lo que se denomina la *configuración* del estado atómico que estamos considerando. La configuración se indica nombrando las subcapas ocupadas e indicando el número de electrones que reside en cada una de ellas por medio de un supraíndice. Por ejemplo $1s^2 2s^2 2p$ designa una configuración en la cual hay dos electrones en la subcapa $1s$, dos en la $2s$ y uno en la $2p$, y para esta configuración la energía total del átomo es entonces $\mathcal{E} = 2E_{1s} + 2E_{2s} + E_{2p}$. Debe quedar claro que cada configuración comprende varios estados diferentes ψ_N de la forma (12.6), todos los cuales tienen la misma energía, y que corresponden a los distintos valores que pueden tener los números cuánticos m_l y m_s de cada electrón.

También debemos notar que los V_i' , y por lo tanto los V_i que figuran en las (12.9) *dependen de la configuración*, porque el efecto sobre el i -ésimo electrón de la repulsión de los demás electrones depende de en qué subcapas residen éstos. Además, todos los V_i' están ligados entre sí, pues la distribución de carga de cada electrón deriva del respectivo $\mathcal{R}_{i,nl}(r)$ y contribuye a determinar la de los demás.

Hasta ahora no hemos dicho nada acerca de como calcular los V_i y por lo tanto como obtener las autofunciones ψ_{k_i} y los autovalores E_{k_i} . Eso lo veremos más adelante. Pero debe quedar claro que tanto la forma de las ψ_{k_i} como los valores de los E_{k_i} dependen de la configuración que estamos considerando, dado que ésta determina el conjunto de los V_i que se deben emplear en las

(12.9). Si se cambia la configuración es preciso entonces calcular de nuevo todas las ψ_{k_i} y los E_{k_i} . Por ejemplo, la energía de un electrón $1s$ no es la misma en la configuración $1s^2 2s^2 2p$ que en la $1s^2 2s 2p^2$ y las correspondientes funciones de onda, si bien son ambas $1s$, son distintas. Asimismo, debemos destacar que las diferentes ψ_{k_i} que se obtienen resolviendo las (12.9) *no son ortogonales*, puesto que son autofunciones de *diferentes* Hamiltonianos.

Por último es importante aclarar que cada uno de los estados ψ_N de la forma (12.6) pertenecientes a una dada configuración tiene una degeneración adicional. En efecto, si indicamos con $\mathcal{P}(\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N)$ una *permutación* de los argumentos de ψ_N , la función

$$\psi_N^{\mathcal{P}}(\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N) = \psi_N(\mathcal{P}(\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N)) \quad (12.11)$$

es también una solución de (12.5), correspondiente al mismo autovalor $\mathcal{E}$ (un ejemplo podría ser la función $\psi'_N(\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_N) = \psi_{k_1}(\mathbf{r}_2)\psi_{k_2}(\mathbf{r}_1)\dots\psi_{k_j}(\mathbf{r}_i)\dots\psi_{k_N}(\mathbf{r}_N)$, obtenida a partir de ψ_N por la transposición de los argumentos $\mathbf{r}_1$ y $\mathbf{r}_2$). Por consiguiente, todas las $N!$ funciones $\psi_N^{\mathcal{P}}$, obtenidas a partir de la (12.6) permutando sus N argumentos de todas las maneras posibles, son soluciones de la (12.5) correspondiente al mismo autovalor. Tendríamos entonces $N!$ autofunciones degeneradas de $\mathcal{H}$ correspondientes al autovalor $\mathcal{E}$. Por lo que sabemos *hasta ahora*, cualquier combinación lineal de ellas podría describir un posible estado del sistema. Veremos, sin embargo, que debido a que los electrones no se pueden distinguir el uno del otro, tan sólo *una particular* combinación lineal entre todas es admisible. En el resto de este Capítulo no intentaremos escribir la ψ_N , dado que podemos avanzar sin necesidad de hacerlo, y recién en el Capítulo 13 volveremos sobre este asunto.

El método del campo autoconsistente

Por el momento concentraremos nuestra atención sobre la dependencia de las funciones de onda ψ_k de un electrón en las variables espaciales, que es la que determina la distribución en el espacio de la carga eléctrica del átomo (y por lo tanto la repulsión Coulombiana entre los electrones). Para ver como se puede llevar adelante nuestro programa, consideremos un problema concreto. Imaginemos el estado fundamental de un átomo de helio ionizado, en el cual el único electrón tiene una función de onda $1s$ dada por la ec. (10.138) con $Z = 2$, cuya parte espacial es

$$\psi_{1,0,0} = Ae^{-2r_1/a_0} \quad , \quad A = \text{cte.} \quad (12.12)$$

donde a_0 es el radio de Bohr. Supongamos ahora que el ion captura otro electrón, en un estado débilmente ligado de momento angular grande, por ejemplo $3d$. Debido a que su momento angular es grande, este electrón tiene una probabilidad despreciable de encontrarse cerca del núcleo, lo cual implica dos cosas:

- puesto que el segundo electrón está casi siempre lejos del núcleo y fuera de la región donde se puede encontrar el primer electrón con una probabilidad apreciablemente diferente de cero, se mueve esencialmente en el campo Coulombiano de *una única carga* (la carga nuclear apantallada por el primer electrón). Por lo tanto, la parte espacial de su función de onda debe ser muy semejante a la correspondiente función hidrogenoide $\psi_{n,l,m_l}(\mathbf{r}_2)$ calculada ahora con $Z = 1$;
- puesto que el segundo electrón no se acerca a la región ocupada por el primer electrón, su presencia perturba muy poco la función de onda de este último, la cual entonces difiere muy poco de la (12.12).

Por lo tanto es razonable suponer que con buena aproximación las funciones de onda del primer y segundo electrón sean, respectivamente,

$$\psi'_{1,1s} \approx \psi_{1,0,0} \quad , \quad \psi'_{2,(n,l,m_l)} \approx \psi_{n,l,m_l} \quad (12.13)$$

Para llegar a este resultado, lo que en realidad hicimos fue sustituir la interacción entre los electrones (dada por los términos (12.2)) por el potencial eléctrico debido a la distribución de carga dada por la función de onda del otro electrón. Además, como el segundo electrón está siempre lejos, hemos supuesto que la carga del primer electrón está concentrada en el origen.

Esto sugiere que se puedan obtener soluciones bastante aproximadas de la (12.1) *para cada configuración* de interés sustituyendo los términos de repulsión entre los electrones por una serie de potenciales $V_i(r)$, uno para cada electrón, calculados a partir de la distribución de carga que resulta de las funciones de onda de los demás electrones. Claramente, el método de solución debe ser iterativo. Se comienza con una hipótesis acerca del potencial que siente cada electrón, como hicimos recién en el ejemplo del átomo de helio excitado. A partir de esos potenciales obtenemos las funciones de onda de cada electrón resolviendo las N ecuaciones de Schrödinger independientes del tiempo (12.9). A continuación calculamos un nuevo potencial para cada electrón a partir de la carga nuclear, más la distribución de carga que resulta de las funciones de onda de los demás electrones que se acaban de calcular. Con estos nuevos potenciales calculamos nuevas funciones de onda, y así seguimos realizando sucesivas iteraciones. Siempre y cuando nuestro punto de partida no haya sido muy errado, después de varias iteraciones encontraremos que el proceso converge, y las nuevas funciones de onda reproducen las distribuciones de carga a partir de la cual fueron calculadas. Esta es la esencia del *método del campo autoconsistente*, o *método de Hartree*, que fue introducido por Douglas Hartree en 1928.

En la práctica resulta que el método del campo autoconsistente es una aproximación de gran valor, pues permite atribuir una función de onda a cada electrón y por lo tanto describir un átomo dando los números cuánticos de las funciones de onda de cada uno de sus electrones. Además, estas funciones de onda no difieren mucho de las funciones de onda de átomos hidrogenoides. La parte angular es, naturalmente, la misma, porque no depende de la forma del potencial (siempre que éste sea central). En cuanto a la parte radial, para un dado n (que determina el número de nodos de $\mathcal{R}_{i,nl}(r)$), es cualitativamente parecida a la de un átomo hidrogenoide.

En resumidas cuentas, podemos concluir lo siguiente:

- En un átomo con varios electrones la interacción de un dado electrón con los demás se puede sustituir, con buena aproximación, por el potencial $V_i(r)$ que resulta de la distribución media de carga de los demás electrones.
- Se puede considerar entonces que cada electrón se mueve independientemente en el potencial $V_i(r)$ que resulta de la carga nuclear más la que se debe a los demás electrones, y que se suele llamar *potencial Coulombiano apantallado*.
- Los niveles de energía del átomo se pueden determinar a partir del conjunto de números cuánticos (n, l) de las N funciones de onda individuales de cada electrón. Este conjunto constituye lo que se llama la *configuración* electrónica del átomo.
- Las funciones de onda individuales de cada electrón no son muy diferentes de las funciones de onda hidrogenoides correspondientes a los mismos (n, l) .

Esperamos que el modelo aproximado que hemos esbozado describa razonablemente bien las propiedades de los átomos con varios electrones. Una de las razones para esperar que el modelo

funcione es que la interacción entre los electrones es una función que varía lentamente con la distancia y por lo tanto al sustituir la interacción real por un promedio no se está cometiendo un error demasiado grande. Veremos, sin embargo, que el modelo del átomo, como lo desarrollamos hasta aquí, ni por asomo puede reproducir las características que se observan en la tabla periódica. Para resolver la dificultad es necesario introducir un nuevo postulado en la teoría.

Propiedades de los elementos

La característica más notable del conjunto de los elementos químicos es que se pueden ordenar en la forma de una tabla periódica, lo cual como ya se dijo en el Capítulo 3 fue hecho por primera vez en 1869 por Dmitri Mendeleev. Existen diferentes versiones o maneras de presentar la tabla periódica, una de las cuales se muestra en la Fig. 12.1. Cuando los elementos se ordenan en la tabla de izquierda a derecha y de arriba abajo por orden de número atómico (Z) creciente, se observa que los elementos que pertenecen a la misma columna (o *grupo*) tienen propiedades químicas semejantes, pero hay una rápida variación de propiedades a medida que se recorren las filas (*períodos* y *series*) de la tabla. Así, todos los elementos del Grupo I (primera columna) son monovalentes, los del Grupo II son bivalentes, y los del Grupo 0 (última columna) son gases nobles, químicamente inertes.

Ahora bien, según la teoría que hemos esbozado, el estado fundamental de cada átomo es aquél en que cada electrón está en el estado de menor energía en el campo eléctrico producido por la carga nuclear y la distribución de carga de los demás electrones. Por otra parte, el estado más bajo de cualquier potencial es siempre un estado s , y por consiguiente las funciones de onda de cada electrón serán todas iguales. En efecto, si las funciones de onda de todos los electrones son idénticas, los potenciales en que se mueve cada electrón son idénticos, y por lo tanto dan lugar a funciones de onda idénticas respetando la autoconsistencia. Por lo tanto la configuración del estado fundamental de un átomo con N electrones sería $1s^N$, en la cual todos los electrones tendrían funciones de onda idénticas, iguales a la función del estado s de más baja energía que se puede presentar en el campo producido por el núcleo y por los demás electrones.

De ser cierta la anterior conclusión resulta lógico pensar que, al pasar de un elemento al siguiente, la naturaleza de esa función de onda cambiará muy poco, ya que sigue siendo del mismo tipo, y el agregado de una carga nuclear y un nuevo electrón no implica una variación sustancial del potencial. Por lo tanto, si bien se puede imaginar que haya cierta variación de las propiedades químicas entre los primeros elementos, a medida que Z aumenta los cambios deberían ser cada vez más pequeños.

Por lo tanto salta a la vista de inmediato que el modelo, así como está, no podrá *nunca* reproducir las características regulares y repetitivas de la tabla periódica, como ser la recurrencia de los metales alcalinos y de los gases nobles, y el aumento constante de la valencia al atravesar un período. Por lo tanto *no sirve* para explicar las características más evidentes de la química de los elementos.

Podemos tener indicios sobre la causa de las fallas del modelo si examinamos las primeras *energías de ionización* E_i de los átomos (Fig. 12.2), que dan la diferencia de energía entre el estado fundamental del átomo neutro y el estado fundamental del átomo ionizado una vez (es decir, que ha perdido un electrón).

Grupo	I	II	III	IV	V	VI	VII	VIII			0	
Período	Serie											
1	1	1 1 H									4 2 He	
2	2	7 3 Li	9 4 Be	11 5 B	12 6 C	14 7 N	16 8 O	19 9 F				20 10 Ne
3	3	23 11 Na	24 12 Mg	27 13 Al	28 14 Si	31 15 P	32 16 S	35 17 Cl				40 18 Ar
4	4	72 19 K	72 20 Ca	72 21 Sc	72 22 Ti	72 23 V	72 24 Cr	72 25 Mn	72 26 Fe	72 27 Co	72 28 Ni	
	5	64 29 Cu	65 30 Zn	70 31 Ga	72 32 Ge	75 33 As	79 34 Se	80 35 Br				84 36 Kr
5	6	85 37 Rb	88 38 Sr	89 39 Y	91 40 Zr	93 41 Nb	96 42 Mo	99 43 Tc	101 44 Ru	103 45 Rh	106 46 Pd	
	7	108 47 Ag	112 48 Cd	115 49 In	119 50 Sn	122 51 Sb	128 52 Te	127 53 I				131 54 Xe
6	8	133 55 Cs	137 56 Ba	57-71 *	178 72 Hf	181 73 Ta	184 74 W	186 75 Re	190 76 Os	192 77 Ir	195 78 Pt	
	9	197 79 Au	201 80 Hg	204 81 Tl	207 82 Pb	209 83 Bi	210 84 Po	210 85 At				222 86 Rn
7	10	223 87 Fr	226 88 Ra	89-103 **	261 104 Rf	262 105 Db	266 106 Sg	264 107 Bh	269 108 Hs	268 109 Mt	271 110 --	

* Lantánidos:

$\begin{smallmatrix} 139 \\ 57 \end{smallmatrix} \text{La}$	$\begin{smallmatrix} 140 \\ 58 \end{smallmatrix} \text{Ce}$	$\begin{smallmatrix} 141 \\ 59 \end{smallmatrix} \text{Pr}$	$\begin{smallmatrix} 144 \\ 60 \end{smallmatrix} \text{Nd}$	$\begin{smallmatrix} 145 \\ 61 \end{smallmatrix} \text{Pm}$	$\begin{smallmatrix} 150 \\ 62 \end{smallmatrix} \text{Sm}$	$\begin{smallmatrix} 152 \\ 63 \end{smallmatrix} \text{Eu}$	$\begin{smallmatrix} 157 \\ 64 \end{smallmatrix} \text{Gd}$	$\begin{smallmatrix} 159 \\ 65 \end{smallmatrix} \text{Tb}$	$\begin{smallmatrix} 162 \\ 66 \end{smallmatrix} \text{Dy}$	$\begin{smallmatrix} 165 \\ 67 \end{smallmatrix} \text{Ho}$	$\begin{smallmatrix} 167 \\ 68 \end{smallmatrix} \text{Er}$	$\begin{smallmatrix} 169 \\ 69 \end{smallmatrix} \text{Tm}$	$\begin{smallmatrix} 173 \\ 70 \end{smallmatrix} \text{Yb}$	$\begin{smallmatrix} 175 \\ 71 \end{smallmatrix} \text{Lu}$
---	---	---	---	---	---	---	---	---	---	---	---	---	---	---

** Actínidos:

$\begin{smallmatrix} 227 \\ 89 \end{smallmatrix} \text{Ac}$	$\begin{smallmatrix} 232 \\ 90 \end{smallmatrix} \text{Th}$	$\begin{smallmatrix} 231 \\ 91 \end{smallmatrix} \text{Pa}$	$\begin{smallmatrix} 238 \\ 92 \end{smallmatrix} \text{U}$	$\begin{smallmatrix} 237 \\ 93 \end{smallmatrix} \text{Np}$	$\begin{smallmatrix} 244 \\ 94 \end{smallmatrix} \text{Pu}$	$\begin{smallmatrix} 243 \\ 95 \end{smallmatrix} \text{Am}$	$\begin{smallmatrix} 247 \\ 96 \end{smallmatrix} \text{Cm}$	$\begin{smallmatrix} 247 \\ 97 \end{smallmatrix} \text{Bk}$	$\begin{smallmatrix} 251 \\ 98 \end{smallmatrix} \text{Cf}$	$\begin{smallmatrix} 252 \\ 99 \end{smallmatrix} \text{Es}$	$\begin{smallmatrix} 257 \\ 100 \end{smallmatrix} \text{Fm}$	$\begin{smallmatrix} 258 \\ 101 \end{smallmatrix} \text{Md}$	$\begin{smallmatrix} 259 \\ 102 \end{smallmatrix} \text{No}$	$\begin{smallmatrix} 262 \\ 103 \end{smallmatrix} \text{Lr}$
---	---	---	--	---	---	---	---	---	---	---	--	--	--	--

Fig. 12.1. La tabla periódica de los elementos. Los elementos se identifican por su símbolo químico, el número atómico Z (subíndice) y el número de masa A (supraíndice), que es el valor del peso atómico expresado en unidades del peso atómico del hidrógeno, redondeado al entero más próximo cuando el elemento tiene más de un isótopo estable. Los número de masa de los elementos con $Z > 83$ corresponden al isótopo más estable. El elemento 110 todavía no tiene nombre ni símbolo aceptado internacionalmente. Los elementos 111 y 112 han sido descubiertos pero no los hemos incluido en la tabla.

Consideremos el estado de menor energía de un electrón en un átomo con Z electrones. Cuando se encuentra lejos del núcleo, la energía potencial del electrón es esencialmente igual a

$$V_{\infty}(r) = -\frac{e^2}{r}, \quad r \rightarrow \infty \quad (12.14)$$

porque la carga positiva Ze del núcleo está apantallada por los $Z - 1$ electrones restantes que lo rodean. Por otra parte, muy cerca del núcleo el apantallamiento debido a los demás electrones es despreciable y la energía potencial de un electrón es

$$V_0(r) = -\frac{Ze^2}{r}, \quad r \rightarrow 0 \quad (12.15)$$

Por lo tanto tendremos que

$$-\frac{Ze^2}{r} < V(r) < -\frac{e^2}{r} \quad (12.16)$$

Claramente, si la energía potencial fuera $V_\infty(r)$ en todas partes, la energía de ionización sería de 13.6 eV como en el átomo de hidrógeno; si en cambio la energía potencial fuese $V_0(r)$ para todo r , tendríamos que $E_i \approx Z \times 13.6$ eV. El hecho que $V_0(r) < V(r) < V_\infty(r)$ implica que

$$13.6 \text{ eV} < E_i < Z \times 13.6 \text{ eV} \quad (12.17)$$

y por lo tanto si las funciones de onda de todos los electrones fuesen la que corresponde a la menor energía, sería razonable pensar que E_i crezca con Z , pero más lentamente que lo que resulta de una relación lineal. Sin embargo mirando la Fig. 12.2 se observa que en lugar de mostrar un constante y paulatino incremento con Z , los valores de E_i tienen el mismo tipo de comportamiento periódico que las propiedades químicas.

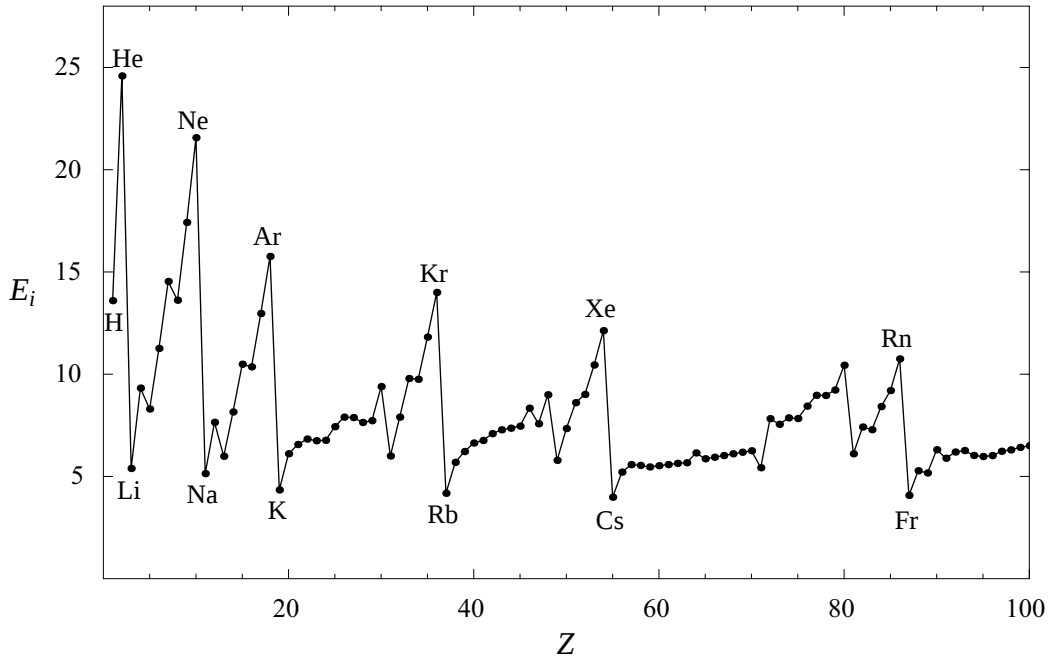


Fig. 12.2. Energías de ionización (en eV) de las diferentes especies atómicas.

Estudiando los valores correspondientes a los primeros elementos se puede entender lo que está ocurriendo. Para el He, $E_i = 24.5$ eV, algo menos del doble que para el hidrógeno. Hasta aquí no hay problemas, pero el elemento siguiente, el litio, tiene $E_i = 4.6$ eV, que es un valor llamativamente pequeño, muy inferior al del hidrógeno, mientras que en base a la (12.17) se esperaría un valor entre 30 y 40 eV. *Es imposible* tener una energía de ionización tan baja si todos los electrones tienen un número cuántico principal $n = 1$. Sin embargo, si uno de los electrones estuviera en un estado con $n = 2$, se entendería mejor el resultado experimental. En efecto, supongamos que el Li tiene dos de sus tres electrones en el estado más bajo ($1s$) y el tercero en el siguiente estado disponible que es el $2s$ (se prefiere este estado y no el $2p$ pues el experimento de Stern-Gerlach con átomos neutros de Li muestra que el estado fundamental es un estado s). La función de onda $2s$ en el campo de una sola carga da una densidad de probabilidad *muy pequeña* para $r < 2a_0$ (Fig. 12.3); por lo tanto el electrón pasa la mayor parte del tiempo a distancias *mayores* que $2a_0$, mientras que los dos electrones $1s$ tienen más del 99% de probabilidad de estar

en $r < 2a_0$. Por consiguiente cabe esperar que el tercer electrón del Li tenga una función de onda muy parecida a la $2s$ del átomo de hidrógeno. Luego su energía de ionización debe ser semejante a la del estado $2s$ del hidrógeno, que es de 3.4 eV, un valor mucho más cercano a los 4.6 eV reales que no los 30-40 eV que estimamos antes. La diferencia entre la nueva estimación y el valor verdadero tiene además el signo correcto, pues el electrón $2s$ del litio tiene una probabilidad pequeña pero *no nula* de estar más cerca del núcleo, lo cual tiende a *aumentar* la energía de ionización por encima de los 3.4 eV. Este modelo del átomo de litio resulta aún más plausible si se considera la *segunda* energía de ionización, esto es, la diferencia de energía entre el ion Li^+ y el ion Li^{++} , que es de 75.3 eV, lo que muestra que los dos electrones del Li^+ están mucho más fuertemente ligados que el tercer electrón del Li neutro, como se debe esperar si están en el estado $1s$.

Veamos ahora el comportamiento de la primera energía de ionización de los ocho átomos siguientes (Be, B, C, N, O, F, Ne y Na). Aunque se observan pequeñas irregularidades, E_i tiende a aumentar con Z , hasta que se llega al Ne. Este comportamiento es el que cabe esperar si los electrones que se van agregando están todos en estados con $n = 2$, puesto que la presencia de los otros electrones con $n = 2$ sólo neutraliza parcialmente el aumento de la carga nuclear. Pero al pasar del Ne al Na, se observa nuevamente una notable reducción de E_i , y con argumentos semejantes a los que usamos para el caso del litio podemos explicar esta disminución suponiendo que el último electrón del Na está en un estado con $n = 3$.

Por consiguiente el comportamiento de las primeras energías de ionización sugiere que en un sistema atómico formado por electrones que se mueven en un campo autoconsistente con funciones de onda caracterizadas por los números cuánticos n, l, m_l, m_s puede haber, cuanto mucho, dos electrones con $n = 1$ y ocho con $n = 2$. ¿De dónde surgen esos números? Es inmediato ver que para $n = 1$ hay 2 estados *diferentes*, los estados $1s$ con $m_s = +1/2$ y $m_s = -1/2$. Asimismo, para $n = 2$ hay exactamente 8 estados *diferentes*: los dos estados $2s$ y los 6 estados $2p$ (ya que si $l = 1$, m_l puede tomar los valores $-1, 0, +1$, y para cada uno de ellos m_s puede valer $-1/2$ ó $+1/2$). Todo esto sugiere algo muy simple, pero al mismo tiempo *nuevo e inexplicable* en base a los conocimientos que se tenían allá en 1920: que *en cada estado* (n, l, m_l, m_s) *sólo cabe*, por así decir, *un electrón*.

El Principio de Exclusión

Durante varios años, la interpretación de la estructura de los átomos con más de un electrón fue causa de perplejidad para los físicos. Ya Bohr, en sus trabajos iniciales, había reconocido que la pregunta de porqué los electrones no se encontraban todos ligados en la capa más profunda constituía un problema fundamental, cuya respuesta no se podía encontrar en la Mecánica Clásica. La respuesta a este interrogante la dio Wolfgang Pauli en 1925, con base en un análisis de los datos de los niveles de energía de los átomos, del tipo que hemos comentado. Está contenida en un nuevo postulado, o principio, que en su enunciación primitiva establecía:

Principio de Exclusión:

- En un átomo multielectrónico, nunca puede existir más de un electrón en cada estado cuántico.

Es interesante mencionar que cuando Pauli formuló este principio aún no se conocía el spin del electrón. Si no se cuenta el spin, los estados de un electrón atómico se caracterizan por *tres* números cuánticos: n, l y m_l . Sin embargo Pauli en su artículo asignó *cuatro* números cuánticos al

electrón. Lo hizo de manera puramente formal, sin basarse en un esquema concreto, simplemente porque le hacía falta para obtener el resultado que buscaba. Fue precisamente al leer ese artículo que Uhlenbeck y Goudsmit, ponderando sobre cuál podría ser el significado físico del misterioso cuarto número cuántico, concibieron la idea del spin (Capítulo 11).

También resulta curioso que Bohr, quien en su momento tuvo la audacia de apartarse de la Mecánica y la Electrodinámica Clásicas para establecer sus célebres postulados, no haya dado el paso que dio Pauli, pese a tener plena conciencia de la importancia del problema a resolver.

Veremos en el próximo Capítulo que el Principio de Exclusión se relaciona con la *indistinguibilidad* de las partículas en la Mecánica Cuántica. Asimismo, como lo demostró 15 años después el mismo Pauli, se puede *deducir* a partir de la Teoría Cuántica de Campos. Pero todo eso no se sabía en 1925. Cuando fue formulado, el Principio de Exclusión no era más que un “decreto”, promulgado con el propósito introducir la regla que estaba faltando en la teoría de la estructura atómica y así legitimar el comportamiento que muestra la experiencia. Pero el de Pauli fue un decreto ciertamente muy inspirado, pues no sólo logró su objetivo original, sino que además puso orden en las propiedades de la materia en todas las escalas. En efecto, gracias al Principio de Exclusión hoy podemos entender desde las propiedades de la materia en el interior de las estrellas hasta la estructura de las partículas del núcleo, pasando por la impenetrabilidad de los sólidos, la razón de porqué ciertos medios conducen la electricidad mientras otros son aislantes, etc..

Antes de analizar la indistinguibilidad de las partículas y sus implicancias, vamos a mostrar que gracias al Principio de Exclusión podemos explicar la estructura de los átomos con varios electrones, y entender la lógica que está detrás de la Tabla Periódica de los elementos.

El Principio de Exclusión y la estructura atómica

Debido al Principio de Exclusión, en el estado fundamental de un átomo los electrones no tienen la misma energía y función de onda, sino que sus números cuánticos n, l, m_l, m_s (y las correspondientes funciones de onda) son todos distintos. Los valores de m_l y m_s dependen de la elección arbitraria de la dirección del eje z ; por lo tanto los electrones que tienen iguales n y l , pero diferentes m_l y m_s , tienen la misma energía $E_{n,l}$, que depende tanto de n como de l pues el campo autoconsistente no es Coulombiano.

Ahora bien, los electrones con diferente n tienen funciones de onda con extensiones espaciales distintas, como se ve en la Fig. 12.3 donde se muestra la distribución radial de probabilidad $P(r) = |u(r)|^2$ de los tres primeros estados s del hidrógeno. La diferencia es aún mayor en átomos con muchos electrones, ya que los electrones con n pequeño pasan más tiempo cerca del núcleo, donde la carga nuclear no está apantallada, luego su distribución radial de probabilidad se distorsiona respecto de la del átomo de hidrógeno, desplazándose hacia r menores.

Para los electrones con mayor n , la carga nuclear está fuertemente apantallada por los electrones con n menor, pero la recíproca, por supuesto, no es cierta. Luego los electrones que están en la capa $1s$ tienen una energía potencial que es prácticamente igual a $V_0(r)$, la que produce la carga Ze del núcleo. Por lo tanto se pueden representar bastante bien por medio de la función de onda

$$\psi_{1,0,0} = Ae^{-Zr/a_0} \quad , \quad A = \text{cte.} \quad (12.18)$$

cuya extensión radial es $1/Z$ veces la que corresponde al hidrógeno. En cambio, los electrones exteriores con el mayor n se mueven en un campo que varía entre $V_0(r)$ para r pequeño y $V_\infty(r)$

para r grande. Luego la parte interior de la función de onda (y por ende la distribución radial de probabilidad) se contrae, pero la parte externa queda casi igual a la del átomo de hidrógeno, y se acentúa la diferencia de energía entre electrones con diferente n . El resultado neto es que los electrones externos de todos los átomos tienen casi exactamente la misma extensión (dos o tres veces el radio de Bohr) como lo indica el valor casi constante de los radios atómicos.

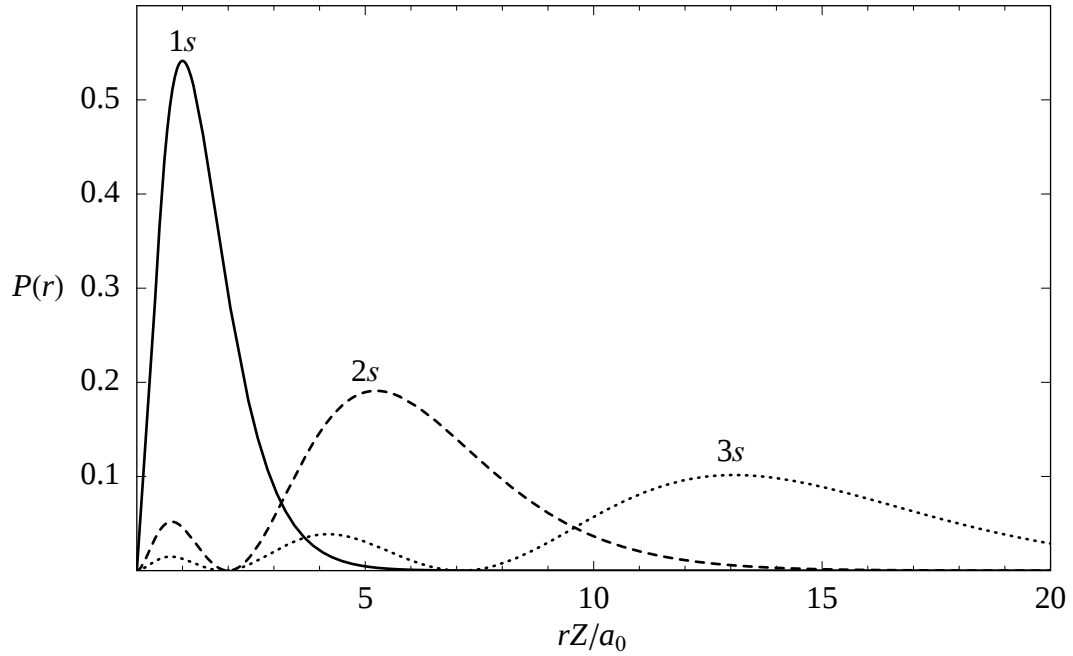


Fig. 12.3. Distribución de probabilidad radial para los tres primeros estados s del hidrógeno.

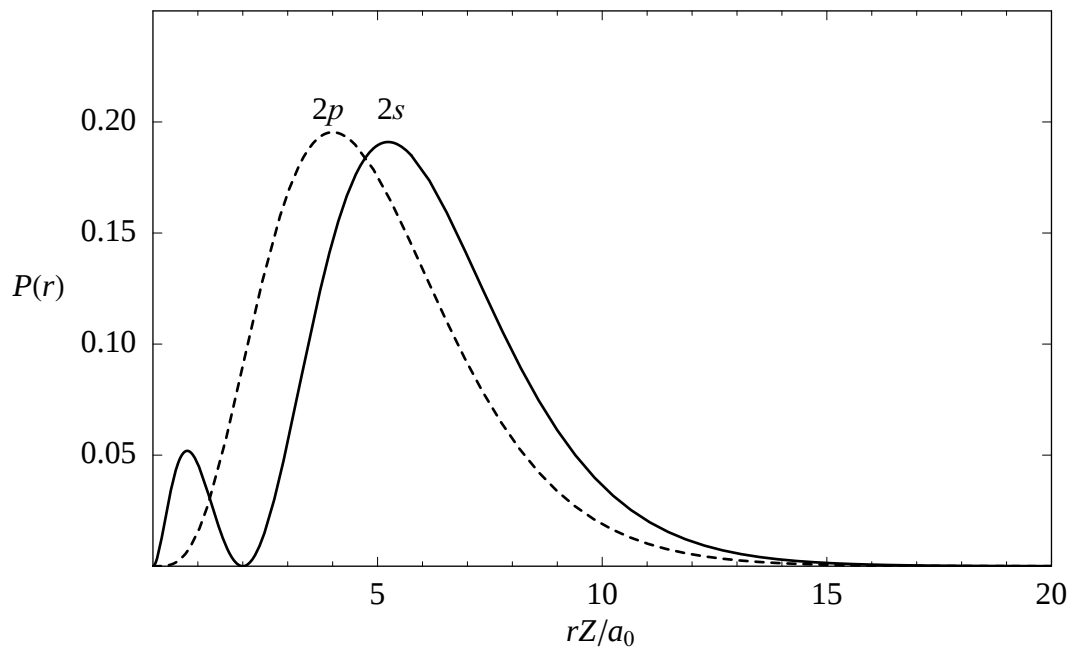


Fig. 12.4. Distribución de probabilidad radial de los estados $2s$ y $2p$ del hidrógeno.

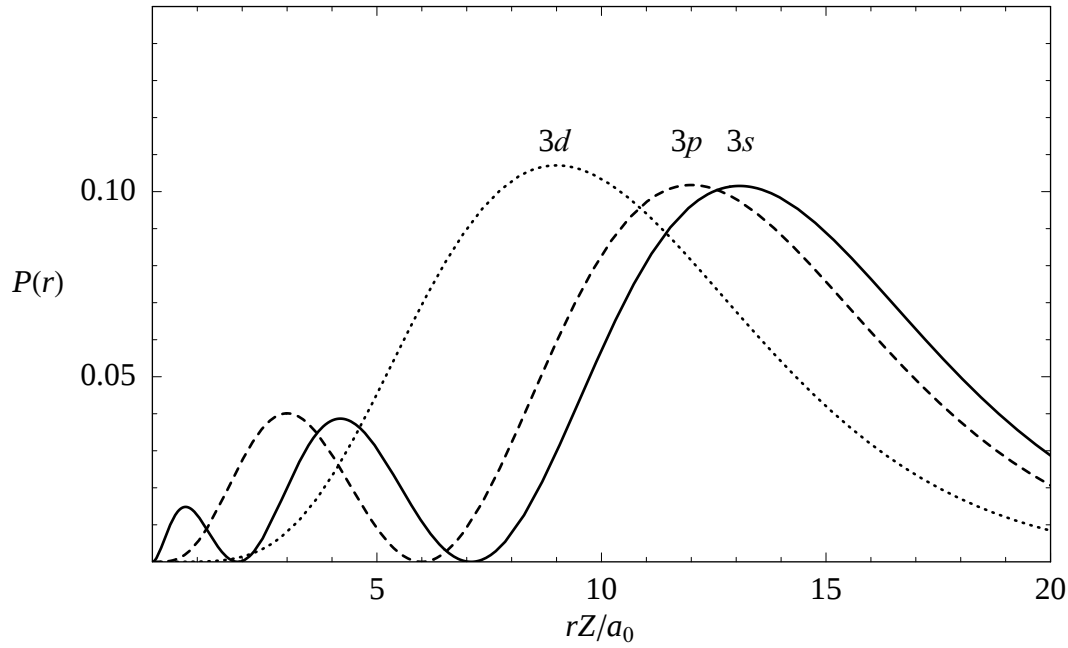


Fig. 12.5. Distribución de probabilidad radial de los estados 3s, 3p y 3d del hidrógeno.

Las Figs. 12.4 y 12.5 muestran las distribuciones de probabilidad radial de los estados de diferente l de las capas con $n = 2$ y $n = 3$ del hidrógeno, respectivamente. Es importante observar el comportamiento de $P(r)$ para r pequeño. Se ve que a medida que l aumenta, la barrera centrífuga empuja al electrón cada vez más lejos del núcleo. Recordemos, en efecto, que para $r \rightarrow 0$ se tiene que $u(r) \sim r^{l+1}$ y entonces $P(r) \sim r^{2l+2}$. Esto implica que a igual n , el apantallamiento de la carga nuclear (debido a los electrones que están más cerca del núcleo) es tanto mayor cuanto más grande es l . Este efecto rompe la degeneración de los niveles con igual n y diferente l que existe para el potencial Coulombiano: la energía de los niveles *aumenta* con l , de modo que

$$E_{2s} < E_{2p} \quad , \quad E_{3s} < E_{3p} < E_{3d} \quad , \quad \dots \quad (12.19)$$

También resulta que

$$E_{4s} < E_{3d} \quad , \quad E_{5s} < E_{4d} \quad , \quad E_{6s} < E_{5d} \quad , \quad \dots \quad (12.20)$$

Esto es, por efecto de la barrera centrífuga los niveles de las capas con diferente n se separan lo suficiente como para intercalarse entre sí.

Los efectos de la extensión espacial de las funciones de onda de diferente n , y de la barrera centrífuga para las funciones de onda del mismo n y diferente l que acabamos de comentar se combinan para producir el ordenamiento de niveles que se indica esquemáticamente en la Fig. 12.6. De acuerdo con el Principio de Exclusión cada *subcapa* n, l puede contener $2(2l + 1)$ electrones. Por lo tanto a medida que se recorre la Tabla Periódica a partir del hidrógeno, las subcapas se van llenando comenzando por la 1s y en orden de energía creciente. En la figura se consignan los elementos cuyos electrones de mayor energía se encuentra en cada subcapa.

Número de electrones que caben en cada subcapa				Z
(2)	(6)	(10)	(14)	
$\frac{7s}{\text{Fr, Ra}}$		$\frac{6d}{104 \rightarrow \dots}$	$\frac{5f}{\text{Ac} \rightarrow \text{Lr}}$	
-----				86
$\frac{6s}{\text{Cs, Ba}}$	$\frac{6p}{\text{Tl} \rightarrow \text{Rn}}$	$\frac{5d}{\text{Hf} \rightarrow \text{Hg}}$	$\frac{4f}{\text{La} \rightarrow \text{Lu}}$	
-----				54
$\frac{5s}{\text{Rb, Sr}}$	$\frac{5p}{\text{In} \rightarrow \text{Xe}}$	$\frac{4d}{\text{Y} \rightarrow \text{Cd}}$		
-----				36
$\frac{4s}{\text{K, Ca}}$	$\frac{4p}{\text{Ga} \rightarrow \text{Kr}}$	$\frac{3d}{\text{Sc} \rightarrow \text{Zn}}$		
-----				18
$\frac{3s}{\text{Na, Mg}}$	$\frac{3p}{\text{Al} \rightarrow \text{A}}$			
-----				10
$\frac{2s}{\text{Li, Be}}$	$\frac{2p}{\text{B} \rightarrow \text{Ne}}$			
-----				2
$\frac{1s}{\text{H, He}}$				

energía ↑

Fig. 12.6. Orden de las subcapas. Los efectos de la extensión espacial de las funciones de onda de diferente n y de la barrera centrífuga de las funciones de onda del mismo n y diferente l hacen que los niveles se ordenen como se indica en la figura (la energía crece de abajo hacia arriba, la escala no es lineal). Cada nivel con n, l dados se denomina *subcapa*. A medida que se recorre la Tabla Periódica a partir del hidrógeno, las subcapas se van llenando en orden de energía creciente empezando por la $1s$. En la figura se mencionan los elementos para los cuales el electrón de mayor energía se encuentra en la subcapa indicada.

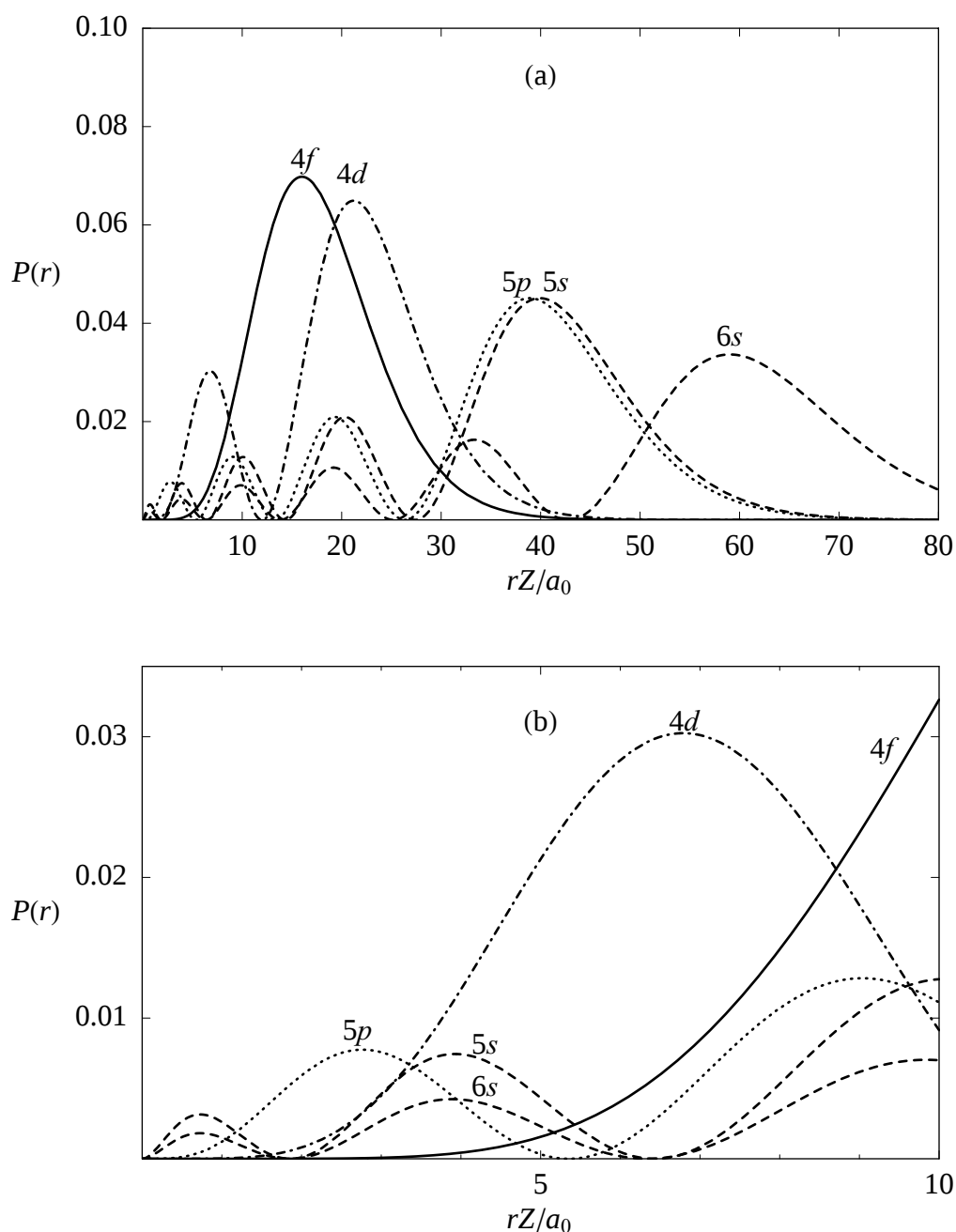


Fig. 12.7. Densidad de probabilidad radial para el nivel $4f$ y para los niveles $5s$, $6s$, $5p$ y $4d$ del átomo de hidrógeno. En (a) se puede apreciar que los electrones $4f$ se mueven en el interior del átomo, pues su extensión es menor que la de los electrones de los demás niveles representados. En (b) se puede apreciar que la barrera centrífuga no permite que los electrones $4f$ se acerquen al núcleo y entonces en los átomos con muchos electrones son fuertemente apantallados por los que ocupan las otras subcapas. Por lo tanto están débilmente ligados y la energía de la subcapa $4f$ es más alta que las de las subcapas $5s$, $6s$, $5p$ y $4d$, como se muestra en la Fig. 12.6. Los lantánidos (el grupo constituido por el La y los siguientes 14 elementos en los cuales se va llenando la subcapa $4f$) tienen propiedades químicas muy semejantes entre sí, porque los electrones $4f$, que se mueven en el interior del átomo, tienen una influencia casi nula sobre ellas.

Por ejemplo, la *configuración* del estado fundamental del K es $1s^2 2s^2 2p^6 3s^2 3p^6 4s^1$. Las líneas horizontales de trazos de la Fig. 12.6, entre las subcapas $1s$ y $2s$, $2p$ y $3s$, $3p$ y $4s$, ... etc., indican que la separación en energía entre cada uno de esos pares de subcapas es siempre grande.

Nuestro diagrama es *cualitativo*, pues en realidad si se quieren hacer predicciones cuantitativas exactas es necesario calcular un esquema como el de la Fig. 12.6 *para cada átomo* (con el método de Hartree). Cuando se hace eso, al pasar de un átomo a otro aparecen pequeñas diferencias respecto del orden de algunas de las subcapas, respecto del que se ve en la Fig. 12.6. Por ejemplo, las subcapas $4f$ y $5d$, que están muy próximas en energía, no siempre se encuentran en el orden indicado en la Fig. 12.6. Lo que sí es *siempre* correcto es el orden (por l creciente) de las diferentes subcapas de una dada capa y el orden (por n creciente) de subcapas del mismo l pero distinto n , que se presentan encolumnadas en la Fig. 12.6.

Los elementos cuya subcapa llena de mayor energía es np (como Ne, Ar, Kr, Xe y Rn) son los *gases nobles*. Puesto que dos de sus electrones exteriores están en la subcapa ns , no proporcionan un apantallamiento eficaz de la carga del núcleo. Por lo tanto los electrones np están ligados muy fuertemente. La energía de ionización es entonces muy alta y esos átomos son químicamente inertes pues no comparten fácilmente sus electrones externos con otros átomos como es necesario para que se formen compuestos. El He también es un gas noble, dado que sus dos electrones $1s$ están muy fuertemente ligados.

Se dice que los gases nobles tienen una configuración de *capas cerradas*, dado que el siguiente electrón se tiene que acomodar en la subcapa cuyo número cuántico principal es $n' = n + 1$, como se advierte observando la Fig. 12.6. Se debe notar, sin embargo, que debido al ordenamiento de energía (12.18), para todo $n > 2$ la capa n se cierra *antes* de que se llenen los niveles con l más alto (es decir, las subcapas nd y nf), como se ve en la Fig. 12.6. Esta circunstancia explica la aparición de los *lantánidos* y los *actínidos*, que son dos grupos de 15 elementos que comienzan, respectivamente, a partir del La y el Ac. En los lantánidos se va llenando la subcapa $4f$, que como se puede apreciar en la Fig. 12.7 se encuentra en el interior del átomo pese a tener más energía que las subcapas *externas* $6s$, $5p$ y $4d$ que están llenas. Por este motivo los elementos del grupo de los lantánidos tienen propiedades químicas casi idénticas, lo cual hace difícil separarlos. El caso de los actínidos es análogo, y en ellos se va llenando la subcapa $5f$, que se encuentra en el interior del átomo igual que la $4f$.

El ordenamiento de niveles de la Fig. 12.6, junto con el Principio de Exclusión, permiten entender las principales características de la Tabla Periódica y las interacciones entre los átomos de diferentes elementos.

La unión química y otras interacciones entre átomos

Las fuerzas interatómicas son sumamente complejas. En estas notas nos limitaremos a una discusión cualitativa pues un tratamiento detallado sería demasiado extenso. La idea básica es que en presencia de otro átomo, debido a las fuerzas electrostáticas y al Principio de Exclusión, las funciones de onda de los electrones de cada átomo se distorsionan, lo cual modifica la energía de los electrones. Si la nueva distribución tiene menor energía, será el estado preferido del sistema y la interacción interatómica será atractiva. Viceversa, si la nueva distribución lleva a un incremento de energía, la fuerza resultante será repulsiva. Los efectos de este mecanismo se manifiestan de diversas maneras en diferentes rangos de distancias, pero a grandes rasgos podemos reconocer tres clases fundamentales de fuerzas interatómicas. Cuando la distancia es muy pequeña todos los átomos se repelen mutuamente. A distancias intermedias prevalecen fuerzas que

dan lugar a enlaces químicos, de resultados de los cuales los átomos se mantienen unidos formando moléculas o estructuras más complejas como cristales. Finalmente, a distancias muy grandes, todos los átomos y moléculas se atraen débilmente. Todas estas interacciones se explican mediante la Mecánica Cuántica de los electrones atómicos. Examinaremos a continuación estas tres clases de interacciones.

La impenetrabilidad de la materia

La repulsión universal que se presenta para distancias muy pequeñas es responsable de que los átomos ocupen un volumen definido, lo que se manifiesta en la escala macroscópica como la *incompresibilidad* de la materia condensada. Esto se debe a que la fuerza repulsiva crece muy rápidamente a medida que disminuye la distancia entre los centros de los átomos, de modo tal que a muchos efectos prácticos los podemos considerar como esferas rígidas con un valor definido del radio. La repulsión a pequeñas distancias tiene dos causas. En primer lugar, cuando dos átomos se procuran interpenetrar, el apantallamiento de las cargas nucleares por parte de los electrones se torna menos efectivo, y la fuerza electrostática entre los núcleos los repele. La segunda causa radica en el Principio de Exclusión. No solamente no está permitido que dos electrones tengan la misma función de onda, sino que también sus funciones de onda deben ser suficientemente diferentes como para construir una función de onda total completamente antisimétrica, como se verá en el Capítulo 13. Cuando dos átomos están tan próximos que sus nubes electrónicas se comienzan a superponer, los electrones de la región de superposición, para satisfacer el Principio de Pauli se ven obligados a pasar parte del tiempo fuera de sus órbitas originales, en orbitales de mayor energía². Esto implica un aumento de la energía del sistema al disminuir la distancia interatómica, que equivale a una fuerza repulsiva.

La unión química

Cuando dos átomos se unen para formar un compuesto químico el rearrreglo afecta algunos de sus electrones externos, y deja al sistema compuesto en un estado cuya energía total es *menor* que la suma de las energías de los estados fundamentales de sus átomos por separado. De resultados de ello el compuesto adquiere una configuración *estable* cuya geometría está bien definida y en la cual los átomos están a distancias *fijas* entre sí. Tal configuración corresponde al estado de mínima energía del sistema, respecto de variaciones de los parámetros que caracterizan la configuración. La energía asociada con los enlaces químicos está típicamente comprendida entre 1 y 10 eV. El rearrreglo de los electrones atómicos puede ocurrir de varias maneras y los enlaces resultantes reciben diferentes denominaciones. El estudio detallado de la unión química es sumamente complejo. Incluso cuando hay un solo electrón en la subcapa más externa, la situación se puede complicar porque a veces los estados de la siguiente subcapa tienen una energía apenas mayor; esto ocurre con las subcapas $4s$ y $3d$ que están muy próximas en energía, y también con la $5s$ y la $4d$, así como con las subcapas $6s$, $4f$ y $5d$.

En forma elemental, la *valencia* se define como el número de átomos de hidrógeno que se combinan con (o que son desplazados por) un átomo del elemento que se está considerando. Como el hidrógeno tiene un solo electrón, cabe esperar que los átomos que tienen un solo electrón fuera de una capa cerrada fuertemente ligada sean *monovalentes*. Así ocurre efectivamente con los

² Esto es, su nueva función de onda es una combinación lineal de la función de onda original y funciones de onda de los estados electrónicos no ocupados, que son de mayor energía.

metales alcalinos (Li, Na, K, Rb, Cs y Fr). Del mismo modo cabe esperar que los átomos que necesitan *un electrón adicional* para cerrar una capa sean monovalentes. Tal es el caso de los *halógenos* (F, Cl, Br, I y At), que preceden a los gases nobles y tienen por lo tanto energías de ionización elevadas; en consecuencia no comparten con facilidad sus electrones con otros átomos, pero pueden *recibir* un electrón adicional (y sólo uno) para cerrar la capa. De manera semejante los átomos con *dos* electrones fuera de una capa cerrada (Be, Mg, Ca, Sr, Ba y Ra) o que necesitan dos electrones más para cerrar una capa (O, S, Se y Te) suelen ser *bivalentes*. Más allá de estos casos sencillos, la variedad de modos en que se pueden ordenar los electrones implica que generalmente puede existir más de una valencia, y que la valencia del elemento que estamos considerando puede depender de la naturaleza de los átomos con los cuales se combina. No obstante, se encuentra que los elementos con el mismo número de electrones externos y con estructura semejante de sus capas internas tienen propiedades químicas parecidas.

Se conocen varias clases de enlaces químicos. Si el rearrreglo produce un desplazamiento neto de parte de la nube electrónica de un átomo a otro, el enlace se denomina *heteropolar*; el caso extremo de una unión de esta clase es el enlace *iónico*, en el cual un electrón de un orbital del primer átomo pasa a ocupar un orbital del segundo átomo. Si no hay transferencia neta de carga de un átomo a otro la unión se denomina *homopolar*, o *covalente*. Discutiremos ahora brevemente los enlaces iónico y covalente.

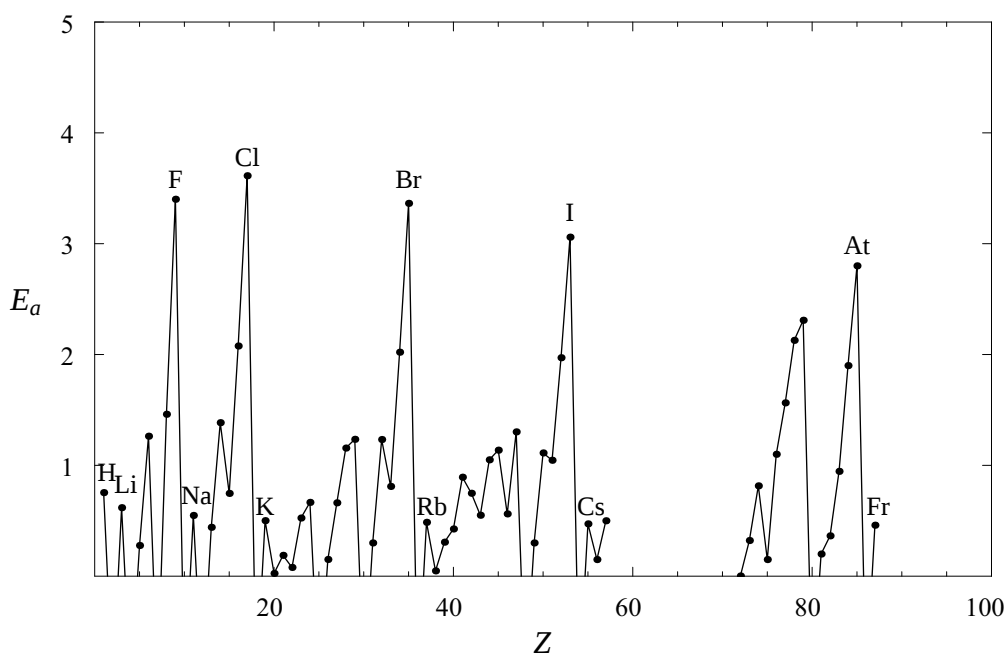


Fig. 12.8. Afinidades electrónicas de los elementos (en eV). La afinidad electrónica se define como la energía E_a liberada cuando el elemento adquiere un electrón adicional para convertirse en un ion negativo. No se dan datos para los lantánidos ($58 \leq Z \leq 71$) ni para los elementos que siguen al Fr ($Z \geq 87$). El Be, N, Mg, Mn, Zn, Cd y Hg y los gases nobles no forman iones negativos (se podría pensar que su afinidad es negativa).

El enlace iónico

La transferencia de un electrón de un átomo a otro deja al primero con una carga neta positiva y al segundo con una carga neta negativa. Los iones así obtenidos se mantienen unidos por la

atracción electrostática entre cargas de signo opuesto. Es interesante recordar que el concepto original de unión química se basó justamente en la atracción entre cargas de signo opuesto.

El enlace iónico se da en los cristales de los haluros alcalinos, por ejemplo en la sal común (NaCl). Los metales alcalinos se ionizan fácilmente, pues su energía de ionización es baja (ver Fig. 12.2). Por otra parte los halógenos tienen una fuerte afinidad para electrones adicionales (ver Fig. 12.8).

Para que una unión iónica en la que el átomo A cede un electrón al átomo B sea energéticamente posible es preciso que la variación ΔE de energía que resulta de ello sea negativa, de modo que

$$\Delta E = E_{i,A} - E_{a,B} + V_{AB} < 0 \quad (12.21)$$

donde V_{AB} es la energía potencial electrostática de la molécula formada por los iones A^+ y B^- .

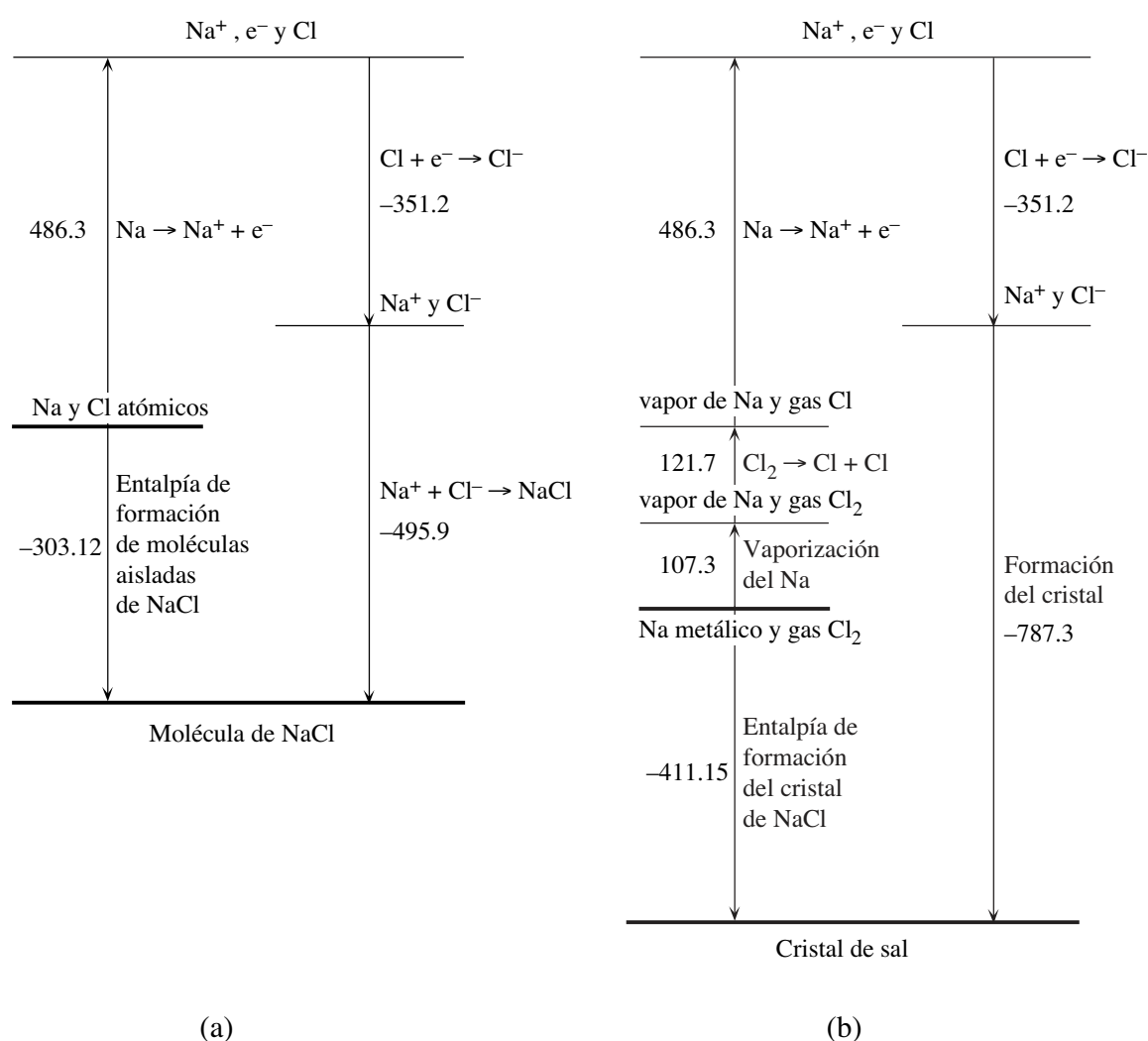


Fig. 12.9. (a) Formación de una molécula aislada de NaCl a partir de átomos aislados de Na y Cl. (b) El ciclo de Born-Fajan-Haber para la formación de un cristal de sal a partir de sodio metálico y gas cloro en condiciones standard. Las entalpías específicas para los diferentes procesos se expresan en KJ/mol.

Consideremos la sal común. Tenemos que $E_{i,\text{Na}} = 5.39$, $E_{a,\text{Cl}} = 3.613$. Para calcular V_{NaCl} necesitamos conocer la distancia d entre los iones, que es igual a la suma de los radios de los iones

Na^+ y Cl^- que valen, respectivamente, 0.99 y 1.81 Å. Luego $d = 2.8$ Å de donde obtenemos $V_{\text{NaCl}} = -5.14$ eV. Resulta entonces $\Delta E = -3.36$ eV, que equivalen³ a -303.12 KJ/mol (Fig. 12.9a). Por otra parte, de acuerdo con las tablas⁴ la entalpía de formación del NaCl es de -411.2 KJ/mol. La diferencia se debe a que se están considerando, en realidad, diferentes procesos. Nuestro cálculo se refiere a la formación de una molécula *aislada* de NaCl a partir de átomos *aislados* de Na y Cl, un proceso de interés puramente teórico dado que no ocurre en la naturaleza ni en el laboratorio. En cambio, la *entalpía de formación* que figura en las tablas se define como el calor absorbido en un proceso real, que a partir del *sodio metálico* y el *cloro gaseoso* en sus estados en las condiciones standard de temperatura y presión, lleva a un cristal de sal en las mismas condiciones de temperatura y presión. En este caso, se puede imaginar una serie de procesos ideales (que representamos en la Fig. 12.9b, donde se consignan las variaciones de entalpía de cada uno de ellos) que permiten calcular la entalpía de formación. Esta serie de procesos (que se denomina ciclo de Born-Fajan-Haber) consiste en: (1) *vaporizar* el sodio metálico, (2) *disociar* las moléculas del gas cloro para tener átomos aislados, (3) *ionizar* el sodio gaseoso para formar los cationes Na^+ y liberar electrones, (4) *combinar* los átomos de cloro con los electrones para formar los aniones Cl^- , y por último (5) *reunir* los cationes Na^+ y los aniones Cl^- para formar el cristal de sal. Al comparar nuestro cálculo anterior con el resultado del ciclo de Born-Fajan-Haber, cabe observar que la liberación de energía en la formación del cristal, *no es igual* al producto de nuestro V_{AB} por el número de pares de cationes y aniones presentes en el cristal. En efecto, en el cristal cada catión (o anión) no interactúa con un único anión (o catión), sino con todos los demás iones del cristal. En tal sentido, se puede pensar que el cristal se comporta como si fuese una única molécula. Los cálculos muestran entonces que la energía potencial electrostática por ion (-4.08 eV) en un cristal de sal es considerablemente mayor que $V_{\text{NaCl}}/2 = -2.57$ eV.

A partir de este ejemplo podemos reconocer dos características básicas del enlace iónico, que lo diferencian del enlace covalente que discutiremos a continuación. En primer lugar, el enlace iónico *no es saturable*, porque la energía de unión por ion varía (en el caso de la sal, entre -2.47 y -4.08 eV) con el número de cationes y aniones que forman el cristal. Esto se debe a que las fuerzas electrostáticas son de largo alcance. En segundo lugar, estas fuerzas son isótropas⁵, pues su intensidad es la misma en todas las direcciones. Por lo tanto la unión iónica *no es direccional*, y en consecuencia no determina la configuración geométrica del cristal. Esta última depende del tamaño de los iones, que a su vez determina de qué forma se deben disponer para minimizar la energía del cristal.

El enlace covalente

En este caso no hay transferencia neta de carga de un átomo a otro, y el rearrreglo de la distribución electrónica responsable del enlace es más sutil. En un enlace covalente participan dos electrones, uno de cada átomo, cuya distribución espacial se desplaza desde la superficie externa de los átomos hacia la región situada entre los centros atómicos. Se podría pensar que esto es con-

³ 1 eV = 96.4853 KJ/mol = 23.0605 kcal/mol.

⁴ Ver por ejemplo Handbook of Physics and Chemistry, Editor en Jefe D. E. Lide (82ª edición 2001-2002, CRC Press).

⁵ En efecto, tanto el catión como el anión son estructuras de capas cerradas y por lo tanto su distribución de carga es esféricamente simétrica.

trario al Principio de Pauli, ya que éste tiende a apartar los electrones el uno del otro. Sin embargo, el Principio de Exclusión se puede satisfacer si los spines de los dos electrones son opuestos, pues en tal caso las partes espaciales de las funciones de onda se combinan de modo que los electrones se acercan.

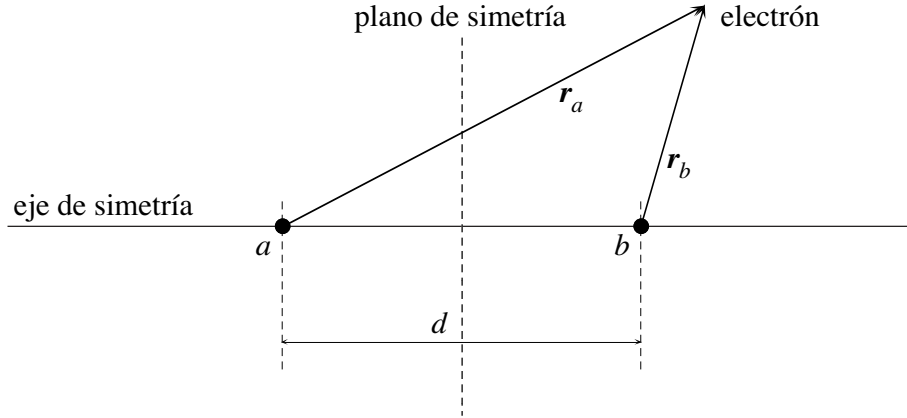


Fig. 12.10. Geometría del problema de la molécula ion hidrógeno H_2^+ .

Para ver los aspectos físicos básicos de la unión covalente, consideremos el caso más simple, la molécula ion hidrógeno H_2^+ . Este sistema consiste de dos protones a y b separados por una distancia d y un único electrón (ver Fig. 12.10). Trataremos a los núcleos como si fueran puntos fijos del espacio, ignorando sus movimientos⁶. El Hamiltoniano del electrón es

$$H = \frac{p^2}{2m} - \frac{e^2}{r_a} - \frac{e^2}{r_b} = H_a - \frac{e^2}{r_b} = H_b - \frac{e^2}{r_a} \quad (12.22)$$

donde r_a y r_b son la distancia entre el electrón y los protones a y b , respectivamente, y

$$H_a = \frac{p^2}{2m} - \frac{e^2}{r_a} \quad , \quad H_b = \frac{p^2}{2m} - \frac{e^2}{r_b} \quad (12.23)$$

son los Hamiltonianos de átomos de hidrógeno ubicados en a y b , respectivamente. Estamos interesados en el estado fundamental del electrón en el campo Coulombiano de ambos protones.

Es evidente que cuando $d \rightarrow \infty$ hay dos estados de mínima energía, equivalentes entre sí, pues el electrón puede estar ligado al protón a o al protón b . Sus funciones de onda normalizadas (ver el Capítulo 10) son, respectivamente

$$\psi_a = \psi_{1,0,0}(r_a, \theta, \varphi) = \left(\frac{1}{\pi a^3}\right)^{1/2} e^{-Zr_a/a_0} \quad , \quad \psi_b = \psi_{1,0,0}(r_b, \theta, \varphi) = \left(\frac{1}{\pi a^3}\right)^{1/2} e^{-Zr_b/a_0} \quad (12.24)$$

En ambos casos la energía del sistema está dada por

⁶ Se puede proceder así dado que su masa es enormemente mayor que la de los electrones, y por lo tanto se mueven mucho más lentamente. Por consiguiente no es preciso considerar también el movimiento de los núcleos y se los puede considerar fijos. Esta importante simplificación del problema se denomina *aproximación de Born-Oppenheimer*.

$$E_a = E_b = E_0 = -\frac{e^2}{2a_0} = -13.6 \text{ eV} \quad (12.25)$$

Supongamos, para fijar ideas, que el electrón está ligado al protón a . Si el protón b se acerca a una distancia finita del protón a , hay una probabilidad no nula que el electrón pase a moverse cerca del protón b , atravesando por efecto túnel la barrera de potencial entre ambos. Por lo tanto su función de onda ya no es ψ_a , sino una superposición de ψ_a y ψ_b . En consecuencia vamos a suponer que cuando d es finito la función de onda del electrón tiene aproximadamente la forma

$$\Psi(\mathbf{r}, t) = c_a(t)\psi_a + c_b(t)\psi_b \quad (12.26)$$

donde c_a y c_b son funciones del tiempo. La (12.26) es solamente una aproximación, pues la descripción exacta del estado del electrón requiere superponer *todas* las autofunciones de H_a y H_b ; sin embargo la inclusión de todos los estados atómicos excitados complica el cálculo sin ayudar para nada a la comprensión física del problema. Nosotros no estamos interesados aquí en obtener un resultado cuantitativo exacto, tan sólo queremos estudiar el problema en forma cualitativa para entender sus aspectos fundamentales.

Sustituyendo la (12.26) en la ecuación de Schrödinger $i\hbar\partial\Psi/\partial t = H\Psi$ resulta

$$i\hbar\dot{c}_a\psi_a + i\hbar\dot{c}_b\psi_b = c_a\left(E_0 - \frac{e^2}{r_b}\right)\psi_a + c_b\left(E_0 - \frac{e^2}{r_a}\right)\psi_b \quad (12.27)$$

Si ahora tomamos el producto escalar⁷ de ψ_a y ψ_b por la (12.27) obtenemos las siguientes ecuaciones diferenciales para los coeficientes c_a y c_b :

$$\begin{aligned} i\hbar\dot{c}_a + i\hbar\dot{c}_b\delta &= c_a(E_0 - W) + c_b(E_0\delta - X) \\ i\hbar\dot{c}_a\delta + i\hbar\dot{c}_b &= c_a(E_0\delta - X) + c_b(E_0 - W) \end{aligned} \quad (12.28)$$

Aquí

$$\delta = (\psi_a, \psi_b) = (\psi_b, \psi_a) = e^{-D}\left(1 + D + \frac{1}{3}D^2\right) \quad , \quad D = d/a_0 \quad (12.29)$$

es un número real menor que la unidad pues las funciones (12.24) son reales, y

$$\begin{aligned} W &= e^2 \int \frac{\psi_a^2}{r_b} dV = e^2 \int \frac{\psi_b^2}{r_a} dV = E_0 \left[\frac{1}{D} - e^{-2D} \left(1 + \frac{1}{D} \right) \right] \\ X &= e^2 \int \frac{\psi_a\psi_b}{r_b} dV = e^2 \int \frac{\psi_a\psi_b}{r_a} dV = E_0 [e^{-D}(1 + D)] \end{aligned} \quad (12.30)$$

son funciones de la distancia d entre los protones, siempre positivas y que disminuyen al aumentar d . Las segundas igualdades en las (12.30) se verifican fácilmente recordando que ψ_a y ψ_b tienen paridad definida por reflexiones en planos perpendiculares al eje de simetría y que pasan por a y b (en nuestro caso son funciones pares, pero igual resultado se obtendría si fuesen impares, como puede ocurrir, para estados p). La cantidad $-W(d)$ es el valor medio de la ener-

⁷ Se debe tener presente que ψ_a y ψ_b no son ortogonales.

gía potencial del electrón de un átomo en a debido a la presencia de un ion en b , o viceversa, la energía potencial del electrón de un átomo en b debido a la presencia de un ion en a . La cantidad $-X(d)$ es una suerte de “energía potencial” que depende del solapamiento de las funciones de onda ψ_a y ψ_b , y se suele llamar *integral de resonancia* o de *intercambio*.

Conviene combinar las (11.28) y obtener ecuaciones separadas para $\dot{c}_a$ y $\dot{c}_b$, para lo cual multiplicamos la segunda por δ y la restamos de la primera, y luego multiplicamos la primera por δ y la restamos de la segunda. Resulta:

$$\begin{aligned} i\hbar\dot{c}_a(1-\delta^2) &= c_a[E_0(1-\delta^2) - W + X\delta] + c_b(-X + W\delta) \\ i\hbar\dot{c}_b(1-\delta^2) &= c_b[E_0(1-\delta^2) - W + X\delta] + c_a(-X + W\delta) \end{aligned} \quad (12.31)$$

Estas ecuaciones muestran que c_a y c_b están acopladas. Podemos obtener dos ecuaciones desacopladas si introducimos las nuevas variables c_+ y c_- definidas por

$$c_+ = c_a + c_b \quad , \quad c_- = c_a - c_b \quad (12.32)$$

En efecto, sumando y restando las (12.30) obtenemos

$$i\hbar\dot{c}_+ = E_+c_+ \quad , \quad i\hbar\dot{c}_- = E_-c_- \quad (12.33)$$

donde

$$E_+(d) = E_0 - \frac{W}{1+\delta} - \frac{X}{1+\delta} \quad , \quad E_-(d) = E_0 - \frac{W}{1-\delta} + \frac{X}{1-\delta} \quad (12.34)$$

Integrando las (12.33) obtenemos

$$c_+ = A_+ e^{-i\hbar E_+ t} \quad , \quad c_- = A_- e^{-i\hbar E_- t} \quad (12.35)$$

donde A_+ y A_- son constantes de normalización.

Las ecuaciones no acopladas (12.33) nos dicen que c_+ y c_- son las amplitudes de dos *estados estacionarios* ψ_+ y ψ_- del Hamiltoniano H del sistema, correspondientes a los autovalores E_+ y E_- . En efecto, si $\Psi_+(\mathbf{r}, t) = c_+(t)\psi_+(\mathbf{r})$, se tiene que

$$i\hbar\dot{c}_+\psi_+ = i\hbar\frac{\partial\Psi_+}{\partial t} = H\Psi_+ = c_+(t)E_+\psi_+ \quad (12.36)$$

que es la primera de las (12.33), y análogamente si $\Psi_-(\mathbf{r}, t) = c_-(t)\psi_-(\mathbf{r})$ se obtiene la segunda de las (12.33). Para encontrar la función de onda ψ_+ , observemos que si el sistema se encuentra en ese estado, entonces se debe tener $c_- = 0$, o sea $c_b = c_a$. Del mismo modo, si el sistema se encuentra en el estado $\Psi_-(\mathbf{r}, t)$ se debe tener $c_+ = 0$, y entonces $c_b = -c_a$. Por lo tanto, reemplazando en la (12.26) obtenemos, respectivamente:

$$\begin{aligned} \Psi_+(\mathbf{r}, t) &= \psi_+(\mathbf{r})e^{-i\hbar E_+ t} \quad , \quad \psi_+(\mathbf{r}) = \frac{1}{\sqrt{2(1+\delta)}}(\psi_a + \psi_b) \\ \Psi_-(\mathbf{r}, t) &= \psi_-(\mathbf{r})e^{-i\hbar E_- t} \quad , \quad \psi_-(\mathbf{r}) = \frac{1}{\sqrt{2(1-\delta)}}(\psi_a - \psi_b) \end{aligned} \quad (12.37)$$

La diferencia de energía entre de los estados estacionarios ψ_+ y ψ_- es

$$\Delta = E_- - E_+ = \frac{2}{1 - \delta^2} (X - \delta W) > 0 \quad (12.38)$$

El estado estacionario *simétrico* ψ_+ tiene *menos* energía que el estado *antisimétrico* ψ_- ya que la energía potencial del electrón en ese estado es menor porque en la región de bajo potencial entre los núcleos $|\psi_+|^2 > |\psi_-|^2$, y su energía cinética es más baja pues en esa misma región $|\nabla\psi_+|^2 < |\nabla\psi_-|^2$, y es casi igual en el resto del espacio, como se puede ver en la Fig. 12.11 en la que mostramos las variaciones espaciales de ψ_+ y ψ_- a lo largo del eje de simetría.

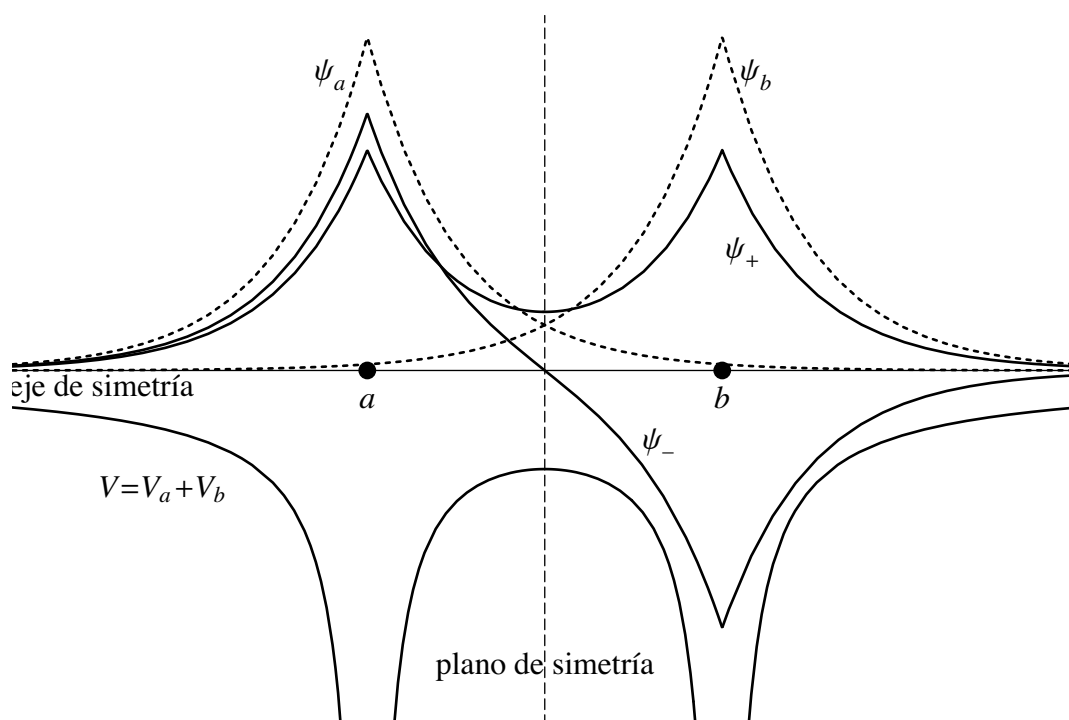


Fig. 12.11. Las funciones de onda ψ_+ y ψ_- de los estados estacionarios del H_2^+ .

Como ya dijimos, nuestro cálculo no es exacto porque en las ecs. (12.31) no consideramos los estados atómicos excitados. Sin embargo los resultados anteriores son una buena aproximación siempre y cuando d no sea muy pequeño. Asimismo, nuestro análisis permite inferir algunos valores límite exactos de la energía, aprovechando las propiedades de simetría de las soluciones (12.36), dado que tales simetrías se deben mantener para todo d . El estado simétrico ψ_+ no tiene nodos, igual que un estado s atómico. Por lo tanto, en el límite $d \rightarrow 0$ cabe esperar que ψ_+ se convierta en el estado $1s$ del ion He^+ ($Z = 2$), cuya energía es $-4E_0 = -54.4$ eV. Por otra parte ψ_- tiene un único plano nodal, que coincide con el plano de simetría del H_2^+ . Esta simetría es la misma que la de un estado atómico $2p$, y por consiguiente cabe esperar que en el límite $d \rightarrow 0$ el estado ψ_- se convierta en el estado $2p$ del ion He^+ , de energía $-E_0 = -13.6$ eV.

En la Fig. 12.12a hemos reproducido el resultado del cálculo exacto⁸ de las energías de los estados simétrico y antisimétrico del H_2^+ , en función de d . Se puede observar que los valores límite que inferimos son correctos.

⁸ El lector interesado puede encontrar una exposición detallada de los refinamientos que permiten realizar un cálculo exacto en el libro *Introduction to Quantum Mechanics* de L. Pauling y E. B. Wilson (Mc Graw-Hill 1935).

Por último, para obtener la energía total del ion, tenemos que sumar a E_+ y E_- la energía potencial debida a la interacción electrostática entre los dos protones⁹, que vale

$$V_{ab} = + \frac{Z^2 e^2}{d} \quad (12.38)$$

Se obtienen así las curvas de la Fig. 12.12.b. Se puede observar que la curva $E_+ + V_{ab}$ muestra un mínimo para $d \approx 2a_0$, que corresponde al estado ligado estable de la molécula H_2^+ . Por este motivo la función de onda simétrica ψ_+ se denomina *orbital ligante*, o *de enlace*. En cambio $E_- + V_{ab}$ no tiene mínimo, y por eso ψ_- se denomina *orbital antiligante*, o *de antienlace*.

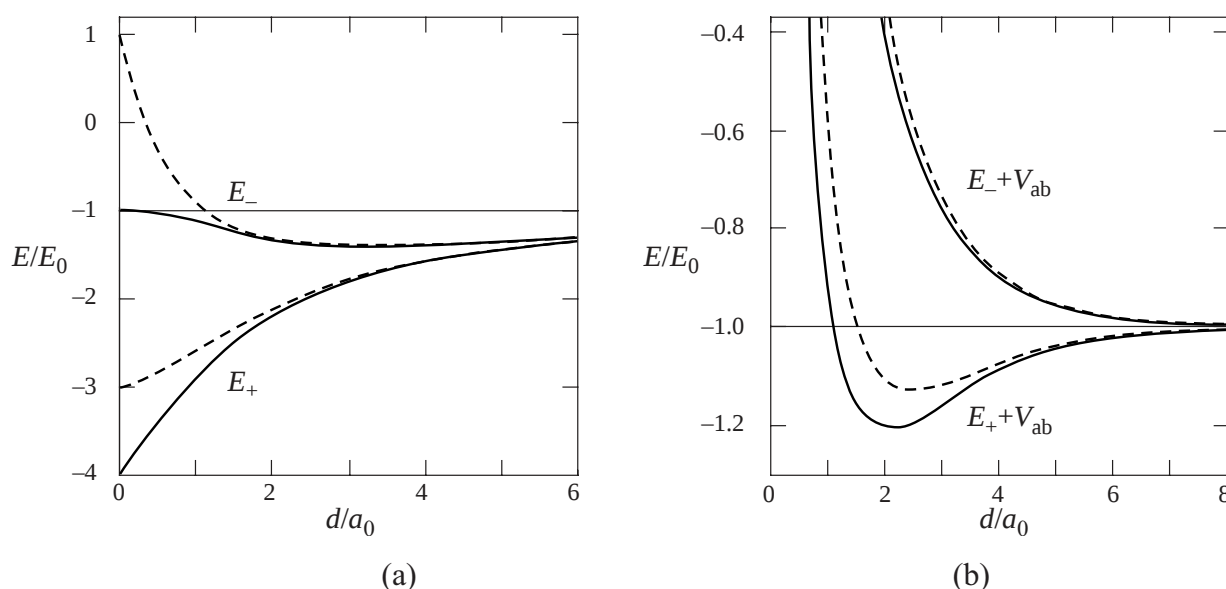


Fig. 12.12. (a) Energías de los estados estacionarios ψ_+ y ψ_- del H_2^+ . (b) Energías de la molécula H_2^+ . Las líneas llenas representan el resultado exacto. Las líneas de puntos muestran el resultado del cálculo aproximado.

El enlace de dos protones por medio de un único electrón que acabamos de estudiar muestra que el origen de la unión covalente es la disminución de la energía debida a la concentración de carga negativa en la región entre los dos núcleos. Sin embargo no es un caso típico, dado que la molécula H_2^+ está cargada. La unión de átomos neutros por medio de electrones compartidos implica que *cada uno de ellos aporta un electrón* para formar el enlace. Consideremos el ejemplo más sencillo, que es la molécula neutra H_2 . Si ignoramos la repulsión Coulombiana entre los electrones, los podemos tratar como independientes¹⁰ y entonces podemos aplicar de inmediato los resultados anteriores para el H_2^+ . El estado de menor energía se obtiene entonces poniendo los dos electrones en el orbital de enlace, lo cual está permitido por el principio de exclusión, siempre y cuando los dos electrones tengan spin opuesto. Si queremos ser precisos, de-

⁹ Esto no se debe hacer cuando $d = 0$, esto es cuando los protones se han unido para formar el núcleo del He^+ .

¹⁰ Esta aproximación de partícula independiente se usa ampliamente en la Física de Sólidos. Es una aproximación bastante razonable, y además es la única tratable en problemas en los que intervienen muchos electrones. Una discusión de la clase de errores a los que da lugar para la molécula de hidrógeno (que permite formarse una idea de cuán confiable es la aproximación) se puede encontrar en R. E. Hall, Física del Estado Sólido (Limusa 1978).

bemos calcular de nuevo la función de onda del orbital de enlace con el método del campo auto-consistente, para así tomar en cuenta la interacción de los electrones.

Por consiguiente, una característica esencial del enlace covalente es que involucra una *pareja* de electrones (uno de cada átomo). Por lo tanto el hidrógeno solo puede formar *un* enlace covalente y, en general, un átomo no puede formar más enlaces covalentes que electrones tiene fuera de capas cerradas. En virtud de esto es que se dice que el enlace covalente es *saturable*.

Un enlace covalente entre dos átomos se puede formar si en cada uno de ellos hay un electrón *no apareado*, lo que a su vez depende de la degeneración de las subcapas externas. Por ejemplo el N tiene 5 electrones externos distribuidos en los 4 orbitales de las subcapas $2s$ y $2p$. Dos de ellos ocupan la subcapa $2s$ formando un par, pero los tres restantes están desapareados en diferentes orbitales $2p$ y por este motivo el N es trivalente.

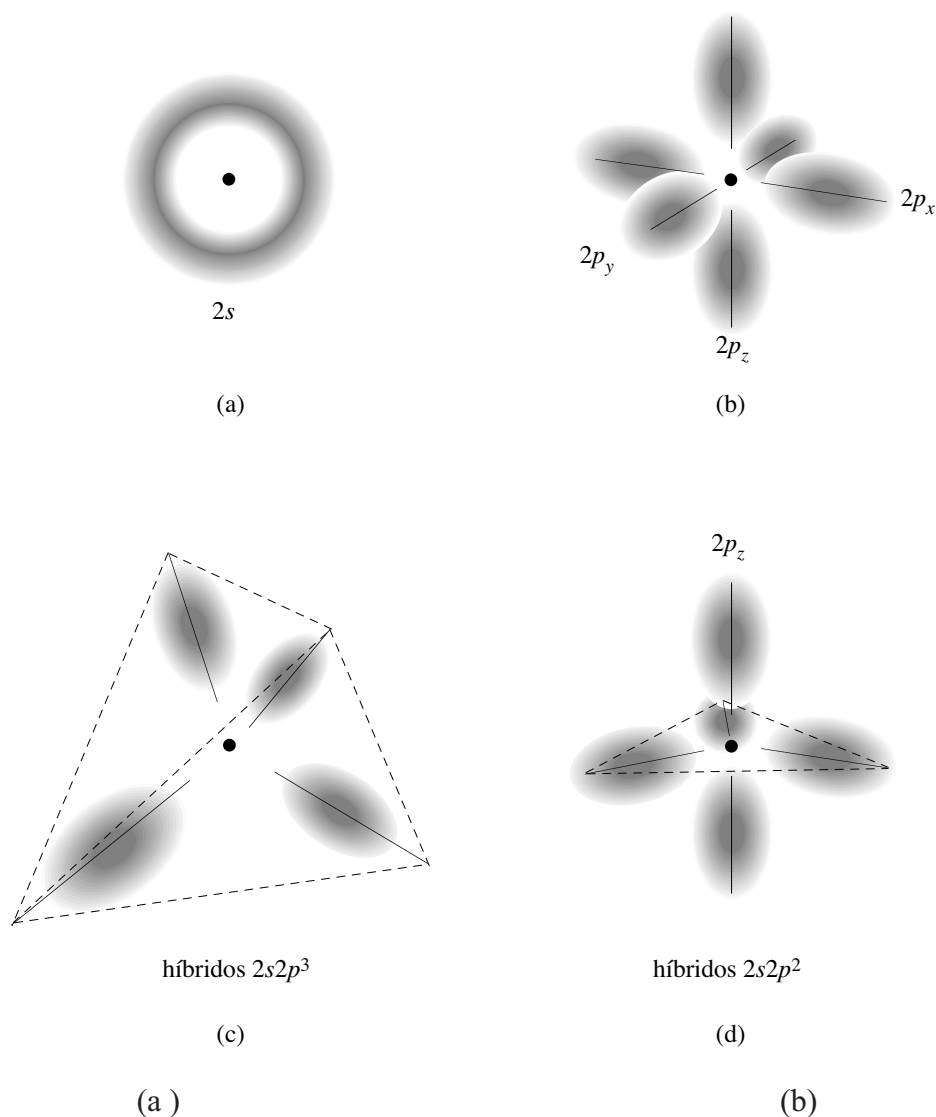


Fig. 12.13. La estructura espacial de las moléculas cuyos átomos están unidos por enlaces covalentes depende de la forma de los orbitales, que tienen lóbulos que apuntan en diferentes direcciones. Aquí hemos representado cualitativamente forma de la distribución espacial de probabilidad para: (a) el orbital $2s$ puro, (b) los tres orbitales $2p$ puros, que son mutuamente ortogonales, (c) los cuatro orbitales híbridos $2s2p^3$, que se disponen en forma de tetraedro, (2) los tres orbitales híbridos $2s2p^2$, que se disponen en el plano (x, y) a 120° entre sí y el restante orbital $2p_z$, que no se modifica.

La geometría de los compuestos químicos, con ángulos definidos entre los enlaces de las moléculas con tres o más átomos, depende de la forma de los orbitales, que tienen lóbulos que apuntan en diferentes direcciones. Si las uniones se forman a partir de orbitales p puros, los enlaces se orientan en direcciones mutuamente ortogonales, como lo hacen las funciones de onda p (Fig. 12.13b).

En realidad, sin embargo, muchos enlaces químicos reales se forman a partir de funciones de onda que son *combinaciones lineales* de orbitales p y s . Este fenómeno se llama *hibridización*, y da lugar a funciones de onda de geometría diferente. Por ejemplo, el carbono, cuyo estado fundamental tiene la configuración $1s^2 2s^2 2p^2$, puede formar hasta *cuatro* enlaces covalentes, pues a partir de la única función de onda $2s$ y las tres funciones $2p_x$, $2p_y$, $2p_z$ se forman cuatro combinaciones lineales independientes (orbitales híbridos sp^3) cada uno de los cuales está ocupado por uno de los cuatro electrones externos (Fig. 12.13c). En el carbono, esta hibridización da lugar a cuatro enlaces dirigidos hacia los vértices de un tetraedro regular. El ángulo entre dos cualesquiera de esos enlaces es de 109.5° . Esta circunstancia es la que da lugar a la estructura cristalina del diamante.

La hibridización sp^3 no es la única posible en el carbono, aunque es la más usual. Por ejemplo, se da también la hibridización sp^2 , en la cual la función de onda $2s$ se combina con dos de las funciones $2p$ para dar *tres* orbitales híbridos equivalentes cuyos lóbulos se disponen en un plano formando entre sí ángulos de 120° (Fig. 12.13d); el restante orbital $2p$ no se modifica y tiene sus lóbulos orientados perpendicularmente a dicho plano. Este tipo de hibridización es responsable de la estructura cristalina del grafito.

De lo dicho se desprende que cuando un átomo forma más de un enlace covalente, estos enlaces forman entre sí ángulos bien definidos. Por lo tanto los enlaces covalentes son *dirigidos*, además de saturables. Estas propiedades son fundamentales porque determinan la geometría de las moléculas y el tipo de estructura cristalina del compuesto.

Los enlaces iónico y covalente son dos casos extremos en lo que hace al comportamiento de los enlaces químicos, y se dan situaciones intermedias. Asimismo, las funciones de onda responsables del enlace no necesariamente están localizadas en el entorno de dos átomos vecinos. Existen *enlaces no localizados*, en los cuales las funciones de onda de los electrones involucrados se extienden sobre varios átomos.

Todas estas propiedades se describen muy bien a partir de la teoría de Hartree. Primero se determinan en forma grosera las posiciones de los centros de los átomos. Después se resuelve la ecuación de Schrödinger para cada electrón, en forma autoconsistente con la densidad de carga debida a los demás electrones. Cuando se han encontrado las funciones de onda de todos los electrones se calcula la energía de la molécula. Luego se repite el mismo procedimiento para otras posiciones de los centros atómicos hasta encontrar la que corresponde al mínimo de la energía de la molécula. La disposición espacial de centros y funciones de onda que se obtiene de este modo es la configuración geométrica de la molécula. Según sea el comportamiento de las funciones de onda de Hartree se pueden encontrar varios tipos de enlace químico. Las predicciones de la geometría de las moléculas y el carácter de los enlaces basadas teoría de Hartree son muy confiables. Sin embargo la teoría no siempre predice con exactitud las energías de unión.

Interacciones atómicas de largo alcance

Cuando los átomos y moléculas han ya establecido entre si sus posibles uniones químicas, subsisten interacciones más débiles y de largo alcance. Dichas fuerzas tienen dos orígenes.

Cuando una molécula tiene una distribución *heteropolar* de electrones (esto es cuando tiene un momento dipolar eléctrico permanente) existe una fuerza de largo alcance debida a los campos eléctricos asociados con dicha distribución. Un ejemplo es la molécula de agua, en la cual hay una transferencia neta de carga desde los átomos de hidrógeno al oxígeno. El campo eléctrico dipolar que resulta de ello atrae tanto a iones cargados positivamente como negativamente, dependiendo de la orientación de la molécula de agua. Esto explica porqué el agua es tan buen solvente para las moléculas polares y las moléculas con uniones iónicas.

Otra fuerza de largo alcance es la *fuerza de van der Waals*, que es una atracción débil que se ejerce entre toda clase de átomos o moléculas no polares. Esta fuerza es la que da cohesión a los líquidos no polares (como el aire líquido y la nafta). Tales líquidos tienen un punto de ebullición bajo, porque las energías de unión debidas a las fuerzas de van der Waals son muy pequeñas (apenas unas décimas de eV). Las fuerzas de van der Waals provienen de un sutil efecto cuántico: la existencia de campos eléctricos fluctuantes fuera de un átomo, pese a que la nube electrónica rodea al núcleo y neutraliza su carga. Estos campos fluctuantes están asociados con las posibles posiciones de los electrones que se mueven en sus orbitales.

No entraremos aquí en más detalles para no extender en demasía este Capítulo, y remitimos al lector interesado a la abundante literatura que existe sobre estos temas. Pero debe quedar claro que el modelo atómico de capas, junto con el Principio de Exclusión, permiten explicar satisfactoriamente la estructura atómica de los elementos y sus propiedades.

13. PARTÍCULAS IDÉNTICAS

En el Capítulo 12 introdujimos el Principio de Exclusión en base a la evidencia empírica, y vimos que permite explicar satisfactoriamente la configuración del estado fundamental de los átomos con más de un electrón y entender las propiedades de los elementos químicos que se expresan en la Tabla Periódica. Esto es posible porque disponemos del modelo del campo autoconsistente, según el cual los electrones atómicos se mueven independientemente y por lo tanto sus estados se pueden caracterizar mediante los números cuánticos n, l, m_l, m_s . Sin embargo nuestro modelo del átomo no es exacto.

Por otra parte, si bien nuestra identificación de las configuraciones atómicas es correcta, esto *no implica* que la función de onda del conjunto de los electrones de un átomo se pueda escribir en la forma (12.6), o sea como

$$\psi_N = \psi_{k1}(\mathbf{r}_1, \sigma_1) \psi_{k2}(\mathbf{r}_2, \sigma_2) \dots \psi_{kN}(\mathbf{r}_N, \sigma_N) = \prod_{i=1}^N \psi_{ki}(\mathbf{r}_i, \sigma_i) \quad , \quad ki \equiv (n, l, m_l, m_s)_i \quad (13.1)$$

En efecto, al escribir ψ_N de esta manera no sólo estamos especificando los estados $(n, l, m_l, m_s)_i$ que están ocupados: también estamos asignando a cada estado un electrón *en particular*. Esto *no es lícito*, porque en la Mecánica Cuántica no es posible distinguir entre sí las partículas idénticas como los electrones.

Por último no está claro todavía cómo aplicar el Principio de Exclusión a otras situaciones en las que no podemos identificar los números cuánticos que corresponden a cada partícula.

Por todas estas razones la formulación del Principio de Exclusión que dimos en el Capítulo 12 no es aún satisfactoria y es necesario encontrar un planteo más general, que se pueda aplicar a todos los casos. Eso es lo que haremos en este Capítulo. Para ello es necesario primero examinar más profundamente las implicancias de la indistinguibilidad de las partículas idénticas, que es una característica fundamental de su descripción cuántica.

La indistinguibilidad y la función de onda de un sistema de varias partículas idénticas

En la Mecánica Clásica se puede siempre (al menos en línea de principio) *identificar una dada partícula y seguirla en su movimiento*, porque mientras no la perdamos de vista conserva su identidad y la podemos *distinguir* de las demás partículas aunque éstas sean *idénticas* a ella. Pero en la Mecánica Cuántica esto *no se puede hacer*, pues debido al principio de incerteza, la extensión espacial de la función de onda que describe nuestra partícula es finita, lo cual conduce inevitablemente a un solapamiento con las funciones de onda de otras partículas idénticas a ella. En esas circunstancias, cuando observamos una partícula *no podemos identificar de cuál de ellas se trata*. Este hecho produce efectos muy importantes, que no tienen un análogo clásico, pues la *indistinguibilidad* de las partículas idénticas es una característica exclusivamente cuántica. Examinaremos ahora las consecuencias de la indistinguibilidad inherente a la descripción cuántica de un sistema de partículas idénticas.

Vamos a suponer que un sistema de N partículas idénticas (por ejemplo N electrones) se describe mediante una función de onda de la forma

$$\Psi_N = \Psi(\xi_1, \xi_2, \dots, \xi_N, t) \quad (13.2)$$

Aquí con ξ_i ($i = 1, 2, \dots, N$) indicamos simbólicamente *todas* las variables que describen una *única* partícula (las tres coordenadas en el caso de partículas sin spin, las tres coordenadas más la variable de spin en el caso de los electrones y otras partículas con spin).

Observemos que esta suposición implica que las mismas variables que sirven para describir una partícula aislada, son también adecuadas para describirla cuando forma parte de un sistema con otras partículas idénticas a ella. En otras palabras, estamos dando por cierto que la partícula mantiene su individualidad cuando pasa a integrar el sistema. No es obvio que esto deba suceder siempre, pero la experiencia indica que así ocurre en muchos casos de interés. Por ejemplo, es un hecho experimental que los electrones que forman parte de un átomo se pueden siempre identificar como electrones, independientemente de cuántos haya en el átomo.

También debemos tener presente que hay una *arbitrariedad* inherente a la definición del término “partícula” tal como lo estamos usando, y que la *identidad* de dos partículas es en gran medida una cuestión de convención. Por ejemplo, los protones y neutrones se suelen considerar como dos especies (que se distinguen por diferencias de masa, carga eléctrica, momento magnético y propiedades de estabilidad). Sin embargo a veces se los describe como dos distintos estados de una misma especie, el *nucleón*; en tal caso todos los nucleones se consideran *idénticos*, y los estados protón y neutrón se caracterizan por diferentes valores de una nueva variable dinámica, el *spin isobárico* o *isospin* (lo cual requiere introducir un número cuántico adicional, que representa el autovalor de la nueva variable dinámica). En ciertos problemas, también, conviene considerar como partículas a sistemas compuestos como núcleos, átomos o moléculas. Su naturaleza compuesta da lugar a *grados de libertad internos* que se deben incluir entre las variables dinámicas que describen dichas partículas. Esta amplitud de la definición de partícula no es un obstáculo para el desarrollo de la Mecánica Cuántica de sistemas de partículas idénticas, y los principios de la teoría se pueden aplicar a cualquier especie de “partícula”, con tal que se utilicen las variables (y los números cuánticos) adecuadas a cada caso.

La interpretación de la (12.1) es que

$$\Psi_N^* \Psi_N d\xi_1 d\xi_2 \dots d\xi_N = |\Psi(\xi_1, \dots, \xi_N, t)|^2 d\xi_1 \dots d\xi_N = P(\xi_1, \dots, \xi_N, t) d\xi_1 \dots d\xi_N \quad (13.3)$$

representa la probabilidad¹ $P(\xi_1, \xi_2, \dots, \xi_N, t)$ de encontrar en un dado instante t una partícula en el intervalo $(\xi_1, \xi_1 + d\xi_1)$, otra partícula en el intervalo $(\xi_2, \xi_2 + d\xi_2)$, ..., y finalmente la última partícula en el intervalo $(\xi_N, \xi_N + d\xi_N)$.

Ahora bien, puesto que las partículas son indistinguibles, debemos tener que

$$P(\mathcal{P}(\xi_1, \xi_2, \dots, \xi_N), t) = P(\xi_1, \xi_2, \dots, \xi_N, t) \quad (13.4)$$

para *cualquier* permutación $\mathcal{P}(\xi_1, \xi_2, \dots, \xi_N)$ de los argumentos de P , dado que siendo idénticas, el *orden* en que se citan los intervalos donde se encuentran es claramente irrelevante.

La (12.22) implica que se debe cumplir

$$\Psi(\mathcal{P}(\xi_1, \xi_2, \dots, \xi_N), t) = e^{i\alpha_{\mathcal{P}}} \Psi(\xi_1, \xi_2, \dots, \xi_N, t) \quad (13.5)$$

¹ El lector debe tener presente que ahora P es una densidad de probabilidad en un espacio de $3N$ dimensiones para cada una de las $N(2s+1)$ componentes spinoriales de Ψ_N y por lo tanto no se puede visualizar en el espacio ordinario de tres dimensiones como en el caso de una única partícula.

donde $\alpha_{\mathcal{P}}$ es un número real. Mostraremos ahora que este requerimiento impone restricciones muy importantes sobre la forma de Ψ .

Para ver esto, recordemos que las permutaciones $\mathcal{P}(\xi_1, \dots, \xi_N)$ de N elementos forman un grupo, y que toda permutación se puede expresar como el producto de *transposiciones*, esto es de intercambios de dos elementos del conjunto $(\xi_1, \dots, \xi_N)$. Por lo tanto no hace falta considerar las $N!$ permutaciones posibles, sino que basta considerar las $N(N-1)/2$ transposiciones.

Consideremos una transposición $\mathcal{T}(j, k)$ que consiste en el intercambio $\xi_j \leftrightarrow \xi_k$ del j -ésimo argumento con el k -ésimo argumento de Ψ . Su efecto sobre Ψ es

$$\mathcal{T}(j, k)\Psi = \Psi(\xi_1, \xi_2, \dots, \xi_k, \dots, \xi_j, \dots, \xi_N, t) = e^{i\alpha_{jk}} \Psi(\xi_1, \xi_2, \dots, \xi_j, \dots, \xi_k, \dots, \xi_N, t) \quad (13.6)$$

Como $\mathcal{T}(j, k)\mathcal{T}(j, k) = 1$ (pues si reiteramos la transposición volvemos al orden inicial), es evidente que $e^{2i\alpha_{jk}} = 1$. Esto da *dos* posibilidades: $\lambda_{jk} = e^{i\alpha_{jk}} = 1$ y $\lambda_{jk} = e^{i\alpha_{jk}} = -1$. Ahora bien, es fácil verificar que si $\mathcal{T}(j, k)$ y $\mathcal{T}(l, m)$ son dos transposiciones cualesquiera, se cumple que

$$\lambda_{jk} = \lambda_{lm} \equiv \lambda \quad (13.7)$$

de modo que se dan solamente dos casos: (a) $\lambda = +1$, y entonces *cualquier* transposición deja invariante a Ψ , (b) $\lambda = -1$ y entonces *cualquier* transposición cambia Ψ en $-\Psi$.

En el primer caso tendremos

$$\Psi(\mathcal{P}(\xi_1, \xi_2, \dots, \xi_N), t) = \Psi(\xi_1, \xi_2, \dots, \xi_N, t) \quad (13.8)$$

y Ψ no se modifica cuando se permutan sus argumentos. Se dice entonces que Ψ es *simétrica* respecto del intercambio de sus argumentos.

En el segundo caso tendremos

$$\Psi(\mathcal{P}(\xi_1, \xi_2, \dots, \xi_N), t) = (-1)^{q_{\mathcal{P}}} \Psi(\xi_1, \xi_2, \dots, \xi_N, t) \quad (13.9)$$

donde $q_{\mathcal{P}}$ es el número de transposiciones que es necesario efectuar para realizar la permutación $\mathcal{P}$. Se dice entonces que Ψ es *antisimétrica* respecto del intercambio de sus argumentos. En particular

$$\Psi(\xi_1, \xi_2, \dots, \xi_j, \dots, \xi_k, \dots, \xi_N, t) = -\Psi(\xi_1, \xi_2, \dots, \xi_k, \dots, \xi_j, \dots, \xi_N, t) \quad (13.10)$$

y Ψ *cambia de signo* cuando se intercambian dos cualesquiera de sus argumentos ξ_j y ξ_k .

Por consiguiente resulta que la función de onda de un sistema de partículas idénticas es o simétrica o antisimétrica respecto del intercambio de sus argumentos. Esto es consecuencia de la indistinguibilidad.

Funciones de onda simétricas y antisimétricas

Vamos a ver ahora que la propiedad que la función de onda es o simétrica o antisimétrica se conserva en el tiempo. La evolución de Ψ está dada por la ecuación de Schrödinger

$$i\hbar \frac{\partial \Psi_N}{\partial t} = \mathcal{H} \Psi_N \quad (13.11)$$

donde el Hamiltoniano del sistema es la suma de las energías cinéticas de las N partículas más las energías potenciales de las fuerzas externas que actúan sobre cada una de ellas, más las energías potenciales V_{jk} de las fuerzas de interacción entre las partículas:

$$\mathcal{H} = \sum_{j=1}^N \left[\frac{\mathbf{p}_j^2}{2\mu} + V(\mathbf{r}_j) \right] + \sum_{\substack{j,k=1 \\ i>j}}^N V_{jk} \quad (13.12)$$

Está claro que $\mathcal{H}$ debe ser *simétrico* respecto del intercambio de la designación de las partículas. Esto implica que $\mathcal{H}\Psi_N$ tiene la misma simetría (o antisimetría) que Ψ_N y por consiguiente de la (13.11) se desprende que $\Psi_N(t+dt)$ tiene *la misma* simetría (o antisimetría) que Ψ_N . Continuando el proceso de integración se encuentra entonces que la simetría o antisimetría se preserva para todo tiempo.

En el caso de un estado estacionario de energía total $\mathcal{E}$ se tiene que

$$\Psi(\xi_1, \xi_2, \dots, \xi_N, t) = e^{-i\mathcal{E}t/\hbar} \Psi(\xi_1, \xi_2, \dots, \xi_N) = e^{-i\mathcal{E}t/\hbar} \psi_N \quad (13.13)$$

y vemos que el carácter de simetría por intercambio de ψ_N es igual al de la función de onda. De resultados de lo que hemos visto, surge la pregunta de si es posible que un sistema de partículas idénticas pueda tener estados de *ambas* clases de simetría por intercambio, es decir, tanto estados simétricos como antisimétricos. La respuesta es que esto *no es posible*. En efecto, sean dos sistemas A y B , ambos formados por partículas idénticas de la misma especie (por ejemplo, dos átomos). El sistema C constituido por el conjunto de A y B debe tener, por lo que vimos recién, una función de onda cuya simetría por intercambio es *definida* (o simétrica, o antisimétrica). Por otra parte, acabamos de ver que la simetría de A y B *no puede cambiar*² por el mero hecho de entrar a formar parte de C . Por lo tanto se deduce que A , B y C deben tener *la misma* simetría de intercambio. Por extensión, inferimos entonces que *todos* los sistemas de partículas de una dada especie deben tener la misma simetría.

En conclusión:

Los estados de un sistema de partículas idénticas o son todos simétricos, o son todos antisimétricos, dependiendo de la clase de partículas de que se trate.

Hasta aquí podemos llegar partiendo de los postulados de la Mecánica Cuántica no relativística de que tratan estas notas. Para avanzar más, y averiguar qué tipo de función de onda (simétrica o antisimétrica) describe un sistema de partículas de una dada clase, es preciso salir del ámbito de nuestra teoría.

Bosones y Fermiones

En el marco de la Teoría Cuántica Relativística de Campos se demuestra que existe una relación entre la simetría o antisimetría de la función de onda que describe un sistema de partículas idénticas y el *spin* de las partículas. La demostración, que no daremos aquí pues excede el ámbito de estas notas, se funda en que dicha relación es necesaria, de lo contrario no es posible formular una teoría cuántica de campos que tenga sentido, y que no lleve a predicciones absurdas.

La relación entre *spin* y simetría por intercambio es la siguiente:

² Esto es verdad independientemente de si A y B interactúan o no.

Relación entre spin y simetría por intercambio:

- los sistemas de *partículas de spin semientero* (electrones, protones, neutrones y otras más) se describen por medio de funciones de onda *antisimétricas*; tales partículas se denominan *Fermiones* pues obedecen a la estadística de Fermi-Dirac;
- los sistemas de *partículas de spin entero* (fotones y otras más) se describen por medio de funciones de onda *simétricas*; tales partículas se denominan *Bosones* porque obedecen a la estadística de Bose-Einstein.

Este resultado (que en el presente contexto tenemos que asumir como un *postulado* adicional) se denomina *relación entre spin y estadística*. Las estadísticas de Bose-Einstein y de Fermi-Dirac se tratarán en el Capítulo 15.

Sistemas de partículas independientes

Consideremos los estados estacionarios de un sistema de N partículas idénticas. Supongamos que la ecuación de Schrödinger independiente del tiempo tiene la forma

$$\mathcal{H}\psi_N = \left(\sum_{j=1}^N H_j(\mathbf{r}_j, \mathbf{p}_j) \right) \psi_N = E\psi_N \quad (13.14)$$

donde $H_i(\mathbf{r}, \mathbf{p})$ es el Hamiltoniano de *una única partícula*³. La función de onda $\psi_k(\xi)$ de un estado estacionario de *una partícula* (recordamos que ξ indica *todas* las variables que la describen, es decir las tres coordenadas *más* la variable de spin) es solución de la ecuación

$$H_i\psi_k(\xi) = E_k\psi_k(\xi) \quad (13.15)$$

y se caracteriza por ciertos números cuánticos, que en conjunto indicamos con k .

Queremos escribir la función de onda ψ_N del sistema, cuando las N partículas ocupan los estados $\psi_{k_1}, \dots, \psi_{k_N}$ con números cuánticos $k_1, \dots, k_N$. Claramente, ψ_N no puede ser de la forma (13.1), que *no tiene* simetría definida por intercambio. La forma correcta de proceder es la siguiente: sea

$$\psi(\xi_1, \xi_2 \dots \xi_N) = \psi_{k_1}(\xi_1)\psi_{k_2}(\xi_2) \dots \psi_{k_N}(\xi_N) \quad (13.16)$$

Sea ahora $\mathcal{P}(\xi_1, \xi_2, \dots, \xi_N)$ una permutación de los argumentos de $\psi(\xi_1, \xi_2 \dots \xi_N)$. Entonces

$$\psi_{Ns} = \frac{1}{\sqrt{N_s}} \sum_{\text{toda } \mathcal{P}} \psi(\mathcal{P}(\xi_1, \xi_2, \dots, \xi_N)) \quad (13.17)$$

donde la suma abarca *todas* las $N!$ permutaciones de los argumentos de ψ , es *simétrica* por el intercambio de dos cualesquiera de sus argumentos. Por lo tanto (a menos del factor $e^{-iEt/\hbar}$) la (13.17) es la función de onda correcta para *Bosones*.

En cambio

³ Esto significa que consideramos que cada partícula se mueve *independientemente* de las demás, esto es, que no interactúa con ellas, o bien que esas interacciones se pueden describir por medio de un potencial efectivo que depende solamente de las variables de la partícula (como ocurre para el modelo del átomo que tratamos en el Capítulo 12).

$$\psi_{Na} = \frac{1}{\sqrt{\mathcal{N}_a}} \sum_{\text{toda } \mathcal{P}} (-1)^{q_{\mathcal{P}}} \psi(\mathcal{P}(\xi_1, \xi_2, \dots, \xi_N)) \quad (13.18)$$

donde $q_{\mathcal{P}}$ es el número de transposiciones que es necesario efectuar para realizar la permutación $\mathcal{P}$, es *antisimétrica* por el intercambio de dos cualesquiera de sus argumentos. Por lo tanto (siempre a menos del factor $e^{-iEt/\hbar}$) la (13.18) es la función de onda correcta para *Fermiones*. Los factores $1/\sqrt{\mathcal{N}_{s,a}}$ se introdujeron en (13.17) y (13.18) para normalizar⁴ las funciones de onda.

El principio de exclusión de Pauli

El hecho que un sistema de Fermiones, como los electrones de un átomo, se describe mediante una función de onda antisimétrica tiene importantísimas consecuencias, entre las cuales se cuenta el Principio de Exclusión.

Consideremos en efecto los estados estacionarios de un átomo con N electrones. Como vimos en el Capítulo 12, el átomo se puede describir por medio del modelo aproximado de Hartree, en el cual la ecuación de Schrödinger independiente del tiempo tiene la forma

$$\mathcal{H}\psi_N = \left(\sum_{i=1}^N H_i(\mathbf{r}_i, \mathbf{p}_i) \right) \psi_N = E\psi_N \quad (13.19)$$

donde

$$H_i(\mathbf{r}, \mathbf{p}) = \frac{\mathbf{p}^2}{2\mu} - V_i(r) \quad , \quad V_i(r) = -\frac{Ze^2}{r} + V_i'(r) \quad (13.20)$$

es el Hamiltoniano de *una única partícula* en un campo de fuerzas centrales dado por $V_i(r)$ (el potencial Coulombiano apantallado, que resulta de la carga nuclear más la que se debe a los demás electrones). La función de onda $\psi_{pi}(\mathbf{r})$ del *i-ésimo electrón* es solución de la ecuación

$$H_i\psi_{pi}(\xi) = \left[\frac{\mathbf{p}^2}{2\mu} - V_i(r) \right] \psi_{pi}(\xi) = E_{pi}\psi_{pi}(\xi) \quad (13.21)$$

y se caracteriza por los cuatro números cuánticos $pi \equiv (n, l, m_l, m_s)_{pi}$.

La *función de onda* ψ_N del átomo, cuyos N electrones ocupan los estados $\psi_{p1}, \psi_{p2}, \dots, \psi_{pN}$ con números cuánticos $p1 \equiv (n, l, m_l, m_s)_{p1}$, $p2 \equiv (n, l, m_l, m_s)_{p2}$, ..., etc. está dada por la (13.18).

Una forma equivalente de escribirla es

$$\psi_N = \frac{1}{\sqrt{\mathcal{N}}} \begin{vmatrix} \psi_{p1}(\xi_1) & \psi_{p1}(\xi_2) & \dots & \psi_{p1}(\xi_j) & \dots & \psi_{p1}(\xi_N) \\ \psi_{p2}(\xi_1) & \psi_{p2}(\xi_2) & \dots & \psi_{p2}(\xi_j) & \dots & \psi_{p2}(\xi_N) \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \psi_{pi}(\xi_1) & \psi_{pi}(\xi_2) & \dots & \psi_{pi}(\xi_j) & \dots & \psi_{pi}(\xi_N) \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \psi_{pN}(\xi_1) & \psi_{pN}(\xi_2) & \dots & \psi_{pN}(\xi_j) & \dots & \psi_{pN}(\xi_N) \end{vmatrix} \quad (13.22)$$

⁴ Si las ψ_k son ortonormales se tiene que $\mathcal{N}_{s,a} = N!$.

En efecto, al desarrollar el determinante (13.22) (que se denomina *determinante de Slater*) vemos que ψ_N es una combinación lineal de productos del tipo (13.1), en la cual figuran *todas* las $N!$ permutaciones de los N argumentos $\xi_1, \xi_2, \dots, \xi_N$; por lo tanto es una solución de la (13.19) correspondiente al autovalor

$$E = E_{p1} + E_{p2} + \dots + E_{pi} + \dots + E_{pN} \quad (13.23)$$

y además es antisimétrica frente a cualquier intercambio $\xi_j \leftrightarrow \xi_k$. Esto último es evidente, pues el intercambio de los argumentos equivale a intercambiar las correspondientes columnas del determinante (13.22).

El determinante de Slater es *idénticamente nulo* si cualquier par de partículas se encuentran en el *mismo estado*. En efecto, si dos conjuntos de números cuánticos p_i, p_k coinciden, el determinante tiene dos filas iguales y por lo tanto es nulo. Luego al postular que ψ_N es completamente antisimétrica se *garantiza* que se cumpla el Principio de Exclusión. Sin embargo, el postulado de la antisimetría de la función de onda es *más general*, ya que se puede aplicar a todo tipo de situación, incluso cuando no se puede asignar números cuánticos a las partículas.

Claramente, ψ_N se anula también cuando dos de sus *argumentos* son iguales (en tal caso dos de las columnas del determinante de Slater son iguales). Esto significa que dos electrones (o dos Fermiones, en general) no pueden estar en el mismo lugar del espacio *si tienen el mismo spin*.

Las interacciones de intercambio

La circunstancia que el spin de las partículas no intervenga en la ecuación de Schrödinger no invalida dicha ecuación ni los resultados que de ella se obtienen. En realidad la interacción eléctrica entre las partículas no depende de su spin⁵. Del punto de vista matemático esto implica que (en ausencia de campos magnéticos) el Hamiltoniano de un sistema de partículas cargadas no contiene operadores de spin, de manera que cuando opera sobre la función de onda no produce ningún efecto sobre las variables de spin. Por ese motivo, cada una de las componentes de la función de onda (ver la ec. (11.80)) satisface la ecuación de Schrödinger por separado de la otra. Por consiguiente la función de onda ψ_N del sistema se puede escribir como el producto de una función φ de las coordenadas por una función χ de las variables de spin, de la forma

$$\psi_N(\mathbf{r}_1, \sigma_1; \mathbf{r}_2, \sigma_2; \dots) = \varphi(\mathbf{r}_1, \mathbf{r}_2, \dots) \chi(\sigma_1, \sigma_2, \dots) \quad (13.24)$$

Llamaremos función de onda *orbital*, o *de las coordenadas*, a φ , y función de onda *de spin* a χ . La ecuación de Schrödinger, en la aproximación no relativística, determina únicamente la función de onda de las coordenadas dejando arbitraria la función de onda de spin. Entonces cuando no nos interesa el spin de las partículas podremos aplicar la ecuación de Schrödinger e identificar la función de onda con la función de onda orbital.

Veremos ahora una importante consecuencia de la indistinguibilidad: pese a que la interacción entre las partículas no involucra su spin, *la energía del sistema depende igualmente del spin total*. Esto se debe a que, por ser idénticas las partículas, ψ_N tiene una simetría de intercambio definida.

⁵ Esto es cierto dentro de la aproximación no relativística. Si se toman en cuenta los efectos relativísticos, la interacción entre partículas cargadas depende de sus spines.

Consideremos para simplificar un sistema de *dos* partículas idénticas. Al resolver la ecuación de Schrödinger, encontraremos una serie de niveles de energía, a cada uno de los cuales le corresponde una cierta función de onda de las coordenadas $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ que es o simétrica o antisimétrica. En efecto, debido a la identidad de las partículas, el Hamiltoniano $\mathcal{H}$ del sistema (y entonces la ecuación de Schrödinger) es invariante frente a permutaciones de las mismas. Entonces, si el nivel de energía E que estamos considerando no es degenerado, al permutar las variables $\mathbf{r}_1$ y $\mathbf{r}_2$ la autofunción $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ no puede cambiar sino por un factor constante. Reiterando la permutación es inmediato ver que ese factor no puede ser sino 1 o -1 . En el primer caso φ es simétrica por intercambio, y en el segundo es antisimétrica. Si hay degeneración, de modo que al nivel E le corresponden dos o más autofunciones $\varphi_i(\mathbf{r}_1, \mathbf{r}_2)$, se pueden siempre elegir oportunas combinaciones lineales de simetría definida. Examinemos ahora algunos casos.

Sistema de dos partículas sin spin

En este caso no hay factor de spin en la ψ , y la función de onda se reduce a la función orbital φ . Puesto que las partículas sin spin son Bosones, $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ debe ser *simétrica*. Esto significa que *no todos* los niveles de energía que se obtienen resolviendo la ecuación de Schrödinger se pueden realizar, porque aquellos que corresponden a funciones φ antisimétricas *no son admisibles*.

Para ver lo que esto implica recordemos lo visto en el Capítulo 10. Para tratar el sistema conviene separar el movimiento relativo de las partículas del movimiento del centro de masa, poniendo $\varphi(\mathbf{r}_1, \mathbf{r}_2) = \varphi_{\text{cm}}(\mathbf{r}_{\text{cm}})\varphi_{\text{rel}}(\mathbf{r})$, donde $\mathbf{r}_{\text{cm}}$ es la posición del centro de masa y $\mathbf{r} = \mathbf{r}_2 - \mathbf{r}_1$ es la posición relativa de las partículas. El momento angular $\mathbf{L}$ del movimiento *relativo* es constante del movimiento y entonces $\varphi_{\text{rel}}(\mathbf{r}) = \mathcal{R}_{n,l}(r)Y_l^m(\theta, \phi)$. Ahora bien, el intercambio de las partículas equivale a la inversión de coordenadas relativas $\mathbf{r} \rightarrow -\mathbf{r}$, y sabemos que la paridad de los armónicos esféricos está dada por $(-1)^l$. Resulta por consiguiente que:

- Un sistema de dos partículas idénticas sin spin sólo puede tener estados de momento angular orbital relativo *par*.

Sistema de dos partículas de spin 1/2

Por tratarse de Fermiones la función de onda del sistema $\psi = \varphi(\mathbf{r}_1, \mathbf{r}_2)\chi(\sigma_1, \sigma_2)$ debe ser antisimétrica ante el intercambio de las partículas. Luego si $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ es simétrica, $\chi(\sigma_1, \sigma_2)$ debe ser antisimétrica, y viceversa. Buscaremos entonces las funciones de spin de simetría definida.

La función de onda de spin de cada electrón es

$$\chi_{m_s}(\sigma) = \begin{cases} \alpha & (m_s = 1/2) \\ \beta & (m_s = -1/2) \end{cases} \quad (13.25)$$

luego $\chi(\sigma_1, \sigma_2)$ debe ser una combinación lineal de

$$\alpha_1\alpha_2, \quad \alpha_1\beta_2, \quad \beta_1\alpha_2, \quad \beta_1\beta_2 \quad (13.26)$$

Definimos ahora el spin total como

$$\mathbf{S} = \mathbf{S}_1 \times 1 + 1 \times \mathbf{S}_2 \quad (13.27)$$

donde la notación indica que el primer factor de los productos cartesianos $\mathbf{S}_1 \times 1$ y $1 \times \mathbf{S}_2$ opera sobre el primer factor, y el segundo factor opera sobre el segundo factor de los (13.26). Puesto

que $\mathbf{S}$ conmuta con $\mathcal{H}$, elegiremos las $\chi(\sigma_1, \sigma_2)$ de modo que sean autofunciones de $\mathbf{S}^2$ y S_z , correspondiente a los autovalores $\hbar^2 S(S+1)$ y $\hbar M_S$ ($M_S = S, S-1, \dots, -S$). Usando las técnicas de operadores de los Capítulos 10 y 11 es fácil verificar que S puede valer 0 o 1.

La única autofunción correspondiente a $S = 0$ es

$$\chi_{0,0} = \frac{1}{2}(\alpha_1\beta_2 - \beta_1\alpha_2) \quad (S = 0, \quad M_S = 0) \quad (13.28)$$

y es *antisimétrica* respecto del intercambio de las variables de spin de las dos partículas. Los estados con $S = 0$ se denominan *singletes*.

Las tres autofunciones correspondientes a $S = 1$ son

$$\begin{aligned} \chi_{1,1} &= \alpha_1\alpha_2 & (S = 1, \quad M_S = 1) \\ \chi_{1,0} &= \frac{1}{2}(\alpha_1\beta_2 + \beta_1\alpha_2) & (S = 1, \quad M_S = 0) \\ \chi_{1,-1} &= \beta_1\beta_2 & (S = 1, \quad M_S = -1) \end{aligned} \quad (13.29)$$

y son *simétricas* respecto del intercambio de las variables de spin de las dos partículas. Los estados con $S = 1$ se llaman *tripletes*.

Llegamos entonces a los siguientes resultados:

- Los niveles de energía que corresponden a funciones orbitales simétricas se realizan cuando el spin total del sistema es nulo (estados singlete).
- Los niveles de energía que corresponden a funciones orbitales antisimétricas se presentan cuando el spin total del sistema es igual a 1 (estados triplete).

Sistema de dos partículas de spin arbitrario s

En este caso se trata de Bosones si s es entero o de Fermiones si es semientero y claramente la simetría de la función de onda total ψ debe ser $(-1)^{2s}$. Por otra parte se puede demostrar que S puede tener los valores $2s, 2s-1, \dots, 0$ (y para cada uno de ellos $M_S = S, S-1, \dots, -S$), y que la simetría de las χ_{S,M_S} está dada por $(-1)^{2s-S}$. Resulta entonces que:

- Para garantizar que ψ tenga la simetría correcta, la simetría de la función orbital $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ debe ser $(-1)^S$.

Por ejemplo si $s = 1$ el spin total S puede valer 2 (quintupletes, con χ_{2,M_S} y $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ simétricas), 1 (tripletes, con χ_{1,M_S} y $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ antisimétricas) o 0 (singletes, con $\chi_{0,0}$ y $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ simétricas).

Correlaciones espaciales en un sistema de dos partículas idénticas

El hecho que la función de onda espacial de un sistema de partículas idénticas (Bosones o Fermiones) tiene simetría de intercambio definida hace que sus posiciones estén correlacionadas. Para ver esto, consideremos dos partículas en estados cuyas funciones orbitales son $\varphi_p(\mathbf{r})$ y $\varphi_q(\mathbf{r})$ (normalizadas pero no necesariamente ortogonales). Las funciones de onda espaciales *normalizadas* simétrica y antisimétrica del sistema son, respectivamente

$$\varphi_s = \frac{1}{\sqrt{2(1 + \delta^* \delta)}} [\varphi_p(\mathbf{r}_1)\varphi_q(\mathbf{r}_2) + \varphi_q(\mathbf{r}_1)\varphi_p(\mathbf{r}_2)] \quad (13.30)$$

$$\varphi_a = \frac{1}{\sqrt{2(1 - \delta^* \delta)}} [\varphi_p(\mathbf{r}_1)\varphi_q(\mathbf{r}_2) - \varphi_q(\mathbf{r}_1)\varphi_p(\mathbf{r}_2)] \quad (13.31)$$

donde el producto escalar $\delta = (\varphi_p, \varphi_q)$ da una medida de la falta de ortogonalidad de φ_p y φ_q . Las correspondientes densidades de probabilidad son

$$P_s(\mathbf{r}_1, \mathbf{r}_2) = \varphi_s^* \varphi_s \quad (13.32)$$

$$= \frac{1}{2(1 + \delta^* \delta)} \left\{ |\varphi_p(\mathbf{r}_1)|^2 |\varphi_q(\mathbf{r}_2)|^2 + |\varphi_q(\mathbf{r}_1)|^2 |\varphi_p(\mathbf{r}_2)|^2 + 2 \operatorname{Re}[\varphi_p^*(\mathbf{r}_1) \varphi_q(\mathbf{r}_1) \varphi_q^*(\mathbf{r}_2) \varphi_p(\mathbf{r}_2)] \right\}$$

y

$$P_a(\mathbf{r}_1, \mathbf{r}_2) = \varphi_s^* \varphi_s \quad (13.33)$$

$$= \frac{1}{2(1 + \delta^* \delta)} \left\{ |\varphi_p(\mathbf{r}_1)|^2 |\varphi_q(\mathbf{r}_2)|^2 + |\varphi_q(\mathbf{r}_1)|^2 |\varphi_p(\mathbf{r}_2)|^2 - 2 \operatorname{Re}[\varphi_p^*(\mathbf{r}_1) \varphi_q(\mathbf{r}_1) \varphi_q^*(\mathbf{r}_2) \varphi_p(\mathbf{r}_2)] \right\}$$

Estas densidades de probabilidad determinan la probabilidad de encontrar una de las partículas en el entorno de $\mathbf{r}_1$, cuando la otra partícula está en el entorno de $\mathbf{r}_2$.

Consideremos ahora el resultado (que podríamos llamar *clásico*) que se obtendría si las dos partículas fuesen *distinguibiles*. En este caso, los estados

$$\varphi_{pq}(\mathbf{r}_1, \mathbf{r}_2) = \varphi_p(\mathbf{r}_1) \varphi_q(\mathbf{r}_2) \quad \text{y} \quad \varphi_{qp}(\mathbf{r}_1, \mathbf{r}_2) = \varphi_q(\mathbf{r}_1) \varphi_p(\mathbf{r}_2) \quad (13.34)$$

se deberían considerar distintos, e igualmente probables. Luego, la probabilidad que el sistema esté en cada uno de ellos valdría 1/2. Pero cuando el sistema está en el estado φ_{pq} , la probabilidad de encontrar una partícula en el entorno de $\mathbf{r}_1$ y la otra en el entorno de $\mathbf{r}_2$ es

$$P_{pq}(\mathbf{r}_1, \mathbf{r}_2) = |\varphi_p(\mathbf{r}_1)|^2 |\varphi_q(\mathbf{r}_2)|^2 \quad (13.35)$$

Del mismo modo cuando el sistema está en el estado φ_{qp} , la probabilidad de encontrar una partícula en el entorno de $\mathbf{r}_1$ y la otra en el entorno de $\mathbf{r}_2$ es

$$P_{qp}(\mathbf{r}_1, \mathbf{r}_2) = |\varphi_q(\mathbf{r}_1)|^2 |\varphi_p(\mathbf{r}_2)|^2 \quad (13.36)$$

Por consiguiente la expresión *clásica* de la densidad de probabilidad de un sistema de dos partículas idénticas pero distinguibles, cuando una de ellas está en el estado φ_p y la otra en el estado φ_q , sin que sepamos a priori en qué estado está cada una de ellas es

$$P_D(\mathbf{r}_1, \mathbf{r}_2) = \frac{1}{2} [P_{pq}(\mathbf{r}_1, \mathbf{r}_2) + P_{qp}(\mathbf{r}_1, \mathbf{r}_2)] = \frac{1}{2} \left\{ |\varphi_p(\mathbf{r}_1)|^2 |\varphi_q(\mathbf{r}_2)|^2 + |\varphi_q(\mathbf{r}_1)|^2 |\varphi_p(\mathbf{r}_2)|^2 \right\} \quad (13.37)$$

Si comparamos esta expresión con las (13.32) y las (13.33) vemos que en general

$$P_s(\mathbf{r}_1, \mathbf{r}_2) \neq P_D(\mathbf{r}_1, \mathbf{r}_2) \neq P_a(\mathbf{r}_1, \mathbf{r}_2) \quad (13.38)$$

y que las diferencias provienen del factor de normalización $f = (1 + \delta^* \delta)^{-1}$ y de los términos de interferencia $I_{s,a}(\mathbf{r}_1, \mathbf{r}_2) = \pm 2 \operatorname{Re}[\varphi_p^*(\mathbf{r}_1) \varphi_q(\mathbf{r}_1) \varphi_q^*(\mathbf{r}_2) \varphi_p(\mathbf{r}_2)]$.

Es importante notar que estas cantidades dependen del *solapamiento* de las funciones de onda orbitales. Si $\varphi_p(\mathbf{r})$ y $\varphi_q(\mathbf{r})$ no se superponen (Fig. 13.1a) es evidente que $\delta = 0$, luego $f = 1$, y que además $I_{s,a}(\mathbf{r}_1, \mathbf{r}_2) = 0$. En este caso, se tiene que

$$P_s(\mathbf{r}_1, \mathbf{r}_2) = P_D(\mathbf{r}_1, \mathbf{r}_2) = P_a(\mathbf{r}_1, \mathbf{r}_2) \quad (13.39)$$

Por lo tanto los efectos de la indistinguibilidad se ponen de manifiesto solamente cuando las funciones de onda de las dos partículas se solapan (como en la Fig. 13.1b). Si no hay solapamiento, esto es, *en el límite clásico, las partículas se comportan como si fuesen distinguibles*.

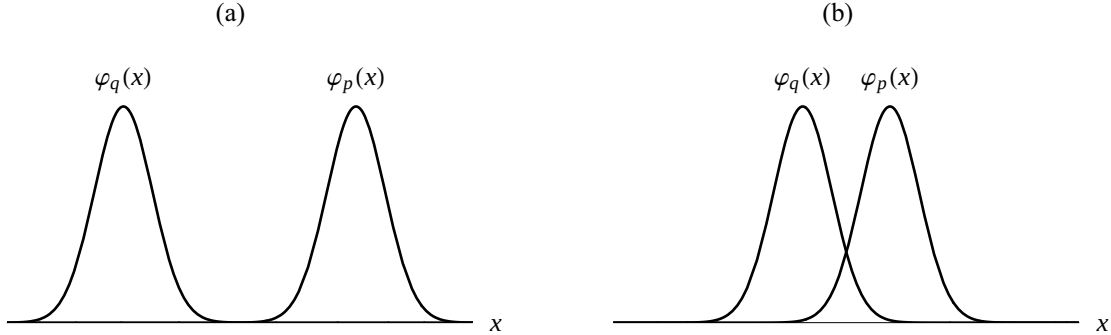


Fig. 13.1. Cuando las funciones de onda de dos partículas idénticas no se solapan, como en (a), φ_p y φ_q son ortogonales y los términos de interferencia son nulos. Por lo tanto las posiciones de las partículas no están correlacionadas, $P_s = P_D = P_a$, y todo ocurre como si fuesen distinguibles. En cambio cuando sus funciones de onda se solapan, como se indica en (b), no es posible distinguir una de otra, y al calcular P_s y P_a aparecen términos de interferencia entre φ_{pq} y φ_{qp} . Debido a ello las posiciones de las partículas están correlacionadas y $P_s \neq P_D \neq P_a$. En comparación con lo que sucedería si fuesen distinguibles, todo ocurre como si las partículas se repelieran cuando $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ es antisimétrica, y se atrajeran cuando es simétrica.

Es interesante ver cuánto vale la probabilidad de encontrar las partículas en la misma posición, es decir cuando $\mathbf{r}_1 = \mathbf{r}_2 = \mathbf{r}$. De las (13.32), (13.33) y (13.37) resulta

$$P_s(\mathbf{r}, \mathbf{r}) = \frac{2}{(1 + \delta^* \delta)} |\varphi_p(\mathbf{r}_1)|^2 |\varphi_q(\mathbf{r}_2)|^2, \quad P_D(\mathbf{r}, \mathbf{r}) = |\varphi_p(\mathbf{r}_1)|^2 |\varphi_q(\mathbf{r}_2)|^2, \quad P_a(\mathbf{r}, \mathbf{r}) = 0 \quad (13.40)$$

Luego cuando la función de onda espacial es antisimétrica, la probabilidad que las partículas se encuentren en el mismo lugar es nula. En cambio, cuando la función de onda espacial es simétrica, la probabilidad que las partículas se encuentren en el mismo lugar es *mayor* que la que se tendría si fuesen distinguibles. Todo ocurre como si las partículas se *repelieran* cuando $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ es antisimétrica, y se *atrajeran* cuando es simétrica, se entiende respecto de lo que sucedería si fuesen distinguibles.

Las densidades de probabilidad (13.32) y (13.33) son funciones de seis variables y no se pueden visualizar en el espacio ordinario. Para que el lector se pueda hacer una idea intuitiva de las correlaciones espaciales entre dos partículas, consideremos un caso unidimensional, en que $\varphi_p(x)$ y $\varphi_q(x)$ son Gaussianas (Fig. 13.2), normalizadas y de igual ancho, separadas por una distancia d (x y d se expresan en múltiplos del ancho de la Gaussiana).

$$\varphi_p = \frac{1}{\pi^{1/4}} e^{-\frac{1}{2}(x-d/2)^2}, \quad \varphi_q = \frac{1}{\pi^{1/4}} e^{-\frac{1}{2}(x+d/2)^2} \quad (13.41)$$

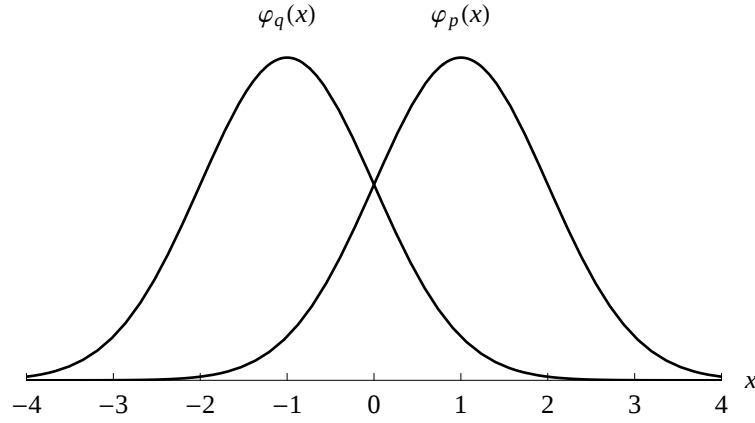


Fig. 13.2. Funciones de onda orbitales de dos partículas idénticas que se mueven en una dimensión.

En la Fig. 13.3 se pueden apreciar las densidades de probabilidad clásicas $P_{pq}(x_1, x_2)$ y $P_{qp}(x_1, x_2)$ dadas por las ecs. (13.35) y (13.36) y correspondientes a los dos estados igualmente probables φ_{pq} y φ_{qp} . En la Fig. 13.4 se pueden apreciar los resultados cuánticos (ecs. (13.32) y (13.33)) para $P_s(x_1, x_2)$ y $P_a(x_1, x_2)$, calculados para diferentes valores de d . Es fácil verificar que el solapamiento de $\varphi_p(x)$ y $\varphi_q(x)$ disminuye exponencialmente con d . Para $d = 4$ el efecto de la indistinguibilidad es insignificante, pero para valores de d próximos a la unidad, o menores, el efecto es muy importante.

La generalización de los resultados de esta Sección a sistemas de más de dos partículas idénticas no es difícil, pero lleva a ciertas complicaciones. Por lo tanto no la presentamos aquí para no extender demasiado este Capítulo⁶.

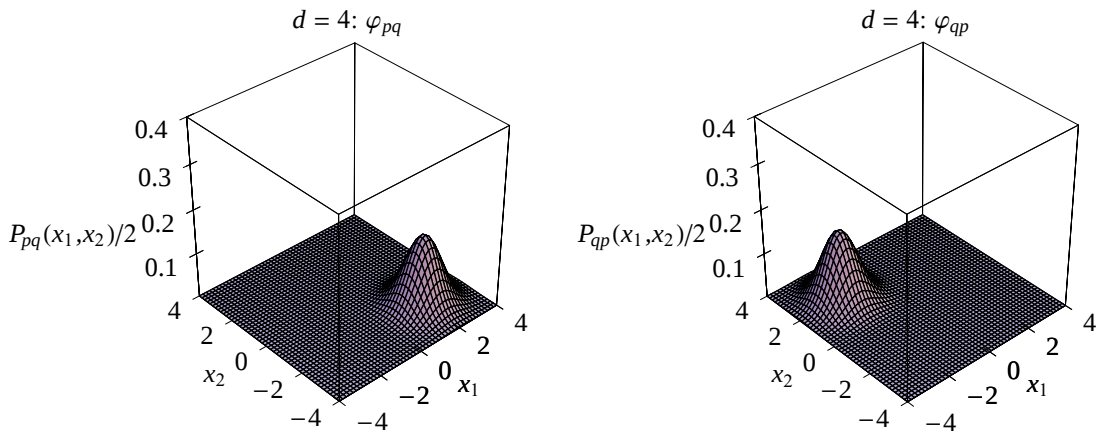


Fig. 13.3. Densidad de probabilidad para un sistema unidimensional de dos partículas idénticas pero distinguibles descritas por funciones de onda orbitales de forma Gaussiana, separadas por una distancia $d = 4$. Los diagramas corresponden a los dos estados igualmente probables φ_{pq} y φ_{qp} .

⁶ El lector interesado puede encontrar una discusión de este tema en el libro Quantum Mechanics de L. D. Landau y L. M. Lifschitz (Pergamon).

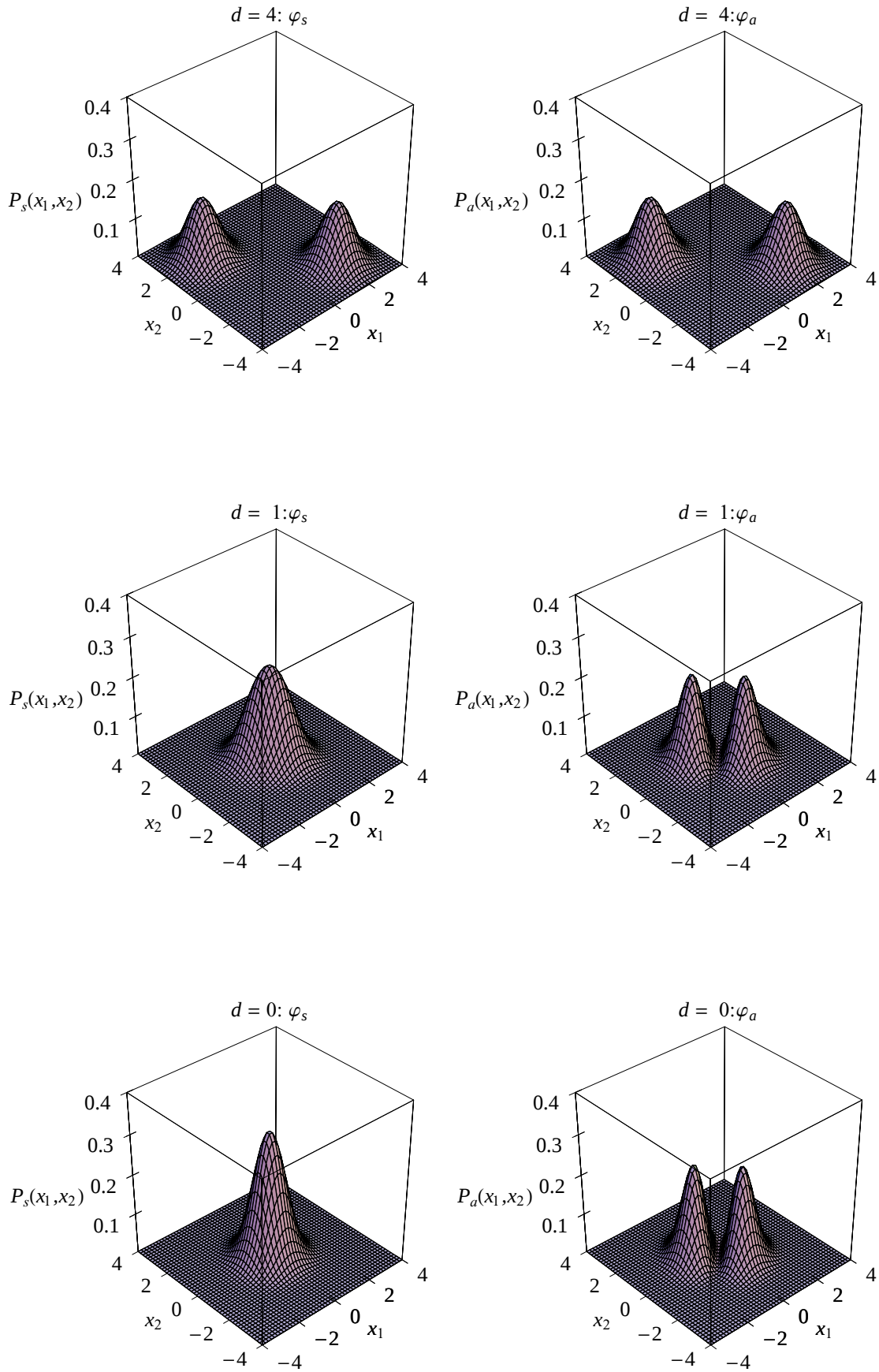


Fig. 13.4. Densidad de probabilidad para un sistema unidimensional de dos partículas descritas por funciones de onda orbitales de forma Gaussiana, separadas por diferentes distancias d . Los diagramas de la izquierda corresponden a funciones simétricas y los de la derecha a funciones antisimétricas.

El átomo de helio

Como aplicación de la teoría que hemos expuesto vamos a estudiar los estados estacionarios del átomo de helio. Observamos que el Hamiltoniano $\mathcal{H}$ del sistema no depende de las variables de spin de las partículas, puesto que hemos despreciado los efectos de spin y otras interacciones magnéticas en los $H_i(\mathbf{r}, \mathbf{p})$. Para simplificar el tratamiento vamos ignorar en un primer momento la repulsión entre los electrones, ya que no queremos aquí hacer un tratamiento exacto, y de esa forma resultará más claro el rol de la indistinguibilidad. Con esta aproximación, que llamaremos de orden 0, la energía total depende solamente de n_1 y n_2 .

La función de onda del sistema se puede escribir como el producto de una función de las variables espaciales $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ por una función $\chi_{S,M}$ de las variables de spin, de la forma

$$\psi(\mathbf{r}_1, \sigma_1, \mathbf{r}_2, \sigma_2) = \varphi_{(n,l,m_l)_1, (n,l,m_l)_2}(\mathbf{r}_1, \mathbf{r}_2) \chi_{S,M} \quad (13.42)$$

Pero ψ debe ser antisimétrica ante el intercambio $(\mathbf{r}_1, \sigma_1) \leftrightarrow (\mathbf{r}_2, \sigma_2)$, y $\chi_{S,M}$ es simétrica si $S = 1$ y antisimétrica si $S = 0$. Luego $\varphi_{(n,l,m_l)_1, (n,l,m_l)_2}$, debe ser *antisimétrica* para los estados triplete y *simétrica* para los estados singlete.

Por lo tanto, si los electrones ocupan los estados $(n,l,m_l)_1$ y $(n,l,m_l)_2$, la función de onda espacial del sistema es

$$\varphi_{\text{triplete}} = \frac{1}{\sqrt{2}} [\varphi_{(nlm_l)_1}(\mathbf{r}_1) \varphi_{(nlm_l)_2}(\mathbf{r}_2) - \varphi_{(nlm_l)_2}(\mathbf{r}_1) \varphi_{(nlm_l)_1}(\mathbf{r}_2)] \quad (13.43)$$

para los estados triplete, y

$$\varphi_{\text{singlete}} = \frac{1}{\sqrt{2}} [\varphi_{(nlm_l)_1}(\mathbf{r}_1) \varphi_{(nlm_l)_2}(\mathbf{r}_2) + \varphi_{(nlm_l)_2}(\mathbf{r}_1) \varphi_{(nlm_l)_1}(\mathbf{r}_2)] \quad (13.44)$$

para los estados singlete. En estas fórmulas

$$\varphi_{nlm_l}(\mathbf{r}) = R_{nl}(r) Y_l^{m_l}(\theta, \varphi) \quad (13.45)$$

donde $R_{nl}(r)$ es la función radial hidrogenoide correspondiente a $Z = 2$ e $Y_l^{m_l}(\theta, \varphi)$ es un armónico esférico.

El hecho que $\varphi(\mathbf{r}_1, \mathbf{r}_2)$ es antisimétrica para los estados triplete y simétrica para los estados singlete trae aparejada una profunda diferencia entre los primeros (que se suelen denominar *orthelio*) y los segundos (*parahelio*). Veremos, en efecto, que la energía de los estados depende del spin, a pesar que las variables de spin no aparecen en $\mathcal{H}$.

Antes de proseguir es oportuno hacer una breve digresión para aclarar al lector en qué consiste la notación espectroscópica, dado que se usa muy frecuentemente en la literatura. Se trata de una forma compacta de designar los *niveles* de energía de los átomos con varios electrones. Para ello se cita primero la *configuración* a la que pertenecen, y a continuación se indica el autovalor del momento angular orbital total $\mathbf{L} = \mathbf{L}_1 + \mathbf{L}_2 + \dots$ con las letras mayúsculas $S, P, D, F, \dots$ que llevan como superíndice el spin total $\mathbf{S} = \mathbf{S}_1 + \mathbf{S}_2 + \dots$ (representado por el número 1 si es singlete, 3 si es triplete, etc.). Queda definido así lo que se denomina un *término* espectral⁷, por ejemplo

⁷ Más precisamente, un término LS , pues hay otras formas de acoplar los momentos angulares orbitales y de spin de los varios electrones

3P es un término caracterizado por $S = 1$ y $L = 1$. Como veremos en seguida, debido a la repulsión Coulombiana entre los electrones los diferentes términos de una configuración tienen distintas energías. Si además se toman en cuenta los efectos relativísticos (efecto spin-órbita), los estados pertenecientes a un mismo término se separan ulteriormente en energía, de acuerdo con el autovalor J de su momento angular total $\mathbf{J} = \mathbf{L} + \mathbf{S}$ (L^2 , S^2 y $\mathbf{J}$ son constantes del movimiento para un átomo aislado que no está sometido a campos magnéticos externos). Por lo tanto, para terminar de identificar el *nivel* de energía, se agrega a la designación del término el valor de J como subíndice, por ejemplo el término 3P da lugar a tres niveles diferentes: 3P_2 , 3P_1 , y 3P_0 . A cada uno de estos niveles le pertenecen $2J + 1$ *estados*, que se identifican por el autovalor M_J ($= J, J - 1, \dots, -J$) correspondiente a la proyección de $\mathbf{J}$ sobre el eje z (arbitrario). Así 1S_0 indica que el nivel tiene spin total nulo, momento angular orbital nulo y momento angular total nulo (siempre que $J = 0$ el nivel tiene un único estado). Del mismo modo 3P_1 indica los tres estados con $L = 1$, $S = 1$, $J = 1$ y $M_J = 1, 0, -1$. En ausencia de campos externos, los $2J + 1$ estados de un dado nivel están degenerados. De esta manera, la notación espectroscópica identifica inequívocamente los niveles de energía atómicos, y las transiciones que ocurren entre ellos que dan lugar a las líneas espectrales.

Hechas estas aclaraciones, volvemos a los estados estacionarios del átomo de helio. El estado fundamental corresponde a la configuración $1s^2$, en la cual ambos electrones están en el estado de partícula individual de menor energía, esto es cuando $(n, l, m_l)_1 = (n, l, m_l)_2 = (1, 0, 0)$. En esta configuración hay un único estado, que es obviamente un *singlete*, y en la notación espectroscópica se designa como $1s^2 \ ^1S_0$ (no hay estados triplete cuando $(n, l, m_l)_1 = (n, l, m_l)_2$).

Los estados excitados de menor energía se originan en configuraciones del tipo $1s2s$, $1s2p$, $1s3s$, $1s3p$, $1s3d$, etc., y como se ve consisten en dejar un electrón del estado fundamental $1s^2 \ ^1S_0$ en el estado $1s$ y excitar el segundo electrón a un estado de mayor energía. Los términos que resultan pueden ser tanto singletes como tripletes. En esos casos, los tripletes tienen *menor* energía que los correspondientes singletes. Esto se puede ver si se calcula el efecto de la repulsión Coulombiana entre los electrones, que hasta ahora no hemos considerado. Podemos estimar fácilmente este efecto calculando el valor medio del potencial de interacción entre los electrones

$$V_i = \frac{e^2}{|\mathbf{r}_1 - \mathbf{r}_2|} \quad (13.46)$$

Si calculamos el valor esperado de V_i para los estados de un triplete obtenemos

$$\bar{V}_{i, \text{triplete}} = \iint \varphi_{\text{triplete}} V_i \varphi_{\text{triplete}} d\mathbf{r}_1 d\mathbf{r}_2 = D - I \quad (13.47)$$

donde hemos introducido la notación

$$\begin{aligned} D &= \iint \varphi_1^*(\mathbf{r}_1) \varphi_2^*(\mathbf{r}_2) V_i \varphi_1(\mathbf{r}_1) \varphi_2(\mathbf{r}_2) d\mathbf{r}_1 d\mathbf{r}_2 \\ I &= \iint \varphi_2^*(\mathbf{r}_1) \varphi_1^*(\mathbf{r}_2) V_i \varphi_1(\mathbf{r}_1) \varphi_2(\mathbf{r}_2) d\mathbf{r}_1 d\mathbf{r}_2 \end{aligned} \quad (13.48)$$

y para abreviar las fórmulas escribimos $(n, l, m_l)_1 \rightarrow 1$ y $(n, l, m_l)_2 \rightarrow 2$.

Del mismo modo para los estados del singlete que proviene de la misma configuración resulta

$$\bar{V}_{i, \text{singlete}} = \iint \varphi_{\text{singlete}} V_i \varphi_{\text{singlete}} d\mathbf{r}_1 d\mathbf{r}_2 = D + I \quad (13.49)$$

Las integrales D e I se denominan respectivamente *integral directa* e *integral de intercambio* y son positivas. La integral directa representa el valor medio clásico de la energía potencial de interacción electrostática entre dos distribuciones de carga de densidades $\rho(\mathbf{r}_1) = -e|\varphi_1(\mathbf{r}_1)|^2$ y $\rho(\mathbf{r}_2) = -e|\varphi_2(\mathbf{r}_2)|^2$ y es el resultado que se tendría si los electrones se pudieran distinguir uno de otro. Su efecto es romper la degeneración entre las configuraciones con el mismo n pero diferente l . La integral de intercambio es un efecto puramente cuántico debido a la indistinguibilidad de los electrones.

La presencia de la integral de intercambio I aumenta la energía potencial de repulsión de los estados del singlete y disminuye la de los pertenecientes al triplete respecto del valor clásico dado por la integral directa. Por tal motivo los tripletes que provienen de una configuración $1snl$ tienen siempre menor energía que los singletes de la misma configuración. Es fácil comprender porqué debe ser así. En efecto, la función de onda orbital de los tripletes es antisimétrica, por lo tanto término medio los electrones están más lejos el uno del otro que en los singletes, que tienen una función de onda orbital simétrica. Por consiguiente el efecto de repulsión es menor en los tripletes que en los singletes. La Fig. 13.5 muestra en forma esquemática como se rompe la degeneración de los diferentes términos espectrales de las configuraciones $1s2s$ y $1s2p$ del átomo de helio.

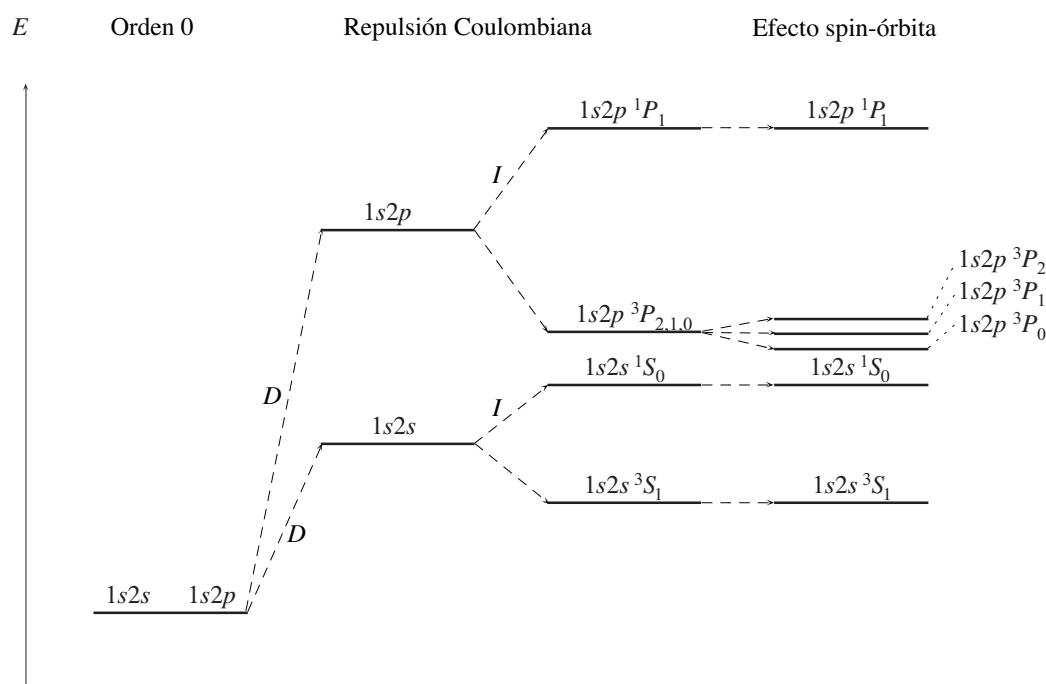


Fig. 13.5. Diagrama cualitativo que muestra como se rompe la degeneración de los niveles de las configuraciones $1s2s$ y $1s2p$ del átomo de helio. Al orden 0, esto es si ignoramos la repulsión entre los electrones, todos estos estados tienen la misma energía. El término directo D aumenta de manera diferente las energías de estas dos configuraciones. El término de intercambio I separa los términos triplete y singlete de cada configuración. Por último el efecto spin-órbita separa los diferentes niveles de cada término.

Las características que acabamos de comentar se observan también en los resultados de cálculos más exactos que el que hemos presentado aquí. En la Fig. 13.6 se muestra un esquema de los primeros términos del átomo de helio (la separación de los niveles de cada término por el efecto spin-órbita es demasiado pequeña para ser apreciada en el diagrama). Los cálculos teóricos más

precisos de los niveles de energía muestran un excelente acuerdo con los datos observacionales (la discrepancia es apenas del orden de una parte en 10^7).

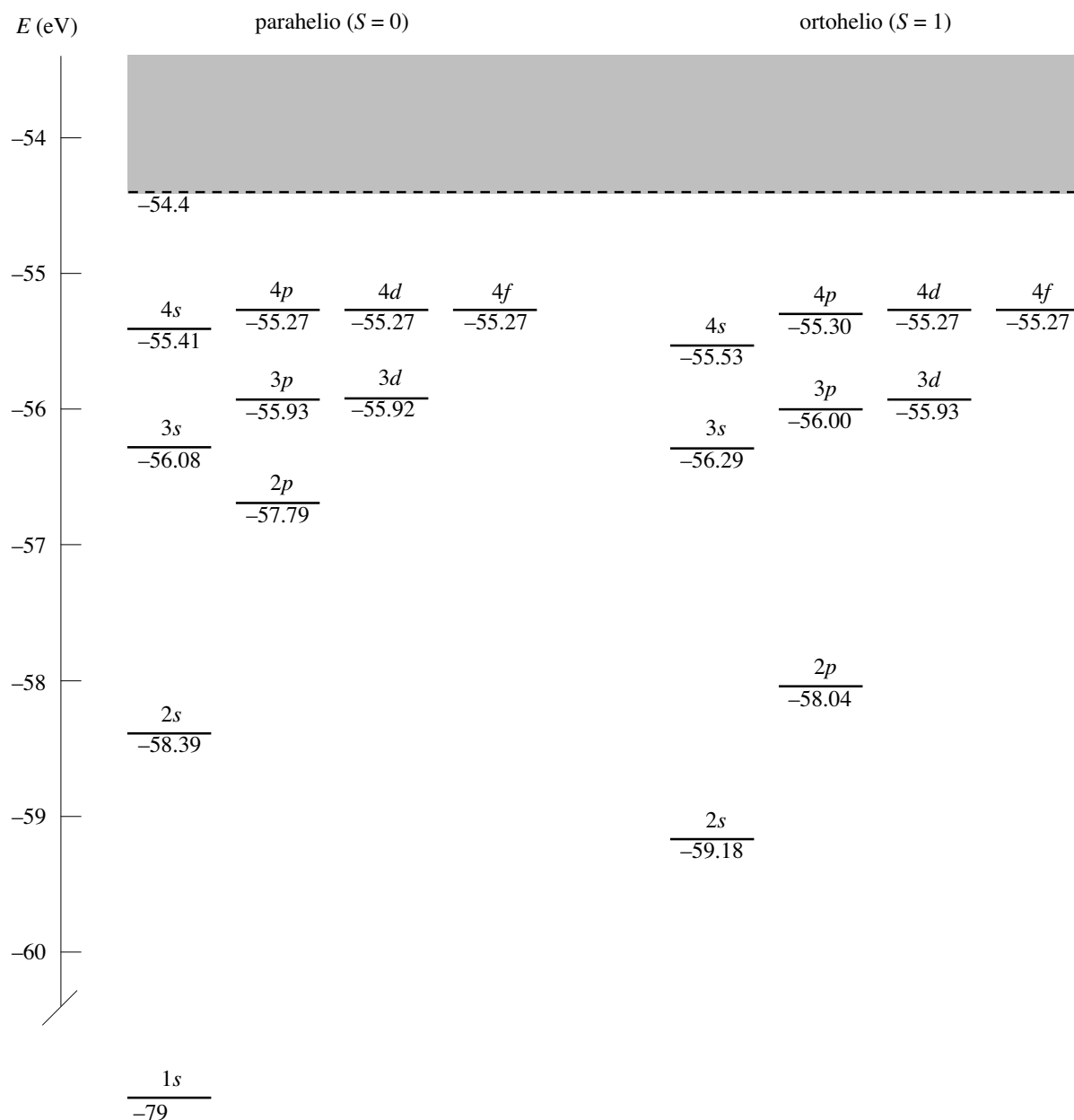


Fig. 13.6. Esquema de los primeros términos espectrales del átomo de helio, provenientes de configuraciones del tipo $1snl$. Para mayor claridad se han colocado en la parte izquierda los singletes (parahelio) y a la derecha los tripletes (ortohelio). Para que el gráfico resulte más legible omitimos escribir en las leyendas que identifican cada término el $1s$ que está presente en todas estas configuraciones, dándolo por sobreentendido. Tampoco hemos escrito el nombre del término, porque es obvio. Por ejemplo el término designado $3d$ del parahelio, cuya energía es de -55.92 eV es claramente 1D y el término del ortohelio designado $4p$ (a -55.3 eV) es 3P . La zona gris por encima de -54.4 eV corresponde al espectro continuo (el potencial de ionización del helio es de 24.6 eV, que sumados a la energía de -79 eV del estado fundamental dan precisamente -54.4 eV).

14. SEGUNDA CUANTIFICACIÓN

En este Capítulo presentaremos las técnicas para tratar sistemas cuánticos de muchas partículas idénticas (Bosones o Fermiones). Cuando se estudian sistemas de muchas partículas idénticas con interacciones arbitrarias es útil emplear un formalismo especial, denominado *segunda cuantificación*, que fue desarrollado en 1927 por Paul A. M. Dirac para los Bosones y extendido a los Fermiones por Eugene Wigner y Pascual Jordan en 1928. La importancia y utilidad de la segunda cuantificación estriba en que permite tomar en cuenta *automáticamente* en los cálculos los aspectos combinatorios que derivan de la estadística apropiada al tipo de partículas del sistema. Además facilita extender la Mecánica Cuántica no relativística a sistemas en los cuales el número de partículas no es una constante del movimiento. Dicha extensión es necesaria para describir los fenómenos que se presentan en el dominio relativístico.

No es imprescindible usar ese formalismo (que de hecho no se suele tratar en los textos introductorios) para presentar las estadísticas cuánticas, pero los argumentos empleados para deducirlas son de todos modos equivalentes a los que se usan para introducir la segunda cuantificación, por lo tanto creemos que se justifica el esfuerzo necesario para aprender esta técnica.

Segunda cuantificación de sistemas de Bosones

Trataremos primero los sistemas de Bosones. Como punto de partida supongamos conocer un sistema completo de *autofunciones ortonormales de una partícula*:

$$\psi_1(\xi), \psi_2(\xi), \dots \quad (14.1)$$

Estas autofunciones pueden corresponder, por ejemplo, a los estados estacionarios de la partícula en un campo de fuerzas externas $v(\xi)$. Destacamos que dicho campo se puede elegir *arbitrariamente*, y no tiene porqué coincidir con el campo *real* al cual están sometidas las partículas del sistema bajo estudio. La variable ξ indica el conjunto de las variables espaciales $\mathbf{r}$ y de spin σ de una partícula, y los subíndices 1, 2, ... representan los números cuánticos que identifican la autofunción.

Consideremos ahora del punto de vista puramente *formal* un sistema de n Bosones que no interactúan y que están sometidos a $v(\xi)$. Si el sistema está en un estado estacionario, cada una de las partículas que lo componen se encuentra en uno de los estados $\psi_1(\xi), \psi_2(\xi), \dots$. Sea n_i el número de partículas que están en el estado ψ_i ; n_i se denomina el *número de ocupación* o la *populación* de dicho estado y puede ser nulo (si n es finito, así ocurre para infinitos i) o tomar cualquier valor entero positivo $\leq n$. Claramente, la enunciación de todos los números de ocupación de los estados ψ_i determina el estado del sistema, que queda especificado por

$$n_1, n_2, \dots \text{ enteros no negativos tales que } \sum n_i = n \quad (14.2)$$

Los usaremos entonces para designar la función de onda del sistema, y escribiremos

$$\Psi_n = \Psi_{n; n_1, n_2, \dots} \quad (14.3)$$

Nos proponemos ahora construir un formalismo en el cual los números de ocupación jueguen el rol de *variables independientes*, en vez de las variables habituales $\xi_1, \xi_2, \dots, \xi_n$ de las partículas.

El espacio de números de ocupación para Bosones

De acuerdo a lo ya visto $\Psi_{n;n_1, n_2, \dots}$ se debe expresar como la suma simetrizada (ec. (13.17)) de los productos de todas las ψ_i cuyo número de ocupación no es nulo. Tendremos $\mathcal{N}_s = n!$ pues las ψ_i son ortonormales. Por lo tanto

$$\Psi_{n;n_1, n_2, \dots} = \frac{1}{\sqrt{n!}} \sum_{\mathcal{P}} \mathcal{P}[\psi_{p1}(\xi_1) \psi_{p2}(\xi_2) \dots \psi_{pn}(\xi_n)] \quad (14.4)$$

donde la sumatoria abarca todas las $n!$ permutaciones de los n argumentos $\xi_1, \xi_2, \dots, \xi_n$. Para que no haya confusión con los subíndices de las ψ y de los números de ocupación, tendremos cuidado con la notación. Esto es engorroso, pero lo precisamos solamente para introducir el formalismo, ya que una vez completada esa tarea quedará una notación simple, sintética y elegante. Con los subíndices $1, 2, \dots, i, \dots$ designamos los infinitos estados (14.1) de una partícula. Con los subíndices $p1, p2, \dots, pn$ indicamos los números cuánticos del estado ocupado por *cada una* de las partículas. Como en un dado estado puede haber más de una partícula, usamos los subíndices $r1, r2, \dots, rq$ ($1 \leq q \leq n$) para designar los q diferentes estados $\psi_{r1}, \psi_{r2}, \dots, \psi_{rq}$ ocupados por las partículas. Entonces en la lista de números de ocupación que identifican a Ψ , aparecen ceros para todos los n_i con $i \neq r1, r2, \dots, rq$. Como el estado rj está ocupado por n_{rj} partículas, ψ_{rj} aparece n_{rj} veces (con diferentes argumentos) en los productos $\psi_{p1}(\xi_1) \psi_{p2}(\xi_2) \dots \psi_{pn}(\xi_n)$ de la (14.4). Por lo tanto tendremos que

$$\begin{aligned} p1 &= p2 = \dots = pn_{r1} \equiv r1 \\ p(n_{r1} + 1) &= p(n_{r1} + 2) = \dots = p(n_{r1} + n_{r2}) \equiv r2 \\ &\dots \\ p(n_{r1} + n_{r2} + \dots + n_{r(q-1)} + 1) &= p(n_{r1} + n_{r2} + \dots + n_{r(q-1)} + 2) = \dots = pn \equiv rq \end{aligned} \quad (14.5)$$

La (14.4) no es todavía una forma útil de escribir $\Psi_{n;n_1, n_2, \dots}$ porque no figuran explícitamente en ella los números de ocupación $n_1, n_2, \dots$, que queremos usar como variables independientes. Además (salvo cuando todos los n_i no nulos son iguales a la unidad) en la (14.4) hay muchos términos iguales, que resultan de permutar entre sí las variables asignadas a cada uno de los q conjuntos $(\psi_{p1}, \psi_{p2}, \dots, \psi_{pn_{r1}})$, $(\psi_{p(n_{r1}+1)}, \psi_{p(n_{r1}+2)}, \dots, \psi_{p(n_{r1}+n_{r2})})$, ... etc. de autofunciones idénticas. En vista de esto escribiremos $\Psi_{n;n_1, n_2, \dots}$ de forma que no figuren términos repetidos. Para eso consideraremos todas las maneras distintas de repartir los argumentos $\xi_1, \xi_2, \dots, \xi_n$ en q grupos de $n_{r1}, n_{r2}, \dots, n_{rq}$ elementos cada uno (sin que importe el orden dentro de cada grupo), para asignar esos grupos como argumentos de los conjuntos de $n_{r1}, n_{r2}, \dots, n_{rq}$ autofunciones idénticas $\psi_{r1}, \psi_{r2}, \dots, \psi_{rq}$ de los q diferentes estados ocupados. Para saber cuántos términos distintos hay en la expresión de $\Psi_{n;n_1, n_2, \dots}$ calcularemos el número de tales reparticiones.

Recordemos algunas nociones de combinatoria. Un conjunto de r objetos elegidos (sin que importe en que orden) de un conjunto de n objetos ($r \leq n$) es una *combinación* $C(n;r)$ de n objetos tomados de a r por vez. El número $\mathcal{N}[C(n;r)]$ de combinaciones diferentes es¹

¹ El primer objeto se puede elegir de n maneras diferentes, el segundo de $n - 1$ maneras, y así siguiendo hasta llegar al r -ésimo que se puede elegir de $n - r + 1$ maneras. Como no importa el orden en que se elijan, hay que dividir el producto $n(n - 1) \dots (n - r + 1)$ por $r!$ y resulta la (14.6). $\mathcal{N}[C(n;r)]$ es igual al *coeficiente binomial*, esto es, el coeficiente de $u^r v^{n-r}$ en la expansión de $(u + v)^n$.

$$\mathcal{N}[C(n;r)] = \frac{n!}{r!(n-r)!} \equiv \binom{n}{r} \quad (14.6)$$

Después de realizar una cualquiera de las posibles $C(n;r)$, nuestro conjunto inicial queda reducido a $n-r$ elementos. Si entre estos elementos elegimos ahora s (sin que importe en que orden) tendremos $\mathcal{N}[C(n-r;s)]$ combinaciones $C(n-r;s)$ diferentes. Del conjunto residual de $n-r-s$ elementos podemos elegir t elementos (siempre sin que importe en que orden) y tendremos $\mathcal{N}[C(n-r-s;t)]$ combinaciones $C(n-r-s;t)$ distintas. De esta forma podemos continuar hasta agotar el último conjunto residual y habremos repartido los n elementos del conjunto inicial entre Q conjuntos de $r, s, t, \dots$ elementos cada uno. Designaremos con $\mathcal{R}(n;r,s,t,\dots)$ a cada una de estas *reparticiones*². Claramente, el número de reparticiones *diferentes* $\mathcal{N}[\mathcal{R}(n;r,s,t,\dots)]$ está dado por

$$\mathcal{N}[\mathcal{R}(n;r,s,t,\dots)] = \mathcal{N}[C(n;r)] \mathcal{N}[C(n-r;s)] \mathcal{N}[C(n-r-s;t)] \dots \quad (14.7)$$

Por lo tanto, usando la (14.6) obtenemos³

$$\mathcal{N}[\mathcal{R}(n;r,s,t,\dots)] = \frac{n!}{r!s!t!\dots} \equiv \binom{n}{r,s,t,\dots} \quad (14.8)$$

Volviendo a $\Psi_{n;n_1, n_2, \dots}$, la escribiremos entonces como

$$\Psi_{n;n_1, n_2, \dots} = \left(\frac{n_{r1}! n_{r2}! \dots n_{rq}!}{n!} \right)^{1/2} \sum_{\mathcal{R}} \mathcal{R}[\psi_{p1}(\xi_1) \psi_{p2}(\xi_2) \dots \psi_{pn}(\xi_n)] \quad (14.9)$$

donde la suma abarca todas las diferentes reparticiones $\mathcal{R}(n;n_{r1}, n_{r2}, \dots, n_{rq})$ de los argumentos $\xi_1, \xi_2, \dots, \xi_n$ en q conjuntos de $n_{r1}, n_{r2}, \dots, n_{rq}$ elementos cada uno (sin que importe el orden), para asignarlos a las $n_{r1}, n_{r2}, \dots, n_{rq}$ autofunciones idénticas $\psi_{r1}, \psi_{r2}, \dots, \psi_{rq}$ correspondientes a los q estados ocupados. Por lo que acabamos de ver, el número de reparticiones diferentes, esto es, el número de términos de la sumatoria en la (14.9) es

$$\mathcal{N}[\mathcal{R}(n;n_{r1}, n_{r2}, \dots, n_{rq})] = \frac{n!}{n_{r1}! n_{r2}! \dots n_{rq}!} \quad (14.10)$$

Por consiguiente $\Psi_{n;n_1, n_2, \dots}$ está correctamente normalizada, pues en virtud de la ortonormalidad de las ψ_i , al calcular $(\Psi_{n;n_1, n_2, \dots}, \Psi_{n;n_1, n_2, \dots})$ los únicos términos que dan una contribución no nula son los cuadrados de los módulos de los sumandos de la (14.9), y cada uno de ellos da una contribución igual a la unidad⁴. La (14.9) es la expresión que necesitamos: tiene la simetría correcta, y en ella figuran explícitamente los números de ocupación $n_1, n_2, \dots$.

El conjunto de las funciones (14.9) para todos los (infinitos) posibles juegos de valores de los números de ocupación, con

² También llamadas *distribuciones*.

³ El número de reparticiones diferentes $\mathcal{N}[\mathcal{R}(n;r,s,t,\dots)]$ es igual al *coeficiente multinomial*, esto es, el coeficiente de $u^r v^s w^t \dots$ en la expansión de $(u + v + w + \dots)^n$.

⁴ En efecto, es fácil ver que los términos de la sumatoria (14.9) son ortogonales.

$$n_1, n_2, \dots \text{ enteros no negativos tales que } \sum n_i = n \quad (14.11)$$

cumple las relaciones de ortonormalidad

$$(\Psi_{n;n_1', n_2', \dots}, \Psi_{n;n_1, n_2, \dots}) = \delta_{n_1', n_1} \delta_{n_2', n_2} \dots \quad (14.12)$$

y es una *base ortonormal completa* en términos de la cual podemos representar cualquier estado de n Bosones. El espacio generado por las funciones básicas $\Psi_{n;n_1, n_2, \dots}$ se llama *espacio de números de ocupación* de n partículas y se indica con $\mathcal{E}_n$.

Procediendo de esta manera podemos construir los espacios de números de ocupación para cualquier valor de n : 0, 1, 2, ..., etc., que corresponden a los estados de sistemas con 0 (el “vacío”), 1, 2, ..., etc. partículas. El espacio $\mathcal{E}$ definido como

$$\mathcal{E} = \mathcal{E}_0 \oplus \mathcal{E}_1 \oplus \mathcal{E}_2 \oplus \dots \oplus \mathcal{E}_n \oplus \dots \quad (14.13)$$

se llama *espacio de números de ocupación*. En este espacio los $n_1, n_2, \dots$ juegan el rol de variables independientes. Los diferentes $\mathcal{E}_n$ son subespacios de $\mathcal{E}$ y claramente son ortogonales entre sí. Las variables ordinarias $\xi_1, \xi_2, \dots, \xi_n$ del espacio de configuraciones no figuran explícitamente en el formalismo basado en el espacio de números de ocupación, aunque naturalmente siguen estando presentes como argumentos de las funciones básicas. Estas últimas, por supuesto, están definidas en el *espacio de configuraciones* de las n partículas.

Elementos de matriz de operadores en el espacio de números de ocupación

Veremos ahora de qué manera se expresan en el espacio $\mathcal{E}_n$ los operadores que representan los observables de un sistema de n partículas.

Consideremos un operador $f^{(1)}$ que representa una variable dinámica f de una partícula (por ejemplo el impulso $\mathbf{p}$). Su efecto sobre cualquier estado de la partícula se puede calcular a partir de sus elementos de matriz calculados con las autofunciones de nuestro sistema ortonormal completo (14.1). Estos elementos de matriz son

$$f^{(1)}_{ik} = (\psi_i(\xi), f^{(1)}\psi_k(\xi)) = \int \psi_i^*(\xi) f^{(1)}\psi_k(\xi) d\xi \quad (14.14)$$

y por lo tanto son cantidades conocidas una vez que se ha fijado el sistema (14.1). Claramente, los elementos diagonales $f^{(1)}_{ii}$ representan los valores esperados de $f^{(1)}$ en los estados $\psi_i(\xi)$ y los elementos fuera de la diagonal $f^{(1)}_{ik}$ ($i \neq k$) representan las transiciones de los estados $\psi_k(\xi)$ a los estados $\psi_i(\xi)$ inducidas por el operador $f^{(1)}$.

Sea $f_a^{(1)}$ al operador que representa f para la a -ésima partícula, esto es, el operador que actúa solamente sobre las variables ξ_a ($a = 1, 2, \dots, n$). El correspondiente operador simétrico respecto de todas las partículas del sistema (por ejemplo el impulso total $\mathbf{P}$) se define como

$$F^{(1)} = \sum_a f_a^{(1)} \quad (14.15)$$

El efecto de $F^{(1)}$ sobre cualquier estado del sistema de n partículas se puede calcular si se conoce su efecto sobre las funciones básicas $\Psi_{n;n_1, n_2, \dots}$. Determinaremos por consiguiente los elementos de matriz de $F^{(1)}$, que se definen como

$$F^{(1)}_{n;n'_1, n'_2, \dots} = (\Psi_{n;n'_1, n'_2, \dots}, F^{(1)}\Psi_{n;n_1, n_2, \dots}) \quad (14.16)$$

Cada uno de los sumandos de la (14.15) opera sobre las variables de *una* partícula, luego tiene efecto solamente sobre *una* de las funciones de los productos $\psi_{p1}(\xi_1)\psi_{p2}(\xi_2)\dots\psi_{pn}(\xi_n)$ que figuran en las Ψ_n . Consideremos la contribución al elemento de matriz (14.16) debida a uno cualquiera de los $f_a^{(1)}$, que opera sobre las variables ξ_a . Claramente los elementos diagonales de $f_a^{(1)}$ contribuyen solamente a los elementos diagonales

$$F^{(1)}_{n;n_1, n_2, \dots} = (\Psi_{n;n_1, n_2, \dots}, F^{(1)}\Psi_{n;n_1, n_2, \dots}) \quad (14.17)$$

que corresponden a que los números de ocupación $n_1, n_2, \dots$ quedan inalterados. En cambio, los elementos fuera de la diagonal de $f_a^{(1)}$, que corresponden a transiciones de la partícula desde uno de los estados $\psi_1, \psi_2, \dots$ a otro, van a contribuir solamente a aquellos elementos de la matriz (14.16) para los cuales el número de ocupación del estado inicial de la partícula disminuye en una unidad y el número de ocupación del estado final aumenta en una unidad.

Es evidente entonces que los únicos elementos de matriz (14.16) no nulos van a ser: (a) los elementos diagonales, y (b) aquellos elementos fuera de la diagonal tales que uno de los números de ocupación aumenta en 1 y otro disminuye en 1. Al hacer las cuentas hay que tener presente que las integrales (14.14) no dependen de como se designe la variable de integración, y por lo tanto $(\psi_i(\xi_a), f_a^{(1)}\psi_k(\xi_a)) = f^{(1)}_k$ para todo a .

El elemento diagonal (14.17) es el valor esperado de $F^{(1)}$ en el estado $\Psi_{n;n_1, n_2, \dots}$. Se obtiene

$$\overline{F^{(1)}} = F^{(1)}_{n_1, n_2, \dots} = \sum_i n_i f^{(1)}_i \quad (14.18)$$

Los elementos no diagonales diferentes de cero son

$$F^{(1)}_{n_i-1, n_k} = \sqrt{n_i n_k} f^{(1)}_k \quad (14.19)$$

Aquí para aligerar la notación hemos omitido escribir en la designación del elemento de matriz los números de ocupación que no cambian, dándolos por sobreentendidos. La (14.19) es la expresión del elemento de matriz de $F^{(1)}$ correspondiente a la transición de una partícula desde el estado ψ_k al estado ψ_i . Por lo tanto el número de ocupación⁵ del estado ψ_k disminuye en una unidad pasando de n_k a $n_k - 1$, el del estado ψ_i pasa de $n_i - 1$ a n_i , y los demás números de ocupación no cambian.

Dada la importancia de estos resultados, mostraremos como se llega a ellos. El cálculo es un tanto engorroso pero simple. Para calcular $\overline{F^{(1)}}$ vamos a escribir $\Psi_{n;n_1, n_2, \dots}$ en la forma

$$\Psi_{n;n_1, n_2, \dots} = \left(\frac{n_{r1}! \dots n_{rq}!}{n!} \right)^{1/2} \sum_{\mathcal{R}} \Phi_{n, \mathcal{R}} \quad , \quad \Phi_{n, \mathcal{R}} = \mathcal{R}[\psi_{p1}(\xi_1) \dots \psi_{pn}(\xi_n)] \quad (14.20)$$

donde $\mathcal{R} = \mathcal{R}(n; n_{r1}, \dots, n_{rq})$. Entonces

⁵ De la forma que hemos escrito los números de ocupación en la (14.19) se tiene que $n = \sum n_i - 1$.

$$\overline{F^{(1)}} = \frac{n_{r1}! \dots n_{rq}!}{n!} \sum_{\mathcal{R}} \sum_{\mathcal{R}'} (\Phi_{n, \mathcal{R}'}, \sum_a f_a^{(1)} \Phi_{n, \mathcal{R}}) = \frac{n_{r1}! \dots n_{rq}!}{n!} \sum_{\mathcal{R}} (\Phi_{n, \mathcal{R}}, \sum_a f_a^{(1)} \Phi_{n, \mathcal{R}}) \quad (14.21)$$

La última igualdad de la (14.21) se cumple porque

$$(\Phi_{n, \mathcal{R}'}, \sum_a f_a^{(1)} \Phi_{n, \mathcal{R}}) = 0 \quad \text{si} \quad \mathcal{R}' \neq \mathcal{R} \quad (14.22)$$

Por otra parte, es fácil verificar que para toda $\mathcal{R}(n; n_{r1}, \dots, n_{rq})$ resulta

$$(\Phi_{n, \mathcal{R}}, \sum_a f_a^{(1)} \Phi_{n, \mathcal{R}}) = \sum_{j=1}^q n_{rj} f_{qj}^{(1)qj} = \sum_i n_i f^{(1)i}_i \quad (14.23)$$

Recordando que $\mathcal{N}[\mathcal{R}(n; n_{r1}, n_{r2}, \dots, n_{rq})]$ está dado por (14.10) se llega al resultado (14.18).

Calculemos ahora el elemento de matriz no diagonal (14.19). Al usar la expresión (14.9) para escribir $\Psi_{n; n'_1, n'_2, \dots}$ debemos recordar que uno de los estados ocupados $\psi_{r1}, \psi_{r2}, \dots, \psi_{rq}$ es ψ_i (digamos, por ejemplo, que $rs = i$), además $n'_k = n_k - 1$, $n'_i = n_i$ y $n'_j = n_j$ para todo $rj \neq i, k$, luego en el coeficiente de normalización figura $(n_k - 1)!$. Del mismo modo al escribir $\Psi_{n; n_1, n_2, \dots}$ debemos poner $rt = k$ (pues el estado ψ_k está ocupado) y recordar que el número de ocupación del estado ψ_i es $n_i - 1$, luego en el coeficiente de normalización figura $(n_i - 1)!$. Por lo tanto

$$F^{(1)n_i, n_k-1}_{n_i-1, n_k} = \frac{n_{r1}! \dots n_{rq}!}{n!} \frac{1}{\sqrt{n_i n_k}} \sum_a (\Phi_{n_i, n_k-1}, f_a^{(1)} \Phi_{n_i-1, n_k}) \quad (14.24)$$

donde hemos escrito

$$\begin{aligned} \Phi_{n_i-1, n_k} &= \sum_{\mathcal{R}'} \mathcal{R}' [\psi_{p'1}(\xi_1) \psi_{p'2}(\xi_2) \dots \psi_{p'n}(\xi_n)] \\ \Phi_{n_i, n_k-1} &= \sum_{\mathcal{R}} \mathcal{R} [\psi_{p1}(\xi_1) \psi_{p2}(\xi_2) \dots \psi_{pn}(\xi_n)] \end{aligned} \quad (14.25)$$

y usamos $'$ para recordar que debido al cambio de estado de una partícula algunos de los números cuánticos de los estados ocupados de $\Psi_{n; n_1, n_2, \dots, n_i-1, \dots, n_k, \dots}$ difieren de los que designan los estados ocupados de $\Psi_{n; n_1, n_2, \dots, n_i, \dots, n_k-1, \dots}$.

Calculemos un término cualquiera de la sumatoria de la (14.24), por ejemplo la contribución de $f_1^{(1)}$, esto es $(\Phi_{n_i, n_k-1}, f_1^{(1)} \Phi_{n_i-1, n_k})$. Puesto que $f_1^{(1)}$ actúa sobre las variables ξ_1 , debemos tener $\psi_{p1'}(\xi_1) = \psi_k(\xi_1)$ y $\psi_{p1}(\xi_1) = \psi_i(\xi_1)$. Asimismo, debido a la ortogonalidad de las ψ_j , las únicas reparticiones $\mathcal{R}$ y $\mathcal{R}'$ de los argumentos de las otras ψ_j ($j \neq i, k$) de Φ_{n_i-1, n_k} y Φ_{n_i, n_k-1} que dan contribuciones no nulas son aquellas para las cuales $\psi_{p'2} = \psi_{p2}$, $\psi_{p'3} = \psi_{p3}$, $\dots$, $\psi_{p'n} = \psi_{pn}$. Estas reparticiones consisten en distribuir los $n-1$ argumentos restantes $(\xi_2, \xi_3, \dots, \xi_n)$ entre $n_i - 1$ estados ψ_i , $n_k - 1$ estados ψ_k , y los demás estados cuyos números de ocupación no han cambiado. El número de reparticiones diferentes que contribuyen es entonces

$$\begin{aligned} \mathcal{N}[\mathcal{R}(n-1; n_{r1}, \dots, n_i-1, \dots, n_k-1, \dots, n_{rq})] &= \frac{(n-1)!}{n_{r1}! \dots (n_i-1)! \dots (n_k-1)! \dots n_{rq}!} \\ &= \frac{n_i n_k (n-1)!}{n_{r1}! \dots n_i! \dots n_k! \dots n_{rq}!} \end{aligned} \quad (14.26)$$

Cada una de ellas aporta el elemento de matriz $f^{(1)l}_k$. Idéntico resultado se obtiene al calcular la contribución de $f^{(1)}_2, f^{(1)}_3, \dots$, etc. (pues el valor de $f^{(1)l}_k$ no depende de como se designe la variable de integración en la (14.14)). Puesto que hay n de estas contribuciones, obtenemos que

$$\sum_a (\Phi_{n_i, n_k-1}, f^{(1)}_a \Phi_{n_i-1, n_k}) = \frac{n_i n_k n!}{n_{r1}! \dots n_{rq}!} f^{(1)l}_k \quad (14.27)$$

Sustituyendo (14.27) en (14.24) obtenemos finalmente la (14.19).

Los resultados (14.18) y (14.19) nos muestran que los elementos de matriz de cualquier operador de una partícula del tipo $F^{(1)}$ entre los estados base de un sistema de n partículas están determinados por los números de ocupación $n_1, n_2, \dots$ y por los elementos de matriz (14.14) entre estados de una partícula, que son cantidades conocidas. Las expresiones (14.18) y (14.19) permiten entonces calcular el efecto de $F^{(1)}$ sobre cualquier estado Ψ_n de n partículas, pues en virtud de la completitud del sistema de estados base es siempre posible expresar Ψ_n como un combinación lineal de las $\Psi_{n; n_1, n_2, \dots}$.

La expresión del valor esperado de $F^{(1)}$ tiene una interpretación sencilla: cada partícula contribuye a $F^{(1)}$ con el valor esperado de $f^{(1)}$ correspondiente al estado en que se encuentra. Puesto que para cada i hay n_i partículas en el estado ψ_i , se obtiene la (14.18).

En cuanto al elemento de matriz no diagonal (14.19), que corresponde a la transición $\psi_k \rightarrow \psi_i$ de una partícula, vemos que depende solamente de los números de ocupación n_k y n_i de los estados involucrados, y es independiente de los demás números de ocupación y del número total de partículas n del sistema. Es, además, proporcional al elemento de matriz de una partícula $f^{(1)l}_k$ correspondiente a esa transición.

La probabilidad $P_{k \rightarrow i}$ que el operador $F^{(1)}$ induzca la transición $\psi_k \rightarrow \psi_i$ es igual al cuadrado del módulo del correspondiente elemento de matriz. Como $f^{(1)}$ es Hermitiano tenemos entonces

$$P_{k \rightarrow i} = n_i n_k \left| f^{(1)l}_k \right|^2 = n_i n_k f^{(1)l}_k f^{(1)k}_i \quad (14.28)$$

Recordemos que n_k es el número de partículas que ocupaban el estado ψ_k *antes* que se produjera la transición y que n_i es el número de partículas que ocupan el estado ψ_i *después* que tuvo lugar. Por lo tanto, la (14.28) nos dice que la probabilidad que se produzca una transición de una partícula del sistema a partir de un estado ψ_k es proporcional al número de partículas que ocupan dicho estado, lo cual es lógico. También es lógico que $P_{k \rightarrow i}$ no dependa de la población de los estados que no están involucrados en la transición. Pero además $P_{k \rightarrow i}$ es proporcional a n_i (que es la cantidad de partículas que quedan en el estado de llegada ψ_i de resultados de la transición: una más de las que había antes). Este resultado (cuyo origen puede parecer misterioso) muestra que *cuanto mayor es la población de un estado, tanto más probable es una transición que agregue una partícula más*. La tendencia de las partículas a acumularse⁶ es característica de los Bosones. El origen de esta tendencia se aclara si se considera la probabilidad $P_{i \rightarrow k}$ de la transición inversa $\psi_i \rightarrow \psi_k$. Evidentemente

$$P_{i \rightarrow k} = n_i n_k \left| f^{(1)k}_i \right|^2 = n_i n_k f^{(1)k}_i f^{(1)i}_k = P_{k \rightarrow i} \quad (14.29)$$

⁶ Se entiende que se trata de una acumulación, o condensación, en el espacio de números de ocupación. El lector no se debe confundir y pensar que es una condensación en el espacio físico ordinario.

de modo que

$$P_{i \rightarrow k} = P_{k \rightarrow i} \quad (14.30)$$

Este último resultado se conoce como la *condición de balance detallado*, y se puede demostrar que debe valer si la dinámica del sistema es invariante frente a la reversión temporal. Comparando (14.28) y (14.29) debe resultar claro entonces que si la probabilidad que se produzca una transición es proporcional al número de partículas que ocupan el estado de partida, necesariamente debe ser también proporcional al número de partículas que quedan en el estado de llegada. Si así no fuera, no se podría cumplir la condición de balance detallado.

Operadores de aniquilación y de creación de Bosones

Para continuar nuestro desarrollo de un formalismo cuyas variables independientes son los números de ocupación, es preciso encontrar la forma de representar las variables dinámicas del sistema por medio de operadores que actúen sobre los números de ocupación (y no sobre funciones de las variables ξ_j del espacio de configuraciones). Para eso vamos a introducir los operadores fundamentales $\hat{a}_i$ de la segunda cuantificación, que operan directamente sobre las variables $n_1, n_2, \dots$ (usamos $\hat{}$ para indicar que el operador actúa sobre los números de ocupación, y distinguirlo de los operadores que actúan sobre las variables del espacio de configuraciones). Los definiremos especificando su efecto sobre las funciones básicas $\Psi_{n;n_1, n_2, \dots}$ del espacio de números de ocupación $\mathcal{E}$ del modo siguiente

$$\hat{a}_i \Psi_{n;n_1, n_2, \dots, n_i, \dots} = \sqrt{n_i} \Psi_{n-1; n_1, n_2, \dots, n_i-1, \dots} \quad (14.31)$$

Es decir: el efecto de $\hat{a}_i$ es reducir en una unidad el número de ocupación del estado i y multiplicar la nueva función por $\sqrt{n_i}$. Por este motivo $\hat{a}_i$ se denomina *operador de aniquilación* (de una partícula en el estado i). Notar que $\hat{a}_i$ reduce en una unidad el número total de partículas, de modo que transforma funciones del subespacio $\mathcal{E}_h$ en funciones del subespacio $\mathcal{E}_{h-1}$. Podemos representar $\hat{a}_i$ por medio de una matriz, cuyos únicos elementos no nulos son

$$\hat{a}_{i n_i}^{n_i-1} = (\Psi_{n-1; n_1, \dots, n_i-1, \dots}, \hat{a}_i \Psi_{n; n_1, \dots, n_i, \dots}) = \sqrt{n_i} \quad (14.32)$$

De acuerdo con la definición (7.92) el operador $\hat{a}_i^\dagger$ adjunto de $\hat{a}_i$ cumple

$$(\hat{a}_i^\dagger \Psi_{n-1; n_1, \dots, n_i-1, \dots}, \Psi_{n; n_1, \dots, n_i, \dots}) = (\Psi_{n-1; n_1, \dots, n_i-1, \dots}, \hat{a}_i \Psi_{n; n_1, \dots, n_i, \dots}) \quad (14.33)$$

Tomando el complejo conjugado de esta expresión encontramos que

$$(\Psi_{n; n_1, \dots, n_i, \dots}, \hat{a}_i^\dagger \Psi_{n-1; n_1, \dots, n_i-1, \dots}) = (\hat{a}_i \Psi_{n; n_1, \dots, n_i, \dots}, \Psi_{n-1; n_1, \dots, n_i-1, \dots}) = \sqrt{n_i} \quad (14.34)$$

Por lo tanto se tiene que los únicos elementos de matriz no nulos de $\hat{a}_i^\dagger$ son

$$\hat{a}_{i n_i-1}^{n_i} = (\Psi_{n; n_1, \dots, n_i, \dots}, \hat{a}_i^\dagger \Psi_{n-1; n_1, \dots, n_i-1, \dots}) = \sqrt{n_i} \quad (14.35)$$

de modo que

$$\hat{a}_i^\dagger \Psi_{n-1; n_1, \dots, n_i-1, \dots} = \sqrt{n_i} \Psi_{n; n_1, \dots, n_i, \dots} \quad (14.36)$$

que con el cambio $n_i \rightarrow n_i - 1$ podemos escribir de la forma

$$\hat{a}_i^\dagger \Psi_{n; n_1, \dots, n_i, \dots} = \sqrt{n_i + 1} \Psi_{n+1; n_1, \dots, n_i+1, \dots} \quad (14.37)$$

Luego el efecto de $\hat{a}_i^\dagger$ es aumentar en una unidad la población del estado i y multiplicar la nueva función por $\sqrt{n_i + 1}$. Por lo tanto $\hat{a}_i^\dagger$ transforma funciones de $\mathcal{E}_h$ en funciones de $\mathcal{E}_{h+1}$. Por este motivo $\hat{a}_i^\dagger$ se denomina *operador de creación* (de una partícula en el estado i).

Consideremos ahora el efecto del operador $\hat{n}_i = \hat{a}_i^\dagger \hat{a}_i$ sobre $\Psi_{n; n_1, n_2, \dots}$. De (14.31) y (14.36) obtenemos de inmediato

$$\hat{a}_i^\dagger \hat{a}_i \Psi_{n; n_1, n_2, \dots, n_i, \dots} = n_i \Psi_{n; n_1, n_2, \dots, n_i, \dots} \quad (14.38)$$

Por lo tanto $\hat{n}_i$ está representado por una matriz diagonal cuyos elementos son n_i . Podemos entonces decir que $\hat{n}_i$ es el *operador número de partículas en el estado i* . Claramente, $\Psi_{n; n_1, n_2, \dots}$ es autofunción de $\hat{n}_i$ y el correspondiente autovalor es n_i . Asimismo, el operador

$$\hat{n} = \sum_i \hat{n}_i = \sum_i \hat{a}_i^\dagger \hat{a}_i \quad (14.39)$$

es el *operador número total de partículas*, y $\Psi_{n; n_1, n_2, \dots}$ es autofunción de $\hat{n}$ con autovalor n . Estos resultados pueden parecer triviales pero son de utilidad cuando se consideran estados más generales del espacio de números de ocupación (por ejemplo superposiciones de diferentes estados básicos) en los cuales los números de ocupación de los estados $\psi_1(\xi), \psi_2(\xi), \dots$ y el número total de partículas no están bien definidos.

De manera semejante a como obtuvimos la (14.38), a partir de la (14.37) y de la ecuación (que no escribimos) obtenida de la (14.31) mediante la sustitución $n_i \rightarrow n_i - 1$, es fácil verificar que

$$\hat{a}_i \hat{a}_i^\dagger \Psi_{n; n_1, n_2, \dots, n_i, \dots} = (n_i + 1) \Psi_{n; n_1, n_2, \dots, n_i, \dots} \quad (14.40)$$

Restando la (14.38) de la (14.40) obtenemos la relación de conmutación entre $\hat{a}_i$ y $\hat{a}_i^\dagger$:

$$[\hat{a}_i, \hat{a}_i^\dagger] = \hat{a}_i \hat{a}_i^\dagger - \hat{a}_i^\dagger \hat{a}_i = 1 \quad (14.41)$$

En cuanto a los operadores de creación y de aniquilación que operan sobre números de ocupación diferentes, es obvio que conmutan, de modo que podemos escribir que para todo i, j

$$[\hat{a}_i, \hat{a}_j^\dagger] = \delta_{ij} \quad , \quad [\hat{a}_i, \hat{a}_j] = [\hat{a}_i^\dagger, \hat{a}_j^\dagger] = 0 \quad (14.42)$$

Dados los operadores $\hat{a}_i^\dagger$ podemos generar todas las funciones básicas de $\mathcal{E}$. En efecto, a partir del estado de vacío $\Psi_0 \equiv \Psi_{0; 0_1, 0_2, \dots}$ (único estado de $\mathcal{E}_0$) podemos obtener las funciones básicas de una partícula (es decir $\mathcal{E}_1$) como $\hat{a}_i^\dagger \Psi_0$ ($i = 1, 2, \dots$), luego las funciones básicas de $\mathcal{E}_2$ como $\hat{a}_k^\dagger \hat{a}_i^\dagger \Psi_0$ ($i, k = 1, 2, \dots$) y así sucesivamente. Las relaciones de conmutación (14.42) garantizan que las funciones básicas de $\mathcal{E}$ tienen la simetría correcta para los Bosones.

Las relaciones de conmutación (14.42) son fundamentales. Nosotros las obtuvimos a partir de la definición (14.31) de los operadores de aniquilación. Pero se podría haber procedido a la inversa, esto es, tomar las (14.42) como punto de partida y postular que $\hat{a}_i^\dagger \hat{a}_i$ es el operador número de partículas del estado i . De esa forma se pueden deducir las (14.31) y (14.37). Es interesante comparar el presente tratamiento con el método que usamos en el Capítulo 9 para tratar el oscilador armónico simple. Se puede observar que los operadores $\hat{a}_i$ y $\hat{a}_i^\dagger$ juegan el mismo rol que los operadores de bajada y subida a y $a^\dagger$ que definimos entonces.

Operadores en el espacio de números de ocupación

Los operadores de creación y de aniquilación permiten expresar cualquier variable dinámica de un sistema de muchas partículas por medio de un operador que actúa sobre los números de ocupación. En efecto, consideremos el operador de una partícula

$$F^{(1)} = \sum_a f_a^{(1)} \quad (14.43)$$

que ya estudiamos previamente. Es inmediato verificar, usando las propiedades (14.32) y (14.35) de los $\hat{a}_i$ y $\hat{a}_i^\dagger$ que el operador

$$\hat{F}^{(1)} = \sum_{i,k} f_k^{(1)i} \hat{a}_i^\dagger \hat{a}_k \quad (14.44)$$

coincide con $F^{(1)}$. En efecto, todos los elementos de matriz de $\hat{F}^{(1)}$ entre los estados base de $\mathcal{E}$ coinciden con los elementos de matriz (14.18) y (14.19) de $F^{(1)}$.

La (14.44) es un resultado fundamental de la segunda cuantificación, porque en la (14.44) los $f_k^{(1)i}$ son simplemente cantidades (números con dimensiones) conocidas. Por lo tanto hemos logrado expresar un operador ordinario (por ahora limitado a la forma (14.43)) que actúa sobre funciones de las coordenadas bajo la forma de un operador que actúa sobre las nuevas variables del espacio de números de ocupación $\mathcal{E}$.

Se puede notar que la expresión (14.44) se parece a la expresión

$$\bar{f} = \sum_{i,k} f_k^i a_i^* a_k \quad (14.45)$$

del valor medio de una variable dinámica f en un dado estado, escrito en términos de los coeficientes a_i del desarrollo de la función de onda como superposición de estados estacionarios. De allí proviene el nombre “segunda cuantificación”.

Es fácil generalizar el resultado (14.44) a operadores de la forma

$$F^{(2)} = \sum_{a>b} f_{ab}^{(2)} \quad (14.46)$$

donde $f_{ab}^{(2)}$ indica el operador que representa una magnitud que depende de un par de partículas (por ejemplo, la energía potencial de interacción entre dos partículas cargadas) y que por lo tanto actúa sobre funciones de ξ_a y ξ_b . Por medio de cálculos análogos a los que hicimos antes se encuentra que $\hat{F}^{(2)}$ (que expresa $F^{(2)}$ en el espacio de números de ocupación) está dado por

$$\hat{F}^{(2)} = \frac{1}{2} \sum_{i,k,l,m} f^{(2)ik}_{lm} \hat{a}_i^\dagger \hat{a}_k^\dagger \hat{a}_l \hat{a}_m \quad (14.47)$$

donde los elementos de matriz

$$f^{(2)ik}_{lm} = \iint \psi_i^*(\xi_1) \psi_k^*(\xi_2) f^{(2)} \psi_l(\xi_1) \psi_m(\xi_2) d\xi_1 d\xi_2 \quad (14.48)$$

son cantidades conocidas.

En efecto, las matrices de $\hat{F}^{(2)}$ y $F^{(2)}$ coinciden. Las fórmulas (14.44) y (14.47) se pueden generalizar para operadores $F^{(n)}$ que involucran cualquier número n de partículas, y que son simétricos respecto de las n partículas.

De esta forma podemos expresar en función de los operadores $\hat{a}_i$ y $\hat{a}_i^\dagger$ el Hamiltoniano $\hat{H}$ del sistema de n partículas en interacción. El operador H es, evidentemente, simétrico respecto de todas las partículas y en la aproximación no relativística no depende de los spines de las partículas. Por lo tanto lo podemos escribir en general en la forma

$$H = \sum_a h_a^{(1)} + \sum_{a>b} v_{ab}^{(2)} + \sum_{a>b>c} v_{abc}^{(3)} + \dots \quad (14.49)$$

Aquí $h_a^{(1)}$ es la parte de H que depende de las coordenadas de la a -ésima partícula y tiene la forma

$$h_a^{(1)} = -\frac{\hbar^2}{2m} \nabla_a^2 + v^{(1)}(\mathbf{r}_a) \quad (14.50)$$

donde $v^{(1)}(\mathbf{r}_a)$ es la energía potencial de la partícula en el campo externo. Los otros términos de la (14.49) corresponden a las interacciones entre las partículas, y hemos separado por comodidad los términos que dependen de las coordenadas de dos, tres, ... etc. partículas.

Podemos entonces usar la fórmulas (14.34), (14.47), y otras análogas para escribir $\hat{H}$:

$$\hat{H} = \sum_{i,k} h^{(1)i}_k \hat{a}_i^\dagger \hat{a}_k + \frac{1}{2} \sum_{i,k,l,m} v^{(2)ik}_{lm} \hat{a}_i^\dagger \hat{a}_k^\dagger \hat{a}_l \hat{a}_m + \dots \quad (14.51)$$

que escrito de esta forma es un operador que actúa sobre las funciones de los números de ocupación. En el caso de partículas que no interactúan el único término de la (14.51) que subsiste es el primero, de modo que

$$\hat{H} = \sum_{i,k} h^{(1)i}_k \hat{a}_i^\dagger \hat{a}_k \quad (14.52)$$

Si las $\psi_1(\xi), \psi_2(\xi), \dots$ son las autofunciones del Hamiltoniano de una partícula, de modo que $h^{(1)}\psi_i = e_i\psi_i$, la matriz $\hat{H}^i_k$ es diagonal y se tiene

$$\hat{H} = \sum_i e_i \hat{a}_i^\dagger \hat{a}_i = \sum_i e_i \hat{n}_i \quad (14.53)$$

de donde reemplazando los operadores $\hat{n}_i$ por sus autovalores n_i obtenemos los niveles de energía del sistema como

$$E_{n;n_1, n_2, \dots} = \sum_i e_i n_i \quad (14.54)$$

como era de esperar.

El espectro continuo y los operadores del campo de Bosones

Hasta aquí hemos supuesto que el sistema completo (14.1) corresponde a un operador con un espectro discreto de autovalores. Es importante extender el formalismo a los sistemas completos de autoestados de una partícula pertenecientes al *espectro continuo* de un observable. El ejemplo más común de una base continua es, naturalmente, el de las autofunciones del operador ξ , definidas por $\xi\psi_{\xi'} = \xi'\psi_{\xi'}$ y que están normalizadas como

$$(\psi_{\xi'}, \psi_{\xi''}) = \delta(\xi' - \xi'') = \delta(\mathbf{r}' - \mathbf{r}'')\delta_{\sigma'\sigma''} \quad (14.55)$$

Aquí con ξ' indicamos los autovalores continuos $\mathbf{r}'$ del operador posición $\mathbf{r}$, que junto con los autovalores discretos σ' de una cualquiera de las componentes del spin de la partícula, identifican las autofunciones básicas. Para evitar confusiones conviene aclarar nuestra notación pese a que su significado es obvio. Por un lado ξ indica el *operador* que representa los observables $\mathbf{r}$ y σ de la partícula. Por otro lado, cuando escribimos $\psi(\xi)$, ξ es la *variable* que representa las coordenadas y el spin en el espacio de configuración de la partícula. Con $\xi', \xi'', \xi''', \dots$ indicamos los *autovalores* de ξ . Por lo tanto al escribir $\psi_{\xi'}(\xi)$ estamos designando un particular estado de la partícula: aquél en el cual los observables $\mathbf{r}$ y σ tienen valores bien definidos $\mathbf{r}'$ y σ' . Al considerar todos los estados de nuestra base ortonormal $\psi_{\xi'}$, ξ' (es decir $\mathbf{r}'$ y σ') toma todos los valores posibles, pero no debemos confundir los autovalores que identifican cada uno de los estados básicos con la variable ξ .

El formalismo que desarrollamos hasta aquí se adapta fácilmente al espectro continuo, pero en este caso no es práctico usar los números de ocupación para designar los estados de un sistema de varias partículas, a menos que dividamos artificialmente el espacio (x, y, z) en un número infinito de celdas, pequeñas pero no infinitesimales (con lo cual, de hecho, volveríamos al espectro discreto). Vamos a ver que la segunda cuantificación se puede aplicar a una base continua introduciendo los *operadores del campo* ψ , definidos por

$$\hat{\psi}(\xi) = \sum_i \psi_i(\xi) \hat{a}_i \quad , \quad \hat{\psi}^\dagger(\xi) = \sum_i \psi_i^*(\xi) \hat{a}_i^\dagger \quad (14.56)$$

Para cada particular valor $\xi = \xi'$ de la variable ξ , las (14.56) definen un par de operadores $\hat{\psi}(\xi')$ y $\hat{\psi}^\dagger(\xi')$. El conjunto de los mismos, para todos los posibles valores de ξ' forma un infinito continuo que designamos con $\hat{\psi}(\xi)$ y $\hat{\psi}^\dagger(\xi)$ (aquí la variable ξ se considera un parámetro, ya que no actúa sobre los números de ocupación). En virtud de la ortonormalidad de la base $\psi_1(\xi), \psi_2(\xi), \dots$ es fácil ver que $\hat{a}_i$ y $\hat{a}_i^\dagger$ se expresan en términos de los $\hat{\psi}(\xi')$ y $\hat{\psi}^\dagger(\xi')$ como

$$\hat{a}_i = \int \psi_i^*(\xi') \hat{\psi}(\xi') d\xi' \quad , \quad \hat{a}_i^\dagger = \int \psi_i(\xi') \hat{\psi}^\dagger(\xi') d\xi' \quad (14.57)$$

Claramente, en virtud de las propiedades de $\hat{a}_i$ y $\hat{a}_i^\dagger$, los operadores $\hat{\psi}(\xi)$ y $\hat{\psi}^\dagger(\xi)$ disminuyen y aumentan en una unidad, respectivamente, el número total de partículas del sistema.

Es fácil ver que $\hat{\psi}^\dagger(\xi')$ crea una partícula en el punto ξ' . Efectivamente, $\hat{a}_i^\dagger$ crea una partícula cuya función de onda es $\psi_i(\xi)$. Luego $\hat{\psi}^\dagger(\xi')$ crea una partícula cuya función de onda es

$$\sum_i \psi_i^*(\xi') \psi_i(\xi) = \delta(\xi' - \xi) \quad (14.58)$$

donde $\delta(\xi' - \xi) = \delta(\mathbf{r}' - \mathbf{r})\delta(\sigma' - \sigma)$, y para escribir el resultado hemos usado la relación de clausura (7.75). Por lo tanto, $\hat{\psi}^\dagger(\xi')$ crea una partícula cuya posición es $\mathbf{r}'$ y cuyo spin es σ' .

Vemos entonces que los operadores del campo $\hat{\psi}(\xi)$ y $\hat{\psi}^\dagger(\xi)$ equivalen al conjunto de todos los operadores $\hat{a}_i$ y $\hat{a}_i^\dagger$ ($i = 1, 2, 3, \dots$) ya que dando a ξ diferentes valores ξ' se puede aniquilar o crear una partícula en cualquier punto del espacio y con cualquier valor del spin.

Las reglas de conmutación de $\hat{\psi}$ y $\hat{\psi}^\dagger$ se obtienen de inmediato a partir de las reglas de conmutación de $\hat{a}_i$ y $\hat{a}_i^\dagger$. Es evidente que

$$[\hat{\psi}(\xi'), \hat{\psi}(\xi'')] = 0, \quad [\hat{\psi}^\dagger(\xi'), \hat{\psi}^\dagger(\xi'')] = 0 \quad (14.59)$$

Por otra parte

$$\hat{\psi}(\xi')\hat{\psi}^\dagger(\xi'') - \hat{\psi}^\dagger(\xi'')\hat{\psi}(\xi') = \sum_i \sum_k \psi_i(\xi') \psi_k^*(\xi'') (\hat{a}_i \hat{a}_k^\dagger - \hat{a}_k^\dagger \hat{a}_i) = \sum_i \psi_i(\xi') \psi_i^*(\xi'') \quad (14.60)$$

de donde resulta

$$[\hat{\psi}(\xi'), \hat{\psi}^\dagger(\xi'')] = \delta(\xi'' - \xi') \quad (14.61)$$

La expresión (14.44) del operador $\hat{F}^{(1)}$ se escribe en términos de los nuevos operadores como

$$\hat{F}^{(1)} = \int \hat{\psi}^\dagger(\xi) f^{(1)} \hat{\psi}(\xi) d\xi \quad (14.62)$$

donde se entiende que $f^{(1)}$ actúa sobre las funciones $\psi_i(\xi)$ del operador $\hat{\psi}(\xi)$. En efecto, sustituyendo las (14.56) en (14.62) resulta

$$\hat{F}^{(1)} = \sum_{i,k} \int \psi_i^*(\xi) f^{(1)} \psi_k(\xi) d\xi \hat{a}_i^\dagger \hat{a}_k = \sum_{i,k} f^{(1)ik} \hat{a}_i^\dagger \hat{a}_k \quad (14.63)$$

que coincide con la (14.44). En particular si $f^{(1)}$ es una función $g(\xi)$, de la (14.62) obtenemos

$$\hat{G}^{(1)} = \int g(\xi) \hat{\psi}^\dagger(\xi) \hat{\psi}(\xi) d\xi \quad (14.64)$$

Por lo tanto $\hat{\psi}^\dagger(\xi) \hat{\psi}(\xi) d\xi$ es el operador número de partículas en el intervalo $(\xi, \xi + d\xi)$.

Del mismo modo que obtuvimos la (14.62) se encuentra que

$$\hat{F}^{(2)} = \frac{1}{2} \sum_{i,k,l,m} f^{(2)iklm} \hat{a}_i^\dagger \hat{a}_k^\dagger \hat{a}_l \hat{a}_m = \frac{1}{2} \iint \hat{\psi}^\dagger(\xi) \hat{\psi}^\dagger(\xi') f^{(2)} \hat{\psi}(\xi') \hat{\psi}(\xi) d\xi d\xi' \quad (14.65)$$

Usando la (14.62), la (14.66), la expresión (14.50) de $h_d^{(1)}$ e integrando por partes (sobre las coordenadas) al término que contiene los Laplacianos, podemos expresar al Hamiltoniano (14.49) en términos de los nuevos operadores:

$$\begin{aligned}
\hat{H} &= \sum_{i,k} h^{(1)}_k \hat{a}_i^\dagger \hat{a}_k + \frac{1}{2} \sum_{i,k,l,m} v^{(2)}_{lm} \hat{a}_i^\dagger \hat{a}_k^\dagger \hat{a}_l \hat{a}_m + \dots \\
&= \int \left\{ \frac{\hbar^2}{2m} \nabla \hat{\psi}^\dagger(\xi) \nabla \hat{\psi}(\xi) + v^{(1)}(\xi) \hat{\psi}^\dagger(\xi) \hat{\psi}(\xi) \right\} d\xi \\
&\quad + \frac{1}{2} \iint \hat{\psi}^\dagger(\xi) \hat{\psi}^\dagger(\xi') v^{(2)}(\xi, \xi') \hat{\psi}(\xi') \hat{\psi}(\xi) d\xi d\xi'
\end{aligned} \tag{14.66}$$

El resultado (14.66) se torna más intuitivo gracias al siguiente argumento. Sea un sistema de partículas, todas descriptas en un dado instante t por una misma función de onda $\psi(\xi)$, que suponemos normalizada de acuerdo con

$$\int |\psi|^2 d\xi = n \tag{14.67}$$

Si en la (14.66) reemplazamos los operadores $\hat{\psi}(\xi)$ y $\hat{\psi}^\dagger(\xi)$ por las funciones $\psi(\xi)$ y $\psi^*(\xi)$, respectivamente, la expresión que resulta es el valor medio de la energía del sistema cuando las n partículas están en el estado $\psi(\xi)$. A partir de esto obtenemos una regla para deducir la forma del Hamiltoniano en el formalismo de la segunda cuantificación: se escribe la expresión del valor medio de la energía, expresándolo por medio de la función de onda de una partícula $\psi(\xi)$, normalizada en la forma (14.65); luego se reemplazan $\psi(\xi)$ por el operador $\hat{\psi}(\xi)$ y $\psi^*(\xi)$ por $\hat{\psi}^\dagger(\xi)$, teniendo cuidado de escribir los operadores en el orden llamado *normal*, esto es con $\hat{\psi}^\dagger(\xi)$ a la izquierda de $\hat{\psi}(\xi)$.

Si el sistema está constituido por varias clases de Bosones, hay que introducir en el formalismo de segunda cuantificación los operadores $\hat{a}$, $\hat{a}^\dagger$ o bien $\hat{\psi}$, $\hat{\psi}^\dagger$ para cada clase de partículas. Es evidente que los operadores correspondientes a partículas de diferente clase conmutan entre sí.

Segunda cuantificación de sistemas de Fermiones

Todas las ideas básicas del método de segunda cuantificación se mantienen para los sistemas constituidos por Fermiones idénticos. En lo que hace a las expresiones concretas de los operadores que representan las magnitudes físicas, y los operadores de aniquilación y de creación, hay cambios, naturalmente.

El espacio de números de ocupación para Fermiones

La función de onda $\Psi_{n;n_1, n_2, \dots}$ de un sistema de n Fermiones tiene la forma

$$\Psi_{n;n_1, n_2, \dots} = \frac{1}{\sqrt{n!}} \begin{vmatrix} \psi_{r_1}(\xi_1) & \psi_{r_1}(\xi_2) & \dots & \psi_{r_1}(\xi_j) & \dots & \psi_{r_1}(\xi_n) \\ \psi_{r_2}(\xi_1) & \psi_{r_2}(\xi_2) & \dots & \psi_{r_2}(\xi_j) & \dots & \psi_{r_2}(\xi_n) \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \psi_{r_i}(\xi_1) & \psi_{r_i}(\xi_2) & \dots & \psi_{r_i}(\xi_j) & \dots & \psi_{r_i}(\xi_n) \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \psi_{r_n}(\xi_1) & \psi_{r_n}(\xi_2) & \dots & \psi_{r_n}(\xi_j) & \dots & \psi_{r_n}(\xi_n) \end{vmatrix} \tag{14.68}$$

donde hemos empleado la notación $\psi_{r_1}, \psi_{r_2}, \dots, \psi_{r_i}, \dots, \psi_{r_n}$ para subrayar que los estados ocupados son todos diferentes. Dada la antisimetría de $\Psi_{n;n_1, n_2, \dots}$ es importante establecer una convención para determinar su signo (esta cuestión no se plantea en el caso de Bosones, pues debido

a la simetría de Ψ , el signo, una vez elegido, se mantiene para cualquier permutación de las ξ_i). Para ello vamos a convenir en numerar consecutivamente *de una vez y para siempre* todos los estados ψ_i de una partícula. En base a tal numeración, llenaremos las líneas del determinante (14.68) en orden de r_i creciente, esto es:

$$r_1 < r_2 < r_3 < \dots < r_n \quad (14.69)$$

y la columnas las llenaremos con funciones de diferentes variables, en el orden $\xi_1, \xi_2, \dots, \xi_n$. No puede haber dos números iguales entre $r_1, r_2, \dots, r_n$, pues el determinante sería nulo. En otras palabras, los números de ocupación n_i pueden valer solamente 0 ó 1.

Elementos de matriz de operadores en el espacio de números de ocupación

Consideremos un operador de la forma (14.15), esto es

$$F^{(1)} = \sum_a f_a^{(1)} \quad (14.70)$$

donde $f_a^{(1)}$ es el operador de una partícula correspondiente a la partícula a , esto es, que opera sobre las variables ξ_a . Por las mismas razones que invocamos antes, los únicos elementos de matriz no nulos de $F^{(1)}$ son: (a) los elementos diagonales, que corresponden a que los números de ocupación n_i no cambian, (b) aquellos elementos fuera de la diagonal tales que uno de los números de ocupación (n_i) pasa de 0 a 1 y otro (n_k) disminuye de 1 a 0 (ponemos los sub-índices para que quede claro a qué estado se refiere cada número de ocupación).

Para los elementos diagonales es fácil ver (no haremos los cálculos) que se obtiene un resultado análogo al que obtuvimos antes para los Bosones, esto es

$$\bar{F}^{(1)} = \sum_i n_i f_i^{(1)} \quad (14.71)$$

La única diferencia de la (14.71) respecto de la (14.18) es que los números de ocupación n_i de los Fermiones pueden tomar solamente los valores 0 y 1.

Consideremos ahora el elemento de matriz no diagonal

$$\begin{aligned} F^{(1)}_{0_i, 1_k}^{1_i, 0_k} &= (\Psi_{n; n_1, \dots, 1_i, \dots, 0_k, \dots}, F^{(1)} \Psi_{n; n_1, \dots, 0_i, \dots, 1_k, \dots}) \\ &= \sum_a (\Psi_{n; n_1, \dots, 1_i, \dots, 0_k, \dots}, f_a^{(1)} \Psi_{n; n_1, \dots, 0_i, \dots, 1_k, \dots}) \end{aligned} \quad (14.72)$$

En la (14.70) $\Psi_{n; n_1, \dots, 1_i, \dots, 0_k, \dots}$ tiene la forma

$$\Psi_{n; n_1, \dots, 1_i, \dots, 0_k, \dots} = \frac{1}{\sqrt{n!}} \begin{vmatrix} \psi_{r_1}(\xi_1) & \psi_{r_1}(\xi_2) & \dots & \psi_{r_1}(\xi_a) & \dots & \psi_{r_1}(\xi_n) \\ \psi_{r_2}(\xi_1) & \psi_{r_2}(\xi_2) & \dots & \psi_{r_2}(\xi_a) & \dots & \psi_{r_2}(\xi_n) \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \psi_i(\xi_1) & \psi_i(\xi_2) & \dots & \psi_i(\xi_a) & \dots & \psi_i(\xi_n) \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \psi_{r_n}(\xi_1) & \psi_{r_n}(\xi_2) & \dots & \psi_{r_n}(\xi_a) & \dots & \psi_{r_n}(\xi_n) \end{vmatrix} \quad (14.73)$$

Por supuesto no existe en (14.73) una fila ψ_k , dado que $n_k = 0$.

Del mismo modo $\Psi_{n;n_1, \dots, 0_i, \dots, 1_k, \dots}$ está dado por

$$\Psi_{n;n_1, \dots, 0_i, \dots, 1_k, \dots} = \frac{1}{\sqrt{n!}} \begin{vmatrix} \psi_{r1}(\xi_1) & \psi_{r1}(\xi_2) & \dots & \psi_{r1}(\xi_a) & \dots & \psi_{r1}(\xi_n) \\ \psi_{r2}(\xi_1) & \psi_{r2}(\xi_2) & \dots & \psi_{r2}(\xi_a) & \dots & \psi_{r2}(\xi_n) \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \psi_k(\xi_1) & \psi_k(\xi_2) & \dots & \psi_k(\xi_a) & \dots & \psi_k(\xi_n) \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \psi_{rn}(\xi_1) & \psi_{rn}(\xi_2) & \dots & \psi_{rn}(\xi_a) & \dots & \psi_{rn}(\xi_n) \end{vmatrix} \quad (14.74)$$

y lógicamente no existe en (14.74) una fila ψ_i .

Al desarrollar estos determinantes conviene hacerlo por las filas i y k , respectivamente. Para ello es importante tener presente los signos de los *adjuntos*⁷ (o *complementos algebraicos*) de los elementos de dichas filas. Sea $A(i, a)$ el *adjunto* del elemento $\psi_i(\xi_a)$ del determinante (14.73), y $B(k, a)$ el adjunto del elemento $\psi_k(\xi_a)$ del determinante (14.74). Es fácil ver que los determinantes que resultan de suprimir la fila i en (14.73) y la fila k en (14.74), y la columna a en ambos son idénticos. Pero los adjuntos $A(i, a)$ y $B(k, a)$ corresponden a elementos que pertenecen a filas diferentes de (14.73) y (14.74). Sea T_i el número de orden de la fila ψ_i en (14.73); claramente T_i es igual al número de estados ψ_{rj} ocupados de la sucesión $\psi_1, \psi_2, \dots$ con $rj < i$ y por lo tanto

$$T_i = \sum_{j=1}^{i-1} n_j \quad (14.75)$$

Luego el signo de $A(i, a)$ está dado por $(-1)^{T_i+a}$ y del mismo modo el signo de $B(k, a)$ está dado por $(-1)^{T_k+a}$. Esto trae una diferencia de signo entre dichos adjuntos dada por $(-1)^{T_k-T_i}$, donde $T_k - T_i$ es igual al número de estados ocupados de la sucesión $\psi_1, \psi_2, \dots$ comprendidos entre ψ_i y ψ_k . En definitiva

$$B(k, a) = (-1)^{T_k-T_i} A(i, a) = (-1)^{T_k+T_i} A(i, a) \quad (14.76)$$

pues es evidente que los signos que se pongan delante de T_i y T_k en (14.74) son irrelevantes.

Calculemos un término de la suma (14.72), por ejemplo el que corresponde a $f_a^{(1)}$. Claramente, las únicas contribuciones no nulas provienen de los términos

$$\frac{1}{\sqrt{n!}} \psi_i(\xi_a) A(i, a) \quad \text{y} \quad \frac{1}{\sqrt{n!}} \psi_k(\xi_a) B(k, a) = \frac{1}{\sqrt{n!}} (-1)^{T_k+T_i} \psi_k(\xi_a) A(i, a) \quad (14.77)$$

de los desarrollos de $\Psi_{n;n_1, \dots, 1_i, \dots, 0_k, \dots}$ y de $\Psi_{n;n_1, \dots, 0_i, \dots, 1_k, \dots}$, respectivamente. La contribución buscada es entonces

$$(\Psi_{n;1_i,0_k}, f_a^{(1)} \Psi_{n;0_i,1_k}) = \frac{(-1)^{T_k+T_i}}{n!} (\psi_i(\xi_a) A(i, a), f_a^{(1)} \psi_k(\xi_a) A(i, a)) = \frac{(-1)^{T_k+T_i}}{n} f^{(1)}_k \quad (14.78)$$

⁷ Sea D un determinante, y $d(f, c)$ el determinante que se obtiene suprimiendo la fila f y la columna c de D . Entonces $a(f, c) = (-1)^{f+c} d(f, c)$ es el adjunto del elemento (f, c) de D .

Para obtener este resultado, tenemos que recordar que el desarrollo de $A(i, a)$ consta de $(n-1)!$ productos de las $\psi_{r_l}(\xi_m)$, todos ortonormales. Puesto que en el elemento de matriz (14.71) hay n contribuciones iguales a la (14.77) obtenemos finalmente

$$F_{0_i, 1_k}^{(1)1_i, 0_k} = (-1)^{T_k + T_i} f^{(1)k}_i \quad (14.79)$$

Operadores de aniquilación y de creación de Fermiones

Para que el operador de una partícula $F^{(1)}$ se pueda escribir en la forma (14.44), es preciso definir los operadores de aniquilación y de creación de Fermiones de manera que se respete el signo de la (14.79). Por lo tanto se debe tomar en cuenta qué lugar ocupa el estado de la partícula que se está aniquilando y cual ocupará el de la partícula que se está creando dentro de las listas $r_1 < r_2 < r_3 < \dots < r_n$ de los estados ocupados del sistema, antes y después de la transición. Esto se logra con las siguientes definiciones:

$$\begin{aligned} \hat{a}_i \Psi_{n; n_1, n_2, \dots, n_i, \dots} &= (-1)^{T_i} \sqrt{n_i} \Psi_{n-1; n_1, n_2, \dots, n_i-1, \dots} \\ \hat{a}_i^\dagger \Psi_{n; n_1, \dots, n_i, \dots} &= (-1)^{T_i} \sqrt{n_i + 1} \Psi_{n+1; n_1, \dots, n_i+1, \dots} \end{aligned} \quad (14.80)$$

Las (14.80) se pueden escribir más simplemente como

$$\begin{aligned} \hat{a}_i \Psi_{n; n_1, n_2, \dots, 1_i, \dots} &= (-1)^{T_i} \Psi_{n-1; n_1, n_2, \dots, 0_i, \dots} \\ \hat{a}_i^\dagger \Psi_{n; n_1, \dots, 0_i, \dots} &= (-1)^{T_i} \Psi_{n+1; n_1, \dots, 1_i, \dots} \end{aligned} \quad (14.81)$$

los únicos elementos no nulos de la matriz que representa $\hat{a}_i$ son, evidentemente,

$$\hat{a}_i^{0_i}_{1_i} = (\Psi_{n-1; n_1, \dots, 0_i, \dots}, \hat{a}_i \Psi_{n; n_1, \dots, 1_i, \dots}) = (-1)^{T_i} \quad (14.82)$$

Recordando la definición (7.92) del adjunto de un operador es inmediato verificar que el operador de creación $\hat{a}_i^\dagger$ definido por la (14.81) es, efectivamente, el adjunto de $\hat{a}_i$, y que sus elementos de matriz no nulos son

$$\hat{a}_i^{\dagger 1_i}_{0_i} = (\Psi_{n; n_1, \dots, 1_i, \dots}, \hat{a}_i^\dagger \Psi_{n-1; n_1, \dots, 0_i, \dots}) = (-1)^{T_i} \quad (14.83)$$

Consideremos ahora los elementos de matriz del operador

$$\hat{F}^{(1)} = \sum_{i,k} f^{(1)k}_i \hat{a}_i^\dagger \hat{a}_k \quad (14.84)$$

entre los estados base del espacio de números de ocupación. Claramente, los únicos elementos no nulos son

$$\hat{F}^{(1)1_i, 0_k}_{0_i, 1_k} = (\Psi_{n; n_1, \dots, 1_i, \dots, 0_k, \dots}, \hat{F}^{(1)} \Psi_{n; n_1, \dots, 0_i, \dots, 1_k, \dots}) = (-1)^{T_i + T_k} f^{(1)k}_i \quad (14.85)$$

y por lo tanto $\hat{F}^{(1)}$ coincide con $F^{(1)}$.

Consideremos ahora el efecto del operador $\hat{n}_i = \hat{a}_i^\dagger \hat{a}_i$ sobre $\Psi_{n; n_1, n_2, \dots}$. Usando las definiciones (14.81) obtenemos de inmediato

$$\hat{n}_i \Psi_{n; n_1, n_2, \dots, n_i, \dots} = \hat{a}_i^\dagger \hat{a}_i \Psi_{n; n_1, n_2, \dots, n_i, \dots} = \begin{cases} \Psi_{n; n_1, n_2, \dots, 1_i, \dots} & \text{si } n_i = 1 \\ 0 & \text{si } n_i = 0 \end{cases} \quad (14.86)$$

es decir

$$\hat{n}_i \Psi_{n; n_1, n_2, \dots, n_i, \dots} = \hat{a}_i^\dagger \hat{a}_i \Psi_{n; n_1, n_2, \dots, n_i, \dots} = n_i \Psi_{n; n_1, n_2, \dots, n_i, \dots} \quad (14.87)$$

Por lo tanto $\hat{n}_i$ es el *operador número de partículas en el estado i*. Asimismo, el operador

$$\hat{n} = \sum_i \hat{n}_i = \sum_i \hat{a}_i^\dagger \hat{a}_i \quad (14.88)$$

es el *operador número total de partículas*, igual que para los Bosones.

Veamos el efecto del operador $\hat{a}_i \hat{a}_i^\dagger$ sobre $\Psi_{n; n_1, n_2, \dots}$. Claramente:

$$\hat{a}_i \hat{a}_i^\dagger \Psi_{n; n_1, n_2, \dots, n_i, \dots} = \begin{cases} 0 & \text{si } n_i = 1 \\ \Psi_{n; n_1, n_2, \dots, 1_i, \dots} & \text{si } n_i = 0 \end{cases} \quad (14.89)$$

Luego podemos escribir simplemente

$$\hat{a}_i \hat{a}_i^\dagger \Psi_{n; n_1, n_2, \dots, n_i, \dots} = (1 - n_i) \Psi_{n; n_1, n_2, \dots, n_i, \dots} \quad (14.90)$$

Además, de la (14.87) y la (14.90) resulta la siguiente relación entre los operadores de aniquilación y de creación del mismo estado de una partícula:

$$\hat{a}_i \hat{a}_i^\dagger + \hat{a}_i^\dagger \hat{a}_i = 1 \quad (14.91)$$

Consideremos ahora el efecto del operador $\hat{a}_i^\dagger \hat{a}_k$ ($i \neq k$) sobre $\Psi_{n; n_1, n_2, \dots}$. Usando las definiciones (14.81) obtenemos de inmediato

$$\begin{aligned} \hat{a}_i^\dagger \hat{a}_k \Psi_{n; n_1, n_2, \dots, 0_i, \dots, 1_k, \dots} &= (-1)^{T_k + T_i} \Psi_{n-1; n_1, n_2, \dots, 1_i, \dots, 0_k, \dots} & \text{si } i < k \\ \hat{a}_i^\dagger \hat{a}_k \Psi_{n; n_1, n_2, \dots, 1_k, \dots, 0_i, \dots} &= -(-1)^{T_k + T_i} \Psi_{n-1; n_1, n_2, \dots, 0_k, \dots, 1_i, \dots} & \text{si } i > k \end{aligned} \quad (14.92)$$

Notar el signo $-$ que aparece para $i > k$, que se debe a que al haber aniquilado la partícula que ocupaba el estado ψ_k , la cantidad de estados ocupados con $rj < i$ ha disminuido en una unidad.

Por último calculemos el efecto del operador $\hat{a}_k \hat{a}_i^\dagger$ ($i \neq k$, es el mismo producto que antes, pero en el orden inverso) sobre $\Psi_{n; n_1, n_2, \dots}$. Encontramos que

$$\begin{aligned} \hat{a}_k \hat{a}_i^\dagger \Psi_{n; n_1, n_2, \dots, 0_i, \dots, 1_k, \dots} &= -(-1)^{T_k + T_i} \Psi_{n-1; n_1, n_2, \dots, 1_i, \dots, 0_k, \dots} & \text{si } i < k \\ \hat{a}_k \hat{a}_i^\dagger \Psi_{n; n_1, n_2, \dots, 1_k, \dots, 0_i, \dots} &= (-1)^{T_k + T_i} \Psi_{n-1; n_1, n_2, \dots, 0_k, \dots, 1_i, \dots} & \text{si } i > k \end{aligned} \quad (14.93)$$

Resulta por consiguiente que

$$\hat{a}_k \hat{a}_i^\dagger + \hat{a}_i^\dagger \hat{a}_k = 0 \quad , \quad i \neq k \quad (14.94)$$

Las dos igualdades (14.91) y (14.94) se pueden resumir en una sola escribiendo para todo i, k :

$$\hat{a}_k \hat{a}_i^\dagger + \hat{a}_i^\dagger \hat{a}_k = \delta_{ik} \quad (14.95)$$

Efectuando cálculos análogos es sencillo verificar también que para todo i, k se cumple que

$$\hat{a}_k \hat{a}_i + \hat{a}_i \hat{a}_k = 0 \quad , \quad \hat{a}_k^\dagger \hat{a}_i^\dagger + \hat{a}_i^\dagger \hat{a}_k^\dagger = 0 \quad (14.96)$$

Las (14.95) y (14.96) se suelen llamar *relaciones de anticonmutación*, para distinguirlas de las relaciones de conmutación (14.42) que se cumplen para los Bosones. Vemos, por consiguiente, que los operadores de aniquilación y creación de Fermiones son *anticonmutativos*, a diferencia de los operadores de aniquilación y creación de Bosones, que son conmutativos.

Es importante destacar otra diferencia entre los operadores de aniquilación y creación de Bosones y de Fermiones. Los operadores $\hat{a}_i, \hat{a}_k$ de Bosones son completamente independientes: cada uno de ellos opera sobre una sola variable n_i y el resultado de su acción no depende de los valores de los demás números de ocupación. Para los Fermiones, en cambio, el efecto del operador $\hat{a}_i$ no depende solamente del número de ocupación n_i , sino que también depende de los números de ocupación de todos los estados ψ_j precedentes ($j < i$), como resulta de la definición (14.81). Por este motivo, la acción de operadores $\hat{a}_i, \hat{a}_k$ distintos *no* se puede considerar independiente.

Operadores en el espacio de números de ocupación y operadores del campo de Fermiones

Habiendo definido las propiedades de los operadores de aniquilación y creación de Fermiones en la forma precedente, todas las fórmulas desde la (14.43) hasta la (14.54) que obtuvimos para los Bosones subsisten íntegramente para los Fermiones. Los operadores del campo de Fermiones se definen (del mismo modo que para los Bosones) como

$$\hat{\psi}(\xi) = \sum_i \psi_i(\xi) \hat{a}_i \quad , \quad \hat{\psi}^\dagger(\xi) = \sum_i \psi_i^*(\xi) \hat{a}_i^\dagger \quad (14.97)$$

donde las variables ξ se consideran como parámetros. Del mismo modo que antes, se ve que $\hat{\psi}^\dagger(\xi')$ crea una partícula en el punto ξ' .

Las reglas de *anticonmutación* de $\hat{\psi}$ y $\hat{\psi}^\dagger$ se obtienen de inmediato a partir de las reglas de anticonmutación de $\hat{a}_i$ y $\hat{a}_i^\dagger$. Es evidente que

$$\hat{\psi}(\xi') \hat{\psi}(\xi'') + \hat{\psi}(\xi'') \hat{\psi}(\xi') = 0 \quad , \quad \hat{\psi}^\dagger(\xi') \hat{\psi}^\dagger(\xi'') + \hat{\psi}^\dagger(\xi'') \hat{\psi}^\dagger(\xi') = 0 \quad (14.98)$$

Por otra parte

$$\hat{\psi}(\xi') \hat{\psi}^\dagger(\xi'') + \hat{\psi}^\dagger(\xi'') \hat{\psi}(\xi') = \delta(\xi'' - \xi') \quad (14.99)$$

Las expresiones de los operadores de las magnitudes físicas de un sistema de Fermiones se expresan entonces en términos de $\hat{\psi}$ y $\hat{\psi}^\dagger$ por medio de las ecuaciones (14.62) a (14.66), que no reproducimos aquí.

Si el sistema está constituido por diferentes especies de Fermiones y Bosones, hay que introducir para cada especie los correspondientes operadores de segunda cuantificación ($\hat{a}_i$ y $\hat{a}_i^\dagger$ o bien $\hat{\psi}$ y $\hat{\psi}^\dagger$). Los operadores de Fermiones conmutan con los operadores de Bosones. En cuanto a los operadores de Fermiones de distinta especie, dentro de los límites de la presente teoría no relativística, es lícito suponer que conmuten, y también que anticonmuten. En ambos casos, la aplicación de la segunda cuantificación lleva a los mismos resultados.

Sin embargo, puesto que eventualmente la teoría se va a extender al caso relativístico, en el cual se admiten las transmutaciones de partículas de una especie en otra, conviene postular que *los operadores de segunda cuantificación de Fermiones de distinta especie anticonmutan*. La conveniencia de proceder así resulta evidente si se tiene presente que a veces dos especies distintas se consideran como dos estados internos distintos de una única especie (como ocurre con el protón y el neutrón, que se consideran como distintos estados de una sola especie, el nucleón).

Con estos resultados el lector puede por fin apreciar la elegancia y sencillez de la segunda cuantificación. Todo lo que hace falta para implementar el formalismo es conocer un sistema completo cualquiera $\psi_1(\xi), \psi_2(\xi), \dots$ de autoestados de una partícula y los elementos de matriz $f^{(1)}_k, f^{(2)}_{lm}, \dots$, etc. de los operadores de interés. A partir de esos datos y mediante los operadores de aniquilación y creación $\hat{a}_i$ y $\hat{a}_i^\dagger$ apropiados (o bien por medio de los operadores del campo $\hat{\psi}$ y $\hat{\psi}^\dagger$), se puede construir el espacio de números de ocupación $\mathcal{E}$ y calcular cualquier cantidad de interés para sistemas con un número arbitrario (que puede ser variable) de partículas idénticas. No hace falta ya preocuparse por los aspectos combinatorios que derivan de la estadística apropiada a Bosones y Fermiones porque los mismos se toman en cuenta *automáticamente* en los cálculos gracias a las relaciones de conmutación fundamentales (14.42), (14.59) y (14.61) para los Bosones y las relaciones de anticonmutación (14.95), (14.96), (14.98) y (14.99) para los Fermiones.

Ecuaciones de movimiento para los operadores de campos de Fermiones y Bosones

Recordemos que en el Capítulo 8 definimos el operador derivada total con respecto del tiempo dF/dt de un operador F por medio de la ec. (8.60), que podemos escribir como

$$i\hbar \frac{dF}{dt} = [F, H] + i\hbar \frac{\partial F}{\partial t} \quad (14.100)$$

donde H es el Hamiltoniano del sistema. La (14.100) se reduce a

$$i\hbar \frac{dF}{dt} = [F, H] \quad (14.101)$$

si F no depende explícitamente del tiempo. La (14.100) o la (14.101), según corresponda, es la ecuación de movimiento de F .

Aplicaremos estas fórmulas para obtener la ecuación de movimiento de los operadores de aniquilación $\hat{a}_i$ y de los operadores del campo $\hat{\psi}(\xi)$. Claramente $\hat{a}_i$ no depende del tiempo, y puesto que las funciones básicas $\psi_i(\xi)$ tampoco dependen del tiempo, $\hat{\psi}(\xi)$ no depende explícitamente del tiempo. Por lo tanto corresponde usar la (14.101). El Hamiltoniano de nuestro sistema se puede escribir en la forma (14.51):

$$\hat{H} = \sum_{j,l} h^{(1)}_l \hat{a}_j^\dagger \hat{a}_l + \frac{1}{2} \sum_{j,k,l,m} v^{(2)}_{lm}{}^{jk} \hat{a}_j^\dagger \hat{a}_k^\dagger \hat{a}_l \hat{a}_m + \dots = \hat{H}_1 + \hat{H}_2 + \dots \quad (14.102)$$

Calculemos el conmutador $[\hat{a}_i, \hat{H}]$, para lo cual recordemos que para todo i, k se cumple

$$\hat{a}_k \hat{a}_i^\dagger \mp \hat{a}_i^\dagger \hat{a}_k = \delta_{ik}, \quad \hat{a}_k \hat{a}_i \mp \hat{a}_i \hat{a}_k = 0, \quad \hat{a}_k^\dagger \hat{a}_i^\dagger \mp \hat{a}_i^\dagger \hat{a}_k^\dagger = 0 \quad (14.103)$$

donde los signos $-$ corresponden a Bosones y los signos $+$ a Fermiones. Aplicando las reglas (14.103) es muy sencillo verificar que tanto para Bosones como para Fermiones resulta que

$$[\hat{a}_i, \hat{a}_j^\dagger \hat{a}_l] = \delta_{ij} \hat{a}_l \quad , \quad [\hat{a}_i, \hat{a}_j^\dagger \hat{a}_k^\dagger \hat{a}_l \hat{a}_m] = \delta_{ij} \hat{a}_k^\dagger \hat{a}_l \hat{a}_m + \delta_{ik} \hat{a}_j^\dagger \hat{a}_m \hat{a}_l \quad (14.104)$$

donde el lector tiene que prestar atención al orden en que se escriben los operadores. Resulta entonces que

$$\hat{a}_i \hat{H}_1 - \hat{H}_1 \hat{a}_i = \sum_l h^{(1)l}_i \hat{a}_l \quad (14.105)$$

y que

$$\hat{a}_i \hat{H}_2 - \hat{H}_2 \hat{a}_i = \frac{1}{2} \sum_{j,l,m} (v^{(2)ij}_{lm} + v^{(2)ji}_{ml}) \hat{a}_j^\dagger \hat{a}_l \hat{a}_m = \sum_{j,l,m} v^{(2)ij}_{lm} \hat{a}_j^\dagger \hat{a}_l \hat{a}_m \quad (14.106)$$

puesto que $v^{(2)ij}_{lm} = v^{(2)ji}_{ml}$, como se verifica fácilmente recordando la (14.48).

De (14.105) y (14.106) resulta entonces la ecuación de movimiento para $\hat{a}_i$

$$i\hbar \frac{d\hat{a}_i}{dt} = \sum_l h^{(1)l}_i \hat{a}_l + \sum_{j,l,m} v^{(2)ij}_{lm} \hat{a}_j^\dagger \hat{a}_l \hat{a}_m \quad (14.107)$$

La ecuación de movimiento para los operadores del campo es

$$\begin{aligned} i\hbar \frac{d\hat{\psi}(\xi)}{dt} &= [\hat{\psi}(\xi), \hat{H}] = \sum_i \psi_i(\xi) [\hat{a}_i, \hat{H}] \\ &= \sum_{i,l} \psi_i(\xi) h^{(1)l}_i \hat{a}_l + \sum_{i,j,l,m} \psi_i(\xi) v^{(2)ij}_{lm} \hat{a}_j^\dagger \hat{a}_l \hat{a}_m \end{aligned} \quad (14.108)$$

Recordando la (14.57), el primer término del miembro derecho es

$$\sum_{i,l} \psi_i(\xi) h^{(1)l}_i \hat{a}_l = \int \left(\sum_{i,l} \psi_l^*(\xi') h^{(1)l}_i \psi_i(\xi) \right) \hat{\psi}(\xi') d\xi' \quad (14.109)$$

Usando la (14.14) y la relación de clausura (7.75) resulta

$$\sum_{i,l} \psi_i(\xi) h^{(1)l}_i \hat{a}_l = h^{(1)}(\xi) \hat{\psi}(\xi) \quad (14.110)$$

Del mismo modo, usando la (14.48) y la relación de clausura (7.75) obtenemos con un poco de paciencia que

$$\sum_{i,j,l,m} \psi_i(\xi) v^{(2)ij}_{lm} \hat{a}_j^\dagger \hat{a}_l \hat{a}_m = \int v^{(2)}(\xi, \xi') \hat{\psi}^\dagger(\xi') \hat{\psi}(\xi') d\xi' \hat{\psi}(\xi) \quad (14.111)$$

de modo que la ecuación de movimiento queda en la forma

$$i\hbar \frac{d\hat{\psi}(\xi)}{dt} = h^{(1)}(\xi) \hat{\psi}(\xi) + \int v^{(2)}(\xi, \xi') \hat{\psi}^\dagger(\xi') \hat{\psi}(\xi') d\xi' \hat{\psi}(\xi) = [h^{(1)}(\xi) + \mathcal{V}_{ef}(\xi)] \hat{\psi}(\xi) \quad (14.112)$$

La (14.112) es una ecuación integrodiferencial en el espacio ordinario y en el tiempo, y se aplica tanto a Bosones como a Fermiones. Debido a la presencia del término de interacción $\mathcal{V}_{ef}(\xi)$ proveniente de $\hat{H}_2$, se trata de una ecuación no lineal. Se puede interpretar que este término describe el efecto sobre una partícula de sus interacciones con las demás partículas. El potencial efectivo que describe dichas interacciones está dado por

$$\mathcal{V}_{ef}(\xi) = \int v^{(2)}(\xi, \xi') \hat{\psi}^\dagger(\xi') \hat{\psi}(\xi') d\xi' \quad (14.113)$$

y se puede calcular solamente si se conoce previamente la solución de la (14.112). Esto sugiere el empleo de procedimientos iterativos para encontrar soluciones autoconsistentes de la ecuación del movimiento. Esas técnicas son, en efecto, de uso frecuente para resolver problemas de muchos cuerpos.

Conexión de la segunda cuantificación con la Teoría Cuántica de Campos

Si no hay interacciones entre las partículas, la (14.112) se reduce a

$$i\hbar \frac{d\hat{\psi}(\xi)}{dt} = h^{(1)}(\xi) \hat{\psi}(\xi) \quad (14.114)$$

Esta ecuación es formalmente idéntica a la ecuación de Schrödinger para una partícula⁸ (ec. (7.10)):

$$i\hbar \frac{d\psi(\xi)}{dt} = h^{(1)}(\xi) \psi(\xi) \quad (14.115)$$

Debemos recordar, sin embargo, que $\hat{\psi}(\xi)$ es un *operador del campo* y no una *función de onda* ordinaria. La ecuación (14.114) se puede considerar como la versión *cuantificada* de la ecuación de Schrödinger (14.115), si interpretamos esta última como la ecuación de un *campo clásico* $\psi(\xi)$ (el campo que describe la onda asociada con una partícula, digamos, por ejemplo, un electrón). Este es precisamente el punto de contacto entre la Teoría Cuántica de Campos y la segunda cuantificación, que hemos desarrollado para tratar problemas de muchos cuerpos.

De la segunda cuantificación a la Teoría Cuántica de Campos hay un solo paso, fundamentalmente de interpretación. Existen dos maneras equivalentes de dar ese paso.

La primera toma como punto de partida las ecuaciones clásicas que describen un campo $\psi(\xi)$ (como la 14.115), e interpreta que dichas ecuaciones rigen la dinámica de los operadores del campo $\hat{\psi}(\xi)$ y $\hat{\psi}^\dagger(\xi)$, operadores que cumplen las reglas de conmutación fundamentales (14.42), (14.59) y (14.61) si el campo describe Bosones o bien las relaciones de anticonmutación (14.95), (14.96), (14.98) y (14.99) si se trata de Fermiones. El campo queda entonces *cuantificado*, pues los operadores $\hat{\psi}(\xi)$ y $\hat{\psi}^\dagger(\xi)$ destruyen y crean partículas. Las partículas, es decir los cuantos del campo, aparecen o desaparecen de a una por vez (y no en fracciones arbitrariamente pequeñas). Que esto sea obvio para el electrón, se debe en realidad a que siempre hemos dado por sentado que se trata de una partícula, incluso cuando lo describimos por medio de una función de onda. No es igualmente obvio, sin embargo, para el campo de la radiación electro-

⁸ La notación que empleamos para los operadores del campo fue elegida precisamente teniendo presente esta identidad formal.

magnética, que estamos acostumbrados a imaginar como una onda, aunque la evidencia experimental que discutimos en el Capítulo 4 muestra que en sus interacciones con la materia se comporta como un sistema de partículas (los fotones). La Teoría Cuántica de Campos toma en cuenta así los dos aspectos (onda y partícula) de los entes de escala atómica y subatómica, que quedan ambos incorporados a la esencia misma del formalismo, implementando de este modo el Principio de Complementaridad del cual hablamos al final del Capítulo 6.

La segunda manera parte de considerar nuestro formalismo en el espacio de los números de ocupación. En esencia, dicho formalismo analiza un sistema de muchas partículas como un sistema mecánico dotado de infinitos grados de libertad, cada uno de los cuales está representado por uno de los estados básicos $\psi_1(\xi), \psi_2(\xi), \dots$ de una partícula. Cada grado de libertad se puede excitar, lo que corresponde a que esté ocupado por una partícula (o más, si se trata de Bosones). De acuerdo con la segunda cuantificación, cada grado de libertad equivale formalmente a un oscilador armónico simple: recordemos en efecto que los estados excitados del oscilador se pueden obtener a partir del estado fundamental por aplicación reiterada del operador de subida $a^\dagger$ (como hicimos en el Capítulo 9), del mismo modo que aplicando repetidas veces el operador de creación $\hat{a}_i^\dagger$ podemos ir poblando con partículas el estado ψ_i .

De acuerdo con este punto de vista, el procedimiento a seguir para cuantificar un campo (por ejemplo el campo de la radiación electromagnética) consiste en analizar el campo en términos de un sistema ortonormal completo de autofunciones, esto es, de modos propios de oscilación (tal como hicimos en el Capítulo 4 para deducir la distribución de Planck), lo que equivale a representar el campo como una superposición de infinitos osciladores. Luego es preciso cuantificar esos osciladores, lo que se consigue postulando oportunas reglas de conmutación (o de anticonmutación) para los operadores de aniquilación y de creación $\hat{a}_i$ y $\hat{a}_i^\dagger$.

Ambas maneras de proceder son equivalentes, y llevan a los mismos resultados. Esto es lógico, dado que en realidad estamos describiendo la misma cosa, esto es, los estados de un sistema de muchas partículas, pero usando bases diferentes. En el primer caso, nuestras funciones básicas pertenecen al espectro continuo (las autofunciones del operador ξ) y en el segundo pertenecen al espectro discreto.

Nuestro tratamiento ha sido no relativístico, lo cual limita el presente formalismo a sistemas de muchas partículas de interés para la Física del Estado Sólido y la Física Nuclear, pero excluye el estudio de procesos de alta energía, de interés para la Física de Partículas. No haremos en estas notas un tratamiento de la Teoría Cuántica Relativística de Campos. Pero se debe señalar que los pasos conceptuales que hay que dar para desarrollar dicha teoría son los mismos que hemos expuesto aquí. Se presentan importantes complicaciones, desde luego, pues es preciso modificar el formalismo para que sea covariante.

En el tratamiento relativístico surge una importante diferencia respecto de lo que hemos visto aquí. En nuestro tratamiento hemos *postulado* la simetría de la función de onda de un sistema de Bosones, y su antisimetría para sistemas de Fermiones. Hicimos así porque es un dato de la realidad que los Fermiones cumplen el Principio de Pauli, mientras que cualquier número de Bosones pueden ocupar el mismo estado. Este postulado nos llevó a las relaciones fundamentales de conmutación para los Bosones y de anticonmutación para los Fermiones. Sin embargo, cabe señalar que en el marco de una teoría no relativística no estamos *obligados* a proceder así: en efecto, es posible formular una teoría *consistente* de muchas partículas postulando relaciones de conmutación (o de anticonmutación) para todas las clases de partículas, sean ellas de spin entero o semientero. La forma correcta de proceder no surge a partir de un requisito de consistencia ló-

gica de la teoría, sino de la evidencia experimental. No es así, en cambio, en la Teoría Cuántica Relativística de Campos. En el marco de la teoría relativística es preciso satisfacer dos requerimientos⁹ sin los cuales la teoría es absurda y se debe descartar:

- Los observables que corresponden a puntos del espacio-tiempo separados por una distancia tipo de espacio¹⁰ deben conmutar, de lo contrario se violaría el principio de causalidad.
- La energía del sistema debe ser semidefinida positiva.

Resulta entonces que si se intenta cuantificar un campo de spin entero por medio de relaciones de anticonmutación, se viola el *primer* requerimiento. Por lo tanto la consistencia de la teoría *obliga* a cuantificar los campos de *spin entero* por medio de *conmutadores*.

También se encuentra que si se trata de cuantificar un campo de spin semientero por medio de relaciones de conmutación, se viola el *segundo* requerimiento. Luego la consistencia de la teoría *obliga* a cuantificar los campos de *spin semientero* por medio de *anticonmutadores*.

Este es pues el origen de la relación entre spin y simetría de intercambio que mencionamos en el Capítulo 13.

El método de Hartree-Fok

En el Capítulo 12 mencionamos el método de Hartree del campo autoconsistente, con el cual se calculan en forma aproximada las energías de los estados fundamentales de átomos con varios electrones. Dicha aproximación tiene el defecto que las funciones de onda de varios electrones usadas en el cálculo no están antisimetrizadas como corresponde. Como un ejemplo de la aplicación de la segunda cuantificación vamos a presentar el método de Hartree-Fok¹¹, que es una mejora del método de Hartree, que toma en cuenta la correcta simetría de intercambio de las funciones de onda del sistema. El método se aplica no sólo a los átomos sino también a todo sistema de muchos Fermiones que interactúan de una forma cualquiera. Por lo tanto nuestra presentación será general, y sólo al final la particularizaremos al caso de los átomos.

Consideremos entonces un sistema de n Fermiones descrito por el Hamiltoniano

$$\hat{H} = \sum_{\alpha, \alpha'} h^{(1)\alpha'}_{\alpha} \hat{b}_{\alpha'}^{\dagger} \hat{b}_{\alpha} + \frac{1}{2} \sum_{\alpha, \beta, \alpha', \beta'} v^{(2)\alpha'\beta'}_{\alpha\beta} \hat{b}_{\alpha'}^{\dagger} \hat{b}_{\beta'}^{\dagger} \hat{b}_{\alpha} \hat{b}_{\beta} \quad (14.116)$$

Aquí $\hat{b}_{\alpha}$ y $\hat{b}_{\alpha}^{\dagger}$ ($\alpha = 1, 2, \dots$) son los operadores de aniquilación y creación de una sistema ortonormal completo de autoestados ϕ_{α} de una partícula, correspondientes a un observable cualquiera. Lo que nos interesa es encontrar una *nueva* base $\psi_1, \psi_2, \dots$ con operadores de aniquilación y creación $\hat{a}_k$ y $\hat{a}_k^{\dagger}$, tales que el estado

$$\Psi_n = \hat{a}_n^{\dagger} \hat{a}_{n-1}^{\dagger} \dots \hat{a}_2^{\dagger} \hat{a}_1^{\dagger} \Psi_0 \quad (14.117)$$

donde Ψ_0 es el estado de vacío, tenga la propiedad que el valor esperado de $\hat{H}$ sea el mínimo. Claramente, el estado fundamental verdadero de $\hat{H}$ no es Ψ_n , pero se puede pensar que será una combinación lineal de estados de n partículas en la cual Ψ_n es el término dominante, y en tal caso el valor medio de $\hat{H}$ en el estado Ψ_n :

⁹ Formulados por W. Pauli. En la teoría no relativística estos requerimientos se satisfacen siempre.

¹⁰ Esto es que ninguno de los dos puntos estén dentro del cono de luz del otro.

¹¹ El físico ruso Vladimir Fok fue quien desarrolló la extensión del método de Hartree de la que nos estamos ocupando aquí.

$$\langle \hat{H} \rangle = \frac{(\Psi_n, \hat{H} \Psi_n)}{(\Psi_n, \Psi_n)} \quad (14.118)$$

será una buena aproximación (por exceso) de la energía del estado fundamental del sistema.

En la nueva base tendremos que

$$\hat{H} = \sum_{k,l} h^{(1)}_l \hat{a}_k^\dagger \hat{a}_l + \frac{1}{2} \sum_{q,r,s,t} v^{(2)}_{st} \hat{a}_q^\dagger \hat{a}_r^\dagger \hat{a}_s \hat{a}_t \quad (14.119)$$

Si Ψ_n es tal que minimiza el valor esperado de $\hat{H}$, entonces para todo $\Psi'_n = \Psi_n + \delta\Psi_n$, donde $\delta\Psi_n$ es una variación *infinitesimal* cualquiera de Ψ_n se debe cumplir que

$$\delta\langle \hat{H} \rangle = \frac{(\Psi'_n, \hat{H} \Psi'_n)}{(\Psi'_n, \Psi'_n)} - \frac{(\Psi_n, \hat{H} \Psi_n)}{(\Psi_n, \Psi_n)} = 0 \quad (14.120)$$

Si despreciamos términos del segundo orden y usamos la propiedad Hermitiana de $\hat{H}$, la (14.120) nos lleva a la condición

$$(\Psi_n, \Psi_n)[(\hat{H} \Psi_n, \delta\Psi_n) + (\delta\Psi_n, \hat{H} \Psi_n)] - (\Psi_n, \hat{H} \Psi_n)[(\Psi_n, \delta\Psi_n) + (\delta\Psi_n, \Psi_n)] = 0 \quad (14.121)$$

Consideremos ahora la variación de Ψ_n . Evidentemente $\delta\Psi_n$ se obtiene de resultados de un cambio desde la base $\psi_1, \psi_2, \dots$ a la base $\psi_1', \psi_2', \dots$, en la cual los nuevos operadores de creación $\hat{a}_k'^\dagger$ están dados por

$$\hat{a}_k'^\dagger = \hat{a}_k^\dagger + \delta\hat{a}_k^\dagger = \sum_j \hat{a}_j^\dagger (\delta_{jk} + i\varepsilon_{jk}) \quad (14.122)$$

donde las ε_{jk} son cantidades pequeñas ($|\varepsilon_{jk}| \ll 1$) y el factor i se introdujo para asegurar que la transformación $\psi_1, \psi_2, \dots \rightarrow \psi_1', \psi_2', \dots$ sea unitaria. De la (14.122) resulta

$$\delta\hat{a}_k^\dagger = i \sum_j \hat{a}_j^\dagger \varepsilon_{jk} \quad (14.123)$$

La variación general $\delta\Psi_n$ se puede expresar entonces como una combinación lineal de variaciones independientes

$$\delta\Psi_{jk} = \varepsilon_{jk} \hat{a}_j^\dagger \hat{a}_k \Psi_n \quad (14.124)$$

donde se debe tener presente que $\hat{a}_k$ aniquila un Fermión de uno de los estados *ocupados* $\psi_1, \dots, \psi_n$ y $\hat{a}_j^\dagger$ crea una partícula en uno de los estados previamente *desocupados* $\psi_{n+1}, \dots$. La unitariedad de la transformación (14.122) requiere que los coeficientes ε_{jk} formen una matriz Hermitiana, pero esto no trae restricciones puesto que por el principio de exclusión resulta que

$$\delta\Psi_{kj} = \varepsilon_{kj} \hat{a}_k^\dagger \hat{a}_j \Psi_n = 0 \quad (14.125)$$

dado que el estado $j = n+1, \dots$ está vacío y el estado $k = 1, \dots, n$ está ocupado. Luego el requerimiento $\varepsilon_{kj} = \varepsilon_{jk}^*$ no restringe las posibles variaciones, cuya independencia está asegurada. Las

variaciones con $j = k$ no cambian el estado y por lo tanto son irrelevantes. En consecuencia basta considerar las variaciones $\delta\Psi_n$ de la forma (14.124) que son ortogonales al “mejor” estado Ψ_n de la forma (14.117).

Volviendo ahora a la condición (14.121), vemos que $(\Psi_n, \delta\Psi_n)$ y $(\delta\Psi_n, \Psi_n)$ son combinaciones lineales de términos de la forma $(\Psi_n, \delta\Psi_{jk})$ y $(\delta\Psi_{jk}, \Psi_n)$, y son nulos puesto que nuestras variaciones son siempre ortogonales a $\delta\Psi_n$. Además $(\Psi_n, \Psi_n) = 1$, y por lo tanto la (14.121) se reduce a las siguientes condiciones

$$(\hat{H}\Psi_n, \delta\Psi_{jk}) + (\delta\Psi_{jk}, \hat{H}\Psi_n) = 0 \quad \text{con } j = n+1, \dots \quad \text{y } k = 1, \dots, n \quad (14.126)$$

Si ahora usamos la expresión (14.119) de $\hat{H}$ en las (14.126) es fácil ver (con un poco de orden y paciencia) que se obtienen las condiciones

$$h^{(1)}_k + v_{\text{ef } k}^{(1)} = 0 \quad \text{con } j = n+1, \dots \quad \text{y } k = 1, \dots, n \quad (14.127)$$

donde los

$$v_{\text{ef } k}^{(1)} = \sum_{t=1}^n (v^{(2)}_{kt} - v^{(2)}_{tk}) \quad (14.128)$$

se pueden pensar como los elementos de matriz de un “potencial efectivo” $v_{\text{ef}}^{(1)}$ de una partícula debido a sus interacciones con las demás. En las (14.127) y (14.128) los índices k y t recorren solamente los estados *ocupados*, mientras que j recorre únicamente los estados *desocupados*.

Conviene introducir el *Hamiltoniano efectivo* de una partícula como el operador

$$h_{\text{HF}} = h^{(1)} + v_{\text{ef}}^{(1)} = h^{(1)} + \sum_{t=1}^n (v^{(2)}_{kt} - v^{(2)}_{tk}) \quad (14.129)$$

En términos de h_{HF} las condiciones (14.127) se escriben en la forma compacta

$$(\psi_j, h_{\text{HF}}\psi_k) = 0 \quad (14.130)$$

Las condiciones (14.130) se satisfacen si las ψ_k y las ψ_j son autofunciones de h_{HF} , es decir si son soluciones del problema de autovalores

$$h_{\text{HF}}\psi_k = e_k\psi_k \quad (14.131)$$

Las (14.131) son las *ecuaciones de Hartree-Fok*. De las mismas resulta de inmediato que

$$(\psi_k, h_{\text{HF}}\psi_k) = h^{(1)}_k + v_{\text{ef } k}^{(1)} = h^{(1)}_k + \sum_{t=1}^n (v^{(2)}_{kt} - v^{(2)}_{tk}) = e_k\psi_k \quad (14.132)$$

Se debe notar que los estados ocupados en (14.127), (14.129) y (14.132) *no son arbitrarios*, pues se deben elegir entre los autoestados (14.131) de manera de minimizar el valor esperado de $\hat{H}$. Frecuentemente la mejor elección se obtiene usando las autofunciones correspondientes a los n autovalores e_k más bajos.

Es interesante ver que el mínimo del valor esperado de $\hat{H}$ *no es igual* a la suma de las energías e_k de partícula individual de los estados ocupados. En la aproximación de Hartree-Fok la energía del estado fundamental está dada por

$$\begin{aligned} E_n &= \langle \hat{H} \rangle = (\Psi_n, \hat{H} \Psi_n) = \sum_{k=1}^n h^{(1)}_k + \frac{1}{2} \sum_{i,k=1}^n (v^{(2)}_{ik} - v^{(2)}_{ki}) \\ &= \frac{1}{2} \sum_{k=1}^n (e_k + h^{(1)}_k) = \sum_{k=1}^n e_k - \frac{1}{2} \sum_{i,k=1}^n (v^{(2)}_{ik} - v^{(2)}_{ki}) \end{aligned} \quad (14.133)$$

Este resultado es razonable pues al sumar los e_k estamos contando dos veces la contribución que proviene de las interacciones entre las partículas.

A los efectos prácticos del cálculo de los elementos de matriz del potencial efectivo $v_{\text{ef}}^{(1)}$, es útil escribirlos en términos del sistema completo original de autoestados ϕ_α ($\alpha = 1, 2, \dots$), en el cual

$$\psi_k = \sum_{\alpha} \phi_{\alpha}(\phi_{\alpha}, \psi_k) \quad (14.134)$$

si en la base original $\hat{b}_{\alpha}^{\dagger}$ la interacción es diagonal, o sea si

$$v^{(2)}_{\alpha\beta}{}^{\alpha'\beta'} = v_{\alpha\beta} \delta_{\alpha'\alpha} \delta_{\beta'\beta} \quad (14.135)$$

es fácil ver que se obtiene

$$\begin{aligned} v_{\text{ef } k}^{(1)j} &= \sum_{t=1}^n (v^{(2)}_{kt}{}^{jt} - v^{(2)}_{tk}{}^{jt}) \\ &= \sum_{t=1}^n \sum_{\alpha, \beta} (\psi_j, \phi_{\alpha})(\psi_t, \phi_{\beta}) v_{\alpha\beta}(\phi_{\beta}, \psi_t)(\phi_{\alpha}, \psi_k) \\ &\quad - \sum_{t=1}^n \sum_{\alpha, \beta} (\psi_j, \phi_{\alpha})(\psi_t, \phi_{\beta}) v_{\alpha\beta}(\phi_{\alpha}, \psi_t)(\phi_{\beta}, \psi_k) \end{aligned} \quad (14.136)$$

Las ecuaciones de Hartree-Fok (14.131) se escriben entonces como

$$\begin{aligned} h_{\text{HF}} \psi_k &= h^{(1)} \psi_k + \sum_{t=1}^n \sum_{\alpha, \beta} \phi_{\alpha}(\psi_t, \phi_{\beta}) v_{\alpha\beta}(\phi_{\beta}, \psi_t)(\phi_{\alpha}, \psi_k) \\ &\quad - \sum_{t=1}^n \sum_{\alpha, \beta} \phi_{\alpha}(\psi_t, \phi_{\beta}) v_{\alpha\beta}(\phi_{\alpha}, \psi_t)(\phi_{\beta}, \psi_k) = e_k \psi_k \end{aligned} \quad (14.137)$$

y multiplicando escalarmente por ϕ_{α} resulta

$$\begin{aligned} \sum_{\beta} h^{(1)}_{\alpha}{}^{\beta}(\phi_{\beta}, \psi_k) + \sum_{t=1}^n \sum_{\beta} (\psi_t, \phi_{\beta}) v_{\alpha\beta}(\phi_{\beta}, \psi_t)(\phi_{\alpha}, \psi_k) \\ - \sum_{t=1}^n \sum_{\beta} (\psi_t, \phi_{\beta}) v_{\alpha\beta}(\phi_{\alpha}, \psi_t)(\phi_{\beta}, \psi_k) = e_k(\phi_{\alpha}, \psi_k) \end{aligned} \quad (14.138)$$

El aspecto de las ecuaciones de Hartree-Fok es engañosamente simple, y la tarea necesaria para resolverlas dista mucho de ser trivial. Las (14.131) aparentan ser un problema común de autovalores, pero los elementos de matriz de $v_{\text{ef}}^{(1)}$, que intervienen en el Hamiltoniano efectivo h_{HF} no se pueden calcular si no se conocen previamente las n autofunciones apropiadas ψ_k , soluciones de las mismas (14.131). Se trata por lo tanto de un sistema de n ecuaciones integrodiferenciales no lineales que se deben resolver simultáneamente en forma iterativa, a partir de un conjunto de n estados “de prueba” ocupados ψ_t . Usando estos estados de prueba se comienza el proceso, calculando los elementos de matriz (14.136) y luego se resuelven las ecuaciones de Hartree-Fok (14.137). Si se tiene la suerte que la elección de las ψ_t fue acertada, habrá n de las soluciones de (14.137) que no diferirán mucho de las funciones de prueba elegidas. Si, como es más probable, las soluciones de las ecuaciones de Hartree-Fok no reproducen las funciones de prueba iniciales, se usan para la siguiente iteración las autofunciones correspondientes a los n autovalores e_k más bajos para calcular nuevamente los elementos de matriz de $v_{\text{ef}}^{(1)}$, y se repite el procedimiento hasta hallar un conjunto de n soluciones autoconsistentes. Generalmente se cuenta con buenos puntos de partida y por lo tanto las iteraciones convergen con razonable rapidez.

Aplicación a un átomo con n electrones

En este caso tenemos que

$$h^{(1)} = \frac{\mathbf{p}^2}{2m} - \frac{Ze^2}{r} \quad (14.139)$$

El potencial de interacción $v^{(2)}$ es diagonal en la representación coordenadas y se escribe

$$v^{(2)}(\mathbf{r}, \sigma; \mathbf{r}', \sigma') = \frac{e^2}{|\mathbf{r} - \mathbf{r}'|} \quad (14.140)$$

Las ecuaciones de Hartree-Fok se escriben entonces en la forma

$$\begin{aligned} \frac{\hbar^2}{2m} \nabla^2 \psi_k(\mathbf{r}, \sigma) - \frac{Ze^2}{r} \psi_k(\mathbf{r}, \sigma) \\ + e^2 \sum_{t=1}^n \sum_{\sigma} \int \psi_t^*(\mathbf{r}', \sigma') \frac{1}{|\mathbf{r} - \mathbf{r}'|} \psi_t(\mathbf{r}', \sigma') \psi_k(\mathbf{r}, \sigma) d\mathbf{r}' \\ - e^2 \sum_{t=1}^n \sum_{\sigma} \int \psi_t^*(\mathbf{r}', \sigma') \frac{1}{|\mathbf{r} - \mathbf{r}'|} \psi_k(\mathbf{r}', \sigma') \psi_t(\mathbf{r}, \sigma) d\mathbf{r}' = e_k \psi_k(\mathbf{r}, \sigma) \end{aligned} \quad (14.141)$$

con $k = 1, \dots, n$. Este sistema de ecuaciones integrodiferenciales acopladas es la aplicación más conocida de la teoría de Hartree-Fok. Si se ignora el último término del miembro izquierdo de la (14.141), que se debe a los elementos de matriz de intercambio de la interacción y se omite el término $t = k$ en la suma sobre los elementos de matriz directos¹², las ecuaciones que quedan son

¹² Ese término representaría la interacción de un electrón consigo mismo. Obviamente, los términos con $t = k$ de la suma directa y la suma de intercambio son idénticos y se cancelan, de modo que no tienen efecto.

$$\begin{aligned}
& \frac{\hbar^2}{2m} \nabla^2 \psi_k(\mathbf{r}, \sigma) - \frac{Ze^2}{r} \psi_k(\mathbf{r}, \sigma) \\
& + e^2 \sum_{t=1, t \neq k}^n \sum_{\sigma} \int \psi_t^*(\mathbf{r}', \sigma') \frac{1}{|\mathbf{r} - \mathbf{r}'|} \psi_t(\mathbf{r}', \sigma') \psi_k(\mathbf{r}, \sigma) d\mathbf{r}' = e_k \psi_k(\mathbf{r}, \sigma)
\end{aligned} \tag{14.142}$$

que no son otra cosa que las *ecuaciones de Hartree*. Antes del advenimiento de las veloces computadoras actuales, estas ecuaciones, más sencillas que las (14.141), tenían mayor interés práctico que las ecuaciones más precisas de Hartree-Fok.

Este ejemplo muestra claramente las virtudes de la segunda cuantificación, su concisión y su elegancia. Piense el lector, en efecto, que sería hartamente engorroso (aunque no imposible) desarrollar los argumentos de esta Sección, si tuviéramos que escribir Ψ_n y $\delta\Psi_n$ explícitamente en forma de determinantes de Slater. Todo ese tedio se evita gracias a nuestro formalismo.

Existen otras importantes e interesantes aplicaciones de la segunda cuantificación a problemas de muchos cuerpos, entre la cuales podemos citar el tratamiento de las interacciones de apareamiento en Física de Sólidos (el método de Bardeen, Cooper y Schrieffer para estudiar los pares de Cooper, que se relacionan con el fenómeno de la superconductividad) y en Física Nuclear (importantes para interpretar el origen de ciertos estados de excitación colectiva del núcleo atómico), la teoría de las ondas de spin y de los magnones, la teoría de perturbaciones en sistemas de muchos cuerpos, etc..

15. LAS ESTADÍSTICAS CUÁNTICAS

El límite clásico

Hemos visto en el Capítulo 13 que cuando las funciones de onda de dos partículas idénticas no se solapan, éstas se comportan como partículas clásicas, esto es *distinguibles*. Esto ocurre porque los efectos cuánticos de la *indistinguibilidad* se ponen de manifiesto solamente cuando hay solapamiento entre las funciones de onda. Consideremos entonces un sistema de n partículas que no interactúan entre sí y que ocupan un volumen V , y supongamos que están en equilibrio térmico a una temperatura T . De acuerdo con el Teorema de Equipartición de la Mecánica Estadística clásica¹, la energía cinética media de traslación de cada partícula de masa m es entonces

$$\bar{\varepsilon} = \frac{3}{2}kT \quad (15.1)$$

donde k es la constante de Boltzmann, y el valor medio del impulso de una partícula es

$$\bar{p} = \sqrt{2m\bar{\varepsilon}} = \sqrt{3mkT} \quad (15.2)$$

Luego su longitud de onda de Broglie, que da la medida de la extensión espacial del paquete de ondas que la describe, vale

$$\lambda_B = \frac{h}{\bar{p}} = \frac{h}{\sqrt{3mkT}} \quad (15.3)$$

Por otra parte, la distancia media ℓ entre las partículas está dada por

$$\ell = (V/n)^{1/3} \quad (15.4)$$

Por lo tanto, las partículas se podrán considerar distinguibles si se cumple que

$$\lambda_B \ll \ell \quad (15.5)$$

esto es, si

$$\frac{nh^3}{V(3mkT)^{3/2}} \ll 1 \quad (15.6)$$

Esta es la condición de validez de la Mecánica Estadística Clásica. Si se cumple la (15.6), las partículas se comportan clásicamente, y se puede aplicar el Teorema de Equipartición. Esto es lo que ocurre con los gases en las condiciones habituales en la naturaleza y en el laboratorio. Sin embargo, cuando la temperatura es muy baja y/o cuando la densidad es muy elevada, la (15.6) no se cumple; entonces los resultados clásicos no son válidos y se manifiestan interesantes efectos cuánticos. Examinemos un poco más la condición (15.6), para entender mejor su significado. La probabilidad que una partícula se encuentre en un particular estado de energía de traslación ε_i está dada por la distribución de Boltzmann

¹ Ver el Capítulo 17 de Termodinámica e Introducción a la Mecánica Estadística.

$$p_i = \frac{1}{Z_{tr}} e^{-\varepsilon_i / kT} \quad (15.7)$$

donde Z_{tr} es la función de partición traslacional

$$Z_{tr} = \sum_i e^{-\varepsilon_i / kT} \quad (15.8)$$

que para una partícula clásica² está dada por

$$Z_{tr} = \frac{V}{h^3} (2\pi mkT)^{3/2} \quad (15.9)$$

Si nuestro sistema consta de n partículas el número de ocupación *medio* $\bar{n}_i$ de cada estado ε_i es

$$\bar{n}_i = np_i = \frac{nh^3}{V(2\pi mkT)^{3/2}} e^{-\varepsilon_i / kT} \quad (15.10)$$

Por lo tanto, la condición (15.6) implica que

$$\bar{n}_i \ll 1 \quad (15.11)$$

Luego en el límite clásico la gran mayoría de los estados de una partícula están vacíos, y sólo unos pocos están ocupados por una sola partícula. La probabilidad que un estado esté ocupado por dos o más partículas es insignificante. En este límite desaparece la diferencia entre Bosones y Fermiones: ambos se comportan como partículas clásicas, idénticas pero distinguibles.

Cuando no se cumple la (15.6) (o lo que es lo mismo la (15.11)) hay que tomar en cuenta los efectos cuánticos, y eso es lo que haremos en este Capítulo. Pero primero queremos señalar dónde está la falla de la Estadística Clásica. La distribución de Boltzmann (15.7) es correcta, pues deriva de consideraciones generales sobre el equilibrio de un sistema bajo determinadas restricciones. El inconveniente está en la función de partición (15.9), que se calculó suponiendo que la probabilidad que una partícula ocupe un determinado estado no depende de si otras partículas están ocupando (o no) ese mismo estado. Sabemos que esto no es cierto, pues debido a la indistinguibilidad cuántica de las partículas idénticas hay correlaciones entre ellas, de distinto carácter según sean Bosones o Fermiones. Por este motivo las distribuciones que se obtienen para Bosones y Fermiones son diferentes. Pero cuando las condiciones del sistema son tales que las correlaciones cuánticas se tornan irrelevantes, ambas distribuciones coinciden con la distribución clásica de Maxwell-Boltzmann.

La función de partición de un sistema de partículas idénticas sin interacción

Para deducir las distribuciones apropiadas para los sistemas cuánticos conviene partir de la *distribución gran canónica* o *distribución de Gibbs*³, que se obtiene de considerar un sistema de volumen fijo V , sumergido en un baño calorífico a la temperatura T , y con el cual puede intercambiar partículas. En estas condiciones la energía E y el número de partículas n del sistema

² La expresión (15.9) se obtiene en el Capítulo 17 de *Termodinámica e Introducción a la Mecánica Estadística*.

³ Ver el Capítulo 16 de *Termodinámica e Introducción a la Mecánica Estadística*.

fluctúan, pero en general dichas fluctuaciones son despreciables para un sistema macroscópico. El sistema se describe entonces en términos de T , V y del potencial químico μ . Usando la notación $\beta = 1/kT$, la distribución de probabilidad correctamente normalizada de encontrar el sistema en un estado con n partículas cuya energía es $E_{n,r}$ está dada por

$$p_{n,r} = \frac{e^{\beta(\mu n - E_{n,r})}}{Z} \quad (15.12)$$

donde Z es la *gran función de partición* del sistema, dada por

$$Z \equiv Z(T, V, \mu) \equiv \sum_{n,r} e^{\beta(\mu n - E_{n,r})} \quad (15.13)$$

A partir de Z se pueden obtener todas las propiedades termodinámicas del sistema. Consideraremos solamente un *gas perfecto cuántico*, es decir un sistema de partículas *sin interacciones*. En este caso el estado del sistema está especificado por los números de ocupación

$$n_1, n_2, \dots \quad (15.14)$$

de los estados de partícula individual $\psi_1(\xi), \psi_2(\xi), \dots$. Los n_i son enteros no negativos tales que

$$\sum_i n_i = n \quad (15.15)$$

Vamos a suponer que los estados ψ_i están ordenados por energía creciente⁴, esto es

$$\varepsilon_1 \leq \varepsilon_2 \leq \dots \leq \varepsilon_i \leq \dots \quad (15.16)$$

Por lo tanto

$$Z = \sum_{n_1, n_2, \dots} e^{\beta[\mu(n_1 + n_2 + \dots) - (n_1 \varepsilon_1 + n_2 \varepsilon_2 + \dots)]} = \prod_i Z_i \quad (15.17)$$

donde hemos escrito

$$Z_i = \sum_{n_i} e^{\beta(\mu - \varepsilon_i)n_i} \quad (15.18)$$

y la sumatoria en Z_i se extiende a todos los valores posibles de n_i . Aquí vemos la ventaja de partir de la distribución de Gibbs, pues en el ensemble gran canónico los n_i no están restringidos por la condición (15.15), y eso permite factorizar Z , porque la distribución estadística para cada estado de una partícula, dada por la (15.18), no depende de lo que ocurre con los demás estados⁵. Esta simplificación tiene su precio, sin embargo: tuvimos que introducir el potencial químico, que no conocemos de antemano y que hay que determinar *a posteriori* imponiendo la condición

⁴ En general debido a la degeneración de los niveles de energía hay muchos términos iguales en la sucesión (15.15).

⁵ Este problema no existe en la estadística de Maxwell-Boltzmann pues al no considerar las correlaciones entre las partículas, la función de partición del ensemble canónico se factoriza sin dificultad.

$$\sum_i \bar{n}_i = \bar{n} \quad (15.19)$$

es decir, la condición (15.15), pero restringida a los valores medios.

Las distribuciones de Bose-Einstein y de Fermi-Dirac

Las fórmulas anteriores valen tanto para sistemas de Bosones como de Fermiones. En lo que sigue vamos a indicar con B las fórmulas que valen para los Bosones y con F las que corresponden a Fermiones. En un sistema de Bosones n_i puede tomar cualquier valor entero positivo a partir de 0. Por lo tanto, sumando la serie geométrica que resulta de (15.18) obtenemos

$$Z_i = \frac{1}{1 - e^{\beta(\mu - \varepsilon_i)}} \quad (n_i = 0, 1, 2, \dots) \quad \text{B} \quad (15.20)$$

En un sistema de Fermiones n_i toma solamente los valores 0 y 1 y la (15.18) se reduce a

$$Z_i = 1 + e^{\beta(\mu - \varepsilon_i)} \quad (n_i = 0, 1) \quad \text{F} \quad (15.21)$$

La probabilidad que el número de ocupación del estado i tenga el valor n_i está dada por

$$p_i(n_i) = \frac{e^{\beta(\mu - \varepsilon_i)n_i}}{Z_i} \quad (15.22)$$

y entonces el número de ocupación *medio* del estado i está dado por

$$\bar{n}_i = \sum_{\text{todo } n_i} n_i p_i(n_i) = kT \left(\frac{\partial \ln Z_i}{\partial \mu} \right)_{T,V} \quad (15.23)$$

Usando la (15.20) y la (15.23) se obtiene la distribución de *Bose-Einstein* para los Bosones:

$$\bar{n}_i = \frac{1}{e^{\beta(\varepsilon_i - \mu)} - 1} \quad \text{B} \quad (15.24)$$

donde el potencial químico se determina pidiendo que

$$n = \sum_i \bar{n}_i = \sum_i \frac{1}{e^{\beta(\varepsilon_i - \mu)} - 1} \quad \text{B} \quad (15.25)$$

En la (15.25) pusimos $n = \bar{n}$ dado que las fluctuaciones se pueden suponer despreciables.

Del mismo modo usando la (15.21) y la (15.23) se obtiene para Fermiones la distribución de *Fermi-Dirac*:

$$\bar{n}_i = \frac{1}{e^{\beta(\varepsilon_i - \mu)} + 1} \quad \text{F} \quad (15.26)$$

y el potencial químico se encuentra a partir de la condición

$$n = \sum_i \frac{1}{e^{\beta(\varepsilon_i - \mu)} + 1} \quad \text{F} \quad (15.27)$$

A continuación vamos a estudiar el significado de estos resultados.

El gas perfecto de Bosones

De acuerdo con los resultados anteriores tenemos que para un gas de Bosones

$$\bar{n}_i = \frac{1}{e^{\beta(\varepsilon_i - \mu)} - 1}, \quad n = \sum_i \frac{1}{e^{\beta(\varepsilon_i - \mu)} - 1} \quad (15.28)$$

Puesto que $\bar{n}_i \geq 0$ siempre, se debe cumplir que $\varepsilon_i > \mu$ para todo i . Por otra parte, el estado de más baja energía de una partícula libre tiene $\varepsilon_1 = 0$. Por lo tanto el potencial químico de un gas de Bosones es siempre *negativo*.

La temperatura de degeneración

Introduciendo la densidad de estados por unidad de intervalo de energía $f(\varepsilon)$, la sumatoria de la segunda de las (15.28) se puede escribir como una integral

$$\sum_i \frac{1}{e^{\beta(\varepsilon_i - \mu)} - 1} = \int_0^\infty \frac{f(\varepsilon)d\varepsilon}{e^{\beta(\varepsilon - \mu)} - 1} = 2\pi V \left(\frac{2m}{h^2} \right)^{3/2} \int_0^\infty \frac{\varepsilon^{1/2} d\varepsilon}{e^{\beta(\varepsilon - \mu)} - 1} \quad (15.29)$$

donde usamos la expresión (9.19) de $f(\varepsilon)$, que se dedujo contando los estados estacionarios de una partícula libre de masa m que se mueve en una caja de volumen V , y hemos supuesto Bosones de spin nulo (si $s \neq 0$ hay que multiplicar (15.29) por $2s + 1$). Tenemos por lo tanto

$$\frac{n}{V} = 2\pi \left(\frac{2m}{h^2} \right)^{3/2} \int_0^\infty \frac{\varepsilon^{1/2} d\varepsilon}{e^{\beta(\varepsilon - \mu)} - 1} \quad (15.30)$$

Como el miembro izquierdo de la (15.30) no depende de T , $\mu(n, T)$ debe ser tal que la integral en (15.30) sea independiente de T . Supongamos variar la temperatura del gas, con n/V constante. Es evidente que a medida que T disminuye, $|\mu|$ debe disminuir, hasta llegar a $\mu = 0$ para una temperatura mínima T_c determinada por la condición

$$\frac{n}{V} = 2\pi \left(\frac{2m}{h^2} \right)^{3/2} \int_0^\infty \frac{\varepsilon^{1/2} d\varepsilon}{e^{\varepsilon/kT_c} - 1} \quad (15.31)$$

T_c se denomina *temperatura de degeneración*, o de *condensación*, por motivos que veremos en breve. Para evaluar la integral hacemos el cambio de variable $\varepsilon = kT_c z$ y queda

$$\frac{n}{V} = \left(\frac{2\pi mkT_c}{h^2} \right)^{3/2} \frac{2}{\sqrt{\pi}} \int_0^\infty \frac{z^{1/2} dz}{e^z - 1} \quad (15.32)$$

De tablas se encuentra que

$$\frac{2}{\sqrt{\pi}} \int_0^{\infty} \frac{z^{1/2} dz}{e^z - 1} = \zeta(3/2) \approx 2.61239 \quad (15.33)$$

donde con $\zeta(q)$ indicamos la función Zeta de Riemann, definida por

$$\zeta(q) = \sum_{k=1}^{\infty} \frac{1}{k^q}, \quad \text{Re}(q) > 1 \quad (15.34)$$

Resulta entonces que

$$T_c = \frac{1}{2\pi mk} \left(\frac{nh^3}{\zeta(3/2)V} \right)^{2/3} \quad (15.35)$$

De la (15.35) obtenemos que

$$\frac{nh^3}{V(3mkT_c)^{3/2}} = \zeta(3/2) \left(\frac{2\pi}{3} \right)^{3/2} \approx 7.92 \quad (15.36)$$

Comparando (15.36) con la condición de validez (15.6) de la estadística clásica vemos que la anulación de μ del gas de Bosones a la temperatura T_c es un fenómeno netamente cuántico.

La condensación de Bose-Einstein

Puesto que $\mu < 0$ siempre, parecería que no se puede enfriar a densidad constante un gas de Bosones hasta $T \leq T_c$, lo que no es cierto pues la ec. (15.29) se cumple para $T > T_c$ pero no para $T \leq T_c$. El origen del problema es que en la (15.29) reemplazamos la suma discreta sobre los estados de una sola partícula por una integral. A medida que la temperatura del gas disminuye y se acerca a T_c el número de ocupación del estado ψ_1 , cuya energía es nula, comienza a aumentar rápidamente. Pero precisamente este estado *no* está tomado en cuenta en la (15.29) porque al reemplazar la suma por una integral, el factor $\epsilon^{1/2}$ del integrando hace que se le asigne un *peso nulo*. Para temperaturas más altas, eso no trae aparejado un error significativo. Pero cuando la temperatura es muy baja no se puede omitir ψ_1 : lo correcto es conservar explícitamente su contribución, y reemplazar los demás términos por la integral⁶. En lugar de (15.30) escribiremos

$$n = \frac{1}{e^{-\beta\mu} - 1} + 2\pi V \left(\frac{2m}{h^2} \right)^{3/2} \int_0^{\infty} \frac{\epsilon^{1/2} d\epsilon}{e^{\beta(\epsilon-\mu)} - 1} \quad (15.37)$$

En esta ecuación

$$n_1 = \frac{1}{e^{-\beta\mu} - 1} \quad (15.38)$$

es el número de partículas en el estado ψ_1 , con energía nula y cantidad de movimiento nulo, y

⁶ Se puede mostrar que siempre que n sea muy grande es suficiente tratar de manera especial solamente el primer término de la suma, y reemplazar todos los demás por la integral.

$$n_{\varepsilon>0} = 2\pi V \left(\frac{2m}{h^2} \right)^{3/2} \int_0^{\infty} \frac{\varepsilon^{1/2} d\varepsilon}{e^{\beta(\varepsilon-\mu)} - 1} \quad (15.39)$$

es el número de partículas cuya energía y cantidad de movimiento no son nulos. Para $T > T_c$, tendremos que $n_1 \ll n$ y entonces lo podemos despreciar. Por lo tanto el potencial químico está dado con excelente aproximación por la ec. (15.30). Para $T \leq T_c$, el potencial químico es negativo y muy próximo a 0, pero nunca se anula exactamente. Si $T \leq T_c$, podemos usar la (15.39) con $\mu = 0$ para calcular $n_{\varepsilon>0}$ con buena aproximación y se obtiene

$$n_{\varepsilon>0} \cong 2\pi V \left(\frac{2m}{h^2} \right)^{3/2} \int_0^{\infty} \frac{\varepsilon^{1/2} d\varepsilon}{e^{\beta\varepsilon} - 1} = V \left(\frac{2\pi mkT}{h^2} \right)^{3/2} \frac{2}{\sqrt{\pi}} \int_0^{\infty} \frac{z^{1/2} dz}{e^z - 1} \quad (15.40)$$

Dividiendo la (15.40) por la (15.32) obtenemos entonces una sencilla fórmula que nos da la fracción de partículas cuya energía y cantidad de movimiento no son nulos:

$$\frac{n_{\varepsilon>0}}{n} = \left(\frac{T}{T_c} \right)^{3/2} \quad (T < T_c) \quad (15.41)$$

Las partículas restantes están en el estado fundamental ψ_1 , y representan una fracción

$$\frac{n_1}{n} = 1 - \left(\frac{T}{T_c} \right)^{3/2} \quad (T < T_c) \quad (15.42)$$

del número total (Fig. 15.1). En síntesis, para $T > T_c$ la fracción de partículas en el estado fundamental ψ_1 es despreciable. Cuando $T < T_c$, n_1/n crece rápidamente al disminuir T y tiende a 1 para $T \rightarrow 0$. Estas partículas tienen energía nula y cantidad de movimiento nula. Luego no contribuyen a la presión⁷, y tampoco a la viscosidad del gas. El fenómeno de concentración de las partículas en el estado fundamental se llama *condensación de Bose-Einstein*. Un gas de Bosones se dice *degenerado*⁸ cuando presenta condensación de Bose-Einstein.

Un gas de Bosones degenerado tiene cierta analogía con un vapor clásico en equilibrio con su líquido (de allí el término “condensación”). Por ejemplo, se puede ver que cuando $T < T_c$ la presión de un gas de Bosones depende solamente de T y es independiente del volumen, de manera semejante a lo que ocurre en el equilibrio líquido-vapor. Se debe notar, sin embargo, que en la condensación de Bose-Einstein no hay una separación de fases *en el espacio*. Pero se puede interpretar que hay una separación de fases *en el espacio de los impulsos*, pues las partículas con energía y cantidad de movimiento nulas (la “fase condensada”) y las partículas con energía y cantidad de movimiento no nulas (la “fase vapor”) tienen propiedades muy diferentes.

⁷ En el cero absoluto la presión del gas de Bosones es nula. También es nula la entropía, de acuerdo con la Tercera Ley de la Termodinámica.

⁸ El término “degenerado” se usa aquí en un sentido diferente del que tiene cuando se habla de la degeneración de un nivel.

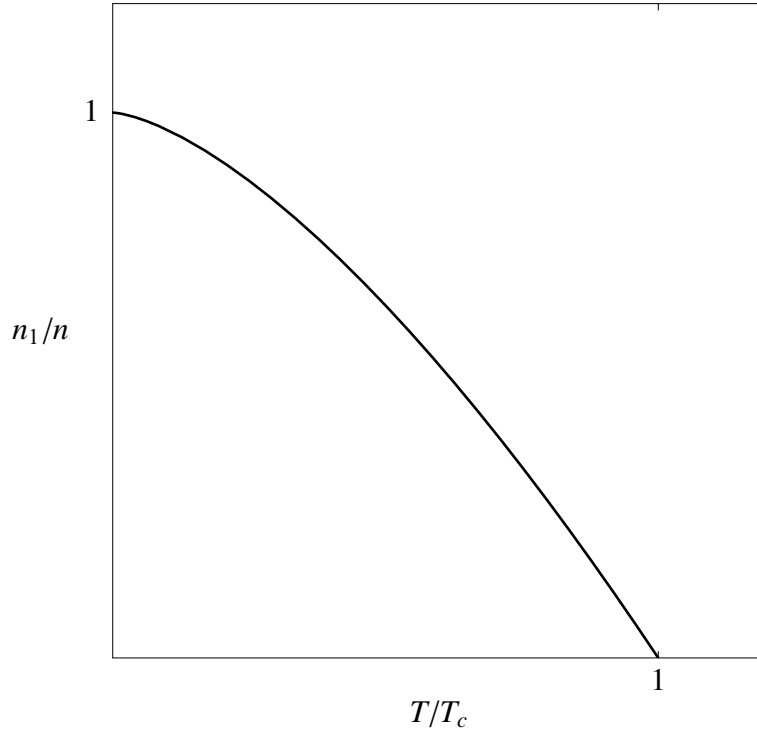


Fig. 15.1. Condensación de Bose-Einstein: cuando $T < T_c$ la fracción de partículas en el estado fundamental crece rápidamente al disminuir T , y tiende a 1 para $T \rightarrow 0$.

Potencial químico de un gas de Bosones

Para calcular las propiedades del gas de Bosones es preciso conocer el potencial químico. Para ello hay que resolver para $\mu(n, T)$ la ec. (15.37), que con el cambio $z = \beta\epsilon$ se escribe como

$$n = \frac{1}{e^{-\beta\mu} - 1} + 2\pi V \left(\frac{2mkT}{h^2} \right)^{3/2} \int_0^\infty \frac{z^{1/2} dz}{e^{z-\beta\mu} - 1} \quad (15.43)$$

La integral en (15.43) es un caso particular de las integrales de la forma general

$$\int_0^\infty \frac{z^{q-1} dz}{t^{-1} e^z \pm 1} = \mp \Gamma(q) \text{Li}_q(\mp t) \quad , \quad \text{Re}(q) > 0 \quad (15.44)$$

Aquí $\text{Li}_q(t)$ es la *función polilogarítmica*, definida por la serie

$$\text{Li}_q(t) = \sum_{k=1}^\infty \frac{t^k}{k^q} \quad (15.45)$$

Algunas propiedades útiles de $\text{Li}_q(t)$ son:

$$\text{Li}_q(1) = \zeta(q) \quad (q > 1) \quad , \quad \text{Li}_q(t) = t + O(t^2) \quad , \quad \frac{d}{dt} \text{Li}_q(t) = \frac{\text{Li}_{q-1}(t)}{t} \quad (15.46)$$

Notar que $\text{Li}_q(t)$ diverge en $t = 1$ para $q \leq 1$.

Volviendo al cálculo de μ , vemos que la (15.43) se convierte en

$$n = \frac{1}{e^{-\beta\mu} - 1} + V \left(\frac{2\pi mkT}{h^2} \right)^{3/2} \text{Li}_{3/2}(e^{\beta\mu}) \quad (15.47)$$

la cual, recordando que

$$\left(\frac{2\pi mkT_c}{h^2} \right)^{3/2} = \frac{n}{V\zeta(3/2)} \quad (15.48)$$

se escribe finalmente como

$$1 = \frac{n_1}{n} + \left(\frac{T}{T_c} \right)^{3/2} \frac{\text{Li}_{3/2}(e^{\beta\mu})}{\zeta(3/2)} \quad , \quad n_1 = \frac{1}{e^{-\beta\mu} - 1} \quad (15.49)$$

Para un sistema macroscópico n es un número enormemente grande (del orden del número de Avogadro), de modo que n_1/n es despreciable frente a la unidad, para todo $T \geq T_c$; por lo tanto, tenemos que, con buena aproximación

$$1 = \left(\frac{T}{T_c} \right)^{3/2} \frac{\text{Li}_{3/2}(e^{\beta\mu})}{\zeta(3/2)} \quad (T > T_c) \quad (15.50)$$

Invirtiendo esta relación obtenemos finalmente el potencial químico (Fig. 15.2). Se puede ver (no damos aquí los detalles) que en el intervalo $0 < T \leq T_c$, μ es negativo y no se anula nunca (salvo para $T = 0$); su valor absoluto es muy pequeño, siendo del orden de kT/n .

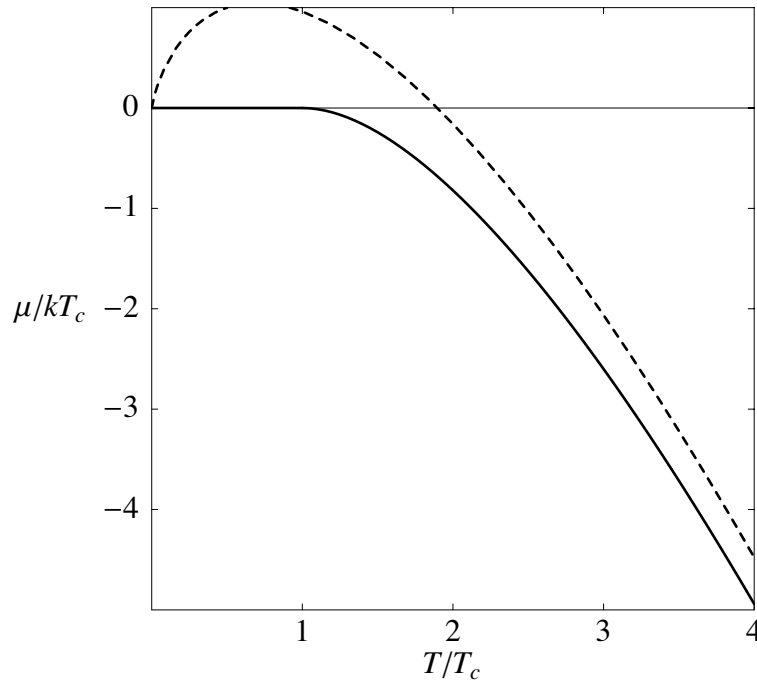


Fig. 15.2. Potencial químico de un gas de Bosones no interactuantes. Con la línea de trazos se muestra el valor clásico de μ , dado por la ec. (15.51).

El límite $e^{\beta\mu} \ll 1$ de la (15.50) permite obtener la expresión de μ para un gas perfecto clásico. Usando la segunda de las (15.46) obtenemos

$$\mu_{\text{clásico}} = -kT \ln \left\{ \frac{1}{\xi(3/2)} \left(\frac{T}{T_c} \right)^{3/2} \right\} = -kT \ln \left\{ \frac{V}{n} \left(\frac{2\pi mkT}{h^2} \right)^{3/2} \right\} \quad (15.51)$$

Los números medios de ocupación

Es interesante calcular los números medios de ocupación, dados por la (14.28). La Fig. 15.3 muestra $\bar{n}(\epsilon)$ para diferentes temperaturas del gas. El gráfico es semilogarítmico pues así se aprecian mejor los apartamientos desde la distribución clásica de Maxwell-Boltzmann, que se representa por medio de rectas de pendiente $1/kT$.

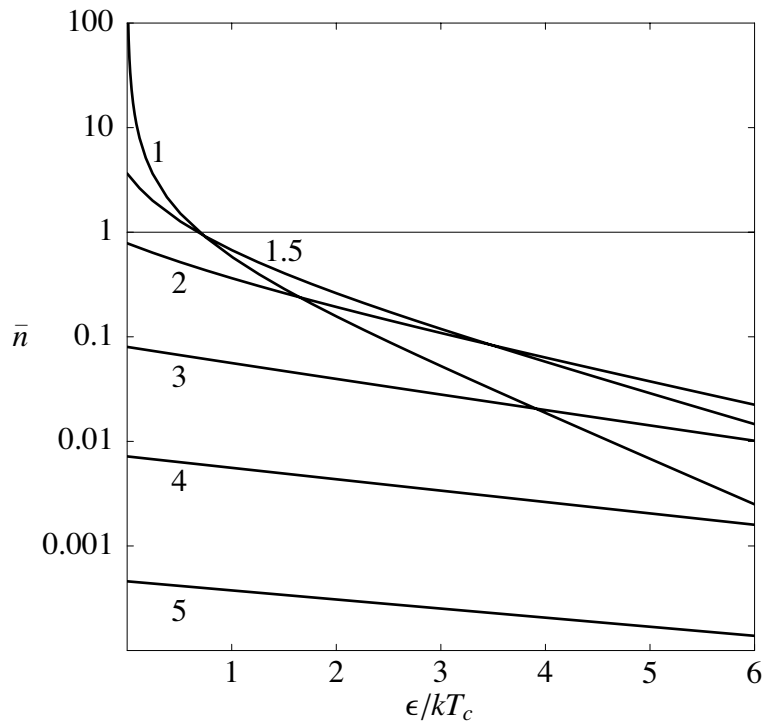


Fig. 15.3. Número de ocupación medio en función de la energía del estado para un gas de Bosones. Las curvas corresponden a $T/T_c = 1, 1.5, 2, 3, 4$ y 5 .

Se puede observar que las curvas para $T/T_c = 3, 4$ y 5 son prácticamente rectas, lo que indica que el comportamiento del gas para esas temperaturas es clásico, como es de esperar porque en esos casos $\bar{n}(\epsilon)$ es siempre mucho menor que la unidad. En cambio, las curvas correspondientes a $T/T_c = 1, 1.5$ y 2 se apartan de las rectas clásicas para ϵ por debajo de $2 - 2.5kT_c$, donde $\bar{n}(\epsilon)$ ya no se puede despreciar frente a la unidad.

El calor específico de un gas de Bosones

Vamos a estudiar ahora el calor específico de un gas de Bosones. Recordando que $nk = NR$, donde N es el número de moles, tenemos que el calor específico molar está dado por

$$\tilde{c}_V = \frac{R}{nk} \left(\frac{\partial E}{\partial T} \right)_V \quad (15.52)$$

Por lo tanto para determinar $\tilde{c}_V$ tenemos que calcular primero la energía interna del gas, dada por

$$E = \sum_i \varepsilon_i \bar{n}_i = \sum_i \frac{\varepsilon_i}{e^{\beta(\varepsilon_i - \mu)} - 1} \cong \int_0^\infty \frac{\varepsilon f(\varepsilon) d\varepsilon}{e^{\beta(\varepsilon - \mu)} - 1} = 2\pi V \left(\frac{2m}{h^2} \right)^{3/2} \int_0^\infty \frac{\varepsilon^{3/2} d\varepsilon}{e^{\beta(\varepsilon - \mu)} - 1} \quad (15.53)$$

Haciendo el cambio de variable $\beta\varepsilon = z$ y usando la (15.35) y la (15.44) resulta

$$E = \frac{3}{2} \frac{nkT}{\zeta(3/2)} \left(\frac{T}{T_c} \right)^{3/2} \text{Li}_{5/2}(e^{\beta\mu}) \quad (15.54)$$

Aquí conviene considerar por separado los casos $T < T_c$ y $T \geq T_c$.

Cuando $T < T_c$, podemos suponer que $\mu \approx 0$ y por lo tanto $\text{Li}_{5/2}(e^{\beta\mu}) \cong \text{Li}_{5/2}(1) = \zeta(5/2)$. Resulta entonces que

$$E = nk \frac{3}{2} \frac{\zeta(5/2)}{\zeta(3/2)} \frac{T^{5/2}}{T_c^{3/2}} \quad (T < T_c) \quad (15.55)$$

y por lo tanto

$$\tilde{c}_V = \frac{15}{4} R \frac{\zeta(5/2)}{\zeta(3/2)} \left(\frac{T}{T_c} \right)^{3/2} \quad (T < T_c) \quad (15.56)$$

Cuando $T > T_c$, obtenemos de la (15.52)

$$\tilde{c}_V = \frac{3}{2} \frac{R}{\zeta(3/2)} \left(\frac{T}{T_c} \right)^{3/2} \left\{ \frac{5}{2} \text{Li}_{5/2}(e^{\beta\mu}) - \text{Li}_{3/2}(e^{\beta\mu}) \left[\frac{\mu}{kT} - \frac{1}{k} \left(\frac{\partial\mu}{\partial T} \right)_n \right] \right\} \quad (15.57)$$

Para obtener la expresión de $(\partial\mu/\partial T)_n$ derivamos la (15.50). Resulta

$$\frac{\mu}{kT} - \frac{1}{k} \left(\frac{\partial\mu}{\partial T} \right)_n = \frac{3}{2} \frac{\text{Li}_{3/2}(e^{\beta\mu})}{\text{Li}_{1/2}(e^{\beta\mu})} \quad (15.58)$$

Obsérvese que tanto μ como $(\partial\mu/\partial T)_n$ tienden a cero para $T \rightarrow T_c$ pues $\text{Li}_{1/2}(t)$ diverge en $t = 1$. Usando la (15.58) obtenemos finalmente

$$\tilde{c}_V = \frac{3}{2} \frac{R}{\zeta(3/2)} \left(\frac{T}{T_c} \right)^{3/2} \left\{ \frac{5}{2} \text{Li}_{5/2}(e^{\beta\mu}) - \frac{3}{2} \frac{\text{Li}_{3/2}(e^{\beta\mu})^2}{\text{Li}_{1/2}(e^{\beta\mu})} \right\} \quad (T \geq T_c) \quad (15.59)$$

Cuando $T = T_c$ tanto la (15.56) como la (15.59) convergen a

$$\tilde{c}_V(T_c) = \frac{15}{4} R \frac{\zeta(5/2)}{\zeta(3/2)} = 1.92567 R \quad (15.60)$$

En la Fig. 15.4 se muestra el comportamiento de $\tilde{c}_V$.

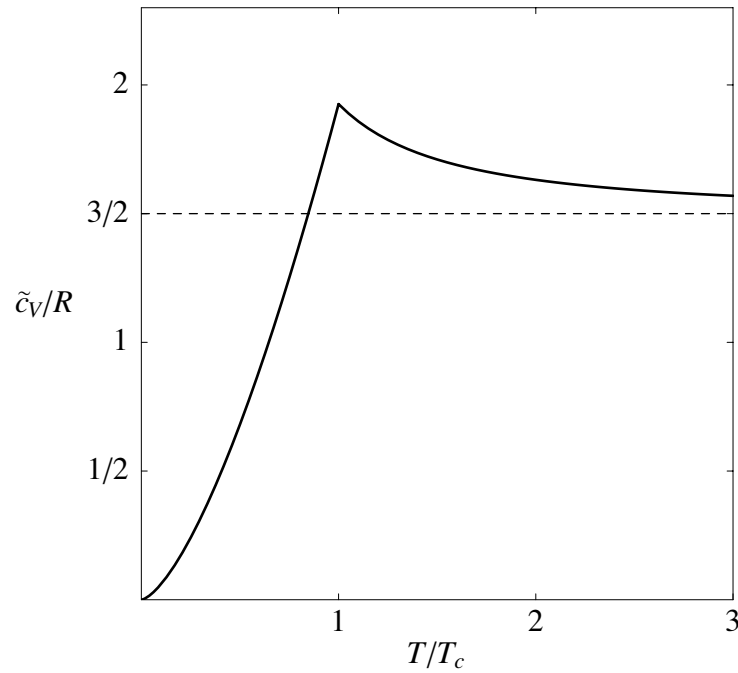


Fig. 15.4. Calor específico molar a volumen constante de un gas de Bosones. La línea de trazos indica el valor clásico $3R/2$. Para $T > T_c$ el calor específico del gas de Bosones es mayor que el valor clásico, y aumenta al disminuir T hasta alcanzar en T_c el máximo dado por la ec. (15.60). Al disminuir la temperatura por debajo de T_c , el calor específico del gas degenerado disminuye rápidamente con T . En $T = T_c$ la derivada $\partial\tilde{c}_V/\partial T$ es discontinua.

Es interesante comparar las propiedades que acabamos de estudiar con el comportamiento del ^4He líquido a bajas temperaturas. Como el ^4He tiene spin cero cumple la estadística de Bose-Einstein y por ser un gas inerte, las fuerzas interatómicas son las de van der Waals, que son muy débiles. Por ese motivo el helio se licúa a una temperatura muy baja (el punto normal de ebullición del ^4He es de 4.2°K), y sólo se solidifica bajo presiones muy grandes. Al enfriar el ^4He líquido en contacto con su vapor, sus propiedades cambian bruscamente a $T = 2.17^\circ\text{K}$. Para $T > 2.17^\circ\text{K}$ se comporta como un líquido normal, llamado helio I. Para $T < 2.17^\circ\text{K}$ presenta propiedades muy llamativas, por ejemplo fluye por capilares muy finos como si no tuviese viscosidad. Esta forma se llama helio II, y sus características se describen por medio del modelo de dos fluidos, que trata el helio II como una mezcla de un fluido normal y otro *superfluido*, sin interacción viscosa entre sí. El fluido normal tiene todas las propiedades familiares de los fluidos. El superfluido, en cambio, tiene características muy curiosas: su entropía es nula, y no tiene viscosidad. El hecho que en un gas de Bosones degenerado hay dos clases de partículas cuyo comportamiento es muy diferente sugiere que las propiedades del helio II se relacionan con el fenómeno de la degeneración. Hay muchas analogías entre el helio II y el gas de Bosones degenerado, además de las que vamos a comentar. En la transición helio I-helio II se observa un comportamiento anormal del calor específico (Fig. 15.5). Como la curva experimental recuerda la letra griega λ (lambda), el paso de helio I a helio II se denomina *transición λ* , y la temperatura de transición se llama *punto λ* . La semejanza entre las Fig. 15.4 y 15.5 llevó a London⁹ a inter-

⁹ Fritz Wolfgang London y Walter Heitler desarrollaron en 1927 el primer tratamiento cuántico de la molécula de hidrógeno. F. London y su hermano Heinz formularon en 1935 la teoría fenomenológica de la superconductividad, un fenómeno estrechamente relacionado con la superfluidad.

pretar que la transición λ del ^4He es una manifestación de la condensación de Bose-Einstein. Usando la ec. (15.35) calculó la temperatura de condensación de un gas perfecto de átomos de helio, para una densidad igual a la del helio I en el punto λ experimental. Con $V/n = 27.6 \text{ cm}^3/\text{mol}$ obtuvo $T_c = 3.13 \text{ }^\circ\text{K}$, lo cual no está demasiado lejos del valor medido $T_\lambda = 2.17 \text{ }^\circ\text{K}$.

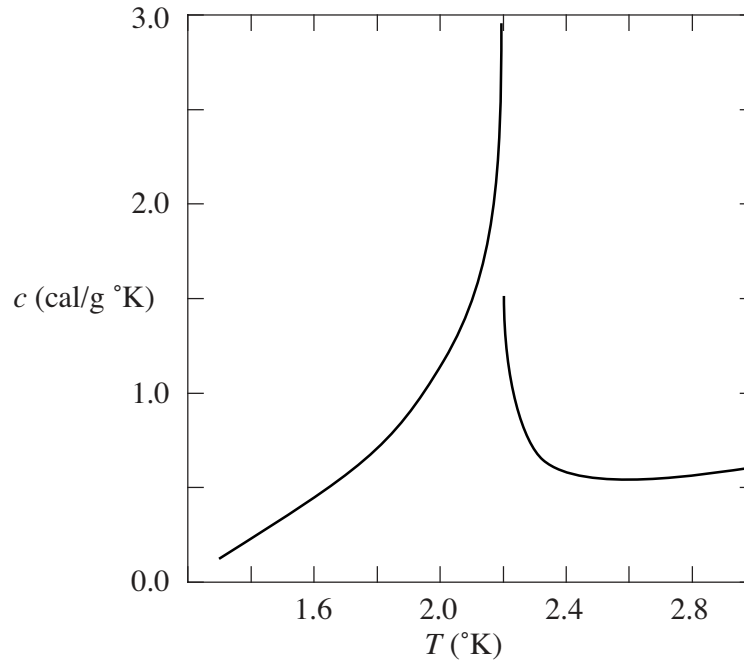


Fig. 15.5. Comportamiento del calor específico del ^4He , tal como resulta de las mediciones. Se puede apreciar la transición λ a $T = 2.17 \text{ }^\circ\text{K}$.

Hay diferencias importantes entre el comportamiento del helio II y el gas de Bosones. Entre ellas se cuenta que la dependencia en la presión de la temperatura de transición es distinta. Además, las mediciones más precisas muestran que en el punto λ , $\tilde{c}_V$ no se mantiene acotado como ocurre para el gas de Bosones, sino que tiene un infinito logarítmico. Si embargo estas diferencias no deben sorprender, pues la teoría que hemos desarrollado corresponde a un gas perfecto, y no se puede aplicar tal cual al helio líquido, en el cual hay interacciones interatómicas, que (evidentemente) no se pueden ignorar. Por consiguiente es razonable concluir que la transición λ del ^4He es lo análogo para un líquido de la condensación de Bose-Einstein de un gas perfecto. Esta conclusión se refuerza si se tiene en cuenta que el ^3He , cuyos átomos tienen spin $1/2$ y por lo tanto son Fermiones, no presenta transición¹⁰ en el punto λ del ^4He .

La superconductividad (que no tratamos por razones de espacio) es un fenómeno análogo a la superfluidez. Según la teoría BCS (formulada en 1957 por John Bardeen, Leon N. Cooper y John R. Schrieffer), a temperaturas muy bajas los electrones de conducción se agrupan de a pares (*pares de Cooper*) formando una suerte de cuasimolécula. Como los pares se comportan como Bosones, se pueden condensar y esto explica las propiedades de los superconductores.

¹⁰ Para el ^3He , la superfluidez se observa a temperaturas inferiores a $3 \times 10^{-3} \text{ }^\circ\text{K}$. Esto se debe a que a esas temperaturas tan bajas los átomos del ^3He se unen formando pares, de igual modo que los electrones en un superconductor. Los pares tienen spin entero y por lo tanto obedecen a la estadística de Bose-Einstein. Tan pronto se forman los pares, éstos sufren una condensación, y por lo tanto todo lo dicho para el ^4He se aplica también al ^3He cuando éste se encuentra a temperaturas para las cuales se han formado los pares.

Durante muchos años se careció de evidencia directa de la existencia de la condensación de Bose-Einstein, aunque se creía que ocurría en el helio líquido. La situación cambió en 1995 cuando Michael H. Anderson, Jason Enscher, Michael Matthews, Carl Wieman y Eric Cornell, observaron la condensación de Bose-Einstein en un gas de ^{87}Rb con una densidad de 2.5×10^{18} átomos/m³, enfriado por debajo de 0.17×10^{-6} °K. El gas estaba confinado por un campo magnético, y para detectar la condensación se desconectaba dicho campo y poco después se registraba la densidad espacial de las partículas. Aquellas cuya velocidad es apreciable se alejan rápidamente de la región de confinamiento, pero los átomos del condensado, cuya cantidad de movimiento es nula, permanecen en dicha región. Si hay muchas partículas en ese estado, se esperaba ver un pico de la densidad en la región donde el gas había estado confinado. Esto fue efectivamente observado, pero sólo si la temperatura está por debajo de T_c . Al enfriar el gas por debajo de T_c se encuentra que el pico se refuerza, tal como está previsto por la teoría. Ese mismo año, en forma independiente, Wolfgang Ketterle y sus colaboradores realizaron experimentos similares pero algo más sofisticados con átomos de sodio y de litio, y también observaron la condensación de Bose-Einstein. Por estos trabajos Cornell, Ketterle y Wieman fueron galardonados con el Premio Nobel para la Física en 2001.

El gas de fotones y la radiación de cuerpo negro

Vamos a estudiar ahora la radiación de cuerpo negro, considerándola como un gas perfecto de *fotones* en equilibrio térmico. De acuerdo con la hipótesis de Einstein (ver el Capítulo 4) vamos a suponer que los fotones son partículas de masa nula y energía dada por la ec. (4.16), esto es

$$\varepsilon_i = h\nu_i \quad (15.61)$$

Como los fotones tienen spin 1, son Bosones y corresponde aplicar la estadística de Bose-Einstein. Al usar la distribución (15.24) hay que tener en cuenta que los fotones pueden ser emitidos y absorbidos por las paredes de la cavidad (este es el mecanismo que garantiza el equilibrio de la radiación de cuerpo negro a T y V fijos), y entonces n no está determinado *a priori*, a diferencia del gas de Bosones que tratamos antes. En este caso n está *determinado* por la condición de equilibrio a T y V fijos¹¹. Esto implica que *el potencial químico del gas de fotones es idénticamente nulo*. En efecto, para todo cambio espontáneo T y V fijos se cumple

$$\Delta F \leq 0 \quad (15.62)$$

donde F es la función de Helmholtz. Luego en el equilibrio F es mínimo, lo que lleva a la condición necesaria $(\partial F / \partial n)_{T,V} = 0$. De aquí se obtiene entonces

$$\mu = \left(\frac{\partial F}{\partial n} \right)_{T,V} = 0 \quad (15.63)$$

Teniendo en cuenta este hecho la distribución (15.24) queda de la forma

$$\bar{n}_i = \frac{1}{e^{\beta \varepsilon_i} - 1} \quad \text{fotones} \quad (15.64)$$

¹¹ ver el Capítulo 8 de Termodinámica e Introducción a la Mecánica Estadística.

A partir de la (15.64) es muy sencillo deducir la distribución espectral de Planck. En efecto, recordemos que la densidad de estados por unidad de intervalo de frecuencia es (ec. (4.1)):

$$N(\nu) = \frac{8\pi V}{c^3} \nu^2 \quad (15.65)$$

Por lo tanto la densidad de energía $\bar{u}(\nu, T)$ (energía por unidad de volumen en el intervalo de frecuencia entre ν y $\nu + d\nu$) está dada por

$$\bar{u}(\nu, T) = V^{-1} N(\nu) \bar{n}(\epsilon) \epsilon \quad (15.66)$$

Sustituyendo (15.63) y (15.64) en (15.66) se obtiene la distribución espectral de Planck (ec. (4.12)).

$$\bar{u}(\nu, T) = \frac{8\pi \nu^2}{c^3} \frac{h\nu}{e^{h\nu/kT} - 1} \quad (15.66)$$

Es fácil verificar que el máximo de $\bar{u}(\nu, T)$ ocurre para

$$\nu = \nu_m = 2.82144 \frac{kT}{h} \quad (15.67)$$

La (15.67) expresa la Ley de desplazamiento de Wien, que fue obtenida originalmente en la forma $\nu_m \sim T$ en base a un argumento puramente termodinámico.

Integrando la (15.66) sobre todas las frecuencias obtenemos la densidad de energía radiante

$$\frac{E}{V} = \int_0^\infty \bar{u}(\nu, T) d\nu = \frac{8\pi h}{c^3} \int_0^\infty \frac{\nu^3}{e^{h\nu/kT} - 1} d\nu = \frac{8\pi^5 k^4}{15c^3 h^3} T^4 \quad (15.68)$$

La (15.68) es la Ley de Stefan-Boltzmann¹² $E/V = (4\sigma/c)T^4$, y podemos entonces expresar la constante de Stefan-Boltzmann como

$$\sigma = \frac{2\pi^5 k^4}{15c^2 h^3} \quad (15.69)$$

Es oportuno comentar brevemente el significado físico de la distribución (15.64) para aclarar algunos aspectos de la radiación mencionados en capítulos anteriores. La Fig. 15.6 muestra la distribución de Planck (15.66) y la Fig. 15.7 los números de ocupación medios (ec. (15.64)) por medio de curvas universales, obtenidas adimensionalizado la energía de los fotones con el factor kT y $\bar{u}(\nu, T)$ con el factor αT^3 , donde $\alpha = 60 h\sigma/\pi^4 k c$. Para comparar, mostramos también en esas figuras la posición de ν_m , y los resultados de la teoría clásica de Rayleigh-Jeans (líneas de trazos) y de la fórmula empírica de Wien mencionada en el Capítulo 4 (líneas de puntos).

Por de pronto, comparando las Figs. 15.3 y 15.7 vemos que el comportamiento de $\bar{n}$ es análogo al de un gas de Bosones a $T = T_c$ (pues en ambos casos $\mu = 0$). Por lo tanto, cualquiera sea su temperatura, el gas de fotones se comporta de manera semejante a un gas de Bosones a $T = T_c$.

¹² Ver el Capítulo 15 de Termodinámica e Introducción a la Mecánica Estadística.

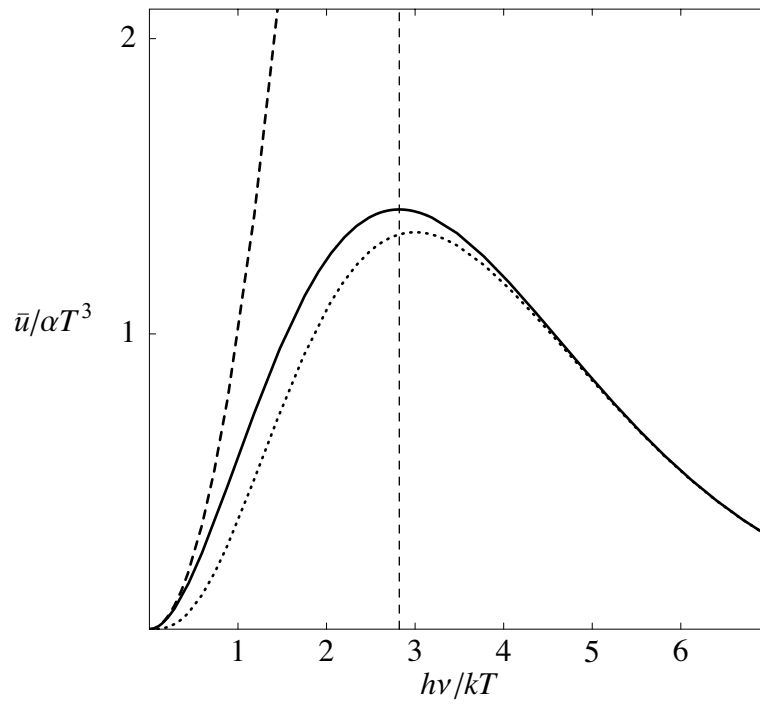


Fig. 15.6. La distribución de Planck. Representamos también la distribución de Rayleigh-Jeans (línea de trazos), que reproduce para bajas frecuencias el resultado de Planck, y la distribución de Wien (línea de puntos), que tiene el comportamiento correcto para frecuencias altas.

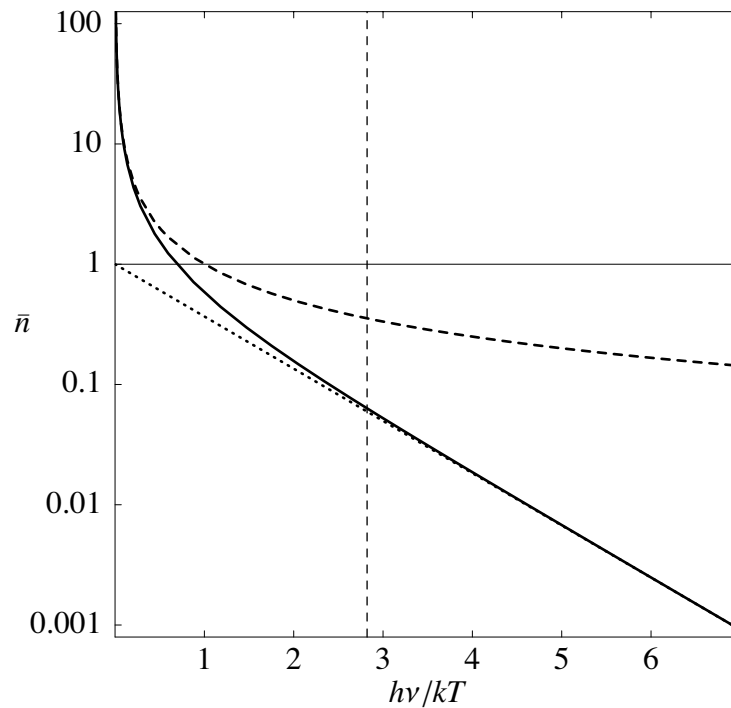


Fig. 15.7. Número medio de ocupación de los estados de fotones para la radiación de cuerpo negro. Las líneas de trazos y puntos corresponden a las fórmulas de Rayleigh-Jeans y de Wien, respectivamente.

Comparando las Figs. 15.6 y 15.7 notamos que el ámbito de validez de la fórmula de Wien corresponde al rango de frecuencias para el cual $\bar{n} \ll 1$, esto es, el rango en el cual los fotones

cumplen la condición (15.11) y entonces se comportan como *partículas clásicas* (esto es, distinguibles). En efecto, la fórmula de Wien no es otra cosa que la distribución *clásica* para un sistema de *partículas* de masa nula y energía $\epsilon = h\nu$. Este rango, que podríamos llamar entonces de partícula clásica, abarca las frecuencias mayores que $\approx 2kT/h$ y comprende un 80% de la energía radiante total. Para comprender lo que esto implica daremos algunos valores. La radiación electromagnética que proviene de casi todas las fuentes naturales es de origen térmico¹³ y su distribución espectral es, *grosso modo*, del mismo tipo que la del cuerpo negro. La luz solar corresponde a 5700 °K, y todos los fotones de longitud de onda más corta que la del cercano infrarrojo están en el rango $\nu > 2kT/h$; para una temperatura de 300°K, lo están todos los fotones de longitud de onda menor que unos 25 μm . Por lo tanto, del infrarrojo en adelante, los fotones provenientes de fuentes naturales se comportan como partículas clásicas, y se describen mediante paquetes o trenes de ondas que *no* se solapan; se trata entonces de radiación *incoherente*¹⁴. En estas condiciones los fenómenos de interferencia y difracción a que dan lugar se deben al comportamiento de *cada fotón individual* (cuyo paquete de ondas, por ejemplo, al pasar por las rendijas del experimento de Young interfiere consigo mismo). Asimismo, estos fotones interactúan con la materia *de a uno por vez*, como vimos en el Capítulo 4 al tratar el efecto fotoeléctrico, el efecto Compton y otros fenómenos.

En cambio la Ley de Rayleigh-Jeans vale para frecuencias bajas, para las cuales $\bar{n} \gg 1$. En este rango el *conjunto* de los fotones que ocupan cada estado de energía $\epsilon_i = h\nu_i$ está descrito por un paquete de ondas que es una superposición *coherente* de los paquetes de onda individuales. Por este motivo, cuando $\bar{n}$ es muy grande los campos eléctricos y magnéticos del conjunto tienen las propiedades de una *onda clásica*, caracterizada por la *fase*, además de la frecuencia y la amplitud. Puesto que en este dominio los cuantos interactúan con la materia en conjuntos coherentes, es decir *colectivamente* y no individualmente, la noción misma de fotón pierde significación. De hecho, en este régimen los fotones son inobservables. El rango de onda clásica¹⁵ abarca las frecuencias por debajo de $\approx 0.01kT/h$ y comprende menos del 3% de la energía radiante total. Para 5700 °K corresponde a $\lambda > 0.025 \text{ cm}$ y para 300°K, $\lambda > 0.5 \text{ cm}$.

Los dos comportamientos que se acaban de comentar representan los casos extremos de altas y bajas frecuencias. La radiación térmica de frecuencias intermedias ($0.01kT/h < \nu < 2kT/h$) tiene parte de las características de onda clásica y parte de las características de partícula clásica, ya que la transición entre ambos regímenes es continua. Vemos así que la descripción cuántica de la radiación incorpora la dualidad onda-corpúsculo de la cual hablamos en el Capítulo 4, y le da un significado más preciso.

Para terminar esta discusión de la radiación de cuerpo negro, observemos que si bien la presente deducción de la (15.66) y la que dimos en el Capítulo 4 son esencialmente equivalentes, hay una diferencia conceptual entre ellas que conviene subrayar. Mientras en el Capítulo 4 consideramos

¹³ Hay algunas excepciones, como por ejemplo la radiación de sincrotrón emitida por fuentes cósmicas (radiogalaxias y quasars).

¹⁴ Aquí estamos empleando los términos “coherente” e “incoherente” en un sentido algo diferente del habitual en la Óptica clásica.

¹⁵ Puesto que éste es el rango donde se ponen de manifiesto los efectos de la indistinguibilidad de los fotones, es aquí donde ocurren los efectos cuánticos de la estadística de los fotones. Desde este punto de vista se puede afirmar que la onda electromagnética clásica es, en realidad, un efecto cuántico debido a la presencia de muchos fotones coherentes.

que cada modo normal del campo de radiación está cuantificado, de modo que su energía toma solamente los valores discretos $0, h\nu, 2h\nu, \dots$, ahora estamos considerando la radiación como un sistema de *partículas* idénticas (los cuantos del campo, es decir los fotones), cada una de las cuales tiene una energía $h\nu$.

La emisión y absorción de fotones

La naturaleza Bosónica de los fotones tiene importantes consecuencias sobre los fenómenos de interacción entre la radiación y la materia. Si bien para el estudio detallado de estos procesos hace falta la Electrodinámica Cuántica (EDC), que no vamos a tratar en estas notas, es posible deducir algunos aspectos fundamentales de la interacción radiación-materia en base a argumentos sencillos, sin necesidad de invocar la EDC. Eso es lo que vamos a hacer aquí.

Emisión y absorción de fotones y la distribución de Planck

Nuestra deducción de la estadística (15.64) del gas de fotones se basó en la distribución de Bose-Einstein, que se obtiene de la Mecánica Estadística considerando la distribución gran canónica de un sistema de Bosones en equilibrio a temperatura y volumen fijos, con el requerimiento adicional que el potencial químico del gas de fotones es nulo (porque la emisión y absorción de fotones por la materia determina el número total de fotones de la cavidad, de modo que éste no está determinado *a priori*). Vamos a mostrar ahora que se puede llegar a la (15.64) directamente, considerando *explícitamente* los procesos de emisión y absorción.

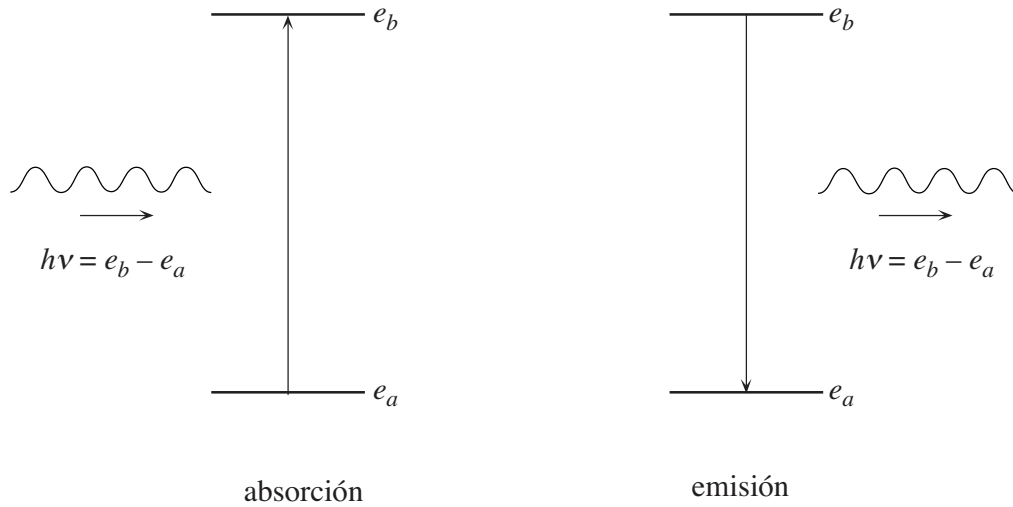


Fig. 15.8. Dos estados estacionarios particulares ψ_a y ψ_b de un átomo. En la transición $a \rightarrow b$ el átomo *absorbe* un fotón de energía $\varepsilon = h\nu = e_b - e_a$. Viceversa, en la transición $b \rightarrow a$ el átomo *emite* un fotón de la misma energía.

Sean dos estados estacionarios particulares ψ_a y ψ_b de un átomo de las paredes de la cavidad, de energías¹⁶ e_a y e_b , con $e_a < e_b$ (Fig. 15.8). Entonces, en la transición $a \rightarrow b$ el átomo *absorbe* un fotón de energía

$$\varepsilon = h\nu = e_b - e_a \quad (15.70)$$

¹⁶ Aquí a y b indican el conjunto de los números cuánticos que identifican el estado. Para cada nivel de energía hay, en general, varios estados degenerados, pero nosotros estamos considerando ahora uno entre ellos en particular.

viceversa, en la transición $b \rightarrow a$ el átomo *emite* un fotón de la misma energía. Ahora bien, hay muchos estados fotónicos ϕ_i con energía ε (su número está dado por la (15.65)), que describen las diferentes direcciones de propagación y polarización del mismo. Nosotros estamos interesados ahora en *uno sólo* de ellos, correspondiente a una particular dirección de propagación y una polarización determinada, que designaremos con ϕ_k . Consideremos entonces una transición $a \rightarrow b$ en la cual el átomo que inicialmente está en el estado ψ_a absorbe un fotón en el estado ϕ_k . La probabilidad $P_{a \rightarrow b}$ que en la unidad de tiempo ocurra esta transición está dada por

$$P_{a \rightarrow b} = p_a T_{a, n_k \rightarrow b, n_k - 1} \quad (15.71)$$

Aquí $p_a = e^{-\beta e_a} / Z$ es la probabilidad que el átomo esté en el estado ψ_a , y

$$T_{a, n_k \rightarrow b, n_k - 1} = |(\psi_b \Phi_{n_k - 1}, \hat{h}_{\text{int}, k} \psi_a \Phi_{n_k})|^2 \quad (15.72)$$

es la probabilidad de transición por unidad de tiempo desde el estado inicial $\psi_a \Phi_{n_k}$ al estado final $\psi_b \Phi_{n_k - 1}$ del sistema (átomo + radiación¹⁷) y $\hat{h}_{\text{int}, k}$ es el término del Hamiltoniano que describe la interacción del átomo con el fotón ϕ_k . Del mismo modo, la probabilidad $P_{b \rightarrow a}$ que en la unidad de tiempo el átomo que se encuentra en el estado ψ_b emita un fotón en el estado ϕ_k es

$$P_{b \rightarrow a} = p_b T_{b, n_k \rightarrow a, n_k + 1} \quad (15.73)$$

donde $p_b = e^{-\beta e_b} / Z$ es la probabilidad que el átomo esté en el estado ψ_b , y

$$T_{b, n_k \rightarrow a, n_k + 1} = |(\psi_a \Phi_{n_k + 1}, \hat{h}_{\text{int}, k} \psi_b \Phi_{n_k})|^2 \quad (15.74)$$

es la probabilidad de transición del estado inicial $\psi_b \Phi_{n_k}$ al estado final $\psi_a \Phi_{n_k + 1}$. No conocemos $\hat{h}_{\text{int}, k}$, pero claramente, es razonable suponer que debe ser de la forma

$$\hat{h}_{\text{int}, k} = h_{\text{int}, k} \hat{a}_k + h_{\text{int}, k}^\dagger \hat{a}_k^\dagger \quad (15.75)$$

donde $\hat{a}_k$ y $\hat{a}_k^\dagger$ son los operadores de creación y de aniquilación de un fotón en el estado ϕ_k , cuyos elementos de matriz (ecs. (14.32) y (14.35)) son, respectivamente

$$\hat{a}_{n_k}^{n_k - 1} = (\Phi_{n_k - 1}, \hat{a}_k \Phi_{n_k}) = \sqrt{n_k} \quad \text{y} \quad \hat{a}_{n_k}^{\dagger n_k + 1} = (\Phi_{n_k}, \hat{a}_k^\dagger \Phi_{n_k}) = \sqrt{n_k + 1} \quad (15.76)$$

y $h_{\text{int}, k}$ y su conjugado Hermitiano $h_{\text{int}, k}^\dagger$ son operadores¹⁸ que actúan sobre las funciones de onda del átomo. La forma exacta de $h_{\text{int}, k}$ y $h_{\text{int}, k}^\dagger$ depende de las variables dinámicas que describen las cargas y corrientes eléctricas atómicas y de la función de onda ϕ_k del fotón que describe los campos electromagnéticos asociados con el mismo; veremos que para nuestros fines no hace falta conocerla en detalle. Sustituyendo entonces (15.75) y (15.76) en (15.72) resulta

$$T_{a, n_k \rightarrow b, n_k - 1} = |(\psi_b, h_{\text{int}, k} \psi_a)|^2 (\Phi_{n_k - 1}, \hat{a}_k \Phi_{n_k})|^2 = n_k |(\psi_b, h_{\text{int}, k} \psi_a)|^2 \quad (15.77)$$

¹⁷ En la designación de Φ damos por sobreentendidos los números de ocupación que no cambian en la transición.

¹⁸ Estos operadores describen la interacción de las cargas y corrientes atómicas con el campo electromagnético del fotón.

Haciendo las mismas sustituciones en (15.74) y observando que

$$|(\psi_a, h_{\text{int},k}^\dagger \psi_b)|^2 = |(h_{\text{int},k} \psi_a, \psi_b)|^2 = |(\psi_b, h_{\text{int},k} \psi_a)|^2 \quad (15.78)$$

obtenemos

$$T_{b,n_k \rightarrow a,n_k+1} = |(\psi_a, h_{\text{int},k}^\dagger \psi_b)|^2 |(\Phi_{n_k+1}, \hat{a}_k^\dagger \Phi_{n_k})|^2 = (n_k + 1) |(\psi_b, h_{\text{int},k} \psi_a)|^2 \quad (15.79)$$

Se puede notar que las (15.77) y (15.79) implican que las probabilidades de transición por unidad de tiempo de un proceso y su inverso son iguales, es decir

$$T_{a,n_k \rightarrow b,n_k-1} = T_{b,n_k-1 \rightarrow a,n_k} \quad (15.80)$$

lo cual significa que en los procesos de emisión y absorción se cumple el principio del balance detallado que mencionamos en el Capítulo 14. Esto es consecuencia de la Hermiticidad de $\hat{h}_{\text{int},k}$, que se advierte en la (15.75). La presencia de los factores n_k y $n_k + 1$ en (15.77) y (15.79) se debe a la naturaleza Bosónica del fotón.

Ahora bien, si la radiación está en equilibrio térmico a la temperatura T se debe cumplir que $P_{a \rightarrow b} = P_{b \rightarrow a}$. Resulta por lo tanto que

$$n_k e^{\beta \epsilon_k} = n_k + 1 \quad (15.81)$$

Es importante observar que puesto que $P_{a \rightarrow b}$ y $P_{b \rightarrow a}$ contienen ambos el factor $|(\psi_b, h_{\text{int},k} \psi_a)|^2$, éste se cancela al imponer la condición $P_{a \rightarrow b} = P_{b \rightarrow a}$, lo cual implica que el mismo resultado (15.81) se obtiene al considerar las transiciones entre cualquier par de los estados atómicos degenerados de energías e_a y e_b que llevan a la emisión y absorción de un fotón ϕ_k .

De la (15.81) obtenemos de inmediato

$$n_k = \frac{1}{e^{\beta \epsilon_k} - 1} \quad (15.82)$$

que es la (15.64). Vemos por lo tanto que la distribución de Planck es una consecuencia de tres hechos: (a) que la dinámica de los procesos de emisión y absorción de fotones cumple con el principio del balance detallado, (b) que los fotones son Bosones y (c) que hay equilibrio térmico.

La física de la interacción entre átomos y fotones: emisión espontánea y estimulada

Imaginemos un átomo que está efectuando una transición del estado ψ_a al estado ψ_b en la cual absorbe un fotón ϕ_k de frecuencia $\nu = (e_b - e_a)/h$ (o viceversa, una transición del estado ψ_b al estado ψ_a en la cual emite un fotón de esa frecuencia). En un instante dado durante la transición su estado está descrito por una función de onda de la forma

$$\Psi(t) = c_a(t) \psi_a e^{-ie_a t/\hbar} + c_b(t) \psi_b e^{-ie_b t/\hbar} \quad (15.83)$$

donde $c_a(t)$ y $c_b(t)$ son ciertos coeficientes que dependen del tiempo (no nos interesa por el momento la normalización de Ψ). La densidad de probabilidad es entonces proporcional a

$$|\Psi(t)|^2 = |c_a(t)|^2 |\psi_a|^2 + |c_b(t)|^2 |\psi_b|^2 + 2 \text{Re}[c_a^*(t) c_b(t) \psi_a^* \psi_b e^{-i(e_b - e_a)t/\hbar}] \quad (15.84)$$

y como se ve tiene una parte que oscila con la frecuencia $\nu = (e_b - e_a)/h$. Del mismo modo, si se calcula la corriente de probabilidad, se encuentra que contiene una parte que oscila con la misma frecuencia ν . Puesto que la densidad de probabilidad y la corriente de probabilidad están relacionadas con las distribuciones espaciales de carga y corriente eléctrica del átomo, vemos que el estado Ψ describe distribuciones de cargas y corrientes que *oscilan* con la misma frecuencia que lo hacen los campos electromagnéticos del fotón que está siendo absorbido (o emitido). Esta observación sugiere que el fenómeno de absorción (o de emisión) de radiación por parte de un átomo es análogo al aumento (o disminución) de la energía de un oscilador que está siendo forzado a oscilar por una fuerza aplicada, en resonancia con su frecuencia natural. Que el oscilador gane o pierda energía en este proceso, depende de la fase relativa entre sus oscilaciones y las de la fuerza excitadora. Esta es la imagen *clásica* de la absorción y emisión de radiación por parte de un átomo. De acuerdo con esta imagen es natural pensar que tanto la absorción como la emisión de radiación son procesos *estimulados* (o *inducidos*) por la presencia del campo de radiación de la frecuencia que corresponde a la transición, y que su probabilidad debe ser proporcional a la intensidad del mismo.

Por otra parte los resultados anteriores (15.77) y (15.79) nos muestran que las probabilidades de transición por unidad de tiempo para la absorción y emisión de un fotón ϕ_k están dadas por

$$\begin{aligned} T_{a,n_k \rightarrow b,n_k-1} &= n_k |(\psi_b, h_{\text{int},k} \psi_a)|^2 && \text{absorción} \\ T_{b,n_k \rightarrow a,n_k+1} &= (n_k + 1) |(\psi_b, h_{\text{int},k} \psi_a)|^2 && \text{emisión} \end{aligned} \quad (15.85)$$

Observemos que la intensidad del campo de radiación de frecuencia ν (y dirección de propagación y polarización correspondientes a ϕ_k) es proporcional a la energía del mismo, esto es, a $n_k h \nu$. Vemos por lo tanto que la tasa de absorción es efectivamente proporcional a la intensidad del campo, tal como esperábamos. No ocurre lo mismo con la emisión, cuya tasa resulta proporcional a $n_k + 1$, lo que implica que el átomo puede emitir un fotón aún cuando la energía del campo de radiación de frecuencia ν es *nula*. Este es un resultado puramente *cuántico*, que proviene del principio de incerteza, que prohíbe que el campo electromagnético sea estrictamente nulo¹⁹. Debido a ello, en ausencia de fotones el campo tiene igualmente fluctuaciones, y son estas fluctuaciones las que provocan la emisión del átomo cuando $n_k = 0$.

Es lógico entonces escribir $T_{b,n_k \rightarrow a,n_k+1}$ como la suma de dos contribuciones: la que se debe a la *emisión espontánea* provocada por las fluctuaciones de vacío del campo electromagnético y la *emisión estimulada* por los fotones ϕ_k presentes, esto es

$$\begin{aligned} T_{b,n_k \rightarrow a,n_k+1} &= T_{b,n_k \rightarrow a,n_k+1}^e + T_{b,n_k \rightarrow a,n_k+1}^i \\ T_{b,n_k \rightarrow a,n_k+1}^e &= |(\psi_b, h_{\text{int},k} \psi_a)|^2 && \text{emisión espontánea} \\ T_{b,n_k \rightarrow a,n_k+1}^i &= n_k |(\psi_b, h_{\text{int},k} \psi_a)|^2 && \text{emisión estimulada} \end{aligned} \quad (15.85)$$

¹⁹ Sin entrar en detalles, al tratar en forma cuántica el campo electromagnético, de acuerdo con los lineamientos expuestos en el Capítulo 14, se deben interpretar las ecuaciones de Maxwell como las ecuaciones de los operadores del campo. De resultas de ello las componentes de los campos $\mathbf{E}$ y $\mathbf{B}$ son operadores que satisfacen ciertas relaciones de conmutación, las cuales tienen como consecuencia (entre otras) que en el estado de vacío los campos eléctrico y magnético no son estrictamente nulos sino que fluctúan.

Por otra parte, como es obvio, la absorción es siempre *estimulada*.

Tal como vimos anteriormente, el grueso del espectro de la radiación que proviene de fuentes naturales es incoherente pues corresponde a números de ocupación $n_k \ll 1$. Por lo tanto la emisión estimulada es inobservable en esas condiciones: la emisión es en su totalidad espontánea. Veremos en breve bajo qué condiciones es posible observar la emisión estimulada y cuales son sus peculiares características.

Los coeficientes de Einstein

Las tasas según las cuales tienen lugar los procesos de absorción y emisión de radiación por parte de un átomo están determinadas, como acabamos de ver, por las probabilidades de transición. Cabe observar, sin embargo, que las probabilidades de transición que consideramos en la Sección anterior se refieren a transiciones entre dos *particulares* estados atómicos ψ_a y ψ_b , elegidos entre los g_a estados degenerados del nivel e_a y los g_b estados degenerados del nivel e_b , asociadas con la emisión o absorción de un fotón de energía $\varepsilon = h\nu = e_b - e_a$ y con una determinada dirección de propagación y una dada polarización, correspondientes a un *particular* estado ϕ_k . Por otra parte el estudioso está muchas veces interesado en conocer la intensidad de la radiación de una dada frecuencia ν absorbida o emitida por un conjunto de muchos átomos debido a transiciones entre los niveles e_a y e_b , y no le importa saber en cuáles estados particulares ψ_a y ψ_b se encuentra cada átomo antes y después de la transición, ni tampoco en conocer la dirección de propagación y la polarización de los fotones involucrados. En este caso, las probabilidades de transición relevantes para la absorción (A) y la emisión (E) son

$$T_A = T_{e_a \rightarrow e_b} \quad \text{y} \quad T_E = T_{e_b \rightarrow e_a} \quad (15.86)$$

donde $T_{e_a \rightarrow e_b}$ es la probabilidad que el átomo que se encuentra en uno cualquiera de los estados del nivel e_a efectúe una transición a uno cualquiera de los estados del nivel e_b , sin que importe ni la dirección de propagación ni la polarización del fotón absorbido y $T_{e_b \rightarrow e_a}$ es la probabilidad que el átomo que está en un estado del nivel e_b efectúe una transición a un estado del nivel e_a , cualesquiera sean la dirección de propagación y la polarización del fotón emitido. Estas probabilidades de transición se relacionan con las $T_{a,n_k \rightarrow b,n_k-1}$ y $T_{b,n_k \rightarrow a,n_k+1}$ dadas por la (15.77) y la (15.79) por medio de

$$T_{e_a \rightarrow e_b} = \frac{2}{g_a} \sum_a \sum_b \int_k T_{a,n_k \rightarrow b,n_k-1} dk \quad \text{y} \quad T_{e_b \rightarrow e_a} = \frac{2}{g_b} \sum_b \sum_a \int_k T_{b,n_k \rightarrow a,n_k+1} dk \quad (15.87)$$

Aquí las sumas comprenden los g_a estados degenerados del nivel e_a y los g_b estados degenerados del nivel e_b , y la integral se extiende sobre todas las direcciones de propagación del fotón; el factor 2 toma en cuenta las dos polarizaciones del fotón, y la división por el factor de degeneración del estado inicial se debe a que hay que promediar sobre los diferentes estados iniciales degenerados puesto que todos ellos son igualmente probables.

Es importante observar que debido a las transiciones originadas por los procesos radiativos y por las colisiones que ocurren en el medio, los estados atómicos no son estrictamente estacionarios, sino que tienen una “vida media” τ_i , que depende del estado ψ_i y de la densidad y temperatura del medio, que determinan la frecuencia de las colisiones. Por este motivo las energías de los estados atómicos están indeterminadas dentro de una incerteza $\Delta e_i \approx h/\tau_i$. Por lo tanto las líneas espectrales (sea de emisión como de absorción) no son estrictamente monocromáticas, pues los

fotones provenientes de distintos átomos tienen energías que difieren en cantidades del orden de Δe_i . Las integrales en (15.87) se tienen que extender entonces sobre un intervalo de energías del fotón del orden de Δe alrededor de ε . Esto da lugar a un *perfil de línea* de ancho no nulo, descrito por una función $f(\nu)$ cuya forma precisa depende los procesos que determinan la vida media del nivel. La función $f(\nu)$ tiene un máximo en $\nu = \nu_0 = \varepsilon/h$ y tiende rápidamente a cero al crecer $|\nu - \nu_0|$; es habitual normalizarla de modo que

$$\int_0^{\infty} f(\nu) d\nu = 1 \quad (15.88)$$

Para tratar estos problemas es útil emplear los *coeficientes de Einstein*, que se relacionan con T_A y T_E , y que se pueden determinar a partir de magnitudes fácilmente medibles. Así, escribiremos que la probabilidad que *en ausencia de radiación* el átomo que está en un estado del nivel e_b emita espontáneamente un fotón de frecuencia comprendida entre ν y $\nu + d\nu$ está dada por

$$A_{b \rightarrow a} f(\nu) d\nu \quad (15.89)$$

donde $A_{b \rightarrow a}$ es el *coeficiente de Einstein de emisión espontánea*, y representa la probabilidad que en la unidad de tiempo el átomo efectúe espontáneamente la transición $b \rightarrow a$, cualquiera sea la frecuencia del fotón emitido.

Supongamos ahora que el átomo está en presencia de radiación, cuya densidad espectral de energía es $u(\nu)$. La presencia del campo de radiación cambia la situación anterior de dos maneras: (a) si el átomo está en el nivel e_b la probabilidad de la transición $b \rightarrow a$ cambia; (b) si el átomo está en el estado e_a se torna posible la transición $a \rightarrow b$ por absorción de un fotón. Puesto que ambos efectos se deben al campo es lógico suponer que en el primer caso la probabilidad de emitir un fotón de frecuencia entre ν y $\nu + d\nu$ sea ahora

$$[A_{b \rightarrow a} f(\nu) + B_{b \rightarrow a} u(\nu) f'(\nu)] d\nu \quad (15.90)$$

y en el segundo caso, la probabilidad de absorción sea

$$B_{a \rightarrow b} u(\nu) f''(\nu) d\nu \quad (15.91)$$

Las nuevas funciones $f'(\nu)$ y $f''(\nu)$ tienen sus máximos en $\nu = \nu_0 = \varepsilon/h$, tienden rápidamente a cero al crecer $|\nu - \nu_0|$ y están normalizadas del mismo modo que $f(\nu)$. *A priori*, las tres funciones f , f' y f'' pueden ser diferentes, pero hay por lo menos un caso, aquél en que el átomo forma parte de un sistema en equilibrio termodinámico, en que $f = f' = f''$ y por ese motivo de supone que son siempre idénticas. Las cantidades $B_{b \rightarrow a}$ y $B_{a \rightarrow b}$ son, respectivamente, los *coeficientes de Einstein de emisión estimulada* y de *absorción*. Los coeficientes de Einstein son *constantes* características de los niveles a y b , y se pueden calcular teóricamente en términos de los elementos de matriz de $h_{\text{int},k}$ usando las expresiones (15.87), o en su defecto se pueden determinar experimentalmente.

Es fácil mostrar que entre los *tres* coeficientes de Einstein deben existir *dos* relaciones²⁰, y por ese motivo una vez que se conoce uno de ellos se pueden calcular de inmediato los otros dos. En

²⁰ Estas relaciones son consecuencia del principio del balance detallado y de la naturaleza Bosónica de los fotones.

efecto, consideremos la radiación en una cavidad y los átomos de sus paredes, en equilibrio termodinámico a la temperatura T . La relación entre el número de átomos en niveles a y b está dada por la distribución de Boltzmann

$$\frac{N_b}{N_a} = \frac{g_b}{g_a} e^{-\beta(e_b - e_a)} = \frac{g_b}{g_a} e^{-\beta\epsilon} \quad (15.92)$$

El número de átomos que en la unidad de tiempo efectúan la transición $b \rightarrow a$ (por emisión tanto espontánea como estimulada, y suponiendo $f = f'$) resulta ser

$$N_b \int_0^{\infty} [A_{b \rightarrow a} + B_{b \rightarrow a} u(\nu)] f(\nu) d\nu = N_b [A_{b \rightarrow a} + B_{b \rightarrow a} u(\nu_0)] \quad (15.93)$$

dado que $u(\nu)$ varía muy lentamente con ν . Del mismo modo, el número de átomos que en la unidad de tiempo efectúan la transición $a \rightarrow b$ (suponiendo $f = f''$) es

$$N_a \int_0^{\infty} B_{a \rightarrow b} u(\nu) f(\nu) d\nu = N_a B_{a \rightarrow b} u(\nu_0) \quad (15.94)$$

En el equilibrio termodinámico estos números deben ser iguales entre sí, de modo que

$$\frac{N_b}{N_a} = \frac{B_{a \rightarrow b} u(\nu_0)}{A_{b \rightarrow a} + B_{b \rightarrow a} u(\nu_0)} \quad (15.95)$$

Usando la (15.92) y omitiendo el subíndice 0, obtenemos

$$\frac{g_b}{g_a} e^{-\beta\epsilon} = \frac{B_{a \rightarrow b} u(\nu)}{A_{b \rightarrow a} + B_{b \rightarrow a} u(\nu)} \quad (15.96)$$

La densidad espectral de energía radiante $u(\nu, T)$ está dada por la (15.66) y entonces tenemos

$$\frac{g_b}{g_a} e^{-\beta\epsilon} = \frac{B_{a \rightarrow b} \frac{8\pi h \nu^3}{c^3}}{(e^{\beta\epsilon} - 1) A_{b \rightarrow a} + B_{b \rightarrow a} \frac{8\pi h \nu^3}{c^3}} \quad (15.97)$$

Esta ecuación vale para toda T ; en particular, para $T \rightarrow \infty$ se tiene que $e^{\beta\epsilon} \rightarrow 1$ y $e^{-\beta\epsilon} \rightarrow 1$ y por lo tanto se debe tener que

$$g_b B_{b \rightarrow a} = g_a B_{a \rightarrow b} \quad (15.98)$$

Sustituyendo este resultado en la (15.97) se obtiene una identidad si

$$A_{b \rightarrow a} = B_{b \rightarrow a} \frac{8\pi h \nu^3}{c^3} \quad (15.99)$$

Las relaciones (15.98) y (15.99) son las relaciones buscadas. Las obtuvimos para el caso del equilibrio termodinámico, pero puesto que se trata de relaciones entre constantes valen siempre.

Relación entre los coeficientes de Einstein y las magnitudes macroscópicas

La variación de la intensidad de un haz de radiación de frecuencia ν que atraviesa un espesor dx de un medio en el cual no tienen lugar procesos de emisión o de difusión es, por definición

$$dI(\nu) = -\kappa u(\nu) dx \quad (15.100)$$

donde κ es el *coeficiente de absorción*. La contribución del haz a la densidad de energía radiante está dada por $4\pi I(\nu)/c$, luego la variación de $u(\nu)$ en el lapso $dt = dx/c$ que emplea el haz en atravesar el espesor dx es

$$du(\nu) = -\kappa u(\nu) c dt \quad (15.100)$$

Por otra parte, si la atenuación se debe a las transiciones $a \rightarrow b$ de los átomos que se encuentran en el nivel e_a , usando la definición (15.91) del coeficiente de Einstein $B_{a \rightarrow b}$, se debe tener que

$$du(\nu) = -h\nu N_a B_{a \rightarrow b} u(\nu) c dt \quad (15.101)$$

puesto que cada transición sustrae del haz un fotón de energía $h\nu$. Comparando (15.100) con (15.101) se obtiene la siguiente relación entre κ y $B_{a \rightarrow b}$:

$$\kappa = \frac{h\nu}{c} N_a B_{a \rightarrow b} \quad (15.102)$$

En nuestro argumento no tomamos en cuenta, sin embargo, que el haz de radiación estimula la emisión de los átomos del nivel e_b , que se manifiesta como una absorción *negativa*. Luego para obtener el verdadero coeficiente de absorción tenemos que restar del miembro derecho de la (15.102) la cantidad $h\nu N_b B_{b \rightarrow a} / c$. Por lo tanto el coeficiente de absorción correcto es

$$\kappa' = \frac{h\nu}{c} (N_a B_{a \rightarrow b} - N_b B_{b \rightarrow a}) = \frac{h\nu}{c} N_a B_{a \rightarrow b} \left(1 - \frac{N_b}{N_a} \frac{B_{b \rightarrow a}}{B_{a \rightarrow b}} \right) = \frac{h\nu}{c} N_a B_{a \rightarrow b} \left(1 - \frac{g_a}{g_b} \frac{N_b}{N_a} \right) \quad (15.103)$$

donde usamos la relación (15.98). Si los átomos están en equilibrio termodinámico a la temperatura T obtenemos entonces

$$\kappa' = \frac{h\nu}{c} N_a B_{a \rightarrow b} (1 - e^{-h\nu/kT}) \quad (15.104)$$

De momento que hemos incluido la emisión estimulada en el coeficiente de absorción, la *emisividad* j está dada solamente por la emisión espontánea. Por lo tanto

$$j = \frac{1}{4\pi} N_b A_{b \rightarrow a} h\nu \quad (15.105)$$

de donde resulta, recordando las relaciones entre los coeficientes de Einstein, la Ley de Kirchhoff-Planck:

$$\frac{j}{\kappa'} = \frac{2h\nu^3}{c^2} \frac{1}{e^{h\nu/kT} - 1} = B_\nu(T) \quad (15.106)$$

Si se ignora la emisión inducida, es decir, si se usa la expresión (15.102) del coeficiente de absorción en lugar de la (15.104), se obtiene en lugar de la (15.106) la expresión

$$\frac{j}{\kappa} = \frac{2h\nu^3}{c^2} e^{-h\nu/kT} = B_\nu(T) \quad (h\nu \gg kT) \quad (15.107)$$

que es una aproximación de la Ley de Kirchhoff-Planck que vale para $h\nu \gg kT$, es decir en el límite en que los fotones se comportan como partículas clásicas.

En el límite opuesto, cuando $h\nu \ll kT$, la emisión estimulada domina al punto que se puede ignorar la emisión espontánea. En este caso, de la (15.104) resulta que

$$\kappa' = \frac{h^2\nu^2}{ckT} N_a B_{a \rightarrow b} \quad (h\nu \ll kT) \quad (15.108)$$

a partir de la cual se obtiene la bien conocida relación de Rayleigh-Jeans para la relación entre la emisividad y el coeficiente de absorción

$$\frac{j}{\kappa'} = 2kT \left(\frac{\nu}{c} \right)^2 = B_\nu(T) \quad (h\nu \ll kT) \quad (15.108)$$

en este límite de onda clásica la constante de Planck no interviene en $B_\nu(T)$.

El efecto laser

Consideremos la expresión (15.03) del coeficiente de absorción, que tiene validez general:

$$\kappa' = \frac{h\nu}{c} N_a B_{a \rightarrow b} \left(1 - \frac{g_a}{g_b} \frac{N_b}{N_a} \right) \quad (15.109)$$

En el equilibrio térmico tenemos que

$$1 - \frac{g_a}{g_b} \frac{N_b}{N_a} = 1 - e^{-h\nu/kT} > 0 \quad (15.110)$$

de manera que $\kappa' > 0$, lo que significa que la absorción domina sobre la emisión estimulada y por lo tanto el coeficiente de absorción es *positivo*, lo que significa que el haz de luz se *atenúa* al atravesar el medio. Sin embargo, en condiciones particulares, puede ocurrir que

$$\frac{g_a}{g_b} \frac{N_b}{N_a} > 1 \quad (15.111)$$

y en este caso se tiene que la emisión estimulada domina sobre la absorción: el coeficiente de absorción es entonces *negativo*. De resultas de ello el haz de luz se *amplifica* al atravesar el medio. Este fenómeno se denomina *efecto laser*²¹ y toda fuente luminosa que aprovecha este efecto se denomina *laser*. Para que tenga lugar el efecto laser es preciso que las poblaciones de los niveles e_b y e_a que intervienen en la transición estén *fuera* del equilibrio térmico, de manera que

²¹ El término laser es un acrónimo derivado de la frase “light amplification by stimulated emission of radiation” que describe en inglés el fenómeno de la amplificación de luz por emisión estimulada de radiación.

se pueda satisfacer la condición (15.110). Dicho en pocas palabras, la población del nivel de más alta energía (e_b) debe ser significativamente mayor que la que correspondería al equilibrio, circunstancia que en la jerga se denomina *inversión de poblaciones*. Hay diferentes técnicas²² para conseguir este resultado, que no vamos a comentar porque el lector interesado las puede estudiar en la literatura especializada. Aquí nos limitaremos a describir brevemente las peculiares características de la luz laser y algunos interesantes y novedosos fenómenos a que da lugar, debido a los cuales el laser es un fértil campo de investigación sea pura como aplicada, además de tener numerosas aplicaciones (que es imposible tratar en el breve espacio de estas notas).

El elemento esencial de todo laser es el *medio activo*, que puede ser sólido, líquido o gaseoso, donde se produce la inversión de poblaciones de los niveles de interés mediante una técnica oportuna (descargas eléctricas, bombeo óptico, reacciones químicas, etc.). El funcionamiento del laser puede ser tanto continuo (como el familiar laser de He-Ne) como pulsado. En ambos casos el medio activo tiene que estar colocado dentro de un resonador óptico (esencialmente un interferómetro de Fabry-Pérot) que selecciona una particular dirección, de modo que los fotones emitidos en esa dirección atraviesan el medio activo muchas veces estimulando la emisión de más fotones en cada pasada. De esa forma la población de un particular estado fotónico ϕ_k se torna muy numerosa, y parte de la misma sale del dispositivo como un haz *coherente* (pues todos los fotones, que ocupan el mismo estado, están en fase entre si) y perfectamente *colimado*. Para muchas aplicaciones este haz es suficiente. Si se necesitan haces de mayor intensidad, se emplean dispositivos más complejos que consisten de *cadena*s cuyo eslabón primario es el laser que acabamos de describir (que se llama *oscilador*) el cual genera un haz que luego pasa sucesivamente por una serie de *amplificadores*, que son unidades cada una de las cuales contiene el medio activo, pero que a diferencia del oscilador no poseen resonador óptico. Como el nombre lo indica la función de estas unidades es amplificar el haz que sale del oscilador hasta alcanzar la intensidad que se desea.

La radiación laser tiene características diferentes de la radiación que proviene de las fuentes naturales (térmicas). Puesto que los fotones están en el mismo estado, son coherentes (mientras que los que provienen de fuentes térmicas son mayormente incoherentes). Esto hace que se comporten como una onda electromagnética clásica (dado que $n_k \gg 1$) lo que permite, entre otras cosas, observar con suma facilidad los fenómenos de interferencia y difracción, que normalmente no se ponen en evidencia con la luz ordinaria a menos de tomar los particulares recaudos descriptos en los textos de Óptica a fin de evitar los efectos de la falta de coherencia de los fotones. La otra consecuencia de la coherencia es que los fotones de un rayo laser pueden interactuar colectivamente (y no sólo individualmente, como ocurre con la luz ordinaria) con la materia. Esto da lugar a fenómenos no lineales peculiares cuando la intensidad del haz laser es muy grande.

Sin entrar en detalles, se puede decir que no hay límite (salvo en lo que se refiere al costo) a la potencia máxima del laser, y que enfocando el haz que produce es posible alcanzar enormes densidades de energía, *imposibles de obtener con fuentes térmicas*. Por ejemplo, con un laser

²² Si bien la emisión estimulada fue reconocida por Einstein en su famoso trabajo de 1917, tuvieron que pasar cuatro décadas antes que Charles Townes, A. I. Schawlow, Nikolay Basov y Aleksandr Prokhorov demostraran la posibilidad de construir un laser óptico y que en 1960 Theodore Maiman desarrollara el primer laser de rubí. Desde entonces el progreso fue vertiginoso, y hoy existe una gran variedad de laser que se emplean en las más variadas aplicaciones, en las que se aprovecha la altísima colimación, la coherencia y la gran potencia de la radiación laser.

pulsado de vidrio-Nd, que emite luz de $\lambda = 1.06 \mu\text{m}$, se han obtenido intensidades de 10^{18} W/cm^2 en el foco²³. Una sencilla cuenta muestra que esta intensidad corresponde a tener una densidad de energía de $200 \text{ eV/\AA}^3$ ($1 \text{ \AA}^3$ es el orden de magnitud del volumen de un átomo), que implica campos eléctricos del orden de 10^{12} V/m .

Sin necesidad de llegar a valores tan extremos, la interacción luz laser-materia con haces de gran intensidad se caracteriza porque se tiene siempre un gran número de fotones coherentes que interactúan *simultáneamente*²⁴ sobre cada átomo iluminado por el haz. Ocurren entonces una variedad de fenómenos no lineales que no se observan con la luz ordinaria. Algunos de ellos han encontrado importantes aplicaciones. Mencionaremos aquí el *efecto fotoeléctrico multifotónico*, la *excitación multifotónica* y la *ionización multifotónica* de niveles atómicos.

En el efecto fotoeléctrico producido con luz ordinaria (incoherente) cada fotoelectrón absorbe un único fotón y se cumple la relación de Einstein (4.18) entre la frecuencia de la radiación y la energía cinética K_{max} de los fotoelectrones

$$K_{\text{max}} = h\nu - w_0 \quad (15.112)$$

donde w_0 es la función trabajo, que es una característica del metal empleado como fotocátodo. Si $\nu < \nu_0 = w_0/h$, los fotones no tienen energía suficiente para extraer fotoelectrones y no hay efecto fotoeléctrico, por grande que sea la intensidad (o sea el flujo de fotones) del haz de luz. Pero cuando los fotones son *coherentes*, y si la intensidad del haz es suficiente, se puede producir el efecto fotoeléctrico *multifotónico*, en el cual cada fotoelectrón absorbe simultáneamente cierto número $\mathcal{N}$ de fotones. En este caso, la (15.112) se debe modificar, y se tiene que

$$K_{\text{max}} = \mathcal{N} h\nu - w_0 \quad (15.113)$$

y entonces la frecuencia mínima necesaria no es ν_0 sino $\nu_{0\mathcal{N}} = \nu_0/\mathcal{N} = w_0/\mathcal{N}h$. Si el haz es muy intenso, $\mathcal{N}$ puede ser un número muy grande, y entonces el efecto fotoeléctrico ocurre para (casi) cualquier frecuencia, tal como sería de acuerdo con la teoría ondulatoria *clásica* del campo electromagnético.

En la excitación multifotónica se provoca una transición atómica desde un nivel e_a a un nivel e_b de mayor energía usando fotones cuya energía es *una fracción* (la mitad, un tercio, ...) de la diferencia $\varepsilon = e_b - e_a$ (Fig. 15.9a). En efecto, si $\mathcal{N}$ fotones idénticos (en el caso de la absorción $\mathcal{N}$ -fotónica) cada uno de los cuales tiene la energía $h\nu_{\mathcal{N}} = \varepsilon/\mathcal{N}$ llegan al átomo dentro de un intervalo de tiempo suficientemente breve, el átomo responde como si fuesen *un único fotón* de la frecuencia correcta $\nu = \varepsilon/h = \mathcal{N}\nu_{\mathcal{N}}$ y puede absorberlos, realizando la transición $e_a \rightarrow e_b$. Se trata de un proceso no lineal, y por lo tanto depende de la intensidad del haz laser. Si el nivel e_b pertenece al continuo (Fig. 15.9b), el proceso se denomina ionización multifotónica. Además del interés puramente científico que presentan estos procesos, la excitación multifotónica tiene hoy una importante aplicación, la *microscopía por fluorescencia multifotónica*, que permite obtener imágenes tridimensionales de muy alta resolución.

²³ Según la ley de Stefan-Boltzmann la radiancia de un cuerpo negro a la temperatura de la superficie del Sol (5700°K) es de $6 \times 10^3 \text{ W/cm}^2$. Para alcanzar una intensidad de 10^{18} W/cm^2 (pero distribuida en todo el espectro) con una fuente térmica, se precisa $T > 2 \times 10^7^\circ\text{K}$. Tan sólo una mínima fracción de la misma se encontraría en el intervalo de frecuencias correspondiente al entorno de $\lambda = 1.06 \mu\text{m}$.

²⁴ Puesto que estos fotones son coherentes, se pueden describir por medio de un campo electromagnético clásico.

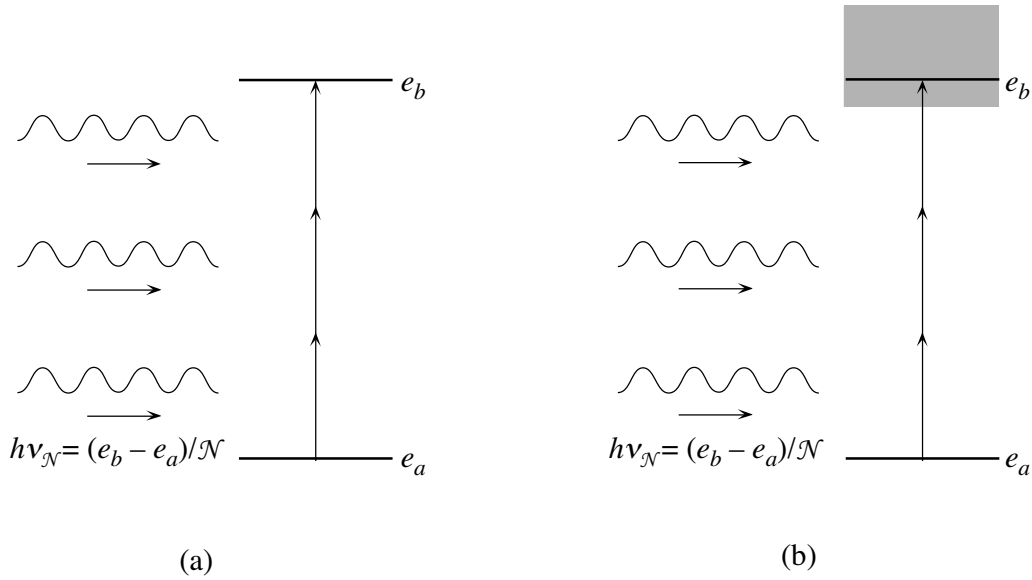


Fig. 15.9. Procesos no lineales en la interacción entre la luz laser y los átomos: (a) excitación multifotónica; (b) ionización multifotónica.

El gas perfecto de Fermiones

De acuerdo con las ecs. (15.26) y (15.27), en un gas de Fermiones que no interactúan se tiene

$$\bar{n}_i = \frac{1}{e^{\beta(\varepsilon_i - \mu)} + 1}, \quad n = \sum_i \frac{1}{e^{\beta(\varepsilon_i - \mu)} + 1} \quad (15.113)$$

Introduciendo la densidad de estados por unidad de intervalo de energía $f(\varepsilon)$, la sumatoria de la segunda de las (15.113) se puede escribir como una integral

$$n = \sum_i \frac{1}{e^{\beta(\varepsilon_i - \mu)} + 1} = (2s + 1) \int_0^\infty \frac{f(\varepsilon) d\varepsilon}{e^{\beta(\varepsilon - \mu)} + 1} = (2s + 1) 2\pi V \left(\frac{2m}{h^2} \right)^{3/2} \int_0^\infty \frac{\varepsilon^{1/2} d\varepsilon}{e^{\beta(\varepsilon - \mu)} + 1} \quad (15.114)$$

donde s es el spin, y usamos la expresión (9.19) de $f(\varepsilon)$, que se dedujo considerando los estados estacionarios de una partícula libre de masa m que se mueve en una caja de volumen V . El factor $2s + 1$ en (15.114) tiene en cuenta los diferentes estados de spin de las partículas. En lo sucesivo vamos a suponer que $s = 1/2$. Haciendo el cambio de variable $z = \beta\varepsilon$, la (15.114) se escribe como

$$n = 4\pi V \left(\frac{2mkT}{h^2} \right)^{3/2} \int_0^\infty \frac{z^{1/2} dz}{e^{z - \beta\mu} + 1} = 4\pi V \left(\frac{2mkT}{h^2} \right)^{3/2} e^{\mu\beta} \int_0^\infty \frac{z^{1/2} dz}{e^z + e^{\mu\beta}} \quad (15.115)$$

Esta es entonces la ecuación que determina el potencial químico μ , que como se ve, es función de T , V y n .

La energía de Fermi y las propiedades del gas de Fermiones en el cero absoluto

Llamaremos *energía de Fermi*, por definición, a

$$\varepsilon_F = \mu(T = 0, V, n) \quad (15.116)$$

Vamos a ver ahora que ε_F debe ser positivo. En efecto, si se tuviera $\varepsilon_F < 0$, en el límite $T \rightarrow 0$ tendríamos que $\mu\beta \rightarrow \varepsilon_F\beta \rightarrow -\infty$, y puesto que

$$\int_0^\infty \frac{z^{1/2} dz}{e^z + e^{\mu\beta}} \rightarrow \int_0^\infty z^{1/2} e^{-z} dz = \frac{\sqrt{\pi}}{2} \quad (15.117)$$

se tendría que en ese límite el miembro derecho de la (15.115) se anularía, lo que es absurdo pues implicaría $n = 0$. Por lo tanto debe ser $\varepsilon_F > 0$. Tenemos entonces que para $T = 0$ los números de ocupación medios de los estados valen

$$\bar{n}_i = \frac{1}{e^{\beta(\varepsilon_i - \varepsilon_F)} + 1} = \begin{cases} 1 & \text{si } \varepsilon_i < \varepsilon_F \\ 0 & \text{si } \varepsilon_i > \varepsilon_F \end{cases} \quad (T = 0) \quad (15.118)$$

Este resultado tiene una clara interpretación: cuando $T = 0$ el gas de Fermiones está en el estado de mínima energía, en el cual las partículas ocupan los n estados ψ_i de menor energía, pues el principio de exclusión no permite que haya más de una partícula en cada estado (las diferentes orientaciones del spin están tenidas en cuenta por el factor $2s + 1$ en la (15.114)). Estos estados son aquellos cuya energía es menor que ε_F , que es la energía del nivel más alto ocupado para $T = 0$. Esto permite determinar ε_F , pues para $T = 0$ la (15.114) (con $s = 1/2$) se reduce a

$$n = 4\pi V \left(\frac{2m}{h^2} \right)^{3/2} \int_0^{\varepsilon_F} \varepsilon^{1/2} d\varepsilon = \frac{8\pi}{3} V \left(\frac{2m}{h^2} \right)^{3/2} \varepsilon_F^{3/2} \quad (T = 0) \quad (15.119)$$

de donde obtenemos

$$\varepsilon_F = \frac{h^2}{2m} \left(\frac{3}{8\pi} \frac{n}{V} \right)^{2/3} \quad (15.120)$$

De la (15.120) se ve que la energía de Fermi es inversamente proporcional a la masa de las partículas, y proporcional a la potencia 3/2 de su concentración media n/V .

La *temperatura de Fermi* se define por medio de $\varepsilon_F \equiv kT_F$ y vale

$$T_F = \frac{h^2}{2mk} \left(\frac{3}{8\pi} \frac{n}{V} \right)^{2/3} \quad (15.121)$$

y la velocidad de Fermi se define por medio de $v_F \equiv \sqrt{2\varepsilon_F/m}$. En la Tabla 15.1 damos los valores de ε_F , T_F y v_F para los electrones de conducción de algunos metales.

Es fácil calcular el valor de la energía interna del gas de Fermiones en el cero absoluto:

$$E = 2 \sum_i \varepsilon_i \bar{n}_i = 4\pi V \left(\frac{2m}{h^2} \right)^{3/2} \int_0^{\varepsilon_F} \varepsilon^{3/2} d\varepsilon = \frac{3}{5} n \varepsilon_F = E_0 \quad (T = 0) \quad (15.122)$$

Se puede mostrar (no lo haremos aquí) que la entropía del gas de Fermiones es *nula* en el cero absoluto, de acuerdo con la Tercera Ley de la Termodinámica. Usando la relación de Euler²⁵

²⁵ Ver el Capítulo 8 de Termodinámica e Introducción a la Mecánica Estadística.

$$E = TS - pV + \mu n \quad (15.123)$$

es fácil calcular que en el cero absoluto la presión vale

$$p = p_F = \frac{2}{5} \frac{n}{V} \varepsilon_F \quad (T = 0) \quad (15.124)$$

Por lo tanto un gas de Fermiones tiene una presión no nula en el cero absoluto.

Tabla 15.1: Valores de ε_F , T_F y v_F para los electrones de conducción de algunos metales.

Elemento	n/V (10^{22} cm^{-3})	ε_F (eV)	T_F (10^4 °K)	v_F (10^8 cm/s)
Li	4.60	4.7	5.5	1.30
K	1.34	2.1	2.4	0.85
Cu	8.50	7.0	8.2	1.56
Au	5.90	5.5	6.4	1.39

El potencial químico del gas de Fermiones

Para calcular el potencial químico conviene expresar el número de partículas en función de la temperatura de Fermi. De la (15.121) obtenemos que

$$n = \frac{8\pi}{3} V \left(\frac{2mkT_F}{h^2} \right)^{3/2} \quad (15.125)$$

Introduciendo esta expresión en la (15.115) y usando la (15.44) obtenemos entonces

$$1 = -\frac{3\sqrt{\pi}}{4} \left(\frac{T}{T_F} \right)^{3/2} \text{Li}_{3/2}(-e^{\beta\mu}) \quad (15.126)$$

Invirtiendo esta relación se llega al resultado buscado, que se muestra en la Fig. 15.10, donde también representamos el valor clásico de μ , que se obtiene en el límite $e^{\beta\mu} \ll 1$ y resulta ser²⁶

$$\mu_{\text{clásico}} = -kT \ln \left\{ 2 \frac{V}{n} \left(\frac{2\pi mkT}{h^2} \right)^{3/2} \right\} \quad (15.127)$$

Para $T \ll T_F$, la dependencia de μ en T es muy débil. En esta región vale²⁷ la aproximación

$$\int_0^\infty \frac{z^{q-1} dz}{e^{z-\eta} \pm 1} = \mp \Gamma(q) \text{Li}_q(\mp e^\eta) \approx \frac{\eta^q}{q} \left(1 + \frac{2q(q-1)\pi^2}{3(3 \pm 1)\eta^2} \right), \quad (\eta \rightarrow \infty) \quad (15.128)$$

Usando la (15.128) se obtiene fácilmente una expresión aproximada del potencial químico que vale para temperaturas bajas:

²⁶ La (15.127) difiere de la (15.51) debido al factor 2 (que proviene del spin) en el argumento del logaritmo.

²⁷ El cálculo no es difícil pero algo engorroso y por eso no lo reproducimos aquí. El lector interesado lo puede encontrar en Landau y Lifschitz, Statistical Physics, Pergamon.

$$\mu \approx \epsilon_F \left[1 - \frac{\pi^2}{12} \left(\frac{T}{T_F} \right)^2 \right] \quad (T \ll T_F) \quad (15.129)$$

Para $T \gg T_F$ el potencial químico tiende a $\mu_{\text{clásico}}$ (ya para $T > 2T_F$ la diferencia es muy pequeña).

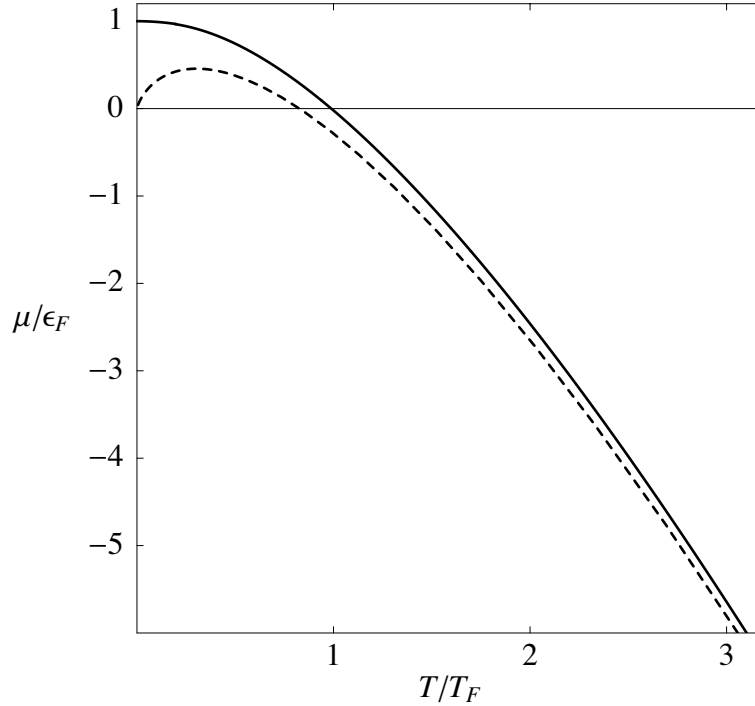


Fig. 15.10. Potencial químico de un gas de Fermiones no interactuantes. Con la línea de trazos se muestra el valor clásico de μ , dado por la ec. (15.127).

Los números medios de ocupación

Es interesante calcular los números medios de ocupación, dados por la primera de las (15.113). La Fig. 15.11 muestra $\bar{n}(\epsilon)$ para diferentes temperaturas del gas. El gráfico es semilogarítmico pues así se aprecian mejor los apartamientos desde la distribución clásica de Maxwell-Boltzmann, que se representa por medio de rectas de pendiente $1/kT$. Se puede observar que a medida que se eleva la temperatura a partir de $T = 0$ °K, las partículas que estaban ocupando los estados de energía más alta, próxima a la energía de Fermi, se excitan y pasan desde los estados con $\epsilon < \mu$ a estados con $\epsilon > \mu$. Para $T \ll T_F$ la mayor parte de estas excitaciones afectan partículas con energía cercana a $\epsilon = \mu$.

Un gas de Fermiones se dice *degenerado* cuando su temperatura es inferior a la temperatura de Fermi; en el cero absoluto el gas está *completamente degenerado*. En un gas degenerado, el principio de exclusión, que no permite más de una partícula por estado, obliga las partículas a ocupar estados de mayor energía. Si la temperatura es muy alta ($T \gg T_F$), el sistema no es degenerado puesto que los números medios de ocupación de los estados son mucho menores que la unidad, y en ese caso el principio de exclusión tiene poco efecto en la práctica. El gas se comporta entonces como un gas clásico. La transición entre el gas completamente degenerado en $T = 0$ y el gas clásico para $T \gg T_F$ es continua, pero la descripción matemática del gas en este rango de temperaturas es complicada.

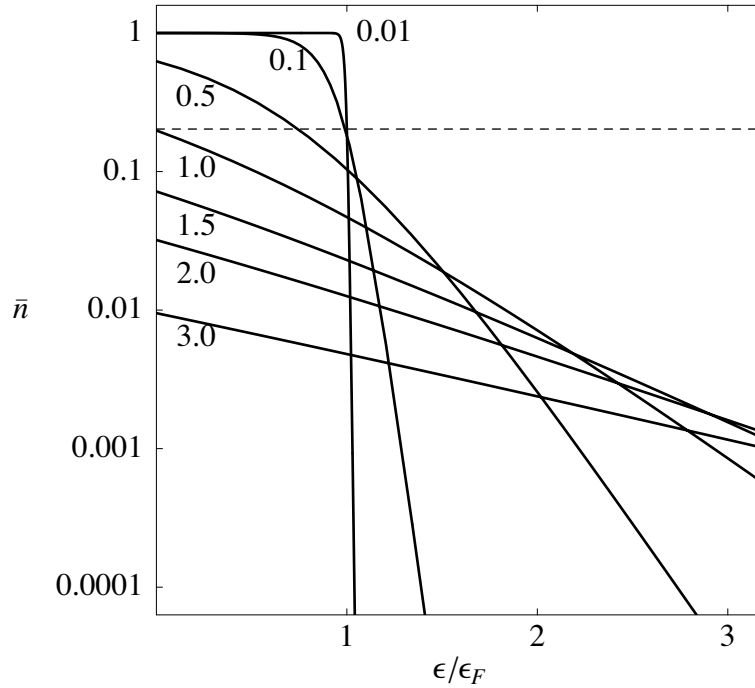


Fig. 15.11. Número de ocupación medio en función de la energía del estado para un gas de Fermiones. Las curvas corresponden a $T/T_F = 0.01, 0.1, 0.5, 1, 1.5, 2$ y 3 . La línea horizontal de trazos corresponde a $\bar{n}(\varepsilon = \mu) = 1/2$. Se puede observar que las curvas para $T/T_F = 1.5, 2$ y 3 son casi rectas, lo que indica que el comportamiento del gas es clásico, como es de esperar porque $\bar{n}(\varepsilon) \ll 1$. En cambio las curvas para $T/T_F = 0.01, 0.1$ y 0.5 difieren de las rectas clásicas donde $\bar{n}(\varepsilon)$ ya no se puede despreciar frente a la unidad.

Calor específico de un gas de Fermiones

Para calcular el calor específico molar a volumen constante procederemos de manera semejante a como lo hicimos para el gas de Bosones. Calculamos primero la energía interna del gas

$$E = 2 \sum_i \varepsilon_i \bar{n}_i = 2 \sum_i \frac{\varepsilon_i}{e^{\beta(\varepsilon_i - \mu)} + 1} \cong \frac{3}{2} n \varepsilon_F^{-3/2} \int_0^\infty \frac{\varepsilon^{3/2} d\varepsilon}{e^{\beta(\varepsilon - \mu)} + 1} \quad (15.130)$$

Usando la (15.44) obtenemos

$$E = -\frac{9\sqrt{\pi}}{8} n k T \left(\frac{T}{T_F} \right)^{3/2} \text{Li}_{5/2}(-e^{\beta\mu}) \quad (15.131)$$

Usando ahora la definición (15.52) de $\tilde{c}_V$, y la expresión

$$\frac{\mu}{kT} - \frac{1}{k} \frac{\partial \mu}{\partial T} = \frac{3}{2} \frac{\text{Li}_{3/2}(-e^{\beta\mu})}{\text{Li}_{1/2}(-e^{\beta\mu})} \quad (15.132)$$

que se obtiene derivando la (15.126) con respecto de T , obtenemos entonces

$$\tilde{c}_V = -\frac{9\sqrt{\pi}}{8} R \left(\frac{T}{T_F} \right)^{3/2} \left[\frac{5}{2} \text{Li}_{5/2}(-e^{\beta\mu}) - \frac{3}{2} \frac{\text{Li}_{3/2}(-e^{\beta\mu})^2}{\text{Li}_{1/2}(-e^{\beta\mu})} \right] \quad (15.133)$$

En la Fig. 15.12 se puede apreciar el comportamiento del calor específico.

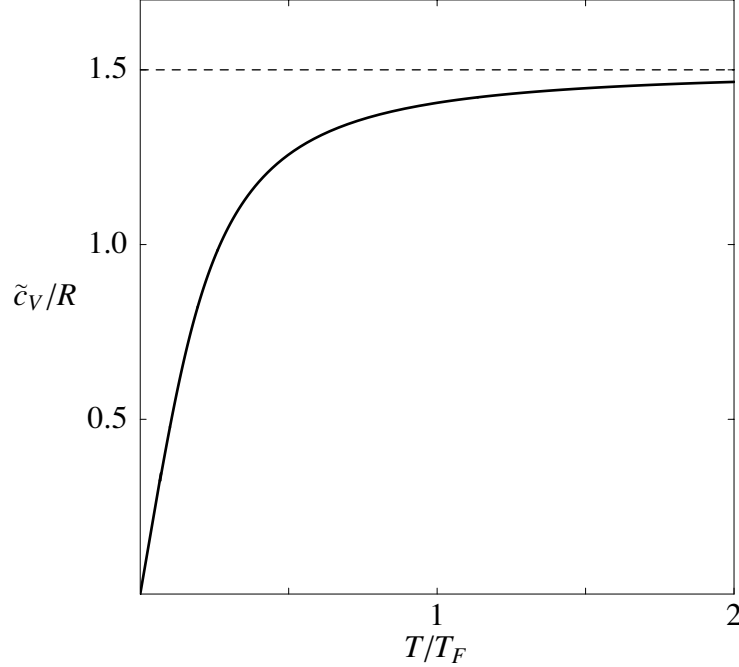


Fig. 15.12. Calor específico molar a volumen constante de un gas de Fermiones. Se ve que $\tilde{c}_V$ es siempre menor que el valor clásico $3R/2$ (representado por la línea de trazos) y para $T < T_F$ disminuye rápidamente con T . Para $T \ll T_F$ la dependencia de $\tilde{c}_V$ en T es lineal.

Usando la (15.128) se obtienen fácilmente las expresiones aproximadas de E y $\tilde{c}_V$ para temperaturas bajas:

$$E \approx \frac{3}{5} n \varepsilon_F \left[1 + \frac{5\pi^2}{12} \left(\frac{T}{T_F} \right)^2 \right] \quad (T \ll T_F) \quad (15.134)$$

y

$$\tilde{c}_V \approx \frac{\pi^2}{2} R \frac{T}{T_F} \quad (T \ll T_F) \quad (15.135)$$

El valor bajo de $\tilde{c}_V$ para $T \ll T_F$ se debe a que para un cambio de temperatura del orden de T , cabe esperar que la energía de una partícula aumente en kT . Pero para la mayoría de las partículas, con excepción de aquellas cuya energía es próxima a la energía de Fermi, los estados hacia los cuales se tendrían que excitar están ocupados (como se puede ver de la curva correspondiente a $T = 0.01 T_F$ en la Fig. 15.11). Por lo tanto, de acuerdo con el principio de exclusión, esas partículas no se pueden excitar. Sólo se pueden excitar las partículas que están en un intervalo de energía del orden de kT alrededor de ε_F , cuyo número es del orden de

$$n_{\text{exc}} \approx f(\varepsilon_F) kT \quad (15.136)$$

Por otra parte

$$f(\varepsilon_F) = \frac{4\pi V}{h^3} (2m)^{3/2} \varepsilon_F^{1/2} \quad (15.137)$$

Dividiendo esta ecuación por la (15.119) resulta entonces

$$\frac{n_{\text{exc}}}{n} \approx \frac{3}{2} \frac{T}{T_F} \quad (15.138)$$

de manera que la fracción de partículas que es excitada es del orden de T/T_F . Puesto que cada una de ellas se lleva una energía del orden de kT , tenemos que

$$E(T) \approx E_0 + kT n_{\text{exc}} = \frac{3}{5} n \varepsilon_F \left[1 + \frac{5}{2} \left(\frac{T}{T_F} \right)^2 \right] \quad (T \ll T_F) \quad (15.139)$$

y el calor específico molar resulta entonces

$$\tilde{c}_V \approx 3R \frac{T}{T_F} \quad (T \ll T_F) \quad (15.140)$$

Comparando la (15.139) y la (15.140) con los resultados exactos (15.134) y (15.135) vemos que nuestra estimación es correcta dentro de un factor del orden de la unidad.

Discutiremos ahora brevemente un par de aplicaciones de la teoría que hemos desarrollado.

El modelo de electrón libre de los metales

En el estado sólido los átomos de un metal forman una red cristalina y los electrones de valencia, que están débilmente ligados, se pueden mover libremente en todo el cristal. Por eso cuando se aplica una diferencia de potencial eléctrico al metal, esos electrones producen una corriente, y por lo mismo se denominan *electrones de conducción*. Los electrones de conducción se mueven como partículas casi libres y su camino libre medio es muy grande (bajo condiciones adecuadas puede ser del orden del cm); en efecto, son poco perturbados por la presencia de los iones positivos y sólo sufren escasas colisiones con los demás electrones de conducción.

Las razones de este comportamiento son de carácter cuántico, y las vamos a mencionar brevemente sin dar demostraciones, que el lector puede encontrar si lo desea en la bibliografía. Si se resuelve la ecuación de Schrödinger para un electrón en el potencial creado por una distribución periódica de iones positivos, se encuentra que la parte espacial de la función de onda de un electrón de energía ε tiene la forma

$$\psi_{\mathbf{k}} \sim g_{\mathbf{k}}(\mathbf{r}) e^{i\mathbf{k} \cdot \mathbf{r}} \quad , \quad k^2 = 2m^* \varepsilon / \hbar^2 \quad (15.141)$$

Aquí la función $g_{\mathbf{k}}(\mathbf{r})$ es un *factor de modulación* que tiene la misma periodicidad espacial que la red, y el parámetro $m^* = m^*(\mathbf{k})$, que se denomina la *masa efectiva* del electrón, depende del vector de onda $\mathbf{k}$ y de la geometría de la red. Si la red es infinita, resulta que no todos los valores de ε están permitidos, sino que el espectro de energía consiste de una serie de tramos continuos, llamados *bandas*, separados casi siempre por intervalos finitos aunque puede ocurrir en determinados casos que dos bandas se superpongan parcialmente. Cada banda surge de la ruptura de la

degeneración de los infinitos niveles que consisten en que el electrón esté ligado en un particular nivel a un determinado ion de la red. Si el cristal no es infinito y consiste de N iones, entonces cada banda no es continua, sino que consiste de un conjunto de N niveles discretos, muy próximos entre sí.

Lo anterior tiene una importante consecuencia en lo referente a la conductividad eléctrica del cristal. En efecto, para que la aplicación de un campo eléctrico externo de intensidad moderada produzca una corriente, es preciso que se pueda producir un movimiento neto de los electrones. Si todos los N niveles de una banda están ocupados, esto no es posible, pues los estados correspondientes a todos los valores de $\mathbf{k}$ están ocupados y no hay forma de que tenga lugar un flujo neto de electrones. Este es el caso de todas las bandas que se originan a partir de los niveles atómicos internos. La única banda que puede dar lugar a conducción de la electricidad es entonces la *banda de valencia* que surge a partir de los niveles atómicos más externos de los átomos, y las propiedades eléctricas del cristal dependen esencialmente de dos circunstancias: (1) cuántos electrones de valencia tiene cada átomo y (2) cuán separada en energía se encuentra la banda de valencia de la banda inmediatamente superior (originada a partir del primer nivel excitado atómico). Se pueden dar dos situaciones: (a) que la banda de valencia esté *parcialmente* llena, (b) que la banda de valencia esté *completamente* llena. En el primer caso, el cristal es un *conductor*, pues la aplicación de un campo eléctrico aún de muy baja intensidad puede excitar algunos electrones a estados vacíos de la banda, de modo tal que haya un flujo neto de electrones; esto es lo que ocurre con los elementos de valencia impar. En el segundo caso, que ocurre para los elementos de valencia par, si la separación en energía de la banda inmediatamente superior es considerable, el cristal es un *aislante*, pues la aplicación de un campo eléctrico moderado no produce corriente; esto sucede hasta que el campo es tan intenso que puede causar que algunos electrones pasen a la banda de energía más alta, que está vacía, en tal caso se produce la *ruptura* (o “breakdown”) del aislante. Pero si la banda inmediatamente superior está *muy próxima* en energía, si bien el cristal es aislante a $T = 0$, cuando la temperatura aumenta algunos electrones de la banda de valencia pueden pasar a la banda superior si $kT \approx \Delta\epsilon$, donde $\Delta\epsilon$ es la separación en energía de las bandas; entonces el cristal puede conducir la electricidad. Este es el caso de los *semiconductores*. Finalmente, si la banda de valencia está llena pero hay una superposición parcial con la banda inmediatamente superior, el cristal es *conductor*; este es el caso de algunos metales bivalentes.

Volviendo ahora a los metales, un electrón de conducción en un cristal perfecto a 0°K no está localizado, sino que se propaga libremente. Sólo las imperfecciones de la red (debidas a la presencia de impurezas, los defectos de la red y las vibraciones de los iones) perturban la propagación de los electrones. Además el fondo de carga positiva debido a los iones compensa en promedio las cargas de los electrones, y las interacciones residuales entre ellos tienen escasa importancia. Esto último es una consecuencia del principio de exclusión. En efecto, consideremos una colisión entre dos electrones que están inicialmente en los estados $\psi_{\mathbf{k}}$ y $\psi_{\mathbf{k}'}$, y que después de la misma están en los estados $\psi_{\mathbf{k}''}$ y $\psi_{\mathbf{k}'''}$. Esta colisión puede ocurrir solamente si los estados $\psi_{\mathbf{k}''}$ y $\psi_{\mathbf{k}'''}$ no están ocupados, pues en el caso contrario este proceso está prohibido por el principio de Pauli. Por otra parte, la mayoría de los estados finales $\psi_{\mathbf{k}''}$ y $\psi_{\mathbf{k}'''}$ que son energéticamente posibles están ocupados. Por lo tanto estos procesos no ocurren y por este motivo es razonable pensar que en primera aproximación los electrones de conducción se comporten como si fuesen libres.

Esta es la justificación del *modelo de electrón libre*, que trata los electrones de conducción de un metal como un gas perfecto de Fermiones, ignorando su interacción con los iones positivos de la red, y las interacciones entre los electrones mismos. Pese a su simplicidad este modelo es muy útil pues explica en forma cualitativa muchas características de los metales.

El modelo de electrón libre de los metales fue desarrollado originalmente por Paul Drude (1900) y Hendrik Antoon Lorentz (1905) usando la estadística clásica de Maxwell-Boltzmann. Se pudo así deducir la ley de Wiedemann-Franz (que da la dependencia con T del cociente entre las conductividades térmica y eléctrica), pero esto se debe a una coincidencia fortuita. Desde el inicio la teoría de Drude enfrentó un grave problema, pues tratados en forma clásica los electrones libres aportarían a la capacidad calorífica del metal una cantidad $\frac{3}{2}k$ por electrón, que se tendría que sumar a la capacidad calorífica debida a las vibraciones de la red cristalina. De resultados de esto, a la temperatura ambiente se debería tener $\tilde{c}_V = (3 + 3/2)R$ para un metal monovalente (con un solo electrón de conducción por átomo). Sin embargo los metales cumplen muy bien la Ley de Dulong y Petit, esto es $\tilde{c}_V = 3R$.

La Mecánica Cuántica permite resolver esta dificultad. En efecto, lo correcto es tratar los electrones como Fermiones, lo cual fue hecho por primera vez por Sommerfeld en 1928. De la Tabla 15.1 vemos que la temperatura de Fermi de los metales es de más de 10^4 °K. Por lo tanto, a temperaturas inferiores a la temperatura de fusión, se tiene que $T \ll T_F$ y entonces *los electrones de conducción están degenerados*. Podemos entonces usar la expresión aproximada (15.135) para calcular la contribución de los electrones a $\tilde{c}_V$. Esta contribución es muy pequeña en comparación con la que proviene de las vibraciones de la red, lo cual explica porqué los metales cumplen muy bien la Ley de Dulong y Petit.

La situación es diferente a temperatura muy baja. Cerca del cero absoluto, la capacidad calorífica de la red cristalina es proporcional a T^3 , mientras que la capacidad calorífica electrónica es proporcional a T . Por lo tanto a temperaturas suficientemente bajas domina la contribución electrónica, cosa que en efecto se verifica experimentalmente.

El equilibrio de las enanas blancas

Se denomina *enana blanca* una estrella que al final de su evolución se ha transformado en un objeto muy denso, pequeño y poco luminoso. En una estrella normal (como el Sol) ocurren una serie de reacciones nucleares cuyo resultado neto es la fusión de 4 núcleos de hidrógeno para producir un núcleo de helio. En cada instante se producen gran número de tales reacciones, y en consecuencia se libera una enorme cantidad de energía. Esto es lo que mantiene a la estrella caliente e impide que se contraiga debido a la gravedad, manteniéndola en equilibrio con una densidad media baja (1.416 g cm^{-3} para el Sol). Cuando la estrella agota su hidrógeno y pierde su fuente de energía, se enfría y finalmente se contrae. Cuando esto ocurre los átomos de helio se comprimen muy fuertemente. Los electrones que orbitan un átomo tienden a quedar más localizados y por lo tanto, debido al principio de incerteza, su energía cinética aumenta. Si la densidad aumenta suficientemente, la energía cinética de los electrones supera la energía potencial que tiende a ligarlos a los núcleos, y se convierten en partículas libres. De resultados de esto cada átomo de helio se transforma en una partícula α más dos electrones libres. Hay entonces una densidad enorme de electrones, que están relativamente fríos pues el combustible nuclear quedó agotado. Este es el estado de enana blanca.

Para estimar la energía cinética por electrón cuando es una partícula libre podemos razonar del modo siguiente: cada electrón ocupa como mínimo un volumen del orden de

$$v_e = \lambda_B^3 = \left(\frac{h}{p}\right)^3 \quad (15.142)$$

donde λ_B es la longitud de onda de Broglie. Si tenemos n electrones no podrán ocupar un volumen menor que $V = nv_e/2$ (dividimos por 2 pues dos electrones con spin opuesto pueden ocupar el mismo lugar) y entonces la energía cinética debe ser mayor que

$$\frac{p^2}{2m_e} = \frac{h^2}{2m_e} \left(\frac{n}{2V}\right)^{2/3} \quad (15.143)$$

Por otra parte la concentración de los electrones es $n/V = Z/d^3$ donde d es la distancia entre los núcleos y Z su carga. Por lo tanto la energía cinética típica de un electrón es

$$\frac{p^2}{2m_e} = \frac{2^{1/3}\pi^2\hbar^2 Z^{2/3}}{m_e d^2} \quad (15.144)$$

el valor típico de la energía potencial es $-2Ze^2/d$. Por lo tanto el valor absoluto de la energía potencial es menor que la energía cinética si

$$\frac{2^{1/3}\pi^2\hbar^2 Z^{2/3}}{m_e d^2} > \frac{2Ze^2}{d} \quad \text{o sea si} \quad d < 2^{-2/3}\pi^2 Z^{-1/3} a_0 \quad (15.145)$$

donde $a_0 = \hbar^2 / m_e e^2$ es el radio de Bohr. Cuando esta desigualdad se satisface (a menos de un factor numérico del orden de la unidad) los electrones no pueden quedar ligados a los núcleos. A partir de la misma podemos estimar qué densidad debe tener la estrella para que los electrones queden libres. Sea m' la masa de la estrella por electrón. Puesto que al ionizarse cada átomo produce Z electrones, la masa de la estrella por núcleo es Zm' ; por lo tanto la densidad tiene que cumplir (a menos de un factor numérico del orden de la unidad) la desigualdad

$$\rho = \frac{Zm'}{d^3} > \frac{Z^2 m'}{a_0^3} \quad (15.146)$$

Para enanas blancas constituidas por helio $Z = 2$, luego $m' = 2m_H = 3.345 \times 10^{-24}$ g y por lo tanto la densidad tiene que ser mayor que 100 g/cm^3 .

Las consideraciones anteriores presuponen que los electrones no son relativísticos, y para que esta sea una aproximación razonable se debe cumplir que $v_F \ll c$, una restricción que implica que la densidad no debe superar unos $2 \times 10^6 \text{ g/cm}^3$. Para estrellas cuya densidad se encuentra entre 100 y $2 \times 10^6 \text{ g/cm}^3$, los electrones se comportan como un gas degenerado de Fermi no relativístico, en el cual las interacciones Coulombianas con los núcleos se pueden despreciar.

La densidad de una típica enana blanca es de 10^6 g/cm^3 , la masa es de 10^{33} g, y la temperatura (interior) es de 10^7 °K. Para estos valores, $T_F \approx 2 \times 10^9$ °K, de modo que los electrones forman un gas de Fermi degenerado. La enana blanca más conocida es Sirio B, cuya masa de 2×10^{33} g es apenas menor que la del Sol; su radio es de 5600 km, algo más pequeño que el de la Tierra.

Para ver cómo queda determinado el radio de una enana blanca, consideraremos primero el caso en que los electrones no son relativísticos. La presión en una estrella de masa M y radio R , debida a la gravitación, es del orden de

$$p_G = -\alpha GM^2 / R^4 \quad (15.147)$$

donde $G = 6.673 \times 10^{-8} \text{ dy cm}^2 / \text{g}$ es la constante universal de la gravitación y α es un factor numérico del orden de la unidad. La estabilidad mecánica requiere que esta presión sea equilibrada por la presión del gas de electrones, dada por la ec. (15.124). El número de electrones es $n = M / m'$ (m' es la masa de la estrella por electrón) y $V = 4\pi R^3 / 3$. Por lo tanto

$$p_F = \alpha' \frac{\hbar^2 M^{5/3}}{m_e m'^{5/3} R^5} \quad , \quad \alpha' = \frac{1}{5} \left(\frac{3^8}{4^5 \pi} \right)^{1/3} \quad (15.148)$$

La condición de equilibrio $p_F + p_G = 0$ requiere entonces que

$$M^{1/3} R = \gamma \frac{\hbar^2}{G m_e m'^{5/3}} \quad , \quad \gamma = \frac{\alpha'}{\alpha} \quad (15.149)$$

Un cálculo preciso muestra que $\gamma \approx 4.5$. La (15.149) muestra que $R \sim M^{-1/3}$, y que es mucho menor que el radio de una estrella normal de la misma masa, de aquí la denominación de *enana* que se da a estas estrellas.

La ley $R \sim M^{-1/3}$ se cumple mientras la densidad de equilibrio sea tal que los electrones no sean relativísticos. Pero $\rho \sim MR^{-3} \sim M^2$, y por lo tanto si la masa es muy grande, al contraerse la estrella los electrones se tornan relativísticos, lo cual invalida nuestro anterior cálculo del radio de equilibrio. Veamos qué sucede entonces, para lo cual conviene considerar el límite ultrarelativístico en el cual $p^2 / m_e \gg m_e c^2$, y por lo tanto la energía cinética está dada por

$$\varepsilon = (p^2 c^2 + m_e^2 c^4)^{1/2} - m_e c^2 \approx cp \quad (15.150)$$

y no por la expresión clásica $p^2 / 2m_e$. En este caso no podemos usar la expresión (9.19) de la densidad de estados por unidad de intervalo de energía, y es preciso calcular nuevamente $f(\varepsilon)$ teniendo en cuenta que $\varepsilon = cp$. Por medio de un razonamiento análogo al que nos permitió obtener la (9.19) es fácil verificar (lo dejamos como ejercicio para el lector) que ahora se obtiene la siguiente expresión de la densidad de estados por unidad de intervalo de energía:

$$f(\varepsilon) = 8\pi \frac{V}{h^3 c^3} \varepsilon^2 \quad (15.151)$$

Usando la (15.151) resulta

$$\varepsilon_F = hc \left(\frac{3}{8\pi} \frac{n}{V} \right)^{1/3} \quad (15.152)$$

La energía cinética total es entonces

$$E = \frac{3}{4} n \varepsilon_F \quad (15.153)$$

La presión se puede calcular a partir de la (15.123) y resulta

$$p = \frac{1}{3} \frac{E}{V} = \frac{1}{4} \frac{n}{V} \varepsilon_F \quad (15.154)$$

Para el caso de la estrella obtenemos

$$p = \gamma' \frac{hc}{R^4} \left(\frac{M}{m'} \right)^{4/3}, \quad \gamma' = \frac{\pi}{2} \left(\frac{9}{32\pi^2} \right)^{4/3} \quad (15.155)$$

Vemos entonces que cuando los electrones son ultrarelativísticos, la situación es completamente diferente al caso no relativístico, pues ahora la presión del gas de electrones es proporcional a R^{-4} , igual que la presión debida a la gravitación. La condición que ambas presiones sean iguales en valor absoluto determina un valor crítico de la masa

$$M^* = \left(\frac{\delta hc}{Gm'^{4/3}} \right)^{3/2}, \quad \delta = \frac{\gamma'}{\alpha} \quad (15.156)$$

tal que si $M < M^*$ el equilibrio es posible, pero si $M > M^*$ la presión del gas de electrones *no puede* nunca equilibrar la presión debida a la gravedad, y la contracción de la estrella continúa. Un cálculo exacto permite determinar que $\delta = 2.13$. La masa crítica M^* se denomina *masa de Chandrasekhar*²⁸, y su valor es de aproximadamente 1.45 masas solares. Toda estrella cuya masa es menor que la masa de Chandrasekhar termina su existencia como una enana blanca estable, una vez que ha agotado su combustible nuclear²⁹.

Si $M > M^*$, cuando la estrella ha agotado el hidrógeno colapsa bajo la acción de la gravedad. A medida que se produce dicho colapso, la densidad electrónica crece y por lo tanto el potencial químico de los electrones aumenta. Eventualmente, llega el momento en que se torna posible la reacción nuclear $e^- + p \rightarrow n + \nu_e$, en la cual un protón captura un electrón para transformarse en un neutrón y se emite un neutrino (esta reacción es la inversa del decaimiento β usual). De resultas de esto la estrella se convierte en una *estrella de neutrones*. La estabilidad de una estrella de neutrones³⁰ se puede estudiar del mismo modo que la de una enana blanca. Mientras la densidad no es demasiado elevada, se llega al equilibrio para un radio dado por una ley del tipo $R \sim M^{-1/3}$, como en el caso de una enana blanca. Pero si la masa es muy grande los neutrones se tornan relativísticos y hay entonces un valor crítico de la masa, más allá del cual no hay equilibrio posible. Los cálculos se complican en este caso, pues es preciso usar la Teoría General de la Relatividad dado que el campo gravitatorio es muy intenso. No existe certidumbre sobre el valor crítico de la masa de una estrella de neutrones, pero se cree que es menor que 5, y más probablemente sea de alrededor de 3 masas solares.

Las estrellas cuya masa es demasiado grande como para formar estrellas de neutrones estables no pueden detener el colapso y acaban su existencia como agujeros negros.

²⁸ Subrahmanyan Chandrasekhar, quien determinó dicho valor límite, obtuvo en 1983 el Premio Nobel por haber formulado la teoría de los estadios tardíos de la evolución de estrellas masivas.

²⁹ En particular, el Sol se convertirá en una enana blanca con un radio de unos 7000 km.

³⁰ En el curso del colapso la estrella puede explotar como una supernova, con lo cual pierde una fracción importante de su masa; en tal caso el residuo que se convierte en estrella de neutrones tiene una masa mucho menor de la que inicialmente tenía la estrella.