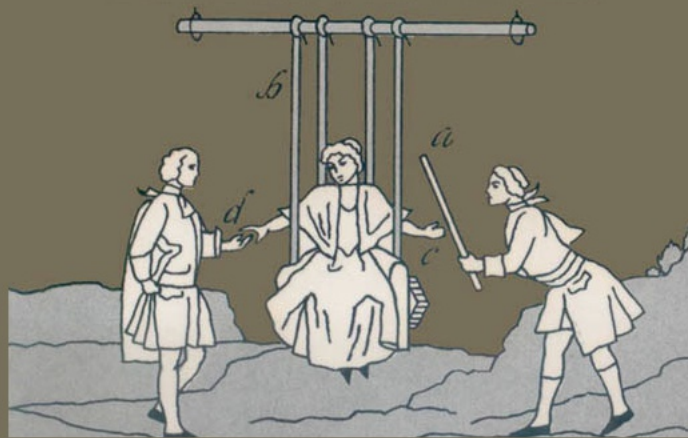


Física para Todos Libro 3

A. I. Kitaigorodski

ELECTRONES



Editorial MIR

Prefacio

El primer libro de la serie «Física para todos» dio a conocer al lector las leyes generales del movimiento de los macrocuerpos y las fuerzas de gravedad. El segundo libro estaba dedicado a la estructura molecular de la materia y al movimiento de las moléculas.

En el presente libro, en el tercero, centramos nuestra atención en examinar la estructura eléctrica de la materia, las fuerzas eléctricas y el campo electromagnético.

El siguiente libro, el cuarto, tratará de los fotones, la estructura del núcleo atómico y las fuerzas nucleares.

Los cuatro libros incluirán datos acerca de todos los conceptos y leyes fundamentales de la física. Los hechos concretos que estos libros exponen se han elegido de una forma tal que prevea la ilustración máximamente patente del contenido de las leyes físicas, la delineación de los métodos característicos para la física de analizar los fenómenos, el conocimiento de los caminos que seguía la física en el curso de su evolución y, finalmente, la demostración, a rasgos más generales, del hecho de que la física es el fundamento de todas las ciencias naturales y de la técnica.

La fisonomía de la física ha cambiado a ojos de una sola generación. Muchos de sus apartados se amplificaron transformándose en ramas independientes de enorme valor aplicado. Me permito sugerir que hoy en día uno no puede considerarse ingeniero erudito conociendo tan sólo los fundamentos de la física. Y la serie de libros con cuya ayuda los representantes de las más distintas profesiones podrán formarse idea sobre los principios de la física y ponerse al tanto de las nuevas que tuvieron lugar en las ciencias físicas durante los últimos decenios, semejante serie, precisamente, debe convertirse en física para todos.

Una vez más hago recordar al lector que el libro que tiene entre las manos no es un manual sino un libro de divulgación científica. Cuando se trata de un manual el volumen concedido a tal o cual material viene condicionado por el grado de dificultad que representa su asimilación. Un libro de divulgación científica no se rige

por esta regla y, por lo tanto sus diferentes páginas no se leen con la misma facilidad. Otra diferencia esencial radica en que en nuestros libros podemos permitirnos el lujo de exponer esquemáticamente una serie de apartados tradicionales, haciendo replegarse el viejo material para dar cabida al nuevo.

Ahora quiero referirme al libro «Electrones». La necesidad de hacer recordar las definiciones de los conceptos elementales mediante los cuales se describen los fenómenos eléctricos, esta necesidad me la he aprovechado en una forma algo peculiar, a saber: procurando dar a conocer al lector el enfoque fenomenológico de la física.

Dos capítulos de los seis están dedicados a la física aplicada. La electrotecnia se presenta en forma compendiada, ya que una descripción detallada de esta disciplina requiere recurrir a dibujos y esquemas. Esta es la razón de que estimamos posible limitarnos a la exposición solamente de los principios básicos de electrotecnia y de los importantes datos que cada uno es susceptible de saber.

Lo mismo atañe al capítulo consagrado a la radio. El pequeño volumen del libro dio la posibilidad de abordar tan sólo la historia del problema y ofrecer una exposición somera de los fundamentos de radiotecnica.

Octubre de 1981

A. I. Kitaigorodski

Capítulo 1

Electricidad

Contenido:

- *Corriente eléctrica*
- *Electricidad inmóvil*
- *Campo eléctrico*
- *Que se debe tomar por base*
- *La evolución de una teoría de la electricidad*

Corriente eléctrica

Basándose en el ejemplo de la teoría de la electricidad es posible (y, también, se debe) dar a conocer al lector que muestra interés por la física el llamado enfoque fenomenológico del estudio de la naturaleza. La palabra «fenómeno» procede del griego «*phainomenon*» que significa «lo que aparece». En cuanto al enfoque de que se trata, éste consiste en lo siguiente. El investigador no se interesa por la «naturaleza de las cosas». Se vale de las palabras únicamente para contar sobre los hechos. La finalidad del investigador no es «explicar», sino tan sólo describir el fenómeno. Casi todos los términos que introduce tienen para él un sentido solamente en el caso de que sea posible indicar el método de evaluar mediante un número de tales o cuales conceptos.

Recorre a algunas denominaciones auxiliares con el único fin de facilitar la exposición verbal de los hechos. Pero el papel que desempeñan estas denominaciones es absolutamente secundario: en lugar de las mismas hubiera sido posible proponer otros nombres o emplear «algo», o bien, «alguna cosa».

El método fenomenológico es de enorme importancia para las ciencias naturales. Y los fenómenos eléctricos, a las mil maravillas, pueden servir de ejemplo para que el lector comprenda la esencia de dicho método.

Al final de este capítulo relataré sucintamente cuál fue la secuencia en el desarrollo de los acontecimientos, mientras tanto, ahora, voy a esbozar cierto esquema ideal de la creación de la teoría fenomenológica de los fenómenos eléctricos.

Reunamos en un personaje fantástico a Carlos Augusto de Coulomb (1730 - 1806), Alejandro Volta (1745-1827), Jorge Simón Ohm (1789 - 1854), Andrés María Ampère (1775 - 1836), Juan Cristian Oersted (1777 - 1851), Emil Jristiánovich Lenz (1804 - 1805) y a algunos otros admirables hombres de ciencia. Figurémonos que a este investigador le es inherente el modo de pensar científico actual y le pongamos en la boca la terminología moderna. Precisamente en nombre de este investigador presentemos nuestro relato.

Empieza su trabajo de formación de la teoría fenomenológica de la electricidad por un examen atento del acumulador. Ante todo, presta su atención al hecho de que el acumulador tiene dos «polos». Al tocarlos simultáneamente con las manos le queda claro, de una vez, que es mejor no hacer tal cosa (porque el golpe es bastante desagradable). Sin embargo, después de esta primera experiencia se le ocurre lo siguiente: por lo visto, algo ha corrido a través de mi cuerpo. Llamemos este «algo» electricidad.

Obrando con máximo cuidado el investigador comienza a conectar los polos mediante diferentes trocitos de alambre, barritas y cordoncitos. Se convence del siguiente hecho: los objetos puestos en contacto con los polos a veces se calientan fuertemente, a veces se calientan poco y en algunos casos no se calientan, en general.

Cuando procede a la elección de palabras idóneas para caracterizar el descubrimiento hecho el investigador se decide a hablar de éste de la siguiente manera. Cuando conecto los polos mediante un alambre por este último fluye la electricidad. Voy a llamar este fenómeno corriente eléctrica. La experiencia ha demostrado que diferentes objetos se calientan de una forma disímil. Aquellos que se calientan bien evidentemente «conducen» bien la electricidad y se denominan conductores.

Muchos cuerpos se calientan mal, por lo visto, «conducen» mal la electricidad o bien crean una gran resistencia a la corriente que fluye. Y aquellos que no se calientan en absoluto se denominan aisladores o dieléctricos.

El investigador comienza a trabajar con los líquidos. Se pone de manifiesto que también en esto caso diferentes sustancias se comportan de distinta manera. Y, finalmente, se llega a un interesante descubrimiento: si se toma como líquido la

solución de vitriolo azul y se sumergen en el baño los electrodos de carbón (este nombre se da a los objetos fijados a los polos), el científico halla en uno de los electrodos el precipitado rojizo de cobre.

Ahora el investigador ya está completamente convencido de que el fenómeno que estudia tiene una relación con la circulación de cierto fluido. Queda claro que vale la pena hablar sobre la dirección de la corriente. Convenimos en marcar con el signo «menos» el electrodo en que se deposita el cobre, considerando que el segundo electrodo es positivo. Por cuanto las expresiones «electrodo negativo» y «electrodo positivo» son largas se proponen los términos «cátodo» y «ánodo», respectivamente. La corriente fluye del «más» al «menos», es decir, del ánodo al cátodo.

Pero el valor del descubrimiento está lejos de agotarse sólo con hacer constancia de este hecho. Se establece que cada segundo en el cátodo se deposita una misma masa de cobre. Seguramente que los átomos de cobre llevan en su seno el fluido eléctrico. Esta es la razón de que el investigador introduce en uso dos nuevos términos. En primer lugar, supone que la masa M del cobre es proporcional a la cantidad q de electricidad que pasó por el circuito, o sea, introduce la definición

$$q = kM$$

donde k es el coeficiente de proporcionalidad. Y, en segundo lugar, propone denominar intensidad de la corriente la cantidad de electricidad que pasa por el circuito en unidad de tiempo:

$$I = q/\tau$$

El investigador se ha enriquecido sustancialmente. Puede caracterizar la corriente por medio de dos magnitudes susceptibles de medirse: por la cantidad de calor que se libera en un tramo determinado del circuito en unidad de tiempo y por la intensidad de la corriente.

Ahora se le ofrece otra posibilidad: comparar las corrientes engendradas por diferentes fuentes. Se mide la intensidad de la corriente I , también se mide la

energía Q que se libera en forma de calor por un mismo trocito de alambre. Repitiendo los experimentos con distintos conductores el investigador averigua que la relación entre la cantidad de calor y la cantidad de electricidad que fluye a través del alambre es diferente para distintas fuentes de corriente. Sólo se requiere proponer un término apropiado para esta relación. Se ha elegido la palabra «tensión». Cuanto más alta es la tensión, tanta mayor cantidad de calor se libera.

El siguiente razonamiento puede tomarse como argumento a favor de la elección de esta palabra. Cuanto mayor es la tensión con que el hombre arrastra una carretilla con carga, tanto más calor siente. De este modo, al designar la tensión con la letra U , obtendremos

$$U = Q/q \text{ o bien } Q = U/\tau$$

Resumamos, hemos hecho los primeros pasos. Se han descubierto dos fenómenos. La corriente, al pasar a través de algunos líquidos, hace precipitar una sustancia, además, la corriente libera calor. El calor sabemos medirlo. El método para medir la cantidad de electricidad se ha dado, es decir, se ha dado la definición de este concepto. Además, se han dado las definiciones de los conceptos derivados: de la intensidad de la corriente y de la tensión.

Se ha escrito una serie de fórmulas elementales. Pero debe prestarse atención a la siguiente circunstancia: estas fórmulas no pueden llamarse leyes de naturaleza. En particular, el investigador dio el nombre de tensión a la relación Q/q pero no halló que Q/q es igual a la tensión.

Y he aquí que ha llegado la hora para buscar la ley de la naturaleza. Para un mismo conductor pueden medirse independientemente dos magnitudes: la intensidad de la corriente y el calor o la intensidad de la corriente y la tensión (que, de principio, es lo mismo).

El estudio de la dependencia entre la intensidad de la corriente y la tensión lleva al descubrimiento de una importante ley. La absoluta mayoría de los conductores está sujeta a la ley:

$$U = IR.$$

La magnitud R puede llamarse resistencia en plena correspondencia con las observaciones cualitativas iniciales. El lector conoce la notación: es la ley de Ohm. Al sustituir en la fórmula anterior el valor de la intensidad de la corriente de la expresión de la ley de Ohm, hallamos:

$$Q = \frac{U^2}{R} \tau$$

Es evidente que se puede escribir la expresión de la energía liberada por el conductor en forma de calor también de otra manera:

$$Q = I^2 R \tau$$

De la primera fórmula se infiere que la cantidad de calor es inversamente proporcional a la resistencia. Cuando se dice esta frase hay que añadir: *a tensión invariable*. Precisamente este caso se tenía en cuenta cuando por primera vez se hizo uso del término «resistencia». Mientras tanto, la segunda fórmula que afirma que el calor es directamente proporcional a la resistencia requiere que se agregue: *a una intensidad constante de la corriente*.

En las expresiones presentadas el lector reconocerá la ley que lleva los nombres de Joule y Lenz.

Después de haber establecido que la tensión y la intensidad de la corriente son proporcionales obteniendo de este modo la posibilidad de determinar la resistencia del conductor, el investigador, como es natural, se plantea la pregunta: ¿de qué manera esta importante magnitud está relacionada con la forma y las dimensiones del conductor y con la sustancia de la cual esto se ha fabricado?

Los experimentos conducen al siguiente descubrimiento. Resulta que

$$R = \rho \frac{l}{S}$$

donde l es la longitud del conductor, y S , su sección transversal. Esta expresión elemental es válida cuando se trata de un conductor lineal de sección invariable por toda su longitud. Si se quiere, recurriendo a unas operaciones matemáticas más complicadas, podemos escribir la fórmula de resistencia para el conductor de cualquier forma. Bueno, ¿y el coeficiente ρ ? ¿Qué significa esto? Dicho coeficiente caracteriza el material del cual está hecho el conductor. El valor de esta magnitud, que recibió el nombre de resistencia específica, o resistividad, oscila dentro de unos límites muy amplios. Por el valor de ρ las sustancias pueden diferenciarse miles de millones de veces.

Realicemos varias transformaciones formales más que nos serán útiles en lo adelante. La ley de Ohm puede anotarse en la forma siguiente:

$$I = \frac{US}{\rho l}$$

A menudo tenemos que ver con la relación entre la Intensidad de la corriente y el área de la sección del conductor. Esta magnitud se denomina densidad de la corriente y se suele designar con la letra j . Ahora la misma ley se escribirá así:

$$I = \frac{1}{\rho} \frac{U}{l} S$$

Al investigador le parece que en lo que concierne a la ley de Ohm todo está claro. Disponiendo de una cantidad ilimitada de conductores cuya resistencia se conoce es posible renunciar a engorrosas determinaciones de la tensión por medio del calorímetro, pues la tensión es igual al producto de la intensidad de la corriente por la resistencia.

Sin embargo, pronto, el científico llega a la conclusión de que esta afirmación necesita ser precisada. Valiéndose de una misma fuente de corriente él cierra sus polos mediante diferentes resistencias. En cada experimento la intensidad de la corriente será, naturalmente, distinta. Pero resulta que también el producto de la intensidad de la corriente por la resistencia, o sea, IR , tampoco queda el mismo. Al

dedicarse al estudio de este fenómeno, por ahora todavía incomprensible para él, el investigador averigua que a medida que aumenta la resistencia el producto IR tiende a cierto valor constante.

Al designar con ξ este límite, hallamos la fórmula que no coincide con la establecida mediante mediciones directas de la intensidad de la corriente y la tensión. La nueva expresión tiene la siguiente forma:

$$\xi = I (R + r)$$

¿Por qué tan extraña contradicción?

Es necesario recapacitar. Ah, claro está, la contradicción es aparente. Es que la medición directa de la tensión empleando el método calorimétrico se refería tan sólo al conductor que cerraba el acumulador. Mientras tanto se ve claramente que el calor se desprende también en el propio acumulador (para cerciorarse de ello es suficiente tocar el acumulador con la mano). El acumulador posee su propia resistencia. El sentido de la magnitud r que aparece en la nueva fórmula es evidente: es la resistencia interna de la fuente de la corriente. En cuanto a ξ esta magnitud requiere una denominación especial. No se puede decir que la misma resultó ser muy acertada: la magnitud ξ se llama fuerza electromotriz (f.e.m.) aunque no tiene significado ni tampoco dimensión de la fuerza.

Las dos fórmulas conservaron (cabe señalar que en este caso se ha observado la justicia histórica) el nombre de leyes de Ohm. Únicamente, la primera fórmula recibió el nombre de ley de Ohm para una porción del circuito, mientras que la segunda se llama ley de Ohm para el circuito total.

Vaya que ahora, al parecer, no quedan ya dudas. Las leyes de la corriente continua están establecidas.

No obstante, el investigador no se ve satisfecho. El empleo del calorímetro resulta engorroso. Por si esto fuera poco, ¡hace falta pesar el cátodo con el precipitado de cobre! No puedo negar que es un método muy incómodo de medir la tensión.

Un buen día, ¡de veras que lo fue! el investigador, por pura casualidad, ubicó junto al conductor por el cual circulaba la corriente una aguja magnética. Él hizo un gran

descubrimiento: la aguja gira cuando pasa la corriente, con la particularidad de que lo hace en distintas direcciones, según sea la dirección de la corriente.

No es difícil determinar el momento de la fuerza que actúa sobre la aguja magnética. Basándose en el fenómeno descubierto es posible crear un instrumento de medida. Únicamente se necesita establecer el carácter de la dependencia del momento de la fuerza respecto a la corriente. El investigador resuelve este problema y construye magníficos instrumentos de aguja que permiten medir la intensidad de la corriente y la tensión.

Sin embargo, nuestro relato sobre aquello que el investigador realizó en la primera mitad del siglo diecinueve al estudiar las leyes de la corriente continua sería incompleto, si no señalásemos que descubrió la interacción de las corrientes: las corrientes que se dirigían en un mismo sentido se atraían, mientras las de dirección diferente se repelían. Se sobreentiende que este fenómeno también se puede utilizar para medir la intensidad de la corriente.

Desde luego, no me limitaré a los últimos párrafos al hablar sobre las leyes del electromagnetismo; a este fenómeno se dedica un capítulo aparte. Pero he considerado indispensable recordar estos importantes datos con el fin de contar cómo se introducen los conceptos cuantitativos fundamentales y las unidades de medida que caracterizan los fenómenos eléctricos: la corriente, la carga y el campo.

Electricidad inmóvil

Demos por sentado que nuestro investigador ideal está enterado de los variados fenómenos que, en los tiempos remotos, habían obtenido el nombre de eléctricos. Las propiedades peculiares del ámbar, de una varilla de vidrio frotada con piel, la aparición de una chispa que saltaba entre dos cuerpos llevados a estado «electrizado» se estudiaban (o, mejor dicho, se aprovechaban para crear efectos) ya hacía mucho. Por esta razón era lógico que el investigador, al abordar el estudio de la corriente eléctrica, se planteara la pregunta: ¿el fluido que circula por el conductor y el fluido que puede permanecer en estado inmóvil sobre cierto cuerpo hasta que no lo «descarguen», es que ambos constituyen el mismo «algo»? Por lo demás, incluso abstrayéndose de la información acumulada anteriormente, ¿acaso

uno no debe poner a sí mismo la siguiente pregunta: si la electricidad es «algo» que fluye a semejanza de un líquido no sería posible «verterlo en un vaso»?

Si el investigador quisiera obtener una respuesta directa a esta pregunta, tendría que proceder de la siguiente manera. Se toma una fuente de corriente con una tensión bastante alta (por ahora no hablamos sobre las unidades de medida, por lo tanto el lector, sin impacientarse, debe esperar la respuesta a la pregunta qué se suele considerar alta tensión, qué es una gran intensidad de la corriente, etc.). Uno de los polos se pone a tierra y sobre el segundo se coloca una pequeña bolita, abalorio hecho de hoja de aluminio muy fina. La bolita se suspende de un hilo de seda, de la misma manera se procede con otra bolita.

Ahora arrimamos estas dos minúsculas bolitas muy cerca una a otra (digamos, a una distancia de 2 mm entre sus centros). El investigador con entusiasmo, con admiración (puede proponer cualquier otro epíteto) observa que las bolitas se repelen. Por el ángulo de desviación de los hilos y conociendo la masa de las bolitas puede calcularse la fuerza que actúa entre las mismas.

El investigador saca la conclusión: si las bolitas están cargadas por contacto con el mismo polo del acumulador éstas se repelen. En cambio, si una bolita recibió la electricidad de un polo y la otra bolita del otro polo, entonces las bolitas se atraerán.

Este experimento corrobora que tenemos el derecho de hablar sobre la electricidad como si fuera un líquido y demuestra que se puede tratar tanto con la electricidad móvil, como con la en reposo.

Por cuanto el investigador sabe determinar la cantidad de electricidad por la masa del cobre depositado en el cátodo, existe la posibilidad de aclarar «cuánto líquido se ha vertido en el vaso», es decir, cuál es la cantidad de electricidad que la bolita «se ha apropiado» del electrodo del acumulador.

En primer lugar el investigador se convence de lo siguiente. Si la bolita cargada «se pone a tierra», es decir, si se conecta mediante un conductor a la Tierra, la bolita pierde su carga. Seguidamente se demuestra que la carga «escurre» por el conductor, o sea, que por el conductor fluye la corriente. Y, finalmente, se tiene la posibilidad de medir la cantidad de cobre que precipita en el cátodo de un aparato

con electrólito interpuesto en el camino a la Tierra, es decir, se puede medir la cantidad de la electricidad inmóvil que se encontraba en la bolita.

Esta cantidad de electricidad el investigador la denomina carga de la bolita y le atribuye un signo: positivo o negativo, según sea el electrodo del cual se ha tomado el fluido eléctrico.

Ahora se puede iniciar la siguiente serie de experimentos. Desde diferentes acumuladores, valiéndose de bolitas de distintas dimensiones, pueden tomarse diferentes cantidades de electricidad. Al colocar las bolitas a diferentes distancias unas de otras es posible medir la fuerza de interacción entre éstas. El investigador halla lo siguiente importante ley de la naturaleza:

$$F = K \frac{q_1 q_2}{r^2}$$

la fuerza de interacción es directamente proporcional al producto de las cargas de las bolitas e inversamente proporcional al cuadrado de distancia entre éstas. El lector reconocerá en la fórmula que acabamos de escribir la ley de Coulomb que fue establecida de una manera absolutamente distinta a la que exponemos. Pero no olvide que nuestro investigador es un personaje extrahistórico.

Campo eléctrico

El investigador conoce fuerzas de dos tipos. Unas fuerzas aparecen durante el contacto directo de un cuerpo con otro. Se presentan en el caso de tracción o de empuje. En cuanto a las fuerzas que actúan a una distancia, hasta el momento el investigador tenía noción solamente de la fuerza de la gravedad o, en mayor escala, la fuerza de gravitación universal.

Ahora a esta fuerza que ya conocía se agregó otra: la de atracción o repulsión culombiana entre dos cuerpos cargados. Esta fuerza se parece mucho a la fuerza de la gravedad.

Hasta las formulas también se hacen recordar mutuamente.

La fuerza de la gravedad que actúa sobre el cuerpo por parte de la Tierra no ponía grandes inconvenientes durante los cálculos. En lo que se refiere a las fuerzas

coulombianas o, como también se llaman, fuerzas electrostáticas, aquí se puede topar con los casos en que las cargas eléctricas están distribuidas en el espacio de una forma muy complicada y, además, desconocida.

Sin embargo, se puede pasar sin conocer la distribución de estas cargas. Ya estamos enterados de que estas cargas «se sienten» unas a otras a una distancia. ¿Por qué no decir: las cargas crean el campo eléctrico? Puede parecer que debe surgir una dificultad a raíz de que no vemos ningún campo eléctrico. Pero mi opinión es, dice el investigador, que el campo eléctrico no debe considerarse como una función matemática que facilita el cálculo. Si sobre una carga situada en cierto punto actúa una fuerza, este hecho significa que dicho punto (del espacio) se encuentra en un estado especial. El campo eléctrico es una realidad física, es decir, existe por sí mismo, aunque no lo vemos. Por supuesto, el investigador que trabaja a principios del siglo XIX no puede demostrar su pensamiento. Pero el futuro manifestaría que él tenía razón.

La ley de Coulomb establece la fórmula con cuya ayuda se puede determinar la acción que una bolita ejerce sobre la otra. Una bolita la podemos dejar fija, mientras que la segunda se colocará en distintos puntos del espacio. En todos los puntos sobre la bolita móvil (de ensayo) actuará una fuerza. Ahora el mismo hecho se enuncia de otra forma: una bolita cargada de electricidad crea en su torno un campo de fuerzas eléctricas, o, más brevemente, un campo eléctrico.

De fuente del campo eléctrico pueden servir cuerpos cargados de cualesquiera formas. En este caso la ley de Coulomb ya no es válida, pero recurriendo a la bolita de ensayo es posible medir el campo eléctrico que rodea el cuerpo cargado y definirlo de un modo completamente exhaustivo, indicando la magnitud y la dirección de la fuerza. Para conseguir que la descripción del campo sea independiente de la elección de la magnitud de la carga de la bolita de ensayo, el campo eléctrico se caracteriza por la magnitud llamada su intensidad:

$$E = F/q$$

donde q es la carga eléctrica de la bolita de ensayo.

Existe un método patente de representación del campo eléctrico por medio de líneas de fuerza (líneas de intensidad). En dependencia de la forma de los cuerpos cargados y de su disposición mutua estos gráficos pueden tener el más variado aspecto. En la fig. 1.1 se muestran los cuadros más simples de los campos.

El sentido de estos cuadros es el siguiente: la tangente a la línea de intensidad de un punto cualquiera indica la dirección de la fuerza eléctrica en este punto. El número de líneas que corresponden a una unidad del área perpendicular a las líneas de intensidad es absolutamente convencional, lo único que se exige es que sea proporcional al valor de E . Y en el caso en que se habla del número de líneas de intensidad sin utilizar los cuadros, se supone sencillamente que este número es igual al valor de E .

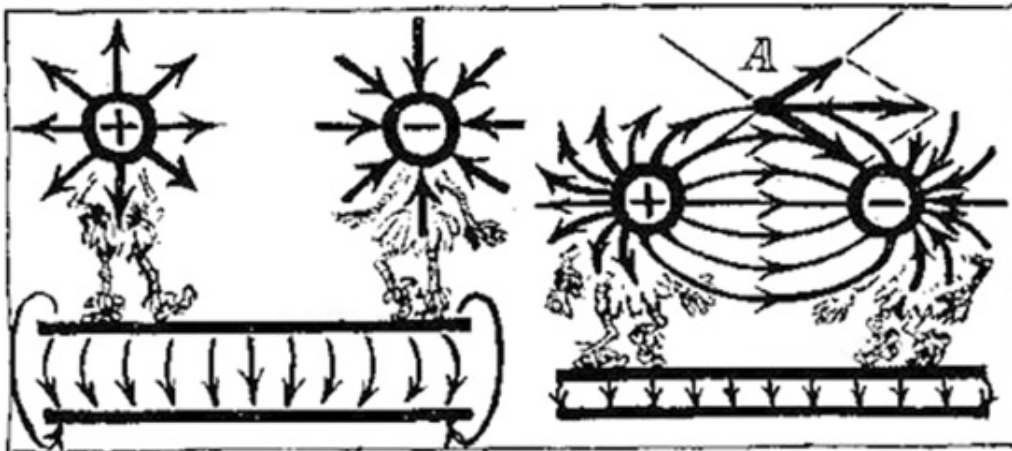


Figura 1.1

Si una carga eléctrica libre se sitúa en el campo eléctrico, dicha carga se desplazará a lo largo de las líneas de intensidad, a menos que tercier otras fuerzas, por ejemplo, las de la gravedad.

El aspecto más sencillo lo tienen los campos eléctricos de los cuerpos de forma esférica. Si estas esferas están muy distanciadas el cuadro de las líneas de intensidad se representa como en la fig. 1.1, a la izquierda. Si dos esferas o dos cargas que se pueden representar como puntos se hacen acercarse unas a otras, los campos se van a superponer. La intensidad del campo resultante se obtendrá por la regla del paralelogramo. Al realizar la construcción mostrada en la fig. 1.1, a la

derecha, se puede esclarecer cuál es la dirección de la línea de intensidad y a qué es igual la intensidad del campo dado en cualquier punto A.

Si los cuerpos cargados tienen la forma de láminas, el aspecto del campo será tal como se muestra en la parte inferior de la figura. Al aproximar las láminas y aumentar su área se puede conseguir una homogeneidad casi ideal del campo; el efecto de frontera será insignificante. Dos cuerpos cargados dispuestos uno cerca del otro se denomina condensador.

Como sabemos, el trabajo de traslación de un cuerpo bajo la acción de una fuerza es igual al producto de la fuerza por la longitud del camino. Para transferir la carga de una placa del condensador a la otra siguiendo a lo largo de la línea de intensidad se requiere un trabajo igual a qEl . El trabajo necesario para la transferencia de una unidad de la cantidad de electricidad es igual a E .

Unamos dos placas del condensador mediante un conductor. Cuando por el conductor se transfiere la cantidad de electricidad q se libera la energía qU . Por cuanto, a todas luces, no se da una diferencia de principio entre el movimiento de la bolita cargada en el campo eléctrico y el desplazamiento del «líquido» eléctrico a lo largo de un conductor metálico, igualamos entre sí estas dos expresiones de la energía invertida por el campo:

$$qEl = qU$$

La validez de la expresión escrita puede comprobarse fácilmente distanciando las placas del condensador y midiendo la fuerza que actúa sobre la carga de ensayo.

Dicha medición puede llevarse a cabo por un método muy elegante sin recurrir para nada a la suspensión de la bolita cargada de un hilo de seda.

Todo el mundo conoce muy bien que en el aire los cuerpos ligeros caen mucho más lentamente que los pesados. Cabe señalar que precisamente por esta causa con anterioridad a los experimentos de Galileo, los sabios de la Antigüedad y del Medioevo suponían que la velocidad de movimiento del cuerpo (y no la aceleración) es proporcional a la fuerza. El carácter erróneo de este punto de vista se demostró espectacularmente sólo cuando se fijaron cómo caían los pedacitos de papel y una bola metálica en un tubo vertical del que se había succionado el aire. Resultó que

todos los cuerpos cobran velocidad con la misma rapidez, es decir, caen a la Tierra con la misma aceleración. Sin embargo, precisamente en esta ocasión para nosotros tiene sentido «incluir» la influencia del aire cuya resistencia implica el hecho de que la ligera bolita metálica hueca valiéndose de la cual demostramos la ley de Coulomb caiga muy lentamente.

Si hacemos caer la bolita cuando esta se encuentra entre las placas del condensador, entonces, variando la tensión entre las placas, es posible elegir un campo que retenga la caída de dicha bolita. El equilibrio se alcanza a condición de que la fuerza de la gravedad sea igual a la fuerza del campo

$$mg = qE.$$

A partir de esta igualdad se puede hallar el valor de la intensidad del campo y confirmar la certeza de nuestros razonamientos teóricos.

El número de líneas de intensidad que pasan a través de cualquier superficie mental o real que se encuentra en el campo eléctrico se denomina flujo de líneas de intensidad. ¿A qué es igual el flujo de líneas de intensidad que atraviesa una superficie cerrada la cual abarca los cuerpos cargados?

Al principio analicemos un caso elemental en que el campo resulta creado por una sola bolita. Describamos alrededor de la bolita una esfera. Si el radio de la esfera es R y entonces, la intensidad en cualquier punto de la superficie de la esfera es igual a Kq/R^2 . El área de la esfera es igual a $4\pi R^2$. En consecuencia, el flujo de líneas de intensidad que atraviesa la esfera será igual a $4\pi Kq$. Sin embargo, está claro que el flujo quedará el mismo si tomamos cualquier otra superficie.

Ahora hagamos más complejo el cuadro, suponiendo que el campo es engendrado por un gran número de cuerpos cargarlos de cualquier forma. Pero es que podemos dividirlos mental mente en porciones ínfimas cada una de las cuales sea equivalente a una carga puntual. Abarquemos el sistema de las cargas con una superficie arbitraria. El flujo procedente de cada carga es igual a $4\pi Kq$. Resulta muy natural la suposición de que los flujos se sumen aritméticamente, y, por consiguiente, el flujo total a través de cualquier superficie cerrada que abarca todas las cargas es proporcional a la carga total de los cuerpos que se encuentran dentro de esta

superficie. Esta afirmación es la ley fundamental que rige los campos electromagnéticos (una de los cuatro ecuaciones de Maxwell, véase el Capítulo 5).

Preste atención a que no hemos deducido ni demostrado esta fórmula. Hemos adivinado que el asunto debe ir así y no de otra manera. Ello, precisamente, significa que tenemos que ver con una ley general de la naturaleza cuya justedad se establece por la confirmación experimental de cualquier corolario que se derive de la ley general.

Es muy importante conocer una regla que sea válida para cualesquiera sistemas.

Con la ayuda de una ley escrita, un ordenador calculará rápidamente el campo eléctrico creado por el más complejo sistema de cuerpos cargados. Entre tanto, nosotros nos satisfacemos con un problema modesto, deduciendo (y demostrando, valiéndose de este caso elemental, los procedimientos de la física teórica) una fórmula de valor práctico para la capacidad del condensador.

Primeramente demos la definición de este concepto difundido. Se denomina capacidad del condensador la relación entre la carga que se acumula en sus placas y la tensión entre las armaduras, es decir,

$$C = q/U$$

El término «capacidad» es acertado. Efectivamente, a una tensión dada la carga que toma el condensador depende tan sólo del tamaño y la forma de las placas.

En el caso del condensador las líneas de intensidad no se dirigen hacia los lados sino salen de la placa positiva y entran en la negativa. Si se desprecia la deformación del campo en los extremos del condensador, el flujo puede expresarse como producto ES . La ley general da la posibilidad de escribir la siguiente igualdad:

$$ES = 4\pi Kq,$$

es decir, la intensidad del campo entre las armaduras es

$$E = 4\pi K \frac{q}{S}$$

Por otra parte, la intensidad del campo del condensador puede anotarse como

$$E = U/d.$$

Igualando estas dos expresiones obtenemos la fórmula para la capacidad del condensador:

$$C = \frac{S}{4\pi K d}$$

Los condensadores técnicos son cintas metálicas que están en estrecho contacto con mica o con papel parafinado. Estas sustancias pertenecen a los dieléctricos. ¿Qué sentido tiene la introducción del dieléctrico entre las armaduras del condensador? La experiencia demuestra que la capacidad del condensador C está relacionada con la capacidad del condensador sin junta C_0 por la fórmula $C = \epsilon C_0$.

La magnitud ϵ lleva el nombre de constante dieléctrica. Los valores de ϵ para el aire, mica, agua y la sal de Seignette son iguales, respectivamente, a 1, aproximadamente 6, 81 y 9000.

Qué se debe tomar por base

La ley de Ohm y la ley de Joule - Lenz vinculan entre sí la energía, la intensidad de la corriente, la tensión y la resistencia. Se puede decir que la tensión es igual al producto de la intensidad de la corriente por la resistencia. También es posible decir: la intensidad de la corriente es la tensión dividida por la resistencia. Sin embargo, estas dos definiciones que se pueden encontrar en los libros de texto llevan implícito un inconveniente: el de ser cómodos únicamente en el caso de que sea válida la ley de Ohm. Pero, como hemos dicho, esta ley no siempre resulta certera. Esta es la razón por la cual lo mejor es proceder de la forma como ya hemos hecho, o sea, considerar, precisamente, que la magnitud derivada es la resistencia del conductor que se define como la relación de la tensión en los extremos del conductor a la intensidad de la corriente que fluye a través de éste.

Por cuanto la energía de la corriente eléctrica puede medirse partiendo de la ley de la conservación de la energía, es decir, basándose en las acciones térmicas y mecánicas de la corriente, queda claro el carácter racional de definir la intensidad de la corriente o la tensión como magnitudes derivadas de energía. Lo más natural es determinar la intensidad de la corriente valiéndose del fenómeno de la electrólisis, y la tensión en los extremos de un tramo del circuito como el cociente de la división de la energía liberada por la cantidad de electricidad.

No obstante, el lector debe darse cuenta, con toda claridad, de que este sistema de definiciones no es el único. En vez de la electrólisis, como base de determinación de la intensidad de la corriente, puede elegirse también cualquier otra acción de ésta: por ejemplo, la acción de la corriente sobre la aguja magnética o sobre otra corriente.

De principio, no hay nada vicioso en el siguiente camino: se elige cierta fuente de corriente normalizada y la tensión de cualquier otra fuente se determina por la cantidad de elementos normalizados equivalentes. No es una fantasía. Semejante proposición tuvo lugar y la fuente normalizada lleva el nombre de pila Weston.

Existe también otra variante: se puede construir el sistema de definiciones y unidades de medida eligiendo cierta resistencia patrón, y midiendo, igual que antes, todas las demás resistencias después de poner en claro cuántos elementos normalizados pueden sustituir el conductor en cuestión. En su tiempo, como tal unidad de resistencia se empleaba una columna de mercurio de longitud y sección prefijadas.

Es útil siempre tener presente que la secuencia en que se introducen los conceptos físicos es una cosa arbitraria. Por supuesto, el contenido de las leyes de la naturaleza no se altera debido a ello.

Hasta el momento teníamos que ver con los fenómenos eléctricos relacionados con la corriente eléctrica continua. Incluso sin rebasar los marcos de este grupo de fenómenos se brinda la posibilidad de construir diferentes sistemas de definiciones de los conceptos y, respectivamente, distintos sistemas de unidades de medida. Pero, en realidad, la posibilidad de elegir resulta ser aún más amplia, ya que los fenómenos eléctricos no se reducen, en modo alguno, a la corriente eléctrica continua.

Hasta la fecha, muchos libros de texto físicos definen el concepto de magnitud de la carga eléctrica (o, que es lo mismo, de la cantidad de electricidad) a partir de la ley de Coulomb, seguidamente, en la escena se presenta la tensión y tan sólo al fin y a la postre, una vez terminada la exposición de la electrostática, el autor introduce los conceptos de intensidad de la corriente y de resistencia eléctrica. Como ha visto el lector, seguimos por otro camino.

Todavía más arbitrio se puede observar en la elección de las unidades de las magnitudes físicas. El investigador tiene el derecho de proceder tal como le parece más conveniente. Solamente no debe olvidar que la elección de las unidades repercutirá en los coeficientes de proporcionalidad que forman parte de diferentes fórmulas.

No hay nada malo en escoger independientemente las unidades de la intensidad de la corriente, de la tensión y de la resistencia. Más en este caso, en la fórmula de la ley de Ohm aparecerá cierto coeficiente numérico que posee dimensión. Hasta el último tiempo, mientras el severo veredicto de la Comisión Internacional no haya expulsado todavía de la física las tan familiares calorías, la fórmula de la ley de Joule - Lenz contenía un coeficiente numérico. La causa de ello residía en el hecho de que las unidades de la intensidad de la corriente y de la tensión se determinaban de una forma completamente independiente con respecto a la elección de la unidad de energía (calor, trabajo).

En los párrafos anteriores he escrito en forma de proporcionalidades y no igualdades solamente dos fórmulas: aquella que relaciona la masa de la sustancia depositada en el electrodo con la cantidad de electricidad y la ley de Coulomb. No lo he hecho casualmente, sino por la sencilla razón de que los físicos, por ahora, muy a desgana pasan al Sistema Internacional, SI, adoptado como ley, y siguen empleando todavía el llamado sistema absoluto de unidades en que el valor de K en la fórmula de Coulomb para la interacción de las cargas en el vacío se toma igual a la unidad. Al obrar de esta manera, predeterminamos el valor de la llamada unidad «absoluta» de la cantidad de electricidad (la carga es igual a la unidad si dos cargas iguales situadas a una distancia unitaria interaccionan con una fuerza unitaria).

De ser consecuentes, entonces, al medir la masa en gramos, tendríamos que calcular el valor del coeficiente h en la ley de la electrólisis, indicando qué cantidad

de sustancia se deposita en el electrodo en una unidad absoluta de carga. Sin embargo, absténgase de hojear las páginas de los manuales, no encontrara semejante valor para dicho coeficiente. Por cuanto los físicos estaban enterados de la categórica oposición de los físicos a renunciar al amperio y culombio, los primeros instituían en la fórmula de la electrólisis aquel número que determinaba la masa de sustancia que precipitaba al pasar a través del líquido un culombio de electricidad. En los libros figuraban dos unidades para una misma magnitud. Con todo, estaba claro que el empleo de una o de otra era conveniente en los casos completamente diferentes, pues un culombio equivalía a tres mil millones unidades absolutas.

Sin duda alguna, es cómodo suponer que K es igual a uno, pero los técnicos prestaban atención a que en las ecuaciones para el flujo de líneas de intensidad, de la capacidad del condensador y en otras fórmulas queda el coeficiente 4π que no hace falta a nadie y afirmaban que sería útil librarse de éste.

Como suele suceder, vencieron aquellos quienes se encontraban más próximos a la práctica, que a la teoría, el sistema adoptado actualmente tomó el camino que hace mucho ya siguieron los técnicos. Los partidarios del sistema SI consiguieron también que emplease una sola unidad de energía en todos los campos de la ciencia, exigiendo, además, que como unidad eléctrica fundamental y única figúrese la intensidad de la corriente.

De este modo entramos en el estudio de la electricidad cuya unidad de energía es julio. Como unidad de la cantidad de electricidad elegimos el culombio igual al amperio-segundo. Proponemos definir el amperio por la intensidad de interacción de las corrientes. Esta definición (la insertaremos en páginas siguientes, en el capítulo dedicado al electromagnetismo) se elige de modo que el coeficiente k en la fórmula de la electrólisis resulte ser el mismo a que todo el mundo se ha acostumbrado hace mucho. No obstante, hay que tener presente que este coeficiente en el sistema SI no define la magnitud del culombio. Si la exactitud de la medición crece, nos veremos obligados a medir esta magnitud de una forma tal que se conserve la definición del amperio (a decir verdad, no creo que este tiempo llegue, ya que no me puedo figurar que la exactitud en la medición de las fuerzas electrodinámicas supere la de la medición de la masa).

En adelante, el sistema SI sigue por el camino que yo había obligado a recorrer a nuestro investigador. Aparece la unidad de tensión, el voltio, igual al julio dividido por el culombio; la unidad de resistencia, el ohmio, igual al voltio dividido por el amperio; la unidad de resistencia específica: el ohmio multiplicado por el metro.

Pero ahora llegamos a la ley de Coulomb y vemos que ya no tenemos el derecho de manipular arbitrariamente con el coeficiente K. La fuerza se mide en newtons; la distancia, en metros, y la carga, en culombios. El coeficiente K se convierte en dimensional y tiene cierto valor que debe determinarse por vía experimental.

A la ley de Coulomb se suele recurrir raras veces, mientras que la expresión de la capacidad del condensador es la fórmula de trabajo en muchos cálculos técnicos. Con el fin de librarse del factor 4π en las fórmulas del flujo de líneas de intensidad, de la capacidad del condensador y en muchas otras los técnicos ya hace mucho han sustituido el coeficiente K por la expresión $1/4\pi\epsilon_0$. Debido a razones completamente comprensibles ϵ_0 puede llamarse permeabilidad dieléctrica del vacío, sin embargo, oficialmente, esta magnitud se denomina constante eléctrica. Esta resulta ser igual a

$$\epsilon_0 = \frac{8,85 \times 10^{-12} \text{ C}^2}{(\text{N m}^2)}$$

De este modo, ahora el flujo de líneas de intensidad se expresa por medio de la fórmula

$$\frac{(q_1 + q_2 + \dots)}{\epsilon_0}$$

y la capacidad del condensador

$$C = \frac{\epsilon \epsilon_0 S}{d}$$

La unidad de capacidad, un faradio, es igual a un culombio dividido por voltio:

$$1 \Phi = 1 \text{ C/V.}$$

La evolución de la teoría de la electricidad

El orden de secuencia en que actuaba nuestro investigador «sintetizado» no tiene nada que ver con la real evolución de la teoría de la electricidad.

Los fenómenos electrostáticos se conocían ya en la remota antigüedad. Es difícil decir si los sabios griegos conocían qué cuerpos, además del ámbar (en griego «electrón» significa «ámbar») adquirían, después de frotarlos, unas propiedades especiales, atrayendo las pajitas. Tan sólo en el siglo XVII William Gilbert demostró que esta extraña propiedad la poseían el diamante, el lacre, el azufre, el alumbre y muchos otros cuerpos. Al parecer, este ilustre hombre de ciencia fue el primero en crear instrumentos con cuya ayuda se podía observar la interacción de los cuerpos electrizados. En el siglo XVIII ya no se ignoraba que algunos cuerpos eran capaces de retener las cargas, mientras que por otros cuerpos las cargas «escurren-». Son pocos los que ponen en tela de juicio el hecho de que la electricidad es algo como un fluido. Se crean las primeras máquinas electrostáticas por cuyo medio pueden generarse chispas y realizarse el siguiente experimento: se hace «estremecerse» una hilera de hombres en la cual cada uno está cogido de la mano del vecino y el primero toca el conductor de la máquina eléctrica en función. La alta sociedad de muchos países visita los laboratorios de los científicos como si éstos fuesen un circo. Y los científicos, a su vez, tratan de impartir a los correspondientes fenómenos un carácter máximamente teatral.

En el siglo XVIII ya se puede hablar sobre la electrostática como de una ciencia. Está fabricada gran cantidad de diferentes electros copios y Coulomb comienza a efectuar mediciones cuantitativas de las fuerzas de interacción de las cargas.

En 1773, Luis Galvani (1737 - 1798) comenzó a investigar las contracciones musculares de la rana operadas bajo la acción de la tensión eléctrica.

Continuando los experimentos de Galvani, Volta, a finales del siglo XVIII llega a comprender que por los músculos de la rana corre un fluido eléctrico. El siguiente paso notorio es la creación de la primera fuente de corriente eléctrica: de una pila galvánica, y más tarde, también de la pila de Volta.

Al despuntar el siglo XIX la noticia sobre el descubrimiento de Volta ya se difundió por todo el mundo científico. Comienza el estudio de la corriente eléctrica. Un descubrimiento sigue tras otro.

Una serie de investigadores estudia la acción térmica de la corriente. También estaba dedicado a estos estudios Oersted que, electivamente, por pura casualidad, descubrió la acción de la corriente sobre la aguja magnética.

Los brillantes trabajos de Ohm y Ampère fueron realizados aproximadamente en un mismo período: en los años 20 del siglo XIX.

Los trabajos de Ampère, rápidamente, le granjearon una gran fama. En cambio, Ohm no tenía suerte. Sus artículos en los cuales el escurioso experimento se compaginaba con cálculos precisos y que se distinguían por su carácter riguroso y la introducción consecuente de los conceptos fenomenológicos, descartando absolutamente la «naturaleza» de las cosas, no acapararon la atención de los contemporáneos.

Es extraordinariamente difícil leer los trabajos originales de los físicos que habían trabajado en aquellos tiempos. Sus hallazgos experimentales se exponen en un lenguaje ajeno para nosotros. En una serie de casos ni siquiera es posible comprender qué entendía el autor empleando tal o cual palabra. Los nombres de los grandes científicos viven en la memoria de los descendientes tan sólo gracias a la atenta labor de los historiadores de la ciencia.

Capítulo 2

Estructura eléctrica de la materia

Contenido:

- *Porción mínima de electricidad*
- *Flujos iónicos*
- *Rayo electrónico*
- *Experimento de Millikan*
- *Modelo del átomo*
- *Cuantificación de la energía*
- *La ley periódica de Mendeleiev*
- *Estructura eléctrica de las moléculas*
- *Los dieléctricos*
- *Conductibilidad de los gases*
- *Descarga autónoma*
- *Sustancia en estado de plasma*
- *Metales*
- *Salida de los electrones del metal*
- *Fenómenos termoeléctricos*
- *Semiconductores*
- *Unión p—n*

Porción mínima de electricidad

Durante un largo período todos los datos que tenían los físicos en lo que concernía a los fenómenos eléctricos se reducían a la seguridad de que la electricidad es algo como un líquido. Aun a finales del siglo XIX estaba en boga la siguiente anécdota. Un examinador deseoso de reírse sobre el estudiante no preparado dice: «Bueno, como usted no pudo dar respuesta a ninguna de mis preguntas, permítame que le ponga una más, de lo más simple: ¿qué es la electricidad?» El estudiante contesta: «Señor profesor, palabra de honor que lo conocía, pero se me olvidó». El examinador exclama. «¡Qué pérdida para la humanidad! Había una persona que sabía qué es la electricidad, pero hasta esa persona lo olvidó».

Las primeras conjeturas acerca de que la electricidad no es un líquido continuo, sino consta de unas partículas especiales, así como la seguridad de que las partículas eléctricas están vinculadas, de cierta forma, con los átomos se han promovido basándose en el estudio de la electrólisis.

Al realizar los experimentos relacionados con la descomposición de las sustancias disueltas en agua durante el paso de la corriente a través de la solución, Miguel Faraday (1791 - 1867) había establecido que una misma corriente eléctrica da lugar a la sedimentación en los electrodos de distinta cantidad de sustancia en dependencia de qué compuesto químico estaba disuelto en agua. Faraday halló que al depositarse un mol de una sustancia monovalente a través del electrolito pasaban 96.500 C, mientras que con la sedimentación de un mol de una sustancia divalente esta cantidad se duplicaba.

¿Usted habrá pensado, tal vez, que, al llegar a este resultado, Faraday exclamara «ieureka!», declarando que esclareció la naturaleza de la electricidad? De ningún modo. El gran experimentador no se permitió tal fantasía. Faraday, en todo caso, en lo que concernía a la corriente eléctrica, se comportaba como el protagonista del capítulo anterior. Estimaba necesario utilizar tan sólo aquellos conceptos que se podían caracterizar con un número.

¿Cómo puede ser?, preguntará, el lector. ¿Acaso no se ha demostrado que $6,02 \times 10^{23}$ (se acuerda de que es el número de Avogadro) átomos transportan 96.500 culombios de electricidad? En consecuencia, al dividir el segundo número por el primero, obtendré la cantidad de electricidad que transporta cada uno de los átomos monovalentes. La operación de división nos da $1,6 \times 10^{19}$ culombios. ¡He aquí la mínima porción de electricidad, o bien «átomo de electricidad», o bien «carga elemental»!

Pero el número de Avogadro fue determinado tan sólo para el año 1870. Solamente entonces (figúrese, inada más que hace un siglo!) los físicos propensos a inventar hipótesis (su temperamento y modo de pensar los diferencia sustancialmente del investigador, a quien no gusta salir de los marcos del fenómeno) llegaron a la conclusión de que es muy probable la siguiente conjetura. A la par de átomos eléctricamente neutros existen partículas que transportan en su seno una o varias cargas elementales de electricidad (positiva o negativa). Los átomos que portan la

carga positiva (los cationes) se depositan durante la electrólisis en el cátodo; los átomos que son portadores de la carga negativa (los aniones) se depositan en el ánodo.

Las moléculas de las sales solubles en agua se disocian en aniones y cationes, por ejemplo, la molécula de la sal común, o sea, del cloruro de sodio, no se disocia en átomos de cloro y átomos de sodio, si no en el ion positivo de sodio y el ion negativo de cloro.

Flujos iónicos

Se sobreentiende que el fenómeno de la electrólisis solamente sugiere al investigador la idea sobre la existencia de las partículas eléctricas.

A finales del siglo XIX y principios del XX se propusieron numerosos procedimientos para transformar las moléculas en fragmentos cargados (dicho fenómeno se denomina ionización), se demostró por qué vía podían crearse flujos dirigidos de partículas cargadas y, por fin, se elaboraron métodos de medición de la carga y de la masa de los iones. El primer conocimiento con los flujos iónicos los físicos lo trabaron conectando al circuito de la corriente eléctrica un tubo de vidrio lleno de gas enrarecido. Si la tensión en los electrodos que están soldados en el tubo es pequeña, la corriente no pasará a través de este tubo. Sin embargo, resulta que no es nada difícil convertir el gas en conductor. La acción de los rayos X, de la luz ultravioleta y de la irradiación radiactiva conduce a la ionización del gas. También se pueden sortear las medidas especiales, pero en este caso es necesario aplicar al tubo de gas una tensión más alta.

¡El gas se convierte en conductor de corriente! Se puede suponer que las moléculas se fragmentan en aniones y cationes. Los aniones se dirigen al electrodo positivo, y los cationes, al negativo. Una etapa importante en el estudio de este fenómeno lo constituía la creación del flujo de partículas. Con este objeto es preciso practicar en el electrodo un orificio y acelerar por medio de un campo eléctrico los iones de un mismo signo que han atravesado dicho orificio. Empleando dos diafragmas es posible crear un haz estrecho de aniones o cationes que se muevan a una velocidad considerable. Si el haz incide sobre una pantalla del tipo de la utilizada en el televisor, observaremos un punto luminoso. Si dejamos pasar el flujo de los iones a

través de dos campos eléctricos mutuamente perpendiculares y cambiamos la tensión en los condensadores que crean estos campos es posible obligar al punto a desplazarse caóticamente por la pantalla.

Valiéndonos de semejante dispositivo podemos determinar el más importante parámetro de la partícula, a saber, la relación de su carga a la masa.

En un campo acelerador los iones cobran energía igual al trabajo de las fuerzas eléctricas, es decir

$$\frac{mv^2}{2} = eU$$

La tensión la conocemos, y la velocidad de las partículas se halla por los más variados procedimientos. Se puede, digamos, medir la desviación del punto luminoso en la pantalla. Está claro que la desviación será tanto mayor cuanto mayor sea el camino recorrido por la partícula y cuanto menor resulte su velocidad inicial. El problema se resuelve de una forma completamente rigurosa. Se parece al cálculo de la trayectoria de una piedra lanzada horizontalmente.

También existen procedimientos para la medición directa del tiempo consumido por el ion para recorrer todo el camino.

De este modo, conocemos la tensión y la velocidad del ion. ¿Qué se puede calcular, entonces, como resultado de semejante experimento? De la ecuación se infiere: se puede calcular la relación de la carga de la partícula respecto a su masa. Y es una pena que por mucho que se cambien las condiciones del experimento, por más variadas que sean las desviaciones y aceleraciones, a que se recurren, de las partículas, no se logra, de ningún modo, separar la magnitud de la carga de la masa. Solamente teniendo en cuenta los datos revelados por los químicos y el valor de la carga elemental obtenido a base de la electrólisis, se consigue sacar la firme conclusión: las cargas de todos los iones monovalentes son idénticas; las cargas de todos los iones divalentes representan el doble de las primeras, y de los iones trivalentes, el triple... Por esta razón, las diferencias en las relaciones de la carga respecto a la masa que se logran medir con una precisión extraordinariamente grande pueden considerarse como método de medición de la masa del ion.

Precisamente por esta causa el instrumento que desempeña un importantísimo papel en la química y la tecnología química y se basa en el principio del sencillo experimento cuya descripción hemos dado, este instrumento lleva el nombre de espectrógrafo de masas (el libro 4), aunque, en esencia, mide la relación de la carga respecto a la masa de los iones.

Rayo electrónico

Dejemos de seguir el intrincado curso de los acontecimientos históricos que condujo a los físicos a la firme convicción de que no solo existía la mínima porción de electricidad, sino de que esta porción tiene su portador material denominado electrón. Presentemos la descripción del experimento que se muestra en las clases escolares.

El instrumento destinado para esta tarea se denominaba, en su tiempo, tubo catódico. Actualmente lleva el nombre de tubo de haz electrónico Si usted, estimado lector, hace mucho ha terminado sus estudios y no ha visto este instrumento, no se aflija. Usted conoce muy bien el tubo de haz electrónico, ya que éste es el elemento principal de su aparato televisor en cuya pantalla el rayo electrónico dibuja cuadros que a veces le complacen y que siempre le permiten matar el tiempo libre.

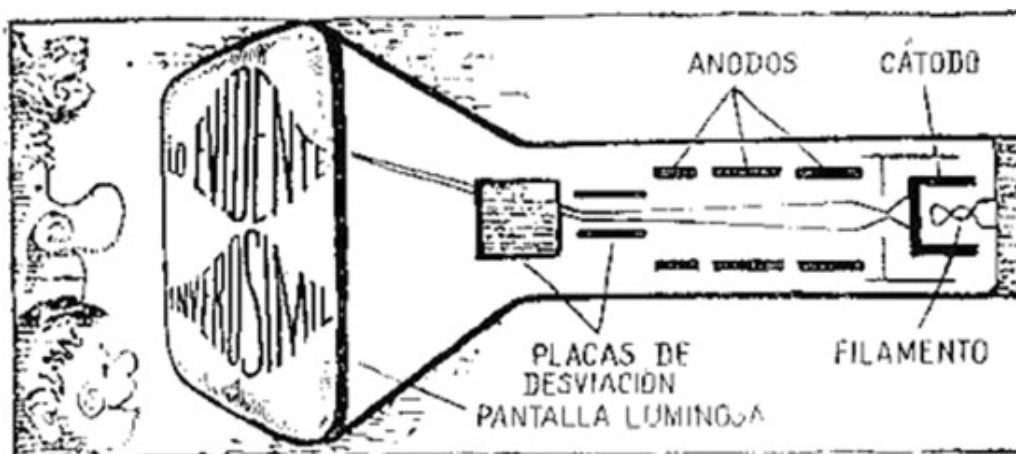


Figura 2.1

Pero volvamos al experimento escolar. El esquema del tubo se representa en la fig. 2.1. El tubo está perfectamente bombeado no hay en éste moléculas susceptibles

de desintegrarse. Sin embargo, al caldear con la corriente (dicha corriente se llama catódica) el filamento metálico y conectar seguidamente el cátodo y el ánodo a los polos respectivos de la fuente de tensión, usted descubrirá en la pantalla del tubo un punto luminoso y, por medio de un instrumento de medida, establecerá que del ánodo al cátodo comenzó a fluir la corriente eléctrica. Es lógico darle el nombre de corriente anódica.

Por cuanto la corriente fluye a través del vacío, no hay otro remedio sino hacer la conclusión de que el filamento incandescente emite partículas cargadas negativamente. El fenómeno lleva el nombre de emisión termoelectrónica o emisión termoiónica. Todo cuerpo incandescente posee esta capacidad.

Las partículas, y no ocultamos del lector que éstas son precisamente los electrones, se dirigen a los ánodos que tienen la forma de vasos con un orificio redondo en sus fondos. Los electrones salen en forma de un haz estrecho. Este haz puede investigarse mediante los mismos procedimientos que acabamos de describir para el haz de iones.

Una vez cerciorados, con ayuda de la pantalla luminosa, de que el filamento incandescente emite electrones, abordamos la determinación, utilizando placas de desviación, de la relación de la carga a la masa. El resultado es como sigue. La relación para el electrón es 1840 veces mayor que la misma relación para el ion más ligero, a saber, para el ion de hidrógeno. De aquí sacamos la conclusión de que el electrón es 1840 veces más ligero que el ion de hidrógeno. Este hecho significa que la masa del electrón es igual a 9×10^{-28} g.

Sin embargo, el lector tiene el derecho de advertir que nos apresuramos demasiado. Pues no se puede, basándose en la medición de la relación de la carga a la masa del electrón, concluir que su masa es menor que la del ion. ¿Y si resulta que las cargas de un ion positivo y del electrón son completamente disímiles?

La primera determinación de la relación de la carga a la masa del electrón se llevó a cabo todavía a finales del siglo XIX, por el brillante físico Joseph John Thomson (1856 - 1940). (Los amigos le llamaban J.J. A lo mejor, el empleo de estas siglas que se encuentran a menudo en las publicaciones de memorias no tanto se debe al amor que los ingleses experimentan a toda clase de abreviaturas como a la circunstancia de que en el siglo XIX había vivido y trabajado otro relevante físico

que tenía el mismo apellido. Se trata de William Thomson a quien por sus méritos científicos fue otorgado el título nobiliario de modo que comenzó a llamarse lord Kelvin). Por supuesto, el tubo catódico utilizado por Joseph Thomson fue mucho menos perfecto que el oscilógrafo moderno. El científico comprendía muy bien que su medición tan sólo hace probable el carácter discreto de la carga eléctrica y la existencia de una porción mínima de electricidad.

Por muy extraño que parezca, a pesar de que muchos físicos observaron el comportamiento de los rayos catódicos y anódicos había aún muchos partidarios de la hipótesis de que estos rayos tenían la naturaleza ondulatoria. Estos investigadores no veían la necesidad de reconocer que las corrientes fluyentes por el alambre metálico y por el líquido, así como aquellas que pasan a través de los gases y el vacío son parientes muy próximos. Insistieron en obtener pruebas más directas. Y, sin duda alguna, podemos comprender sus razones: para que una hipótesis se convierta en hecho no bastan argumentos indirectos.

De este modo, en primer término era necesario corroborar esta seguridad mediante mediciones directas de la carga de la partícula. Y en los primeros años del siglo XX el propio Thomson y sus discípulos abordaron esta tarea, y sus intentos, no se vieron frustrados, de ningún modo. Las mediciones más exactas fueron realizadas en 1909 por Roberto Millikan.

Experimento de Millikan

La idea sobre el carácter discreto de la electricidad se conceptúa como muy audaz, mientras que el cálculo de la carga elemental recurriendo al método con cuya exposición iniciamos el capítulo puede interpretarse también de otra forma. ¿Por qué no decir, por ejemplo, que los aniones existen realmente, mientras que la electricidad negativa es un líquido que es arrastrado por el ion positivo? Un ion capta una cantidad de este líquido; otro ion capta otra cantidad, y el experimento da cierto valor medio. Es una explicación muy sensata.

Como acabamos de decir, los experimentos de Thomson resultaron ser un argumento fuerte pero no decisivo a favor de la existencia del electrón. Por esta causa huelga demostrar cuán importante fue para la física el experimento en que la existencia de la carga elemental de electricidad quedó probada con tal grado de

evidencia que todas las dudas al respecto fueron eliminadas inmediatamente. Tal experimento fue puesto en 1909 por el físico norteamericano Roberto Millikan. No voy a referirme a otros trabajos de esto científico. Pero dicha investigación, por sí sola, hubiera bastado para que su nombre figurase en todos los manuales de física. La idea de este maravilloso experimento se basa en un hecho simple. De la misma manera como una varilla de vidrio que, frotada con piel, adquiere propiedades eléctricas, se comportan también otros cuerpos. Este fenómeno se denomina electrización por fricción. Mas, hablando con propiedad, ¿por qué se debe pensar que dicha propiedad es inherente tan sólo a los cuerpos sólidos? Puedo ser que también se van a electrizar las gotitas de aceite que inyectaremos a una cámara cualquiera, pues, al pasar por la boquilla del pulverizador el aceite se someterá a fricción. Resulta que todo sucede precisamente así. Para cerciorarse de ello es necesario montar un dispositivo muy sencillo, de principio: dirigir el chorro de las gotitas de aceite al espacio entre las armaduras de un condensador dispuestas horizontalmente y adaptar un microscopio que permita observar el movimiento de las gotitas. Mientras el campo eléctrico no se haya aplicado, las gotitas, como es natural, caerán por acción de la fuerza de la gravedad. Las gotitas son ligeras y por esta razón la fuerza de la gravedad, casi de inmediato; se equilibrará por la fuerza de la resistencia del aire, de modo que dichas gotitas caerán uniformemente. Sin embargo, apenas a las placas se aplica la tensión el cuadro cambia. El movimiento de la gota o se acelera, o bien, se hace más lento en dependencia de la dirección del campo eléctrico. Millikan optó una dirección tal del campo que hizo a la gota moverse más lentamente. Aumentando poco a poco el campo él logró, por decirlo así, dejar la gota suspendida en el aire. Durante horas enteras el investigador observó una sola gota. Valiéndose del campo pudo, según su deseo, dirigir su movimiento y dejarla inmóvil.

¿Qué se puede calcular, entonces, por medio de semejante experimento? En primer término examinemos los datos que se obtendrán por observación en ausencia del campo. La igualdad de las fuerzas de la gravedad y de la resistencia del campo puede escribirse en la siguiente forma:

$$mg = av$$

Lo densidad del aceite se halla con facilidad por experimentos independientes y el diámetro de la gota se mide con el microscopio. Siendo así, no cuesta trabajo calcular la masa de la gota. La caída de la gota es lenta, en consecuencia, haciendo marcas en el vidrio del microscopio y empleando un cronómetro determinemos con bastante precisión la velocidad de caída v de la gota. Entonces, a partir de la ecuación escrita anteriormente, se halla el coeficiente de resistencia a .

Ahora conectemos el campo. Lo más conveniente es conseguir tal estado de cosas, que la gota comience a ascender uniformemente. A las dos fuerzas que actuaban ya se añadió la tercera, la cual por parte del campo eléctrico cuya intensidad E nos es conocida (la relación de la tensión a la distancia entre las placas del condensador) El movimiento uniforme hacia arriba significa que las tres fuerzas se han equilibrado. La condición de este equilibrio tendrá la siguiente forma:

$$qE - mg - av' = 0$$

El nuevo valor de la velocidad v' se mide empleando el mismo microscopio. De este modo resulta que se conocen todas las magnitudes que forman parte de la ecuación, salvo la carga de la gota. Calculemos el valor de esta carga y lo anotamos en el cuaderno que, indispensablemente, debe llevar cualquier escrupuloso experimentador.

Ahora sí que nos hemos aproximado a la principal invención. La corriente en el electrolito. Discurría Millikan, se transporta por iones de distinto signo. Pero los iones pueden formarse también en el gas. El aire se ioniza empleando los más diversos procedimientos. Por ejemplo, todo el dispositivo puede situarse junto a un tubo de rayos X. Estos rayos ionizan el aire; hecho perfectamente conocido también en aquella época. Pero si la gota está cargada, entonces, atraerá los iones de signo contrario. Apenas el ion se adhiera a la gota la carga de esta última cambiará. Y en cuanto la carga se modifique también la gota varía su velocidad que puede hallarse inmediatamente por medio de una nueva medición.

Las observaciones demostraron que esta idea resultó ser cierta. Después de haber conectado el tubo de rayos X, diferentes gotas cada vez comenzaron a cambiar a salto su velocidad. Sin quitar los ojos de una misma gota, el observador medía la

diferencia de velocidad antes y después de conectar el tubo de rayos X. A base de la fórmula que hemos insertado se calculaba, de una vez el valor de q .

¿No ha comprendido todavía cuál es la finalidad de este proceder? Recapacítelo mejor. Si la carga eléctrica elemental existe, entonces en el caso de que a la gota se ha unido un ion monovalente los valores medidos deben ser iguales a esta carga, o bien, en el caso de adherirse a la gota varios iones, dichos valores deben ser múltiplos de la magnitud de la carga elemental.

A medida que Millikan llevaba a cabo sus experimentos para las gotas de aceite, de mercurio y de glicerina, cambiando, al mismo tiempo los signos de la carga de las gotas, su cuaderno se llenó de números que correspondían a los valores de q , y todos estos valores resultaron ser múltiplos de una misma magnitud, precisamente de aquella que había sido hallada debido a las investigaciones de la electrólisis.

Después de que Millikan publicó sus resultados, hasta los escépticos dejaron de poner en tela de juicio el que la carga eléctrica se encontrase en la naturaleza en forma de porciones discretas. Aunque, de hecho, hablando con rigor, tampoco los experimentos de Millikan demostraban directamente la existencia del electrón como partícula.

Sin embargo, las hipótesis adelantan a los hechos. Ya a principios del siglo XIX había quien estuviera seguro de la naturaleza granular de la electricidad. G. J. Stoney, en 1891, por primera vez calculó la carga del ion; también fue él quien prepuso el término «electrón», mas no para la partícula, sino para la carga de un ion negativo monovalente. Los experimentos de Thomson hicieron creer en la existencia del electrón como partícula a la absoluta mayoría de los físicos. Paul Drude fue el primero quien, sin ambages, definió el electrón como partícula que porta una carga elemental de electricidad negativa.

De este modo resulta que el electrón obtuvo reconocimiento antes de haber sido «visto».

En cuanto a la prueba directa de existencia del electrón de ésta sirvieron los experimentos finos realizados más tarde. Se hace incidir sobre la pantalla un haz débil de partículas y cada una de estas partículas puede calcularse. Cada electrón produce un centelleo en la pantalla luminosa. Desde luego, hace mucho ya que con este fin no se utilizan pantallas luminosas, sino contadores especiales que por el

nombre de su inventor se denominan contadores de Geiger. En pocas palabras, la idea de este contador consiste en que un electrón, a guisa de gatillo de un revólver, da inicio a un fuerte impulso de corriente que es fácil de registrar. De esta manera el físico tiene la posibilidad de establecer el número de electrones que caen en una determinada trampa por un segundo. Si en calidad de dicha trampa se toma un matraz metálico a cuyo interior van a parar electrones, el matraz se cargará paulatinamente de una cantidad de electricidad suficiente para que se la pueda medir con precisión. Para hallar la carga del electrón lo único que se requiere es dividir la cantidad de electricidad por el número de electrones captados.

Solamente después de estos experimentos podemos decir: la existencia del electrón dejó de ser una hipótesis. Es un hecho.

Pasamos sin detenernos, a velocidad de un coche de carreras, al lado de los descubrimientos que pusieron los cimientos de la física moderna. ¡Más, qué se puede hacer! ¡Ésta es su suerte! Los nuevos acontecimientos hacen replegarse al segundo plano a los viejos y hasta los eventos cruciales que tuvieron lugar durante la edificación del templo de la ciencia pasan a ser incumbencia de los historiadores. Ahora, tal vez, podemos contestar a la pregunta de qué es la electricidad. El fluido eléctrico es el flujo de partículas eléctricas (los sinónimos son: cargadas, poseedoras de carga eléctrica). El cuerpo es eléctricamente cargado si el número de partículas de un signo supera el de partículas de otro signo.

- Vaya una explicación - dice indignado el lector. ¿Qué es una partícula eléctrica?

- ¿Acaso no le parece claro? Las partículas se denominan eléctricas si interaccionan de acuerdo a la ley de Coulomb.

- ¿Y nada más? - pregunta con perplejidad el lector.

- Nada más - responde el físico. Nada más en lo que atañe a la respuesta a su pregunta. Pero, más adelante, le esperan respuestas a muchas otras interesantes cuestiones. No olvide que todavía no hemos hecho mención de las ocasiones en que nos aguardan los encuentros con la partícula elemental de la electricidad positiva. Además, en perspectiva, nos enteraremos de que las partículas eléctricas se caracterizan no sólo por la carga y la masa, sino también por otras propiedades.

Sin embargo, en primer término, nuestra conversación se referirá a la estructura del átomo.

Modelo del átomo

¿Cómo está estructurado el átomo de partículas eléctricas? La respuesta se ha obtenido con la ayuda de los rayos emitidos por el elemento radio. Hablaremos sobre esta maravillosa sustancia, así como sobre la gran familia de los elementos radiactivos naturales y artificiales en el cuarto libro. Mientras tanto, nos basta con saber que el radio emite ininterrumpidamente una radiación electromagnética dura (rayos gamma), el flujo de electrones (que en su tiempo se denominaba rayos beta) y rayos alfa que no son sino iones del átomo de helio doblemente cargados.

El relevante físico inglés (nació en Nueva Zelanda) Ernesto Rutherford (1871 - 1937) propuso en 1911 el llamado modelo planetario del átomo, idea a la cual llegó basándose en las escrupulosas investigaciones de la dispersión de las partículas alfa mediante diferentes sustancias. Rutherford realizaba experimentos con el pan de oro cuyo espesor constituía tan sólo una décima de micrón. Resultó que entre 10.000 partículas alfa solamente una se desviaba en un ángulo superior a 10° y que raras veces también se encontraron partículas alfa desviándose en un ángulo obtuso alcanzando éste incluso 180° .

En estos experimentos sorprendentes por su sencillez se fijaba el paso de cada partícula por separado. Se sobreentiende que la técnica moderna permite efectuar las mediciones de una manera totalmente automatizada.

Así, pues, inmediatamente se pone en claro que los átomos constan principalmente del... vacío. Los raros choques frontales deben entenderse de la siguiente forma. Dentro del átomo existe un núcleo cargado positivamente. Cerca del núcleo se disponen los electrones. Estos son muy ligeros y, por lo tanto, no forman un obstáculo serio para la partícula alfa. Los electrones licúan la partícula alfa, pero el choque con cada electrón solitario no puede desviar la partícula de su camino.

Rutherford supuso que las fuerzas de interacción entre el núcleo del átomo y la partícula alfa de cargas homónimas son fuerzas coulombianas. Al suponer, luego, que la masa del átomo está concentrada en su núcleo calculó la probabilidad de desviación de las partículas a un ángulo dado y obtuvo una coincidencia brillante entre la teoría y el experimento.

Es así cómo los físicos comprueban los modelos producto de su imaginación.

- ¿El modelo predice los resultados del experimento?
- Sí.
- Entonces, ¿resulta que refleja la realidad?
- ¿Para qué hablar tan categóricamente? El modelo explica una serie de fenómenos, por consiguiente, es válido. Y precisarlo es asunto del futuro...

Los resultados de los experimentos de Rutherford no dejaban dudas en cuanto a la certeza de la siguiente afirmación: los electrones, por acción de las fuerzas de Coulomb se mueven alrededor del núcleo.

De la teoría derivaban también algunas evaluaciones cuantitativas que más tarde se vieron confirmadas. Las dimensiones de los núcleos atómicos más pequeños resultaban iguales a 10^{-13} cm, aproximadamente, mientras que las del átomo eran del orden de 10^{-8} cm.

La confrontación de los resultados del experimento con los cálculos dio la posibilidad de evaluar también las cargas de los núcleos en colisión. Dichas evaluaciones desempeñaron un gran papel - incluso se puede decir que un papel prioritario - en la interpretación de la ley periódica de la estructura de los elementos.

Bueno, el modelo del átomo está construido. Pero, inmediatamente, surge la siguiente pregunta. ¿Por qué los electrones (partículas cargadas negativamente) no caen al núcleo (cargado positivamente)? ¿Por qué el átomo es estable?

¿Pero qué hay aquí de incomprensible? - dirá el lector. Es que los planetas tampoco caen al Sol. Una fuerza de procedencia eléctrica, al igual que la de la gravedad, es fuerza centrípeta, asegurando el movimiento circular de los electrones alrededor del núcleo.

Sin embargo, la cuestión reside precisamente en que la analogía entre el sistema planetario y el átomo reviste tan sólo carácter somero. Como se pondrá de manifiesto más tarde, desde el punto de vista de las leyes generales del campo electromagnético el átomo está obligado a emitir ondas electromagnéticas. Desde luego, no es indispensable conocer la teoría del electromagnetismo. La materia, es decir, los átomos, es capaz de irradiar luz y calor.

Siendo así, el átomo pierde energía y, por consiguiente, el electrón debe caer al núcleo.

¿Cuál es, entonces, la salida de esta situación? Es muy «simple»: hay que avenirse a los hechos y elevarlos al rango de la ley de la naturaleza. Este paso lo emprendió, precisamente, en 1913 el gran físico de nuestro siglo Niels Bohr (1885 - 1962).

Cuantificación de la energía

Al igual que todos los primeros pasos este también fue relativamente tímido. Procedamos a la exposición de una nueva ley de la naturaleza que no sólo salvó el modelo del átomo propuesto por Rutherford, sino también nos obligó a llegar a la conclusión de que la mecánica de los macrocuerpos es inaplicable a las partículas de masa pequeña.

El carácter de la naturaleza es tal que algunas magnitudes como, por ejemplo, el momento de impulso y la energía no pueden tener una serie continua de valores para cualquier sistema de partículas que están en interacción. En cambio, el átomo de que hablamos ahora o el núcleo atómico de cuya estructura hablaremos más tarde tienen su secuencia de niveles de energía inherente tan sólo al sistema dado. Se tiene el nivel más bajo (el nivel cero). La energía del sistema no puede ser menor de este valor. En el caso del átomo esto quiere decir que existe un estado tal en que el electrón se encuentra a cierta distancia mínima del núcleo.

La variación de la energía del átomo puede efectuarse solamente a salto. Si el salto resulta hacia «arriba» este hecho significa que el átomo absorbió energía. Si el salto ocurre hacia «abajo», entonces, el átomo irradió energía.

La ley que acabamos de formular lleva el nombre de ley de cuantificación de la energía. También se puede decir que la energía reviste carácter cuántico.

Cabe señalar que la ley sobre la cuantificación tiene un carácter absolutamente universal. Es aplicable no sólo al átomo, sino también a cualquier objeto que consta de miles de millones de átomos. Sin embargo, cuando se trata de grandes cuerpos, con frecuencia, meramente podemos «pasar por alto» la cuantificación de la energía. La cuestión radica en que, hablando en rasgos generales, en un objeto que consta de un millón de billones de átomos el número de niveles de energía aumenta un millón de billones de veces. Los niveles de energía se situarán tan cerca uno al otro que, prácticamente, se confundirán. Esta es la razón por la cual pasaremos por alto el carácter discreto de los valores posibles de la energía. De este modo, cuando

se trata de macrocuerpos, aquella mecánica que hemos expuesto en el libro 1 prácticamente no cambia.

En el libro 2 hemos aclarado que la transmisión de la energía de un cuerpo a otro puede efectuarse en forma de trabajo y en forma de calor. Ahora estamos en condiciones de explicar la diferencia entre estas dos formas de transmisión de la energía. Durante la acción mecánica (digamos, durante la compresión) los niveles de energía del sistema se desplazan. Dicho desplazamiento es sumamente insignificante y se pone de relieve sólo por experimentos finos y sólo en el caso de que las presiones sean lo suficientemente grandes. En lo que concierne a la acción térmica, ésta consiste en la traslación del sistema de un nivel de energía más bajo a otro nivel más alto (calentamiento), o bien, de un nivel más alto al nivel más bajo (enfriamiento).

La cuantificación de la energía, igualmente que de otras magnitudes mecánicas, es una ley general de la naturaleza de la cual derivan, con rigor, los más variados corolarios que encuentran su confirmación en la experiencia.

Es posible que usted quisiera preguntar: ¿por qué la energía se cuantifica? A esta pregunta no hay respuesta.

¡La naturaleza está estructurada así! Cualquier explicación es la reducción de un factor particular a otro más general. Por el momento no conocemos ni un solo enunciado tan general que de éste, como corolario, derive la cuantificación de la energía. Por supuesto, no queda excluido que en el futuro se descubrirán las leyes tan globales que los principios de la mecánica cuántica resulten sus corolarios. Sea como fuere, hoy en día la ley de cuantificación representa una de las pocas magnas leyes de la naturaleza que no necesitan una motivación lógica. La energía se cuantifica porque... se cuantifica.

En esta forma general dicha ley fue establecida en los años 1925 - 1927 gracias a los trabajos del físico francés Luis de Broglie, el físico austríaco Erwin Schrödinger y el físico alemán Werner Heisenberg. La ciencia cuya base la constituye el principio de cuantificación (un momento, se me olvidó decir que «quantum» de que procede la palabra «cuanto» puedo traducirlo como «porción») recibió el nombre de mecánica cuántica u ondulatoria. ¿Y por qué ondulatoria? se lo daremos a conocer más tarde.

La ley periódica de Mendeleiev

En 1868 el gran químico ruso Dmitri Ivanovich Mendeleiev (1834 - 1907) publicó la ley periódica de secuencia de los elementos químicos que él había descubierto. No vamos a insertar aquí la tabla de Mendeleiev la cual el lector puede encontrar en un libro de texto escolar de química. Hagamos recordar que, al disponer los elementos conocidos en una serie según sus pesos atómicos, Mendeleiev prestó atención a que las propiedades químicas y algunas particularidades físicas de los elementos varían periódicamente en dependencia del peso atómico.

En la tabla compuesta por Mendeleiev cada uno de los elementos pertenece a uno de los nueve grupos y a uno de los siete períodos. Los elementos pertenecientes a un mismo grupo, Mendeleiev los dispuso en forma de columnas de modo que aquellos cuyos símbolos se ubicaban uno debajo de otro poseyeran propiedades químicas similares. Resultó que este propósito se podía lograr solamente en el caso de suponer que existían elementos todavía no descubiertos. Para estos elementos hipotéticos Mendeleiev dejó en su tabla «casillas vacías». La clarividencia del gran sabio se manifestó también en el hecho de que emplazó el átomo de níquel en su «debido» lugar tras el cobalto, a pesar de que el peso atómico del cobalto era algo mayor.

Algunas «casillas vacías» se llenaron ya durante la vida de Mendeleiev, con lo que se granjeó la fama mundial pues a todos quedó claro que la composición de esta tabla no es un simple acto formal, sino el descubrimiento de una gran ley de la naturaleza.

El sentido del número de orden que la tabla atribuye al elemento químico devino patente tan sólo después de que los físicos ya se libraron de dudas en lo que se refería a la validez del modelo planetario del átomo propuesto por Rutherford y de la ley de cuantificación de la energía. ¿Cuál es, entonces, este sentido? La respuesta resulta extraordinariamente sencilla: el número de orden es igual a número de electrones que giran alrededor del núcleo. También se puede decirlo de otra forma: el número de orden del elemento es la carga positiva de su núcleo expresada en unidades de la carga del electrón.

La ley periódica de Mendeleiev, el principio de cuantificación de la energía y el estudio de los espectros característicos ópticos y de rayos X de los átomos (hablaremos sobre estos últimos más tarde) permitieron comprender el porqué de la similitud del comportamiento químico de los átomos que se encuentran en una columna de la tabla de Mendeleiev

La energía del átomo es la energía de interacción de los electrones con el núcleo (la aportación de la energía de interacción de los electrones entre sí, en el razonamiento que sigue, resulta insustancial). Por cuanto la energía se cuantifica sería lógico suponer que los electrones de cada átomo pueden disponerse en una serie según sus energías. El primer electrón está ligado al núcleo de la forma más fuerte, la ligazón del segundo electrón con el núcleo es más débil; del tercero, más débil todavía, y así sucesivamente, de modo que los electrones del átomo se disponen según escalones de energía. La lógica no nos falla, pero la experiencia lleva a precisar este cuadro. En primer lugar, resulta que cada escalón de energía lo pueden ocupar no uno sino dos electrones. Es cierto que estos electrones se distinguen uno del otro por una propiedad llamada «espín». Esta propiedad es vectorial. Los aficionados a la representación patente pueden trazar en su imaginación que en el escalón lleno se encuentran dos «puntitos» provistos de flechas con la particularidad de que una flecha mira hacia «abajo» y la otra hacia «arriba».

El propio término «espín» apareció de la siguiente manera. Procede de la palabra inglesa «spin» que significa «girar rápidamente». Para que uno se forme la idea de cuál es la diferencia entre dos electrones sentados en el mismo peldaño se inducía a pensar que un electrón gira en sentido horario, y otro, en sentido anti horario, haciéndolo alrededor de su propio eje. Dicho modelo fue sugerido a raíz del parecido superficial entre el átomo y el sistema planetario. Como el electrón es una especie de planeta, ¿por qué no permitirle que gire alrededor de su eje? Una vez más debo decepcionar ni lector: representar patentemente el espín del electrón es una tarea imposible. Pero en el siguiente capítulo hablaremos sobre cómo medirlo.

Sin embargo, ésta no es la única conclusión importante (a la cual nos llevó el estudio escrupuloso de los espectros de los átomos). La segunda conclusión

consistía en que los escalones de energía estaban separados unos de otros por unas distancias desiguales y pueden dividirse en grupos.

Tras el primer escalón llamado nivel K sigue una ruptura de energía, después se tiene un grupo de 8 electrones designado con la letra L; luego sigue el grupo de 18 electrones designado con la letra M... Nos abstenemos de exponer aquí la disposición de los niveles y el orden de su llenado para todos los átomos. El cuadro resulta lejos de ser tan simple y su descripción hubiera requerido mucho espacio. En nuestro pequeño libro los detalles no importan y he hecho mención de los escalones con el único fin de explicar con mayor claridad en qué, precisamente, consiste el parecido entre los átomos que se encuentran uno debajo del otro en la tabla de Mendeleiev. Resulta que tienen el mismo número de electrones en el grupo superior de los escalones.

De este modo se esclarece el concepto químico de valencia del átomo. Por ejemplo, el litio, sodio, potasio, rubidio, cesio y francio tienen un electrón en el grupo superior de peldaños, el berilio, magnesio, calcio, etc. poseen dos electrones en este grupo. Los electrones de valencia (los externos) están ligados al átomo con menor fuerza que los demás. Por esta razón, durante la ionización de los átomos que se encuentran en la primera columna se forman más fácilmente las partículas con una sola carga. Los iones de berilio, magnesio, etc. portan sobre sí dos cargas, y así sucesivamente.

Estructura eléctrica de las moléculas

Los químicos denominan molécula al representante mínimo de una sustancia. Los físicos, las más de las veces, hacen uso de este vocablo tan sólo en el caso de que este representante mínimo exista realmente como un cuerpo pequeño y aislado.

¿Es que existe la molécula de la sal común? Claro que sí, contesta el químico y escribirá la fórmula: NaCl. La sal común es el cloruro de sodio. Su molécula consta de un átomo de sodio y un átomo de cloro. Sin embargo, esta respuesta es correcta solamente desde el punto de vista formal. En realidad, ni en un pequeño cristal de sal común, ni en su solución acuosa, ni en el vapor de cloruro de sodio, en ninguna parte, podemos descubrir una pareja de átomos que se comporte como un todo único. Como ya mencionamos en el segundo libro, en un cristal cada átomo de

sodio está rodeado de seis cloros vecinos. Todos estos vecinos gozan de igualdad de derechos y no se puede decir, de ningún modo, cuál de éstos «pertenece» al átomo dado de sodio.

Hagamos disolverse la sal común en agua. Resultará que la solución es un magnífico conductor de la corriente. Por medio de rigurosos experimentos, de los cuales ya hemos hablado, es posible demostrar que la corriente eléctrica representa el flujo de átomos de cloro cargados negativamente que se mueven en una dirección y el flujo de átomos de sodio cargados positivamente que se mueven en dirección opuesta. De este modo vemos que durante la disolución los átomos de cloro y de sodio tampoco forman una pareja fuertemente ligada.

Una vez establecido el modelo de átomo se pone de manifiesto que el anión del cloro no es sino un átomo de cloro con un electrón «sobrante», mientras que al catión, en cambio, le «falta» un electrón.

¿Se puede, a base de estos hechos, sacar la conclusión de que también el sólido está construido de iones y no de átomos? Sí. Esta circunstancia se demuestra por numerosos experimentos en cuya descripción no nos detenemos.

¿Y qué se puede decir en cuanto al vapor del cloruro de sodio? En el vapor tampoco encontramos moléculas. El vapor del cloruro de sodio consta de iones o de diferentes grupos - sumamente inestables - de iones. Se puede hablar de moléculas de los compuestos iónicos solamente en el sentido químico de esta palabra.

Los compuestos iónicos, obligatoriamente, se disuelven en agua. Tales soluciones, cuyo ejemplo clásico lo representan las sales simples de los metales a semejanza de cloruro de sodio, acusan buena conductividad denominándose debido a ello electrólitos fuertes.

Ahora aduzcamos algunos ejemplos de sustancias construidas de verdaderas moléculas, o sea, de moléculas en el sentido físico de esta palabra. Se trata de oxígeno, nitrógeno, dióxido de carbono, hidrocarburos, carbohidratos, esteroides, vitaminas... La lista puede continuarse prolongadamente.

Cualesquiera clasificaciones siempre son hasta cierto grado convencionales. Por esta razón debo advertir al lector que, a veces, arrostramos también los casos cuando en un estado de agregación la sustancia consta de moléculas físicas, y en otros, no. A estas sustancias pertenece la sustancia tan importante como es el agua. Las

moléculas del vapor de agua son, sin duda alguna, cuerpecitos aislados. En cambio, en los cristales de hielo ya no es fácil «delimitar el contorno» de una molécula y decir que este preciso átomo de hidrógeno está ligado solamente con aquel átomo de oxígeno, y no con otro.

Sea como sea, la caso de cristales moleculares resulta ser muy amplia.

En el segundo libro ya hablamos de cómo están estructurados los cristales moleculares. Recordemos que en el cristal de dióxido de carbono cuya fórmula es CO_2 el átomo de carbono tiene dos vecinos de oxígeno muy cercanos. También en todos los demás casos, al estudiar la estructura de un cristal molecular, vemos inmediatamente que existe la posibilidad de dividir el cristal en grupos de átomos dispuestos estrechamente.

Como están dispuestos estrechamente, esto significa que están ligados por grandes fuerzas. Así os, precisamente. Hablando en rasgos generales, las fuerzas que ligan los átomos pertenecientes a una misma molécula son cien veces - y, en ocasiones, incluso mil veces - mayores que las fuerzas que actúan entre los átomos de moléculas vecinas.

¿En qué consiste, entonces, la ligazón intramolecular? Queda bastante claro que en este caso no son suficientes las representaciones sobre la atracción de los iones eléctricamente cargados negativos y positivos. Es que existen moléculas de oxígeno, nitrógeno o hidrógeno construidas de átomos equivalentes. Es imposible suponer que un átomo pierde y otro adquiere el electrón. ¿A santo de qué el electrón debe preferir encontrarse junto a uno de los dos átomos equivalentes?

La explicación de la esencia de la ligazón intramolecular llegó tan sólo junto con la mecánica cuántica. Acabamos de decir al lector que la energía de cualquier sistema se cuantifica y le hemos dado a conocer que en un nivel de energía pueden encontrarse dos electrones con «espines» de dirección opuesta. Prosigamos: de las hipótesis principales de la mecánica cuántica deriva un interesante corolario. Resulta (y ya no es una hipótesis, sino una deducción matemática rigurosa que no insertamos debido a su complejidad) que el valor más bajo de energía que puede tomar el electrón se determina por las dimensiones de la zona en cuyo seno se desplaza. Cuanto mayores son estas dimensiones tanto más baja es la energía de este «nivel cero».

Ahora figúrese que dos átomos de hidrógeno van acercándose uno al otro. Si éstos se unen en un sistema, entonces, el «apartamento» para cada electrón llegará a ser dos veces mayor, aproximadamente. En un mismo apartamento pueden convivir pacíficamente dos electrones con espines de dirección opuesta. Por consiguiente, tal convivencia resulta provechosa. Para ambos electrones se acrecentó la zona de existencia. Resulta que la energía total del sistema después de que los dos átomos se unieron en un todo único disminuyó. Y el hecho de que cualquier sistema tienda a pasar al estado con mínima energía – con tal de que semejante posibilidad exista – este hecho lo conocemos perfectamente. Por esta misma causa, una bola dejada a su «libre albedrío» rueda cuesta abajo.

De este modo, la formación de un enlace químico significa la colectivización de los electrones. Existe cierta cantidad de electrones (se llaman internos) que giran alrededor de los núcleos de los átomos, sin embargo, algunos electrones (llevan el nombre de externos) abarcan en su recorrido por lo menos un par de los más próximos átomos y, a veces, hasta migran por todos los átomos de la molécula.

Una sustancia construida de moléculas se reconoce por sus propiedades eléctricas. La solución de semejante sustancia no conduce corriente. Las moléculas no se disgregan en partes y la molécula entera es eléctricamente neutra. En los líquidos y vapores las moléculas conservan su estructura: todo el grupo de átomos se mueve como un todo único, se desplaza en movimiento progresivo y gira. Los átomos pertenecientes a una molécula sólo pueden oscilar alrededor de sus posiciones de equilibrio.

La molécula neutra no porta en sí una carga eléctrica. Pero no se apresure a hacer la conclusión de que dicha molécula no crea un campo eléctrico. Si la molécula no es simétrica, los centros de gravedad de sus cargas positiva y negativa seguramente no van a coincidir. Intuitivamente, queda claro que la coincidencia de los centros de gravedad de las cargas de ambos signos tendrá lugar en tales moléculas como los de oxígeno o de nitrógeno que constan de dos átomos iguales. De la misma manera, no es difícil creer que en una molécula como, por ejemplo, la de monóxido de carbono CO, estos centros pueden encontrarse desplazados uno respecto al otro. Si esta dislocación existe, se dice acerca de la molécula que ésta posee un momento dipolar.

El término tiene la siguiente procedencia: una molécula «dipolar» se comporta como un sistema de dos cargas puntuales (un punto es el centro de gravedad de las cargas negativas, y otro punto, el centro de gravedad de las cargas positivas). Un dipolo se caracteriza por la carga y por el «brazo» del dipolo, es decir, por la distancia entre los centros.

No me exijan la demostración de que una molécula no simétrica posee el momento dipolar eléctrico. Podemos ahorrarnos el tiempo para los razonamientos teóricos porque la existencia de un momento dipolar constante (o, como también se dice, duro) se demuestra sin dificultad experimentalmente.

Los dieléctricos

Entre los conceptos de «dieléctrico», «no conductor de la corriente» y «aislador» podemos interponer signos de igualdad.

Pertenecen a los dieléctricos los gases moleculares, los líquidos moleculares y las soluciones de los sólidos contruidos a partir de moléculas. Son dieléctricos sólidos los vidrios tanto los orgánicos, como los inorgánicos (a base de silicatos, boratos, etc.), sustancias polímeras contruidas de macromoléculas, los plásticos, los cristales moleculares, así como los cristales iónicos.

En el Capítulo 1 hemos recordado al lector que la capacidad del condensador aumenta sí en el espacio entre las placas se introduce cualquier dieléctrico. Figúrense que el condensador está conectado a una fuente de tensión constante. La capacidad ha aumentado, pero la tensión sigue siendo la misma. Ello quiere decir que a las armaduras del condensador ha llegado una carga suplementaria. Al parecer, en este caso debería incrementar la intensidad del campo. Sin embargo, la intensidad del campo no ha cambiado, por la simple razón de que es igual al cociente de división de la tensión por la distancia entre las placas. ¿Cómo se puede salir de esta contradicción? Por la única vía, o sea, admitiendo que en el aislador se ha engendrado un campo eléctrico de dirección opuesta; este fenómeno lleva el nombre de polarización del dieléctrico.

¿Qué tienen de particular, entonces, estas cargas, las que surgen en el seno del dieléctrico? ¿Cómo se puede comprender el fracaso de los intentos de «derivar» a Tierra la carga del dieléctrico? Incluso de no conocer nada sobre la estructura

eléctrica de la materia podemos decir que estas cargas resultan «atada» y no libres como en el metal.

Y si disponemos de datos acerca de la estructura de las moléculas podremos explicar, de una forma exhaustiva, la esencia del fenómeno de polarización y esclarecer el mecanismo de formación del «anticampo» teniendo en cuenta que este último, en igualdad de las demás condiciones, es tanto mayor cuanto mayor es ϵ .

Ante todo, es preciso contestar a la pregunta de qué puede hacer el campo eléctrico con el átomo y la molécula. Por impacto del campo eléctrico los electrones del átomo neutro y del ion pueden desplazarse en el sentido opuesto al campo. El átomo o el ion se convierten en un dipolo y se engendra un campo de dirección contraria. De este modo, la polarización de la sustancia está condicionada por la polarización de los átomos, iones o moléculas de los cuales está construida.

El mecanismo de polarización al cual nos hemos referido, se denomina proceso de creación de dipolos blandos.

Si no hay campo tampoco se dan dipolos. Cuanto mayor es el campo, tanto mayor es el desplazamiento del centro de gravedad de los electrones, tanto mayor es el momento dipolar «inducido» y tanto mayor es la polarización.

La formación de los dipolos blandos no puede depender de la temperatura. La experiencia demuestra que existen dieléctricos sobre los cuales la temperatura no ejerce influencia. Por lo tanto, para éstos es justo el mecanismo descrito.

Bueno, ¿y qué se puede idear para los casos en que se tiene una dependencia explícita de la constante dieléctrica (de la permitividad) respecto a la temperatura? Un estudio atento de la relación existente entre la estructura de la molécula y el comportamiento de la sustancia en el campo eléctrico, así como el carácter de la dependencia térmica de la polarización siempre disminuye con el crecimiento de la temperatura nos lleva al siguiente pensamiento. Si las moléculas, también en ausencia del campo, poseen un momento dipolar (dipolos «duros») y pueden variar su orientación, resulta que este hecho explicará la dependencia térmica de la constante dieléctrica.

Efectivamente, sin campo, las moléculas están dispuestas «como les da la gana». Los momentos dipolares se suman de modo geométrico. Por esta causa, para un volumen que contiene muchas moléculas el momento resultante será igual a cero.

El campo eléctrico hace como si «pasara el peine» por las moléculas obligándolos a mirar preferentemente a un mismo lado. Dos fuerzas entran en lid: el movimiento térmico que impone un desorden en la disposición de las moléculas y la acción ordenadora del campo. Se sobreentiende que cuanto más alta es la temperatura, tanto más difícil es para el campo «superar» el movimiento térmico de las moléculas. De aquí, precisamente, se infiere que la constante dieléctrica de tales sustancias debe disminuir con el descenso de la temperatura.

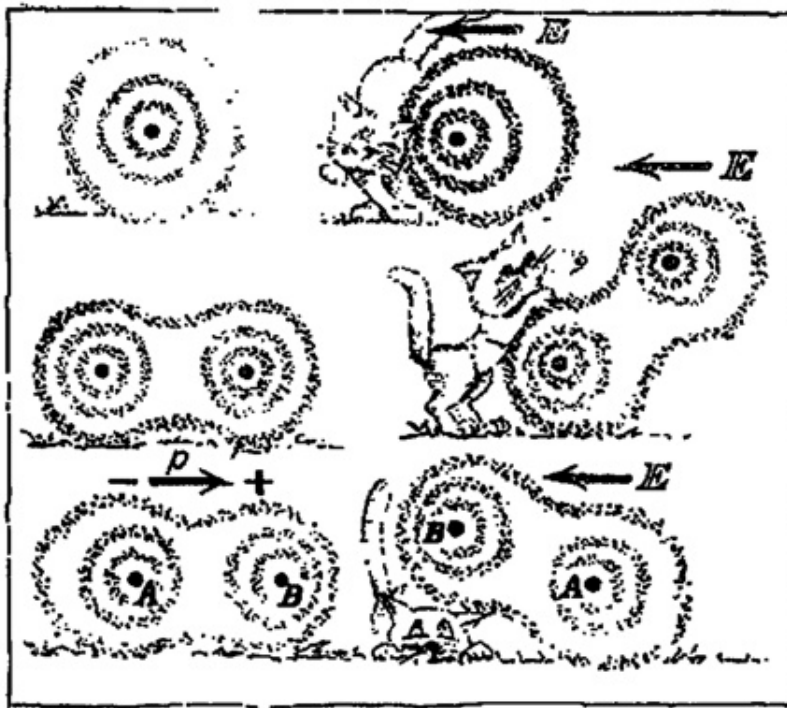


Figura 2.2

Para retener mejor en la memoria lo expuesto insertamos la fig. 2.2. El dibujo en la parte superior de la figura muestra que la polarización del átomo se reduce a la dislocación y deformación de las envolturas electrónicas. Cuanto mayor es la distancia que separa el electrón del átomo tanto más repercute en él la acción del campo. Las capas representadas en estos dibujos esquemáticos por medio de puntos simbolizan los lugares en que se hallan los electrones. Hay que recordar que el cuadro reviste un carácter sumamente convencional ya que los diferentes electrones tienen en las moléculas unas zonas de existencia distintas por su forma.

En la parte media de la figura se ilustra el comportamiento de una molécula diatómica simétrica. En ausencia del campo ésta no posee momento. El campo induce el momento eléctrico. Dicho momento puede ser diferente por su valor, en función del ángulo bajo el cual la molécula se dispone respecto al campo. El momento se forma debido a la deformación de las envolturas electrónicas.

Finalmente, en el esquema inferior se representa el comportamiento de una molécula que posee el momento dipolar también en ausencia del campo. En nuestro esquema la molécula sólo ha girado. Sin embargo, en el caso general, las sustancias cuyas moléculas poseen el momento en ausencia del campo acusan ambos mecanismos de polarización: a la par de girar las moléculas pueden también tener lugar las dislocaciones de los electrones. No es difícil separar estos dos efectos realizando mediciones a temperaturas muy bajas cuando, prácticamente, falta la influencia del movimiento térmico. Si este modelo es certero, no debemos observar la dependencia respecto a la temperatura de la constante dieléctrica de las sustancias cuyas moléculas son simétricas, por ejemplo, como las moléculas de oxígeno o de cloro. En cambio, si la molécula diatómica se compone de dos átomos distintos, como, por ejemplo, la molécula del monóxido de carbono CO, en este caso debe tener lugar la dependencia de ϵ respecto a la temperatura. Y así sucede efectivamente. A las moléculas con el momento dipolar muy considerable pertenece el nitro benceno.

¿Qué ocurrirá con un dieléctrico ordinario al aumentar la intensidad del campo eléctrico? Evidentemente, debe incrementar la polarización de la sustancia. Este fenómeno tiene lugar a costa de la dilatación de los dipolos: en el átomo es la dislocación de la nube electrónica respecto al núcleo, y en la molécula puede manifestarse en forma de alejamiento recíproco de dos iones. Sea como sea, es natural poner la siguiente pregunta; ¿hasta qué momento el electrón llevado por el campo lejos del núcleo sigue siendo, como antes, electrón del átomo, y dos iones que se encuentran ya bastante distanciados uno del otro continúan formando una molécula? Indiscutiblemente, el límite sí que existe, y a una intensidad suficiente tiene lugar la llamada perforación o ruptura del dieléctrico. El orden de dicha intensidad es de varios mil kilovoltios por metro. En todo caso, la ruptura está relacionada con la liberación de electrones o de iones, es decir, con la formación de

portadores libres de corriente. El dieléctrico deja de ser tal y por el mismo fluye la corriente eléctrica.

Con mayor frecuencia, el fenómeno de ruptura se nos presenta cuando en el aparato de radio o el televisor falla el condensador. Sin embargo, conocemos también otros ejemplos de ruptura: se trata de descargas eléctricas en los gases. El tema de descarga eléctrica en los gases lo expondremos especialmente. Mientras tanto, trabajaremos conocimiento con dos importantes miembros de la familia de los dieléctricos; los piezoeléctricos y los ferroeléctricos.

El representante principal de la clase de los piezoeléctricos es el cuarzo. Los miembros de esta clase (a la cual, además del cuarzo pertenecen, por ejemplo, el azúcar y la turmalina) deben poseer una determinada simetría.



Figura 2.3

En la fig. 2.3 se representa un cristal de cuarzo. El eje principal de este cristal es un eje de simetría de 3^{er} orden. En el plano perpendicular se encuentran tres ejes de 2^{do} orden.

Mediante el procedimiento señalado en la figura, del cristal se corta una placa de cerca de 2 cm de espesor. Vemos que ésta es perpendicular al eje principal y que los ejes de 2^{do} orden se encuentran en su plano. Seguidamente, de esta placa gruesa, perpendicularmente a uno de los ejes de 2^o orden se corta una placa delgada de cerca de 0,5 mm de espesor. Con esta delgada placa piezoeléctrica obtenida de tal manera (en la figura a la derecha está desplazada hacia abajo) pueden realizarse interesantes experimentos.

Vamos a comprimir la placa a lo largo de la dirección A, perpendicular a los ejes de simetría, conectando al mismo tiempo a los planos laterales de la placa un electrómetro, o sea, el instrumento que detecta la carga eléctrica (para que exista el contacto eléctrico dichos planos deben platearse). Resulta que por efecto de la compresión en las caras de la placa aparecen cargas de distinto signo. Si en vez de la compresión se recurre a la tracción las cargas cambian de signo: donde durante la compresión aparecía la carga positiva durante la tracción se crea la carga negativa, y viceversa. Precisamente este fenómeno, es decir, la aparición de las cargas eléctricas bajo la acción de la presión o la tracción, recibió el nombre de piezoelectricidad.

Los dispositivos piezoeléctricos de cuarzo son extraordinariamente sensibles: los instrumentos eléctricos permiten medir las cargas que se generan en el cuarzo a una fuerza despreciable que no se puede medir por otros procedimientos. El piezocuarzo es capaz también marcar las variaciones muy rápidas de la presión, hecho inaccesible para otros instrumentos de medición. Esta es la razón de que el fenómeno descrito tiene enorme importancia práctica como método de registro eléctrico de toda clase de acciones mecánicas, incluyendo el sonido. Basta con soplar ligeramente sobre la placa piezoeléctrica de cuarzo para que el instrumento eléctrico haga eco de esta acción.

Las placas piezoeléctricas de cuarzo se aplican en medicina para auscultar los ruidos cardíacos del hombre. De la misma forma estas placas se aplican en la técnica al comprobar el trabajo de las máquinas: por si existen cualesquiera ruidos «sospechosos».

El cuarzo como fuente del efecto piezoeléctrico se emplea en los lectores de los tocadiscos. El movimiento de la aguja por el surco del disco proporciona la compresión del piezocristal la cual, a su vez, lleva a la generación de la señal eléctrica. La corriente eléctrica se amplifica, se suministra al altoparlante electrodinámico y se transforma en sonido.

Hasta el momento hemos hablado de sustancias cuya polarización se crea mediante el campo eléctrico, como asimismo (raras veces) por medio de la deformación mecánica. Si se quita el impacto externo, la sustancia resulta eléctricamente neutra. No obstante, a la par de este comportamiento difundido tenemos que vérselas con

cuerpos especiales que poseen un momento eléctrico total en ausencia de fuerzas externas. Está claro que no encontraremos cuerpos semejantes entre los líquidos ni gases, ya que el movimiento térmico al cual no se opone la acción ordenadora del campo conducirá, inminentemente, a un caos en la disposición de las moléculas dipolares.

Sin embargo, podemos figurarnos unos cristales en los cuales la disposición de los átomos es tal que los centros de gravedad de los aniones y cationes dentro de cada celdilla elemental (celdilla unidad) están dislocados de igual manera. Entonces, todos los momentos dipolares están orientados en un mismo sentido. En este caso, se podría esperar una polarización máximamente posible y, por consiguiente, una magnitud colosal de la constante dieléctrica.

Semejantes cristales sí que existen. Las sustancias que poseen las propiedades en cuestión se denominan ferroeléctricas. Cabe señalar, sin embargo, que, por cuanto dicho fenómeno fue descubierto por primera vez en los cristales de la sal de Seignette, en algunos idiomas (por ejemplo, en ruso y en alemán), para esta clase de sustancias se emplea también el término «seignettoeléctricos».

Entre los ferroeléctricos una enorme importancia práctica la tiene el titanato de bario. En su ejemplo, precisamente, analizaremos el comportamiento extraordinariamente peculiar de esta clase de sustancias.

La celdilla elemental de su cristal se representa en la fig. 2.4. El vértice de la celdilla se ha elegido en los átomos de bario. Los pequeños circulitos claros son los aniones de oxígeno y el círculo grande en el centro corresponde al catión de titanio.

La celdilla en el dibujo viene representada de tal modo como si fuese cúbica. Efectivamente, una celdilla estrictamente cúbica existe, pero tan sólo tan solo a la temperatura superior a 120 °C. Está claro que la celdilla cúbica es simétrica y no puede poseer el momento dipolar. A raíz de ello, por encima de esta temperatura, que se denomina punto de Curie, desaparecen las propiedades peculiares del titanato de bario. Por encima de esta temperatura dicha sustancia se comporta como un dieléctrico común y corriente.

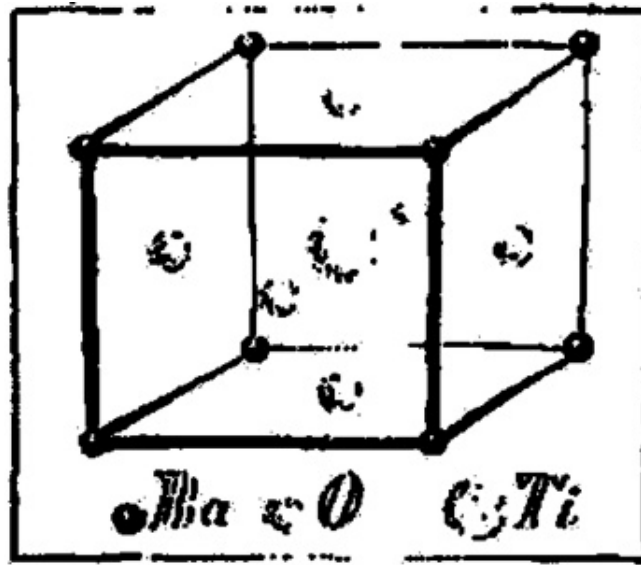


Figura 2.4

Cuando la temperatura desciende por debajo de $120\text{ }^{\circ}\text{C}$ tiene lugar el desplazamiento de los iones de oxígeno y de titanio hacia los lados opuestos a una magnitud del orden de $0,01\text{ nm}$. La celdilla adquiere el momento dipolar.

Fíjense en la siguiente circunstancia importantísima. Dicho desplazamiento puede realizarse, con igual éxito, en tres sentidos: a lo largo de los tres ejes del cubo. Las dislocaciones dan lugar a las deformaciones de las celdillas. A raíz de ello, no toda partición del cristal en dominios dentro de los cuales los momentos dipolares están orientados el mismo lado, resulta ventajosa.

En la fig. 2.5 se muestran las posibles particiones del cristal en zonas idealmente polarizadas (éstas llevan el nombre de dominios). Paralelamente al caso en que todo el cristal representa un solo dominio, o sea, el caso que conduce al máximo campo eléctrico, son posibles variantes menos ventajosas y, finalmente, incluso los casos del dibujo de extremo derecho de la figura) en que el campo que sale fuera de los límites del ferroeléctrico resulta igual a cero.

¿Cómo se comporta el ferroeléctrico al superponerle un campo eléctrico externo? Resulta que el mecanismo de polarización consiste en el crecimiento del dominio orientado en sentido «necesario» por medio de desplazamiento de los límites. Los dominios cuyos momentos están orientados de modo que forman un ángulo agudo

con el campo «se comen» a los dominios cuyo ángulo de orientación respecto al campo, es obtuso.

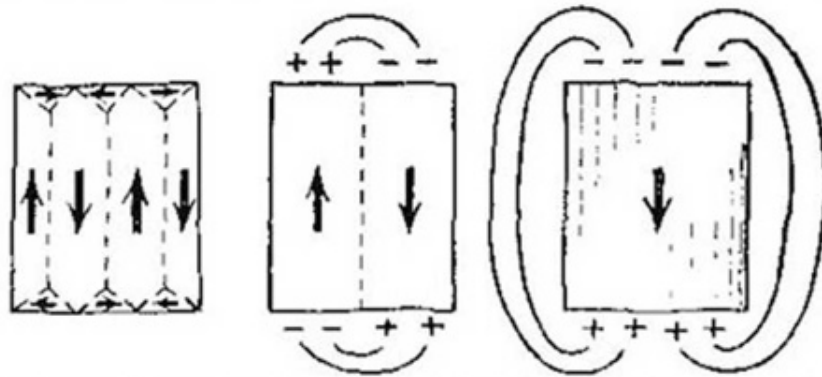


Figura 2.5

Para los campos muy grandes puede observarse también el vuelco de los dominios.

El titanato de bario es el ferroeléctrico industrial básico. Se obtiene en la sintetización de dos polvos: del dióxido de titanio y del carbonato de bario. Se forma una especie de materia cerámica.

Los ferroeléctricos cerámicos han encontrado una amplia aplicación en la electrotecnia y radiotecnia.

Y no es solamente el que aumentan de una manera ostensible la capacidad de los condensadores. Para estas sustancias, como queda claro a partir de la descripción del mecanismo de polarización, el valor de ϵ incrementará con el crecimiento de la intensidad del campo eléctrico externo. El condensador se convierte en «variconde», o sea, un condensador variable mediante el cual se realiza con máxima facilidad la modulación de frecuencia. Es un proceso que se opera en cualquier receptor de radio y televisor.

En muchas ocasiones la cerámica a base de materias ferroeléctricas sustituye el cuarzo. Valiéndose de ésta es posible engendrar un sonido más fuerte. Igualmente, en este caso es mayor el coeficiente de amplificación del ultrasonido. El campo en que el cuarzo no tiene rivales es la estabilización de radiofrecuencia.

La absoluta mayoría de los capítulos sobre la electricidad comienza con el relato acerca de las cargas eléctricas que se crean por fricción en las varillas de vidrio o de

ebonita. Pero se suele sortear la explicación de este fenómeno. ¿Cuál es la razón de ello?

En primer término se debe subrayar que la electrización de los dieléctricos por fricción no está relacionada (en todo caso, directamente) con la polarización de la cual acabamos de hablar. Efectivamente, el fenómeno de polarización es la formación de cargas eléctricas ligadas cuya singularidad consiste precisamente en que éstas no pueden «eliminarse» del dieléctrico. Las cargas engendradas en la superficie del vidrio o de la ebonita por fricción con la piel de gato son, sin duda alguna, cargas libres y, por supuesto, son electrones.

En sus rasgos generales el cuadro está más o menos claro, y nada más. Por lo visto, aquella miserable cantidad de electrones libres de que dispone el aislador, en diferentes dieléctricos está enlazada con sus moléculas por distintas fuerzas. Por lo tanto si dos cuerpos se ponen en estrecho contacto, los electrones pasarán de uno de éstos al otro. Tendrá lugar la electrización. Si a embargo, el «estrecho contacto» es la puesta de las superficies a una distancia igual a la interatómica. Por cuanto en la naturaleza no existen superficies lisas a nivel atómico la fricción ayuda a eliminar salientes de todo género y aumenta, por decirlo así, el área de verdadero contacto.

El paso de electrones de un cuerpo al otro tiene lugar para cualquier par de cuerpos: metales, semiconductores y aisladores. Sin embargo, se logra electrizar tan sólo los aisladores ya que solamente en estos cuerpos las cargas creadas quedan en los sitios en que se trasladaron al pasar de un cuerpo al otro.

No puedo decir que esta teoría deje a uno profundamente satisfecho. No está esclarecido que especiales méritos tienen la ebonita, el vidrio y la piel de gato. Puede ponerse, además, un sinnúmero de otros interrogantes para los cuales no existe una respuesta clara.

Conductibilidad de los gases

Si se llena de gas un tubo de vidrio y se sueldan a éste electrodos a los cuales se aplica una tensión, tendremos a nuestra disposición una simple instalación con cuya ayuda se puede proceder al estudio de la conductibilidad de los gases. Se puede variar la sustancia a través de la cual pasa la corriente, cambiar la presión del gas y la tensión.

Las investigaciones referentes a la conductibilidad de los gases desempeñaron un enorme papel en el desarrollo de nuestras ideas sobre la estructura eléctrica de la materia. Los trabajos principales en este campo se llevaron a cabo en el siglo XIX.

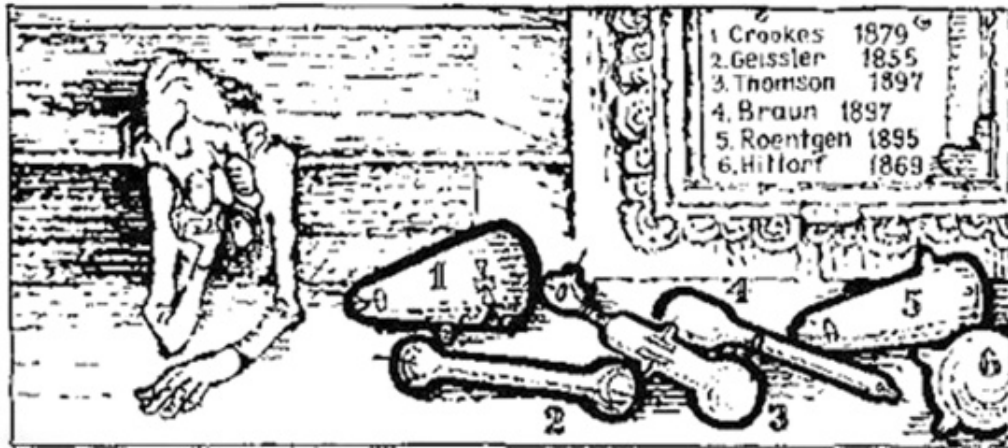


Figura 2.6

En la fig. 2.6 se representan tubos de diferente forma valiéndose de los cuales los científicos estudiaban los fenómenos que trataríamos a continuación. Por cuanto las antiguas esculturas y cuadros hace tiempo que desaparecieron de la venta, los anticuarios pasaron a hacer comercio de equipos de laboratorio, de modo que hoy día, en las tiendas de antigüedades del Occidente es posible adquirir (y, además, por un precio bastante grande) uno de los raros ejemplares de los tubos mostrados en la figura.

La causa de que en los gases se engendre la corriente eléctrica reside en que las moléculas neutras se rompen en aniones y cationes. Además, de las moléculas o de los átomos puede desprenderse un electrón. La corriente se crea por el haz de iones positivos y por los haces de iones negativos y electrones que se mueven en dirección contraria.

Para lograr que el gas se convierta en conductor de corriente es necesario transformar las moléculas neutras o los átomos en partículas cargadas. Dicho proceso puede efectuarse por acción de un ionizador externo, así como debido al choque de las partículas del gas. Como ya se ha mencionado, las fuentes externas de la ionización pertenecen los rayos ultravioleta, los rayos X, como asimismo los

rayos cósmicos y radiactivos. La alta temperatura también causa la ionización del gas.

El paso de la corriente a través de los gases viene acompañado, con frecuencia, de efectos luminosos. En dependencia de la sustancia, la presión y la intensidad la luminiscencia acusa distinto carácter. El estudio de esta luminiscencia también ha desempeñado un gran papel en la historia del desarrollo de la física, más precisamente, ha servido de fuente de los datos sobre los niveles de energía de los átomos y las leyes de la radiación electromagnética.

La conductibilidad de los gases no está sujeta a la ley de Ohm. Aquella viene definida por la curva de dependencia de la intensidad de la corriente respecto a la tensión. Dicha curva se denomina (no sólo en el caso de los gases, sino también para cualesquiera sistemas conductores que no se rigen por la ley de Ohm) característica tensión-corriente.

Examinemos los fenómenos característicos para todo gas que tienen lugar al aumentar la tensión que se aplica al tubo de descargas de gases. El comportamiento del gas que pasa a ser objeto de nuestra descripción tiene lugar en un amplio intervalo de presiones. Hacemos caso omiso solamente de aquellas presiones pequeñas para las cuales el recorrido libre de las moléculas deviene conmensurable con las dimensiones del tubo de descarga de gases. Tampoco sometemos a nuestro examen las presiones tan grandes que la densidad de los gases se aproxime a la de los líquidos.

De este modo, apliquemos al tubo de descarga de gases una pequeña tensión. Si el ionizador está ausente la corriente no fluirá a través del tubo. En presencia del ionizador en el gas se encuentran partículas cargadas: iones y electrones. Al superponer un campo éste enviará las partículas a los electrodos. La velocidad con que las partículas cargadas migrarán en dirección a los electrodos depende de muchas circunstancias y, en primer término, de la intensidad del campo y la presión del gas.

Sobre el movimiento ordenado de los iones y electrones acontecido por acción de una fuerza eléctrica constante se superpone un movimiento caótico. La partícula acelerada por el campo eléctrico recorre una distancia corta.

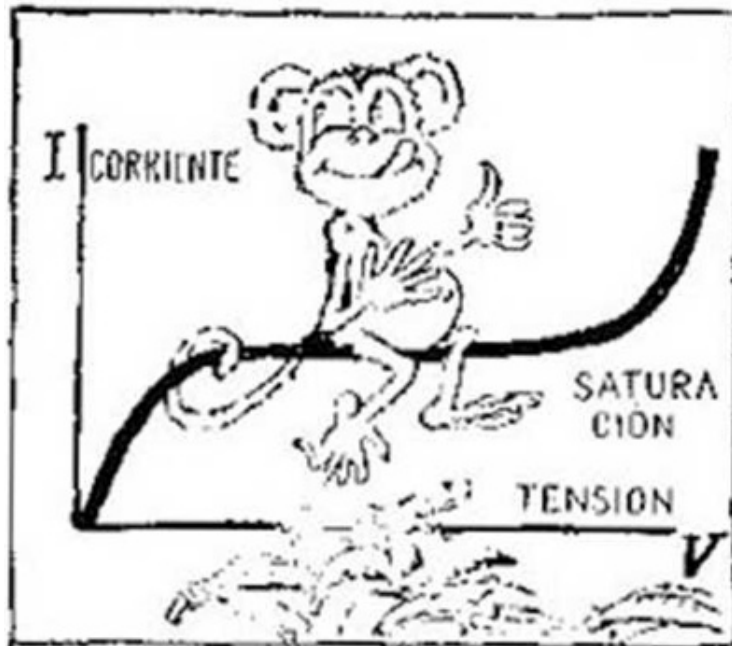
Y este corto recorrido, inminentemente, termina en una colisión. Cuando las velocidades de movimiento son pequeñas estas colisiones se producen según la ley de choque elástico.

La longitud media del recorrido libre se determina, ante todo, por la presión del gas. Cuanto más alta es la presión, tanto más corto es el recorrido libre y menor la velocidad media de movimiento ordenado de la partícula. La tensión aplicada al tubo de descarga de gases actúa en sentido inverso, o sea, aumenta la velocidad media de movimiento ordenado de las partículas.

Si no se hubiera aplicado la tensión al tubo, en el gas se habrían desarrollado los siguientes acontecimientos: el ionizador habría creado iones y los iones de distinto signo, al encontrarse unos con otros, se habrían reunificado o, como se suele decir, se habrían recombinado. Puesto que durante la recombinación tiene lugar el encuentro de un par de partículas, resulta que la velocidad de recombinación será proporcional al cuadrado del número de partículas. Si el ionizador trabaja continuamente, entre ambos procesos se establecerá un equilibrio. Así, precisamente, ocurre en la ionosfera que rodea nuestro globo terráqueo. Según sea la hora del día o de la noche y la estación del año el número de partículas ionizadas en un centímetro cúbico oscila desde un millón hasta cien millones de iones y electrones. De este modo, el grado de ionización es una magnitud del orden de uno por ciento (recuerde cuántas moléculas del aire se hallan en una unidad de volumen a grandes alturas).

Retornemos al gas ionizado que se encuentra en el tubo bajo la tensión eléctrica. Se sobreentiende que ésta altera el equilibrio, ya que una parte de los iones llega a los electrodos sin tener tiempo para recombinarse.

A medida que aumenta la tensión la parte cada vez mayor de los iones engendrados en unidad de tiempo alcanza los electrodos: la corriente eléctrica a través del gas incrementa. Así se desarrollan los acontecimientos hasta que para la recombinación, en general, no queda tiempo; en esto caso la totalidad de los iones creados por los ionizadores llega a los electrodos. Está claro que el subsiguiente ascenso de la tensión no puede aumentar la corriente (corriente de saturación, meseta en la fig. 2.7).

*Figura 2.7*

Cuanto menor es la densidad del gas tanto menores son las tensiones del campo para las cuales se alcanza la corriente de saturación.

La intensidad de la corriente de saturación es igual a la carga de los iones formados por el ionizador en un segundo en el volumen del tubo. Habitualmente, las corrientes de saturación no son grandes: son del orden de microamperios o menores. Por supuesto, su magnitud depende del número de los proyectiles destructores que el gas recibe del ionizador.

Si se trabaja en el régimen de la característica tensión-corriente que no sale de los límites de la corriente de saturación y se protege el gas contra la acción del ionizador externo, la corriente cesa. En este caso se habla de la descarga de gas no automantenida.

Con un consiguiente aumento de la presión aparecen nuevos fenómenos. Llegar un momento en que la velocidad de los electrones se hace suficiente para expulsar los electrones de los átomos y moléculas neutros. En este caso la tensión, en el tubo debe alcanzar un valor tal que el electrón tenga tiempo para acumular a lo largo de su recorrido libre una energía suficiente para ionizar la molécula. La aparición de la ionización de choque, influye en la curva de dependencia de la corriente respecto a la tensión: la corriente comienza a incrementar por cuanto el aumento de la tensión

significa el crecimiento de la velocidad de movimiento del electrón. A su vez, el crecimiento de la velocidad implica el acrecentamiento de la capacidad ionizante del electrón y, por lo tanto, la creación de un gran número de parejas de iones y el ascenso de la intensidad de la corriente. La curva de la característica tensión-corriente va rápidamente hacia arriba. En comparación con la corriente de saturación la intensidad de la corriente aumenta centonares y miles de veces, se da comienzo a la luminiscencia del gas.

Si la acción del ionizador externo se elimina ahora, la corriente no cesa. Pasamos a la región de descarga autónoma. La tensión para la cual se produce este cambio cualitativo se denomina tensión disruptiva (tensión de perforación), o bien, tensión de encendido de la descarga de gas.

El brusco aumento de la corriente después de pasar este límite crítico se explica por el crecimiento en avalancha del número de cargas.

Un solo electrón formado destruye la molécula neutra dando lugar a la creación de dos cargas de energía tan grande que éstas son capaces de desintegrar otro par de moléculas que encuentran en su camino. De dos cargas se forman cuatro; de cuatro, ocho... No se puede negar que el nombre de «avalancha» está aquí completamente justificado.

Se ha creado una teoría cuantitativa que presagia bastante bien la forma de las características tensión-corriente de los gases.

Descarga autónoma

Existen muchas variantes de esta descarga. Nos detenemos solamente en algunas de ellas.

Descarga por chispas. No es difícil observar en los experimentos más elementales la chispa que salta a través del aire entre dos electrodos. Para ello basta acercar los alambres que se encuentran bajo la tensión a una distancia lo suficientemente pequeña. ¿Qué quiere decir «suficientemente»? Si se trata del aire, entonces, es preciso crear una intensidad del campo igual a 30.000 V/cm. Esto significa que a la pequeña distancia de 1 mm nos satisface la diferencia de potencial de 300 V. Cada uno de los lectores habría observado reiteradas veces en la práctica cotidiana, las pequeñas chispas ya sea ocupándose de tendido eléctrico averiado, o bien,

acercando por casualidad uno a otro dos cables derivados de un acumulador (en este caso los cables deben aproximarse a una distancia igual al espesor de una hoja de afeitar).

La tensión disruptiva depende de la densidad del gas. También es importante la forma de los electrodos.

La chispa perfora no sólo el gas, sino también los líquidos dieléctricos y cuerpos sólidos. Al técnico electricista le es importante conocer las tensiones disruptivas de todos los materiales con que trabaja.

Hoy día nos parece completamente evidente que el rayo no es sino una chispa que salta entre dos nubes cargadas de electricidad de distinto signo. Sin embargo, en su tiempo, a los físicos [Mijaíl Vasílievich Lomonósov (1711 - 1765) y Benjamín Franklin (1706 - 1790)] costó mucho trabajo encontrar la demostración de esta hipótesis. Y Georg Richman (1711 - 1753) que trabajó junto con Lomonósov perdió la vida en el intento de poner el rayo a la Tierra por medio de cordel conductor de corriente que representaba la cola de una cometa lanzada al aire durante la tormenta.

Se pueden citar interesantes cifras que caracterizan la descarga por chispas en el rayo. La tensión entre la nube y la Tierra constituye de 10^8 a 10^9 V, la intensidad de la corriente oscila desde decenas hasta con centenares de miles de amperios y el diámetro del canal resplandeciente es de 10 a 20 cm.

La duración del resplandor del rayo es corta: del orden de microsegundo. No es difícil calcular a grandes rasgos que las cantidades de electricidad que pasan por el canal del rayo son relativamente pequeñas.

Filmando las chispas celestes se logró estudiarlas bien. Con gran frecuencia el rayo representa una serie de descargas por chispas que siguen un mismo camino. El rayo dispone de una especie de «líder» que desbroza para las cargas eléctricas el camino más conveniente y siempre caprichosamente ramificado.

Frecuentemente se han observado relámpagos esféricos. Lamentablemente, este fenómeno no se deja reproducir en las condiciones de laboratorio. Los relámpagos esféricos son bolas brillantes de plasma de gas de 10 a 20 cm de diámetro. Se desplazan lentamente y, a veces, permanecen inmóviles. Su existencia dura varios segundos o, de cuando en cuando, minutos, después de lo cual desaparecen lo que

se acompaña con una fuerte explosión. Se debe reconocer que hasta la fecha no se ha propuesto una teoría exhaustiva de este interesante fenómeno.

Descarga en arco. Esta fue obtenida por primera vez por V. V. Petrov en 1802. Con este fin ponía en contacto dos carbones conectados a una potente fuente de tensión y luego separaba los electrodos. Este procedimiento sigue en uso hasta el día de hoy. Es cierto que ahora se utilizan carbones especiales fabricados de polvo de grafito prensado. El carbón positivo se quema más rápidamente que el negativo. Por esta causa, el aspecto exterior permite determinar inmediatamente cuál de los carbones está conectado al polo positivo, ya que en el extremo de este electrodo se forma un hoyo: un cráter. La temperatura del aire en el cráter durante la presión ordinaria alcanza 4.000 grados, si la presión aumenta la temperatura del arco puede llevarse hasta casi 6.000 grados, es decir, hasta la temperatura de la superficie del Sol. El arco entre electrodos metálicos produce una llama cuya temperatura es mucho más baja.

Para mantener la descarga en arco se necesita una tensión limitada, del orden de 40 a 50 V. La intensidad de la corriente puede alcanzar cientos de amperios, por cuanto la resistencia de la columna brillante, de gas no es grande.

¿Cómo explicar, entonces, la gran conductibilidad eléctrica del gas para las tensiones tan pequeñas? Los iones moleculares se aceleran hasta velocidades no grandes y sus colisiones no pueden influir en la aparición de una corriente intensa. La explicación es la siguiente. En el primer momento en el punto de contacto tiene lugar un fuerte calentamiento. Debido a ello se inicia el proceso de emisión termoiónica: el cátodo expulsa una gran cantidad de electrones. A propósito, de aquí se infiere que la alta temperatura es importante tan sólo para el cátodo, y el ánodo puede ser frío.

El mecanismo de la descarga en arco de este tipo es completamente disímil al de la descarga por chispas.

Huelga recordar al lector cuán grande es la importancia de este fenómeno para la práctica. La descarga en arco se emplea en la soldadura y el corte de los metales, así como en la electrometalurgia.

Descarga luminiscente. Este tipo de descarga autónoma también es de gran valor práctico ya que se produce en los tubos de descarga luminiscente o, como también

se llaman, en las lámparas de luz diurna. El tubo se construye y se llena de gas (siendo la presión considerablemente menor que la atmosférica) teniendo previsto asegurar su trabajo en las condiciones de tensión superior a la de encendido. La corriente eléctrica en los tubos de descarga luminiscente se engendra por ionización de las moléculas con ayuda de los electrones, así como debido a que los electrones se expulsan del cátodo del tubo. El tubo de descarga luminiscente no se enciende de una vez. La causa de ello, por lo visto, radica en que el primer impacto debe obtenerse de un número limitado de partículas cargadas que siempre están presentes en cualquier gas.

Descarga silenciosa (efecto corona). Esta se observa a presión atmosférica en un campo en alto grado heterogéneo, por ejemplo, en la proximidad de alambres o puntas. Las intensidades deben ser altas: del orden de millones de voltios por metro. Es indiferente que polo está conectado a la punta. De este modo, puede existir tanto la corona positiva, como negativa. Por cuanto la intensidad del campo disminuye a medida que se aleja de la punta, la corona desaparece a una pequeña distancia. Se puede decir que el efecto corona es la perforación no completa del espacio gaseoso. La corona se origina debido a las avalanchas electrónicas que se dirigen a la punta, o bien, desde la punta al espacio exterior. Se sobreentiende que en la zona de la corona paralelamente a los electrones existen iones positivos y negativos: productos de destrucción de las moléculas neutras del aire. La corona produce luminiscencia sólo en aquel pequeño espacio cerca de la punta dentro del cual subsiste la avalancha electrónica.

En la aparición de la corona influyen las condiciones atmosféricas y, en primer lugar, la humedad.

El campo eléctrico de la atmósfera puede provocar la luminiscencia de las cimas de los árboles y de los mástiles de los barcos. En los tiempos pasados este, fenómeno recibió el nombre de fuegos de San Telmo. Su aparición se consideraba de mal agüero. A este prejuicio se puede encontrar una explicación racional: es completamente posible que la luminiscencia surja en aquel preciso instante en que se aproxima una tormenta o huracán.

Hace muy poco ha ocurrido una historia muy aleccionadora. Dos investigadores aficionados, los esposos Kirlian, durante muchos años estudiaron el siguiente

fenómeno. El hombre coloca su mano que está conectada una fuente de alta tensión sobre la película fotográfica separada del segundo electrodo de este circuito de corriente por una capa de aislador. Al conectar la tensión en la película fotográfica aparece un cuadro borroso de la palma de la mano y de los dedos.

Estos experimentos atrajeron suma atención de los especialistas en el campo de la llamada parapsicología, doctrina que la mayoría absoluta de los físicos y psicólogos considera como pseudociencia. La atención despertada se explica por el hecho de que los autores del descubrimiento y sus adeptos relacionaban el aspecto de la foto con el estado psíquico del sujeto.

Una amplia propaganda de esta interpretación tan extravagante de los experimentos en cuestión impulsó a un grupo de físicos y psicólogos que trabajaron en las universidades de los EEUU a que decidieran a someter dichos experimentos a una meticulosa comprobación para encontrar una explicación más sencilla del indudable hecho que el aspecto de las fotos obtenidos por este procedimiento, efectivamente, resulta distinto para diferentes personas e, incluso, para una misma persona, siendo disímiles las condiciones de obtención de la imagen fotográfica.

Los investigadores llegaron a la siguiente conclusión: «Las fotografías obtenidas por el método de Kirlian representan, generalmente, la imagen de la descarga luminiscente (efecto de corona) que tiene lugar durante la exposición. La mayoría de las diferencias en las fotos se explica por la humedad de la mano, así como por el contenido de agua en los tejidos. Durante la exposición la humedad pasa a la emulsión de la película fotográfica, cambia el campo eléctrico y el aspecto de la foto».

Los investigadores suponen utilizar en lo ulterior esta técnica, la cual prefieren llamar «fotografía de la descarga luminiscente» para el «descubrimiento y determinación de la cantidad de humedad tanto en los objetos vivos, como en los no animados».

Este interesante hecho que fue publicado en la revista «Scientific American» del diciembre de 1976, permite sacar dos conclusiones. Primero, cualquier fenómeno real merece que se fije en él y es completamente posible que éste resulte útil para la práctica. Segundo, el investigador que descubre un nuevo hecho debe superar, ante todo, la tentación de darle una interpretación que no cuadra dentro de los

marcos de las ideas científicas de nuestro tiempo. Solamente después de que se haya demostrado de un modo exhaustivo que las teorías existentes son incapaces de explicar el nuevo hecho, se podrá presentar el descubrimiento al juicio de los especialistas.

A los hechos reales a los que se da una explicación falsa les podemos llamar, recordando una vieja anécdota, experimentos con cucarachas. Su contenido es como sigue: a una cucaracha se le arrancan las patitas; esta cucaracha mutilada se pone sobre una mesa junto a otra sana y se dan unos golpecitos en la mesa. La cucaracha sana se echa a correr, mientras que la «mutilada», como es lógico, no se mueve de su lugar. Se saca la conclusión: el oído la cucaracha lo tiene en las patitas.

Cada año en la prensa aparecen varios artículos describiendo «experimentos con cucarachas». Es conveniente prevenir al lector respecto de este hecho.

Sustancia en estado de plasma

El termino «*plasmonzustand*» fue propuesto por primera vez todavía en 1939 por dos científicos alemanes cuyo artículo el autor de estas líneas había traducido para la revista «*Uspeji fizícheskij nauk*» («Éxitos de las ciencias físicas»). El nombre parece acertado. En efecto, plasma no es cuerpo sólido, ni líquido, ni tampoco gas. Es un estado especial de la sustancia.

La ionización térmica de los gases, es decir, la separación de los electrones de los átomos y la destrucción de la molécula neutra dando lugar a la formación de los iones comienza a las temperaturas superiores a 5000 a 6000 grados. ¿Vale la pena, entonces, discutir este problema? Es que no existen materiales capaces de resistir una temperatura más alta.

Sin duda alguna, sí, vale la pena. La aplastante mayoría de los cuerpos celestes, aquellos como, por ejemplo, nuestro Sol, se encuentran en estado de plasma; de ejemplo de plasma puedo servir la ionosfera. Valiéndose de campos magnéticos, las llamadas botellas magnéticas, también en las condiciones de laboratorio, se puede retener el plasma en un volumen limitado. Además, podemos hablar del plasma de la descarga en un gas. El grado de ionización del gas depende no sólo de la temperatura, sino también de la presión. El hidrógeno a una presión de 1 mm de Hg

resultará ionizado prácticamente por completo a la temperatura de 30 mil grados. En estas condiciones un átomo neutro corresponde a 20 mil partículas cargadas.

El estado de plasma del hidrógeno representa una mezcla de partículas, que se mueven desordenadamente y chocan entre sí, de dos gases: el «gas» de los protones y el «gas» de los electrones.

El plasma que se formó de otras sustancias es una mezcla de muchos «gases». En éste podemos encontrar electrones, núcleos desnudos, diferentes iones, así como una cantidad insignificante de partículas neutras.

El plasma cuya temperatura es de decenas y centenares de miles de grados se llama frío. Y el plasma caliente significa millones de grados.

Sin embargo, el concepto de temperatura del plasma hay que tratarlo con cuidado. Como conoce el lector, la temperatura se determina de manera unívoca por la energía cinética de las partículas. En un gas que consta de partículas pesadas y ligeras el equilibrio se establece solamente en el caso de, que tanto las partículas ligeras como las pesadas adquieren la misma energía cinética media. Esto quiere decir que en un gas que vive prolongadamente en condiciones estables el movimiento de las partículas pesadas es lento, mientras que el de las ligeras es rápido. El plazo en que se establece el equilibrio depende de aquello que había «en el principio». No obstante, en igualdad de las demás condiciones, el equilibrio se establece tanto más tarde cuanto mayor sea la diferencia entre las masas de las partículas.

Precisamente con tal situación tenemos que ver en el plasma. Es que las masas del electrón y del más ligero de los núcleos se diferencian casi dos mil veces. En cada colisión el electrón transfiere al núcleo o al ion tan sólo una pequeña parte de su energía. Solamente después de un gran número de encuentros las energías cinéticas medias de todas las partículas del plasma se igualan. Semejante plasma se denomina isotérmico. Tal es, por ejemplo, el plasma que se halla en las profundidades del Sol y de las estrellas. La velocidad con que se establece el equilibrio en el plasma caliente se encuentra entre los límites de partes de segundo a segundos.

De otra forma se desarrollan los asuntos en el plasma de descarga de gases (chispa, arco, etc.). En este caso, las partículas no sólo se mueven de forma

desordenada, sino también crean la corriente eléctrica. El electrón de movimiento rápido en su viaje entre los electrodos meramente no tiene tiempo para ceder la mayor parte de su energía a los iones de andar perezoso. Por esta razón en la descarga de gases el valor de la velocidad media de movimiento de los electrones es mucho más alto que el correspondiente valor para los iones. Este plasma se denomina no isotérmico y se debe caracterizar por dos temperaturas (y hasta por tres, si se toman en consideración las partículas neutras). Como es natural, la temperatura electrónica es mucho más alta que la iónica. Por ejemplo, en la descarga en arco la temperatura electrónica representa de 10 mil a 100 mil grados, mientras que la temperatura iónica es próxima a 1000 grados.

El comportamiento de las partículas en el plasma puede describirse empleando las mismas magnitudes que se utilizan en la teoría cinética de los gases. Se han elaborado muchos métodos que permiten determinar directa o indirectamente la longitud del recorrido libre de las partículas, el tiempo del recorrido libre y la concentración de las partículas de diferentes clases.

Para que el lector tenga noción de los órdenes de las magnitudes con que se tiene que operar insertemos algunos números que caracterizan el plasma de hidrógeno de alta concentración (10^{20} iones por metro cúbico). Resulta que en el plasma frío (la temperatura de diez mil grados) la longitud del recorrido libre es igual a 0,03 cm y el tiempo del recorrido libre es de 4×10^{-10} s. Este mismo plasma pero calentado hasta cien millones de grados se caracteriza por los números de 3×10^6 cm y 4×10^{-4} s, respectivamente.

Al citar estos datos es obligatorio añadir de qué colisiones se trata. Hemos mencionado los valores para los encuentros de los electrones con iones.

A todas luces es evidente que el volumen que contiene muchas partículas será eléctricamente neutro. Pero a nosotros nos puede interesar el comportamiento del campo eléctrico en algún punto del espacio. El mismo varía de una forma rápida y brusca ya que cerca de este punto, ora pasan los iones, ora vuelan raudos los electrones. Se puede calcular la rapidez de estos cambios, también es posible calcular el valor medio del campo. El plasma, con una enorme precisión, satisface la condición de neutralidad. El espíritu riguroso exige que utilicemos la palabra «cuasi neutralidad», es decir, casi neutralidad. ¿Pero qué significa este «casi»?

Un cálculo no muy complicado demuestra lo siguiente. Tracemos en el plasma un segmento de un centímetro de longitud. Calculemos la concentración de los iones y electrones en cada punto de este segmento. Cuasi neutralidad significa que estas concentraciones deben ser «casi» iguales. Ahora figurémonos que en un centímetro cúbico existe un lote «sobrante» de electrones no neutralizado por iones positivos. Resultará que a la densidad de las partículas igual a la del aire atmosférico junto a la superficie de la Tierra en el segmento examinado se forma un campo de 1000 V/cm, aproximadamente, si la diferencia en las concentraciones de los iones y electrones es igual a una milmillonésima parte de un tanto por ciento! He aquí lo que significa la palabra «casi».

Sin embargo, incluso esta alteración ínfima de la igualdad de las cargas de dos signos durará un instante brevísimo. El campo formado irá a expulsar las partículas sobrantes. Dicho automatismo opera ya en las zonas que se miden por milésimas partes de centímetro.

También en el libro 4 hablaremos sobre el plasma y las botellas magnéticas. El lector, seguramente, ha oído hablar y, puede ser que ha encontrado la descripción de las instalaciones del tipo «Tokamak». Un gran destacamento de científicos trabaja en su perfeccionamiento. El asunto reside en que la posibilidad de crear un plasma de alta temperatura pueda llevar a la fisión de los núcleos atómicos ligeros, proceso que se acompaña de liberación de una energía colosal. Los físicos consiguieron realizar este proceso en las bombas. Sin embargo, la importante cuestión de si se logra crear un plasma que posee una temperatura lo suficientemente alta, así como el suficiente tiempo de vida para que resulte realizable en la práctica una central eléctrica termonuclear, esta importante cuestión encontrará su respuesta a finales del siglo XX.

Metales

La división de los sólidos en distintas clases de acuerdo con el valor de su resistencia eléctrica se determina por la movilidad de los electrones.

La corriente eléctrica es un flujo de partículas cargadas en movimiento. Cuando se trata de los flujos de iones o electrones, literalmente vemos la corriente eléctrica. Durante su paso a través de los líquidos la corriente eléctrica también se manifiesta

de una forma completamente patente, por cuanto en los electrodos tiene lugar la sedimentación de la sustancia. En cambio, en lo que se refiere a los cuerpos sólidos debemos juzgar sólo de manera indirecta sobre aquello que representa la corriente eléctrica que por éstos circula.

Existe una serie de hechos que permiten afirmar lo siguiente. En todos los sólidos los núcleos atómicos no se desplazan. La corriente eléctrica es el flujo de electrones. Los electrones se mueven impulsados por la energía suministrada por la fuente de corriente. Esta fuente le genera dentro del cuerpo sólido el campo eléctrico.

La fórmula que relaciona la tensión y la intensidad del campo eléctrico es válida para cualesquiera conductores. Esta es la razón de que aunando las fórmulas insertadas en páginas anteriores podemos escribir la ley de Ohm para un conductor sólido en la forma

$$j = \sigma E$$

($\sigma = 1/\rho$ se denomina conductancia específica).

Los electrones del sólido pueden dividirse en ligados y libres. Los electrones ligados pertenecen a átomos determinados y los electrones libres forman un gas electrónico sui generis. Estos electrones pueden desplazarse por el cuerpo sólido. En ausencia de tensión eléctrica, el comportamiento de los electrones libres carece de orden. Cuanto más obstaculizado ven los electrones libres su movimiento, cuanto más frecuentes son sus choques con los átomos inmóviles y entre ellos mismos, tanto mayor es la resistencia eléctrica del cuerpo.

En los dieléctricos, la absoluta mayoría de los electrones tiene su amo: un átomo o una molécula. El número de electrones libres es despreciable.

En los metales cada átomo entrega uno o dos electrones para el uso común. Precisamente este gas electrónico sirve de transportador de corriente.

Partiendo de un modelo muy burdo podemos determinar de una forma aproximada, el valor de la conducción eléctrica y comprobar este modelo.

Análogamente como hicimos cuando razonábamos acerca del gas de las moléculas supongamos que cada electrón logra recorrer, sin sufrir una colisión, cierto trecho.

La distancia entre los átomos del metal es igual a varios angstrom. Es lógico suponer que la longitud del recorrido libre de los electrones por el orden de su magnitud es igual a $10 \text{ \AA}$, o sea. 10^{-7} cm .

Por acción de la fuerza aceleradora eE el movimiento del electrón hasta la colisión se prolonga durante el tiempo l/v . Basándose en los datos tomados de las investigaciones de la emisión termoiónica se puede evaluar la velocidad caótica y de los electrones. El orden de esta velocidad es de 10^8 cm/s .

Para determinar la velocidad del movimiento ordenado de los electrones, es decir, la velocidad del movimiento que engendra la corriente, es necesario multiplicar la aceleración eE/m por el tiempo del recorrido libre. Con este proceder se admite que cada colisión hace parar el electrón después de lo cual éste comienza a acelerarse de nuevo. Una vez efectuada la multiplicación obtenemos el valor de la velocidad del movimiento dirigido de los electrones que, precisamente, crea la corriente;

$$u = \frac{eEl}{mv}$$

Ahora planteémonos el problema de calcular la resistencia específica del metal. Si el orden de la magnitud resulta correcto, entonces, esto significa que nuestro modelo «funciona».

Dejemos atinar al lector que la densidad j de la corriente puede escribirse como producto del número de electrones en una unidad de volumen por la carga del electrón y la velocidad del movimiento ordenado: $j = neu$. Al sustituir en esta fórmula el valor de la velocidad del movimiento ordenado de los electrones, obtenemos

$$j = \frac{ne^2 l}{mv}$$

es decir, la conductancia específica es igual a

$$\sigma = \frac{ne^2 l}{m v}$$

Si se considera que cada átomo entrega para el uso común un solo electrón, resultará que el conductor tiene la conductancia específica del orden de $10^{-5} \Omega\text{-m}$. ¡Es una magnitud muy razonable! Esto confirma tanto la validez de nuestro burdo modelo, como la elección acertada de los valores de los parámetros de nuestra «teoría». Pongo la palabra «teoría» entre comillas por la única causa de que es burda y elemental. No obstante, este ejemplo ilustra el típico enfoque físico en la interpretación de los fenómenos.

De acuerdo con la teoría del gas electrónico libre la resistencia eléctrica debe disminuir con la caída de la temperatura. Pero cuidado con apresurarse a vincular esta circunstancia con la variación de la velocidad caótica de movimiento de los electrones. El asunto no estriba en ésta. Dicha velocidad depende poco de la temperatura. La disminución de la resistencia está relacionada con el hecho de que la amplitud de las oscilaciones de los átomos llega a ser menor debido a lo cual se acrecienta la longitud del recorrido libre de los electrones.

El mismo hecho puede formularse también con las siguientes palabras: al aumentar las amplitudes de las oscilaciones de los átomos los electrones se dispersan en mayor grado por distintos lados. Por supuesto, debido a ello la componente de la velocidad en la dirección de la corriente debe disminuir, es decir, la resistencia debe crecer.

El que la resistencia del metal (y no sólo del metal) crece al añadirle impurezas también se explica por el aumento de la dispersión de los electrones. En efecto, los átomos de las impurezas desempeñan el papel de defectos de la estructura cristalina y, en consecuencia, contribuyen a la dispersión de los electrones.

La energía eléctrica se transmite por alambres. Debido a la resistencia eléctrica en los alambres se consume la energía de la fuente de corriente. Las pérdidas de energía son enormes y la lucha contra las mismas representa una tarea técnica de suma importancia.

Se abriga la esperanza de que esta tarea pueda resolverse ya que existe el admirable fenómeno de superconductividad.

El físico holandés Kamerlingh-Onnes, en 1911, descubrió que a temperaturas próximas a cero absoluto algunos cuerpos, a salto, pierden prácticamente por completo su resistencia eléctrica. Si en el anillo de un superconductor se excita la corriente eléctrica, ésta fluirá en el conductor durante días enteros sin atenuarse. Entre los metales puros es el niobio el que posee la más alta temperatura (9 K) a la cual se manifiestan las propiedades de superconductividad. Huelga decir con qué tenacidad el enorme destacamento de científicos se dedica a la búsqueda de superconductores que puedan adquirir esta maravillosa propiedad a una temperatura más alta. Sin embargo, por ahora los éxitos no son muy grandes; se ha encontrado una aleación que se transforma en superconductora a la temperatura de cerca de 20 K.

En todo caso, hay motivo para suponer que este límite podrá elevarse (y, probablemente, hasta llevar a la temperatura ambiente). La búsqueda se lleva a cabo entre sustancias polímeras especiales y entre materiales laminares compuestos en los cuales se alternan el dieléctrico y el metal. Es difícil sobreestimar la significación de este problema. Yo asumo la responsabilidad de considerarla uno de los problemas importantísimos de la física moderna

Los trabajos dedicados a la búsqueda de los superconductores que adquieren dicha propiedad a temperaturas lo suficientemente altas tomaron una gran envergadura después de haberse confirmado la teoría de este fenómeno. La teoría sugirió las vías de búsqueda de los materiales necesarios.

Es característico que entre el descubrimiento del fenómeno y su explicación transcurrió un largo lapso. La teoría fue creada en 1917. Cabe señalar que las leyes de la física cuántica con cuya ayuda está construida la teoría de la superconductividad se habían establecido todavía en 1920. De aquí se infiere que la teoría del fenómeno distaba mucho de ser simple. En este libro sólo puedo dar la explicación, por decirlo así, desde la mitad de la historia. Resulta que a medida que las oscilaciones de la red atómica se van atenuando, algunos electrones logran «aparearse». Semejante «par» se comporta de una manera acorde. Cuando tiene lugar la dispersión del par en los átomos (y, como hemos mencionado con anterioridad, precisamente esta dispersión sirve de causa de la resistencia), el rebote de uno de los miembros del par a un lado se compensa por el

comportamiento de su «amigo». Se compensa en el sentido de que el impulso total del par de electrones se mantiene invariable. De este modo, la dispersión de los electrones sigue en pie, mas deja de influir en el paso de la corriente.

Paralelamente a los electrones apareados en el superconductor existe también el gas electrónico ordinario. De este modo, simultáneamente, parece como si existieran dos líquidos: uno, común y corriente, y otro, superconductor. Si la temperatura del superconductor comienza a elevarse desde cero, entonces, el movimiento térmico romperá cada vez mayor número de «pares» de electrones y la parte correspondiente al gas electrónico ordinario irá en aumento. Por fin, llegará la temperatura crítica a la cual desaparecerán los últimos electrones apareados.

En el libro 2, valiéndonos del modelo de dos líquidos, el ordinario y el especial, hemos explicado el fenómeno de superfluidez observado en el helio líquido. Esos dos fenómenos guardan un parentesco próximo: la superconductividad es la superfluidez del líquido electrónico

El par de electrones del que acabamos de hablar tiene el espín total cero. Las partículas cuyo espín es igual a cero o a un número entero se denominan bosones. En determinadas condiciones los bosones pueden acumularse en gran cantidad en un mismo nivel de energía. En este caso su movimiento se hace idealmente acorde y nada puede impedir su traslación. En el libro 4 volveremos a hablar sobre este fenómeno.

Salida de los electrones del metal

Por cuanto parte de los electrones se comporta a semejanza del gas formado de partículas rápidas es natural esperar que los electrones sean capaces de salir a la superficie del metal. Para que el electrón pueda abandonar el metal aquél debe superar las fuerzas de atracción de los iones positivos. El trabajo que el electrón debe invertir para lograr este objetivo se denomina trabajo de salida.

Cuanto más alta es la temperatura del metal, tanto mayor es la velocidad cinética de movimiento de los electrones. Si el metal se pone incandescente, entonces logrará abandonarlo un número considerable de electrones.

El fenómeno de emisión termoiónica -así se llama la salida del metal de los electrones - puede investigarse por medio de un sencillo experimento. En una

bombilla eléctrica se suelda un electrodo adicional. Mediante un instrumento sensible es posible medir el valor de la corriente eléctrica que se creará debido a que una parte de los electrones que se «evaporan» va a parar al electrodo (una parte y no todos por la simple razón de que los electrones salen del filamento de la lámpara bajo diferentes ángulos).

Si se quiere estimar el trabajo de salida conviene recurrir a la tensión «de bloqueo» (o de corte), es decir, conectar el electrodo soldado al polo negativo del acumulador. Aumentando paulatinamente la tensión llegaremos a tal valor suyo para el cual los electrones ya no pueden alcanzar el electrodo.

El trabajo de salida de los electrones para el wolframio es igual a 5 eV, aproximadamente. Si es necesario, empleando recubrimientos especiales se puede bajar el valor de este trabajo hasta 1 eV.

¿Y qué representa esta unidad de trabajo, el electronvoltio? No es difícil comprender, basándose en su denominación, que es igual a la energía que adquiere el electrón al recorrer el tramo del camino que se encuentra bajo la tensión de 1 V, Un electronvoltio vale $1,6 \times 10^{-19}$ J. Aunque las velocidades térmicas de los electrones son bastante considerables, la masa del electrón es muy pequeña. Por esta causa, la altura mencionada de la barrera es muy elevada. La teoría y la experiencia muestran que la salida de los electrones depende ostensiblemente de la temperatura. El ascenso de la temperatura desde 500 hasta 2000 K implica el incremento de la corriente de emisión en miles de veces.

La salida de los electrones del metal debido al movimiento térmico es, por decirlo así, un proceso natural. Pero también se puede expulsar el electrón del metal.

En primer término, se puede hacerlo bombardeando el metal también con los electrones. El fenómeno lleva el nombre de emisión electrónica secundaria. Este se utiliza para la multiplicación de los electrones en los instrumentos técnicos.

Inconmensurable más esencial es la expulsión de los electrones a partir de cuerpos sólidos mediante la luz. Dicho fenómeno lleva el nombre de efecto fotoeléctrico.

Fenómenos termoeléctricos

Hace mucho tiempo (para la evolución de la humanidad no es sino un instante, mas para el desarrollo de la ciencia casi una eternidad), más de 150 años atrás, se había

descubierto un hecho simple. Si se forma un circuito eléctrico de un trozo de alambre de cobre y otro de alambre de bismuto, soldándolos en dos puntos, por el circuito fluirá la corriente. Pero fluirá sólo en el caso de que una de las uniones soldadas se encuentre a una temperatura más alta que la otra. Dicho fenómeno se llama, precisamente, termoelectricidad.

¿Qué es, entonces, lo que obliga los electrones a desplazarse por el conductor compuesto? El fenómeno resulta ser no tan simple. La fuerza electromotriz se engendra debida a dos circunstancias. Primero, el campo eléctrico de contacto; segundo, el campo eléctrico de temperatura.

Acabamos de señalar que para la salida del electrón fuera del metal se necesita un trabajo. Es lógico suponer que este trabajo de salida A no es idéntico para distintos metales. Siendo así, entre dos metales unidos por soldadura aparece una tensión igual a

$$\frac{1}{e}(A_1 - A_2)$$

Por vía experimental se puede cerciorar de la existencia de la tensión de contacto. Pero, de por sí, ésta no puede ser causa de la corriente eléctrica en un circuito cerrado. En efecto, el circuito cerrado consta de dos uniones soldadas y las tensiones de contacto se extinguirán mutuamente. ¿Por qué, entonces, la diferencia de temperatura de las uniones soldadas engendra la fuerza electromotriz? Es la lógica la que sugiere la respuesta. Evidentemente, la tensión de contacto depende de la temperatura. El calentamiento de una de las uniones soldadas torna desiguales las tensiones provocando la aparición de la corriente. Sin embargo, hace falta tomar en consideración también otro fenómeno. Es completamente natural suponer que entre los extremos del conductor exista un campo eléctrico, si estos extremos se encuentran a diferentes temperaturas. Es que a una temperatura más alta los electrones se mueven con mayor rapidez. Si es así, comenzará la difusión de las cargas eléctricas que se prolongará hasta que se forme un campo que compense la tendencia a la distribución uniforme.

Los experimentos no dejan duda respecto a que ambos fenómenos están presentes simultáneamente y que los dos deben tenerse en cuenta en la creación de la teoría. Las fuerzas termoelectromotrices no son grandes: son del orden de varios milivoltios para la diferencia de temperatura de 100 grados. Pero estas tensiones se miden fácilmente. Por esta causa, el efecto termoelectromotriz se utiliza para la medición de las temperaturas. Es que no se puede sumergir un termómetro de vidrio en la masa fundida de un metal. Precisamente, en estos casos el par termoeléctrico (así se denomina el elemento térmico empleado para la medición de la temperatura) resulta ser un magnífico instrumento. Desde luego, el par termoeléctrico (o, simplemente, termopar) posee también otros muchos méritos). ¡Cuán esencial es la posibilidad de medir la temperatura a grandes distancias! ¡Y la sensibilidad! Las mediciones eléctricas son en alto grado exactas y resulta que, empleando un termopar, pueden medirse las diferencias de temperatura iguales a millonésimas partes de grado.

Esta alta sensibilidad permite utilizar los termoelementos para medir flujos térmicos que llegan desde objetos lejanos. El lector por sí mismo puede apreciar aproximadamente las posibilidades de un termoelemento. Basta decir que para éste no representan un límite décimas partes de ergio por segundo.

De la misma forma que los acumuladores los termoelementos, a veces, también se reúnen en baterías. Si no se requiere una energía muy grande, semejante batería puede servir de generador de energía, encontrando su aplicación para la comunicación por radio.

Semiconductores

Muchas sustancias, tanto elementos como compuestos químicos, en correspondencia con los valores de su conductibilidad llenan el amplísimo intervalo entre los conductores y los dieléctricos. Hace mucho tiempo ya se daban cuenta de la existencia de tales sustancias. Sin embargo, apenas unos seis lustros atrás alguien difícilmente había podido prever que la física de los semiconductores haría nacer una rama de la industria cuya trascendencia es imposible sobreestimar. No hay semiconductores, esto significa que no existen ordenadores modernos, ni

tampoco televisores ni magnetófonos. Sin los semiconductores es inconcebible la radiotecnica moderna.

La conductibilidad de los aisladores se encuentra entre 10^{-8} a $10^{-18} \Omega^{-1} \text{ m}^{-1}$, la de los metales tiene los valores entre 10^2 y 10^4 de estas mismas unidades. La conductibilidad específica de los semiconductores se ubica entre estos dos intervalos. Sin embargo, nos enteramos de que se trata de un semiconductor no solamente por el valor de su conductibilidad.

Igualmente que en el caso de los metales, al fluir la corriente por los semiconductores no observamos cambios químicos algunos. Este hecho significa que los iones de estas sustancias los cuales forman la armazón de la red cristalina no se trasladan por acción del campo. Por lo tanto, al igual que en los metales, debemos atribuir la conductibilidad al movimiento de los electrones.

Aunque parece que esta circunstancia es axiomática, sin embargo, en los albores del estudio de los semiconductores los físicos decidieron comprobar, por si acaso, que cargas son transportadores de corriente. En el caso de los sólidos esta comprobación puede realizarle por medio del efecto Hall.

En el siguiente capítulo voy a recordar al lector que por acción del campo magnético las partículas positivas y negativas se desvían y, además, en diferentes sentidos.

Si el cuerpo sólido en cuyo interior se mueven las cargas se confecciona en forma de banda y se coloca en un campo magnético correspondientemente orientado, entre los extremos de la banda aparece la tensión. EL esquema del experimento se ilustra en la fig. 2.8.

Cuál fue el asombro de los físicos al descubrir que se debe tratar con los cuerpos los cuales durante su investigación, de acuerdo con el esquema señalado, se comportan en unas ocasiones como si por el alambre se desplazaran partículas positivas, mientras que en otras ocasiones parece como si los transportadores de la electricidad tuviesen signo negativo.

No es difícil dar un nombre a esta actitud suya. En el primer caso hablaremos sobre la conductibilidad positiva (tipo p), y en el segundo, sobre la negativa (tipo n). Pero el quid de la cuestión no reside en el nombre, sino en la explicación de la médula del asunto. No hay duda alguna que en el seno del semiconductor se mueven

electrones. ¿Cómo, entonces, salir de la contradicción? ¿De qué modo explicar la conductibilidad positiva?

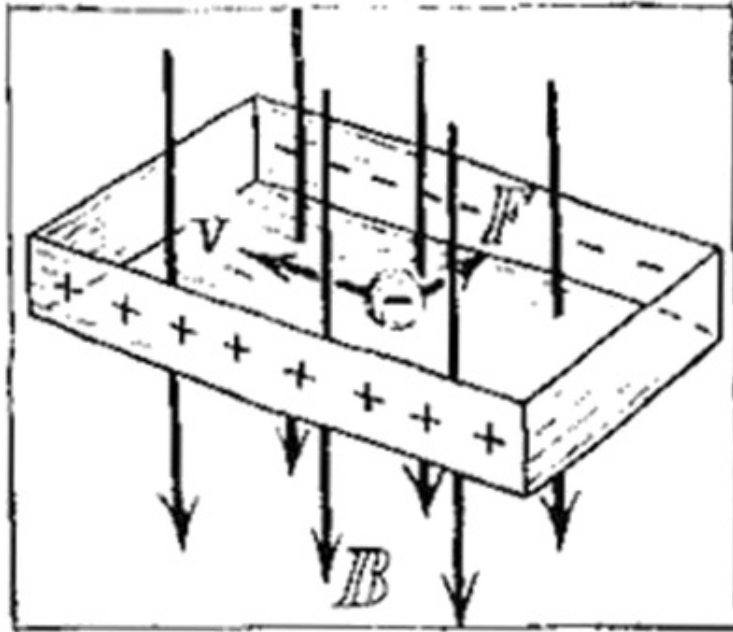


Figura 2.8

Figúrense una fila de deportistas. De pronto, un hombre la abandona por cualquier causa. Queda un sitio vacante. Aunque esto suena no muy bello, digamos así: se formó un «hueco». Para alinear la fila se da la orden al vecino del «hueco» que se traslade al sitio libre. Pero entonces, como queda completamente claro, se forma un nuevo sitio vacante. Este también puede llenarse mandando al siguiente deportista ocupar la plaza correspondiente al «hueco». Si los deportistas se trasladan de derecha a izquierda, el «hueco» se desplazara de izquierda a derecha. Precisamente este esquema explica la conductibilidad positiva de los semiconductores.

La concentración de los electrones libres en los semiconductores es muy pequeña. Por eso, ya el propio valor de la conductibilidad (acuérdense de la fórmula que hemos deducido hace poco para la densidad de la corriente) nos sugiere que la mayoría de los átomos del semiconductor no son iones, sino átomos neutros, sin embargo, el semiconductor, a pesar de todo, no es dieléctrico. Esta circunstancia significa que el número pequeño de electrones se ha dejado en libertad. Dichos electrones se desplazarán como en un metal creando la conductibilidad negativa, o

sea, electrónica. Pero el ion positivo rodeado de átomos neutros se encuentra en estado inestable. Apenas al cuerpo sólido se aplica el campo eléctrico este ion positivo trata de «sonsacar» el electrón de su vecino; de una forma absolutamente análoga obrará el átomo vecino. El ion positivo es completamente análogo al «hueco». La intercepción de los electrones puede sobreponerse al movimiento de los electrones libres. De este modo surge la conductibilidad positiva, llamada también conducción por huecos.

Pero, ¿si este modelo no lo gusta? Puedo ofrecerle otro. Como hemos dicho, la energía de las partículas se cuantifica. Esta es la ley fundamental de la naturaleza. Todos los fenómenos que se desarrollan en los semiconductores se explican perfectamente si admitimos que, al igual que en el átomo, también en el sólido los electrones están distribuidos por los niveles de energía. Sin embargo, puesto que en el sólido los electrones son muchísimos, resulta que ahora los niveles parece como si confluyeran en bandas de energía (o, llamándoles de otra forma, en zonas de energía).

La interacción de los electrones internos entre sí es débil, y, por lo tanto, el ancho de la banda será no grande. Sobre los átomos internos no influye, prácticamente, el hecho de que los átomos a los cuales dichos electrones pertenecen, entran en la composición del cuerpo sólido.

En cambio, cuando se trata de los electrones externos el asunto es otro. Sus niveles forman las bandas. Para diferentes cuerpos el ancho de estas bandas y las «distancias» entre éstas son distintos (en este caso se debe decir «intervalos de energía» y la palabra «distancias» empleada en este contexto es el argot físico).

El cuadro presentado explica de una forma excelente la división de los cuerpos sólidos, de acuerdo con su conductibilidad eléctrica, en metales, semiconductores y aisladores (fig. 2.9).

Cuando la banda está por completo llena de electrones y la distancia hasta la banda superior vacía es grande, el cuerpo es un aislador. Si la banda superior está llena de electrones parcialmente semejante cuerpo es un metal, puesto que un campo eléctrico, por muy pequeño que sea, es capaz de hacer pasar el electrón a un nivel de energía un poco más alto.



Figura 2.9

Un semiconductor se caracteriza por el hecho de que su banda superior está separada de la más próxima inferior por un pequeño espacio. A diferencia de los aisladores y metales, en el caso de los semiconductores el movimiento térmico es capaz de trasladar los electrones de una banda a otra. En ausencia del campo el número de semejantes transiciones hacia arriba y hacia abajo es idéntico. El aumento de la temperatura sólo lleva aparejado el hecho de que la concentración de los electrones en la banda superior crece.

¿Pero qué ocurrirá cuando sobre el semiconductor se superponga un campo?

En este caso el electrón libre que se halla en la banda superior comenzará a desplazarse haciendo una aportación a la conductibilidad negativa. Pero se alterará el equilibrio de las transiciones hacia abajo y hacia arriba. Por esta causa, en la

banda inferior se forma un «hueco» que por acción del campo se desplazará hacia el sentido contrario. Tales semiconductores se denominan de conductibilidad mixta (positivo-negativa o conductibilidad p - n).

La teoría de bandas de los semiconductores es una teoría armoniosa. El lector no debo suponer que el modelo descrito es artificial o ideado. Este explica de una forma sencilla y clara la principal diferencia entre el metal y el semiconductor, a saber, su comportamiento especial relacionado con el cambio de la temperatura. Como se ha señalado en el párrafo anterior, con el ascenso de la temperatura disminuye la conductibilidad eléctrica de los metales, ya que los electrones más a menudo chocan contra obstáculos. En los semiconductores la elevación de la temperatura implica el aumento del número de electrones y huecos y, en consecuencia, el aumento de la conductibilidad. Este efecto, como lo demuestran los cálculos, supera considerablemente la disminución de la conductibilidad debido a los impactos contra los obstáculos.

Para la técnica el valor principal lo tienen los conductores con impurezas. En este caso se logra crear cuerpos poseedores de conductibilidad solamente positiva o solamente negativa. La idea es extraordinariamente sencilla.

Los semiconductores más difundidos son el germanio y el silicio. Estos elementos son tetravalentes. Cada átomo está enlazado con cuatro vecinos. El germanio perfectamente puro representará un semiconductor de tipo mixto. El número de huecos y electrones correspondientes a 1 cm^3 es muy pequeño, a saber, es igual a $2,5 \times 10^{13}$ a temperatura ambiente. Esto significa un electrón libre y un hueco, aproximadamente, por mil millones de átomos.

Sustituyamos ahora uno de los átomos de germanio por un átomo de arsénico. El arsénico es pentavalente. Cuatro de sus electrones están destinados para enlazarse con los átomos del «amo», o sea, del germanio, mientras que el quinto será libre. El material acusará conductibilidad electrónica (negativa) puesto que la aparición del átomo de germanio, como se sobreentiende, no dará lugar a la formación de huecos.

Si la impureza de arsénico es completamente despreciable constituyendo un átomo por millón, la conductibilidad del germanio aumentará ya mil veces.

Ya queda absolutamente claro qué se requiere para convertir el germanio en conductor del tipo p. Para lograr este objetivo es necesario sustituir el átomo de germanio por un átomo trivalente, digamos, por el átomo de indio.

En este caso la situación toma el siguiente cariz. El átomo de germanio que se encuentra al lado del «huésped» se transforma en ion positivo por cuanto, de grado o por fuerza, tendrá que formar enlace con el átomo de indio al que falta un electrón. Pero ya conocemos que el ion positivo hace las veces de hueco. Por impacto del campo el «hueco» se desplazará; en cambio, no habrá movimiento de electrones libres.

No es de asombrarse que la industria de los semiconductores ejerciera una enorme influencia en la técnica de formación de cristales puros. ¡Y no puede ser de otro modo si unas millonésimas partes de impurezas juegan un papel decisivo!

Sería erróneo formarse la idea de que en los conductores del tipo n falla la conducción por huecos. Los huecos existen, pero su número es sustancialmente menor que el de los electrones libres. En el caso de los semiconductores del tipo n los electrones son portadores principales de corriente, mientras que los huecos, que constituyen una minoría, se denominan portadores secundarios de corriente. Por el contrario, en los semiconductores del tipo p los portadores principales son los huecos, y secundarios, los electrones.

Unión p-n

Después que ha quedado esclarecido qué representan los semiconductores p y los semiconductores n podemos llegar a comprender un efecto muy interesante y de gran importancia para la electrónica moderna. El efecto aparece en la zona de transición entre los semiconductores p y n que están unidos estrechamente entre sí (unión p - n). La palabra inglesa «transition» (transición, paso) sirvió de base para nombrar toda una clase de instrumentos cuyo trabajo tiene por fundamento la unión p - n. ¿Qué ocurrirá, entonces, si se toman dos barras de igual sección con caras frontales pulidas muy, pero muy esmeradamente, con la particularidad de que una barra está hecha de germanio con impureza de indio (semiconductor del tipo p), y la otra barra, de germanio con impureza de arsénico (semiconductor del tipo n), y si, seguidamente, estas barras se juntan entre sí apretando estrechamente una a

otra sus caras frontales? Obtenemos, de hecho, un cristal único de germanio con la única diferencia de que en una mitad suya habrá un exceso de electrones libres, y en la otra, un exceso de huecos.

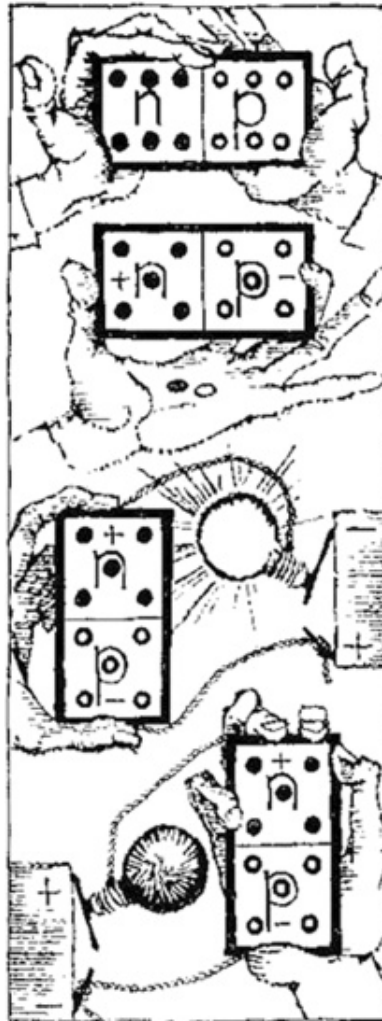


Figura 2.10

Para evitar que las explicaciones sean muy complicadas olvidemos de los portadores de corriente secundarios. En el momento inicial de tiempo (véase la fig. 2.10, por arriba) ambas mitades del cristal son eléctricamente neutras. Pero en la parte n existe (a pesar de la neutralidad eléctrica) un número «sobrante» de electrones (puntos negros), mientras que en la parte derecha, la parte p, del «sándwich» «sobran» los huecos (circulitos).

Tanto los electrones, como los huecos pueden pasar libremente a través de la frontera. La causa de la transición es completamente igual a la que tenemos en el caso de mezclado de dos gases cuando se comunican los recipientes que los contienen. Sin embargo, a diferencia de las moléculas de gas, los electrones y los huecos son aptos para la recombinación.

Hemos tenido a la izquierda seis puntitos negros, y a la derecha, seis circulitos. Apenas hubo comenzado la transición el circulito y el punto se aniquilaron mutuamente. El siguiente esquema muestra que en la parte izquierda quedaron menos electrones de lo necesario para que esta mitad del sándwich fuese eléctricamente neutra; la parte derecha también tiene un círculo menos.

Al quitar el electrón a la parte izquierda la cargamos positivamente; por la misma causa la parte derecha adquirió una carga negativa.

Para los subsiguientes huecos y electrones el paso a través de la frontera ya está obstaculizado. Tienen que moverse contra el campo eléctrico formado. La transición durará cierto lapso, mientras el movimiento térmico sea capaz de superar la barrera energética cada vez creciente, luego llegará el equilibrio dinámico.

¿Y qué será si al sándwich p-n se aplica la tensión y, además, de tal forma como se representa en el tercer esquema desde arriba? Evidentemente, en este caso a los portadores de corriente se comunica una energía suplementaria que les permite salvar la barrera.

Por el contrario, si a la parte n se conecta el polo positivo, la transición de los huecos y electrones a través de la frontera sigue siendo imposible.

De este modo resulta que la unión p-n posee propiedades rectificadoras

Actualmente, en los más diferentes campos de la técnica se utilizan rectificadores (válvulas, diodos, nombres que, en esencia, son sinónimos), cuyo principio de funcionamiento acabamos de explicar.

Nuestro esquema es sumamente aproximado. El análisis hemos hecho caso omiso de cualesquiera detalles del comportamiento de los huecos y electrones capaces de saltar a través de la frontera sin la recombinación y. lo que es lo primordial, se ha despreciado la corriente de los portadores secundarios, circunstancia que tiene por consecuencia el que la rectificación de la corriente por el sándwich p-n no es

completa. En la realidad, a pesar de todo, al aplicar la tensión de acuerdo con el esquema inferior tiene lugar una corriente débil.

Ahora caractericemos con mayores pormenores los acontecimientos que ocurren en la frontera cuando se establece el equilibrio dinámico.

Renunciemos a la simple conjetura admitida con anterioridad, o sea, acordémonos de la existencia de los portadores secundarios.

El cuadro de establecimiento del equilibrio dinámico será el siguiente. Del seno del cristal p, cada vez más cerca de la frontera, comenzará a crecer la corriente por huecos. La aportación a ésta la hacen los huecos que tendrán tiempo para llegar hasta la unión p- n y salvarla sin recombinarse con los electrones.

Por supuesto, estos huecos deben poseer, además, una energía suficiente como para saltar sobre la barrera de potencial.

Después de atravesar la zona de transición esta corriente comienza a extinguirse poco a poco debido a la recombinación con los electrones. Al mismo tiempo en la parte n desde la profundidad, en dirección opuesta crece la corriente por huecos. En esta zona la cantidad de huecos es mucho menor, pero, en cambio, éstos no deben salvar la barrera para ir a parar a la zona p, se puede decir que la barrera se adapta de tal forma, que las corrientes directa e inversa se compensan una a otra.

Todo lo expuesto se refiere también a la corriente electrónica. Es cierto que los valores de la corriente por huecos y de la electrónica pueden diferenciarse considerablemente debido a que las zonas p y n están de distinto modo enriquecidas de impurezas y, por consiguiente, de portadores libres. Si, por ejemplo, en la zona p hay mucho más huecos que electrones en la zona n, la corriente por huecos será mucho mayor que la electrónica. En este caso la zona p se denomina emisor (radiador) de portadores libres de corriente y la zona n lleva el nombre de base.

Esta descripción más detallada de los acontecimientos que tienen lugar en la frontera p-n nos permite comprender que la rectificación de la corriente no puede ser completa.

En efecto, si el polo positivo se conecta al cristal p la barrera se tornará más baja. La tensión instiga los electrones. Si el polo positivo está conectado a la parte n, entonces, el campo eléctrico creado por el manantial de la alimentación coincide por

su dirección con el campo de la barrera. El campo en la transición aumentará. Ahora disminuirá el número de electrones capaces de salvar la barrera, al igual que el de huecos aptos para pasar al lado opuesto. De aquí el crecimiento de la resistencia en la zona de transición que conduce a la llamada característica no simétrica tensión-intensidad.

Resumiendo, podemos decir que el análisis más profundo dilucida patentemente cuál es la causa de que la rectificación operada en la capa de transición no será completa.

Capítulo 3

Electromagnetismo

Contenido:

- *Medida del campo magnético*
- *Acciones del campo magnético homogéneo*
- *Acciones del campo magnético no homogéneo*
- *Corrientes amperianas*
- *La nube electrónica del átomo*
- *Momentos magnéticos de las partículas*
- *Inducción electromagnética*
- *Dirección de la corriente de inducción*
- *Algunas palabras acerca de la historia del descubrimiento de la ley de la inducción electromagnética*
- *Corrientes de inducción en torbellino*
- *Choque de inducción*
- *Permeabilidad magnética del hierro*
- *Dominios*
- *Cuerpos diamagnéticos y paramagnéticos*
- *El campo magnético de la Tierra*
- *Campos magnéticos de las estrellas*

Medida del campo magnético

Desde tiempos muy remotos los hombres se daban cuenta de la interacción de las varillas y agujas hechas de algunas menas de hierro. Estos objetos se distinguían por una propiedad singular: uno de los extremos de las agujas indicaba el norte. De este modo a la aguja se podía atribuir la posesión de dos polos: el polo norte y el polo sur. Con facilidad se demostraba que los polos homónimos se repelían y los de diferente signo se atraían.

Un estudio metódico del comportamiento de estos cuerpos peculiares que recibieron el nombre de imanes o cuerpos magnéticos lo realizó William Gilbert

(1544 - 1603). Se esclarecieron tanto las leyes generales de su comportamiento en los distintos puntos del globo terráqueo, como las reglas de su recíproca acción.

El 21 de julio de 1820 el físico danés Oersted publicó, haciéndole gran propaganda, su trabajo que llevaba un título bastante extraño: «*Experimenta circa effectum conflictus electrici in acum magneticam*» («*Experimentos referentes al efecto del conflicto eléctrico sobre la aguja magnética*»). Este pequeño escrito - tan sólo de cuatro páginas - daba a conocer al lector que Oersted (y, si queremos ser más exactos, un oyente de Oersted) prestó atención a que la aguja magnética se desviaba si ésta se disponía cerca del alambre por el cual circulaba la corriente eléctrica.

Inmediatamente tras este descubrimiento siguió otro. El relevante físico francés Andrés María Ampère (1775 - 1836) descubrió que las corrientes eléctricas interaccionan entre sí.

De este modo resulta que los imanes actúan sobre otros imanes y corrientes, mientras que las corrientes influyen en otras corrientes o imanes.

Para caracterizar estas interacciones al igual de las eléctricas, es conveniente introducir el concepto de campo. Digamos que las corrientes eléctricas, así como los imanes naturales y artificiales, engendran campos magnéticos.

Cabe subrayar que solamente por la investigación de los campos alternos se demuestra la existencia real de los campos eléctricos y magnéticos, o, en otras palabras, el hecho de que el campo es una forma de la materia. Entre tanto, el campo es para nosotros tan sólo un concepto cómodo, y nada más. En efecto, los manantiales del campo magnético pueden encontrarse ocultos detrás de un biombo, pero nosotros estamos en condiciones de juzgar sobre su presencia debido a las acciones que esto produce.

Los mismos sistemas que originan el campo magnético reaccionan a su presencia, es decir, el campo magnético actúa sobre las agujas magnéticas y las corrientes eléctricas. La primera tarea que se plantea ante el investigador que estudia el magnetismo es la «palpación» del espacio en el cual existe el campo magnético. Al caracterizar el campo magnético determinamos en cada punto del campo la fuerza que afectaba la carga unitaria. ¿Y cómo conviene proceder para describir el campo magnético?

En el caso general, el comportamiento de una pequeña aguja magnética es bastante complicado. Esta girará de modo determinado, pero, a veces, también realizará un movimiento de avance. Para poder caracterizar el campo magnético hay que impedir a la aguja que se desplace. En primer lugar es necesario poner en claro en qué dirección mira su polo norte (os decir, aquel extremo suyo que en ausencia de corriente y de objetos magnéticos mira en el sentido del Norte).

Con anterioridad, hemos señalado que un procedimiento gráfico adormido para describir el campo eléctrico es la introducción en uso de líneas de intensidad. La dirección de estas líneas indicaba adonde se desviaba la carga positiva. La densidad de las líneas correspondía al valor de la intensidad. De una manera análoga se puede proceder también al caracterizar el campo magnético. El extremo de la aguja magnética que gira libremente indicará la dirección de las líneas de fuerza del campo magnético las cuales, en la actualidad, se suelen llamarse líneas de inducción.

¿Y qué se debe tomar por la medida de «intensidad» del campo magnético? Por supuesto, se puede medir, empleando un simple dispositivo, el momento de fuerza que actúa sobre la aguja magnética. Sin embargo, tal vez, valga la pena buscar otro método. Es que la aguja magnética es una especie de «cosa en sí». Al realizar los experimentos con la aguja magnética, tenemos que buscar, simultáneamente, tanto la medida de «intensidad» del campo magnético, como la medida que caracteriza la aguja. Los físicos prefieren evitar tal situación. Mejor es matar primero un pájaro y sólo después el otro.

De este modo, por ahora conservemos para la aguja magnética la función de determinar el perfil de las líneas de inducción. Y para introducir la medida cuantitativa de «intensidad» del campo magnético recurramos a uno de los experimentos de Ampère, quien todavía en 1820 descubrió que un cuadro plano con corriente se comporta de una forma muy parecida a la aguja magnética. Resulta que el cuadro gira en el campo magnético, con la particularidad de que la normal a su plano mira en la misma dirección que la aguja magnética, es decir, a lo largo de las líneas de inducción. Hace las veces del polo norte aquel lado del cuadro en que, al observarlo, vemos la corriente ir en el sentido antihorario.

A diferencia de la aguja magnética el cuadro con corriente no es un objeto inconcebible para caracterizarlo. Las propiedades del cuadro con corriente, o sea, de un circuito con corriente, se definen unívocamente por la intensidad de la corriente, el área y la dirección de la normal al área. Es de suponer que este cuadro será un instrumento bastante bueno para la «palpación» del campo magnético.

Entonces, resulta que decidimos tomar por la medida de «intensidad» del campo magnético, el momento de rotación que actúa sobre el circuito con corriente. No conviene pensar que semejante instrumento es menos conveniente que la aguja magnética. Un experimentador hábil puede confeccionar un cuadro de área minúscula o inventará un sencillo método para compensar el giro que realiza el campo, recurriendo a la compresión de un muelle graduado.

Ante todo tenemos que averiguar el comportamiento de diferentes circuitos de prueba en un punto determinado del campo magnético invariable.

El resultado de esta investigación es el siguiente: el momento de fuerza es proporcional al producto de la intensidad de la corriente por el área. Esto significa que el circuito de prueba no se caracteriza por la intensidad de la corriente y el área tomados independientemente, sino por su producto.

Además de este producto necesitamos saber cómo está dispuesta la normal del circuito respecto a la dirección del campo. Es que el circuito se comporta a semejanza de una aguja magnética. Por esta razón, si disponemos el circuito de tal forma que su

normal positiva (o sea, el vector que sale por el lado norte) se oriente a lo largo de las líneas de inducción, el circuito quedará justamente en esta posición (el momento de fuerza es igual a cero) (fig. 3.1, abajo).

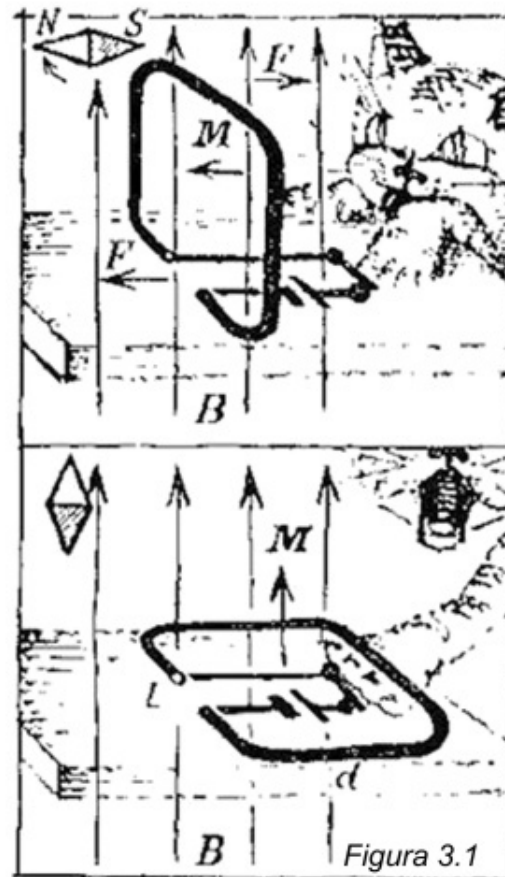


Figura 3.1

Si el circuito se sitúa de modo que su normal queda perpendicular a las líneas de inducción, el momento de fuerza será máximo (fig. 3.1, arriba).

De todo lo expuesto se infiere lo racional que es introducir un nuevo concepto. Un concepto muy importante, como lo comprenderemos más tarde. Vamos a caracterizar el circuito con corriente mediante el vector M al que damos el nombre de momento magnético (véase la fig. 3.1.).

La magnitud del momento magnético se toma igual al producto de la intensidad de la corriente I por el área del circuito $S = Id$:

$$M = IS$$

Al vector S se comunica la dirección de la normal positiva al plano del contorno.

De este modo poseemos la magnitud con cuya ayuda podemos medir el campo. Lo más conveniente es medir el momento máximo de fuerza que actúa sobre el circuito de prueba.

En el caso de pasar de un punto del campo al otro o modificar el campo a costa de desplazamiento de sus fuentes, o bien, cambiando las intensidades de las corrientes que crean el campo, obtendremos cada vez distintos valores del momento del par de fuerzas F que actúan sobre el circuito de prueba. El momento máximo del momento de fuerza lo podemos escribir así:

$$N = BM$$

donde B es una magnitud que, precisamente, tomamos por la medida del campo. Esta magnitud lleva el nombre de inducción magnética. A base de lo expuesto podemos decir que la inducción magnética es igual al momento máximo de fuerza que actúa sobre el circuito de prueba con el momento magnético unitario. Y la densidad de las líneas de inducción, es decir, su número que recae en una unidad de área, la tomaremos, justamente, proporcional a la magnitud B . El vector B está dirigido a lo largo de las líneas de inducción.

El momento magnético, la inducción magnética y el momento de fuerza que es nuestro viejo conocido, son vectores. Sin embargo, al recapacitar un poco

tendremos que convenir en que estos vectores se diferencian de los de desplazamiento, velocidad, aceleración, fuerza ... Efectivamente, el vector, digamos, de velocidad de movimiento indica en qué sentido se mueve el cuerpo; los vectores de aceleración y de fuerza señalan en qué dirección acciona la atracción o la repulsión. La flechilla con que terminamos el segmento, símbolo del vector, tiene en estos ejemplos un sentido totalmente objetivo y real. Pero, en lo que respecta a nuestros nuevos conocidos y al momento de fuerza, los asuntos toman otro cariz. Los vectores están dirigidos a lo largo del eje de rotación. Está claro que la flecha que corona uno u otro extremo del segmento el cual define el eje de rotación reviste un carácter completamente convencional. No obstante, es necesario ponerse de acuerdo en cuanto a la dirección del vector. La flecha puesta en el «extremo» del eje de rotación no posee un sentido. Pero la dirección de la rotación sí que tiene un sentido objetivo. Y esta dirección la tenemos que caracterizar. Se conviene en poner la flecha en el eje de rotación de tal modo que, al mirar en contra del vector, observar la rotación en el sentido horario, o bien, en el sentido antihorario. Los físicos se han acostumbrado a la segunda variante.

Estos dos tipos de vectores llevan nombres expresivos que hablan de por sí: vectores polares y axiales (del latín «axis», eje).

Las mediciones de los campos de diferentes sistemas nos llevan a las siguientes reglas. En los imanes siempre descubrimos dos polos: el polo norte del cual parten las líneas de fuerza, y el polo sur en que éstas terminan. Y lo que sucede con las líneas de inducción en el seno del imán, esto, naturalmente, no lo podemos determinar por vía experimental.

En lo que se refiere a los campos magnéticos de las corrientes (fig. 3.2). aquí se revela la siguiente regularidad: las líneas de inducción del campo magnético envuelven la corriente.

En este caso se debe tener presente que, si miramos a lo largo de la corriente, las líneas de inducción tendrán la dirección en que se mueve la aguja del reloj. El punto y la crucecita en los dibujos significan (y esto esta comúnmente admitirlo) que la corriente se dirige hacia nosotros o se aleja de nosotros, respectivamente.

El momento magnético, como resulta evidente de la fórmula, se mide en amperios multiplicados por metro cuadrado.

Hasta hace poco la unidad de inducción magnética ha sido el gaussio. Un gaussio es igual a 1 V-s/m^2 , sin embargo, puesto que el centímetro se ha expulsado, se ha propuesto otra unidad, la tesla (T): 1 T es igual a 1 V-s/m^2 .

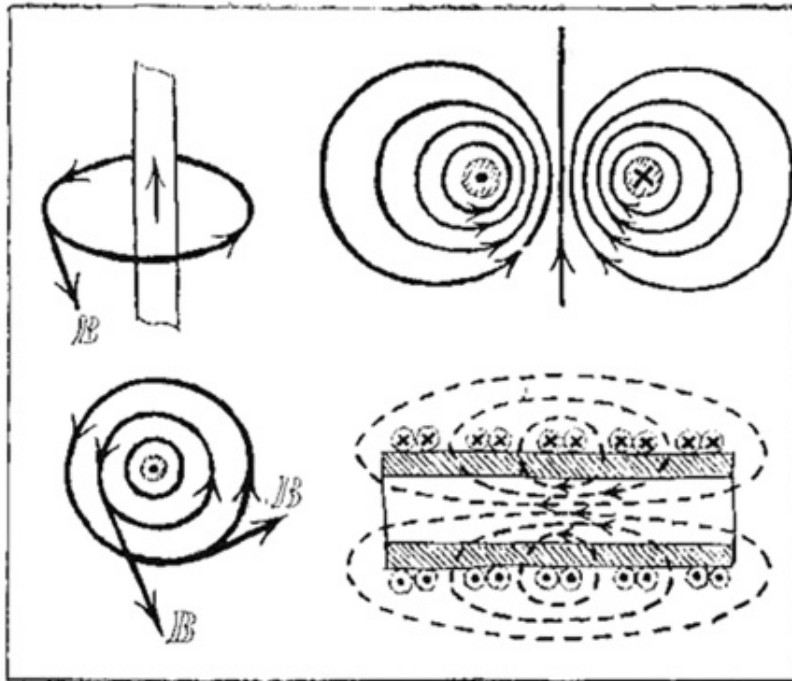


Figura 3.2

La procedencia de la dimensión queda completamente clara a partir de la fórmula para la fuerza electromotriz de inducción insertada en páginas anteriores.

Sin embargo, por ahora, los intentos de la Comisión Internacional de retirar del uso el gaussio sufren un fracaso: los campos magnéticos, como antes, se valoran por el número de gaussios.

Los campos magnéticos se engendran por las corrientes y los imanes permanentes. Los campos magnéticos, a su vez, ejercen influencia sobre las corrientes y los imanes permanentes. Si por cualquier causa el investigador no quiere recurrir al concepto de campo magnético, puede dividir todos los tipos de interacciones en que toman parte los campos magnéticos en cuatro grupos; magnéticas, o sea, las acciones del imán sobre el imán; electromagnéticas, es decir las acciones de las corrientes sobre el imán; magnetoeléctricas, es decir, las acciones del imán sobre la

corriente, y, por fin, electrodinámicas, es decir las acciones de la corriente sobre la corriente.

Esta terminología la utilizan, fundamentalmente, los técnicos. Por ejemplo, dicen que un instrumento es magnetoeléctrico cuando el imán resulta fijo y el cuadro con corriente es móvil.

Las interacciones electrodinámicas se han tomado por base de la definición actual de la unidad de intensidad de la corriente. Esta definición suena así: el amperio es la intensidad de una corriente invariable que, al pasar por dos conductores paralelos rectilíneos de infinita longitud y sección circular infinitamente pequeña dispuestos a una distancia de 1 m uno del otro en el vacío, engendra entre estos conductores una fuerza igual a 2×10^{-7} N por 1 m de longitud.

En el sistema SI adoptado actualmente por todo el mundo, la unidad de intensidad de la corriente es fundamental. El culombio, en correspondencia, se define como amperio-segundo. Tengo que confiar al lector que a mí me gusta más el sistema de unidades en el cual la unidad de cantidad de electricidad es fundamental y viene expresada por medio de la masa de plata depositada durante la electrólisis. Sin embargo, es difícil discutir con los metrologos. Por lo visto, la definición citada anteriormente tiene algunos méritos, aunque me parece que la medición práctica de la fuerza electrodinámica con alto grado de precisión es una tarea que dista mucho de ser simple.

Al conocer cómo determinar la dirección del campo magnético, así como las reglas para hallar la dirección de la fuerza que actúa sobre la corriente por parte del campo magnético (de lo que hablaremos un poco más tarde) el lector estará en condiciones de averiguar el mismo que las corrientes que fluyen paralelamente se atraen y las dirigidas en sentidos opuestos se repelen.

Acciones del campo magnético homogéneo

Es homogéneo aquel campo magnético cuya acción sobre cualesquiera indicadores del campo es idéntica en sus diferentes puntos.

Se logra crear semejante campo entre los polos de un imán. Es natural que cuanto más cerca uno del otro se dispongan los polos y cuanto mayor sea la superficie plana de las caras frontales del imán, tanto más homogéneo será el campo

Ya se ha hablado acerca de la acción del campo magnético homogéneo sobre la aguja magnética y el cuadro con corriente: si falta el muelle compensador, éstos se sitúan en el campo de una manera tal que su momento magnético coincida con la dirección del campo. El «polo norte» mirará al «polo sur» del imán. El mismo hecho puede expresarse con las siguientes palabras: el momento magnético se orientará a lo largo de las líneas de inducción del campo magnético.

Analicemos ahora la acción del campo magnético sobre las cargas en movimiento. Es sumamente fácil cerciorarse de que semejante acción existe y, además, es bastante imponente: es suficiente acercarse al rayo electrónico originado por el cañón electrónico, el más común y corriente imán escolar. El punto luminoso de la pantalla se despinzará y cambiará de lugar en esta pantalla en dependencia de la posición del imán.

De la demostración cualitativa del fenómeno se puede pasar a la investigación cualitativa, y, en este caso, resultará que la fuerza que actúa sobre el electrón por parte del campo magnético B , teniendo en cuenta que el electrón se mueve en dicho campo con la velocidad v y forma un ángulo recto con las líneas de inducción, es igual a

$$F = ceB$$

donde e es la carga de la partícula (desde luego, la ley es válida no sólo para los electrones, sino también para cualesquiera otras partículas cargadas).

En cambio, si la partícula se mueve a lo largo de las líneas de inducción, entonces, en efecto, el campo no actúa sobre ésta! Al lector que tiene nociones de trigonometría no es difícil atinar cómo escribir la expresión de la fuerza para el caso del movimiento bajo cierto ángulo respecto al campo. Y nosotros no recargaremos el texto con fórmulas que no necesitaremos posteriormente.

Pero, todavía no se ha dicho nada respecto a la dirección de la fuerza. Y ello es muy importante. La experiencia demuestra que la fuerza es perpendicular tanto a la dirección del movimiento de la partícula, como a la dirección de la inducción. O, en otras palabras: es perpendicular al plano que atraviesa los vectores v y B . Sin embargo, tampoco haciendo constancia de esto hecho lo hemos expresado todo.

Cada medalla tiene dos caras. ¿En qué consiste su diferencia? En la dirección del giro que hace coincidir un vector con el otro. Si vemos que el giro del vector v al vector B en un ángulo menor de 180° se opera contra el sentido de las agujas del reloj, entonces, esta cara la llamamos positiva.

Los elementales esquemas vectoriales representados en la fig. 3.3, a la izquierda, demuestran que una partícula cargada positivamente se desvía hacia el lado de la normal positiva. El electrón se desvía hacia el lado opuesto.

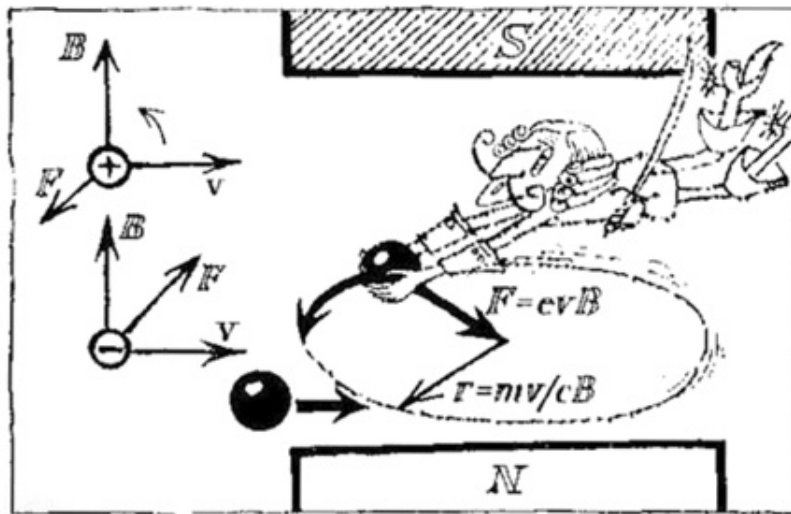


Figura 3.3

Ahora fijémonos a qué interesante resallado conduce esta ley en el caso del electrón que entró con pequeña velocidad formando un ángulo recto en el campo magnético permanente (fig. 3.3 a la derecha). Adivine, ¿qué trayectoria describirá el electrón? Ya lo tenemos por supuesto que se moverá describiendo una circunferencia. La fuerza del campo magnético es una fuerza centrípeta y podemos calcular inmediatamente el radio de la circunferencia igualando entre sí mv^2/r y evB . De este modo, el radio de la trayectoria es igual a

$$r = mv/eB$$

Presten atención al hecho de que por la conducta de la partícula podemos calcular sus propiedades. Sin embargo, otra vez topamos con la misma relación que se presentó cuando estudiarnos el movimiento de la partícula en el campo eléctrico. ¡No se logra determinar por separado la carga eléctrica y la masa de la partícula! También en este caso el experimento nos lleva al valor de la relación e/m .

De este modo, la partícula se mueve por una circunferencia si su velocidad está dirigida bajo un ángulo recto respecto al campo magnético; la partícula se mueve según una recta si su velocidad está dirigida a lo largo del campo magnético. ¿Y qué tenemos en el caso general? Su respuesta, claro está, ya la tiene preparada. La partícula se mueve siguiendo una línea helicoidal cuyo eje es la línea de inducción. Dicha línea helicoidal se compondrá de espiras arrolladas, espaciada o apretadamente, en dependencia del ángulo inicial de entrada del electrón en el campo magnético.

Por cuanto el campo magnético actúa sobre la partícula en movimiento, también debe ejercer su influjo en cada trocito de conductor por el cual fluye la corriente. Examinemos una «porción» del rayo electrónico de longitud l . Supongamos que en esta porción caben n partículas. La fuerza que actúa sobre un conductor de igual longitud por el cual fluye el mismo número de partículas con idéntica velocidad, esta fuerza será igual a $nevB$. La intensidad de la corriente es igual a la carga total que pasa a través del conductor en unidades de tiempo. El tiempo τ durante el cual los electrones examinados recorrerán el camino l es igual a

$$\tau = l/v.$$

Es decir, podemos anotar la intensidad de la corriente de la siguiente forma:

$$I = \frac{ne}{\tau} = \frac{nev}{l}$$

Al sustituir la velocidad

$$\frac{h}{me}$$

tomada de esta expresión en la fórmula para la fuerza que actúa sobre la «porción» del rayo electrónico, hallaremos precisamente la fuerza que ejerce su acción en el conductor de longitud l . He aquí esta expresión:

$$F = I l B$$

Esta expresión es válida sólo en el caso de que el conductor es perpendicular al campo.

La dirección de desviación del conductor por el cual circula la corriente puede determinarse con la ayuda del esquema representado en la fig. 3.3.

En señal de respeto a los investigadores que habían trabajado en el siglo XIX inserto la fig. 3.4.

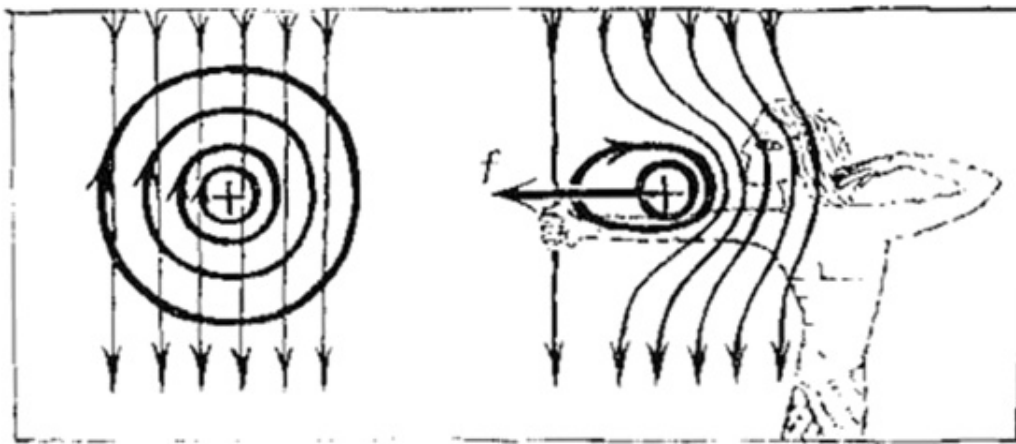


Figura 3.4

Este dibujo, desde luego, no sólo reviste interés académico. Ayuda a recordar la regla de desviación de las corrientes. La figura muestra cómo el campo propio de la corriente (que se dirige «desde nosotros») se sumará con el campo exterior. El resultado de dicha adición se representa a la derecha.

Si líneas de inducción se conciben como las tensiones de la materia del éter (semejante punto de vista se había difundido ampliamente en el siglo XIX), la dirección del desplazamiento del conductor obtiene una interpretación patente: el conductor, meramente, se expulsa por el campo.

Demostremos ahora que la acción del campo magnético sobre la carga en movimiento y un «segmento» de la corriente es el mismo fenómeno por el cual comenzamos el análisis de las acciones del campo magnético.

Examinemos otra vez la fig. 3.1, por arriba. En la figura se representan las fuerzas que actúan sobre el circuito con corriente. Las fuerzas no influyen sobre los tramos del conductor que van a lo largo de las líneas de inducción; sobre otros los tramos actúan el par de fuerzas y en la figura se ve que el momento de este par es igual, precisamente, al producto de la fuerza por el brazo:

$$N = I l B d = I S B = M B$$

De este modo, la expresión para el momento de fuerza como producto del momento magnético del circuito por la inducción magnética deriva directamente de la fórmula de la fuerza que actúa sobre la carga.

A propósito, la fórmula $F = evB$ con la cual comenzamos este párrafo lleva el nombre de Lorentz (Hendrick Antoon Lorentz, 1853 - 1928, físico holandés, propuso esta fórmula en 1805).

Acciones del campo magnético no homogéneo

No presenta ninguna dificultad crear un campo magnético no homogéneo. Por ejemplo, se puede dar a los polos del imán la forma curvilínea (fig. 3.5).

Entonces, el curso de las líneas de inducción del campo será tal como se muestra en la figura.

Supongamos que los polos están lo suficientemente separados uno del otro y coloquemos cerca de uno de éstos la aguja magnética. Como hemos mencionado de paso, en el caso general la aguja magnética no sólo gira sino también puede realizar el movimiento de avance.

Un movimiento giratorio, únicamente, de la aguja magnética (o del cuadro con corriente) se observa en el caso de que el campo sea homogéneo. Mientras que, en un campo no homogéneo tendrán lugar los dos movimientos.

La aguja girará de tal modo que se sitúe a lo largo de las líneas de inducción y, seguidamente, comenzará a atraerse al polo (véase la fig. 3.5). La aguja es atraída

a aquella zona donde el campo es más fuerte. (Desde luego, el pintor puso demasiado ahínco en su trabajo: difícilmente se puede creer que hasta un campo muy fuerte romperá la brújula.)

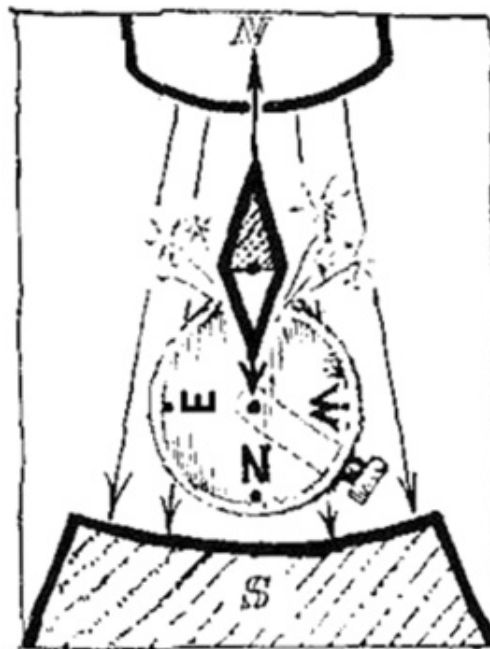


Figura 3.5

¿En qué reside la causa de semejante comportamiento?

Evidentemente, consiste en que en un campo no homogéneo sobre la aguja actúa no un solo par de fuerzas. Las «fuerzas» que inciden en el polo norte y el polo sur de la aguja situada en un campo no homogéneo no son iguales. Aquel extremo suyo que se encuentra en el campo más fuerte se somete a la acción de una fuerza mayor. Por esta razón, después del giro, el cuadro de las fuerzas presenta el aspecto mostrado en la figura: queda en exceso la fuerza actuante, en el sentido del campo más fuerte.

Verdad es que el comportamiento de un circuito con corriente de espesor ínfimo será absolutamente análogo. De este modo, cuando he comenzado por el modelo de la aguja con dos «polos» solamente he rendido tributo a la tendencia para la representación patente.

¿Cuál es, entonces, la ley de la naturaleza? ¿A qué es igual la fuerza? Un experimento y los cálculos demuestran que para cualquier sistema poseedor del

momento magnético M , dicha fuerza es igual al producto del momento del sistema por la curvatura del aumento del campo.

Sea que la aguja magnética se quedó situada a lo largo de las líneas de inducción. Los valores del campo en los puntos donde se encuentran el polo norte y el polo sur de la aguja magnética se diferencian entre sí. Tracemos el gráfico del campo a lo largo de la línea que pasa a través de los polos. Para mayor sencillez sustituyamos el tramo de la verdadera curva del campo entre los polos por una recta, lo que se puede hacer con tanta mayor precisión cuanto menor es la aguja, es decir, cuanto más próximos uno del otro son sus polos. La pendiente, o sea. la tangente del ángulo que esta recta forma en el gráfico con el eje horizontal, se expresará como el cociente de división de la diferencia de los campos por la longitud de la aguja. La fórmula tendrá el siguiente aspecto:

$$F = M \frac{(B_N - B_S)}{l}$$

donde l es la longitud de la aguja, y B_N y B_S representan la inducción del campo en los extremos norte y sur de la aguja. (No se asombren que la tangente del ángulo resulte ser una magnitud dimensional.)

Si en lugar de la fracción escrita ponemos el valor de la tangente del ángulo de la tangente a la curva que representa el curso del campo en el punto en que se encuentra la partícula que nos interesa, los «polos desaparecerán» y la fórmula será válida para cualquier partícula o sistema de partículas.

Resumiendo, podemos decir que en un campo no homogéneo el sistema o la partícula que poseen el momento magnético se atraen a los polos del imán o se repelen de éstos en dependencia de cómo está dirigido el momento magnético a lo largo o en contra de las líneas de inducción.

¿Acaso el momento magnético puede situarse en contra de la dirección del campo? ¡Sí, puede! Pero en qué casos, de ello hablaremos más tarde.

Corrientes amperianas

Hasta el siglo XIX no fue nada difícil crear teorías físicas. El cuerpo se ha calentado, esto significa que contiene mayor cantidad de calórico. Una medicina permite conciliar más rápidamente el sueño, por consiguiente, en esta se encierra una fuerza somnífica. Ciertas varillas fabricadas de menas de hierro señalan el norte. Un comportamiento raro, pero podemos comprenderlo inmediatamente si decimos que semejantes varillas y agujas poseen «alma» magnética. Como se conoce, los agujas magnéticas desde tiempos muy remotos sirvieron fielmente a los navegantes. Sin embargo, a veces, se encaprichaban. No es de extrañar, el asunto queda muy claro: la culpa la tienen los espíritus malignos! Tampoco es de extrañar que resultó posible imantar hierro, así como acero y algunas otras aleaciones. Sencillamente, éstos son cuerpos aptos para acoger fácilmente el «alma» magnética.

Después de los descubrimientos de Oersted y Ampère se puso de manifiesto que es posible tender un puente entre los fenómenos eléctricos y magnéticos. Había una época en la cual estuvieron en boga, con la misma amplitud, dos teorías. Según el primer punto de vista todo se dilucidaba al admitir que el conductor por el cual circula el fluido eléctrico se convierte en imán. Otro punto de vista lo mantenía Ampère. Este afirmaba que el «alma» magnética de las barras de hierro constaba de corrientes eléctricas microscópicas.

A muchos el punto de vista de Ampère parecía más lógico. Sin embargo, a esta teoría no se le atribuía ninguna importancia seria, no obstante en la primera mitad del siglo XIX apenas si había alguien quien pensara en la posibilidad de descubrir realmente estas corrientes, sin hablar ya de que se ponía en tela de juicio el hecho de que el mundo está construido de átomos y moléculas.

Solamente cuando en el siglo XX los físicos, con una serie de brillantes experimentos, demostraron que el mundo que nos rodea, en efecto, está construido de átomos y moléculas y los átomos constan de electrones y núcleos atómicos, se comenzó a creer en las corrientes amperianas como en un hecho real basándose en el cual se puede tratar de comprender las propiedades magnéticas de la sustancia. La mayoría de los científicos convino en que las «corrientes moleculares» imaginadas por Ampère se originan debido al movimiento de los electrones en torno a los núcleos atómicos.

Parecía que, valiéndose de estas ideas, se lograría explicar los fenómenos magnéticos. Efectivamente, el electrón que se mueve alrededor del núcleo puede asemejarse a la corriente eléctrica, tenemos derecho de atribuir a este sistema un momento magnético y enlazarlo con el momento de impulso de una partícula cargarla en movimiento.

Esta última afirmación se demuestra de una manera simplísima.

Supongamos que el electrón gira por una circunferencia de radio r . Puesto que la intensidad de la corriente es igual a la carga transportada en unidad de tiempo, resulta que el electrón que gira puede asemejarse a la corriente cuya intensidad es $I = Ne$, donde N es el número de revoluciones por segundo. La velocidad de la partícula puede relacionarse con el número de revoluciones mediante la expresión $v = 2\pi rN$; en consecuencia, la intensidad de la corriente es igual a

$$I = \frac{ve}{2\pi r}$$

Es lógico que el momento magnético del electrón que se mueve alrededor del núcleo se denomine orbital, este será igual a

$$M = IS = \frac{ve}{2\pi r} \pi r^2 = \frac{1}{2} e v r$$

Haciendo recordar al lector (véase el libro 1) que el momento de impulso de una partícula es igual a $L = mvr$, pondremos en claro que entre el momento de impulso y el momento magnético tiene lugar la siguiente relación, de suma importancia para la física atómica:

$$M = \frac{e}{2m} L$$

De aquí se infiere que los átomos deben poseer momentos magnéticos.

Valiéndose de diferentes procedimientos en que no nos detendremos es posible obtener el gas atómico de las más variadas sustancias. Por medio de dos ranuras en

la cámara de gas se originan haces de átomos neutros de hidrógeno, litio, berilio... Estos pueden dejarse pasar a través de un campo magnético no homogéneo y en la pantalla aparecerán las huellas del haz. El interrogante que planteamos ante la naturaleza consiste en lo siguiente: ¿se desviarán los flujos de átomos en el campo magnético de la vía recta, y, si lo hacen, entonces, de qué modo, precisamente?

El átomo posee el momento orbital y, en consecuencia, se comporta a semejanza de una aguja magnética. Si el momento magnético está dirigido a lo largo del campo, el átomo debe desviarse en el sentido del campo fuerte; en el caso de disposición antiparalela, debe desviarse a la zona del campo débil. El valor de la desviación puede calcularse por la fórmula similar a la expresión para la fuerza que actúa sobre la aguja magnética la cual insertamos anteriormente.

Lo primero que nos ocurre es que los momentos magnéticos de los átomos están dispuestos al azar. Y, siendo así, estamos esperando que el haz se ensanche.

No obstante, la experiencia aportó resultados completamente distintos. El haz de átomos nunca se ensancha, éste se desintegra en dos, tres, cuatro y más componentes, en dependencia de la clase de los átomos. La desintegración siempre es simétrica. En algunos casos entre los componentes del haz está presente el rayo no desviado, a veces, el rayo no desviado falta, y, finalmente, también ocurre que el haz no se desintegra en general.

De este experimento que, sin duda, es uno de los más importantes entre los realizados por los físicos en cualquier época se infiere, en primer término, que el movimiento de los electrones en torno al átomo se puede asemejar, en efecto, a la corriente eléctrica cerrada. Asemejar en un sentido estrecho y completamente determinado: al igual que a las corrientes cerradas a los átomos también puede atribuirse el momento magnético. Continuamos: los momentos magnéticos de los átomos pueden formar solamente ciertos ángulos discretos con la dirección del vector de la inducción magnética. En otras palabras, las proyecciones de los momentos magnéticos sobre esta dirección se cuantifican.

El hecho de que los datos habían sido vaticinados en todos los detalles resultó ser un gran triunfo de la física teórica. De la teoría se desprende que el momento de impulso y el momento magnético del electrón que deben su origen al movimiento de los electrones atómicos en el campo del núcleo (estos momentos se denominan

orbitales¹) son antiparalelos, y sus proyecciones sobre la dirección del campo pueden anotarse en la forma

$$L_z = m \frac{\hbar}{2\pi} \text{ y } M_z = m\mu_B$$

Aquí m es un número entero que puede tomar los valores 0, 1, 2, 3...; $\hbar/2\pi$ es el valor mínimo de la proyección del momento de impulso, y el valor mínimo de la proyección del momento magnético. Las magnitudes $\hbar$ y μ_B se hallan de los experimentos:

$$\hbar = 6,62 \times 10^{-34} \text{ J}\cdot\text{s}$$

$$\mu_B = 0,93 \times 10^{-25} \text{ J/T}$$

Cabe añadir, además, que estas magnitudes constantes de tanta importancia para la física llevan los nombres de los grandes científicos que colocaron los cimientos de la física cuántica: $\hbar$ se denomina constante de Planck. y μ_B , magnetón de Bohr.

Sin embarco, los postulados de la mecánica cuántica resultaron ser insuficientes para poder comprender el diferente carácter de la desintegración de los haces de los átomos de distintos elementos. Hasta los átomos más simples, los átomos de hidrógeno, se comportaban de una manera inesperada. Surgió la necesidad de añadir a las leyes de la mecánica cuántica una hipótesis de extraordinaria trascendencia, la cual ya hemos mencionado de paso. Hay que atribuir al electrón (y, como se averiguó más tarde, también a cualquier partícula elemental) el momento propio de impulso y, en correspondencia, el momento magnético propio (espín). Para comprender que es inevitable asemejar el electrón a la aguja magnética tenemos que, al principio, conocer con más detalle el carácter del movimiento de los electrones atómicos.

La nube electrónica del átomo

¹ Esta denominación se estableció debido a causas históricas, ya que la teoría del átomo comenzó con la hipótesis de que el átomo se parece al Sistema Solar.

Es imposible advertir el movimiento del electrón. Más aún, no se puede esperar que el progreso de la ciencia nos conduzca a que veremos el electrón. La causa es bastante clara. Para «ver» hay que «iluminar». Pero «iluminar» significa actuar sobre el electrón con la energía de un rayo cualquiera. Mientras tanto, el electrón es tan pequeño y posee una masa tan minúscula que toda intromisión con ayuda de un instrumento o aparato para la observación conducirá inevitablemente a que el electrón abandone el lugar en que se encontraba antes.

No solamente aquellos datos módicos acerca de la estructura de los átomos que vamos a comunicar ahora al lector, sino también toda la armoniosa doctrina referente a la estructura electrónica de la materia son fruto de teoría y no del experimento. No obstante, estamos seguros de su carácter fidedigno gracias a la infinita cantidad de resultados observados en el experimento que se deducen de la teoría recurriendo a razonamientos lógicos rigurosísimos. Restablecemos el cuadro de estructura electrónica que no se puede ver con el mismo grado de seguridad con que Sherlock Holmes, guiándose por las huellas dejadas por el criminal, reconstruía el cuadro del delito.

El propio hecho de que el cuadro de estructura electrónica se vaticina partiendo de las mismas leyes de la física cuántica que se establecen por otros experimentos es de por sí un gran manantial de confianza hacia la teoría.

Ya hemos hecho mención de que el número atómico (el número de orden) del elemento químico en la tabla de Mendeleiev no es sino la carga de su núcleo o, lo que es lo mismo el número de electrones que pertenecen al átomo neutro. El átomo de hidrógeno posee un solo electrón; el átomo de helio los tiene dos; el de litio, tres; el de berilio, cuatro, etc.

¿Cómo, en fin de cuentas, se mueven todos estos electrones? La respuesta a esta pregunta no es nada simple y sólo reviste un carácter aproximado. La complejidad del problema consiste en que los electrones interaccionan no solamente con el núcleo, sino también uno con otro.

Afortunadamente, a la repulsión mutua - fenómeno que impulsa los electrones a evitar encuentros entre ellos - a este fenómeno, en todo caso, corresponde un papel menos importante que al movimiento cuya existencia se debe a la interacción del electrón con el núcleo. Precisamente esta circunstancia - y sólo ésta - permite sacar

las conclusiones acerca del carácter del movimiento de los electrones en diferentes átomos.

La naturaleza concedió a cada electrón una zona espacial dentro de la cual éste se mueve. De acuerdo con la forma de estas zonas, los electrones se dividen en categorías que se designan con las letras latinas s, p, d y f.

El más simple es el «apartamento» del electrón s. Es una capa esférica. La teoría señala que con mayor frecuencia el electrón se halla en el centro de dicha capa. De este modo, resulta una simplificación burda hablar sobre la órbita circular de semejante electrón.

La zona del espacio en que deambula el electrón p es completamente distinta. Recuerda por su forma una haltera de gimnasia. Otras categorías de electrones tienen zonas de existencia todavía más complicadas.

Para cada uno de los átomos de la tabla de Mendeleiev, la teoría (en este caso ya atrayendo datos experimentales) puede indicar cuántos electrones de tal o cual clase contiene éste.

Surge la pregunta si esta distribución de los electrones de acuerdo con los tipos de movimiento guarda relación con su distribución por los niveles de energía K, L, M y... de la que hemos hablado en el capítulo anterior. Sí, guarda la más directa relación. La teoría y la experiencia demuestran que los electrones pertenecientes al nivel K pueden ser sólo del tipo s; los que pertenecen al nivel L, de los tipos s y p, los del nivel M, de los tipos s, p y d, etc.

No examinemos con mucho detalle la estructura electrónica de los átomos, limitándonos con la exposición de la estructura electrónica de los primeros cinco elementos de la tabla. Los átomos de hidrógeno, helio, litio y berilio tienen solamente electrones s. El átomo de boro posee cuatro electrones s y un electrón p. La simetría esférica de la zona del espacio en que viaja el electrón s pone en tela de juicio nuestros razonamientos acerca del momento magnético del átomo que contiene un solo electrón. En efecto, si el momento de impulso puede tomar valores idénticos y dirigidos con igual probabilidad en todos los sentidos, resulta que, en promedio, el momento rotacional y, por consiguiente, también el momento magnético de semejante sistema deben ser iguales a cero. A esta deducción natural

Llega también la física cuántica: los átomos que contienen solamente electrones no pueden poseer momento magnético.

Pero, si es así, entonces, los haces de átomos de los primeros cuatro elementos de la tabla de Mendeleiev no deben desviarse en el campo magnético no homogéneo. ¿Y cómo resulta en la realidad? En la realidad este pronóstico no se cumple para los átomos de hidrógeno y litio. Los haces de estos átomos se comportan de una forma sumamente extraña. En ambos casos el flujo de átomos se desdobra en dos componentes desviadas en direcciones contrarias y a iguales distancias respecto a la dirección inicial. ¿Es incomprensible?

Momentos magnéticos de las partículas

El espín del electrón hizo su aparición en las tablas en 1925. La necesidad de introducirlo entre el número de participantes en los acontecimientos desarrollados en el micromundo la demostraron Abraham Goudsmit y George Uhlenbeck. Al enunciar la hipótesis de que el electrón posee el propio momento de impulso, estos investigadores pusieron de manifiesto que se resolvían con naturalidad todos los desatinos que se habían acumulado para aquel período durante la interpretación de los espectros atómicos.

Un poco más tarde se realizaron experimentos para el desdoblamiento de los haces atómicos. Y cuando se averiguó que también aquí tan sólo basándose en el concepto de espín se lograba ofrecer una explicación exhaustiva a los hechos observados, únicamente entonces los físicos dieron crédito al espín.

Transcurrió un corto período y se puso al descubierto que el propio momento de impulso, o sea, espín, era una propiedad inherente no sólo al electrón, sino también a todas las partículas elementales.

Ya hemos mencionado que la denominación «espín» es testimonio de la tendencia natural a la representación palmaria. Por cuanto el momento de impulso entró en la física como característica de un sólido en rotación, resultó que, después de haberse percatado de que para salvar la ley de la conservación del momento de impulso, era necesario atribuir a las partículas elementales cierto valor del momento de impulso, muchos físicos, inmediatamente, recurrieron a un cuadro patente de rotación de la partícula alrededor de su eje. Esta ingenua representación no soporta crítica alguna:

Se puede hablar sobre la rotación de una partícula elemental alrededor de su eje con mayor derecho que razonar acerca de la rotación alrededor de su eje de un punto matemático.

Los adeptos de la representación patente, partiendo de ciertas razones indirectas, lograron evaluar las dimensiones del electrón, o, más exactamente, establecer que incluso en el caso de que este concepto es aplicable al electrón, el tamaño del mismo debe ser menor que una magnitud determinada. La magnitud del espín se conoce: insertamos su valor más adelante en este párrafo. Al suponer que el electrón tiene forma, se puede calcular con qué velocidad giran los «puntos de su superficie». Resulta que esta velocidad es mayor que la de la luz. La persistencia nos hubiera llevado a la necesidad de abandonar la teoría de la relatividad.

Probablemente, el argumento más mortífero contra la representación patente sea el hecho de que el neutrón que no porta sobre sí la carga eléctrica posee espín. Mas, ¿por qué razón este argumento es decisivo? Júzguenlo ustedes mismos.

Si la partícula se hubiera podido figurar en forma de esfera cargada, su rotación en torno al eje habría originado algo como la corriente amperiana. Pero si sucede que una partícula neutra tiene un momento de impulso, así como también un momento magnético (en el libro 4 diremos varias palabras acerca de estas propiedades del neutrón), no se puede ni hablar sobre una analogía con la corriente amperiana.

Por supuesto, no conviene tomar la postura de profeta vaticinando que nunca se logrará explicar el espín y el momento magnético de las partículas elementales partiendo de cierta ley más general que todavía no está descubierta (esta tarea se resuelve parcialmente por la teoría del brillante físico inglés Pablo Dirac, pero acerca de ésta no podemos dar al lector ni siquiera una noción general por ser demasiado abstracta). Sin embargo, hoy en día debemos considerar las «flechillas» que representan el momento de impulso y el momento magnético de una partícula como conceptos primarios (que no se reducen a algo más simple).

Hace unos cincuenta años la mayoría de los físicos sostenían el punto de vista de Einstein quien escribió: *«Toda teoría física debe ser tal que, además de cualesquiera cálculos, se la podía ilustrar con la ayuda de las más sencillas imágenes»*. Lamentablemente, la opinión del gran hombre resultó ser errónea. Y ya durante muchos años los físicos operan tranquilamente con teorías en las cuales figuran

magnitudes susceptibles de medirse, sin que se les pueda poner en correspondencia una imagen visual.

El electrón y otras partículas elementales no tienen «polos». En una serie de casos, con seguridad, hablamos de estas partículas como de puntuales, reconocemos que a las partículas elementales el concepto de forma es inaplicable, mas, a pesar de todo, tenemos que atribuir a las partículas dos propiedades vectoriales: el momento de impulso (el espín) y el momento magnético. Estos dos vectores siempre se sitúan a lo largo de una línea. En unos casos son paralelos, y en otros, antiparalelos.

La experiencia demuestra que las fórmulas generales para las proyecciones del momento de impulso y el momento magnético que aducimos en páginas anteriores son válidas también para los momentos propios. Todos los experimentos, tanto los espectrales, como los concernientes al desdoblamiento de los haces atómicos en un campo magnético no homogéneo se interpretan irreprochablemente si, para el electrón, al número m en la fórmula para la proyección del momento de impulso se le permite tomar dos valores: $\pm\frac{1}{2}$. En lo que se refiere a la fórmula para la proyección del momento magnético, aquí el número m puede tomar dos valores: ± 1 .

El espín del electrón tiene el valor numérico igual a $\frac{1}{2}$ ($h/2\pi$) y puede disponerse tan sólo en dos direcciones: a lo largo del campo y en contra del campo. En cuanto al momento magnético del electrón, éste, siguiendo al espín, también puede tener únicamente dos orientaciones en el campo, y su valor numérico es igual a un magnetón de Bohr.

Ahora pasemos a explicar los resultados de los experimentos con los haces atómicos. Demostremos cuán fácilmente se esclarecen todas las particularidades del desdoblamiento de los haces atómicos valiéndose del concepto de espín.

En efecto, ¿cómo se puede comprender el fenómeno de que los haces de los átomos de helio y de berilio no se desdoblan? Este fenómeno se comprende así. Los electrones de estos átomos carecen de momento orbital, y la causa de ello reside en que pertenecen a la «clase» s . Y en lo que respecta a los espines de los electrones éstos miran hacia lados opuestos. Hablando con propiedad, esta afirmación no deriva de ninguna parte, aunque, intuitivamente, se nos representa por completo

natural. El principio según el cual el par de electrones en el átomo se dispone de tal manera que las direcciones de sus espines sean opuestas se denomina principio de exclusión y lleva el nombre de Wolfgang Pauli (1900 - 1958) quien lo formuló.

¡Cuántas hipótesis...! Si, no son pocas. Pero todas ellas, en conjunto, forman el armonioso edificio de la física cuántica de la cual derivan tantas consecuencias que, ni asomo de duda queda, en cuanto a la justeza de la tesis de que al electrón se debe atribuir el espín, que el valor del número de espín debe tomarse igual a $1/2$ y que los espines del par de electrones se deben subordinar al principio de exclusión de Pauli; sí, no hay físico alguno que tenga la menor duda al respecto. La suma de estas hipótesis refleja la estructura del micromundo.

Volveremos a nuestros haces atómicos. Hemos explicado por qué no se desdoblan los haces de los átomos de helio y de berilio.

Bueno, ¿y cómo se comportan el hidrógeno y el litio?

El hidrógeno tiene un solo electrón. El momento orbital de este es igual a cero, por cuanto es un electrón s . La proyección del espín del electrón puede tomar tan sólo dos valores: $\pm 1/2$, es decir, el espín puede disponerse en contra o a lo largo de la dirección del campo magnético. Precisamente por esta razón el flujo de átomos se desintegrará en dos componentes. Lo mismo ocurrirá a los átomos de litio debido a que dos electrones «equilibrarán» sus espines, mientras que el tercero se comportará de la misma manera que el único electrón del átomo de hidrógeno.

Absolutamente análoga será la conducta de los átomos de otros elementos que contienen en la envoltura superior un electrón no apareado.

Tendría que mencionar, sin demostración, algunos otros teoremas que se demuestran en la física cuántica con el fin de explicar, para los átomos de otros elementos, la desintegración en un gran número de componentes. Teniendo en cuenta el hecho de que sólo los electrones s no poseen momento orbital y que el espín del electrón se manifestará únicamente en el caso de que el electrón se encuentre solitario en su nivel de energía, los físicos lograron explicar, en su totalidad el comportamiento de los flujos de átomos de todas las especies. Al estudiar este sugestivo capítulo de la física incluso el más empedernido escéptico se cerciorará de que todas las hipótesis admitirlas por la física cuántica son leyes generales de la naturaleza.

Temo que muchos lectores no queden satisfechos con estas frases. Está claro que tan sólo los experimentos referentes a la desviación de los haces atómicos en el campo magnético no homogéneo no bastan, de por sí para introducir el concepto tan «raro» como es el espín. Pero nuestro libro es demasiado chico para que yo pueda aducir la colosal cantidad de datos los cuales requieren conceder al espín la carta de naturaleza.

¿Qué vale, por ejemplo, el fenómeno de resonancia magnética que no tiene nada que ver con lo expuesto anteriormente? En este fenómeno las ondas radioeléctricas de diapasón centimétrico se absorben por la sustancia cuando éstas tienen que volver el espín. La energía de interacción entre el momento magnético del electrón y el campo magnético constante en que se coloca la sustancia en los experimentos de resonancia magnética y, por consiguiente, también la diferencia de dos energías (disposición paralela y disposición antiparalela) se calculan sin dificultad. Esta diferencia es igual al cuanto de la onda electromagnética que se absorbe. Determinamos con enorme precisión el valor de la frecuencia de onda a partir del experimento y nos convencemos de la absoluta coincidencia de este valor con aquel que hemos calculado al conocer la inducción del campo y el momento magnético del electrón.

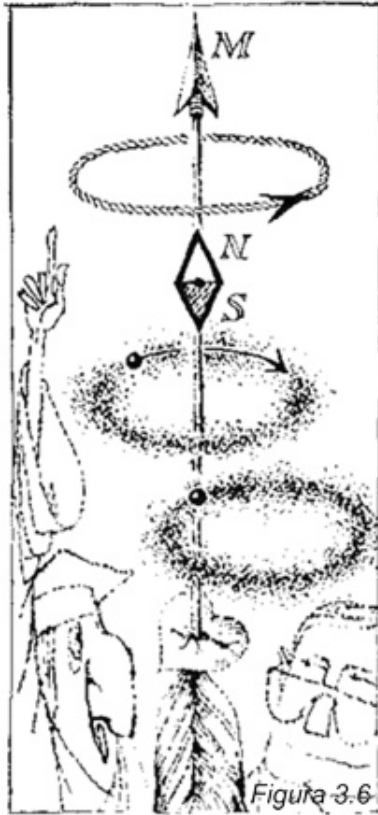
Es admirable que los mismos acontecimientos, pero, naturalmente en otro diapasón de las longitudes de onda se observen para los núcleos atómicos. La resonancia magnética nuclear es el método importantísimo de estudio de la estructura química de la sustancia.

Antes de seguir adelante sería, tal vez, útil hacer el balance de la totalidad de los hechos concernientes a los sistemas que crean el campo magnético y responden a la presencia del campo magnético.

Subrayemos otra vez que la hipótesis de Ampère resultó ser justa sólo parcialmente: los campos magnéticos son engendrados no sólo por las cargas eléctricas en movimiento. Otro manantial del campo magnético son las partículas elementales y, en primer lugar, los electrones que poseen el momento magnético propio. La clasificación técnica de interacciones dada en páginas anteriores resulta imperfecta. Los campos magnéticos se engendran por imanes naturales y artificiales, por corrientes eléctricas (incluyendo los flujos de partículas eléctricas en

el vacío) y por partículas elementales. Estos mismos sistemas, así como las partículas están sujetos a la acción del campo magnético.

La magnitud principal que caracteriza el campo magnético y sus acciones es el



vector del momento magnético, en el caso de las corrientes este vector viene determinado por la forma del circuito con corriente. El momento de la aguja guarda una relación compleja con la estructura atómica de la materia, mas no es difícil medirlo. Los electrones que se mueven en el campo del núcleo poseen el momento magnético «orbital» como si (fíjense, por favor, en este «como si») su movimiento alrededor del núcleo crease la corriente eléctrica. Y, finalmente, el momento magnético propio es una propiedad primaria que caracteriza las partículas elementales.

Para que estos datos fundamentales se graben mejor en su memoria se inserta la fig. 3.6. Esta figura es el balance de nuestros conocimientos de hoy sobre el «alma» magnética, o bien, si quiere, sobre el corazón

magnético. (Es que en francés la palabra «aimant», imán, significa, al mismo tiempo, «amante».) Desafortunadamente, no es sino un juego de palabras y, en la realidad, es un vocablo de procedencia griega que tiene parentesco con la palabra «diamante». Examinen este dibujo con atención, desde arriba hacia abajo. Su objetivo es ayudar a entender que la corriente macroscópica, el imán en forma de barra, el movimiento orbital del electrón y el propio electrón, todos ellos se caracterizan por un mismo concepto físico: el momento magnético.

Inducción electromagnética

La experiencia demuestra que el haz de electrones que se mueven en un campo magnético se desvía del trayecto rectilíneo. Como hemos dicho anteriormente la fuerza en cuestión que recibió el nombre de fuerza de Lorentz está dirigida perpendicularmente a las líneas de inducción y al vector de velocidad de los

electrones. Su magnitud se determina por la fórmula $F = evB$. Es la más simple expresión de la fuerza de Lorentz válida para el caso en que la velocidad de los electrones y la dirección del campo magnético forman un ángulo recto.

Si a esto hecho se añade nuestra seguridad de que en un conductor metálico se contienen electrones libres, entonces, por medio de sencillos razonamientos llegamos a la conclusión de que con ciertos movimientos de los conductores en el campo magnético en éstos debe aparecer la corriente eléctrica.

Dicho fenómeno que, podemos decirlo, forma la base de toda la técnica moderna, lleva el nombre de inducción electromagnética. Ahora vamos a deducir la ley a la que aquélla se subordina.



Figura 3.7

En la fig. 3.7 se representa un circuito conductor que no es sino una barra metálica AC de longitud l que se desliza por los cables metálicos y puede desplazarse entre los polos del imán sin alterar el carácter cerrado del circuito. Si la barra se desplaza perpendicularmente a las líneas de inducción del campo magnético entonces, sobre los electrones del conductor actuará la fuerza y por el circuito fluirá la corriente eléctrica.

Llegamos a la conclusión cuya importancia es imposible sobreestimar: la corriente eléctrica puede surgir en un conductor cerrado aunque al circuito no está conectado un acumulador u otro manantial de corriente.

Calculemos la fuerza electromotriz, es decir, el trabajo necesario para trasladar una unidad de carga a lo largo del circuito cerrado. El trabajo es igual el producto de la fuerza por el recorrido y se efectúa tan sólo en el tramo que se desplaza en el campo. La longitud del recorrido es igual a l y la fuerza por unidad de carga es igual a vB .

La fuerza electromotriz aparecida se denomina f.e.m. de inducción. Su valor se determina por la fórmula

$$\xi^{\text{ind}} = vBl$$

Es deseable generalizar esta fórmula de modo que ésta sea apropiada para cualquier movimiento de cualesquiera circuitos conductores. Llegaremos a esta generalización de la siguiente manera. En el tiempo τ la barra conductora se ha desplazado en la longitud x , siendo la velocidad de movimiento v igual a x/τ . El área del circuito conductor ha disminuido en la magnitud $S = xl$. La expresión de la f.e.m. de inducción adquiere la forma

$$\xi^{\text{ind}} = BS/\tau$$

¿Pero, cuál es el sentido del numerador de la fórmula? Este es completamente evidente: BS es la magnitud en que cambió el flujo magnético (el número de líneas de inducción) que atraviesa el circuito.

Por supuesto, nuestra demostración se da para un caso muy sencillo. Pero el lector me debe creer de palabra que esta demostración puede llevarse a cabo con absoluta rigurosidad para cualquier ejemplo. La fórmula obtenida tiene el valor más general y la ley de la inducción electromagnética se formula así: la f.e.m. de inducción aparece siempre en el caso de que varía el número de líneas de inducción que atraviesa el circuito. En este caso el valor de la f.e.m. de inducción es numéricamente igual a la variación del flujo magnético por unidad de tiempo.

Existen también tales desplazamientos del circuito en el campo magnético para los cuales la corriente no aparece. No habrá corriente si el circuito se mueve en un campo homogéneo paralelamente a las líneas de inducción. En cambio, si el circuito gira en un campo magnético homogéneo, entonces, la corriente se engendra. La corriente también aparecerá si el circuito se acerca o se aleja respecto al polo de un imán en forma de barra.

La experiencia demuestra, además, que la generalización que hemos formulado es todavía más significativa. Hasta ahora se trataba de los casos en que el circuito de corriente y el manantial del campo magnético cambiaban su disposición recíproca. La última fórmula que hemos deducido no dice nada acerca del movimiento. Esta versa tan sólo sobre la variación del flujo magnético. Sin embargo, es que la variación del flujo magnético a través de un circuito conductor no requiere obligatoriamente una traslación.

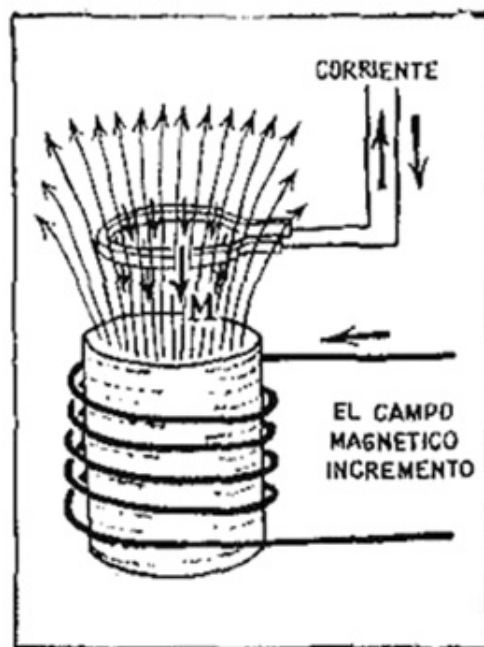


Figura 3.8

Efectivamente, se puede tomar como manantial del campo magnético no un imán permanente, sino un carrete dejando pasar por éste corriente eléctrica proveniente de cualquier fuente extrínseca. Valiéndose de un reóstato o empleando cualquier otro procedimiento es posible variar la intensidad de la corriente en este carrete

primario que es la fuente del campo magnético. En este caso, el flujo magnético que atraviesa el circuito cambiará, siendo invariable la disposición de la fuente del campo magnético (fig. 3.8.).

En estas circunstancias, ¿será también válida nuestra generalización? La experiencia da respuesta a esta pregunta. Y la respuesta resulta positiva. Independientemente de la forma en que cambia el número de las líneas de inducción la fórmula de la f.e.m. sigue en pie.

Dirección de la corriente de inducción

Ahora vamos a señalar que existe una sencilla regla universal concerniente a la dirección de las corrientes de inducción que se engendran. Examinemos varios ejemplos para sacar luego una conclusión general.

Al volver a la fig. 3.7 prestemos atención al siguiente hecho. Si disminuimos el área del circuito el flujo magnético que lo atraviesa también disminuye. La dirección de la corriente mostrada en la figura es tal que el momento magnético de la corriente aparecida está orientado a lo largo de las líneas de inducción. Este hecho significa que el campo propio de la corriente inducida está dirigido de tal modo que tiende a «impedir» la reducción del campo magnético.

Llegaremos a la misma conclusión también para el caso inverso. Si el área del circuito aumenta, también aumentará el flujo que atraviesa el circuito. Pero ahora el momento magnético del circuito mirará contra las líneas de inducción. Es decir, otra vez el campo de la corriente de inducción engendrada obstaculiza la acción que la llamó a la vida.

Otro ejemplo. Supongamos que nuestro circuito se sitúa entre los polos del imán de tal modo que el flujo que lo atraviesa es igual a cero. Comencemos a girar el circuito en el sentido de las agujas del reloj, así como en el sentido contrario. Los dos casos se ilustran en la fig. 3.9.

La línea llena designa la proyección del circuito en la posición inicial, y la línea de trazos representa las proyecciones del circuito en la posición de giro, cuando apareció la corriente eléctrica. Utilizando el esquema vectorial dado en la fig. 3.3, a la izquierda, hallamos la dirección de las corrientes de inducción que se originan en ambos casos.

Durante el giro en el sentido de las agujas del reloj el momento magnético de la corriente de inducción mira hacia abajo, y al girar en el sentido antihorario, mira hacia arriba.

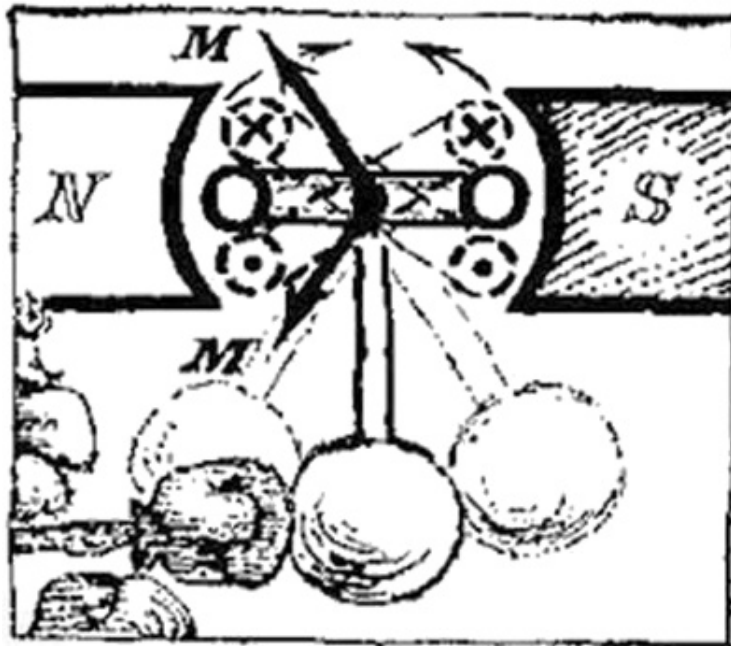


Figura 3.9

A medida que aumenta el ángulo de giro el campo magnético propio del circuito hace disminuir cada vez en mayor grado el campo que fue la causa de la inducción. De este modo vamos que también aquí funciona la misma regla.

Ahora veamos cómo se comportarán los circuitos en los campos no homogéneos. Retornemos a la fig. 3.8. Supongamos que la intensidad de la corriente en la bobina es invariable y analicemos qué sucederá durante el desplazamiento del circuito. Si el circuito se acerca al polo norte, el momento magnético mirará en contra de las líneas de fuerza. Si el circuito se alejase, el campo propio de la corriente inducida intensificaría a éste. Semejante comportamiento puede demostrarse valiéndose del mismo esquema vectorial.

¿Y cómo se desarrollarán los acontecimientos en el caso de campos magnéticos originados por corrientes alternas? El aumento o la disminución de la intensidad de la corriente en la bobina primaria, da lugar a la variación del flujo de las líneas de

inducción del campo magnético. En el circuito (véase otra vez la fig. 3.8) aparece la f.e.m.

¿Y de qué manera puede determinarse la dirección de la corriente? En este caso ya no se puede hacer uso del esquema vectorial, por cuanto no hay movimiento. Es aquí, precisamente, donde nos servirá nuestra generalización. Resulta que también en esta ocasión la dirección de la corriente de inducción engendrada debido a la disminución o el aumento del número de líneas de inducción del campo magnético que atraviesan el circuito se subordinará a aquella misma regla: la corriente de inducción engendrará un campo que parece como si compensara la variación del campo magnético que fue la causa de la inducción.

Algunas palabras acerca de la historia del descubrimiento de la ley de la inducción electromagnética

El descubrimiento del fenómeno de la inducción electromagnética se incluye entre aquellos acontecimientos que se pueden contar con los dedos y que ejercieron una influencia decisiva en el progreso de la humanidad. Esta es la razón de que sería imperdonable no detenernos en la historia de este descubrimiento. El mismo fue hecho mucho antes de que se investigó el comportamiento del haz de electrones en el campo magnético, y el curso histórico de los acontecimientos no coincide, en modo alguno, con el orden de exposición que hemos elegido en el párrafo anterior: la lógica y la secuencia del pensamiento no están obligadas, ni mucho menos, a ir en paralelo al desenvolvimiento histórico de los sucesos.

Para el momento en que Faraday comenzó sus experimentos que llevaron al descubrimiento de la inducción electromagnética la situación en el estudio teórico de los campos eléctricos y magnéticos tomó el siguiente cariz.

La obtención de la corriente continua y las leyes generales de su comportamiento en los circuitos eléctricos ya no representaban para los físicos serios problemas. Se estableció ya la acción de la corriente en el imán permanente y la interacción de las corrientes entre sí. Quedó claro que la corriente continua engendra a su alrededor un campo magnético que puede medirse tanto con la ayuda de un imán, como haciendo uso de otra corriente. Entonces, cabía preguntar: ¿y si existe un fenómeno inverso? ¿no es que el campo magnético engendra corriente en un conductor?

En 1821 Faraday deja en su diario el siguiente apunte: «Convertir el magnetismo en electricidad». El gran sabio necesitó un lapso de diez años para lograr el éxito.



Miguel Faraday (1791 - 1867), gran físico inglés. Le pertenece el descubrimiento del fenómeno de la inducción electromagnética (1831). Este descubrimiento Faraday lo hizo no por casualidad, lo buscó. Las leyes de la inducción electromagnética de Faraday forman la base de la electrotecnia. Es difícil sobrestimar el valor de las leyes de la electrólisis, establecidas por Faraday. El gran hombre de ciencia introdujo en uso y dilucidó los términos tan comunes hoy en día como ánodo, cátodo, anión, catión, ion y electrólito. Faraday demostró que el medio influye en la interacción eléctrica. No se puede dejar de mencionar el descubrimiento de la rotación magnética del plano de polarización. El hecho de que todos los cuerpos pertenecen ya sea a los paramagnéticos, o bien, a los diamagnéticos también lo estableció Faraday. El mundo no conoció otro físico experimentador tan grande como fue Faraday.

Los muchos años de mala suerte encontraron la explicación en el hecho de que Faraday trataba de obtener la corriente colocando el conductor en un campo permanente. En 1831 los esfuerzos tenaces del científico se coronaron con éxito. El fragmento que citamos a continuación tomado de un artículo de Faraday perteneciente al año 1831 representa la primera descripción del fenómeno descubierto:

«Sobre un ancho carrete de madera fue arrollado alambre de cobre de 203 pies de longitud y entre sus espiras se arrolló alambre de igual longitud aislado del primero por hilo de algodón. Una de estas espirales estaba conectada al galvanómetro y la otra con una potente pila... Al cerrar el circuito se lograba ver una acción súbita, pero bastante débil sobre el galvanómetro, y lo mismo se observaba al interrumpirse la corriente. En cambio, durante el paso ininterrumpido de la corriente a través de una de las espirales no se ponía de manifiesto una acción sobre el galvanómetro, ni, en general, acción alguna sobre otra espiral, a pesar de que el calentamiento de toda la espiral conectada a la pila, así como la brillantez de la chispa que saltaba entre los carbones eran testimonio de la potencia de la pila».

El descubrimiento del fenómeno de la inducción electromagnética fue la primera etapa de la labor de Faraday que duró cuatro lustros y cuya finalidad era hallar la ligazón única entre todos los fenómenos eléctricos y magnéticos.

Sin embargo, al hablar sobre la inducción electromagnética es necesario hacer mención también de los nombres de otros físicos relevantes. Al norteamericano Joseph Henry (1797 - 1878) pertenece el descubrimiento del fenómeno de autoinducción. Si la corriente que fluye por el carrete varía, varía también el campo magnético producido por esta corriente, varía el flujo del campo que pasa a través del propio carrete y se induce la f.e.m. en su «propio» circuito.

¿Y quién fue el que descubrió la ley de la dirección de la f.e.m. de inducción? La más completa respuesta a esta pregunta se puede encontrar en los trabajos de Lenz. La regla de Lenz determina la dirección de la corriente de inducción: «Si un conductor metálico se desplaza en la proximidad de una corriente o de un imán, entonces, en este conductor se engendra la corriente galvánica. La dirección de dicha corriente es tal que un alambre en reposo, por efecto de la misma, se hubiera puesto en movimiento directamente opuesto a la traslación real. Se supone que el alambre puede desplazarlo en dirección del movimiento real o en el sentido directamente contrario».

Después de 1840, paulatinamente, paso a paso, se crea el cuadro único del electromagnetismo. El descubrimiento de las ondas electromagnéticas es el último y, quizás, el más brillante trazo en este cuadro.

Corrientes de inducción en torbellino

Por cuanto las corrientes de inducción pueden surgir en conductores de alambre es también bastante lógico que éstas aparezcan en trozos de metal macizos y hechos de una pieza. Cada trozo de metal contiene electrones libres. Si el metal se desplaza en un campo magnético permanente, entonces, sobre los electrones libres actuará la fuerza de Lorentz. Los electrones describirán trayectorias circulares, es decir, engendrarán corrientes en torbellino. Este fenómeno lo descubrió en 1855 el físico francés León Foucault (1819-1868).

Las leyes de la inducción electromagnética son igualmente justas en tanto en el caso de que el flujo magnético varíe debido al desplazamiento relativo del metal y la fuente del campo, como en el caso de que la variación del campo magnético ocurra a raíz de la variación de la corriente eléctrica que origina el campo. Por esta razón, las corrientes de Foucault se manifiestan no sólo cuando tiene lugar el movimiento relativo (el experimento más brillante es el con una moneda a la que se hace caer entre los polos de un fuerte imán; la moneda no cae con una aceleración ordinaria, sino de tal forma como si la caída tuviese lugar en un aceite viscoso: el sentido del experimento es evidente: en la moneda aparecen las corrientes de Foucault cuya dirección, de acuerdo con la regla de Lenz es tal que su interacción con el campo magnético primario frena aquel movimiento que origina la inducción), sino también cuando el campo magnético varía en el tiempo.

Entre las aplicaciones útiles de las corrientes de Foucault pueden mencionarse las siguientes. En primer lugar, las corrientes de Foucault se utilizan en los llamados hemos de inducción para un fuerte calentamiento y hasta la fundición de los metales. En segundo lugar, en muchos instrumentos de medida estas corrientes proporcionan la «atenuación magnética».

Un invento ingenioso (es ya en tercer lugar) es el contador de la energía eléctrica. Por supuesto, usted ha visto que su parte principal es un disco giratorio. Cuantas

más bombillas u hornillos eléctricos se conectan con tanta mayor velocidad gira el disco.

El principio de la estructura del contador consiste en que se crean dos corrientes. Una de éstas fluye por el circuito paralelo a la carga, y la otra, en el circuito de la corriente de carga. Estas dos corrientes circulan por los carretes asegurados sobre núcleos de hierro y se denominan precisamente así: «carrete de Volta» y «carrete de Ampère». La corriente alterna imanta los núcleos de hierro. Puesto que la corriente es alterna los polos de los imanes eléctricos cada vez se alternan. Entre éstos parece como si corriese el campo magnético. Los carretes se sitúan de modo que el campo magnético en movimiento producido por los dos, forme en el seno del disco corrientes en torbellino. La dirección de estas corrientes en torbellino será tal que el campo magnético en movimiento arrastra en pos suyo el disco.

La velocidad de rotación del disco dependerá de los valores de la corriente en los dos carretes. Esta velocidad, como puede demostrarse por medio de cálculos exactos, será proporcional al producto de la intensidad de la corriente por la tensión y por el coseno del desfase, en otras palabras, a la potencia consumida. No nos detendremos en los sencillos procedimientos mecánicos que permiten enlazar el disco en rotación con el indicador de cifras.

Sin embargo, en la mayoría de los casos se trata de deshacerse de las corrientes de Foucault. Es uno de los problemas de desvelo de los diseñadores de las máquinas eléctricas de cualquier tipo. Al igual que todas las corrientes las de Foucault absorben la energía del sistema. En este caso, además, las pérdidas pueden ser tan considerables que surge la necesidad de recurrir a toda clase de artimañas. El método más simple de lucha contra las corrientes de Foucault es la sustitución en las máquinas eléctricas de los trozos macizos de metal por material en chapas. En este caso las corrientes parásitas no tienen dónde «desplegarse», su intensidad se reduce considerablemente y, en correspondencia, disminuyen las pérdidas en calor. No hay duda de que el lector se fijaba en el calentamiento de los transformadores. Y este calentamiento por lo menos al cincuenta por ciento se debe a las corrientes en torbellino.

Choque de inducción

Valiéndose del fenómeno de inducción electromagnética se pueden elaborar métodos sumamente perfectos de medición del campo magnético. Hasta el momento recomendamos utilizar para este fin la aguja magnética o un circuito magnético de prueba por el cual fluye la corriente eléctrica continua de intensidad conocida. La inducción magnética se determinaba por el momento de fuerza la cual actúa sobre el circuito de prueba o la aguja cuyo momento magnético es igual a la unidad.

Ahora procedamos de otra manera. Conectemos a un instrumento de medición una diminuta espira del conductor bajo corriente. Coloquemos la espira en la posición perpendicular a las líneas de inducción y, seguidamente, con un movimiento rápido le damos vuelta a 90°. Durante el tiempo de giro por la bobina fluirá la corriente de inducción y pasará una cantidad completamente determinada de electricidad Q que puede medirse.

¿De qué modo esta cantidad de electricidad estará relacionada con la magnitud del campo en el punto en que colocamos la espira de prueba?

El cálculo es bastante sencillo. De acuerdo con la ley de Ohm la intensidad de la corriente I es igual al cociente de la división de la f.e.m. de inducción por la resistencia,

o sea

$$I = \frac{1}{R} \xi^{\text{ind}}$$

Si hacemos uso de la expresión para la ley de la inducción $\xi^{\text{ind}} = BS/\tau$ y tomamos en consideración que $Q = I\tau$ (I se toma como invariable), la inducción magnética resultará igual a

$$B = \frac{\xi^{\text{ind}} \tau}{S} = \frac{I \tau R}{S} = \frac{QR}{S}$$

Volvemos a repetir que esta fórmula, por supuesto, es válida en el caso de que en la posición final la línea de inducción no atraviese la espira y en la posición inicial

corta el área de la espira formando un ángulo recto. Desde luego, es completamente indiferente que resulta ser la posición final y qué la inicial. Debido al cambio sólo variará en dirección de la corriente, pero no la cantidad de electricidad que pasó a través del circuito.

La sensibilidad de este método de medición aumentará n veces si en lugar de una espira se toma una bobina. La cantidad de electricidad será proporcional al número de espiras n . Los experimentadores hábiles se las ingenian para hacer bobinas de un milímetro de tamaño, de modo que, empleando el método de choque de inducción, se puede sondear el campo con bastante detalle.

Sin embargo, la mayor significación este método la tiene, probablemente, en la medición de la permeabilidad magnética de los cuerpos de hierro. Ahora hablaremos sobre esta importante propiedad del hierro.

Permeabilidad magnética del hierro

Como hemos esclarecido en el capítulo anterior los átomos acusan propiedades magnéticas. Los electrones solitarios poseen el momento magnético y el movimiento de los electrones alrededor del núcleo origina momentos magnéticos orbitales. Los núcleos de los átomos tienen momentos magnéticos. Por esta razón, la introducción de un cuerpo en el campo magnético debe afectar el aspecto del campo y, por el contrario, la presencia del campo magnético influirá en uno u otro grado sobre el comportamiento de las sustancias sólidas, líquidos y gaseosas.

Unas propiedades magnéticas extraordinariamente relevantes las acusan el hierro, algunas de sus aleaciones y ciertas sustancias afines al hierro. Esta pequeña clase de sustancias lleva el nombre de ferromagnéticos. Pueden realizarse, por ejemplo, los siguientes experimentos: colgar en un hilo pequeñas varillas, de tamaño de un fósforo, y acercarlos un imán. Cualesquiera que sean otras materias utilizadas para hacer las varillas: madera, vidrio, plásticos, cobre, aluminio..., el acercamiento a éstas de un imán no ayudará a revelar las propiedades magnéticas de estas sustancias. Para demostrar la existencia de propiedades magnéticas en cualesquiera sustancias deben llevarse a cabo unos experimentos muy finos y escrupulosos, de los que hablaremos más tarde.

Pero los cuerpecitos de hierro se comportarán de una forma completamente diferente. Se moverán, obedientes, tras el más débil imán escolar en forma de barra.

Para que el lector pueda juzgar cuán sensibles son los cuerpos constituidos por hierro frente a la presencia del campo magnético, voy a contar una historia aleccionadora en todos los sentidos cuyo protagonista fue yo mismo.

Hace algunos años me pidieron que tomase conocimiento de los experimentos de un «taumaturgo» checo que se granjeó la fama mundial y a quien los gacetilleros norteamericanos ávidos de noticias sensacionales llamaron «Merlín checo», este individuo tenía en su repertorio varias decenas de experimentos los cuales, supuestamente, no se podían explicar de modo racional. Entre tanto, el propio Merlín checo atribuía los resultados de estos experimentos a su fuerza de sugestión. Uno de sus números cumbre fue la magnetización de una cerilla de madera. Primeramente, mostraba que la cerilla de madera suspendida de un hilo no se desviaba por el imán. Acto seguido, comenzaba a «hipnotizar» la cerilla haciendo ciertos pases misteriosos. Un elemento indispensable de este espectáculo era poner la cerilla en contacto con un «ídolo» metálico que, como explicaba Merlín, servía de receptor de su energía psíquica.

Dedicando un par de semanas a este asunto demostré que todos los experimentos, sin excepción, del mago checo tenían una explicación racional. Pero ¿cómo conseguía imantar la cerilla? ¿Cómo lograba que después de todos sus pases otra vez colgada al misino hilo, la cerilla comenzaba a seguir obedientemente al imán?

El asunto residía en lo siguiente. Al entrar en contacto con el «ídolo» metálico, al extremo de la cerilla pasaba una cantidad ínfima de polvo de hierro. Demostré que bastaba una treinta millonésima parte de gramo de hierro para que la cerilla adquiriese propiedades magnéticas notables. He aquí el segundo caso de «experimentos con cucarachas».

Este ejemplo muestra con bastante claridad que, en primer término, no se debe dar crédito a los «milagros» que contradicen las leyes de la naturaleza, y que, en segundo término - y es precisamente esto hecho el que nos interesa - las propiedades magnéticas del hierro son completamente singulares.

El experimento clásico que permite definir las propiedades magnéticas del hierro se lleva a cabo de la siguiente manera. Se compone un circuito eléctrico a base de dos bobinas puestas una sobre otra. La bobina primaria está conectada al circuito del acumulador, y la secundaria se conecta al aparato que mide la cantidad de electricidad. Si se cierra el circuito primario, el flujo magnético a través de la bobina secundaria cambiará desde el valor cero hasta cierto valor límite Φ_0 . El flujo puede medirse con gran precisión empleando el método de choque de inducción.

El estudio de las propiedades magnéticas de la sustancia se realiza, precisamente, con la ayuda de la instalación cuya descripción acabamos de presentar. Se prepara una barra que se introduce dentro de la bobina. Se comparan los resultados de dos mediciones: sin la barra y con la barra. En el caso de que la barra está confeccionada de hierro o de otros materiales ferromagnéticos la cantidad de electricidad medida con el aparato aumenta varios miles de veces.

Como característica de las propiedades magnéticas del material puede tomarse la relación entre los flujos magnéticos medidos en presencia de la barra y cuando ésta falta. Dicha relación, o sea, $\mu = \Phi/\Phi_0$ lleva el nombre de permeabilidad magnética de la sustancia.

Resumiendo, podemos decir: el cuerpo de hierro aumenta bruscamente el flujo de líneas de inducción. Este hecho puede tener una sola explicación: el propio cuerpo de hierro añade al campo magnético de la corriente eléctrica de la bobina primaria su propio campo magnético.

La diferencia $\Phi - \Phi_0$ se designa habitual mente con la letra J. De este modo,

$$J = (\mu - 1) \Phi_0$$

es un flujo magnético adicional creado por la propia sustancia.

Una vez terminado el experimento para medir la permeabilidad magnética y sacada la barra de la bobina descubrimos que la barra de hierro conserva la magnetización. Esta será menor que J, sin embargo, sigue siendo bastante considerable.

El magnetismo remanente de la barra de hierro puede eliminarse. Con este fin la barra en cuestión debe introducirse otra vez en nuestra instalación, pero ya de tal forma que el campo propio del metal y el campo de la corriente eléctrica de la

bobina primaria están dirigidos hacia los lados opuestos. Siempre se logra elegir una corriente primaria tal que con la ayuda del choque de inducción de dirección contraria se consigue quitar las propiedades magnéticas del hierro llevándolo al estado inicial. Debido a razones históricas en las cuales no nos detendremos, la magnitud del campo desmagnetizante lleva el nombre de fuerza coercitiva.

Esta propiedad peculiar de los materiales ferromagnéticos de conservar el magnetismo en ausencia de la corriente y la posibilidad de eliminar este magnetismo remanente por medio de la corriente eléctrica de correspondiente dirección se llama histéresis. ¿Cuál es la procedencia de esta palabra? No es difícil comprenderlo. Procede de la palabra griega «*hysteresis*» que significa «retraso».

En dependencia de los requisitos técnicos surge la necesidad de materiales magnéticos con diferentes propiedades. El valor de de la aleación magnética «permalloy» se aproxima a 100 000 y los valores máximos de μ para el hierro dulce son cuatro veces menores.

La posibilidad de aumentar el flujo de líneas de inducción un enorme número de veces, al colocar el cuerpo de hierro dentro de una bobina de alambre, lleva a la creación de imanes eléctricos. Se sobreentiende que la potencia del imán eléctrico, es decir, su capacidad de atraer y retener cuerpos de hierro de masa grande crece con el aumento de la corriente que se hace pasar por el devanado del mismo. Este proceso no es ilimitado: existe el fenómeno de saturación, cuando los recursos internos del núcleo ferromagnético resultan agotados.

Al alcanzar cierta frontera de temperatura (por ejemplo, 767 °C para el hierro y 360 °C para el níquel) las propiedades ferromagnéticas desaparecen y la permeabilidad magnética se hace próxima a la unidad, como en todos los demás cuerpos.

Dominios

Las peculiaridades de los ferromagnéticos se explican por su estructura de dominios. El dominio es una zona magnetizada hasta el límite. Dentro del dominio todos los átomos están alineados de tal modo que sus momentos magnéticos resultan ser paralelos unos a otros.

El comportamiento de los dominios magnéticos es completamente análogo al de los dominios eléctricos en los materiales ferroeléctricos. Las dimensiones lineales de los

dominios magnéticos no son tan pequeñas, a saber, son del orden de 0,01 mm. Esta es la razón de que con un artificio no complicado se logra ver los dominios valiéndose de un microscopio común y corriente.

Para lograr que los dominios magnéticos se hagan visibles, a la superficie pulida de un monocristal ferromagnético se aplica una gota de una suspensión coloidal que no es sino una sustancia ferromagnética escrupulosamente triturada del tipo de magnetita. Las partículas coloidales se concentran cerca de las fronteras de los dominios, puesto que a lo largo de éstas los campos magnéticos son particularmente intensos (de una manera completamente análoga los imanes ordinarios agrupan las partículas magnéticas en las zonas próximas a sus polos).

Al igual que en los ferroeléctricos, los dominios en los materiales ferromagnéticos existen no sólo en presencia del campo magnético externo, sino también cuando la muestra no está imantada.

Los dominios en el monocristal no magnetizado se disponen de tal forma que el momento magnético total resulta ser igual a cero. Sin embargo, de aquí no se infiere que la disposición de los dominios es caótica. Otra vez, en completa analogía con lo expuesto en la pág. 68, el carácter de la estructura cristalina impone algunas direcciones en que para los momentos magnéticos resulta más fácil alinearse. Los cristales del hierro tienen la celdilla elemental cúbica y las direcciones de la magnetización más fácil son los ejes del cubo. Para otros metales ferromagnéticos los momentos magnéticos se orientan a lo largo de las diagonales del cubo. En un cristal no imantado se tiene igual cantidad de dominios con los momentos magnéticos dirigidos hacia un lado y de dominios con los momentos magnéticos orientarlos hacia el lado opuesto. Ya hemos aducido ejemplos de estructura de dominios en la fig. 2.5.

La magnetización, al igual que la polarización, consiste en que se «devoran» los dominios cuyos momentos están situados formando un ángulo obtuso con el campo. La lucha de las tendencias al orden y al desorden en la disposición de los átomos es una particularidad indispensable de cualquier estado de la sustancia. Sobre el particular se trata detalladamente en el libro del mismo autor editado en ruso: «El orden y el desorden en el inundo de los átomos».

Como ya se ha aclarado en el libro 2 de la «Física para todos», la tendencia al orden es tendencia al mínimo de energía. Si el movimiento térmico es insignificante, entonces, las partículas dejadas a seguir su libre albedrío forman una maravilla de la arquitectura atómica que es el cristal. El cristal es el símbolo del orden perfecto en el mundo de los átomos. Y la tendencia al desorden se dicta por la ley del crecimiento de la entropía.

Cuando crece la temperatura, las tendencias entrópicas vencen, y el desorden llega a ser la forma dominante de existencia de la materia.

En el caso de los materiales ferromagnéticos los asuntos se presentan de la siguiente manera. A medida que asciende la temperatura los momentos magnéticos comienzan a oscilar. Al principio, esta oscilación transcurre al unísono, sin alterar el orden, luego, ora uno, ora otro átomo se revuelcan ocupando una posición «incorrecta» El número de estos átomos que «abandonaron la fila» crece de un momento al otro y, finalmente, a una temperatura estrictamente determinada (temperatura o punto de Curie) tiene lugar la completa «fusión» del orden magnético.

En este libro a mí me es difícil explicar por qué una cantidad tan insignificante posee propiedades ferromagnéticas. ¿Qué detalles, precisamente, de la estructura de los átomos pusieron estas sustancias en una posición excepcional? Pero, el lector se mostraría demasiado exigente frente al autor si hubiera deseado obtener en este sucinto libro de divulgación las respuestas a todas las preguntas.

Pasemos a la descripción del comportamiento de otras sustancias.

Cuerpos diamagnéticos y paramagnéticos

Como ya hemos dicho, a excepción de una pequeña clase de ferromagnéticos, todas las demás sustancias tienen la permeabilidad magnética muy próxima a la unidad. Los cuerpos cuyo valor de μ la supera un poco la unidad se denominan paramagnéticos y aquellos cuya permeabilidad magnética es menor que la unidad llevan el nombre de diamagnéticos. Damos a continuación ejemplos de sustancias de ambas clases indicando los valores de su permeabilidad magnética:

Paramagnéticos	μ	Diamagnéticos	μ
Aluminio	1,000023	Plata	0,999981
Volframio	1,000175	Cobre	0,999912
Platino	1,000253	Bismuto	0,999324

A pesar de que las desviaciones respecto a la unidad son muy pequeñas se logra realizar mediciones muy exactas. Hablando con propiedad, para alcanzar este fin se puede hacer uso del método de choque de inducción, aquel mismo con que iniciamos nuestro relato sobre las mediciones magnéticas de las propiedades de la sustancia. Sin embargo, los resultados más exactos se obtienen empleando balanza magnética.

En uno de los platillos de la microbalanza analítica (la cual, como es sabido, es capaz de medir las fuerzas con una precisión de diezmillonésimas partes de gramo) se practica un orificio y a través de éste se deja pasar un hilo del cual se cuelga la muestra colocada entre los polos de un imán. Los terminales del imán deben hacerse de tal manera que el campo sea no homogéneo. En este caso el cuerpo ya sea que se atrae dentro de la zona del campo intenso, o bien, se expulsa de esta zona. El cuerpo se atrae cuando el momento magnético de la muestra tiende a disponerse a lo largo del campo, y se expulsa en el caso contrario. La fórmula de la fuerza se ha dado en páginas anteriores.

La muestra se equilibra mediante pesas en ausencia del campo magnético. Pero después de que la muestra resulta afectada por el campo el equilibrio se altera. En el caso de sustancias paramagnéticas será necesario añadir pesas, y en el caso de sustancias diamagnéticas habrá que aligerar el platillo de la balanza. No es difícil calcular que una buena balanza nos ayudará a cumplir esta difícil tarea, ya que (en el caso fácilmente realizable cuando la no homogeneidad del campo es del orden de centésimas de tesla por centímetro) sobre 1 cm^3 de sustancia actuará una fuerza de cerca de un miligramo.

Las dos propiedades, la paramagnética y la diamagnética, se explican de una forma bastante sencilla.

El diamagnetismo es la consecuencia directa de la circunstancia de que en un campo magnético cada electrón describirá una circunferencia. Estas corrientes circulares engendran sus propios momentos magnéticos dirigidos en contra del campo que produjo la rotación.

El diamagnetismo es una propiedad común para todas las sustancias.

El paramagnetismo, sin hablar ya del ferromagnetismo, «eclipsan» las propiedades diamagnéticas de la sustancia.

Pertenecen a los paramagnéticos las sustancias cuyos átomos o iones poseen el momento magnético. El momento puede originarse por el movimiento orbital de los electrones, por el espín de un electrón solitario o por las dos causas tomadas juntas.

En ausencia del campo magnético los átomos de las sustancias diamagnéticas no tienen momento magnético. Los átomos de las sustancias paramagnéticas sí que poseen momentos magnéticos, no obstante, debido al movimiento térmico éstos están dispuestos en completo desorden, absolutamente así como es característico para los cuerpos ferromagnéticos por encima del punto de Curie. Al superponerse un campo comienza la lucha entre la fuerza ordenadora del campo y el desorden impuesto por el movimiento térmico. A medida que disminuye la temperatura un número cada vez mayor de átomos se orienta de modo que su momento magnético forma un ángulo agudo con la dirección del campo. De aquí queda completamente claro que con el descenso de la temperatura aumenta la permeabilidad magnética de los cuerpos paramagnéticos.

El campo magnético de la Tierra

El hombre de hoy está acostumbrado a que cualquier aparato o instrumento aparece como consecuencia del desarrollo de alguna teoría física. Una vez creado el instrumento, de éste se preocupan los ingenieros, y los físicos ya no tienen nada que hacer con él; la naturaleza del fenómeno en que se basa el funcionamiento de dicho instrumento se ha comprendido antes de su creación.

Sin embargo, tratándose de la brújula las cosas eran por completo distintas. Probablemente, hubiera aparecido en China en el siglo XI y durante varios siglos se utilizó como el principal instrumento de navegación, sin embargo, con todo, nadie

entendía de hecho el principio de su funcionamiento. ¿Por qué un extremo de la aguja señalaba siempre el Norte? La mayoría de los sabios explicaba el comportamiento de la aguja debido a la influencia de las fuerzas extraterrestres, digamos, debido a la atracción del extremo de la aguja por parte de la Estrella Polar.

En 1600 salió a la luz la brillante obra de William Gilbert intitulada «*De magnete, magnetisque corporibus, et de magno magnete tellure*» (Sobre el imán terrestre). El riguroso enfoque científico permitió al sabio comprender con mayor profundidad la esencia de los fenómenos magnéticos. Gilbert talló de un trozo de mena magnética una esfera y estudió meticulosamente la orientación de la aguja magnética colgada bajo las distintas partes de dicha esfera, advirtiendo el total parecido con la orientación de la aguja magnética en diferentes zonas de la Tierra. De aquí se sacó la siguiente conclusión; el funcionamiento de la brújula se puede explicar perfectamente en el caso de suponer que la Tierra es un imán permanente cuyo eje está dirigido a lo largo del eje terrestre.

A partir de este momento el estudio del geomagnetismo pasa a un nuevo nivel. Una investigación más escrupulosa demostró que la aguja magnética no está orientada del todo exactamente desde el Norte hacia el Sur. La desviación de la dirección de la aguja respecto al meridiano trazado a través del punto dado se denomina declinación magnética. Los polos magnéticos están desplazados respecto al eje de rotación de la Tierra en $11,5^\circ$ (fig. 3.10).

La aguja no se encuentra justamente en el plano horizontal, sino gira hacia el horizonte formando cierto ángulo denominado ángulo de inclinación magnética. Al investigar la inclinación magnética en diferentes puntos se puede hacer la conclusión de que el «dipolo» magnético se ubica en las profundidades de la Tierra. Esto engendra un campo no homogéneo el cual, en los polos magnéticos, alcanza el valor de $0,6 \times 10^{-4}$ Tl y en el ecuador es igual a de $0,3 \times 10^{-4}$ Tl.

Bueno, ¿qué imán es éste el que se encuentra en el seno de la Tierra? El «dipolo» magnético se dispone en el núcleo de la Tierra que consta de hierro en estado fundido. Pero incluso en este estado el hierro sigue siendo un buen conductor de la electricidad, y para explicar el campo magnético de la Tierra puede proponerse el modelo de una especie de «dínamo magnético». No describiremos este modelo,

limitándonos con indicar que el «imán terrestre» se crea por las corrientes que circulan en el interior del hierro fundido.

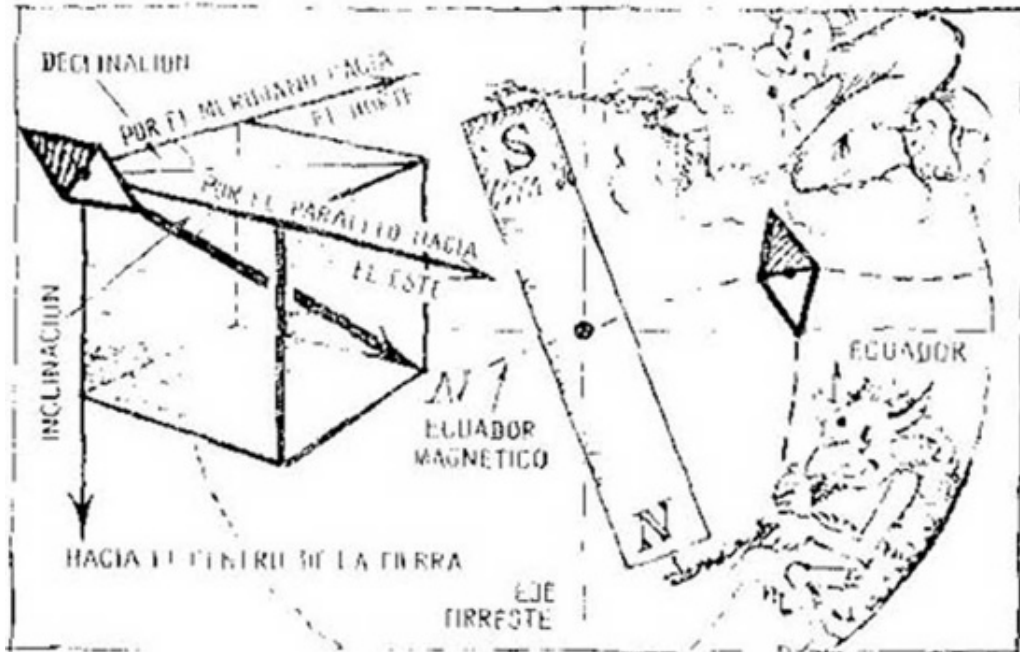


Figura 3.10

El campo magnético de la Tierra varía. Los polos magnéticos se desplazan con una velocidad de 5 a 6 km al año. A escala de la Tierra este desplazamiento es despreciable; el fenómeno apenas si se nota con facilidad durante un centenar de años, por esta razón recibió el nombre de variación secular del campo magnético.

Huelga demostrar cuán esencial es el conocimiento exacto de los elementos del magnetismo terrestre en cualquier sitio de la superficie de nuestro planeta. Hasta la fecha la brújula magnética sirve a los navegantes. Siendo así, éstos deben disponer de mapas de declinaciones o inclinaciones magnéticas. En las proximidades de los polos el extremo norte de la aguja magnética, como se ve en el dibujo, ya, plenamente, deja de mirar al Norte. También en las cercanías del ecuador es difícil pasar sin el mapa del campo magnético. El ecuador magnético no coincide, ni mucho menos, con la línea de las latitudes cero.

El conocimiento del campo magnético de la tierra firme también reviste enorme interés, ya que sirve a los fines de prospección geológica. No podemos detenernos

en estos problemas. La física geológica es un capítulo importante y vasto de la ciencia que merece una conversación especial.

Cabe decir varias palabras sobre las llamadas investigaciones paleomagnéticas que permiten juzgar cuál fue el campo magnético de la Tierra en los tiempos muy remotos. Estas investigaciones se basan, fundamentalmente, en el estudio de la magnetización remanente de las rocas, etc.

He aquí, por ejemplo, en qué consiste la esencia de los métodos que estudian el período prehistórico. Un ladrillo y un jarrón de barro acusan una pequeña magnetización remanente que surge en la arcilla caliente durante la cocción. La dirección del momento magnético corresponde a la del campo magnético en el momento de fabricación y enfriamiento del objeto. A veces se puede juzgar con bastante seguridad sobre la posición de dicho objeto en el momento de su confección.

Citemos otro ejemplo de semejantes investigaciones: se estudia la dirección geográfica del momento magnético de la mena, mientras que su edad se determina basándose en la cantidad de isótopos radiactivos.

Las investigaciones paleomagnéticas son la demostración más rigurosa de la deriva de los continentes. Resultó que las magnetizaciones de los yacimientos de hierro surgidos varios centenares de millones de años atrás en diferentes continentes pueden dirigirse a lo largo de las líneas de inducción del campo magnético de la Tierra, si estos continentes se agrupan en un supercontinente único llamado Gondvana. Más tarde este continente se fragmentó en África, Australia, Antártida y América del Sur.

Hasta este momento hemos hablado solamente sobre la procedencia ultraterrestre del magnetismo y, en efecto, es su fuente principal. Sin embargo, algunas variaciones del campo magnético tienen lugar debido a las partículas cargadas que llegan desde fuera. Se trata, de modo fundamental, de las corrientes de protones y neutrones emitidos por el Sol. Las partículas cargadas se arrastran por el campo hacia los polos donde giran describiendo una circunferencia impulsados por las fuerzas de Lorentz. Dicha circunstancia conduce a dos fenómenos. En primer lugar, las partículas cargadas en movimiento producen un campo magnético complementario, o sea, las tempestades magnéticas. En segundo lugar, ionizan las

moléculas de los gases atmosféricos y como consecuencia aparece la aurora boreal. Las tempestades magnéticas fuertes tienen lugar periódicamente (a intervalos de 11,5 años). Este período coincide con el de procesos intensos que se operan en el Sol.

Mediciones directas realizadas con la ayuda de los aparatos cósmicos han demostrado que los cuerpos más próximos a la Tierra, la Luna y los planetas Venus y Marte, no poseen un campo magnético propio semejante al terrestre. Entre los demás planetas del Sistema Solar solamente Júpiter y, por lo visto, Saturno poseen campos propios. En Júpiter se han descubierto campos hasta 10 Gs y una serie de fenómenos característicos (tempestades magnéticas, radioemisión sincrotrónica, etc.).

Campos magnéticos de las estrellas

No sólo los planetas y las estrellas frías poseen magnetismo, sino también los cuerpos celestes incandescentes.

Por cuanto el más próximo a nosotros es el Sol, conocemos más sobre su campo magnético que sobre los de otras estrellas. El campo magnético del Sol puede observarse visualmente durante los eclipses solares. A lo largo de las líneas de inducción se alinean partículas de materia solar que poseen momento magnético, perfilando el cuadro del campo magnético. Se ven nítidamente los polos magnéticos y es posible evaluar la magnitud del campo magnético que, en algunas regiones cuya superficie es del orden de diez mil kilómetros, supera miles de veces la intensidad del campo magnético de la Tierra. Estas porciones se denominan manchas solares. Puesto que las manchas son más oscuras que los demás lugares del Sol resulta claro que la temperatura aquí es más baja, a saber, es 2000 grados menor que la temperatura «normal» del Sol.

Es indudable que la baja temperatura y los valores elevados del campo magnético guardan una relación entre sí. Sin embargo, no existe una teoría adecuada que enlace estos dos hechos.

Bueno, ¿y cómo van los asuntos en otras estrellas? Los alcances de la astrofísica en los últimos años son tan considerables que resultó posible establecer la existencia de campos magnéticos en las estrellas. Se dejó constancia de que las «manchas

magnéticas estelares» tienen la temperatura de cerca de 10.000 grados y durante varios meses pueden cambiar su posición y hasta desaparecer absolutamente. Resulta más sencillo explicar este cambio si se admite que no son las manchas las que cambian de posición en las estrellas, sino que toda la estrella gira.

Se juzga sobre la presencia de los campos magnéticos basándose en los intensidades anómalas de algunas líneas espectrales. Según parece, las estrellas magnéticas poseen en el ecuador magnético un contenido elevado de hierro.

Los campos magnéticos en el cosmos son muy pequeños (millonésimas partes de gaussio). Este hecho ni siquiera necesita una explicación puesto que en el cosmos reina un altísimo vacío. Cuando de los átomos diseminados por el Universo se forman estrellas» la condensación de la materia estelar viene acompañada con la «condensación» del campo magnético. Pero, siendo así, ¿por qué no todas las estrellas poseen un campo magnético?

La Tierra existió miles de millones de años. De aquí se infiere que el campo magnético de la Tierra se mantiene todo el tiempo por las corrientes eléctricas que fluyen por sus entrañas. Algunas estrellas carentes del campo magnético, evidentemente, se han enfriado hasta tal punto que en su seno cesaron las corrientes eléctricas. No obstante, esta explicación, difícilmente, puede considerarse universal.

Se conocen campos magnéticos de las estrellas que superan 10 000 Gs.

Y unos campos fantásticamente intensos de 10^{15} Gs (¡!) deben encontrarse cerca de las estrellas neutrónicas.

Capítulo 4

Recapitulación de la electrotecnia

Contenido:

- *Fuerza electromotriz sinusoidal*
- *Transformadores*
- *Máquinas generadoras de la corriente eléctrica*
- *Motores eléctricos*

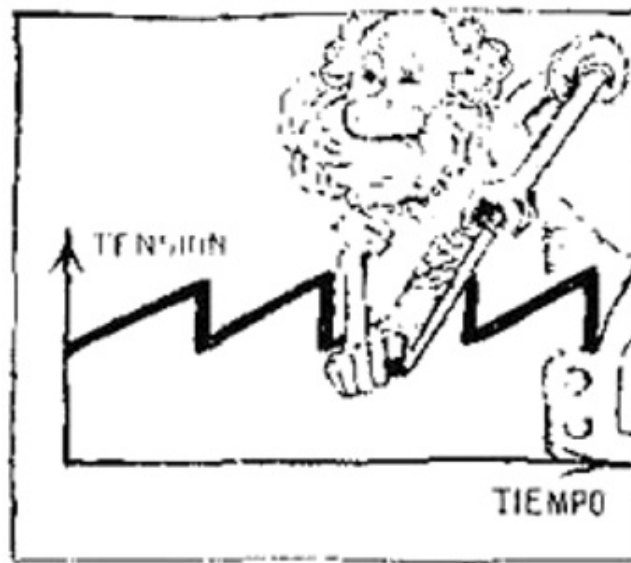
Fuerza electromotriz sinusoidal

El acumulador y la pila son fuentes de corriente continua. Y la red eléctrica nos proporciona la corriente alterna. Las palabras «continua» y «alterna» se refieren a la tensión, a la f.e.m. y a la intensidad de la corriente. Si en el proceso de paso de la corriente todas estas magnitudes quedan invariables la corriente es continua, y si se modifican, entonces, la corriente resulta ser alterna.

El carácter de variación de la corriente respecto al tiempo puede ser diferente en dependencia de la instalación productora de la corriente. Por medio de un tubo de rayos catódicos puede obtenerse la curva que caracteriza la variación de la corriente eléctrica. El haz electrónico se desvía por los campos de dos condensadores planos perpendiculares entre sí. Aplicando a las placas del condensador distintas tensiones es posible lograr que el punto luminoso dejado en la pantalla por el haz deambule por toda la superficie de la misma.

Para obtener el cuadro de la corriente alterna se obra de la siguiente manera. A un par de placas se conecta la llamada tensión en diente de sierra cuya curva se muestra en la fig. 4.1.

Si el haz electrónico está sujeto tan sólo a la acción de ésta, el punto luminoso se desplaza uniformemente por la pantalla para luego, a salto, volver a la posición inicial. La posición del punto luminoso proporciona datos sobre el momento del tiempo. Si sobre el otro par de placas se superpone la tensión alterna estudiada, ésta se «explora» de una forma absolutamente análoga a la de la «exploración» de la oscilación mecánica que se consignó mediante un sencillo dispositivo representado en el libro 1.

*Figura 4.1*

Al decir la palabra «oscilación» no me he equivocado. En la mayoría de los casos las magnitudes que caracterizan la corriente alterna oscilan siguiendo la misma ley armónica de la senoide a la que están subordinadas las desviaciones del péndulo respecto al equilibrio. Para cerciorarse de ello basta conectar al oscilógrafo la corriente alterna urbana.

Por la vertical pueden trazarse la corriente o la tensión. Las características de la corriente son las mismas que los parámetros de la oscilación mecánica. El intervalo de tiempo después del cual se repite el cuadro de las variaciones, como se sabe, lleva el nombre de período T ; la frecuencia de la corriente ν - que es una magnitud inversa al período - para la corriente urbana es igual, ordinariamente, a 50 oscilaciones por segundo.

Cuando se somete al examen una sola senoide es indiferente la elección del punto de referencia del tiempo. En cambio, si dos sinusoides se superponen de una forma tal como se muestra en la fig. 4.2, entonces, es necesario señalar en qué parte del período están desfasadas. Se denomina fase el ángulo $\varphi = 2\pi\tau/T$. De este modo, si las curvas están desplazadas una respecto a la otra un cuarto de período se dice que están desfasadas 90° ; si el desplazamiento constituye una octava parte del período, el desfase es de 45° , etc.

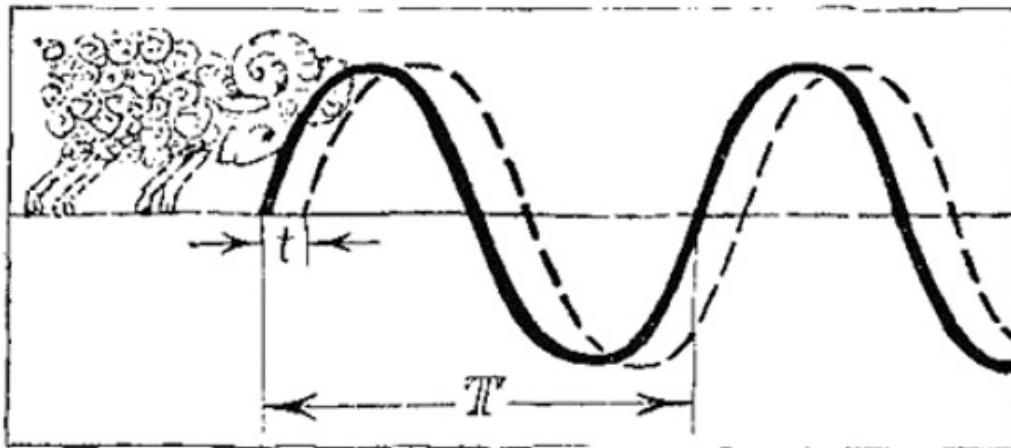


Figura 4.2

Cuando se trata de varias sinusoides desfasadas, los técnicos hablan sobre los vectores de corriente o de tensión. La longitud del vector corresponde a la amplitud de la senoide, y el ángulo entre los vectores, al desfase. Muchos dispositivos técnicos no nos dan una corriente sinusoidal simple, sino una corriente cuya curva representa la suma de varias sinusoides desfasadas.

Demostremos que una corriente sinusoidal simple se engendra si el cuadro conductor gira en el campo magnético homogéneo a velocidad constante.

Cuando la dirección del cuadro respecto a las líneas de inducción es arbitraria, el flujo magnético que atraviesa el circuito es igual a

$$\Phi = \Phi_{\max} \sin \varphi$$

φ es el ángulo entre el plano de la espira y la dirección del campo. Dicho ángulo varía en función del tiempo de acuerdo con la ley $\varphi = 2\pi t/T$.

La ley de la inducción electromagnética permite calcular la f.e.m. de inducción. Anotemos las expresiones de los flujos magnéticos para dos instantes que se diferencian en un intervalo de tiempo sumamente pequeño τ :

$$\Phi = \Phi_{\max} \sin \frac{2\pi}{T} t, \quad \Phi = \Phi_{\max} \sin \frac{2\pi}{T} (t + \tau),$$

La diferencia de estas expresiones es:

$$2\Phi_{\text{ins}} \cos \frac{2\pi}{T} \left(t + \frac{\tau}{2} \right) \sin \left(\frac{2\pi\tau}{T} \right),$$

Puesto que el valor de τ es muy pequeño, resultan válidas las siguientes igualdades aproximadas:

$$\sin \left(\frac{2\pi\tau}{T} \right) \approx \frac{2\pi\tau}{T}, \quad \cos \frac{2\pi}{T} \left(t + \frac{\tau}{2} \right) \approx \cos \frac{2\pi}{T} t,$$

La f.e.m. de inducción es igual a esta diferencia referida al tiempo. Por consiguiente,

$$\xi^{\text{ind}} = \frac{2\pi}{T} \Phi_{\text{ins}} \cos \frac{2\pi}{T} t = \frac{2\pi}{T} \Phi_{\text{ins}} \sin \left(\frac{2\pi}{T} t - \frac{\pi}{2} \right)$$

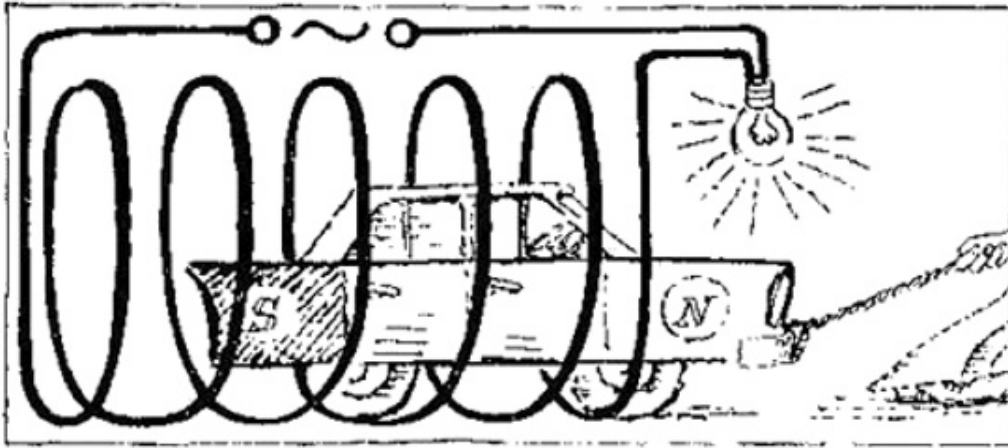
Hemos demostrado que la f.e.m. de inducción se expresa por una senoide desfasada 90° respecto a la senoide del flujo magnético. En cuanto al valor máximo de la f.e.m. de inducción, o sea, de su amplitud, esto es proporcional al producto de la amplitud del flujo magnético por la frecuencia de rotación del cuadro. La ley para la intensidad de la corriente se obtiene si se divide la f.e.m. de inducción por la resistencia del circuito. Pero cometeremos un error grave si ponemos el signo de igualdad entre la resistencia a la corriente alterna que figura en el denominador de la expresión

$$I_{\text{alt}} = \frac{\xi^{\text{ind}}}{R_{\text{alt}}}$$

y la resistencia óhmica, o sea, la magnitud con la cual hemos tratado hasta este momento. Resulta que R_{alt} (no sólo se determina por la resistencia óhmica, sino que

depende también de otros dos parámetros del circuito: de su inductancia y de las capacidades incluidas en éste.

El hecho de que la ley de Ohm se hace más complicada al pasar de la corriente continua a la alterna lo evidencia el siguiente experimento simple.



En la fig. 4.3 se representa el circuito de la corriente que fluye a través de una bombilla eléctrica y una bobina en que puede introducirse el núcleo de hierro. Al principio conectemos la bombilla a la fuente de corriente continua. Reiteradas veces vamos a introducir el núcleo de hierro en la bobina y, respectivamente, sacarlo. ¡No hay efecto alguno! La resistencia del circuito no cambia, por consiguiente, también queda invariable la intensidad de la corriente. Pero repetimos este mismo experimento para el caso de que el circuito está conectado a la corriente alterna. Ahora, si el núcleo está fuera de la bobina la bombilla emite una luz brillante, mas, al introducir el hierro en la bobina, la luz se torna débil. ¿No es verdad que hemos obtenido un resultado ostensible?

De este modo, siendo invariables tanto la tensión externa, como la resistencia óhmica (que depende solamente del material de los conductores, su longitud y sección), la intensidad de la corriente varía en dependencia de la posición del núcleo de hierro en la bobina.

¿Qué significa este hecho?

Recordamos que el núcleo de hierro hace aumentar bruscamente (miles de veces) el flujo magnético que atraviesa la bobina. En el caso de la f.e.m. alterna el flujo

magnético en la bobina cambia constantemente. Pero, si bien, sin el núcleo de hierro el mismo cambiaba desde cero hasta cierta unidad convencional, en presencia del núcleo su variación será desde cero hasta varios miles de unidades.

Con la variación del flujo magnético las líneas de inducción atravesarán las espiras de «su» bobina. En ésta aparecerá la corriente de autoinducción. De acuerdo con la regla de Lenz esta corriente estará dirigida de tal forma que debilite el efecto que la ha engendrado: la f.e.m. externa se encuentra con un obstáculo especial que no existía cuando la corriente era continua. En otras palabras, la corriente alterna tiene una resistencia complementaria debida a que el campo magnético, al atravesar los conductores de su circuito, engendra una f.e.m. especial llamada f.e.m. de autoinducción la cual debilita la intensidad media de la corriente. Esta resistencia complementaria se denomina inductiva.

La experiencia testimonia (y, sin duda alguna, esta circunstancia parecerá al lector muy natural) que el flujo magnético que atraviesa la bobina (o, hablando de una forma más general, que atraviesa todo el circuito de la corriente) es proporcional a la intensidad de la corriente: $\Phi = LI$. En cuanto al coeficiente de proporcionalidad L que se denomina inductancia, éste depende de la geometría del circuito conductor, así como de los núcleos que abarca. Como resulta evidente de la fórmula, el valor numérico de la inductancia es igual al flujo magnético para la intensidad de la corriente de un amperio. La unidad de inductancia L es el henrio ($1 \text{ H} = 1 \Omega \text{ s}$)

Se puede deducir teóricamente y confirmar en el experimento que la resistencia inductiva R_L viene expresada por la fórmula

$$R_L = 2\pi\nu L$$

Si la resistencia óhmica (que ya conocemos) y la resistencia capacitiva (con la cual vamos a entablar conocimiento un poco más tarde) son pequeñas, la intensidad de la corriente en el circuito es igual a

$$I = \frac{\xi}{R_L}$$

Para estar en condiciones de juzgar qué es «pequeño» a qué es «grande» calculemos aproximadamente el valor de la resistencia inductiva para la frecuencia de la corriente urbana y la inductancia de 0,1 H. Resulta unos 30 Ω .

Bueno, ¿y qué representa una bobina cuya inductancia es de 1 H? Para evaluar la inductancia de las bobinas y bobinas de choque (estas últimas son bobinas con núcleo de hierro) se emplea la siguiente fórmula que insertamos sin deducción:

$$L = \mu_0 \mu \frac{n^2}{l} S, \quad \mu_0 = 4\pi \cdot 10^{-7} \text{ J/(A}^2 \text{ m)}$$

aquí n es el número de espiras; l la longitud de la bobina, y S , la sección transversal. De tal manera, por ejemplo, 0,002 H nos proporcionará una bobina con los siguientes parámetros: $L = 15 \text{ cm}$, $n = 1500$, $S = 1 \text{ cm}^2$. Si introducimos un núcleo de hierro con $\mu = 1000$, la inductancia será de 2 H. La existencia de la f.e.m. de cualquier origen y, en consecuencia, también de la f.e.m. de autoinducción significa que se realiza un trabajo. Como sabemos, este trabajo es igual a ξI . Si la corriente es alterna, entonces, tanto ξ , como I a cada instante cambian sus valores. Supongamos que en el momento t sus valores son ξ_1 e I_1 y en el momento $t + \tau$ son de ξ_2 e I_2 , respectivamente. El flujo magnético que atraviesa las espiras de la bobina con la inductancia L es igual a LI . En el momento t , su valor fue L/I_1 y en el momento $t + \tau$, LI_2 . ¿A qué es igual, entonces, el trabajo necesario para incrementar la corriente desde su valor I_1 , hasta el valor de I_2 ? La f.e.m. es igual a la variación del flujo magnético referido al tiempo de variación:

$$\xi = L \frac{(I_2 - I_1)}{\tau}$$

Para obtener el trabajo $\xi I \tau$ es necesario multiplicar la expresión dada por el tiempo y la intensidad de la corriente. ¿Y qué intensidad se tiene en cuenta? Se tiene en cuenta su valor medio, es decir, $(I_1 + I_2)/2$. Llegamos a la conclusión de que el trabajo de la f.e.m. de autoinducción es igual a

$$\frac{L}{2}(I_2 + I_1)(I_2 - I_1) = \frac{L}{2}I_2^2 - \frac{L}{2}I_1^2$$

Este resultado aritmético puede expresarse de la siguiente forma: el trabajo de la f.e.m. es iguala la diferencia de la magnitud $LI^2/2$ en dos momentos de tiempo. Esto significa que en la resistencia inductiva la energía no se disipa, ni tampoco se transforma en calor, hecho que tiene lugar en los circuitos con resistencia óhmica, sino pasa a la «reserva». Precisamente por esta razón es lógico llamar la magnitud $LI^2/2$ energía magnética de la corriente.

Examinemos ahora cómo repercutirá en la resistencia del circuito a la corriente alterna la intercalación de un condensador.

Si al circuito de corriente continua se conecta un condensador la corriente no fluirá, ya que intercalar un condensador es lo mismo que interrumpir el circuito. Pero el mismo condensador en el circuito de corriente alterna no anula la corriente.

Se sobreentiende que a nosotros nos interesa la causa de semejante diferencia. La explicación es sencilla. Después de conectar el circuito al manantial de la corriente alterna la carga eléctrica comienza a acumularse en las armaduras del condensador. A una armadura llega la carga positiva, y a la otra, la negativa. Supongamos que tanto la resistencia inductiva, como la óhmica son pequeñas. La carga continuará hasta que la tensión en las armaduras del condensador llegue a ser máxima e igual a la f.e.m. del manantial. En este instante la intensidad de la corriente es igual a cero.

Al medir, valiéndose de algún instrumento, la intensidad - media por un período - de la corriente en un circuito con condensador, podemos cerciorarnos de que ésta será distinta en dependencia de dos magnitudes. En primer lugar, se demuestra (experimentalmente, así como por medio de razonamientos teóricos) que la corriente disminuye a medida que cae la frecuencia. En consecuencia, la resistencia capacitiva es inversamente proporcional a la frecuencia. Este resultado es completamente natural, pues cuanto menor es la frecuencia en tanto mayor grado la corriente alterna se aproxima, por decirlo así, a la continua.

Modificando los parámetros geométricos del condensador, o sea, la distancia entre las placas y el área de éstas, nos convenceremos de que la resistencia capacitiva es inversamente proporcional también a la capacidad del condensador.

La fórmula para la resistencia capacitiva tiene el siguiente aspecto:

$$R_C = \frac{1}{2\pi f C}$$

El condensador cuya capacidad es de 30 uF para la frecuencia de la corriente urbana da una resistencia de cerca de 100 Ω.

No me propongo contar al lector cómo se calcula la resistencia de los circuitos complejos de corriente compuestos de resistencias óhmicas, inductivas y capacitivas. Lo único que quiero advertir es que la resistencia total del circuito no equivale a la suma de resistencias aisladas.

La intensidad de la corriente eléctrica y la tensión en el tramo del circuito que incluye la resistencia óhmica, el condensador y la bobina de inductancia pueden medirse de una manera ordinaria empleando un oscilógrafo (tubo de rayos catódicos). Tanto la corriente, como la tensión las veremos en la pantalla en forma de sinusoides. No nos asombrará descubrir que estas sinusoides estén desfasadas una respecto a la otra formando cierto ángulo de fase φ . (El lector comprenderá muy pronto que así, precisamente, debe ocurrir, al recordar que, digamos, en un circuito con condensador la corriente es igual a cero cuando la tensión en el condensador es máxima.)

El valor del desfasaje φ reviste gran importancia. No olvide que la potencia de la corriente es igual al producto de la intensidad de la corriente por la tensión. Si las sinusoides de la corriente y de la tensión coinciden dicho valor será máximo, mientras que si éstas están desfasadas de un modo propio para un circuito que posee una sola resistencia capacitiva o una sola resistencia inductiva la potencia será igual a cero. No es difícil verlo al trazar dos sinusoides desfasadas 90°, multiplicar sus ordenadas y sumar estos productos correspondientes a un período. Se puede demostrar rigurosamente que, en el caso general, durante un periodo, en promedio, la potencia de la corriente alterna es igual a

$$W = IU \cos \varphi$$

El aumento del $\cos \varphi$ es tarea del ingeniero electricista.

Transformadores

Usted ha comprado una nevera. El vendedor le ha prevenido que la nevera está calculada para la tensión de la red de 220 V. Pero la tensión de la red en su casa es de 127 V. ¿Es éste un callejón sin salida? De ningún modo. Meramente, tendrá que gastar cierta suma complementaria para adquirir un transformador.

Un transformador es un dispositivo muy sencillo que permite tanto aumentar, como disminuir la tensión. Este consta de un núcleo de hierro que lleva dos devanados (bobinas). El número de espiras en las bobinas es distinto.

Conectamos a una de las bobinas la tensión de la red. Mediante un voltímetro podemos cerciorarnos de que en los extremos de otro devanado aparecerá la tensión diferente de la de la red. Si el devanado primario tiene ω_1 espiras y el secundario ω_2 , la relación de las tensiones será

$$\frac{U_1}{U_2} = \frac{\omega_1}{\omega_2}$$

De este modo, el transformador aumentará la tensión si la tensión primaria está conectada a la bobina con menor número de espiras, y la disminuirá en el caso contrario.

¿Por qué resulta así? Se trata de que todo el flujo magnético, en la práctica, pasa a través del núcleo de hierro. Por consiguiente, ambas bobinas están atravesadas por igual número de líneas de inducción. El transformador funcionará solamente en el caso de que la tensión primaria sea alterna. La variación sinusoidal de la corriente en la bobina primaria originará una f.e.m. sinusoidal en la bobina secundaria. La espira de la bobina primaria y la de la secundaria se encuentran en iguales condiciones. La f.e.m. de una espira de la bobina primaria es igual a la f.e.m. de la red dividida por el número de espiras de esta bobina, U_1/ω_1 y la f.e.m. de la bobina

secundaria es igual al producto del valor U_1/ω_1 por el número de espiras ω_2 . De principio, cada transformador puede utilizarse ya sea como transformador elevador, o bien, como reductor, según sea la bobina a la que está conectada la tensión primaria.

En la práctica cotidiana, con frecuencia, es necesario tratar con los transformadores (fig. 4.4).



Figura 4.4

Además de los transformadores que utilizamos, querámoslo o no, debido a que los aparatos industriales están calculados para una tensión, mientras que la red urbana utiliza otra, además de los mismos, tenemos que vérselas con las bobinas del automóvil. La bobina es un transformador elevador. Para producir la chispa que enciende la mezcla de combustible y aire se necesita una alta tensión la que obtenemos, precisamente, del acumulador del automóvil, transformando previamente la corriente continua del acumulador en corriente alterna con la ayuda de un interruptor. No es difícil comprender que, con la exactitud hasta las pérdidas de energía que se consume para calentar el transformador, al aumentar la tensión disminuye la intensidad de la corriente, y viceversa.

Para las máquinas de soldar se necesitan transformadores reductores. La soldadura requiere corrientes muy fuertes y el transformador de una máquina de soldar tiene tan sólo una espira de salida.

El lector, de seguro, prestó atención a que el núcleo del transformador se fabrica de finas chapas de acero. Esto se hace para no perder energía durante la transformación de la tensión. Como ya señalamos antes, en el material en chapas el papel perteneciente a las corrientes de Foucault es menos importante que en el material macizo.

En casa nos encontramos con transformadores pequeños. En lo que se refiere a los transformadores potentes, éstos representan estructuras muy grandes. En semejantes casos el núcleo junto con los devanados está sumergido en un recipiente lleno de aceite refrigerante.

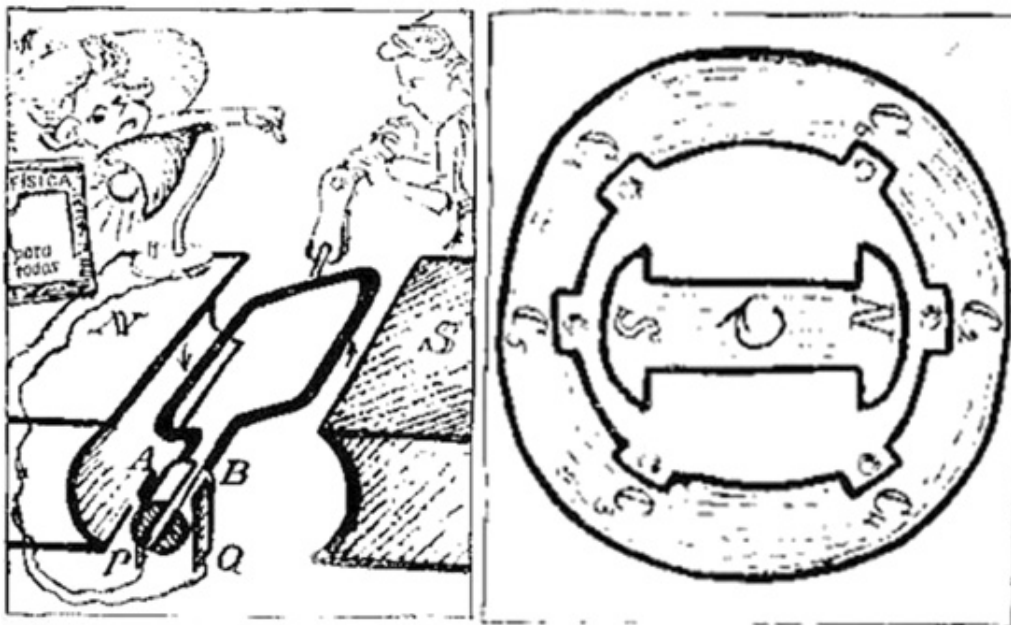
Máquinas generadoras de corriente eléctrica

Las máquinas que transforman el movimiento mecánico en corriente eléctrica se crearon tan sólo unos ciento cincuenta años atrás. El primer generador de corriente fue la máquina de Faraday en la cual una espira de alambre giraba en el campo de imanes permanentes. Poco tiempo se necesitó para que naciera la idea (pero no en la cabeza de Faraday) de sustituir una espira por una bobina y, de este modo, sumar todas las f.e.m. engendradas en todas las espiras. Solamente en 1851 los imanes permanentes fueron sustituidos por los electroimanes, es decir, por bobinas puestas sobre núcleos de hierro. Surgió el término «excitación de la máquina» ya que, para conseguir que ésta comenzara a producir corriente era preciso «animar» el imán eléctrico. Al principio, para excitar la máquina al devanado del imán eléctrico se suministraba corriente de una fuente externa de alimentación.

La consiguiente etapa fue el descubrimiento del principio de autoexcitación de la máquina conforme al cual para excitar los imanes eléctricos no es obligatorio tener una fuente complementaria de alimentación. Es suficiente conectar el devanado de excitación de los electroimanes, empleando uno u otro procedimiento, con el devanado principal de la máquina. Para finales de los años 80 del siglo XIX, la máquina eléctrica adquirió sus rasgos fundamentales que se conservaron hasta la fecha. En la fig. 4.5 se representa el modelo elemental del generador de corriente

continua. Si en el campo de imanes permanentes se hace girar un cuadro con corriente, en este se inducirá la f.e.m. sinusoidal.

En el caso de que se desea obtener corriente continua a partir de la alterna, hace falta proveer la máquina de un dispositivo especial llamado colector. Este consiste de dos semianillos A y B aislados uno del otro y puestos sobre el cilindro común (fig. 4.5). El cilindro gira junto con el cuadro. A los semianillos están aplicados los contactos P y Q (escobillas) mediante los cuales la corriente se deriva al circuito exterior.



Figuras 4.5 y 4.6

Para cada semirrevolución del cuadro sus extremos pasan de una escobilla a otra. Por esta razón, a pesar de la variación de la dirección de la corriente en el propio cuadro la corriente del circuito exterior no cambia su sentido. Puesto que la parte rotativa de una máquina real consta de un gran número de cuadros, o sea, secciones desfasadas una respecto a la otra en un ángulo determinado, y el colector está constituido por un número correspondiente de láminas llamadas delgas, resulta que en las escobillas de la máquina se obtiene, prácticamente, una tensión constante.

Actualmente, se construyen generadores de corriente continua con una potencia desde partes de kilovatio hasta varios miles de kilovatios. Generadores grandes se emplean para la electrólisis en la industria química y la metalurgia no ferrosa (producción de aluminio y de cinc). Los mismos están calculados para grandes corrientes y tensiones relativamente bajas (120 a 200 V, 1000 a 20.000 A). Las máquinas de corriente continua se utilizan también para la soldadura eléctrica.

Sin embargo, los generadores de corriente continua no son productores principales de energía eléctrica. En la URSS, para la producción y distribución de la energía eléctrica está adoptada la corriente alterna de 50 Hz de frecuencia. Un generador de corriente alterna se diseña de tal forma que permita obtener simultáneamente tres f.e.m. de igual frecuencia pero desfasadas una respecto a la otra en un ángulo de fase de $2\pi/3$

En la fig. 4.6 se da una representación esquemática del generador trifásico de dicho tipo. En esta figura cada una de las bobinas viene sustituida por una espira. Los hilos de una de las espiras se designan con $C_1 - C_4$, de la segunda $C_2 - C_5$ y de la tercera, $C_3 - C_6$. Si la corriente entra en C_1 , entonces, ésta sale de C_4 , etc. (Se sobreentiende que en los momentos respectivos a distintas posiciones del rotor y del estator cualquiera de los extremos pueden servir de entrada o de salida de la corriente.) La f.e.m. en las espiras inmóviles del devanado del estator se induce como resultado de su intersección por el campo magnético del imán eléctrico giratorio, o sea, del rotor. Cuando el rotor gira a velocidad constante en los devanados de las fases del estator aparecen f.e.m. que varían periódicamente de igual frecuencia, pero desfasadas una respecto a otra en el ángulo de 120° debido a su desplazamiento espacial.

Tres espiras de la bobina pueden unirse entre sí de diferente manera: en estrella o en triángulo (en delta). Estos esquemas de conexión estuvieron elaborados e introducidos en la práctica por Mijail Osipovich Dolivo-Dobrovolski (1862 - 1910) a principios de la década del 90 del siglo XIX. Para la conexión en estrella los extremos de todos los devanados del generador C_4 , C_5 y C_6 se unen con un punto que se denomina punto neutro (fig. 4.7, a).

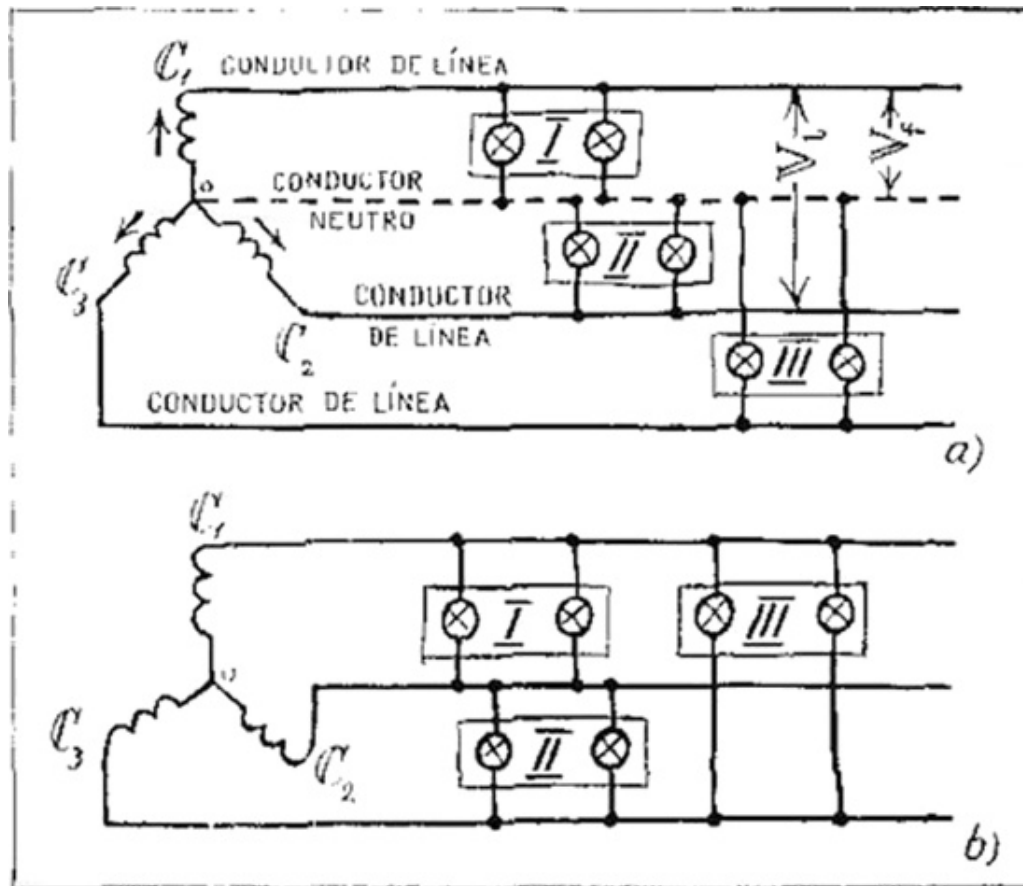


Figura 4.7

El generador se conecta a los receptores de energía mediante cuatro conductores: tres conductores «de línea» que van desde los comienzos de los devanados C_1 , C_2 y C_3 , y un conductor neutro que va del punto neutro del generador. Dicho sistema se denomina con cuatro conductores o hilos.

La tensión entre el punto neutro y el comienzo de la fase se llama fásica. La tensión entre los comienzos de los devanados se denomina de línea. Estas tensiones están relacionadas por medio de la expresión

$$U_l = \sqrt{3}U_f$$

Si las cargas (I, II, III) de todas las tres fases son idénticas, la corriente en el conductor neutro es igual a cero. En este caso es posible prescindir del conductor

neutro, pasando al sistema trifilar. En la fig. 4.7, b se representa el esquema de conexión en estrella para tal caso.



Mijaíl Osipovich Dolivo-Dobrovolski (1862 - 1949) relevante científico o ingeniero ruso, creador del sistema de corriente trifásica que sirve de fundamento para toda la electrotecnia contemporánea. Elaboró todos los elementos, sin excepción, de los circuitos trifásicos de corriente alterna. En 1888 construyó el primer generador trifásico de corriente alterna con el campo magnético giratorio.

La conexión en triángulo también permite la conducción trifilar. En este caso el extremo de cada devanado está unido con el comienzo del siguiente de modo que éstos forman un triángulo cerrado. Los conductores de línea están conectados a los vértices del triángulo. Aquí, la tensión de línea es igual a la fásica, mientras que las corrientes están relacionadas mediante la expresión

$$I_L = \sqrt{3} I_f$$

Los circuitos trifásicos presentan las siguientes ventajas: transmisión más económica de energía en comparación con los circuitos monofásicos y la posibilidad de obtener las tensiones, la fásica y la de línea, en una instalación.

El generador de corriente alterna que hemos descrito pertenece a la clase de máquinas eléctricas sincrónicas. Este nombre lo llevan las máquinas cuya frecuencia de rotación del rotor coincide con la de rotación del campo magnético del estator. Los motores sincrónicos son productores principales de energía y tienen algunas variedades constructivas según sea el método de puesta en rotación de su rotor.

El lector puede preguntar: ¿si se emplea la denominación de «máquinas sincrónicas», por consiguiente, deben existir también las asincrónicas? ¡Justamente! Pero éstas se utilizan como motores y hablaremos de ellas en el siguiente párrafo. En el mismo párrafo nos detendremos también en la cuestión de por qué gira el campo magnético en una máquina trifásica de corriente alterna.

Motores eléctricos

Más de la mitad de toda la energía eléctrica producida se transforma con la ayuda de los motores eléctricos en energía mecánica que se consume para las distintas necesidades de la industria, agricultura, transporte y fines domésticos. La mayor difusión la obtuvo el motor asincrónico, sencillo, seguro, barato y poco exigente en el servicio, inventado en 1880 por el mismo talentoso ingeniero Dolivo-Dobrovolski. Este motor, hasta la fecha, conserva sus rasgos principales. El motor asincrónico se utiliza para el accionamiento de diferentes máquinas-herramientas, en las instalaciones de bombeo y compresión, en la maquinaria de forja y prensado, así como para el equipo elevador y de transporte y para otros mecanismos.

Se debe considerar como prototipo del motor asincrónico el modelo creado por Domingo Francisco Arago (1786 - 1853). En 1824, en la Academia de Ciencias de París, Arago hizo exposición de un fenómeno que llamó «magnetismo de rotación». Demostró que un disco de cobre se pone en rotación al introducirlo en el campo de un imán permanente que está girando. Dolivo-Dobrovolski hizo uso brillante de esta idea, combinándola con las particularidades del sistema trifásico de corrientes que

da la posibilidad de obtener la rotación del campo magnético sin cualesquiera dispositivos adicionales.

Examinemos los esquemas de la fig. 4.8. Para la máxima simplificación aquí vienen representadas tres espiras (en la realidad, por supuesto, la máquina utiliza bobinas con gran número de espiras).

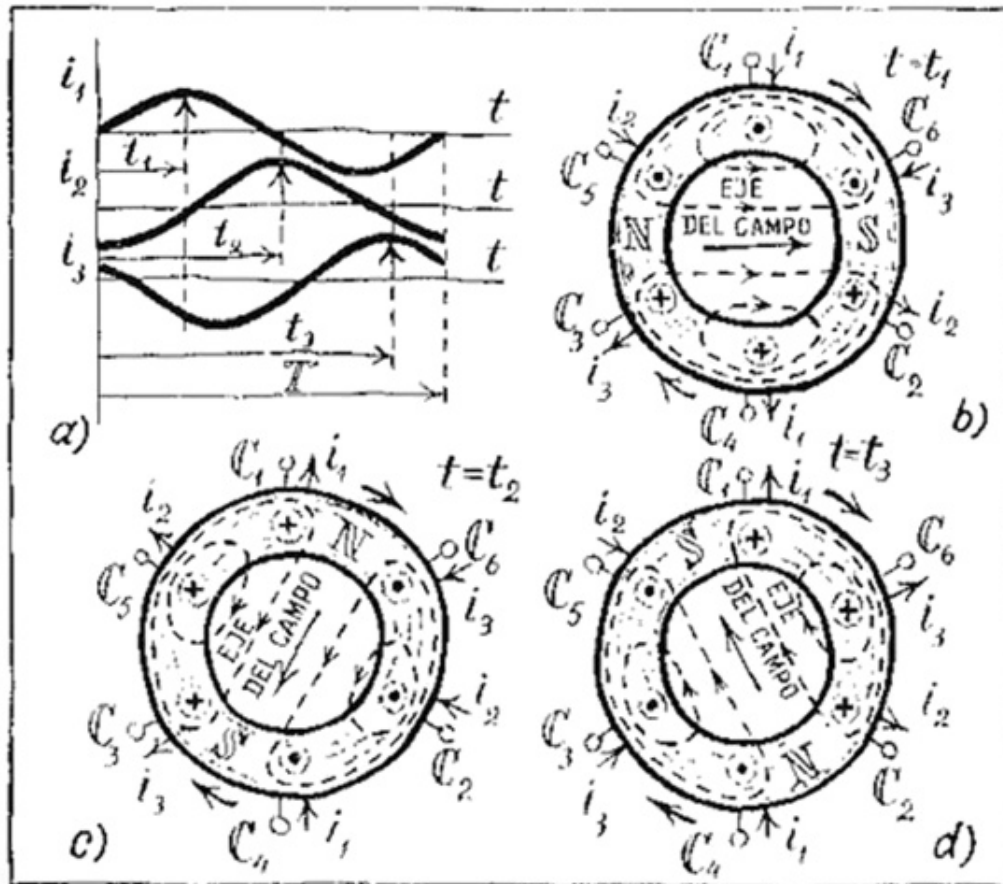


Figura 4.8

El punto negro y la crucecita señalan la entrada y la salida de la corriente en cada espira en cierto momento determinado de tiempo. Estas tres espiras forman entre sí ángulos de 120° . En la fig. 4.8 a, se dan las relaciones fásicas de tres corrientes i_1 , i_2 e i_3 que fluyen por las espiras. A nosotros nos interesa el campo magnético resultante de estas tres bobinas. En la fig. 4.8 b, se representan las líneas de inducción del campo resultante para el momento t_1 (entradas en C_5 , C_1 y C_6 y en las

figs. 4.8 c y d, se han hecho los mismos esquemas para los momentos de tiempo t_2 t_3 .

De este modo, vemos que el campo que nos interesa gira (fíjense en las posiciones de las crucecitas), ¡gira, en el pleno sentido de la palabra! El eje del campo en el centro del sistema se dispone según el eje de aquella espira (fase) cuya corriente es máxima en el momento dado de tiempo.

La figura que acabamos de analizar permite formar una idea de cómo está distribuido el devanado trifásico de corriente alterna en el estator del motor asincrónico trifásico.

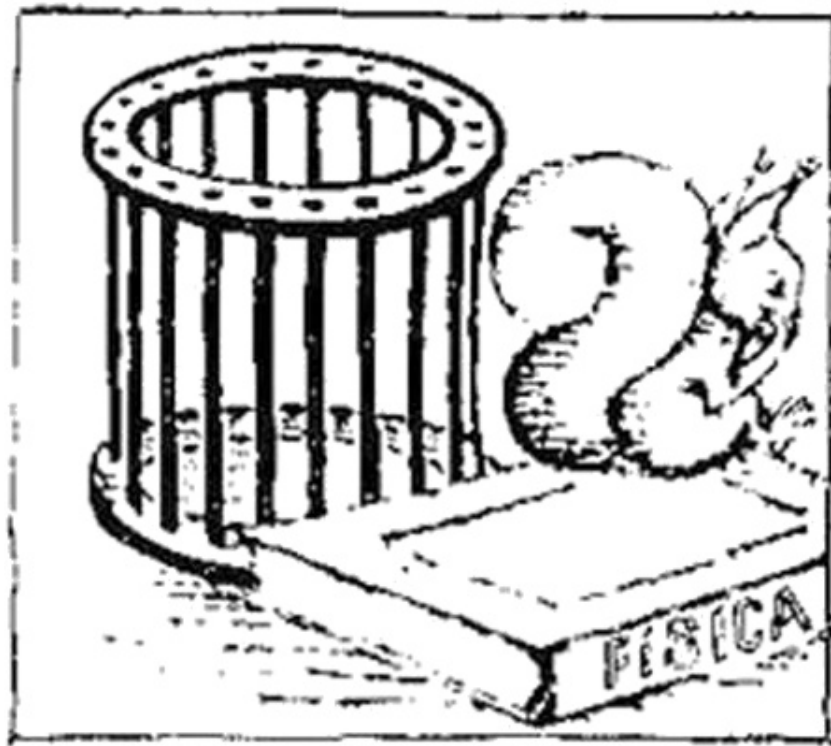


Figura 4.9

El rotor (fig. 4.9) que es arrastrado hasta ponerse en movimiento por el campo magnético giratorio es un rotor cortocircuitado, es decir, no vemos ni los comienzos ni los extremos del devanado. Este rotor se parece a una jaula de ardilla: una fila de barras y anillos que las cierran. Compáreselo con la máquina de corriente continua. ¡Cuánto más sencillo es! Conectamos al estator la corriente alterna trifásica. En la máquina se engendra un campo magnético giratorio. Las líneas del campo

magnético cortan las barras del rotor induciendo en éstas las corrientes. Como resultado de interacción de la barra por la cual circula la corriente y del campo magnético el rotor comienza a girar con una velocidad próxima a la del campo, pero no la alcanza. Debe suceder precisamente así, ya que, de lo contrario, no hubiera tenido lugar la intersección por las barras del rotor de las líneas de inducción del campo giratorio del estator ni se hubiera producido la rotación de la jaula de ardilla. Esta es la razón de que las máquinas de este tipo se denominan asincrónicas. El retardo del rotor se llama deslizamiento.

Los motores asincrónicos abarcan un gran diapasón de potencias: desde partes de vatio hasta centenares de kilovatios. Existen también motores asincrónicos más potentes: hasta 6000 kW para la tensión de 6000 V,

Las micromáquinas asincrónicas se utilizan en los dispositivos de automática como servomotores para transformar la señal eléctrica que reciben en traslación mecánica del árbol, así como en calidad de generadores tacométricos que transforman la rotación mecánica en señal eléctrica.

También pueden servir de motores eléctricos las máquinas sincrónicas que hemos examinado autos y las máquinas de corriente continua. Este hecho deriva del principio evidente de reciprocidad de la máquina eléctrica el cual consiste en que cualquier máquina eléctrica puede trabajar tanto como generador, como en calidad de motor.

Por ejemplo, en la composición del centro hidroeléctrico de Kiev en el río Dniéper entra una estación hidroacumuladora equipada con grupos reversibles que pueden trabajar ya sea como bombas, o bien, como turbinas. Cuando en el sistema energético hay energía eléctrica en exceso, las bombas-turbinas elevan el agua al depósito acumulador. En este caso, la máquina sincrónica que forma parte del grupo trabaja como motor. Cuando el consumo de energía eléctrica llega a ser máximo el grupo hace «entrar en servicio» el agua acumulada.

En las plantas metalúrgicas, minas e instalaciones frigoríficas los motores sincrónicos ponen en funcionamiento bombas, compresores, ventiladores y otros mecanismos que trabajan a velocidad invariable. En los dispositivos automáticos se emplean ampliamente micromotores sincrónicos con una potencia desde partes de vatio hasta varias centenas de vatios. Por cuanto la frecuencia de rotación de estos

motores está ligada, rígidamente, con la frecuencia de la red de alimentación, dichos motores se utilizan en los lugares donde se requiere mantener la velocidad constante de rotación: en los mecanismos de relojes eléctricos, en los mecanismos de arrastre de la cinta de los instrumentos auto registradores y aparatos cinematográficos de proyección, en los aparatos de radio, dispositivos de programación, así como en los sistemas de comunicación sincrónica cuando la velocidad de rotación de los mecanismos se gobierna por la variación de la frecuencia de la tensión de alimentación.

Por su estructura de principio el motor de corriente continua no se diferencia en nada del generador de corriente continua. La máquina tiene un sistema inmóvil de polos cuyo devanado de excitación está conectado por medio de uno u otro procedimiento con el devanado del inducido (en serie o en paralelo). La máquina puede excitarse también de un manantial independiente de corriente. El inducido tiene un devanado distribuido por las ranuras que se conecta a la fuente de corriente continua. El motor, al igual que el generador, está provisto de un colector cuya destinación consiste en que «endereza» el momento rotatorio, es decir, obliga a la máquina girar prolongadamente en un solo sentido.

El motor de corriente continua con excitación en serie está especialmente apropiado para la tracción eléctrica, para las grúas y elevadores. En estos casos se requiere que, a grandes cargas, la frecuencia de rotación disminuya de un modo brusco, mientras que la tracción aumento en grado considerable. Precisamente estas propiedades las acusa el motor de corriente continua con excitación en serie.

Los primeros experimentos para la tracción eléctrica no autónoma fueron realizados en Rusia por Fiodor Apollonovich Pirotski (1845 - 1898). Ya en 1876 adaptó para la transmisión de energía eléctrica la vía ferroviaria de carriles común y corriente, y en agosto de 1880 realizó la puesta en marcha del tranvía eléctrico en la línea experimental de la zona del parque Rozhdestvenski del tranvía de caballo en Petersburgo. Como primer vagón eléctrico se tomó el vagón de dos pisos del tranvía de caballo a cuya carrocería se fijó un motor eléctrico.

El primer tranvía de Rusia, el de Kiev, fue inaugurado para el uso común en 1892. La alimentación de su motor eléctrico se efectuó del conductor de contacto superior. Cabe señalar, además, que la comisión para la construcción se avino a la existencia

del tranvía eléctrico sólo después de haberse cerciorado, mediante cálculos, en la ventaja técnica de la tracción eléctrica frente a la de caballo en las condiciones de pesado perfil de las calles de Kiev, que resultó ser superior a las fuerzas tanto de la tracción de caballo, como a la tracción a vapor.

Los primeros experimentos en el campo de «electro-navegación» fueron realizados por Boris Semiónovich Jacobi (1801 - 1874) que, en 1838, exhibió en el río Neva un bote eléctrico con cabida para catorce personas. El bote se ponía en movimiento por medio de un motor eléctrico de 550 W de potencia. Para la alimentación de este motor, Jacobi utilizó 320 pilas galvánicas. Fue la primera aplicación, en la historia, del motor eléctrico para fines de tracción.

En los últimos años en la prensa comenzó a aparecer la denominación «buque con propulsión turbo-eléctrica». El sentido de dicha denominación se aclara de una forma simple: en un buque de este tipo el vapor pone en movimiento potentes generadores de corriente continua y las hélices se emplazan en los árboles de los motores eléctricos. ¿No es ésta una complicación que sobra? ¿Por qué no se puede colocar la hélice directamente sobre el árbol de la turbina?

La cuestión radica en que la turbina de vapor desarrolla la potencia máxima sólo para un número de revoluciones estrictamente determinado. Turbinas potentes hacen 3000 r.p.m. Cuando la rotación se hace más lenta, la potencia disminuye. Si las hélices se hubieran encontrado directamente en el árbol de las turbinas, entonces, el barco provisto de semejante grupo propulsor habría presentado cualidades de marcha bastante malas. En cambio, el motor eléctrico de corriente continua tiene una característica de tracción ideal: cuanto mayores son las fuerzas de resistencia tanto mayor esfuerzo de tracción éste desarrolla, con la particularidad de que semejante motor puede entregar una potencia grande con un número de revoluciones pequeño en el momento de arranque.

De este modo, el generador y el motor de corriente continua intercalados entre la turbina y la hélice del buque con propulsión turbo-eléctrica hacen las veces de caja de cambio de velocidades automática progresiva y sumamente perfecta. Puede parecer que este sistema es algo engorroso, sin embargo teniendo en cuenta la gran potencia de los modernos buques con propulsión turbo-eléctrica, cualquier otro sistema sería igualmente voluminoso, pero menos seguro.

Existe otro camino para perfeccionar de una manera considerable la instalación de fuerza de un buque con propulsión turbo-eléctrica: es muy provechoso sustituir las voluminosas calderas de vapor por un reactor atómico. En este caso se alcanza una enorme economía en el volumen del combustible que es necesario tomar en el viaje. Se granjeó la fama mundial el primer rompehielos atómico soviético «Lenin». La instalación nuclear de fuerza de este buque con propulsión turbo-eléctrica asegura la navegación autónoma durante un periodo mayor que un año.

Los motores de corriente continua están instalados en locomotoras eléctricas de servicio de línea, en trenes eléctricos suburbanos, en vagones del tranvía y trolebuses. La energía para su alimentación se suministra a partir de estaciones eléctricas estacionarias. En la URSS, para la tracción eléctrica se utiliza la corriente continua y la corriente alterna monofásica de frecuencia industrial de 50 Hz. En las subestaciones de tracción han encontrado un amplio uso los rectificadores de silicio. Cuando se trata del transporte ferroviario la rectificación de la corriente puede tener lugar tanto en las subestaciones, como en los propios trenes eléctricos.

Capítulo 5

Campo electromagnético

Contenido:

- *Las leyes de Maxwell*
- *Modelos mecánicos de irradiación*
- *Dos aspectos del campo electromagnético*
- *Efecto fotoeléctrico*
- *Los experimentos de Hertz*
- *Clasificación de la radiación electromagnética*

Las leyes de Maxwell

Para los años cincuenta del siglo XIX se habían acumulado muchos datos sobre la electricidad y el magnetismo. Sin embargo, estos datos parecían sueltos, a veces contradictorios y, en todo caso, no se enmarcaban en un esquema armonioso.

No obstante, se conocían ya no pocas cosas. En primer término, los físicos sabían que las cargas eléctricas en reposo engendran el campo eléctrico, en segundo término, era de su conocimiento que las corrientes eléctricas crean campos magnéticos, y, en tercer término, se publicaron y obtuvieron reconocimiento universal los resultados de los experimentos de Faraday quien demostró que el campo magnético variable engendra corriente eléctrica.

No cabe duda que en aquella época una serie de científicos y, en primer lugar, Faraday, formaron la idea de que en el espacio que rodea las corrientes y cargas eléctricas se desarrollan ciertos acontecimientos. Este grupo de investigadores suponía que las fuerzas eléctricas y magnéticas se transmiten de un punto a otro. Había muchos intentos de representar en el papel un esquema similar al sistema de engranajes enlazados que demuestre patentemente en qué consiste el mecanismo de transmisión de energía eléctrica. Sin embargo, algunos hombres de ciencia propugnaban la teoría de «teleacción», suponiendo que no se da ningún proceso físico de transmisión de fuerzas eléctricas y magnéticas.



James Clerk Maxwell (1831 - 1879), célebre hombre de ciencia inglés, fundador de la electrodinámica teórica. Las ecuaciones de Maxwell describen el comportamiento de las ondas electromagnéticas y del campo electromagnético independientemente de su procedencia. Maxwell es fundador de la teoría electromagnética de la luz. De sus ecuaciones derivaba automáticamente el valor de la velocidad de propagación de la luz. De la teoría de Maxwell se infería la relación entre la constante dieléctrica y la permeabilidad magnética por una parte y el índice de refracción por otra, la ortogonalidad de los vectores eléctrico y magnético en la onda, así como la existencia de la presión de la luz. También es grande la aportación que hizo Maxwell a la teoría cinética de los gases. Le pertenece la deducción de la distribución de las moléculas del gas por las velocidades.

Sostenían que los conceptos del campo y de las líneas de fuerza deben considerarse tan sólo como imágenes geométricas a las cuales no se pone en correspondencia realidad alguna.

Como suele suceder a menudo en la historia de la ciencia, la verdad resultó encontrarse en cierto punto intermedio: se revelaron como inconsistentes las tentativas de reducir los fenómenos electromagnéticos a los movimientos de una clase especial de materia, o sea, del «éter», mas tampoco tenían razón aquellos

investigadores quienes sugerían que las interacciones electromagnéticas se transferían de una carga o corriente a otras instantáneamente.

El científico inglés James Clerk Maxwell (1831 - 1879) a la edad de 26 años, solamente, publicó su trabajo «Sobre las líneas de fuerza de Faraday». En realidad, este trabajo ya encerraba en sí las leyes que él descubrió. Sin embargo, necesitó varios años más para, abandonando las ideas mecánicas, formular las leyes del campo electromagnético en una forma que no requiriese ilustración gráfica. El propio Maxwell dijo al respecto: «Para el bien de los hombres con diferente mentalidad, la verdad científica debe presentarse en distintas formas y debe considerarse igualmente científica independientemente de si se presenta en la clara forma y los vivos colores de una ilustración física o tomando el aspecto sencillo y deslucido de una expresión simbólica».

Las leyes de Maxwell pertenecen a las leyes fundamentales, leyes universales de la naturaleza. Estas no se deducen por vía de razonamientos lógicos ni tampoco de cálculos matemáticos. Las leyes universales de la naturaleza son una sintetización de nuestro conocimiento. Las leyes de la naturaleza se descubren, se hallan... Para un historiador de ciencia y un psicólogo presenta gran interés seguir el camino de ocurrencias y revelaciones creativas que llevan a los hombres geniales al descubrimiento de las leyes de la naturaleza. Pero es un amplio tema para un libro especial. Y en cuanto a nosotros, no nos queda otra cosa que, junto con el lector, examinar cierto esquema de conjeturas consecutivas que conducen a las leyes de Maxwell.

¿De qué datos disponía Maxwell cuando planteó ante sí la tarea de expresar sucintamente en forma simbólica las leyes que rigen el comportamiento de los campos eléctricos y magnéticos?

En primer término, conocía que cualquier punto del espacio en las cercanías de la carga eléctrica puede caracterizarse por el vector de fuerza eléctrica (intensidad), y todo punto cerca de la corriente eléctrica, por el vector de fuerza magnética. ¿Pero son las cargas inmóviles las únicas fuentes del campo magnético? ¿Y son las corrientes eléctricas las únicas fuentes del campo magnético?

Maxwell da una respuesta negativa a estas dos preguntas y en su búsqueda de las leyes del campo electromagnético emprende el siguiente camino de conjeturas.

Faraday demostró que en un circuito de alambre atravesado por un flujo variable de líneas de inducción se engendra la corriente eléctrica. Pero la corriente aparece cuando sobre las cargas eléctricas actúa la fuerza eléctrica. Siendo así, es posible expresar la ley de Faraday también con la siguiente frase: en un circuito de alambre por el cual pasa un flujo magnético variable aparece el campo eléctrico.

Pero, ¿acaso es esencial el hecho de que el flujo magnético se abarca por un circuito de alambre? ¿No le da lo mismo al campo eléctrico dónde aparecer: en un conductor metálico o en el espacio vacío? Supongamos que es lo mismo. Entonces, también resulta justa la siguiente confirmación: cerca del flujo variable de las líneas de inducción aparece la línea cerrada de intensidad eléctrica.

Las dos primeras leyes de Maxwell concernientes al campo eléctrico ya están formuladas. Hacemos constar que el campo eléctrico se engendra por dos caminos: por cargas eléctricas (en este caso las líneas de intensidad tienen su comienzo en las cargas positivas y terminan en las negativas) y por el campo magnético variable (en este caso la línea de intensidad eléctrica está cerrada y abarca el flujo magnético cambiante).

Ahora pasamos a la búsqueda de las leyes referentes al campo magnético. Este último se engendra por las corrientes, y Maxwell lo conoce. La corriente continua sirve de manantial del campo magnético permanente, mientras que la corriente alterna forma el campo magnético alterno. Pero, es que la corriente alterna se crea en el conductor por el campo eléctrico alterno. ¿Y si no se da un conductor, sino un campo eléctrico variable que existe en el vacío? ¿Acaso no es lógico suponer que cerca del flujo variable de líneas de intensidad aparece la línea cerrada de inducción? El cuadro es sumamente atrayente por su simetría: el flujo magnético variable engendra el campo eléctrico, y el flujo eléctrico variable engendra el campo magnético.

De este modo, a las dos leyes que atañen al campo eléctrico se añaden otras dos que rigen el comportamiento del campo magnético. El campo magnético no tiene fuentes (no hay cargas magnéticas): ésta es la tercera ley; el campo magnético se crea por las corrientes eléctricas y el campo eléctrico variable: y ésta es la cuarta ley.

Las cuatro leyes de Maxwell pueden presentarse en una forma extraordinariamente elegante por medio de ecuaciones matemáticas (fig. 5.1).

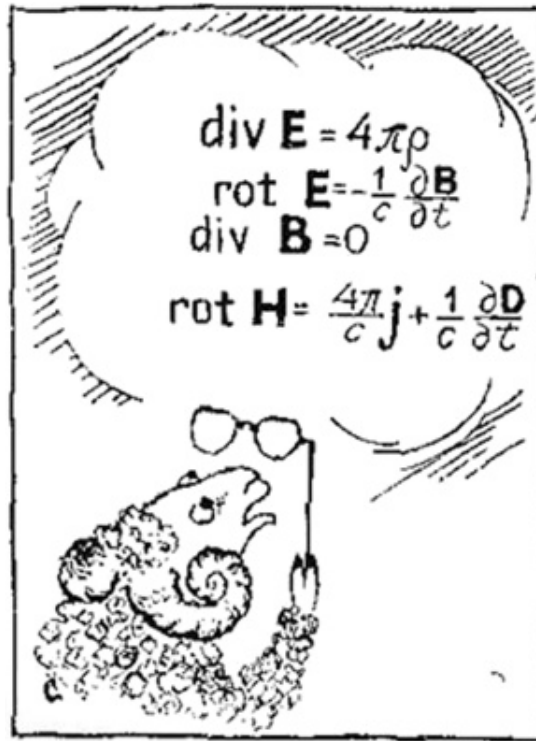


Figura 5.1

Siento no poder dar a conocer al lector cómo se plasma en estas ecuaciones el sentido físico. Se requieren conocimientos matemáticos serios.

Las leyes de Maxwell nos enseñan que no puede existir el campo magnético variable sin el campo eléctrico, ni el campo eléctrico alterno sin el campo magnético. Esta es la razón de que estos dos adjetivos no se separan por una coma. El campo electromagnético es un ente único.

Apartándonos de las cargas que son las fuentes del campo electromagnético, tenemos que tratar con la materia electromagnética, por decirlo así, en forma pura. No es obligatorio examinar los flujos de líneas de fuerza. Las leyes de Maxwell pueden escribirse en la forma aplicable a un punto del espacio. Entonces su formulación es muy simple: en un punto donde varía en función del tiempo el vector eléctrico existe y también varía en función del tiempo el vector del campo magnético.

¿Y no es que todo lo expuesto resulta ser pura fantasía? - preguntará el lector. No olvide que la medición en un punto de los valores de los vectores - que cambian rápidamente - de los campos eléctrico y magnético es una tarea prácticamente irrealizable.

¡Una observación justa! Mas se juzga sobre la grandeza de las leyes de la naturaleza por los corolarios que de éstas derivan. Y los corolarios son innumerables. No voy a exagerar, en modo alguno, si digo que toda la electrotecnia y la radiotecnica se contienen en las leyes de Maxwell.

Sea como fuera, es indispensable contar sobre una deducción importantísima que se desprende de las ecuaciones de Maxwell. Por medio de cálculos absolutamente estrictos se puede demostrar que debe existir el fenómeno de radiación electromagnética.

Supongamos que en una porción limitada del espacio hay cargas y corrientes. En semejante sistema pueden tener lugar diversas transformaciones energéticas. Fuentes mecánicas o químicas crean corrientes eléctricas, las corrientes, a su vez, pueden poner en movimiento unos mecanismos y crear calor que se libera en los conductores. Calculemos las ganancias y las pérdidas. ¡El balance no coincide! El cálculo demuestra que cierta parte de energía de nuestro sistema se fue al espacio. ¿Puede la teoría decir algo sobre esta energía «irradiada»? Resulta que sí, puede. Cerca del manantial la solución de la ecuación tiene una forma compleja, en cambio, a unas distancias que superan sustancialmente las dimensiones del sistema «radiante», el cuadro se torna muy preciso y, lo que es lo principal, susceptible de comprobarse en el experimento.

A grandes distancias, se puede caracterizar la radiación electromagnética - así denominaremos aquel déficit energético que se crea en el sistema de las cargas en movimiento - en cada punto del espacio por la dirección de propagación. La energía electromagnética se desplaza en esta dirección a una velocidad de cerca de 300.000 km/s.

¡Esto valor se deduce de la teoría!

La segunda deducción de la teoría: los vectores eléctrico y magnético son perpendiculares a la dirección de propagación de la onda y son perpendiculares entre sí. Y, en tercer lugar, la intensidad de la radiación electromagnética (la

energía que corresponde a una unidad de superficie) disminuye inversamente proporcional al cuadrado de distancia.

Puesto que se conocía que la luz se propaga precisamente con la velocidad de 300.000 km/s calculada para la radiación electromagnética, así como había datos lo suficientemente exhaustivos acerca de la polarización de la luz, los cuales hacían pensar que la energía luminosa acusa ciertas propiedades «transversales», Maxwell llega a la conclusión de que la luz es una forma de radiación electromagnética.

Al cabo de unos diez años después de la muerte de Maxwell, a finales de los años 80 del siglo XIX, el relevante físico alemán Enrique Hertz (1857 - 1894) confirmó por vía experimental todas las deducciones de la teoría de Maxwell. Después de estos experimentos, las leyes de Maxwell se afirmaron por los siglos de los siglos como una de las piedras angulares, una de aquellas que se pueden contar con los dedos, sobre las cuales descansa el edificio de las ciencias naturales contemporáneas.

Modelos mecánicos de radiación

Los modelos mecánicos se contraponen a los matemáticos. Los primeros se pueden realizar por medio de bolitas, muelles, cuerdas, cordones de goma, etc. El modelo mecánico contribuye a que el fenómeno se haga «palmario». Al construir un modelo mecánico y mostrar su funcionamiento ayudamos a comprender el fenómeno, señalando: he aquí que esta magnitud se comporta a semejanza de aquel desplazamiento. No a cada modelo matemático, ni mucho menos, puede ponerse en correspondencia un modelo mecánico.

Antes de hablar sobre la radiación electromagnética cuya existencia real se establece por un sinnúmero de experimentos y se deduce, con lógica férrea, a partir de las ecuaciones de Maxwell, tenemos que decir algunas palabras acerca de los posibles modelos mecánicos de radiación.

Hay dos modelos de este tipo: el corpuscular y el ondulatorio.

Se puede fabricar un juguete que «radiará» en todas las direcciones flujos de partículas pequeñas: guisantes o semillas de amapola. Este será justamente el modelo corpuscular, ya que la palabra «corpúsculo» significa «partícula».

Una partícula que vuela con cierta velocidad y posee una masa debe comportarse de acuerdo con las leyes de la mecánica. Las partículas son capaces de chocar entre sí cambiando la dirección de su movimiento, no obstante, deben hacerlo obligatoriamente de tal modo que la colisión se subordine a las leyes de la conservación de la energía e impulso. Ciertos cuerpos pueden resultar impermeables para las partículas, en este caso las partículas deben reflejarse de aquéllos según la ley, *el ángulo de incidencia es igual al ángulo de reflexión*. Las partículas pueden absorberse por el medio. Si en un medio para las partículas es más fácil moverse que en el otro, no es difícil explicar el fenómeno de refracción. Al pasar por un orificio en una pantalla no transparente el flujo de partículas que tiene su origen en una fuente puntual debe migrar al interior de un cono. Realmente, es posible una disipación insignificante por cuanto una pequeña parte de partículas puede reflejarse de los bordes del orificio. Pero, por supuesto, estas «reflexiones» pueden ser tan sólo caóticas y no formarán ningún dibujo regular que saiga de los límites de la sombra geométrica.

El modelo ondulatorio se representa, por lo común, mediante un baño de agua. No es difícil hacer oscilar de una manera periódica el agua en un punto cualquiera. Desde este punto como de una piedra arrojada al agua se extienden círculos. La superficie ondulada del agua se advierte a simple vista. La energía se difundirá por todas las partes y una astilla que se encuentra lejos comenzará a oscilar con la frecuencia del punto al cual suministramos la energía.

Es algo más difícil hacer patentes las oscilaciones sonoras. Sin embargo, es posible poner experimentos absolutamente convincentes los cuales demostrarán que la propagación del sonido es la transferencia de un punto a otro de desplazamientos mecánicos del medio.

La propagación del sonido y de la luz, al igual que cualquier otra radiación electromagnética, así como la propagación de las partículas tales como los electrones o los neutrones se puede interpretar mediante los dos modelos. Se dan casos en que ambos modelos explican los fenómenos. A veces, estas explicaciones resultan ser completamente exactas y, a veces, solamente uno de los modelos se presenta como adecuado.

Depende de muchas circunstancias en qué casos es necesario recurrir a uno u otro cuadro patente y qué precisión se puede alcanzar empleando tal o cual modelo. Lo más esencial es la relación entre la longitud de onda de la radiación y el tamaño de los orificios u obstáculos que se encuentran en el camino del flujo de radiación. También es importante el carácter del manantial y del receptor de energía. ¿Se permite considerarlos puntuales? ¿Qué distancias separan el manantial de radiación, el obstáculo y el receptor?

El modelo corpuscular exige que la energía se propague por líneas rectas, o sea, por rayos. No se admite que los rayos se encorven.

Si la longitud de onda de la radiación es mucho menor que el orificio, la radiación (el sonido, la luz, etc.) se propaga por rayos. Se puede demostrar que así debe comportarse la onda; se sobreentiende que así se comporta el flujo de partículas. En semejantes casos ambos modelos son buenos.

Sin embargo, también en los casos contrarios ambos modelos pueden resultar idóneos en igual grado. Si examinamos el paso del rayo luminoso a través de un cristal, o del sonido a través de una red con mallas cuyo tamaño es del orden de milímetros no descubriremos desviaciones del curso de los rayos respecto a la línea recta. En estos casos la longitud de onda de la radiación supera considerablemente las dimensiones de las ranuras. Las ondas «dejan inadvertidos» los obstáculos.

También es fácil demostrar tanto por medio de razonamientos, como mediante experimentos en baño de agua que la ley de la reflexión desde las paredes cuyas irregularidades son menores que la longitud de onda se observa para el modelo ondulatorio.

El lector conoce perfectamente cómo refleja el sonido o cualquier otra onda una superficie lisa plana. Interesantes problemas surgen en los casos cuando la superficie reflectora tiene forma encorvada.

He aquí uno de semejantes problemas. ¿Cuál debe ser la superficie para volver a concentrar la onda que partió de una fuente puntual en un solo punto? La forma de la superficie reflectora debe ser tal que los rayos que inciden sobre esta partiendo de un mismo punto pero formando distintos ángulos se reflejen concentrándose otra vez en un punto. ¿Qué superficie debo ser ésta?

Recordemos al lector las propiedades de una admirable curva que lleva el nombre de elipse. La distancia desde un foco de la elipse hasta cualquier punto de la curva más la distancia desde otro foco hasta este mismo punto es la misma para todos los puntos de la elipse. Figúrense que la elipse gira alrededor del diámetro mayor. La curva que gira barrerá una superficie llamada elipsoidal. (La forma del elipsoide recuerda un huevo.) La elipse posee la siguiente propiedad. Si trazamos un ángulo que se apoya en uno de los puntos y cuyos lados pasan a través de los focos de la elipse, la bisectriz de este ángulo será normal a la elipse. En consecuencia, si una onda o un flujo de corpúsculos salen de un foco de la elipse, resultará que, al reflejarse de la superficie de ésta los mismos llegarán al otro foco.

Para las ondas sonoras la superficie del techo es lisa. Y, si el techo resulta abovedado, en el local puede observarse un caso especial de reflexión del sonido: por cuanto la bóveda por su forma se aproxima a una superficie elipsoidal, el sonido que ha salido de uno de sus focos llegará al otro. Esta propiedad de las superficies abovedadas se conocía ya en la antigüedad.

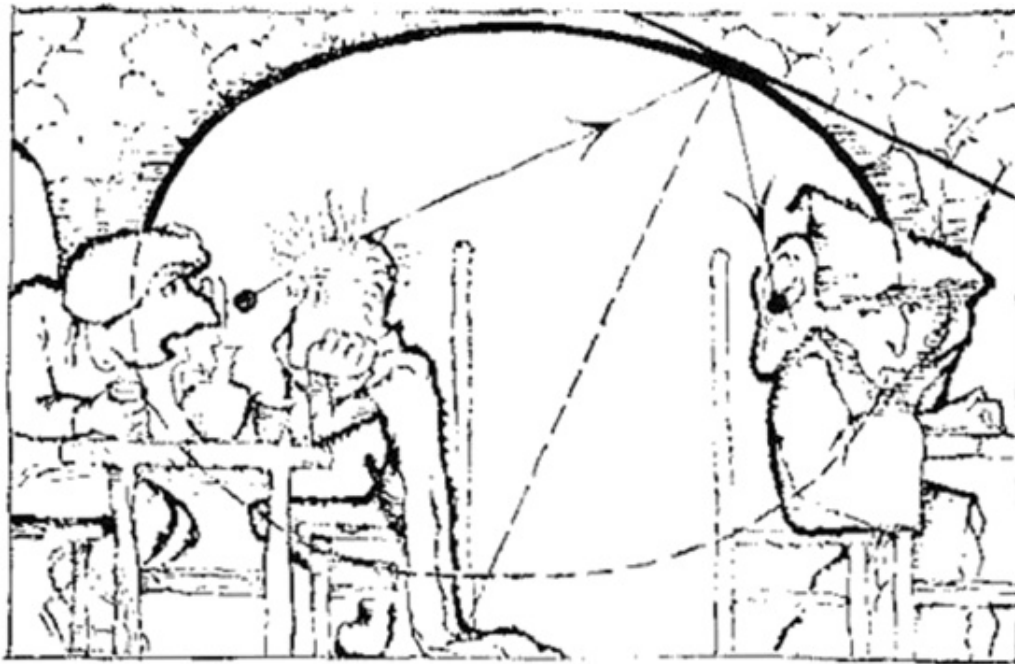


Figura 5.2

En la Edad Media, en el período de la inquisición, dicha propiedad se utilizó para escuchar a hurtadillas las conversaciones. Dos hombres que con voz velada, confiaban uno a otro sus pensamientos, ni siquiera sospechaban que, sigilosamente, les escuchaba el monje que parecía dormitar en otro rincón de la posada (fig. 5.2).

Tanto el modelo corpuscular, como el ondulatorio son de la misma forma convenientes para explicar este fenómeno. Sin embargo, el modelo ondulatorio es incapaz de explicar los fenómenos del tipo de colisiones de las bolas de billar.

Por otra parte, existen algunos hechos importantísimos cuya explicación es superior a las fuerzas del modelo corpuscular.

En primer término, se trata de la interferencia, o sea, de la adición cuando la suma puede resultar menor que los sumandos o, incluso, igual a cero. Si dos ondas llegan a un punto y se suman, el papel cardinal lo desempeña la diferencia de sus fases en este punto.

Si la cresta de una onda corresponde a la cresta de otra onda, entonces, las ondas se suman. En cambio, si la cresta de una onda coincide con la depresión de la otra y si en este caso las amplitudes de las ondas son idénticas, entonces, la adición llegará a... cero: las ondas concurridas en un punto se extinguirán recíprocamente. Al superponerse un campo ondulatorio sobre el otro, en unos sitios tendrá lugar su adición aritmética, y en otros, sustracción. En ello, de una forma precisa, consiste el fenómeno de interferencia. He aquí el primer fenómeno absolutamente imposible de interpretarse en el lenguaje de los flujos de partículas. En caso de comportarse la radiación como un flujo de guisantes, los campos superpuestos deberían intensificar uno a otro siempre y en todas las partes.

El segundo fenómeno importante es la difracción, es decir, el contorno de los obstáculos. El flujo de partículas no puede comportarse de tal manera, pero la onda debe proceder precisamente así. En las escuelas el fenómeno de difracción se muestra excitando ondas en un recipiente lleno de agua. En el camino que sigue la onda se coloca un tabique con orificio y a simple vista se deja ver cómo se contornea el ángulo. La causa de esta conducta es completamente natural. Es que en el plano del orificio las partículas de agua se han puesto en estado de oscilación. Cada punto que se encuentra en el plano del orificio da origen a una onda con el

mismo derecho que el manantial primario de radiación. Y no hay nada que impida a esta onda secundaria «doblar la esquina».

Los fenómenos de interferencia y de difracción se representan sin dificultad si se observa una condición: la longitud de onda debe ser conmensurable con las dimensiones del obstáculo o del orificio. En el siguiente libro precisaremos esta condición y hablaremos más detalladamente sobre la difracción y la interferencia.

Mientras tanto, nos detendremos en la variación de la frecuencia de la onda percibida por el observador estando en movimiento la fuente de radiación. Ya en los albores de la física teórica Cristian Doppler (1803 - 1853) demostró el hecho de que este fenómeno es una consecuencia indispensable del modelo ondulatorio.

Realicemos la deducción de la fórmula de Doppler que nos servirá más tarde. Para que el cuadro sea más patente supongamos que un automóvil va aproximándose a una orquesta en movimiento. El número de condensaciones del aire que, en una unidad de tiempo, alcanzan el oído del chófer será mayor que cuando el coche permanezca inmóvil en una relación $(c + u)/u$, donde c es la velocidad de propagación de la onda, y u , la velocidad relativa del manantial y del receptor de la onda. Por consiguiente,

$$v' = v (1 + u/c)$$

Esto significa que la frecuencia percibida v' se acrecienta para el acercamiento del automóvil y de la orquesta (el tono del sonido es más agudo, $u > 0$) y disminuye cuando éstos se alejan (el tono del sonido es más grave, $u < 0$). Adelantándonos podemos decir que para la onda luminosa esta deducción suena así: al alejarse la fuente tiene lugar el «desplazamiento rojo». El lector apreciará la importancia de esta deducción cuando hablemos sobre las observaciones de los espectros de las lejanas estrellas.

Desde tiempos remotos hasta los años 20 del siglo XX, los pensadores discutían a menudo el problema de si tal o cual transmisión de energía tiene carácter ondulatorio o corpuscular. La experiencia ha demostrado que cualquier irradiación presenta dos aspectos. Y únicamente la combinación de estos dos aspectos refleja la realidad de una forma fidedigna. La teoría elevó este hecho al rango de la ley

fundamental de la naturaleza. Mecánica ondulatoria, mecánica cuántica, física cuántica, todas éstas son denominaciones equivalentes de la teoría moderna del comportamiento de los campos y las partículas.

Dos aspectos del campo electromagnético

En algunos fenómenos la radiación electromagnética se comporta a semejanza de una onda, y en otros, como un flujo de partículas.

En este sentido las leyes de Maxwell acusan un «defecto». Nos presentan tan sólo el aspecto ondulatorio de la radiación electromagnética.

En plena concordancia con los experimentos la solución de las ecuaciones de Maxwell conduce a la conclusión que siempre es posible concebir la radiación electromagnética como suma de ondas de diferentes longitudes e intensidades. Si el sistema radiador representa una corriente eléctrica de frecuencia estrictamente fijada, entonces, la radiación será una onda «monocromática» (de un solo color).

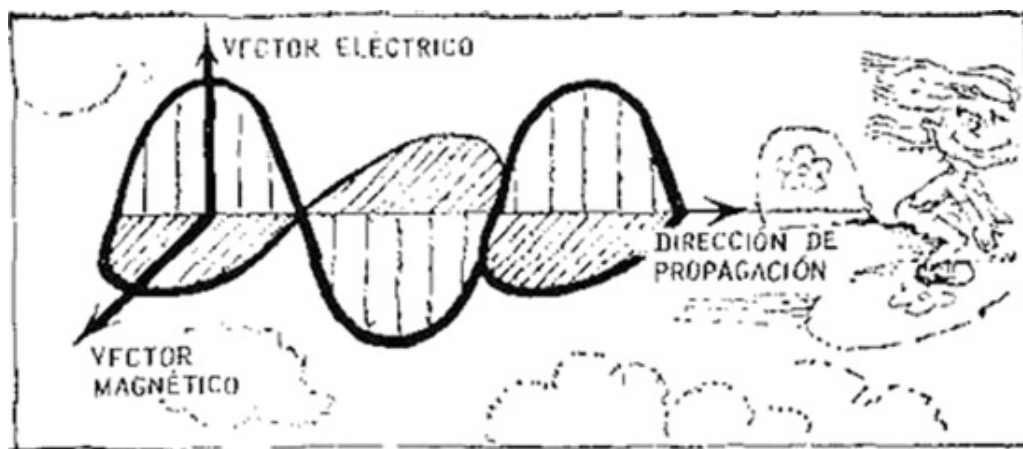


Figura 5.3

La onda electromagnética se ilustra en la fig. 5.3. Con el fin de imaginarnos los cambios que se operan en el espacio durante la propagación de la energía electromagnética, hay que «extender» nuestro cuadro como un todo rígido en dirección del eje de abscisas.

Dicho cuadro es el resultado de la solución de las ecuaciones de Maxwell. Precisamente aquel nos permite hablar sobre las ondas electromagnéticas. Sin embargo, al hacer uso de este término y recurrir a la analogía entre la onda

electromagnética y la onda que se propaga en el agua debido a echarle una piedra, conviene a proceder con suma precaución. Los cuadros patentes, con facilidad, pueden inducirnos en un error. La onda en el agua no es sino un modelo de la onda electromagnética. Esto significa que solamente en algunos aspectos las ondas electromagnéticas y las en el agua se comportan de una forma idéntica.

Bueno, pero, ¿no es cierto que la onda electromagnética representada en la fig. 5.3 y la onda del mar que ora hace subir, ora descender una astilla arrojada al agua se parecen como dos gotas de agua? ¡Nada semejante! Reflexione bien sobre la esencia del dibujo. ¡Por el eje vertical está marcado el vector del campo eléctrico y no, de ningún modo, el desplazamiento espacial!

Cada punto en el eje horizontal muestra que en el caso de colocar en el punto una carga eléctrica, sobre ésta habría ejercido su acción la fuerza representada por el valor de la ordenada. Hablando con propiedad, durante la migración de la onda electromagnética no hay cosa alguna que abandone su lugar. Mientras tanto, hasta para las oscilaciones muy lentas, en la práctica es absolutamente imposible realizar un experimento, que, de una forma palmaria, hubiera demostrado al lector cómo cambia el valor de la onda electromagnética en uno u otro punto.

De este modo resulta que las ideas sobre la onda electromagnética revisten carácter teórico. Hablamos con seguridad sobre la existencia de la onda electromagnética por la sencilla razón de que escuchamos la radio. No dudamos, ni mucho menos, de que la onda electromagnética posee una frecuencia determinada debido a que para la recepción de tal o cual estación es necesario sintonizar el aparato receptor a una frecuencia determinada. Estamos seguros de que el concepto de longitud es aplicable a la onda electromagnética no sólo a causa de que podemos medir la velocidad de la onda y calcular la longitud empleando la ecuación $c = \nu\lambda$ que relaciona la frecuencia, la longitud de onda y la velocidad de su propagación, sino también porque estamos en condiciones de juzgar sobre la longitud de la onda electromagnética estudiando el fenómeno de dirección, es decir, de contorno de los obstáculos, con la particularidad de que los principios de esta medición son los mismos que para las ondas que se propagan en agua.

La causa de que es absolutamente necesario prevenir al lector para que no procure formarse una idea palmaria acerca de la onda electromagnética reside en que, como

ya hemos dicho al principio del párrafo, la radiación electromagnética «se parece» no sólo a la onda, sino también, en una serie de casos, nos «recuerda» el comportamiento de un flujo de partículas. Y representarnos en nuestra imaginación algo que, simultáneamente, se parece al flujo de partículas y a la onda, eso sí que resulta completamente imposible. Se trata de procesos físicos los cuales no pueden dibujarse con tiza en el encerado.

Esta circunstancia no significa, claro está, que no podemos conocer, de un modo exhaustivo, el campo electromagnético. En este caso se debe tener presente que los cuadros que nos ofrecen una representación patente no son sino material didáctico, un método para retener mejor en la memoria las leyes. Y no se debe olvidar que el cuadro ondulatorio es tan sólo el modelo de la radiación electromagnética. ¡Y no más! En los casos convenientes hacemos uso de este modelo, pero no nos debe asombrar, en modo alguno, que en una u otra ocasión dicho modelo sólo nos inducirá en error.

Análogamente, tampoco se puede observar siempre el aspecto corpuscular del campo electromagnético.

Sea como sea, es más simple, (desde luego, es mejor decir, fue más simple anteriormente) observar el aspecto corpuscular de la radiación electromagnética para los casos de ondas cortas. En la cámara de ionización y en otros aparatos análogos se puede observar el choque del cuanto de la radiación electromagnética con el electrón o con otra partícula. La colisión puede transcurrir como el encuentro de bolas de billar. Se excluye completamente la posibilidad de comprender tal comportamiento recurriendo a la ayuda del aspecto ondulatorio de la radiación electromagnética.

Examinemos el nacimiento de la radiación electromagnética en el lenguaje de la teoría de Maxwell. El sistema de cargas oscila con cierta frecuencia. Al ritmo de estas oscilaciones varía el campo electromagnético. La velocidad de propagación de este campo igual a 300.000 km/s dividida por la frecuencia de oscilación y nos proporciona el valor de la longitud de onda de la radiación.

Si pasamos al lenguaje de la física cuántica el mismo fenómeno se describirá ya de la siguiente manera. Existe un sistema de cargas para el cual es característico un sistema de niveles discontinuos de energía. Debido a cierta causa este sistema llegó

al estado de excitación, no obstante, su vida en este estado duró un lapso corto, pasando al nivel más bajo. La energía que se liberó en esto caso

$$E_2 - E_1 = h\nu$$

se radia en forma de partícula que lleva el nombre de fotón. La constante h ya nos es conocida. Es la misma constante de Planck.

Si los niveles de energía del sistema están dispuestos muy cerca unos de otros, el fotón posee pequeña energía, pequeña frecuencia y, por consiguiente, una gran longitud de onda. En este caso el aspecto corpuscular cuántico del campo electromagnético es poco perceptible, poniéndose de manifiesto tan sólo en los fenómenos de absorción relacionados con las variaciones sumamente pequeñas de energía de los electrones o de los núcleos atómicos (resonancia magnética). Para ondas de gran longitud no se logra observar colisiones entre el fotón y las partículas, similares al impacto de las bolas de billar.

Exponemos de una manera sucinta los hechos que, por decirlo así, pusieron a los físicos entre la espada y la pared, obligándoles a reconocer que la teoría ondulatoria (en la que durante muchos decenios se había creído como en una verdad plena y exhaustiva) es incapaz de explicar todos los hechos concernientes a los campos electromagnéticos. Estos hechos son muy numerosos, mas por ahora nos limitamos al fenómeno que lleva el nombre de efecto fotoeléctrico. Después de que el lector se ponga de acuerdo de que el cuadro del campo electromagnético no puede crearse sin el aspecto corpuscular, nos dirigiremos a los admirables experimentos de Hertz en cuyo terreno maduró toda la radiotecnica, y señalaremos de qué modo el aspecto ondulatorio del campo electromagnético fue trazado no sólo en sus rasgos generales, sino también en los detalles.

Efecto fotoeléctrico

El sonoro y bello nombre de «fotón» apareció algo más tarde que el producto de la constante de Planck h por la frecuencia de la onda electromagnética ν . Como hemos dicho con anterioridad, la transición de los sistemas de un estado de energía al otro viene acompañada con la absorción o radiación de una porción de energía $h\nu$. A esta

conclusión llegó en la divisoria entre el siglo XX y el siglo XIX el relevante físico alemán Max Planck. El mismo demostró que sólo por medio de este procedimiento se lograba explicar la radiación de los cuerpos incandescentes. Sus razonamientos se referían a las ondas electromagnéticas obtenidas por un procedimiento no radiotécnico. En aquel entonces todavía no se había demostrado ni universalmente reconocido que aquello que era válido para la luz era cierto también para las ondas radioeléctricas, aunque las leyes de Maxwell indicaban con toda certidumbre que entre las ondas radioeléctricas y otras ondas electromagnéticas, incluyendo la luz, no hay ninguna diferencia de principio. La comprensión y las demostraciones experimentales de la validez universal de la afirmación de Planck llegaron más tarde.

En el trabajo de Planck se trataba de la emisión de la luz por porciones, es decir, por cuantos. Sin embargo, en este trabajo no se señalaba que el carácter cuántico de la radiación hace inevitable la introducción en el análisis del aspecto corpuscular del campo electromagnético. Sí, se decía en aquel tiempo, el campo se radia por porciones, mas la porción es cierto tren de ondas.

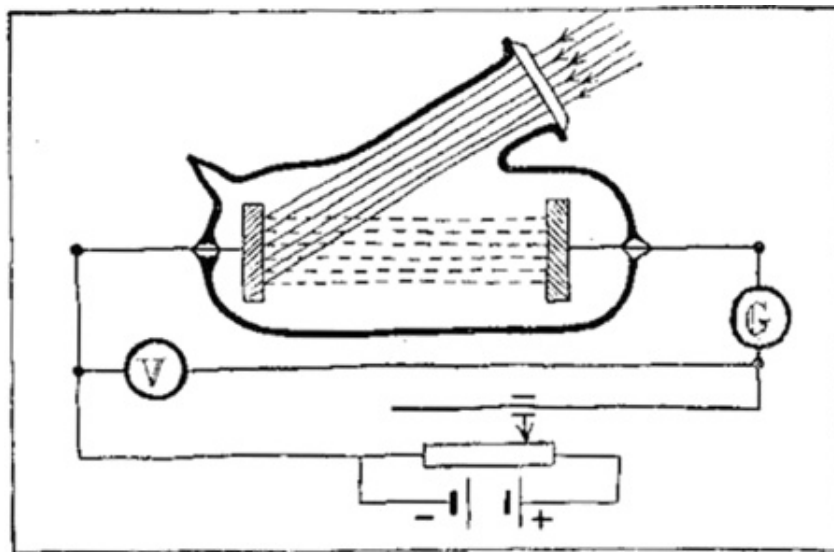


Figura 5.4

El paso más importante, es decir, el reconocimiento del hecho de que la porción emitida de energía $h\nu$ es la energía de la partícula la cual, de un modo inmediato, fue bautizada con el nombre de fotón, este paso lo hizo Einstein quien demostró que

únicamente valiéndose de las ideas corpusculares era posible explicar el fenómeno del efecto fotoeléctrico, o sea, la expulsión de los electrones a partir de los cuerpos sólidos por impacto de la luz.

En la fig. 5.4 se representa el esquema que a finales del siglo XIX dio origen al estudio detallado del fenómeno que recibió el nombre de efecto fotoeléctrico externo.

Por lo visto, Enrique Hertz, en 1888, indicó por primera vez que la luz afecta de cierto modo los electrodos del tubo de vacío. Trabajando simultáneamente, Svante Arrhenius (1859 - 1927). Guillermo Hallwachs (1850 - 1922), Augusto Righi (1850 - 1920) y el notable físico ruso Alexandr Grigóricvich Stolólov (1839 - 1896) demostraron que la iluminación del cátodo daba lugar a la aparición de la corriente. Si un tubo mostrado en la figura (este tubo se denomina célula fotoeléctrica) no se aplica la tensión, entonces, sólo una parte insignificante de electrones que la luz arranca del cátodo llegará al electrodo opuesto. Una débil tensión «instigadora» («menos» en el fotocátodo) aumentará la corriente. Al fin y al cabo, la corriente alcanzará la saturación: todos los electrones (cuyo número a la temperatura dada es completamente determinado) llegan al ánodo.

La intensidad de la corriente fotoeléctrica es estrictamente proporcional a la intensidad de la luz. Esta última se determina de modo unívoco por el número de fotones. En el acto a uno se le ocurre (y los cálculos rigurosos, así como los experimentos confirman esta idea) que un fotón expulsa de la sustancia un electrón.

La energía del fotón se consume para arrancar el electrón al metal y comunicarle velocidad. Precisamente de esta manera se entiende la ecuación que por primera vez fue escrita por Einstein (1905). De aquí esta ecuación:

$$h\nu = \frac{mv^2}{2} + A$$

donde A es el trabajo de salida.

La energía del fotón debe ser, en todo caso, mayor que el trabajo de salida de los electrones a partir del metal.

Y este hecho significa que para el fotón de cada energía (y la energía está relacionada unívocamente con la «cromaticidad») existe su frontera del efecto fotoeléctrico.

Las células fotoeléctricas que utilizan el efecto fotoeléctrico externo descrito por nosotros están difundidas ampliamente. Estas encuentran aplicación en los relés fotosensibles, en la televisión y en el cine sonoro.

La sensibilidad de las células fotoeléctricas puede elevarse al llenarlas de gas. En este caso, la corriente se intensifica debido a que los electrones rompen las moléculas neutras del gas, integrándolas a la corriente fotoeléctrica.

El efecto fotoeléctrico, aunque no aquel que acabamos de describir, sino el llamado efecto fotoeléctrico interno que tiene lugar en los semiconductores en el límite de la capa p - n, cumple un papel de primordial importancia en la técnica moderna. No obstante, para no interrumpir la exposición, aplacemos la conversación sobre el valor aplicado del efecto fotoeléctrico hasta el siguiente libro. Por ahora, teníamos la necesidad de examinar este fenómeno únicamente para demostrar inevitabilidad del reconocimiento de las propiedades corpusculares en el campo electromagnético.

Durante largo tiempo los fotones se encontraban en la situación de hijastros de la física que no sabían dónde meterse. La causa de ello residía en que la demostración de la existencia del fotón, así como la investigación de las leyes del efecto fotoeléctrico en 20 ó 30 años se adelantaron al proceso de afirmación de la física cuántica. Solamente a finales de los años 20 (del siglo XX), cuando estas leyes se establecieron, se pudo comprender por qué la misma constante numérica, o sea, la constante de Planck h , aparecía tanto en la fórmula de energía del fotón, como en la fórmula de la cual habíamos hablado en páginas anteriores y que determinaba los posibles valores del momento de impulso de las partículas.

El valor de esta constante se determina a partir de los más diversos experimentos. El efecto fotoeléctrico, el llamado efecto Compton (variación de la longitud de onda de los rayos X durante la dispersión), la aparición de una radiación para la aniquilación de las partículas, todos estos experimentos, así como muchos otros

Los experimentos de Hertz

Ahora nos referiremos al hecho de cómo fueron demostradas las hipótesis concernientes al aspecto ondulatorio del campo electromagnético.

La lógica y la matemática sacan a partir de las leyes de Maxwell unas deducciones. Estas deducciones podrían ser ciertas, pero, podrían no confirmarse en el experimento. Una teoría física ocupa su lugar en la ciencia tan sólo después de su comprobación experimental. El camino del devenir de la teoría del campo electromagnético, desde hechos sueltos hacia hipótesis generales, desde las hipótesis hacia sus consecuencias, y, la última etapa, el experimento que pronuncia su fallo decisivo, es el único derrotero justo del naturalista. Y este derrotero» se observa con especial nitidez precisamente en el ejemplo de las leyes del campo electro magnético.

Esta es la razón de que nos detendremos de una forma detallada en los experimentos de Hertz que incluso hoy día ayudan al profesor a mostrar al escolar o al estudiante cómo se conforma la seguridad del científico en el carácter certero de las leyes de la naturaleza.

Tenemos que comenzar nuestra historia desde el año 1853 cuando el famoso físico inglés Kelvin demostró matemáticamente que durante la descarga del condensador a través de una bobina de autoinducción en el circuito se engendran oscilaciones eléctricas: la carga en las armaduras del condensador, la tensión en cualquier tramo del circuito y la intensidad de la corriente, todas estas magnitudes cambiarán según la ley de oscilación armónica. Si considerarnos que la resistencia en el circuito es ínfima, estas oscilaciones durarán eternamente.

La fig. 5.5 representa el cuadro que esclarece los fenómenos acaecidos en este, así llamado circuito oscilante. En el instante inicial de tiempo el condensador está cargado. Apenas el circuito se haya cerrado en éste fluirá la corriente. Dentro de un cuarto de período el condensador resultará por completo descargado. Su energía $q^2/2C$ se transformará en energía del campo magnético de la bobina.

En este instante la intensidad de la corriente será máxima. La corriente no se interrumpirá, sino seguirá fluyendo en la misma dirección disminuyendo poco a poco su intensidad.

Dentro de un semiperíodo la intensidad de la corriente se anulará y desaparecerá la energía magnética $LI^2/2$, mientras que el condensador se cargará por completo,

recuperando su energía. Sin embargo, la tensión cambiará de signo. Seguidamente, el proceso se repetirá, por decirlo así, en el sentido contrario. Al pasar cierto tiempo T (período de oscilación) todo regresará al estado inicial y el proceso comenzará de nuevo.

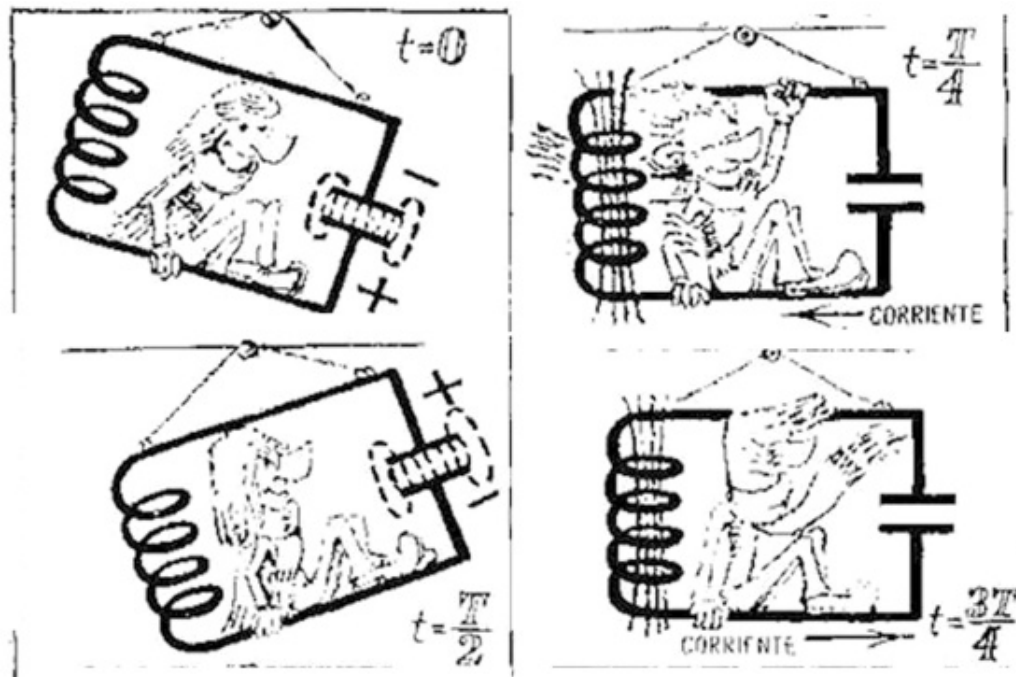


Figura 5.5

Las oscilaciones eléctricas habrían durado infinitamente, si no hubiera sido por la inevitable resistencia a la corriente. Gracias a esta resistencia, en cada período la energía se irá perdiendo, y las oscilaciones, disminuyendo su amplitud, se amortiguarán de una manera rápida.

La ostensible analogía con las oscilaciones de un peso fijado en un muelle nos permite prescindir del análisis algebraico del proceso, formándonos una idea sobre cuál será el período de oscilación en un circuito de este tipo. (El lector debe refrescar en la memoria las correspondientes páginas del libro 1.) En efecto, resulta bastante evidente que la energía eléctrica del condensador es equivalente a la energía potencial de un muelle comprimido, y la energía magnética de la bobina, a la energía cinética del peso.

Equiparando las magnitudes análogas, «deducimos» la fórmula del período de las oscilaciones eléctricas que tienen lugar en el circuito: $q^2/2C$ tiene como análogo $kx^2/2$; $LI^2/2$ tiene como análogo $mv^2/2$; k es análogo de $1/C$, y L , de m . Resulta que la frecuencia de oscilación es

$$\nu = \frac{1}{2\pi\sqrt{LC}}$$

ya que para la oscilación mecánica la correspondiente expresión tiene la siguiente forma

$$\nu = \frac{1}{2\pi} \sqrt{\frac{k}{m}}$$

Ahora tratemos de adivinar el curso de razonamientos de Hertz quien se planteó el problema de demostrar, sin salir del laboratorio, la existencia de las ondas electromagnéticas que se propagan con la velocidad de 300.000 km/s. De este modo, resulta necesario obtener una onda electromagnética con una longitud del orden de 10 m. Si Maxwell tenía razón, entonces, para lograr este objetivo, es preciso que los vectores eléctrico y magnético oscilasen con la frecuencia de 3×10^8 Hz... perdón, s^{-1} . Es que en aquel entonces Hertz no conocía que su nombre sería immortalizado con la denominación de la unidad de frecuencia.

Entonces, ¿por dónde empezar? En primer término, como quiera que las oscilaciones sean amortiguadas, hay que crear un dispositivo capaz de renovar el proceso después de haber cesado la corriente. Esto se realiza sin dificultad. El esquema se representa en la fig. 5.6. Al devanado primario del transformador T se suministra la tensión alterna. Apenas ésta haya alcanzado su valor de perforación entre las esferas conectadas al devanado secundario saltará una chispa.

Justamente esta chispa cierra el circuito oscilante K haciendo las veces de llave, y en el circuito con una frecuencia más o menos alta pasará una decena de oscilaciones con amplitud cada vez menor.

Pero la frecuencia debe ser alta. ¿Qué se debe hacer para conseguirlo? Disminuir la autoinducción, así como disminuir la capacidad. ¿De qué modo?

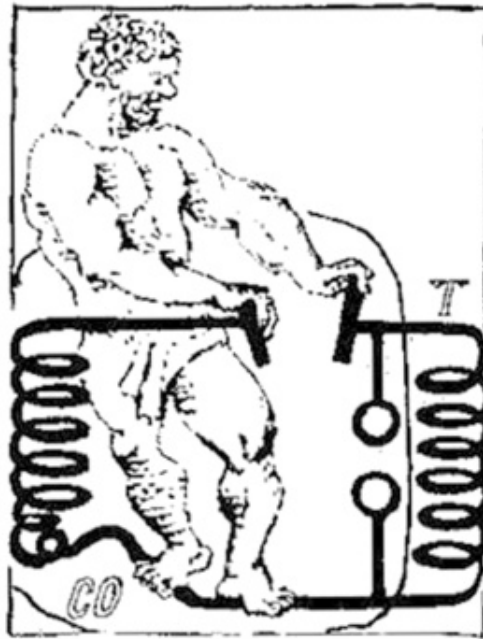


Figura 5.6

Sustituimos la bobina por un conductor recto, y en cuanto a las placas del condensador comenzamos a aumentar la distancia entre éstas y aminorar su área. ¿En qué degenera, entonces, el circuito oscilante? Meramente, de éste no queda nada: sólo dos varillas que terminan en esferas entre las cuales salta la chispa.

De esta forma, Hertz llegó a la idea de su vibrador u oscilador que puede servir tanto de fuente, como de receptor de las ondas electromagnéticas.

A Hertz le fue difícil vaticinar de antemano a qué serían iguales la inductancia y la capacidad de este «circuito» sui géneris, del cual, por decirlo así, no quedaron ni los huesos. La inductancia y la capacidad del vibrador no están concentradas en un punto del circuito, sino están distribuidas a lo largo de las varillas. Se necesita otra teoría.

Sin embargo, el examen de este nuevo enfoque de los circuitos eléctricos en los cuales fluyen corrientes de frecuencia muy alta, nos conduciría demasiado lejos. El lector puede creernos de palabra que en el vibrador de Hertz, en efecto, se engendran oscilaciones de corriente de alta frecuencia.

El «transmisor» y el «receptor» de ondas empleados por Hertz eran, prácticamente, idénticos. En el «transmisor» las oscilaciones se excitaban por una chispa que saltaba de un modo periódico entre las esferas en correspondencia con el trabajo del transformador que suministraba tensión al vibrador. El espacio de chispa, o sea, la distancia disruptiva, podía variarse mediante un tornillo micrométrico.



Enrique Hertz (1857 - 1894), relevante físico alemán que demostró experimentalmente, por medio de “vibrador” y “resonador” que la descarga oscilante engendra en el espacio ondas electromagnéticas. Hertz puso de manifiesto que las ondas electromagnéticas se reflejan, se refractan y se polarizan, con lo que confirmó la teoría de Maxwell. Reconociendo el mérito de Enrique Hertz cuyos experimentos pusieron los cimientos de la radiotecnia, el inventor de la radio Alexander S. Popov en su primera radiotransmisión en 1896 transmitió dos palabras: “Enrique Hertz”.

De transmisores servían ya sea una espira rectangular del conductor interrumpida por el espacio de chispa, o bien, dos varillas que, según se deseaba, podían aproximarse hasta distancias de partes de milímetro.

En su primer trabajo publicado en 1885 Hertz demostró que por el método descrito con anterioridad podían obtenerse oscilaciones de muy alta frecuencia y que estas oscilaciones, efectivamente, creaban en el espacio circundante un campo alterno cuya existencia se permitía apreciar debido a la chispa que saltaba en el «transmisor». El vibrador receptor Hertz lo llamó resonador. Desde el mismo comienzo al científico le era evidente el principio de detección del campo electromagnético que forma la base de la radiotecnia contemporánea.

A propósito, cabe señalar que en los trabajos de Hertz, así como varios decenios después, no estaban en boga las palabras «ondas electromagnéticas» ni «ondas radioeléctricas». Se habló de ondas eléctricas o de ondas de fuerza electrodinámica.

En su siguiente trabajo Hertz demuestra que, en plena correspondencia con las exigencias de la teoría de Maxwell, el medio dieléctrico (una barra de azufre o de parafina) influye en la frecuencia del campo electromagnético. Al recibir este artículo, el editor de la revista, Helmholtz, contestó a Hertz enviándole una tarjeta: *«El manuscrito se ha obtenido. Bravo. El jueves lo mando a la imprenta»*.

Una enorme impresión a los físicos de aquel tiempo la produjo el trabajo de Hertz en que éste demostró la reflexión de las ondas electromagnéticas. Las ondas se reflejaban de una pantalla de cinc de 4 m x 2 m. El vibrador se encontraba a una distancia de 13 m de la pantalla y a una altura de 2,5 m desde el suelo. El resonador sintonizado se colocó a la misma altura y se trasladaba entre el vibrador y la pantalla. Al disponer el resonador a diferentes distancias respecto a la pantalla y observar la intensidad de la chispa, Hertz estableció la existencia de máximos y de mínimos característicos para el cuadro de la interferencia que lleva el nombre de onda estacionaria. La longitud de onda resultó próxima a 9,6 m.

Quisiera subrayar que en aquel entonces nadie podía decir qué material debía servir de espejo para las ondas electromagnéticas. Ahora conocemos que las ondas de tal longitud no penetran en el metal, sino se reflejan de éste.

En sus intentos de obtener demostraciones complementarias de la teoría de Maxwell, Hertz, disminuye las dimensiones geométricas de sus aparatos llevando la

longitud de onda hasta 60 cm. En 1888 publica el trabajo «*Acerca de los rayos de fuerza eléctrica*». Logró concentrar la energía electromagnética por medio de espejos parabólicos. En el foco de los espejos se situaban las varillas del vibrador y del resonador. Valiéndose de este receptor y transmisor de espejo, Hertz demostró que las ondas electromagnéticas no pasaban a través de los metales, mientras que las pantallas de madera no retenían estas ondas.

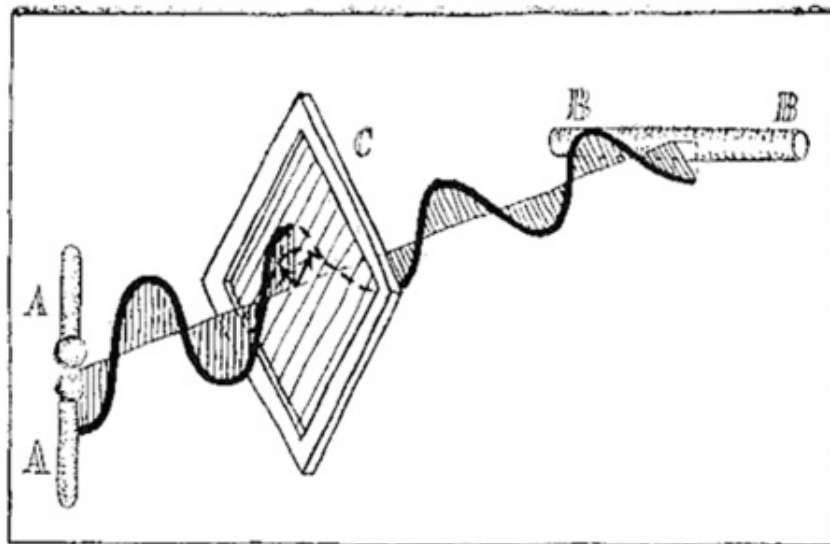


Figura 5.7

En la fig. 5.7 se muestra cómo Hertz demostró la polarización de las ondas electromagnéticas. El camino del rayo electromagnético proporcionado por el vibrador A A, lo interceptaba la rejilla C hecha de alambres de cobre. Haciendo girar la rejilla Hertz mostró que la intensidad de la chispa en el resonador BB variaba. La onda eléctrica incidente en la rejilla se desdoblaba en dos componentes: una perpendicular y otra paralela a los alambres de cobre; fue la componente perpendicular la que salvaba el obstáculo. Estos mismos razonamientos pueden referirse también al vector magnético. El carácter transversal de la onda electromagnética quedó demostrado.

Finalmente, teniendo por objeto el estudio de la refracción de las ondas, Hertz prepara a partir de masa asfáltica, un prisma de más de una tonelada de peso. Para las ondas con una longitud de 60 cm se podía medir con gran precisión el coeficiente de refracción del asfalto. Dicho coeficiente resultó ser igual a 1,69.

¡Demostración de la existencia de las ondas electromagnéticas, medición de su longitud, establecimiento de las leyes de su reflexión, refracción y polarización!

Y todo lo enumerado es resultado de un trabajo de tres años, solamente. Sobran razones para quedarse admirado.

Clasificación de la radiación electromagnética

Los físicos tienen que tratar con la radiación electromagnética de diapasón enorme. La radiación electromagnética de la corriente de frecuencia urbana es absolutamente despreciable. La posibilidad práctica de captar la radiación electromagnética comienza desde las frecuencias del orden de decenios de kilohercios, es decir, desde las longitudes de onda iguales a centenares de kilómetros. Las ondas más cortas tienen la longitud del orden de diezmilésimas partes del micrómetro, o sea, se trata de frecuencias del orden de miles de millones de gigahercios.

Se da el nombre de ondas radioeléctricas a aquella radiación electromagnética que se crea por los medios electrotécnicos, es decir, a costa de oscilaciones de las corrientes eléctricas. Las ondas radioeléctricas más cortas son aquellas cuya longitud es de centésimas partes de milímetro.

Desde varias centenas de micrómetros y más abajo se extiende la zona de las longitudes de onda de la radiación que se origina debido a las transiciones de energía en el seno de las moléculas, en el seno de los átomos y en el seno de los núcleos atómicos. Como vemos, este diapasón se solapa sustancialmente con el radioeléctrico.

La irradiación de luz visible ocupa un sector estrecho. Sus límites son de 0,38 a 0,74 μm . Entre las radiaciones obtenidas por métodos no radiotécnicos la de longitudes de onda más largas se denomina radiación infrarroja, y la de longitudes de onda más cortas, ultravioleta, hasta las longitudes de cerca de 0,1 μm .

La radiación electromagnética obtenida en los tubos de rayos X se solapa parcialmente con la zona de las ondas ultravioletas llegando hasta 0,01 μm donde, a su vez, se solapa con la zona de los rayos λ . Estos últimos tienen su origen en la desintegración nuclear, las reacciones nucleares y las colisiones entre las partículas elementales.

La característica principal de cualquier radiación electromagnética es su espectro. Se denomina espectro el gráfico en que, por la vertical, se marca la intensidad (es decir, energía que recae en unidad de tiempo por una unidad de superficie), y por la horizontal, la longitud de onda o la frecuencia. El espectro más simple lo posee la radiación monocromática («de un solo color»).

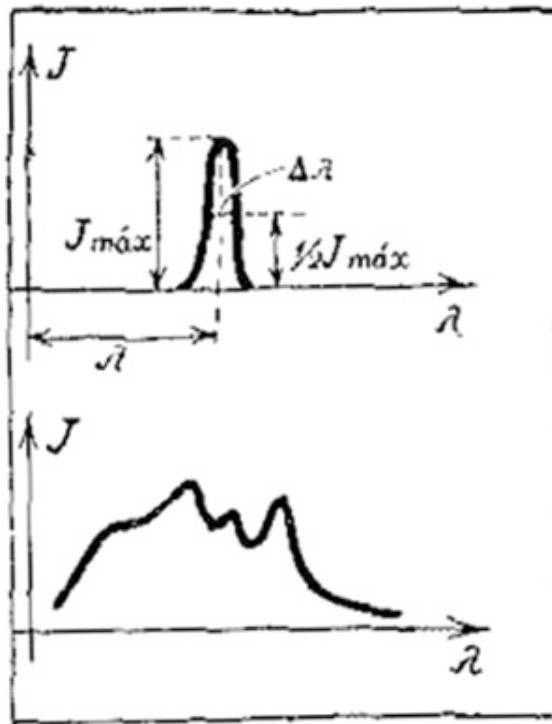


Figura 5.8

Su gráfico consiste en una sola línea de ancho muy pequeño (fig. 5.8, arriba). El grado de monocromaticidad de la línea se caracteriza por la relación $\lambda/\Delta\lambda$. Las estaciones de radio dan una radiación casi monocromática. Por ejemplo, una estación de ondas cortas que trabaja en el diapason de 30 m, la relación $\lambda/\Delta\lambda$ es igual a 1000, aproximadamente.

Los átomos excitados, por ejemplo, los átomos de los gases en las lámparas de luz diurna (la excitación tiene lugar a costa de colisiones de partículas cargadas positiva y negativamente que migran hacia el ánodo y el cátodo) ofrecen un espectro compuesto de una gran cantidad de líneas monocromáticas cuya anchura relativa es

de $(100.000)^{-1}$. En la resonancia magnética se observan líneas cuya anchura es hasta de 10^{-10} .

Estrictamente hablando, no existen espectros continuos. No obstante, si las líneas se solapan, el resultado lleva a la curva de intensidad representada en la parte inferior de la misma figura.

Los datos sobre el espectro electromagnético pueden obtenerse tanto investigando la radiación, como estudiando la absorción. Hablando con propiedad, ambos experimentos son portadores de la misma información. Esto se pone de manifiesto a partir de la ley principal de la física cuántica. En el caso de radiación, el sistema pasa del nivel superior de energía al inferior, y en el caso de absorción, del nivel inferior al superior. Pero la diferencia de energía que determina la frecuencia de la radiación o de la absorción en ambos casos resultará ser idéntica. Y es solamente una cuestión de conveniencia la de qué espectro investigar: el de radiación o el de absorción.

Al caracterizar el espectro de radiación podemos valernos, claro está, tanto del lenguaje ondulatorio, como del corpuscular. Si utilizamos el aspecto ondulatorio de la radiación, decimos que la intensidad es proporcional al cuadrado de amplitud de la onda. Al analizar la radiación en concepto de flujo de partículas, calculamos la intensidad como el número de fotones.

Volvemos a repetir, una vez más, que no nos debe desconcertar, de ningún grado, la utilización alternativa de ambos aspectos del campo electromagnético. La radiación no se parece ni a la onda, ni al flujo de partículas. Los dos cuadros no son modelos propicios para utilizar cuando se explican tales o cuales fenómenos.

No insertamos la escala de las ondas electromagnéticas, pero señalamos con bastante claridad que las denominaciones de sus diferentes sectores son hasta cierto punto convencionales y podemos encontrarnos con los casos en que las ondas de igual longitud tendrán nombres distintos en dependencia del método de su obtención.

A la vez, la escala de las ondas electromagnéticas es continua. No hay porciones que no se puedan obtener por uno u otro procedimiento.

Hay que tener presente, sin embargo, que sólo relativamente hace poco se ha puesto al descubierto el hecho de que las ondas infrarrojas se solapan con las

radioeléctricas, los rayos γ y con los rayos X, etc. Durante mucho tiempo existió una laguna entre las ondas radioeléctricas cortas y las infrarrojas. Ondas de 6 mm de longitud las obtuvo en 1895 el notable físico ruso Piotr Nikoláievich Lébedev, mientras que las ondas térmicas (las infrarrojas) hasta de 0,34 mm de longitud deben su descubrimiento a Rubens.

En 1922, A. A. Glagóleva-Arkádieva eliminó también esta laguna, obteniendo por un método no óptico ondas electromagnéticas cuya longitud fue de 0,35 mm a 1 cm.

Hoy en día las ondas de esta longitud se obtienen por los radiotécnicos sin dificultad. Sin embargo, en aquel entonces se necesitó una buena dosis de ingeniosidad e inventiva para construir el aparato que la investigadora bautizó de radiador de masa. De fuente de radiación electromagnética sirvieron limaduras metálicas que se encontraban en suspensión en el aceite para transformadores. A través de esta mezcla se dejaba pasar una descarga por chispa.

Capítulo 6

La radio

Contenido:

- *Unas páginas de historia*
- *Válvula tríodo y transistor*
- *Radiotransmisión*
- *Radorrecepción*
- *Propagación de las ondas radioeléctricas*
- *Radiolocalización*
- *Televisión*
- *Microcircuitos electrónicos*

Unas páginas de historia

De la misma manera como Faraday no sospechaba que el descubrimiento de la inducción electromagnética conduciría a la creación de la electrotecnia y Rutherford consideraba que las conversaciones acerca de la posibilidad de extraer, en el futuro, energía a partir del núcleo atómico no eran sino una cháchara ignorante, de una forma análoga, Enrique Hertz que descubrió las ondas electromagnéticas y demostró que era posible captarlas a una distancia de varios metros, no sólo no pensó en la comunicación por radio, sino que hasta negó su posibilidad. ¿No es verdad que son tres hechos curiosos? Pero recapitular sobre éstos es tarea de un psicólogo. En cuanto a nosotros, limitándonos a constatar esta sorprendente circunstancia, vamos a ver cómo se desarrollaron los acontecimientos después de la prematura muerte de Hertz que tuvo lugar en 1894.

Los experimentos clásicos de Hertz que hemos descrito bastante detalladamente atrajeron la atención de los científicos de todo el mundo. El profesor de la Universidad de San Petersburgo N. G. Egórov los imitó exactamente. La chispa en el resonador apenas si se percibía. Sólo se podía distinguir en plena oscuridad y, además, recurriendo a un cristal de aumento.

Alexandr Stepánovich Popov (1859 - 1906), un modesto profesor de electrotecnia en una escuela militar de la ciudad de Kronstadt², en 1889, a la edad de 30 años comenzó a perfeccionar los experimentos de Hertz. Las chispas que obtenía en sus resonadores resultaron ser mucho más fuertes que aquellos que lograron producir otros investigadores.

En otoño de 1894 en una revista inglesa apareció el artículo del conocido físico Oliverio Lodge quien daba a conocer que era posible perfeccionar el resonador de Hertz si se empleaba el tubo de Branly. El científico francés Eduardo Branly estudiaba la conductibilidad de las limaduras metálicas. Descubrió que dichas limaduras no siempre oponían igual resistencia a la corriente eléctrica. Resultó que la resistencia de las limaduras echadas en un tubo disminuía bruscamente si éste estaba situado cerca del vibrador de Hertz. La causa de ello residía en la aglutinación de las limaduras. La resistencia de las limaduras podía restituirse, pero, para lograrlo, el tubo se debía sacudir.

Precisamente de esta propiedad de las limaduras metálicas se aprovechó Lodge. Este compuso un circuito a base del tubo de Branly (que recibió el nombre de cohesor, del latín «*cohoesum*», unido), una pila y un galvanómetro sensible. En el momento de paso de las ondas electromagnéticas la aguja del aparato se desviaba. Lodge consiguió detectar las ondas radioeléctricas hasta las distancias de cerca de 40 m.

La inconveniencia de este sistema consistía en que el cohesor fallaba de una manera inmediata, apenas puesto en funcionamiento. Era necesario inventar un procedimiento que permitiría hacer volver las limaduras aglutinadas (soldadas) a su estado inicial y, además, componer el circuito tal que el sacudimiento se realizase como si «por sí mismo».

Fue precisamente esta tarea la que resolvió Popov. Ensayó muchas variantes diversas del cohesor, parándose, en fin de cuentas, en la siguiente construcción. «Dentro de un tubo de vidrio, en sus paredes, están pegadas dos cintas AB y CD de fino platino en chapas, extendiéndose casi a toda la longitud del tubo. Una cinta sale a la superficie exterior por un extremo del tubo, y la otra, por el extremo opuesto. Las cintas de platino se encuentran con sus bordes a una distancia de 2 mm, siendo

² En las cercanías de San Petersburgo; N. del T.

su ancho de 8 mm; los extremos interiores de las cintas B y C no llegan a los tapones que cierran el tubo, para que el polvo colocado en este último no pueda, llenando el tapón, formar hilos conductores no destruibles por sacudimiento, como ocurrió en algunos modelos.

Basta que la longitud de todo el tubo sea de 6 a 8 cm, siendo el diámetro de 1 cm. El tubo, durante su funcionamiento, se dispone horizontalmente, de modo que las cintas se encuentran en su mitad inferior y el polvo metálico las cubre. El mejor funcionamiento se logra cuando el tubo se llene no más que en una mitad».

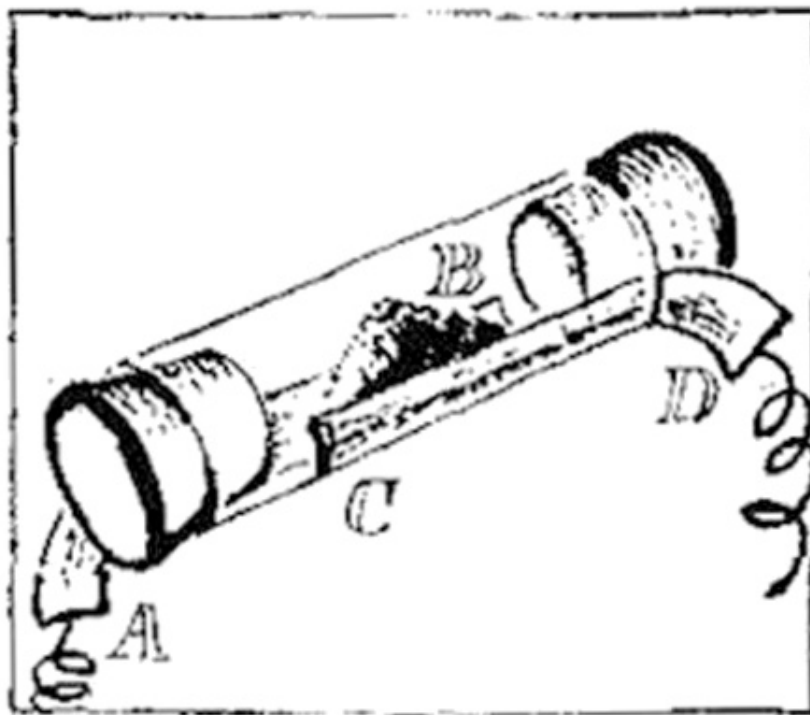


Figura 6.1

El esquema del cohesor de Popov que se caracteriza con sus propias palabras se representa en la fig. 6.1. Popov empleaba polvo de hierro o de acero.

Sin embargo, la tarea principal no consistía en perfeccionar el cohesor, sino en inventar el procedimiento para volverlo al estado inicial después de la recepción de la onda electromagnética. En el primer receptor de Popov cuyo esquema se reproduce en la fig. 6.2 este trabajo lo cumplía el timbre eléctrico común y

corriente. El timbre sustituye la aguja del galvanómetro y su martillo golpea el tubo de vidrio cuando retorna a la posición inicial.

¡Qué solución tan sencilla del rompecabezas! Y de veras que es sencilla. Que el lector aprecie en su justo valor la idea principal, la que no ocurrió a los físicos tan excelentes como Hertz y Oliverio Lodge.

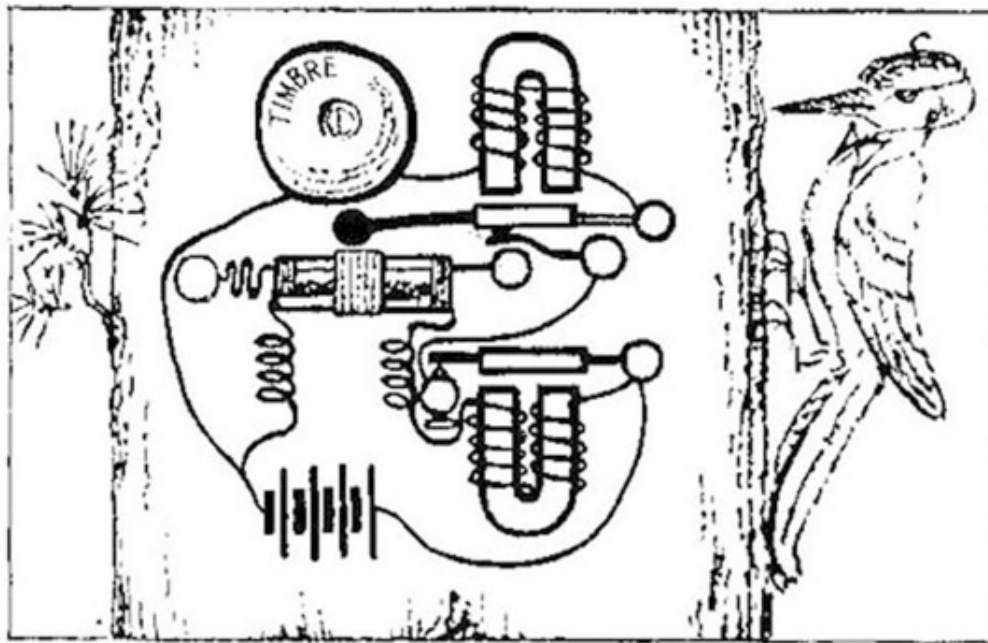


Figura 6.2

He aquí que en este simple esquema se utiliza por primera vez lo que los técnicos llaman circuito de relé. La ínfima energía de las ondas radioeléctricas no actúa directamente, sino se aprovecha para gobernar el circuito de una corriente. El timbre y el cohesor eran los primeros amplificadores.

En los días primaverales de 1895, Popov decidió realizar su experimento en el jardín. Comenzaron a alejar el receptor del vibrador; 50 metros, el timbre responde a la chispa del vibrador; 60 metros, si, funciona; 80 metros, el timbre calla.

Entonces Popov arrima al receptor un rollo de alambre de cobre, echa uno de sus extremos al árbol y el extremo inferior lo conecta al cohesor. El timbre empezó a sonar. Así fue cómo apareció la antena, primera en el mundo.

El 7 de mayo celebramos el Día de la radio. En 1895, este día, Popov, en la sesión ordinaria de la Sociedad fisicoquímica rusa lee su informe que lleva el modesto título de *«Acerca de la relación de los polvos metálicos a las oscilaciones eléctricas»*.



Alexandr Stepánovich Popov (1859 - 1906), físico y electrotécnico ruso, inventor de la radio. Los trabajos de A. S. Popov fueron altamente apreciados por los contemporáneos. En 1900, en la Exposición Mundial de París, por su invento le fue concedida la medalla de oro.

Entre los circunstantes, muchas personas presenciaron unos años atrás los experimentos de Hertz, aquellos mismos experimentos en que la diminuta chispa debía observarse a través de una lupa. Sin embargo, al oír el animado tintineo del timbre del receptor de Popov todos comprendieron que nació el telégrafo sin hilos y apareció la posibilidad de transmitir señales a distancia.

El 24 de marzo de 1896 Popov transmite el primer radiograma en el mundo. De un edificio a otro, a una distancia de 250 metros, cerrando la llave para intervalos

breves y largos, se transmiten las palabras «Enrique Hertz» que se marcan en una cinta telegráfica.

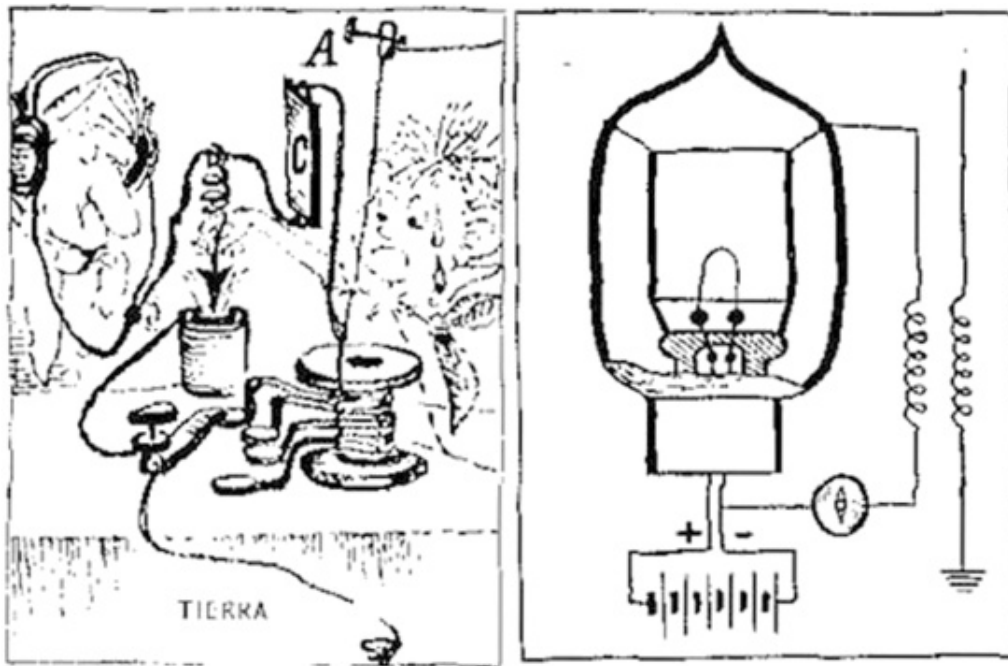
Para 1899 la distancia de la radiocomunicación entre los barcos del destacamento-escuela de minas alcanza ya 11 km. Ni siquiera las mentes más escépticas ponen en tela de juicio el valor práctico del telégrafo sin hilos. El inventor italiano Guillermo Marconi comenzó sus experimentos un poco más tarde que Popov. Seguía atentamente todos los alcances en el campo de electrotecnia y de estudio de las ondas electromagnéticas utilizándolos hábilmente para mejorar la calidad de la radorrecepción y radiotransmisión. Su gran mérito atañe no tanto a la parte técnica del asunto, como a la organizativa. Pero ello no es poca cosa, y el nombre de Marconi se debe pronunciar con respeto, sin olvidar, al mismo tiempo, que la prioridad indiscutible y reconocida en todo el mundo en la creación de la radio pertenece al modesto científico ruso. Marconi no hacía mención de Popov en sus artículos e intervenciones. Pero no todos conocen que en 1901 invitó al profesor A. S. Popov a entrar a trabajar en la sociedad anónima cuyo presidente fue. La distancia de radorrecepción crecía a ritmo rápido. En 1899 Marconi efectuó la radiocomunicación entre Francia e Inglaterra, y en 1901 la radio enlaza América y Europa.

Bueno, ¿qué innovaciones técnicas contribuyeron a este éxito y al nacimiento de la radiodifusión?

Desde el año 1899 la radiotecnica comienza a despedirse del cohesor. En vez de detectar las ondas radioeléctricas a costa de la caída de la resistencia en el circuito que tiene lugar por impacto de la onda electromagnética, se puede hacer uso de un procedimiento completamente distinto. Una onda electromagnética pulsatoria rectificada puede recibirse por un ordinario auricular telefónico.

Empieza la búsqueda de distintos rectificadores. El detector de contacto sólido ampliamente difundido que se empleó hasta los años 20 del siglo XX era un cristal con conductibilidad unilateral. Cristales de este tipo se conocían desde 1874. A éstos pertenecen sulfuros de metales, piritas de cobre y centenares de diversos minerales. Mis coetáneos recuerdan semejantes receptores y el fastidioso procedimiento de la búsqueda de un «buen contacto» por medio de una aguja muelleante el cual aparecía cuando la punta del muelle se apoyaba en el punto

«propicio» del cristal (fig. 6.3). En aquella época ya trabajaban muchas estaciones de radio, por esta razón se debía sintonizar el receptor con la onda, hecho que se conseguía mediante la conmutación por contacto, si se tenía en cuenta la recepción de un número contado de estaciones, o bien, cambiando suavemente la capacidad del condensador, lo que se realiza también en los dispositivos modernos.



Figuras 6.3 y 6.4

Era muy difícil, si no imposible, obtener grandes potencias de las estaciones de radio por chispa ya que el descargador se calentaba. Estas estaciones de radio fueron relevadas por el arco voltaico y la máquina de alta frecuencia. La potencia comenzó a calcularse ya a centenas de kilovatios.

Sin embargo, una verdadera revolución de la radiocomunicación que permitió pasar de la telegrafía sin hilos a la transmisión del habla humana y de la música la produjo la utilización de la válvula electrónica.

En octubre de 1907, el ingeniero electricista inglés John Fleming (1849 - 1945) demostró que la corriente eléctrica de alta frecuencia podía rectificarse con la ayuda de un tubo de vacío que consta de un filamento caldeado por la corriente y rodeado por un cilindro metálico. En la fig. 6.4 se expone el esquema del mismo. Fleming se

daba cuenta de la significación del diodo de vacío para transformar las oscilaciones eléctricas en sonido (también denominaba esta lámpara «audión», de la palabra latina «audio» que significa escucho), sin embargo, no pudo conseguir que su detector encontrara una amplia aplicación.

Los laureles de inventor de la válvula (a veces se llama también tubo o lámpara) electrónica los cosechó el norteamericano Lee de Forest (1873 - 1961). Fue el primero en transformar la lámpara en tríodo (1907). El receptor de lámparas de Lee de Forest recibía las señales a la rejilla de la lámpara, las rectificaba y daba la posibilidad de escuchar las señales telegráficas en el teléfono.

Para el ingeniero norteamericano eran patentes las posibilidades de la lámpara en concepto de amplificador. Pero se necesitaron seis años más para que el ingeniero alemán Meissner, en 1913, introdujese el tríodo en el circuito de generador.

Todavía antes de utilizar la válvula electrónica como generador se hacían intentos de transmitir el habla, es decir, de modular la onda electromagnética. Pero las dificultades eran colosales y no se podía hacer amplia la banda de frecuencias de la modulación. Se consiguió transmitir de cualquier modo la voz humana, pero no la música. Solamente en los años 20 los radiotransmisores y radorreceptores que trabajaban a base de válvulas electrónicas permitieron percatarse de las posibilidades, verdaderamente inagotables, de la radiocomunicación como emisión que abarca todo el diapasón de las audiofrecuencias.

El siguiente salto revolucionario tuvo lugar hace muy poco tiempo, cuando los elementos semiconductores desalojaron de los circuitos de radio los tubos electrónicos. Nació una nueva rama de la física aplicada que examina un enorme conjunto de problemas relacionados con la transmisión, recepción y el almacenamiento de la información.

Válvula tríodo y transistor

Las válvulas tríodo produjeron una revolución en la radiotecnica. Pero la técnica envejece más rápidamente que los hombres. Y hoy en día podemos llamar viejecita la válvula electrónica, y si uno entra en la tienda que vende televisores, seguramente oír cómo los impacientes clientes preguntan sobre la producción de

los televisores a base de semiconductores. No dudo que para el momento en que este libro haya visto la luz en la venta habrá gran cantidad de estos televisores.

No obstante, es más fácil explicar los principios de las dos aplicaciones fundamentales de las válvulas y los transistores, a saber, de la amplificación y generación de ondas de frecuencia determinada, en el ejemplo de la válvula electrónica. Por esta razón más tiempo analizaremos su acción de una forma más detallada que el funcionamiento del transistor.

En la ampolla de una válvula de tres electrodos, además del ánodo y del filamento catódico caldeado por la corriente, está soldado también el tercer electrodo denominado rejilla. Los electrones pasan libremente a través de la rejilla. Sus orificios superan tantas veces las dimensiones de los electrones cuantas veces la Tierra es mayor que una partícula de polvo.

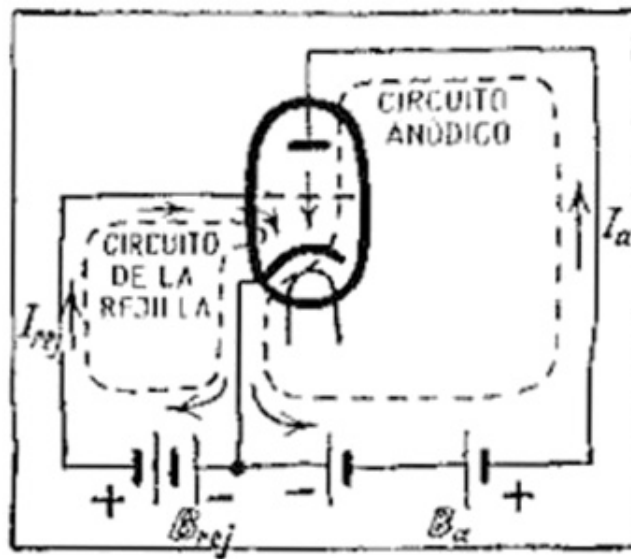


Figura 6.5

La fig. 6.5 ilustra cómo la rejilla permite controlar la corriente anódica. Resulta evidente que una tensión negativa en la rejilla hará disminuir la corriente anódica, y la positiva la aumentará.

Realicemos un sencillo experimento. Apliquemos entre el cátodo y el ánodo una tensión igual a 100 V.

Después comencemos a cambiar la tensión en la rejilla, digamos de tal forma como se muestra en la fig. 6.6, o sea, dentro de los límites de -8 V a $+4\text{ V}$, aproximadamente.

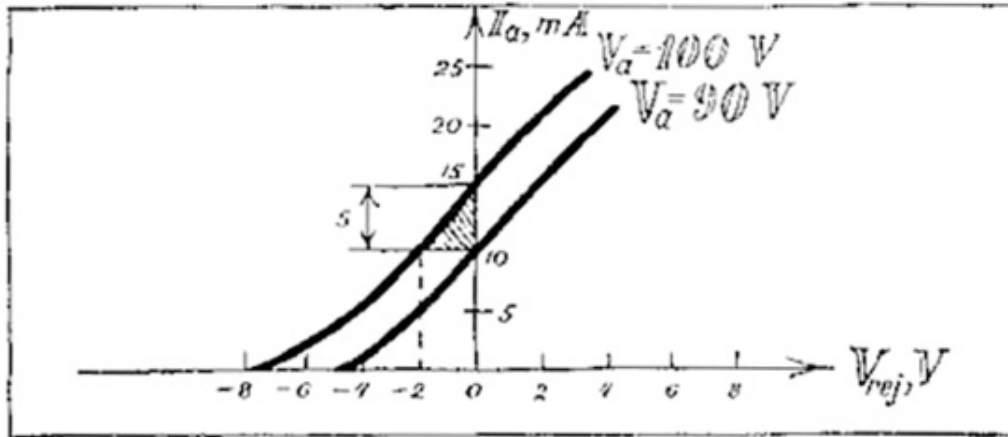


Figura 6.6

Utilizando un amperímetro mediremos la corriente que fluye a través del circuito anódico. Se obtendrá una curva que se representa en la figura. Esta se denomina característica de la lámpara. Vamos a repetir el experimento tomando esta vez la tensión anódica de 90 V . Obtendremos una curva parecida.

Ahora fíjense en el siguiente resultado notable. Como evidencia el triángulo rayado en el dibujo se puede lograr la amplificación de la corriente anódica en 5 mA por dos métodos: ya sea aumentando la tensión anódica en 10 V , o bien, elevando en 2 V la tensión en la rejilla. La introducción de la rejilla convierte la válvula tríodo en amplificador. El factor de amplificación en el ejemplo que hemos examinado es igual a 5 (diez dividido por dos). En otras palabras la tensión en la rejilla actúa sobre la corriente anódica con una fuerza cinco veces mayor que la tensión anódica.

Ahora veremos cómo el tríodo permite generar ondas de determinada longitud.

El esquema correspondiente, simplificado hasta el límite, se representa en la fig. 6.7.

Al conectar la tensión anódica comienza a cargarse el condensador C_{osc} del circuito oscilante a través de la lámpara. La armadura inferior se cargará positivamente.

Inmediatamente, se iniciará el proceso de descarga del condensador a través de la bobina de autoinducción L_b , se generarán oscilaciones libres.

Estas se atenuarían si no fuese por el hecho de que de la lámpara, constantemente, llegase energía. ¿Y cómo conseguir que este apoyo energético se opere coincidiendo en el ritmo, y el circuito oscilante se intensifique a guisa de vaivén de un columpio?

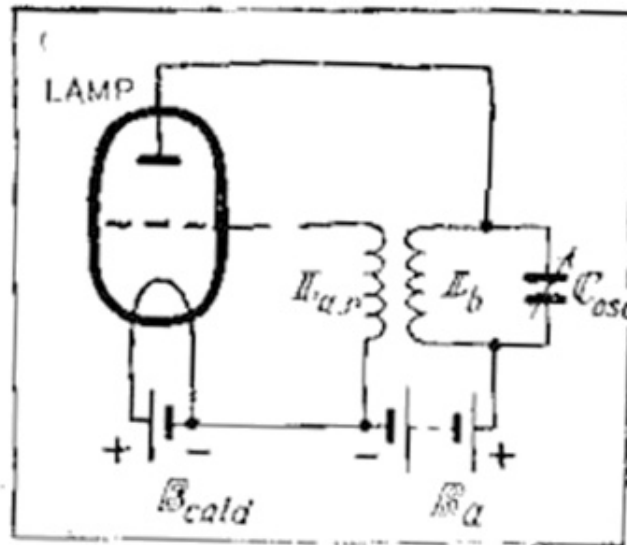


Figura 6.7

Para conseguir este objetivo es necesario el llamado *acople de reacción*. En la bobina L_{ar} la corriente del circuito oscilante induce la f.e.m. de inducción de la misma frecuencia que la de las oscilaciones libres. De este modo, la rejilla crea en el circuito anódico una corriente pulsatoria que intensificará el circuito con su frecuencia propia.

Los dos geniales principios que hemos expuesto, de generación y amplificación de las oscilaciones electromagnéticas han creado la radiotecnica y las ramas conjugadas a ésta. La válvula electrónica abandona las tablas y cede su lugar al transistor, pero las ideas de amplificación y generación permanecen las mismas.

Pero no hemos dicho todavía qué representa el transistor. Este tiene tres electrodos. El emisor corresponde al cátodo; el colector, al ánodo, y la base, a la rejilla. La borna del emisor sirve de entrada y la del colector, de salida.

Como se ve en la fig. 6.8, el transistor consta de dos uniones del tipo p-n. Es posible que la capa p se encuentre en el centro, asimismo, se puede tener transistores n-p-n.

Al emisor siempre se suministra la tensión de polarización positiva de modo que éste puede proporcionar gran cantidad de portadores principales de carga. Cuando el circuito de emisor de baja resistencia varía la corriente en el circuito de colector de alta resistencia, entonces, en este caso obtenemos la amplificación.

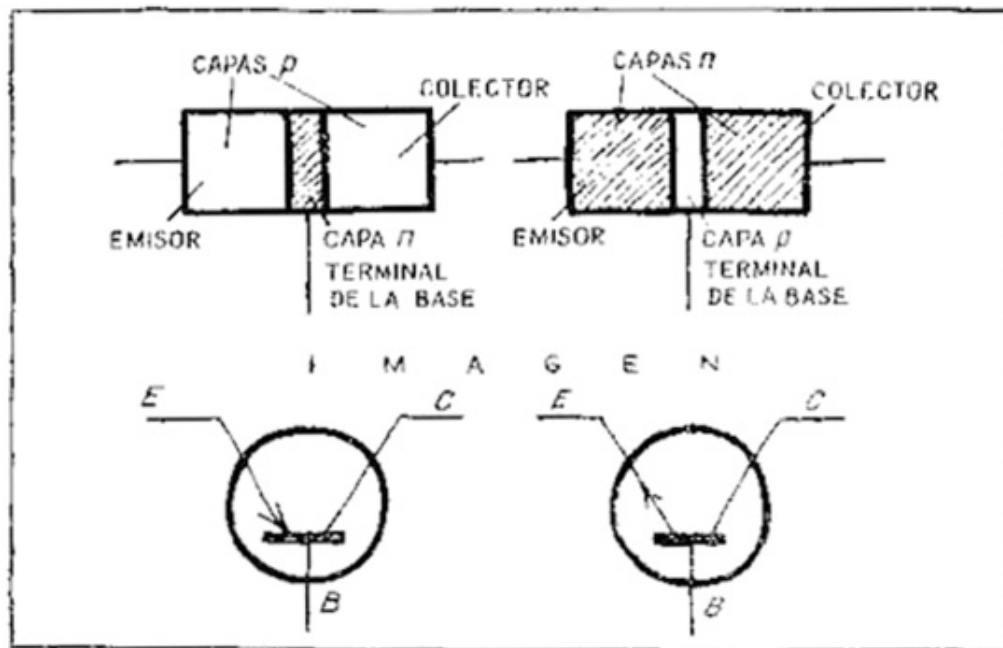


Figura 6.8

En el transistor, ni igual que en la válvula tríodo, una potencia pequeña en el circuito de entrada puede controlar una potencia grande en el circuito de salida. Existe una diferencia en el carácter del control. La corriente anódica de la lámpara, como acabamos de ver, depende de la tensión de rejilla, mientras que la corriente del colector depende de la del emisor.

Los esquemas de conexión de los transistores y su utilización en calidad de amplificadores y generadores, en general, son análogos a los principios de funcionamiento de la válvula tríodo. Pero aquí, nos abstenemos de discutir esta importantísima rama de la física moderna.

Radiotransmisión

Los tipos de radiotransmisiones pueden clasificarse de acuerdo con la potencia de las estaciones. Las estaciones grandes envían al éter potencias cuyo valor llega a un

megavatio. Un diminuto radiotransmisor con cuya ayuda el funcionario de la inspección automovilística comunica a su colega que el coche con la matrícula MMC 35 - 69 cruzó a la luz roja y tiene que ser detenido emite la potencia del orden de un milivatio. Para ciertas finalidades son suficientes las potencias incluso menores.

Son esenciales las diferencias en la estructura de las estaciones que trabajan en ondas de más de varios metros de longitud y de los dispositivos transmisores que emiten ondas ultracortas cuya longitud es de docenas de centímetros y, a veces, hasta de fracciones de centímetro. No obstante, también dentro de cada diapasón de longitudes de onda y potencias, el ingeniero que diseña una estación puede elegir entre un enorme número de esquemas y esta elección viene condicionada por la localidad, los objetivos específicos, razones económicas y, meramente, por la inventiva técnica.

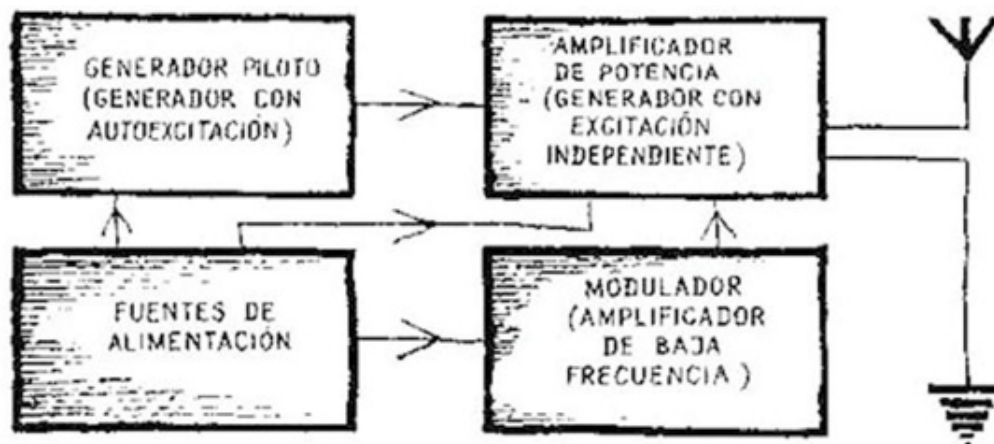
El núcleo de un radiotransmisor lo forma el generador de las ondas radioeléctricas. ¿En qué generador quiere usted centrar su atención? Existen por lo menos cinco posibilidades. Se puede tomar un generador de tubos termoiónicos. Su diapasón es extraordinariamente amplio. Las potencias pueden oscilar entre fracciones de vatio y centenares de kilovatios; las frecuencias, desde decenas de kilohercios hasta varios gigahercios. Pero si usted necesita potencias pequeñas, del orden de unas décimas de vatio, sólo le convendrá un generador transistorizado. Por el contrario, si se requieren potencias mayores que varias centenas de vatios, por ahora (pero, probablemente, no para mucho tiempo) se tiene que renunciar a los transistores. Y cuando se trata de potencias para las cuales pueden servir ambos tipos de generadores, el diseñador, a lo mejor, dará preferencia a la variante con transistor. Sin duda alguna, ganará la elegancia de la solución ingenieril. Un transmisor a base de transistores ocupará un volumen mucho menor y, si es necesario, es más fácil hacerlo portátil que cuando se trata de un transmisor con generador de tubos termoiónicos.

Una aplicación más especializada la tienen los generadores de magnetrón y de klistrón. El primero puede resultar útil en sumo grado, si es necesario enviar al espacio impulsos de varios megavatios de potencia. El diapasón de frecuencias para las cuales sirve el generador de magnetrón es mucho más estrecho, encontrándose, aproximadamente, entre 300 MHz y 300 GHz.

Para este mismo diapasón de ondas ultracortas se utilizan también los klistrones. Pero éstos revisten interés solamente cuando se trata de bajas potencias no superiores a varios vatios, en el diapasón centimétrico y varios miles de vatios en el diapasón milimétrico.

Estos dos últimos generadores, al igual que el quinto tipo, el generador cuántico, son muy específicos y requieren un examen especial. En lo que se refiere a los dispositivos transmisores a base de transistores y lámparas, los mismos resultan ser parecidos. Existen reglas radiotécnicas precisas orientándose por las cuales se puede sustituir una lámpara por un transistor equivalente.

Sin embargo, la acertada elección del generador de oscilaciones electromagnéticas dista mucho de resolver todos los problemas. Hay que decidir de qué modo amplificar la potencia creada por el generador maestro (o, como también se dice, generador piloto). También es preciso elegir el procedimiento para modular la onda portadora por medio de la audiofrecuencia. Hay muchas variantes de transmisión de energía al campo de antena. Además, la propia organización del campo de antena ofrece amplios horizontes para la inventiva del ingeniero.



Ficha 6.9

En radiotecnica se recurre muy a menudo a los llamados diagramas sinópticos. En el dibujo se presentan varios rectángulos con inscripciones. Y a medida que surge la necesidad se explica lo que se halla en cada recuadro. En la fig. 6.9 se muestra el diagrama sinóptico de una estación de radio. El generador piloto engendra

oscilaciones no amortiguadas casi armónicas de aquella misma frecuencia y longitud de onda para la cual usted sintoniza su receptor si desea captar la transmisión de dicha estación. El segundo rectángulo-bloque representa el amplificador de potencia. El nombre habla por sí mismo, y en cuanto a su estructura, no la explicaremos aquí. La misión del bloque denominarlo «modulador» es transformar las oscilaciones acústicas en eléctricas y superponerlas en la onda portadora de la estación de radio.

La modulación puede realizarse por distintos procedimientos. Lo más simple es explicar cómo se lleva a cabo la modulación de frecuencia. En una serie de casos el micrófono no es sino un condensador cuya capacidad varía debido a la presión acústica, ya que la capacidad depende de la distancia entre las placas. Figúrense ahora que un condensador de este tipo está conectado al circuito oscilante que genera la onda. Entonces, la frecuencia de la onda cambiará en correspondencia con la presión acústica.

Por cuanto nosotros nos hemos «entremetido» con el micrófono en el circuito oscilante, resulta que al éter se envía no una frecuencia estrictamente determinada, sino cierta banda de frecuencias. Es bastante evidente que, idealmente, este ensanchamiento debe abarcar todo el intervalo acústico de frecuencias que, como sabemos, es igual a 20 kHz, aproximadamente.

Si la radioemisión se realiza en ondas largas a las que corresponden las frecuencias del orden de 100 Hz, entonces, la banda de transmisión (pasante) constituye la quinta parte de la frecuencia portadora. Está claro que será imposible asegurar en las ondas largas el trabajo de un gran número de estaciones que no se recubren. Los asuntos toman un cariz del todo distinto cuando se trata de ondas cortas. Para la frecuencia de 20 MHz el ancho de la banda constituirá ya fracciones de tanto por ciento respecto al valor de la frecuencia portadora.

En la Unión Soviética, a lo mejor, no se dé casa alguna en que no haya enchufe para la radio. Estas transmisiones se reciben de la llamada red de radiodistribución, que también se denomina radiodifusión por hilo.

En Moscú, la primera red de difusión por hilo destinada para un programa apareció en 1925. La transmisión se realizaba simultáneamente a través de 50 altavoces.

La radiodifusión de un programa se lleva a cabo en las audiofrecuencias. De los estudios de radio el programa se transmite por hilos a la estación amplificadora central. De la estación central, otra vez por hilos, las oscilaciones acústicas se transmiten a los puntos de base donde se amplifican una vez más y por las líneas de alimentación principales se transmiten a las subestaciones de transformadores. De cada subestación los hilos vuelven a distribuirse por las subestaciones de siguiente categoría, por decirlo así. En dependencia de las dimensiones de la ciudad o de la región el número de eslabones-de la red y, respectivamente, el número de reducciones de la tensión puede ser distinto. En las líneas de abonado, la tensión es igual a 30 V.

Desde 1962 en las ciudades de la Unión Soviética se introduce la radiodifusión por hilo de tres programas. La transmisión de los dos programas suplementarios se efectúa por el método de modulación de amplitud con las frecuencias portadoras 78 y 120 kHz. En su casa, usted demodulará estas dos transmisiones (es decir, separará el sonido y «seleccionará» la alta frecuencia) girando la manija del receptor enchufable para tres programas.

De este modo, en la radiodifusión de tres programas por un mismo hilo pasan simultáneamente tres programas: uno - el principal - en audiofrecuencias, y dos programas no demodulados. Por esta causa, sus transmisiones no se interfieren recíprocamente. Una idea sencilla, y el resultado, ¡excelente! El carácter económico, la seguridad y la alta calidad de la transmisión permiten considerar que ante la radiodifusión por hilo se abren grandes perspectivas, incluyendo la creación de redes de teledistribución.

Radiorrecepción

Existe un sinnúmero de construcciones de los radiorreceptores. La radioelectrónica se desarrolla con extraordinaria rapidez y, por lo tanto, los radiorreceptores, además, envejecen muy pronto, de modo que cada año aparecen en la venta nuevos modelos mejores.

¿Qué quiere decir la palabra «mejor» aplicada a un radiorreceptor? La respuesta la conoce cada uno de los lectores, incluso aquel que no está iniciado en la física. Un buen receptor debe seleccionar del caos de las ondas radioeléctricas que alcanzan la

antena solamente aquellas señales que se requieren. Esta propiedad lleva el nombre de selectividad. El receptor debe ser sensible en el máximo grado, es decir, debe captar las más débiles señales. Y, al fin, debe reproducir la música y el habla de la estación que hemos sintonizado sin distorsiones de ningún género. Para reproducir satisfactoriamente el habla de los locutores basta la banda de frecuencias desde 100 Hz hasta 1 kHz. Y el moderno jazz sinfónico requiere una banda desde 30 Hz hasta 20 kHz. La creación de una banda pasante tan amplia es un problema técnico difícil.

Resumiendo, podemos mencionar la sensibilidad, la selectividad y la precisión de la reproducción. Podemos añadir, tal vez, un deseo más: el receptor debe trabajar bien en todos los diapasones de ondas.

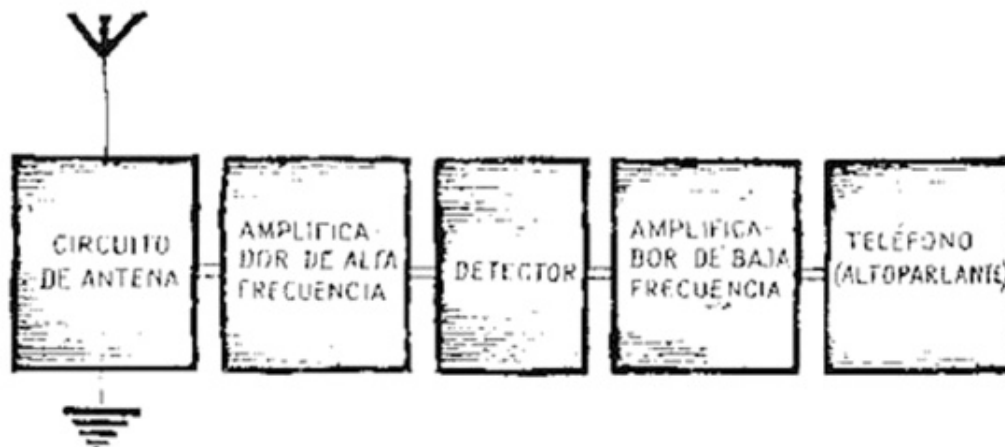


Figura 6.10

El diagrama sinóptico de un radioreceptor de amplificación directa es bastante evidente (fig. 6.10). Ante todo, hay que seleccionar la longitud necesaria de onda y amplificar las oscilaciones de alta frecuencia generadas en la antena por la onda de la estación que nos interesa. Luego, es necesario realizar la detección, o la demodulación: así es como se denomina el proceso de «rechazo» de la frecuencia portadora y de la separación de la corriente eléctrica de la información que trae el sonido. Después, es menester instalar un amplificador más, esta vez ya para las oscilaciones de baja frecuencia. Y la etapa final consiste en transformar estas oscilaciones eléctricas en acústicas, lo que se realiza por medio de un altavoz

electrodinámico o unos auriculares telefónicos que utilizan las personas delicadas que no quieren molestar a los vecinos.

De ordinario, la antena del radioreceptor está ligada de una manera inductiva con los circuitos oscilantes de varios diapasones.

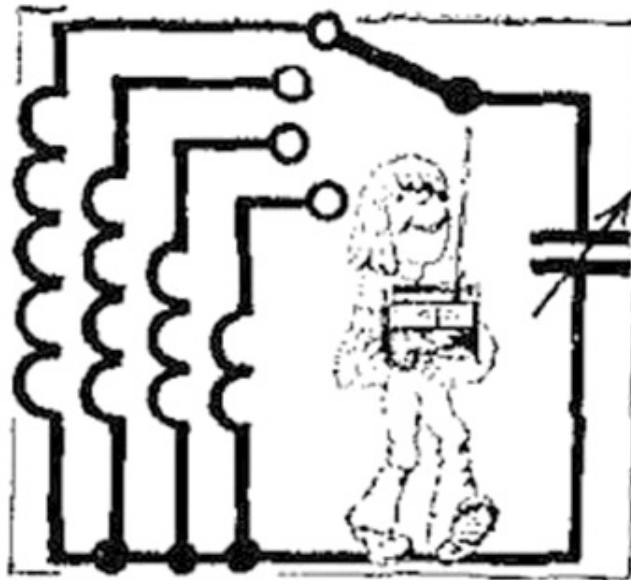


Figura 6.11

Cuando giramos la manija de los diapasones, realizamos la operación cuya representación esquemática se da en la fig. 6.11. Habitualmente, efectuamos la sintonización variando la capacidad del condensador del circuito oscilante receptor dentro de los límites de cada diapason. La capacidad del receptor de seleccionar la frecuencia de una forma óptima se determina por la curva de resonancia del circuito oscilante.

Tengo ante mí el certificado del receptor para automóviles. En el certificado se dice que para la sintonización a la frecuencia de mi estación predilecta, la señal perturbadora donde la estación que trabaja con la frecuencia desplazada a 20 kHz resultará debilitada en 60 dB, es decir, 1000 veces. En los mejores receptores modernos se consigue la selectividad de 120 dB (la reducción de la interferencia un millón de veces).

La sensibilidad de un receptor se caracteriza por la magnitud mínima de la f.e.m. en la antena del receptor que da la posibilidad de escuchar con suficiente claridad (de

20 a 30 dB por encima del nivel de los ruidos) la transmisión. En el radiorreceptor para automóviles la sensibilidad para las ondas largas es no peor de $175 \mu\text{V}$, y para el diapasón de ondas ultracortas, no peor de $5 \mu\text{V}$. En un automóvil es difícil montar una antena cuya longitud es más de 2 m. De aquí, es fácil hallar la intensidad de umbral del campo eléctrico de las ondas radioeléctricas buscadas. Si esta intensidad es menor, digamos, que $2 \mu\text{V/m}$, la señal útil se «ahogará» en los ruidos.

Los ruidos suelen ser de dos tipos: los extrínsecos, que incluyen ruidos industriales o los atmosféricos, y los intrínsecos, debidos a las fluctuaciones de las corrientes en los circuitos de entrada de los radiorreceptores. Cuando la señal que captamos es débil, disminuimos el ancho de la banda pasante. En este caso, los ruidos intrínsecos se reducen proporcionalmente a la raíz cuadrada del ancho de la banda pasante.

Propagación de las ondas radioeléctrica

El caso más simple es de propagar de la onda radioeléctrica en el espacio libre. Ya a una pequeña distancia del radiotransmisor éste puede considerarse como un punto. Y siendo así, el frente de la onda radioeléctrica puede tomarse por esférico. Si trazamos mentalmente varias esferas rodeando el radiotransmisor, queda claro que, en ausencia de absorción, la energía que pasa a través de las esferas permanecerá invariable. Ahora bien, la superficie de la esfera es proporcional al cuadrado del radio. Por consiguiente, la intensidad de la onda, o sea, la energía que corresponde a una unidad de superficie en unidad de tiempo disminuirá a medida de alejarse de la fuente y esta disminución será inversamente proporcional al cuadrado de distancia.

Por supuesto, esta importante regla es válida en el caso de que no se hayan tomado medidas especiales para crear un flujo estrechamente dirigido de ondas radioeléctricas.

Existen diferentes procedimientos técnicos para crear haces dirigidos de ondas radioeléctricas. Uno de los métodos de solución de este problema consiste en utilizar la correcta red de antenas. Las antenas deben disponerse de tal manera que las ondas que éstas emiten se dirijan en el sentido necesario «cresta con cresta». Con el mismo fin se emplean espejos de distintas formas.

Las ondas radioeléctricas que «viajan» por el espacio cósmico se desviarán de la dirección rectilínea, se reflejarán, se disiparán, se refractarán, en el caso de encontrar en su camino obstáculos conmensurables con la longitud de onda y hasta algo menores.

Para nosotros el mayor interés lo representa el comportamiento de las ondas que pasan cerca de la superficie terrestre. En cada caso separado el cuadro puede ser muy peculiar dependiendo de cuál será la longitud de onda.

El papel cardinal lo desempeñan las propiedades eléctricas de la tierra y la atmósfera. Si la superficie es capaz de conducir la corriente, entonces ésta «no permite» que las ondas radioeléctricas se aparten de ella. Las líneas de fuerza eléctricas del campo electromagnético se acercan al metal (y, más ampliamente, a cualquier conductor) formando un ángulo recto.

Ahora figúrense que la radiotransmisión tiene lugar cerca de la superficie del mar. El agua marina contiene sales disueltas, es decir, es un electrólito. El agua del mar es un magnífico conductor de la corriente. Esta es la razón por la cual «retiene» la onda radioeléctrica obligándola a moverse a lo largo de la superficie del mar.

Pero también las llanuras y los territorios cubiertos de bosques son buenos conductores para las corrientes de frecuencias no muy altas. En otras palabras, para las ondas largas el bosque y la llanura se comportan como un metal.

Por lo tanto, las ondas largas se retienen por toda la superficie terrestre y son capaces de dar vuelta al globo terráqueo. A propósito, utilizando este método se puede determinar la velocidad de las ondas radioeléctricas. Los radiotécnicos conocen que la onda radioeléctrica necesita 0,13 s para dar vuelta al globo terráqueo. ¿Y las montañas? Bueno, no olvide que para las ondas largas las montañas ya no son tan altas, y una onda radioeléctrica de un kilómetro de longitud es capaz, más o menos, de contornear una montaña.

En cuanto a las ondas cortas, resulta que la posibilidad de realizar la radiorrecepción a grandes distancias en estas ondas se debe a la existencia sobre la Tierra de la ionosfera. Los rayos solares poseen la capacidad de destruir las moléculas del aire en las zonas superiores de la atmósfera. Las moléculas se transforman en iones y a las distancias de 100 a 300 km de la Tierra forman varias capas cargadas. De este

modo, resulta que para las ondas cortas el espacio en que se mueve la onda es una capa de dieléctrico apretado entre dos superficies conductoras.

Por cuanto las áreas llanas y forestales no son buenos conductores para las ondas cortas, éstas no están en condiciones de retenerlas. Las ondas cortas emprenden un viaje libre, pero topan con la ionosfera que las refleja a guisa de superficie metálica. La ionización de la ionosfera no es homogénea y, por supuesto, es distinta de día y de noche. Debido a ello los caminos recorridos por las ondas radioeléctricas cortas pueden ser las más diversas. Pueden llegar a su radiorreceptor incluso después de haberse reflejado reiteradas veces por la Tierra y la ionosfera. El destino de la onda corta depende del ángulo bajo el cual incide en la capa ionosférica. Si este ángulo es próximo al recto, no se produce la reflexión y la onda se irá al espacio cósmico. Sin embargo, con mayor frecuencia tiene lugar una reflexión total y la onda retorna a la Tierra.

La ionosfera es transparente para las ondas ultracortas. Por esta razón, cuando se trata de dichas longitudes de onda, es posible la radiorrecepción dentro de los límites de la visibilidad directa o por medio de los satélites. Al dirigir la onda al satélite, podemos captar las señales reflejadas de éste a enormes distancias.

Los satélites abrieron una nueva época en la técnica de la radiocomunicación asegurando la posibilidad de la radiorrecepción y la recepción de televisión en ondas ultracortas.

Unas posibilidades interesantes nos depara la transmisión en ondas centimétricas, milimétricas y submilimétricas. Las ondas de esta longitud pueden absorberse por la atmósfera. Sin embargo, resulta que existen «ventanas», y, al elegir de una manera adecuada la longitud de onda es posible utilizar las ondas que se introducen en el diapasón óptico. Y en lo que se refiere a los méritos de estas ondas, ya los conocemos: en un intervalo pequeño de ondas se puede «meter» un enorme número de transmisiones que no se recubren.

Radiolocalización

Los principios de la radiolocalización son bastante sencillos. Enviamos una señal, ésta se refleja del objeto que nos interesa y retorna adonde estamos nosotros. Si el objeto se encuentra a la distancia de 150 m la señal regresará dentro de 1 μ s, y si

la distancia es de 150 km, el intervalo hasta su regreso será de 1 ms. La dirección en que se envía la señal es la dirección de la línea en que se hallaba el avión, el cohete o el automóvil en el momento en que el haz de radio lo encontró.

Se entiende que la onda radioeléctrica debe ir en un haz filiforme y el ángulo de abertura en que se concentra la parte principal de la potencia del haz debe ser del orden de un grado.

El principio, efectivamente, no es complicado, pero la técnica dista mucho de ser simple. Comenzamos con que se plantean requerimientos altos ante el generador. En los diapasones métrico y decimétrico (las ondas más largas, con toda evidencia, no sirven) se emplean generadores de tubos termoiónicos, y en el diapason centimétrico, klistrones y magnetrones.

Como más natural se presenta el método de trabajo por impulsos. Se envían periódicamente al espacio impulsos cortos. En las estaciones de radar contemporáneas la duración del impulso se encuentra en los límites de 0,1 a 10 μ s. La frecuencia con que el impulso se repite debe elegirse de tal forma que la señal reflejada tenga tiempo para volver durante la pausa.

La distancia máxima a que se puede detectar el avión o el cohete está limitada tan sólo por la condición de visibilidad directa. Sin duda alguna, el lector está enterado de que los radares modernos son capaces de recibir las señales reflejadas de cualesquiera planetas de nuestro Sistema Solar. Se sobreentiende que en este caso deben utilizarse las ondas que, sin obstáculos, pasan a través de la ionosfera. Es una circunstancia afortunada que el acortamiento de la longitud de onda influye también de una manera directa en el aumento de la distancia de la visibilidad de localización, por cuanto esta distancia es proporcional no sólo a la energía del impulso enviado, sino también a la frecuencia de radiación.

En la pantalla de un oscilógrafo (un tubo de rayos catódicos) se pueden observar saltos de los impulsos enviado y reflejado. Si el avión se acerca, entonces, la señal reflejada se desplazará en el sentido de la enviada.

No es obligatorio que los radares trabajen en el régimen de impulsos. Supongamos que el avión se mueve en la dirección de la antena con la velocidad v . Desde este avión se refleja continuamente el haz de radio. El efecto Doppler lleva a que la frecuencia de la onda recibida está relacionada con la frecuencia de la onda enviada

mediante la ecuación

$$v_{rec} = v_{env} \left(1 + \frac{2v}{c} \right)$$

Los métodos radiotécnicos permiten determinar las frecuencias con gran precisión. Al realizar el reglaje de resonancia hallamos v_{ref} y por su valor calculamos la velocidad. Si, por ejemplo, la frecuencia de la señal enviada es igual a 10^9 Hz y el avión o el cohete avanzan hacia la antena del radar con la velocidad de 1000 km/h, resulta que la frecuencia recibida será mayor que la transmitida en una magnitud de 1850 Hz.

La reflexión del haz de radio desde un avión, un cohete, desde una motonave o automóvil no es la reflexión desde un espejo. Las longitudes de onda son conmensurables o bien mucho menores que las dimensiones del objeto reflector que tiene una forma complicada. Al reflejarse desde distintos puntos del objeto entre los rayos habrá interferencia recíproca y estos se disiparán por los lados. Estos dos fenómenos darán lugar a que la superficie reflectora efectiva del objeto se diferenciará considerablemente de su superficie real. El cálculo en este caso es complicado y sólo la experiencia del operador que hace uso del radar lo ayuda a decir qué objeto se ha encontrado en el camino del rayo.

El lector, de seguro, ha visto las antenas del radar, o sea, espejos esféricos de alambre que todo el tiempo se encuentran en movimiento ya que están sondeando el espacio. Es posible obligar al espejo del radar a realizar los más diversos movimientos, por ejemplo, tales que el rayo se mueva barriendo el espacio con líneas o circunferencias. Con este trabajo se puede no sólo determinar la distancia hasta el avión, sino también seguir la trayectoria de su movimiento.

Empleando este método el avión se dirige al aterrizaje en condiciones de ausencia de visibilidad. Semejante tarea puede encomendarse tanto al hombre, como a un autómatas.

Al radar se lo puede «engañar». En primer lugar, el objeto puede empantallarse con materiales que absorben las ondas radioeléctricas. Para alcanzar esta finalidad sirven el polvo de carbón y el caucho. En este caso, además, para disminuir el

coeficiente de reflexión, los recubrimientos se hacen ondulados, y este método permite conseguir que la parte leonina de radiación se disipe de una forma caótica en todas las direcciones.

Si desde el avión se arrojan en paquetes tiras de hoja de aluminio o de fibra metalizada, el radar resultará completamente desorientado. Por primera vez este método lo aplicaron los ingleses aún en los años de la segunda guerra mundial. El tercer método consiste en llenar el éter de radioseñales falsas.

La radiolocalización es un capítulo interesantísimo de la técnica que encuentra un amplio uso para muchos fines pacíficos y sin la cual, en la actualidad, es imposible figurar los medios de defensa.

Los principios de radiolocalización forman la base del enlace entre las naves cósmicas y la Tierra. Los radiotelescopios se sitúan de tal modo que la nave no se pierda de vista. Sus antenas tienen enormes dimensiones: hasta centenares de metros. El hecho de que se requieren antenas tan grandes se explica por la necesidad de enviar una señal muy fuerte y captar una señal débil del radiotransmisor. Es natural que sea de suma importancia tener un haz de radio estrecho. Si la antena trabaja a la frecuencia de 2.200 millones de oscilaciones por segundo (la longitud de onda es de 4 cm, aproximadamente), entonces, en la distancia hasta la Luna, el haz se ensancha tan sólo hasta el diámetro de 1000 km. Ahora bien, cuando el haz llegue a Marte (300×10^6 km), su diámetro ya será igual a 700 000 km.

Televisión

Por cuanto 99 lectores de los 100 pasan diariamente una que otra hora frente al televisor sería injusto no decir varias palabras acerca de este gran invento. Actualmente, hablaremos sólo de los principios de la transmisión televisiva.

La idea de transmitir las imágenes a distancia se reduce a lo siguiente. La imagen que se transmite se divide en cuadrados pequeños. El fisiólogo sugerirá cuál debe ser el tamaño del cuadrado para que el ojo deje de advertir las variaciones del brillo dentro de ésta. La energía luminosa de cada porción de la imagen puede transformarse en señal eléctrica valiéndose del efecto fotoeléctrico. Hay que inventar el método con cuya ayuda sería posible leer estas señales. Por supuesto, la

operación se lleva a cabo en un orden estrictamente determinado, como durante la lectura de un libro. Estas señales eléctricas se superponen en la onda magnética portadora mediante el procedimiento completamente análogo a aquel que se emplea durante la radiotransmisión. A continuación, los acontecimientos se desarrollan de una forma absolutamente idéntica a la radiocomunicación. Las oscilaciones moduladas se amplifican y detectan. El televisor debe transformar los impulsos eléctricos en imagen óptica.

Los tubos transmisores de televisión llevan el nombre de vidicones. Con ayuda de una lente la imagen se proyecta en el fotocátodo. Los fotocátodos más difundidos son el de oxígeno-cesio y de antimonio-cesio. El fotocátodo se monta en una ampolla de vacío junto con el fotoánodo.

De principio, la imagen podría transmitirse proyectando alternativamente el flujo luminoso de cada elemento de la misma. En este caso, la corriente fotoeléctrica debe circular solamente durante un corto tiempo, mientras dura la transmisión de cada elemento de la imagen. Sin embargo, semejante funcionamiento sería incómodo, y en el tubo transmisor se utiliza no una sola célula fotoeléctrica, sino un número grande de éstas igual al número de elementos en que se descompone la imagen transmitida. La placa receptora se denomina blanco y se confecciona en forma de mosaico.

El mosaico es una placa fina de mica por uno de cuyos lados está aplicado un gran número de granitos de plata, aislados y cubiertos de óxido de cesio. Cada granito es una célula fotoeléctrica. Por otro lado de la placa de mica está aplicada una película metálica. Entre cada granito del mosaico y el metal parece como si se formase un condensador pequeño que se carga por los electrones arrancados al cátodo. Está claro que la carga de cada condensador será proporcional al brillo del correspondiente punto de la imagen transmitida.

De este modo, en la placa metálica se forma una especie de imagen eléctrica latente del objeto. ¿Cómo puede ésta tomarse de dicha placa? Valiéndose del haz electrónico que debe obligarse a recorrer la placa de la misma manera como el ojo recorre las líneas de un libro. El haz electrónico hace las veces de llave que cierra instantáneamente el circuito eléctrico a través de un microcondensador. La corriente en este circuito creado momentáneamente estará vinculada de modo unívoco con el

brillo de la imagen. Cada señal puede y debe amplificarse múltiples veces por medio de procedimientos ordinarios empleados en la radiotecnica. Durante la transmisión de la imagen el ojo no debe advertir que el haz electrónico recorre consecutivamente los diferentes puntos de la pantalla luminosa. La imagen total se obtiene en la pantalla del tubo receptor por un ciclo de movimiento del haz electrónico. Se debe crear tal frecuencia de imagen que a costa de la persistencia de las imágenes en la retina no se observe el efecto de parpadeo (llamado también centelleo).

¿Y qué frecuencia de imagen debe tomarse?

La sensación de continuidad del movimiento surge cuando la frecuencia de imagen es de cerca de 20 Hz, por esta causa, en la televisión la frecuencia de imagen se toma igual a 25 Hz; sin embargo, con esta frecuencia el efecto de parpadeo todavía es perceptible, por cuya razón los técnicos han recurrido a un método interesante: a saber, la exploración entrelazada. Se ha dejado la frecuencia de 25 Hz, pero el haz electrónico recorre primero las líneas impares y, luego, las líneas pares. La frecuencia con que cambian las semiimágenes se hace igual a 50 Hz, de modo que el centelleo de la imagen se suprime.

Las frecuencias de exploración de imágenes y líneas deben estar estrictamente sincronizadas. En la Unión Soviética la imagen suele dividirse en 625 líneas. Y el número de elementos de la imagen en cada línea es tantas veces mayor que el número de líneas cuantas veces la longitud de la línea supera la altura de la pantalla. En los televisores soviéticos esta relación es igual a $4/3$. De este modo, tenemos que transmitir 25 veces por segundo $4/3$ de los 625^2 elementos de la imagen. Un período basta para transmitir dos elementos. De aquí se infiere que la imagen de la televisión exige una banda con un ancho no menor de 6,5 MHz. Por consiguiente, la frecuencia portadora del emisor no puede ser menor que 40 - 50 MHz, ya que la frecuencia de la onda portadora debe ser por lo menos 6 a 7 veces mayor que el ancho de la banda de las frecuencias transmitidas. Ahora el lector comprenderá por qué para las transmisiones de televisión se pueden utilizar sólo ondas ultracortas y por qué, como consecuencia, el alcance de estas transmisiones está limitado por la visibilidad directa.

Pero he dicho mal: estaba limitado. El acontecimiento revolucionario que da la posibilidad de efectuar la transmisión de televisión a cualesquiera distancias es la utilización de los satélites de comunicación que se ven tanto del punto de recepción, como del de transmisión. La Unión Soviética fue el primer país que utilizó los satélites con esta finalidad. Actualmente, los vastos espacios de la URSS están abarcados por una red de comunicación realizada por una serie de satélites.

Sin tocar la cuestión acerca de la estructura de las potentes estaciones de televisión aducimos tan sólo cifras interesantes que caracterizan las enormes posibilidades de la radiotecnica moderna en el campo de amplificación de las señales. Una videoseñal corriente antes de amplificarla tiene la potencia hasta 10^{-3} W; el amplificador de la potencia la amplifica un millón de veces. La potencia de 10^3 W se suministra a la antena parabólica cuyo diámetro es del orden de 30 m. Esta antena engendra un haz estrechamente dirigido que será reflejado por el satélite. Después de que la onda electromagnética recorra unos 35.000 km hasta el satélite su potencia será igual a 10^{-11} W.

El amplificador montado en el satélite aumenta la potencia de esta señal muy débil hasta de 10 W, aproximadamente. A su vez, la señal reflejada del satélite volverá a la Tierra con una potencia de 10^{-17} W. La amplificación hará retornar la potencia de la videoseñal a su valor de partida igual a 10^{-3} W.

Yo pienso que hace dos lustros ni siquiera un ingeniero con criterios más optimistas habría dado crédito a estas cifras. La eficacia de los dispositivos de recepción aplicados se determina por el producto de las dimensiones de la antena por la amplificación útil del receptor, la cual oscila desde un millón hasta cien millones. En el límite superior, para reducir los ruidos internos, se tiene que recurrir al enfriamiento de la primera etapa de amplificación hasta la temperatura de helio líquido.

Microcircuitos electrónicos

No se puede terminar el capítulo dedicado a la radiotecnica sin decir aunque sea varias palabras sobre la nueva revolución que se desarrolla ante nuestros ojos.

Se trata de la fantástica miniaturización de todos los aparatos radiotécnicos que llegó a ser posible debido al paso desde los aparatos constituidos por elementos

aislados: resistencias, transistores, etc., conectados entre sí con la ayuda de alambres hacia los circuitos eléctricos «dibujados» mediante una técnica especial en un pedacito de silicio de varios milímetros de tamaño.

La nueva tecnología (una de sus variantes) consiste en que, al utilizar estarcidos de distintos tipos y una serie de productos químicos se podría introducir en los puntos necesarios del cristal de silicio o germanio impurezas p e impurezas n. Con este fin se aplica el tratamiento con haces iónicos.

Un circuito eléctrico que consta de decenas de miles de elementos (!) se ubica en un área con dimensiones lineales de cerca de dos milímetros. Cuando hemos dicho «dibujar» el circuito en la mente del lector podía crearse la impresión de que nos referimos tan sólo al tratamiento de la superficie de un pedacito de semiconductor. Pero no es así. El asunto es más complicado. Cada elemento radiotécnico posee una estructura tridimensional. En una diminuta porción de silicio es necesario crear varias capas que contienen diferentes cantidades de impurezas.

Entonces, ¿qué es preciso hacer para conseguir este objetivo? En primer lugar, en la superficie del silicio se crea una capa de óxido. Sobre esta capa se aplica un material fotosensible. El «bizcocho» obtenido se irradia con luz ultravioleta a través del estarcido de forma calculada. Después de revelar, en la superficie del pedacito de silicio, en los puntos en que la luz pasó a través del estarcido, se forman hoyos.

La siguiente etapa consiste en tratar el futuro circuito radioeléctrico con ácido fluorhídrico. Dicho ácido quita el óxido de silicio sin afectar la superficie primitiva (es decir, el silicio), ni tampoco la capa fotosensible. Ahora queda el último paso: tratamiento con otro disolvente que eliminará la capa fotosensible. Como resultado, nuestro pedacito queda recubierto de aislador, de óxido de silicio, allí donde lo requiere el diseño. Y el hoyo con forma necesaria es el silicio dejado al descubierto. Precisamente éste se somete al tratamiento con el haz iónico para introducir la cantidad necesaria de impurezas.

La creación de los microcircuitos electrónicos es, hoy en día, una de las ramas más activas de la técnica.

Las nuevas ideas y descubrimientos en el ámbito de la física de los semiconductores demuestran que los resultados fantásticos alcanzados para el día de hoy distan mucho de representar el límite.

F I N