

Richard P. Feynman

Electrodinámica cuántica



La electrodinámica cuántica es una de las teorías físicas más precisas pero también más complicadas. En este libro («una aventura que, por lo que sabemos, nunca se había intentado», señala en el prefacio Ralph Leighton), RICHARD P. FEYNMAN (1918-1988), premio Nobel de Física en 1965 por sus contribuciones al desarrollo de la electrodinámica cuántica, presenta esa teoría con la claridad, la precisión y la exhaustividad que le hicieron famoso. Suponiendo escasos conocimientos científicos en los lectores y profundizando en el contenido intuitivo y visual de la teoría, Feynman —uno de los físicos más geniales de nuestro siglo— describe la interacción entre luz y electrones, «absurda» desde el punto de vista del sentido común pero que se encuentra en la base de prácticamente todo lo que observamos en el mundo físico. Las páginas de ELECTRODINÁMICA CUÁNTICA explican satisfactoriamente fenómenos tan familiares como la luz reflejándose en un espejo o curvándose cuando pasa del aire al agua. Dos aspectos de este libro serán especialmente apreciados por todos aquellos interesados en la física moderna: por un lado, la forma en que Feynman introduce sus célebres «diagramas», una herramienta absolutamente fundamental para el estudio y aplicación de la electrodinámica cuántica; por otro, su utilización de los «camino posibles», con los que construyó su conocida interpretación de la mecánica cuántica. Finalmente, Feynman explica cómo la electrodinámica cuántica ayuda a comprender los quarks, gluones y otros elementos fundamentales para la física actual.



Richard Phillips Feynman

Electrodinámica cuántica

La extraña teoría de la luz y la materia

ePub r1.3

Piolin 14.08.2018

Título original: *QED. The Strange Theory of Light and Matter*

Richard Phillips Feynman, 1985

Traducción: Ana Gómez Antón

Diseño de portada: Daniel Gil

Editor digital: Piolin

Primer editor (r1.0): Antwan

ePub base r1.2



PRÓLOGO

Las conferencias en memoria de Alix G. Mautner fueron concebidas en honor de mi esposa Alix, que murió en 1982. Aunque realizó su carrera en literatura inglesa, Alix tenía un interés grande y permanente en muchos campos científicos. De modo que parecía adecuado crear una fundación, en su nombre, que patrocinase una serie de conferencias anuales con objeto de comunicar a un público inteligente e interesado el espíritu y los logros de la ciencia.

Estoy encantado de que Richard Feynman haya aceptado dar las primeras conferencias. Nuestra amistad se remonta a cincuenta años atrás, a nuestra niñez en Far Rockaway, Nueva York. Richard conoció a Alix a lo largo de veintidós años, y durante mucho tiempo deseó que desarrollara una explicación de la física de las partículas pequeñas que fuese inteligible para ella y otras personas que no fuesen físicos.

Como nota suplementaria, me gustaría expresar mi aprecio a todos los que contribuyeron a la Fundación Alix G. Mautner y que de esta forma ayudaron a hacer posible estas conferencias.

Leonard Mautner

Los Angeles, California.

Mayo 1983

PREFACIO

Richard Feynman es legendario en el mundo de la física por la manera en que ve el mundo: nunca da nada por sabido y siempre piensa las cosas por sí mismo —a menudo alcanza un conocimiento nuevo y profundo del comportamiento de la naturaleza— con un estilo elegantemente sencillo y refrescante de describirlo.

También es conocido por su entusiasmo al explicar física a sus estudiantes. Después de rechazar innumerables ofertas para pronunciar discursos en prestigiosas sociedades y organizaciones, Feynman siempre está ávido del estudiante que llega a su despacho y le pide que dé una charla en el club de física del instituto local.

Este libro es una aventura que, por lo que sabemos, nunca se ha intentado. Es una explicación directa, honesta, de un tema muy difícil —la teoría de la electrodinámica cuántica— para una audiencia no experta. Está concebido para dar al lector interesado una apreciación del tipo de pensamiento al que los físicos han acudido a fin de explicar cómo se comporta la Naturaleza.

Si usted ha planeado estudiar física (o ya lo está haciendo) no hay nada en este libro que tenga que ser «olvidado»: es una descripción completa, exacta en cada detalle, de un marco al que se pueden incorporar conceptos más avanzados sin necesidad de modificación. Para aquellos que ya hayan estudiado física, ¡es una revelación de lo que estuvieron haciendo *realmente* cuando estuvieron desarrollando todos aquellos cálculos complicados!

De niño, Richard Feynman se sintió inclinado a estudiar cálculo al leer un libro que comentaba, «Lo que puede hacer un loco, otro lo puede hacer». A él le gustaría dedicar este libro a los lectores con unas palabras similares «Lo

que puede entender un loco, otro lo puede entender».

Ralph Leighton.

Pasadena, California.

Febrero 1985.

AGRADECIMIENTO

Este libro se supone que son las actas de las conferencias que sobre electrodinámica cuántica di en UCLA, transcritas y editadas por mi buen amigo Ralph Leighton. En realidad, el manuscrito ha experimentado considerables modificaciones. La experiencia de Mr. Leighton en la docencia y en el arte de escribir han sido de valor considerable en este intento de presentar esta parte importante de la física a una audiencia más amplia.

Muchas exposiciones «populares» de la ciencia logran una simplicidad aparente al describir algo diferente, algo considerablemente distorsionado, de lo que pretenden estar describiendo. El respeto hacia nuestro tema no nos ha permitido hacer esto. A lo largo de muchas horas de discusión, hemos intentado alcanzar el máximo de claridad y de sencillez sin compromisos con la distorsión de la verdad.

Capítulo 1

INTRODUCCIÓN

Alix Mautner sentía una gran curiosidad por la física y, a menudo, me pedía que le explicara cosas. Lo hacía sin problemas, tal como lo hago con un grupo de estudiantes de Caltech que acuden a mí cada jueves una hora, pero en ocasiones fallé en lo que considero la parte más interesante: siempre nos quedábamos atascados en las locas ideas de la mecánica cuántica. Le decía que no podía explicarle esas ideas en una hora o en una velada —requerían mucho tiempo— pero le prometí que algún día prepararía una serie de conferencias sobre el tema.

Preparé algunas conferencias y fui a Nueva Zelanda a ensayarlas ¡probablemente porque Nueva Zelanda está tan alejada que si no tenían éxito no importaría! Bien, la gente de Nueva Zelanda pensó que eran correctas, de manera que supuse que lo eran ¡al menos para Nueva Zelanda! Por consiguiente, aquí están las conferencias que preparé en realidad para Alix pero que desafortunadamente, ahora, no puedo pronunciárselas a ella directamente.

De lo que quiero hablar es de una parte de la física que es *conocida*, no de una parte desconocida. La gente siempre está preguntando por los últimos desarrollos en la unificación de esta teoría con aquella otra, y no nos da la oportunidad de explicarles nada sobre las teorías que conocemos bastante bien. Siempre quieren conocer cosas que no sabemos. De manera que en lugar de confundirles con un montón de teorías a medio hacer y parcialmente analizadas, me gustaría hablarles de un tema que ha sido completamente analizado. Me gusta este área de la física y pienso que es maravillosa: es lo que denominamos electrodinámica cuántica, o QED para abreviar.

Mi objetivo principal en estas conferencias es describir, de la forma más precisa posible, la extraña teoría de la luz y la materia —o de manera más específica, la interacción de la luz y los electrones—. Va a llevar bastante tiempo explicar todo lo que quiero. Sin embargo, hay cuatro conferencias, de modo que me tomaré el tiempo necesario y todo saldrá bien.

La física tiene un historial de sintetizar muchos fenómenos en unas cuantas teorías. Por ejemplo, al principio existían fenómenos de movimiento y fenómenos de calor; había fenómenos de sonido, de luz y de gravedad. Pero pronto se descubrió, después de que Sir Isaac Newton explicase las leyes del movimiento, que algunas de las cosas aparentemente distintas eran aspectos de la misma cosa. Por ejemplo, el fenómeno del sonido podría comprenderse completamente mediante el movimiento de los átomos en el aire. De manera que el sonido no se consideró nunca más como algo separado del movimiento. También se descubrió que los fenómenos caloríficos se explicaban con facilidad a partir de las leyes del movimiento. De esta forma, amplias esferas de la teoría física se sintetizaron en una teoría simplificada. Por otro lado, la teoría de la gravitación no se entendía mediante las leyes del movimiento, e incluso en la actualidad permanece aislada de otras teorías. Hasta ahora, la gravitación no puede ser explicada en términos de otros fenómenos.

Tras la síntesis de los fenómenos de movimiento, sonido y calor, tuvo lugar el descubrimiento de un número de fenómenos que denominamos eléctricos y magnéticos. En 1873 estos fenómenos fueron unificados con los fenómenos luminosos y ópticos en una única teoría por James Clerk Maxwell, quien propuso que la luz es una onda electromagnética. De manera que entonces estaban las leyes del movimiento, las leyes de la electricidad y magnetismo, y las leyes de la gravedad.

Alrededor de 1900 se desarrolló una teoría para explicar lo que era la materia. Se denominó la teoría electrónica de la materia, y decía que existían pequeñas partículas cargadas dentro de los átomos. Esta teoría evolucionó gradualmente hasta incluir un núcleo pesado con electrones girando en torno a él.

Los intentos para comprender el movimiento de los electrones alrededor del núcleo utilizando leyes mecánicas —de forma análoga a como Newton utilizó las leyes del movimiento para descubrir cómo la tierra giraba alrededor del sol— fueron un verdadero fracaso: todo tipo de predicciones resultaron erróneas. (Incidentalmente, la teoría de la relatividad que todos Vds. consideran ser una gran revolución en la física, fue también desarrollada en

esta época. Pero comparada con el descubrimiento de que las leyes del movimiento de Newton no eran válidas para los átomos, la teoría de la relatividad era sólo una pequeña modificación). El elaborar otro sistema que reemplazase a las leyes de Newton llevó largo tiempo porque los fenómenos a nivel atómico eran bastante extraños. Había que perder el sentido común para percibir lo que estaba ocurriendo a nivel atómico. Finalmente, en 1926, se desarrolló una teoría «insensata» para explicar «la nueva forma de comportamiento» de los electrones en la materia. Parecía disparatada, pero en realidad no lo era: se denominó la teoría de la mecánica cuántica. La palabra «cuántica» hace referencia a ese peculiar aspecto de la naturaleza que va en contra del sentido común. Es de este aspecto del que voy a hablarles.

La teoría de la mecánica cuántica también explicaba todo tipo de detalles, como el por qué se combina un átomo de oxígeno con dos de hidrógeno para formar agua, y demás cosas. La mecánica cuántica suministró así la teoría a la química. De modo que la química teórica fundamental es realmente física.

Debido a que la teoría de la mecánica cuántica podía explicar toda la química, y las distintas propiedades de las sustancias, fue un éxito tremendo. Pero aún quedaba el problema de la interacción de la luz y la materia. Es decir, la teoría de Maxwell de la electricidad y el magnetismo tenía que ser modificada de acuerdo con los nuevos principios de la mecánica cuántica. De modo que una nueva teoría, la teoría cuántica de la interacción de la luz y la materia, que es conocida con el horrible nombre de «electrodinámica cuántica» fue finalmente desarrollada por un número de físicos en 1929.

Pero la teoría tenía problemas. Si se realizaban cálculos groseros, se obtenía una respuesta razonable. Pero si se intentaban hacer cálculos más precisos se encontraba que la corrección que se pensaba que iba a resultar pequeña (el siguiente término de la serie, por ejemplo) era de hecho muy grande —de hecho ¡era infinito!—. Por lo que resultó que no se podía calcular *nada* que excediese una cierta precisión.

Por cierto, lo que acabo de esbozar es lo que yo llamo «una historia de la física por un físico», que no es nunca correcta. Lo que les estoy relatando es una especie de historia-mito convencional que los físicos cuentan a sus estudiantes y estos estudiantes se la repiten a los suyos, y no está necesariamente relacionada con el desarrollo histórico real, ¡el cuál desconozco!

De cualquier manera, para proseguir con esta «historia», Paul Dirac,

utilizando la teoría de la relatividad, elaboró una teoría relativista del electrón que no tenía totalmente en cuenta todos los efectos de la interacción del electrón con la luz. La teoría de Dirac establece que el electrón tiene un momento magnético —algo similar a la fuerza de un pequeño imán— que posee un valor exactamente igual a 1, en determinadas unidades. Luego, alrededor de 1948, se descubrió a través de los experimentos que el número real era próximo a 1,00118 (con una incertidumbre de alrededor de 3 en el último dígito). Se sabía, naturalmente, que los electrones interaccionan con la luz, por lo que se esperaba una pequeña corrección. También se esperaba que esta corrección fuese explicable con la nueva teoría de la electrodinámica cuántica. Pero cuando se calculó, en lugar de 1,00118 el resultado era infinito ¡lo que era incorrecto, experimentalmente!

Bien, este problema de cómo calcular las cosas en electrodinámica cuántica, fue solventado por Julián Schwinger, Sin-Itiro Tomonaga, y yo mismo alrededor de 1948. Schwinger fue el primero en calcular esta corrección utilizando un nuevo «juego de capas», su valor teórico era aproximadamente 1,00116, lo suficientemente próximo al valor experimental como para demostrar que estábamos en el buen camino. ¡Al fin, teníamos una teoría cuántica de la electricidad y el magnetismo con la que se podían realizar cálculos! Esta es la teoría que voy a describirles.

La teoría de la electrodinámica cuántica lleva en vigor más de cincuenta años, y ha sido ensayada con precisión cada vez mayor en un rango cada vez más extenso de condiciones. En la actualidad puedo decir orgullosamente ¡que *no existe diferencia apreciable* entre teoría y experimento!

Para darles idea de cómo esta teoría ha sido puesta a prueba, les daré algunos números recientes: los experimentos habían dado para el número de Dirac un valor de 1,00115965221 (con una incertidumbre de 4 en el último dígito); la teoría lo coloca en 1,00115965246 (con una incertidumbre como mucho cinco veces superior). Para que capten la precisión de estos números les diré algo como que: si se midiese la distancia de Los Angeles a Nueva York con semejante precisión, su valor diferiría del correcto en el espesor de un cabello humano. Este es el grado de sutileza con que ha sido probada la electrodinámica cuántica durante los últimos cincuenta años —tanto de manera teórica como experimental—. Por cierto que sólo he escogido un valor como muestra. Existen otras cosas en la electrodinámica cuántica que se han medido con una precisión comparable y que también concuerdan perfectamente. Se han comprobado cosas a escalas de distancia que varían

desde cien veces el tamaño de la Tierra hasta la centésima parte del tamaño del núcleo atómico. ¡Estos números son para intimidarles, para que crean que la teoría probablemente no está muy descaminada! Antes de que terminemos, les describiré cómo se realizan los cálculos.

Me gustaría impresionarles de nuevo con el amplio rango de fenómenos que la teoría de la electrodinámica cuántica describe: Es muy fácil decirlo en retrospectiva: la teoría describe *todos* los fenómenos del mundo físico excepto el efecto gravitacional, el que les mantiene a Vds. en su asiento (en realidad, es una cuestión de gravedad y cortesía, me temo), y los fenómenos radiactivos, que implican al núcleo desplazándose en sus niveles energéticos. De modo que si dejamos a un lado la gravitación y la radiactividad (más exactamente, la física nuclear), ¿qué nos queda? La gasolina que se quema en los automóviles, la espuma y las burbujas, la dureza de la sal o del cobre, la rigidez del acero. De hecho, los biólogos están intentando interpretar el mayor número posible de hechos de la vida en términos de química, y, como ya he explicado, la teoría que se esconde tras la química es la electrodinámica cuántica.

Debo aclarar algo: Cuando digo que todos los fenómenos del mundo físico se pueden explicar mediante esta teoría, en realidad no lo sabemos. La mayoría de los fenómenos que nos son familiares implican un número tan *tremendo* de electrones que es difícil para nuestras pobres mentes el seguir tal complejidad. En semejantes situaciones, podemos usar la teoría para hacernos una idea aproximada de lo que tiene que ocurrir y esto es lo que ocurre, aproximadamente, bajo esas circunstancias. Pero si disponemos en el laboratorio de un experimento que implique sólo unos *pocos* electrones bajo circunstancias *sencillas*, podemos calcular lo que puede ocurrir en forma muy precisa, y medirlo de manera también muy precisa. Cada vez que realizamos estos experimentos, la teoría de la electrodinámica cuántica funciona muy bien.

Los físicos siempre estamos comprobando si hay algo mal en la teoría. Este es el juego, porque si *hay* algo mal, ¡es interesante! Pero hasta ahora, no hemos encontrado nada equivocado en la teoría de la electrodinámica cuántica. Por tanto, yo diría que es la joya de la física —la posesión de la que estamos más orgullosos.

La teoría de la electrodinámica cuántica es también el prototipo de las nuevas teorías que intentan explicar los fenómenos nucleares, las cosas que tienen lugar dentro del núcleo de los átomos. Si se considerase el mundo

físico como un escenario, los actores no sólo serían los electrones, que están fuera del núcleo de los átomos, sino también los quarks y gluones y todos los demás —docenas de tipos de partículas— que están en el interior del núcleo. Y aunque estos «actores» parecen muy distintos entre sí, todos tienen un cierto estilo de actuar —un estilo extraño y peculiar—, el estilo «cuántico». Al final les hablaré un poco de las partículas nucleares. Entre tanto, voy a hablarles de los fotones —partículas de luz— y de los electrones, para hacerlo más sencillo. Porque es *la manera* en que actúan lo que importa, y esta manera es muy interesante.

De modo que ya saben de qué voy a hablarles. La siguiente pregunta es ¿entenderán Vds. lo que voy a contarles? Todos los que acuden a una conferencia científica saben que no van a entenderla, pero quizás el conferenciante tenga una bonita corbata coloreada a la que mirar. ¡Este no es el caso! (Feynman no lleva corbata). Lo que les voy a contar es lo que enseñamos a nuestros estudiantes de física en el tercer o cuarto curso de nuestra escuela graduada, y ¿Vds. creen que yo se lo voy a explicar de manera que lo entiendan? No, Vds. no van a ser capaces de comprenderlo. Entonces ¿por qué molestarles con todo esto? ¿Por qué van a permanecer ahí sentados todo este tiempo, cuando van a ser incapaces de entender lo que les voy a decir? Es mi deber convencerles de que no se vayan porque no lo entiendan. Verán, mis estudiantes de física tampoco lo entienden. Nadie lo entiende.

Me gustaría hablarles un poco más acerca del entendimiento. Cuando nos dan una conferencia existen muchas razones para no comprender al orador. Una, su lenguaje es malo —no sabe lo que quiere decir, o lo dice de forma desordenada— y es difícil de entender. Esto es una cuestión trivial y yo haré lo imposible para evitar en el mayor grado mi acento de Nueva York.

Otra posibilidad, especialmente cuando el conferenciante es físico, es que utilice palabras comunes con un sentido curioso. Los físicos utilizan con frecuencia palabras como «trabajo» o «acción» o «energía» o incluso, como verán «luz» con propósitos técnicos. Así, cuando hablo de «trabajo» en física, no quiero expresar lo mismo que cuando hablo de «trabajo» en la calle. Durante esta conferencia podría usar alguna de estas palabras sin darme cuenta de que lo estoy haciendo en el sentido inhabitual. Intentaré evitarlo —es mi deber—, pero es un error fácil de cometer.

La siguiente razón que pudieran pensar para explicar la ininteligibilidad de lo que les estoy diciendo es que mientras yo les estoy describiendo *cómo* funciona la Naturaleza, Vds. no entenderán *por qué* funciona así. Pero nadie

lo entiende. No puedo explicar por qué la Naturaleza se comporta de esta forma peculiar.

Finalmente, existe esta posibilidad: que después de decirles algo, Vds. no se lo crean. No puedan aceptarlo. No les gusta. Un velo cae sobre Vds. y ya no escuchan más. Voy a describirles cómo es la Naturaleza, y si no les gusta, esto va a interferir con su forma de comprender. Es un problema que los físicos han aprendido a manejar: han aprendido a percibir que el que les guste o no una teoría *no* es el punto esencial. Más bien lo que importa es si la teoría proporciona o no predicciones en consonancia con los experimentos. No es cuestión de si la teoría es una delicia filosófica, o es fácil de entender, o es perfectamente razonable desde el punto de vista del sentido común. La teoría de la electrodinámica cuántica describe a la naturaleza de manera absurda desde el punto de vista del sentido común. Y concuerda completamente con los experimentos. De manera que espero que acepten la Naturaleza como es —absurda.

Me voy a divertir hablándoles de esta absurdidad porque la encuentro deliciosa. Por favor, no abandonen porque no puedan creer que la Naturaleza sea tan extraña. Escúchenme hasta el final, y espero que cuando acabemos estén tan encantados como yo.

¿Cómo voy a explicarles las cosas que no explico a mis alumnos hasta el tercer año de carrera? Déjenme que se lo exponga mediante una analogía. Los Mayas estaban interesados en el amanecer y la puesta de Venus como «estrella» matutina y como «estrella» vespertina —estaban muy interesados en saber cuándo aparecía—. Después de varios años de observación, notaron que los cinco ciclos de Venus se aproximaban mucho a ocho «años nominales» de 365 días (estaban al corriente de que el año verdadero de estaciones era distinto y también hicieron cálculos sobre él). Para realizar los cálculos, los Mayas habían inventado un sistema de barras y puntos que representaba los números (incluido el cero) y tenían reglas con qué calcular no sólo los amaneceres y ocasos de Venus, sino también otros fenómenos celestiales como los eclipses lunares.

En aquellos días, sólo unos cuantos sacerdotes Mayas podían realizar semejantes cálculos tan elaborados. Supongamos que preguntásemos a uno de ellos cómo dar el primer paso en el proceso de predicción de la siguiente aparición de Venus como estrella matutina —restando dos números—. Y supongamos que, al contrario de lo que ocurre en la actualidad, no hubiésemos ido a la escuela y no supiésemos restar. ¿Cómo nos explicaría el

sacerdote lo que es una substracción?

Podría enseñarnos los números representados por las barras y puntos y las reglas de «substracción», o podría decirnos lo que estaba haciendo realmente: «Supongamos, que queremos restar 236 a 584. Primero, contemos 584 judías y pongámoslas en un puchero. Luego quitemos 236 judías y dejémoslas a un lado. Finalmente, contemos las judías del puchero. Ese número es el resultado de substraer 236 a 584».

Podrían decir, «¡Por Quetzalcoatl! ¡Qué *aburrimiento* —contar judías, ponerlas, sacarlas—, qué trabajo!». A lo que el sacerdote respondería, «Esta es la razón por la que tenemos reglas para las barras y puntos. Las reglas son intrincadas, pero son mucho más eficaces como medio de obtener la respuesta que contar judías. Lo importante es que no hay diferencia en lo que al *resultado* se refiere: podemos predecir la aparición de Venus contando judías (lo que es lento pero sencillo de entender) o utilizando las intrincadas reglas (que es mucho más rápido pero que requiere años en la escuela para aprenderlas)».

Comprender *cómo* funciona la substracción —en tanto que no tenga Vd. que hacerlo realmente— no es tan difícil. Esta es mi postura: Voy a explicarles lo que los físicos *hacen* cuando predicen cómo se comporta la naturaleza, pero no voy a enseñarles ningún truco para que puedan hacerlo de manera *eficaz*. Descubrirán que para realizar cualquier predicción razonable con este nuevo esquema de la electrodinámica cuántica, tendrán que hacer un montón de flechas en un papel. Nos lleva siete años —cuatro como estudiante universitario y tres como licenciado— el entrenar a nuestros estudiantes de física para que lo realicen con habilidad y eficiencia. Y he aquí cómo vamos a saltamos siete años de educación en física: explicándoles la electrodinámica cuántica en términos de lo que *realmente estamos haciendo*, ¡espero que sean capaces de comprenderlo mejor que algunos de nuestros estudiantes!

Avanzando un escalón más con el ejemplo de los Mayas, podríamos preguntar al sacerdote por qué cinco ciclos de Venus son casi equivalentes a 2920 días, u ocho años. Habría todo tipo de teorías acerca del *porqué*, tal como «20 es un número importante en nuestro sistema de contar, y si divide 2920 entre 20 obtiene 146 que es una unidad más del número que se puede representar por la suma de dos cuadrados en dos maneras distintas» y cosas por el estilo. Pero esta teoría no tendría nada que ver con Venus realmente. En los tiempos modernos hemos encontrado que las teorías de este tipo no son útiles. De modo que, de nuevo, no vamos a tratar el *por qué* la naturaleza se

comporta de la forma peculiar en que lo hace, no existen buenas teorías que lo expliquen.

Lo que he hecho hasta ahora es ponerles en disposición adecuada para que me escuchen. De otra manera, no hubiese tenido oportunidad. De modo que aquí estamos ¡dispuestos a lanzarnos!

Comencemos con la luz. Cuando Newton empezó a considerar la luz, lo primero que notó es que la luz blanca es una mezcla de colores. Descompuso la luz blanca mediante un prisma, en varios colores, pero cuando hizo pasar luz de color —rojo, por ejemplo— a través de otro prisma, encontró que no podía descomponerla más. Así encontró Newton que la luz blanca es una mezcla de diferentes colores, cada uno de los cuales es puro en el sentido de que no se puede descomponer más.

(De hecho, un color particular de luz puede desdoblarse una vez más, en sentido diferente, de acuerdo con la denominada «polarización». Este aspecto de la luz no es vital para entender el carácter de la electrodinámica cuántica, por lo que en beneficio de la sencillez lo dejaré a un lado —a expensas de no tener una descripción absolutamente completa de la teoría—. Esta ligera modificación no impedirá, en ningún sentido, un entendimiento real de lo que les hablaré. Pero, debo de tener cuidado y mencionarles todo lo que dejo a un lado).

Cuando digo «luz» en estas conferencias, no me refiero solamente a la luz que vemos, del rojo al azul. Ocurre que la luz visible es sólo una parte de una larga escala que es análoga a la escala musical, en la que hay notas más altas y más bajas de las que se pueden oír. La escala de luz puede describirse mediante números —denominados frecuencias— y cuando los números se hacen más grandes, la luz va del rojo al azul, al violeta y al ultravioleta. No podemos ver la luz ultravioleta, pero puede afectar las placas fotográficas. Es luz —solo que su número es diferente—. (No debemos ser tan provincianos: lo que podemos detectar directamente con nuestro propio instrumento, el ojo, ¡no es lo único que existe en el mundo!). Si continuamos cambiando el número, llegamos a los rayos-X, rayos gamma y así sucesivamente. Si cambiamos el número en la otra dirección, vamos de las ondas azules a las rojas, a las infrarrojas (calor) y luego a las ondas de televisión y de radio. Para mí todo esto es «luz». Voy a utilizar la luz roja para la mayoría de mis ejemplos, pero la teoría de la electrodinámica cuántica se extiende a todo el rango que he descrito, y es la teoría que está detrás de todos estos diversos fenómenos. Newton pensó que la luz estaba hecha de partículas —a las que

llamó «corpúsculos»— y tenía razón (pero el razonamiento que utilizó para llegar a tal conclusión era erróneo). Sabemos que la luz está formada de partículas porque podemos tomar un instrumento muy sensible que hace clicks cuando la luz incide sobre él, y si la luz se hace más tenue, los clicks se mantienen igual de sonoros sólo que hay menos. Luego la luz es algo como las gotas de lluvia —cada pequeño pedacito de luz se denomina fotón— y si la luz es de un único color, todas las «gotas de lluvia» tienen el mismo tamaño.

El ojo humano es un instrumento muy bueno: sólo requiere de cinco a seis fotones para activar una célula nerviosa y llevar un mensaje al cerebro. Si hubiésemos evolucionado un poquito más de forma que pudiésemos ver con una sensibilidad diez veces mayor, no tendríamos que tener esta discusión — todos hubiésemos visto una luz muy tenue de un sólo color como una serie de destellos intermitentes de la misma intensidad.

Pueden preguntarse cómo es posible detectar un único fotón. Al instrumento que puede hacerlo se le denomina fotomultiplicador, y describiré brevemente cómo funciona: Cuando un fotón incide sobre la placa metálica A en la parte inferior del dibujo (ver Figura 1) hace que se libere un electrón de uno de los átomos de la placa. El electrón liberado se ve fuertemente atraído por la placa B (que está cargada positivamente), e incide sobre ella con fuerza suficiente como para liberar tres o cuatro electrones. Cada electrón arrancado de la placa B se ve atraído por la placa C (que también está cargada) y en su choque con esta placa libera más electrones a su vez. Este proceso se repite diez o doce veces, hasta que miles de millones de electrones, suficientes como para originar una corriente eléctrica apreciable, inciden sobre la última placa, L. Esta corriente se puede aumentar mediante un amplificador regular y enviarse a un altavoz que produce clicks audibles. Cada vez que un fotón de un color determinado incide sobre el fotomultiplicador, se escucha un click de volumen uniforme.

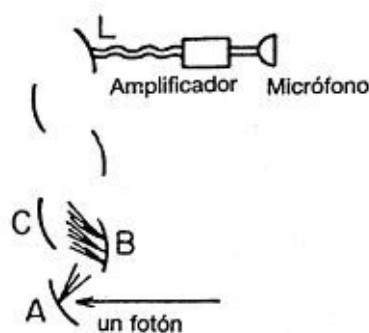


FIGURA 1. Un fotomultiplicador puede detectar un único fotón. Cuando un fotón incide sobre la

placa A, libera un electrón que es atraído por la placa B, cargada positivamente, arrancando más electrones a su vez. Este proceso continúa hasta que miles de millones de electrones inciden sobre la última placa, L, y producen una corriente eléctrica, que es aumentada por un amplificador. Si se conecta un altavoz al amplificador, se oyen clicks de volumen uniforme cada vez que un fotón de un color determinado incide sobre la placa A.

Si colocamos un gran número de fotomultiplicadores por los alrededores y dejamos que una luz muy tenue brille en varias direcciones, la luz va a uno u otro de los multiplicadores y hace un click de intensidad total. Es todo o nada: si un fotomultiplicador se dispara en un momento determinado, ningún otro se dispara simultáneamente (excepto en el caso raro de que dos fotones abandonen la fuente luminosa simultáneamente). No hay desdoblamiento de luz en «medias partículas» que vayan a lugares diferentes.

Quiero resaltar que la luz llega de esta forma —partículas—. Es muy importante saber que la luz se comporta como partículas, especialmente para aquellos de Vds. que han ido a la escuela, en donde probablemente les dijeron algo acerca de la luz comportándose como ondas. Les diré la forma en que *realmente* se comporta —como partículas.

Podrían decir que es el fotomultiplicador el que detecta la luz como partículas, pero no, cada instrumento diseñado con la sensibilidad suficiente para detectar luz débil siempre ha terminado descubriendo lo mismo: que la luz está formada por partículas.

Voy a suponer que están familiarizados con las propiedades de la luz en las circunstancias ordinarias —cosas como que la luz se propaga en línea recta, que se desvía cuando incide sobre el agua; que cuando se refleja en una superficie especular, el ángulo con que incide en la superficie es igual al ángulo con que la abandona; que la luz puede descomponerse en colores; que se pueden observar colores muy bonitos en un charco cuando éste tiene un poquito de aceite; que una lente focaliza la luz, y así sucesivamente—. Voy a utilizar estos fenómenos que les son familiares a fin de ilustrar el comportamiento verdaderamente extraño de la luz; voy a explicarles estos fenómenos familiares en términos de la teoría de la electrodinámica cuántica. Les he hablado sobre el fotomultiplicador para ilustrarles un fenómeno esencial con el que podrían no estar familiarizados —el que la luz está constituida por partículas— ¡pero con el que espero que ahora también se hayan familiarizado!

Bien, creo que todos conocen el fenómeno por el que la luz se ve parcialmente reflejada por algunas superficies tales como la del agua. Hay muchas pinturas románticas de la luz de la luna reflejándose en un lago (¡y

son muchas las veces en las que se han metido en problemas *a causa de* la luz de la luna reflejándose en un lago!). Cuando se mira desde arriba hacia el agua se puede ver lo que está debajo de la superficie (especialmente durante el día), pero también se observa el reflejo de la superficie. El cristal proporciona otro ejemplo: si tiene una lámpara encendida en la habitación y mira hacia afuera a través de una ventana durante el día, puede ver las cosas del exterior a través del cristal junto con un tenue reflejo de la lámpara en la habitación. Luego la luz es reflejada parcialmente por la superficie del cristal.

Antes de continuar, quiero que se den cuenta de la simplificación que voy a realizar y que corregiré posteriormente: Cuando hablo de la reflexión parcial de la luz por el cristal, voy a pretender que la luz se refleja sólo por la *superficie* de cristal. En realidad, un trozo de cristal es un monstruo terrible de complejidad —enormes cantidades de electrones están en movimiento.

Cuando llega un fotón, interacciona con los electrones *del* cristal, no sólo con los de la superficie. El fotón y los electrones realizan una especie de baile, cuyo resultado final es equivalente al fotón incidiendo sólo sobre la superficie. Así que déjenme realizar esta simplificación por un rato. Posteriormente, les enseñaré lo que ocurre en realidad dentro del cristal de forma que puedan comprender por qué el resultado es el mismo.

Ahora me gustaría describirles un experimento y explicarles sus sorprendentes resultados. En este experimento una fuente luminosa emite unos cuantos fotones del mismo color —es decir, luz roja— en la dirección de un bloque de cristal (ver Fig. 2). En A, por encima del cristal, se ha colocado un fotomultiplicador para captar cualquier fotón que sea reflejado por la superficie superior. Para medir cuántos fotones atraviesan esta superficie, se ha colocado otro fotomultiplicador en B, dentro del cristal. No importan las dificultades obvias de colocar un fotomultiplicador dentro de un bloque de cristal; ¿cuáles son los resultados de este experimento?

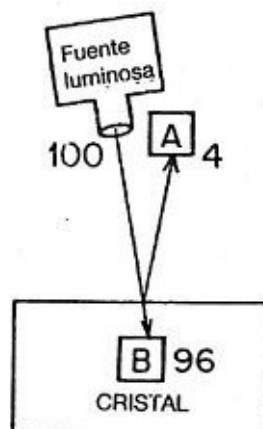


FIGURA 2. Un experimento para medir la reflexión parcial de la luz por una única superficie de cristal. Por cada 100 fotones que abandonan la fuente de luz, 4 son reflejados por la superficie frontal y terminan en el fotomultiplicador en A, mientras que los otros 96 son transmitidos por la superficie frontal y acaban en el fotomultiplicador en B.

De cada 100 fotones que van directos hacia el cristal, formando un ángulo de 90° con él, una media de 4 llegan a A y 96 a B. Luego «reflexión parcial» en este caso significa que el 4 por 100 de los fotones son reflejados por la superficie frontal del cristal, mientras que el 96 por 100 restante es transmitido. Ya estamos ante una gran dificultad: ¿cómo puede ser la luz reflejada *parcialmente*? Cada fotón acaba en A o en B —¿cómo «decide» el fotón si debe ir a A o a B?— (Se ríe la audiencia). Puede sonar a chiste, pero no podemos reír, ¡tenemos que explicar esto en términos de una teoría! La reflexión parcial es ya un gran misterio, y fue un problema muy difícil para Newton.

Existen varias teorías posibles que pueden explicar la reflexión parcial de la luz por el cristal. Una es que el 96 por 100 de la superficie del cristal está formada por «agujeros» que dejan pasar la luz, mientras que el otro 4% de la superficie está cubierta de pequeñas «manchas» de material reflectante (ver Fig. 3). Newton se percató de que esta explicación no era posible^[1]. En un momento encontraremos un extraño rasgo de la reflexión parcial que les volverá locos si se adhieren a la teoría de «agujeros y manchas» —¡o a cualquier otra teoría razonable!

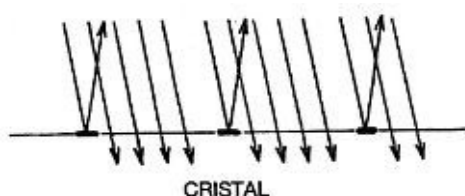


FIGURA 3. Una teoría para explicar la reflexión parcial por una única superficie implica una superficie hecha principalmente de «huecos» que dejan pasar la luz, junto con unas pocas «manchas» que la reflejan.

Otra teoría posible surge de la consideración de que los fotones tengan algún tipo de mecanismo interno —«ruedas» y «engranajes» interiores que giren de alguna manera— de modo que cuando se «apunta» el fotón adecuadamente, pasa a través del cristal, y cuando no se apunta correctamente, se refleja. Podemos probar esta teoría tratando de filtrar los fotones que no están dirigidos en la dirección correcta colocando unas cuantas láminas más de cristal entre la fuente y la primera capa del cristal. Después de atravesar los filtros, los fotones que llegasen al cristal deberían estar *todos* dirigidos correctamente y ninguno debería ser reflejado. El problema de esta

teoría es que no concuerda con el experimento: incluso después de atravesar muchas capas de cristal, el 4% de los fotones que alcanzan una superficie dada son reflejados.

Por mucho que intentemos inventar una teoría razonable que pueda explicar cómo un fotón «decide» si atraviesa un cristal o retrocede, es imposible predecir en qué dirección irá un fotón. Los filósofos han dicho que si las mismas circunstancias no producen siempre los mismos resultados, las predicciones resultan imposibles y la ciencia colapsaría.

He aquí una circunstancia —fotones idénticos llegando siempre en la misma dirección al mismo trozo de cristal— que produce resultados diferentes. No podemos predecir si un fotón dado llegará a A o a B. Todo lo que podemos predecir es que de cada 100 fotones que llegan, una media de 4 serán reflejados por la superficie frontal. ¿Significa esto que la física, ciencia de gran exactitud, se ha reducido a calcular sólo la *probabilidad* de un suceso, y no de predecir de manera exacta lo que va a ocurrir? Sí. Es una retirada, pero así es como es: la naturaleza sólo nos permite calcular probabilidades. Y, sin embargo, la ciencia aún no ha colapsado.

Mientras que la reflexión parcial por una sola superficie es un profundo misterio y un problema difícil, la reflexión parcial por dos o más superficies es un absoluto quebradero de cabeza. Déjenme decirles el porqué. Realizaremos un segundo experimento, en el que mediremos la reflexión parcial de la luz por dos superficies. Reemplazamos el bloque de cristal por una lámina muy delgada —con sus dos superficies exactamente paralelas entre sí— y colocamos el fotomultiplicador debajo de la lámina de cristal, en línea con la fuente de luz. Esta vez, los fotones pueden ser reflejados por la superficie frontal o por la superficie posterior —y finalizar en A—, los demás acabarán en B (ver Fig. 4). Podríamos esperar que la superficie frontal reflejase al 4% de la luz y la superficie posterior el 4% del 96% restante, haciendo un total del 8%. Así, deberíamos encontrar que de cada 100 fotones que salen de la fuente de luz, alrededor de 8 llegasen a A.

Lo que realmente ocurre, bajo estas condiciones experimentales cuidadosamente controladas, es que el número de fotones que llega a A es raramente 8 de cada 100. Con algunas láminas de cristal, se obtiene de forma consistente una lectura de 15 o 16 fotones —¡el doble del valor esperado!—. Con otras láminas de cristal, obtenemos consistentemente sólo 1 o 2 fotones. Otras láminas de cristal dan una reflexión parcial del 10%; ¡algunas eliminan totalmente la reflexión parcial! ¿Qué puede explicar estos locos resultados?

Después de probar la calidad y uniformidad de las distintas laminas de cristal, descubrimos que sólo difieren ligeramente en el espesor.

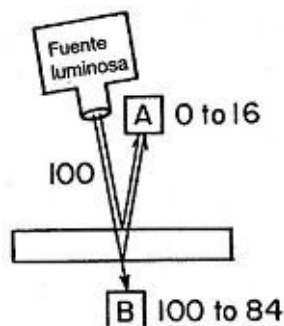


FIGURA 4. Un experimento para medir la reflexión parcial de la luz por dos superficies de cristal. Los fotones pueden alcanzar el fotomultiplicador en A al reflejarse bien en la superficie frontal o en la superficie posterior de una lámina de cristal; alternatively, pueden atravesar ambas superficies y acabar alcanzando el fotomultiplicador en B. Dependiendo del espesor del cristal, de 0 a 16 fotones de cada 100 llegan al fotomultiplicador en A. Estos resultados plantean dificultades a cualquier teoría razonable, incluyendo la de la Figura 3. Parece como si la reflexión parcial pudiese ser «reducida» o «ampliada» por la presencia de una superficie adicional.

Para comprobar la idea de que la cantidad de luz reflejada por las dos superficies depende del espesor del cristal, realicemos una serie de experimentos: comenzando con la lámina de cristal más fina posible, contaremos cuántos fotones inciden en el fotomultiplicador en A cada vez que 100 fotones abandonan la fuente de luz. Luego reemplazaremos la lámina de cristal por otra ligeramente más gruesa y contaremos de nuevo. Después de repetir el proceso una docena de veces ¿cuáles son los resultados?

Con la lámina de cristal más delgada posible, encontramos que el número de fotones que llega a A es casi siempre cero —a veces es 1—. Cuando reemplazamos esta lámina por otra ligeramente más gruesa, encontramos que la cantidad de luz reflejada es más alta —más cercana al esperado 8%—. Después de unos cuantos cambios más, el número de fotones que llega a A aumenta sobrepasando la marca del 8%. Si continuamos substituyendo láminas de cristal cada vez más gruesas —estamos ahora en los alrededores de 5 millonésimas de pulgada—, la cantidad de luz reflejada por las dos superficies alcanza un máximo del 16% y luego decrece pasando por el 8% hasta el valor cero —si la lámina de cristal es del espesor adecuado, no existe reflexión—. (¡Consiga esto mediante «manchas»!).

Con láminas de cristal gradualmente más gruesas, la reflexión parcial aumenta de nuevo al 16% y retorna luego a cero —un ciclo que se repite una y otra vez (ver Fig. 5)—. Newton descubrió estas oscilaciones y realizó un experimento que se podía interpretar de manera correcta ¡sólo si las

oscilaciones continuasen durante al menos 34 000 ciclos! Hoy, con los láseres (que producen una luz monocromática muy pura), se puede ver que este ciclo continúa después de más de 100 000 000 de repeticiones —lo que corresponde a un cristal de más de 50 metros de espesor—. (No vemos este fenómeno todos los días porque la fuente luminosa no es normalmente monocromática).

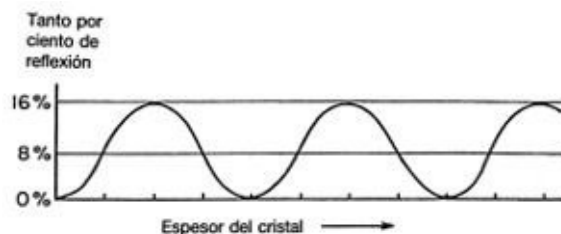


FIGURA 5. Los resultados de un experimento que mide cuidadosamente la relación entre el espesor de una lámina de cristal y la reflexión parcial, demuestran la existencia de un fenómeno llamado «interferencia». Al aumentar el espesor del cristal, la reflexión parcial experimenta un ciclo de repetición desde 0 al 16%, sin síntomas de extinción.

Resulta, por consiguiente, que nuestra predicción del 8% es correcta como media (puesto que la cantidad real varía de manera regular desde cero al 16%) pero es exactamente correcta solamente dos veces por ciclo —como un reloj parado (que tiene la hora correcta dos veces al día)—. ¿Cómo podemos explicar este extraño rasgo de la reflexión parcial que depende del espesor del cristal? ¿Cómo puede reflejar la superficie frontal un 4% de luz (según confirma nuestro primer experimento) cuando, colocando una segunda superficie debajo, a la distancia adecuada, podemos de alguna manera «apagar» la reflexión? Y colocando esta segunda superficie a una distancia ligeramente distinta, ¿podemos «amplificar» la reflexión hasta un 16%! ¿Podría ser que la superficie posterior ejerciese algún tipo de influencia o algún efecto en la habilidad de la superficie frontal para reflejar la luz? ¿Qué ocurriría si colocásemos una *tercera* superficie?

Con una tercera superficie, o cualquier número superior de superficies, el valor de la reflexión parcial cambia de nuevo. Nos encontramos preguntándonos, después de perseguir superficie tras superficie con esta teoría, si habremos alcanzado finalmente la última superficie. ¿Tiene que hacer esto un fotón para «decidir» si se va a reflejar en la superficie frontal? Newton llevó a cabo discusiones ingeniosas referentes a este problema^[2], pero al final se dio cuenta de que no había desarrollado una teoría satisfactoria.

Durante muchos años después de Newton, la reflexión parcial por dos

superficies se explicaba felizmente mediante una teoría de ondas^[3], pero cuando los experimentos se realizaron con luz muy débil incidiendo en fotomultiplicadores, la teoría ondulatoria colapsó: según se iba haciendo la luz más tenue, los fotomultiplicadores seguían haciendo clicks de igual intensidad, sólo que cada vez en menor número. La luz se comportaba como partículas.

La situación actual es que no tenemos un buen modelo para explicar la reflexión parcial por dos superficies; sólo calculamos la probabilidad de que un fotomultiplicador determinado sea alcanzado por un fotón reflejado por una lámina de cristal. He escogido este cálculo como nuestro primer ejemplo del método que ha proporcionado la teoría de la electrodinámica cuántica. Voy a mostrarles «cómo contamos judías» —lo que los físicos hacen para obtener la respuesta correcta—. No voy a explicar cómo los fotones «deciden» en realidad si retroceden o continúan hacia adelante, esto no se conoce. (Probablemente el problema no tiene sentido). Sólo les voy a mostrar cómo calcular la *probabilidad* correcta de que la luz sea reflejada por un cristal de espesor dado, ¡porque esto es lo único que los físicos saben hacer! Lo que hacemos para tener la respuesta a este problema es análogo a lo que tenemos que hacer para tener la respuesta a *cualquier otro* problema explicado por la electrodinámica cuántica.

Prepárense para enfrentarse con ello, no porque sea difícil de entender, sino porque es absolutamente ridículo: Todo lo que hacemos es dibujar pequeñas flechas en una hoja de papel —¡esto es todo!

Pero ¿qué relación tiene una flecha con la probabilidad de que ocurra un suceso en particular? De acuerdo con las reglas de «cómo contamos judías», la probabilidad de un suceso es igual al cuadrado de la longitud de la flecha. Por ejemplo, en nuestro primer experimento (cuando estábamos midiendo la reflexión parcial sólo de la superficie frontal), la probabilidad de que un fotón llegase al fotomultiplicador situado en A era del 4%. Esto se corresponde con una flecha cuya longitud es 0,2 porque el cuadrado de 0,2 es 0,04 (ver Fig. 6).

En nuestro segundo experimento (cuando estábamos reemplazando láminas delgadas de cristal por otras más gruesas), los fotones rebotados por la superficie frontal y la posterior llegaban a A. ¿Cómo dibujar una flecha que representa esta situación? La longitud de la flecha debe variar de cero a 0,4 para representar las probabilidades del 0 al 16%, dependiendo del espesor del cristal (ver Fig. 7).

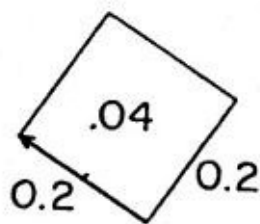


FIGURA 6. Los extraños rasgos de la reflexión parcial por dos superficies han forzado a los físicos a renunciar a efectuar predicciones absolutas para realizar meros cálculos de probabilidad de su suceso. La electrodinámica cuántica proporciona un método para hacerlo —dibujando pequeñas flechas en una hoja de papel—. La probabilidad de un suceso viene representada por el área del cuadrado generado por una flecha. Por ejemplo, una flecha representando una probabilidad del 0,04 (4%) tiene una longitud de 0,2.

Empezaremos por considerar los distintos *caminos* que el fotón puede llevar desde la fuente hasta el fotomultiplicador en A. Puesto que estoy haciendo la simplificación de que la luz rebota, bien en la superficie frontal o en la posterior, existen dos caminos posibles para que un fotón pueda llegar a A. Lo que hacemos en este caso es dibujar *dos* flechas —una por cada forma en que puede ocurrir— y combinarlas en una «flecha final» cuyo cuadrado representa la probabilidad del suceso. Si existieran tres caminos distintos en los que pudiese haber ocurrido el suceso, dibujaríamos tres flechas separadas antes de combinarlas.

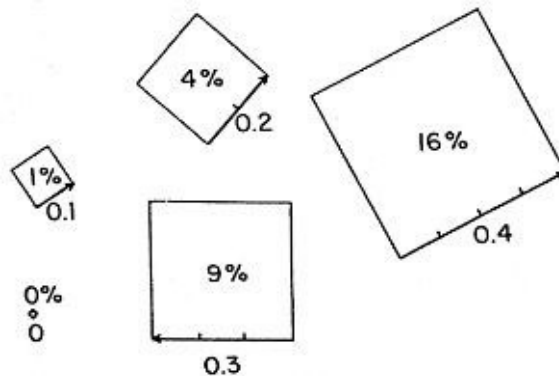


FIGURA 7. Flechas representando probabilidades del 0% al 16% tienen longitudes de 0 a 0,4.

Ahora, déjenme mostrarles cómo combinamos las flechas. Digamos que queremos combinar la flecha x con la flecha y (ver Fig. 8).

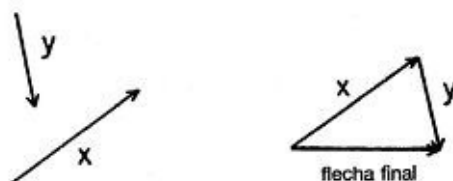


FIGURA 8. Se han dibujado las flechas que representan cada camino posible por el que un suceso puede tener lugar, y luego se han combinado («sumado») de la siguiente manera: Se une la punta de una flecha al extremo posterior de otra —sin cambiar las direcciones de ninguna— y se dibuja una «flecha final» desde la cola de la primera flecha a la cabeza de la última.

Todo lo que hay que hacer es colocar la cabeza de x junto a la cola de y (sin cambiar la dirección de ninguna), y dibujar la flecha final desde la cola de x a la cabeza de y . Esto es todo lo que hay que hacer. Podemos combinar cualquier número de flechas de esta manera (técnicamente, se llama «sumar flechas»). Cada flecha dice cuánto, y en qué dirección, hay que moverse en el baile. La flecha final dice el *único* movimiento a realizar para terminar en el mismo punto (ver Fig. 9).



FIGURA 9. Cualquier número de flechas puede sumarse en la forma descrita en la Figura 8.

Bien, ¿cuáles son las reglas específicas que determinan la longitud y dirección de cada flecha que combinamos a fin de obtener la flecha final? En este caso particular, combinaremos dos flechas —una representando la reflexión por la superficie *frontal* del cristal, y la otra representando la reflexión por la superficie *posterior*.

Consideremos primero la longitud. Como vimos en el primer experimento (donde pusimos el fotomultiplicador dentro del cristal), la superficie frontal refleja alrededor del 4% de los fotones que le llegan. Esto significa que la flecha de «reflexión frontal» tiene una longitud de 0,2. La superficie posterior del cristal también refleja el 4%, luego la longitud de la flecha de «reflexión posterior» también es 0,2.

Para determinar la dirección de cada flecha, imaginemos que tenemos un cronógrafo que puede seguir a un fotón en su movimiento. Este cronógrafo imaginario tiene una única manecilla que gira muy rápidamente. Cuando un fotón sale de la fuente, ponemos en marcha el cronógrafo. Mientras que el fotón está en movimiento la manecilla del cronógrafo gira (alrededor de 36 000 veces por pulgada para la luz roja); cuando el fotón llega al fotomultiplicador detenemos el reloj. La manecilla señala en una cierta dirección. Esta es la dirección en la que dibujaremos la flecha.

Necesitamos una regla más para poder calcular correctamente la respuesta: Cuando estemos considerando el camino de un fotón reflejado por la superficie *frontal* del cristal, invertiremos el sentido de la flecha. En otras palabras, mientras que dibujamos la flecha de la reflexión *posterior* señalando

en el mismo sentido que la manecilla del cronógrafo, dibujaremos la flecha de la reflexión *frontal* en sentido *opuesto* al señalado.

Ahora, dibujemos las flechas para el caso de que la luz se refleje en una lámina de cristal extremadamente delgada. Para dibujar la flecha de reflexión frontal, imaginamos un fotón abandonando la fuente luminosa (la manecilla del cronógrafo empieza a girar), reflejándose en la superficie frontal, y llegando a A (la manecilla del cronógrafo se detiene). Dibujaremos una pequeña flecha de longitud 0,2 en el sentido opuesto al de la manecilla del cronógrafo (ver Fig. 10).

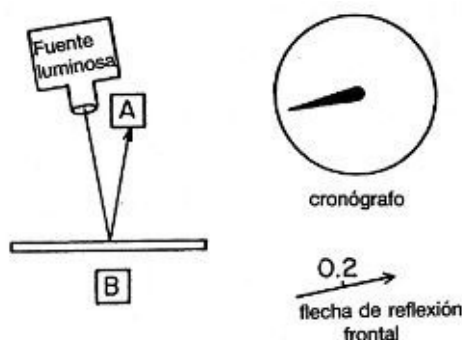


FIGURA 10. En un experimento que mide la reflexión por dos superficies, podemos decir que un único fotón puede llegar a A por dos caminos —por vía de la superficie frontal o de la posterior—. Para cada camino se dibuja una flecha de longitud 0,2, con su dirección determinada por la manecilla de un «cronógrafo» que cronometra al fotón cuando se mueve. La flecha de «reflexión frontal» se dibuja en el sentido opuesto al que señala la manecilla del cronógrafo cuando se detiene.

Para dibujar la flecha de reflexión posterior, imaginamos un fotón saliendo de la fuente de luz (la manecilla del cronógrafo empieza a girar), atravesando la superficie frontal y reflejándose en la superficie posterior, y llegando a A (la manecilla del cronógrafo se detiene). Ésta vez, la manecilla está señalando casi en la misma dirección, porque el fotón reflejado en la superficie posterior del cristal invierte un tiempo ligeramente superior en llegar a A —atravesando dos veces la extremadamente delgada lámina de cristal—. Ahora dibujamos una pequeña flecha de longitud 0,2 en la misma dirección y sentido que señala la manecilla del cronógrafo (ver Fig. 11).

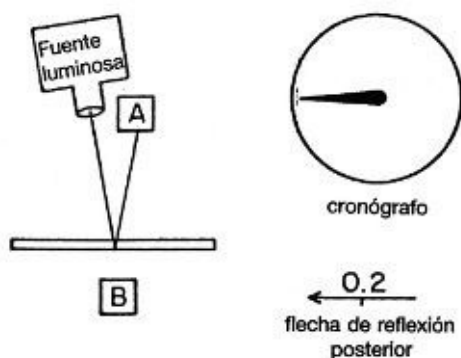


FIGURA 11. Un fotón rebotando en la superficie posterior de una lámina de cristal delgada tarda un poquito más en llegar a A. En consecuencia, la manecilla del cronógrafo señala al final una dirección ligeramente distinta de la que señaló cuando siguió al fotón de la superficie frontal. La flecha de la «reflexión posterior» se dibuja en el mismo sentido que tiene la manecilla al detenerse.

Combinemos ahora las dos flechas. Puesto que las dos tienen la misma longitud pero apuntan en sentidos casi opuestos, la flecha final tiene una longitud cercana a cero, y su cuadrado es más próximo a cero aún. Por tanto, la probabilidad de que la luz sea reflejada por una lámina de cristal infinitamente delgada es esencialmente nula (ver Fig. 12).

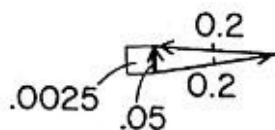


FIGURA 12. La flecha final, cuyo cuadrado representa la probabilidad de reflexión por una lámina de cristal extremadamente delgada, se dibuja sumando la flecha de reflexión frontal con la flecha de reflexión posterior. El resultado es casi nulo.

Cuando reemplazamos la lámina de cristal más delgada posible por otra ligeramente más gruesa, el fotón reflejado por la superficie posterior tarda un poquito más que en el primer ejemplo en alcanzar A; la manecilla del cronógrafo, en consecuencia, gira un poquito más antes de detenerse, y la flecha de reflexión posterior finaliza con un ángulo ligeramente mayor con respecto a la flecha de reflexión frontal. La flecha final es un poquito más larga, y su cuadrado también (ver Fig. 13).

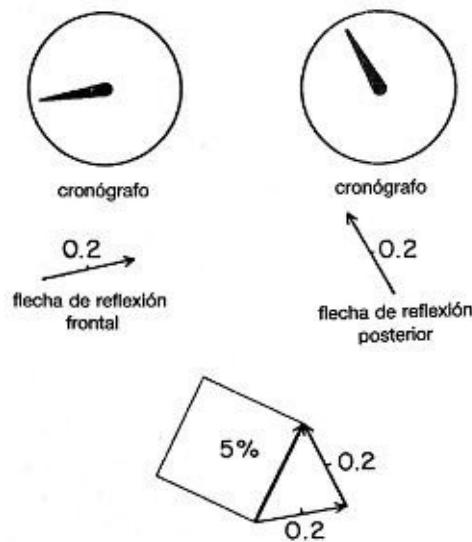


FIGURA 13. La flecha final para una lámina de cristal ligeramente más gruesa es un poco más larga, debido al mayor ángulo formado por las flechas de reflexión frontal y posterior. Esto se debe a que, comparado con el ejemplo anterior, el fotón rebotado por la superficie posterior tarda un poco más en alcanzar A.

Como otro ejemplo, consideremos el caso en que el cristal es lo suficientemente grueso como para que la manecilla del cronógrafo gire media vuelta más al cronometrar al fotón reflejado por la superficie posterior. Esta vez, la flecha de reflexión posterior acaba señalando exactamente en la misma dirección y sentido que la flecha de reflexión frontal. Cuando combinamos las dos flechas, obtenemos una flecha final de longitud 0,4, cuyo cuadrado es 0,16 representando una probabilidad del 16% (ver Fig. 14).

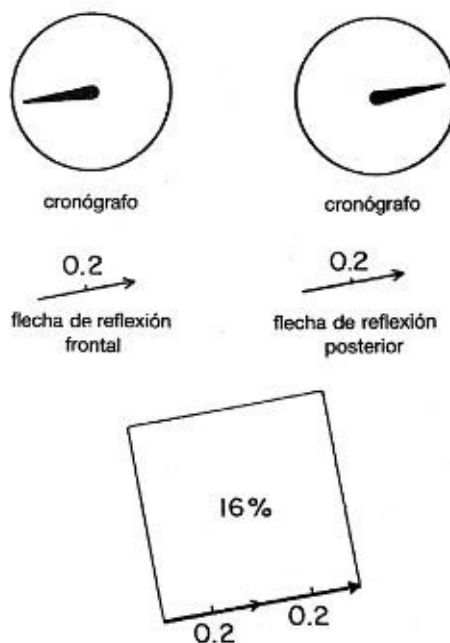


FIGURA 14. Cuando la lámina de cristal es lo suficientemente gruesa como para que la manecilla del cronógrafo dé media vuelta más, las flechas de reflexión frontal y posterior señalan de longitud 0,4, que representa una probabilidad del 16%.

Si aumentamos el espesor del cristal justo para que la manecilla del cronógrafo recorra una vuelta *entera* más, mientras cronometra el camino desde la superficie posterior, nuestras dos flechas señalarían de nuevo sentidos opuestos y la flecha final sería cero (ver Fig. 15). Esta situación ocurre una y otra vez, siempre que el cristal tenga el espesor suficiente para que la manecilla del cronógrafo que sigue la reflexión por la superficie posterior recorra una vuelta entera más.

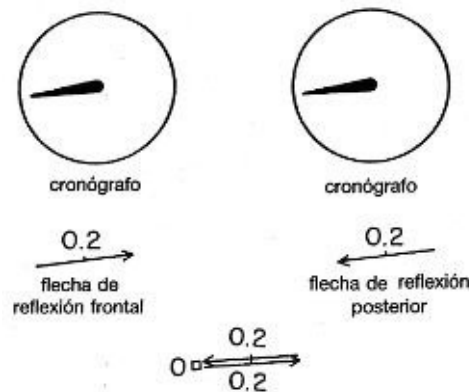


FIGURA 15. Cuando la lámina de cristal tiene el espesor adecuado para que la manecilla del cronógrafo que sigue al fotón de reflexión posterior dé una vuelta entera más, la flecha final es de nuevo cero, y no existe reflexión.

Si el espesor, del cristal es tal que la manecilla del cronógrafo que sigue a la reflexión por la superficie posterior realiza $1/4$ o $3/4$ de vuelta extra, las dos flechas finalizan formando ángulo recto. La flecha final es, en este caso, la hipotenusa de un triángulo rectángulo y, de acuerdo con Pitágoras, el cuadrado de la hipotenusa es igual a la suma de los cuadrados de los otros dos lados. Aquí tenemos el valor correcto «dos veces al día» — $4\% + 4\%$ hace 8% — (ver Fig. 16).

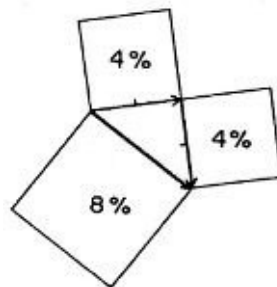


FIGURA 16. Cuando las flechas de reflexión frontal y posterior forman un ángulo recto, la flecha final es la hipotenusa de un triángulo rectángulo. Por tanto, su cuadrado es la suma de los otros dos cuadrados — 8% .

Nótese que al aumentar gradualmente el espesor del cristal, la flecha de reflexión frontal siempre señala en la misma dirección, mientras que la flecha de reflexión posterior cambia gradualmente de dirección. El cambio en la

dirección relativa de las dos flechas hace que la flecha final vaya de 0 a 0,4 en un ciclo que se repite; por tanto, el *cuadrado* de la flecha final recorre un ciclo que se repite y que va de cero al 16%, tal y como observamos en nuestros experimentos (ver Fig. 17).

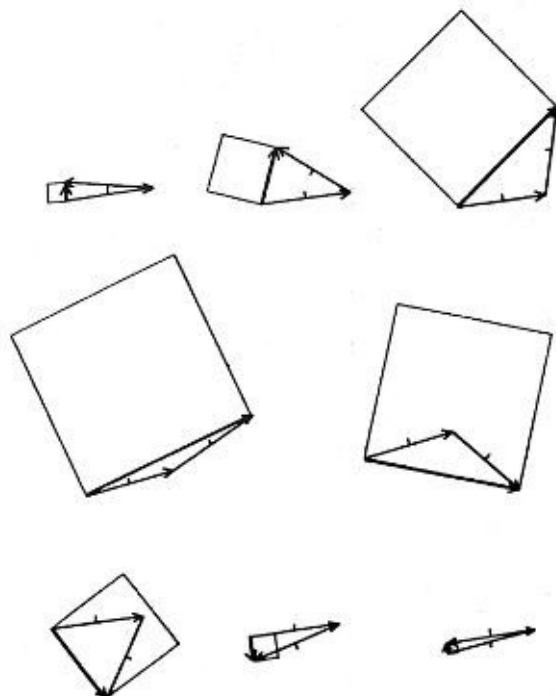


FIGURA 17. Al ser reemplazadas las láminas delgadas de cristal por otras ligeramente más gruesas, la manecilla del cronógrafo, que sigue al fotón que se refleja en la superficie posterior, gira ligeramente más y el ángulo formado por las flechas de reflexión frontal y posterior cambia. Esto hace que la longitud de la flecha final se modifique y que su cuadrado varíe de 0 al 16%, y vuelva al 0, una y otra vez.

Acabo de mostrarles cómo se puede calcular, de manera precisa, este extraño rasgo de la reflexión parcial, dibujando algunas condenadas flechitas en una hoja de papel. La palabra técnica para estas flechas es «amplitud de probabilidad» y yo me siento más dignificado cuando digo que estamos «calculando amplitudes de probabilidad para un suceso». Prefiero, no obstante, ser más honesto y decir que estamos intentando encontrar la flecha cuyo cuadrado representa la probabilidad de algo que está ocurriendo.

Antes de que termine esta primera conferencia, me gustaría hablarles de los colores que ven en las pompas de jabón. O mejor aún, si su coche pierde aceite en un charco, cuando mira al aceite amarronado de ese sucio charco con barro, verá preciosos colores en su superficie. La delgada película de aceite flotando en el charco embarrado es similar a una lámina muy delgada de cristal —refleja luz de un color, desde cero a un máximo, dependiendo de su grosor—. Si hacemos incidir luz roja sobre la película de aceite, veremos

manchas de luz roja separadas por bandas estrechas de negro (donde no existe reflexión) porque el espesor de la película de aceite no es exactamente uniforme. Si hacemos incidir luz azul veremos borrones de luz azul separados por bandas estrechas de negro. Si hacemos incidir luz roja y azul sobre el aceite, veremos zonas que tienen el espesor adecuado para reflejar intensamente sólo la luz roja, otras con el espesor adecuado para reflejar sólo la luz azul, y aún otras áreas con un espesor que permite reflejar intensamente la luz roja y azul simultáneamente (que nuestros ojos ven como violeta), mientras que otras áreas tienen el espesor adecuado para cancelar todas las reflexiones y aparecer como negras.

Para comprender esto mejor, necesitamos saber que el ciclo de cero al 16% de la reflexión parcial por dos superficies se repite más rápidamente para la luz azul que para la roja. Por consiguiente, para ciertos espesores, uno o el otro o ambos colores son intensamente reflejados, mientras que para otros espesores, la reflexión de ambos colores se ve cancelada (ver Fig. 18). Los ciclos de reflexión se repiten con frecuencia distinta porque la manecilla del cronógrafo gira más deprisa cuando sigue a un fotón azul que cuando sigue a uno rojo. De hecho, ésta es la *única* diferencia entre un fotón rojo y un fotón azul (o un fotón de cualquier otro color, incluyendo ondas de radio, rayos X y demás) —la velocidad de la manecilla del cronógrafo.

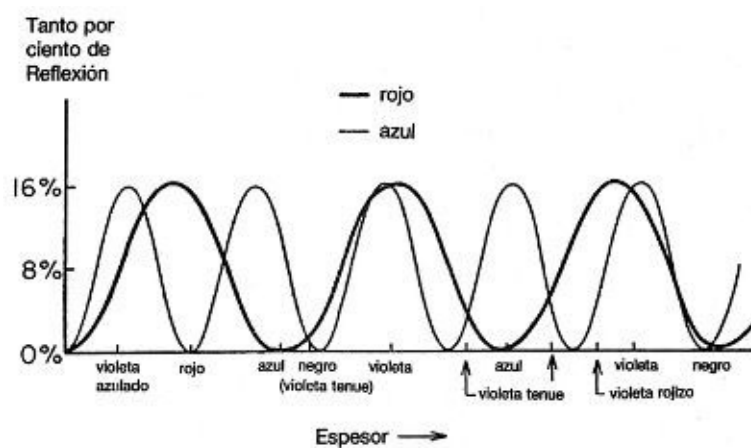


FIGURA 18. Al aumentar el espesor de la lámina, las dos superficies producen una reflexión parcial de la luz monocromática cuya probabilidad fluctúa cíclicamente desde el 0% al 16%. Puesto que la velocidad de la manecilla imaginaria del cronógrafo es distinta para los distintos colores de luz, el ciclo se repite a velocidades diferentes. Así cuando dos colores puros como rojo y azul inciden sobre la lámina, un espesor determinado reflejará sólo rojo, sólo azul, azul y rojo en proporciones diferentes (lo que produce varios matices de violeta), o ningún color (negro). Si la lámina es de espesor variable, como una gota de aceite extendiéndose en un charco fangoso, todas las combinaciones tendrán lugar. Con la luz del Sol, que contiene todos los colores, tendrá lugar todo tipo de combinaciones, lo que produce montones de tonalidades.

Cuando incide luz roja y azul sobre una lámina de aceite, aparecen dibujos

rojos, azules y violetas separados por bordes negros. Cuando la luz del sol, que contiene luz roja, amarilla, verde y azul, brilla sobre un charco de barro con aceite en su superficie, las zonas que reflejan intensamente cada uno de estos colores se superponen y producen todo tipo de combinaciones que nuestros ojos ven como colores diferentes. Al extenderse la película de aceite y moverse sobre la superficie del agua, cambiando su espesor en varios puntos, los dibujos de color cambian constantemente (Si, por otro lado, mirasen el mismo charco por la noche, con una de esas luces de sodio de las farolas incidiendo sobre él, verían sólo bandas amarillentas separadas por otras negras —porque esas luces callejeras emiten sólo luz de un color—).

Este fenómeno de colores producido por la reflexión parcial de la luz blanca por dos superficies se denomina iridiscencia, y se puede encontrar en muchos lugares. Quizás se pregunten cuál es el origen de los colores brillantes de los colibríes y pavos reales. Ahora lo saben. Cómo se desarrollaron esos colores brillantes es también una pregunta interesante. Cuando admiramos un pavo real, debemos agradecer a las generaciones de hembras deslustradas el haber sido tan selectivas con sus machos. (El hombre se introdujo en este asunto posteriormente para hacer más eficaz el proceso de selección de los pavos reales).

En la próxima conferencia les mostraré cómo este absurdo proceso de combinar pequeñas flechas dé la respuesta correcta a esos otros fenómenos que les son familiares: la luz viaja en línea recta, se refleja en un espejo con el mismo ángulo con que llega («el ángulo de incidencia es igual al ángulo de reflexión»), las lentes focalizan la luz, y similares. Este nuevo marco les describirá todo acerca de la luz.

Capítulo 2

FOTONES: PARTÍCULAS DE LUZ

Esta es la segunda de una serie de conferencias sobre electrodinámica cuántica y como está claro que ninguno de Vds. estaba aquí la última vez (puesto que les dije que no iban a entender nada) voy a resumir brevemente la primera conferencia.

Estábamos hablando de la luz. El primer rasgo importante de la luz es que parece que son partículas: cuando luz monocromática (luz de un color) muy tenue incide sobre un detector, el detector emite clicks de la misma intensidad, aunque cada vez con menor frecuencia, según se va haciendo más tenue la luz.

El otro rasgo importante de la luz discutido en la primera conferencia es el de la reflexión parcial de la luz monocromática. Una media del 4% de los fotones que inciden sobre una *única* superficie de cristal es reflejada. Esto es ya un profundo misterio, puesto que es imposible predecir qué fotones se reflejarán y cuáles pasarán a través de la lámina. Con una *segunda* superficie, los resultados son extraños: en lugar del esperado 8% de reflexión por las dos superficies, la reflexión parcial puede ampliarse hasta un elevado 16% o desaparecer, dependiendo del espesor del cristal.

Este extraño fenómeno de la reflexión parcial por dos superficies puede explicarse para una luz intensa mediante una teoría de ondas, pero la teoría ondulatoria no puede explicar cómo el detector emite clicks de igual intensidad cuando la luz se atenúa. La electrodinámica cuántica «resuelve» esta dualidad onda-partícula estableciendo que la luz se compone de partículas (como Newton pensó inicialmente), pero el precio de este gran

avance de la ciencia es una retirada de la física a la posición de ser capaz de calcular sólo la *probabilidad* de que un fotón incida en el detector, sin ofrecer un buen modelo de cómo ocurre realmente.

En la primera conferencia describí cómo los físicos calculan la probabilidad de que un suceso particular tenga lugar. Dibujaré algunas flechas sobre una hoja de papel según unas reglas, que son las siguientes:

- GRAN PRINCIPIO: la probabilidad de un suceso es igual al cuadrado de la longitud de una flecha denominada «amplitud de probabilidad». Una flecha de longitud 0,4, por ejemplo, representa una probabilidad de 0,16 o 16%.
- REGLA GENERAL para dibujar flechas si un suceso puede ocurrir por caminos alternativos: Dibujar una flecha para cada camino, y combinarlas («sumarlas») uniendo la cabeza de una a la cola de la siguiente. Se dibuja luego una «flecha final» desde la cola de la primera flecha hasta la cabeza de la última. La flecha final es tal que su cuadrado da la probabilidad del suceso completo.

Existen también algunas reglas específicas para dibujar flechas en el caso de la reflexión parcial por el cristal (se pueden encontrar en las páginas 36 y 37). Todo lo que precede es un resumen de la primera conferencia.

Lo que me gustaría hacer ahora es mostrarles cómo este modelo del mundo, que es tan completamente distinto de cualquiera que hayan visto antes (que quizá esperen no volver a verlo nunca más), puede explicar todas las propiedades elementales que conocen de la luz: cuando la luz se refleja en un espejo, el ángulo de incidencia es igual al ángulo de reflexión; la luz se curva cuando pasa del aire al agua; la luz viaja en línea recta; la luz se puede focalizar por una lente, y así sucesivamente. La teoría también describe muchas otras propiedades de la luz con las que probablemente no estén familiarizados. De hecho, la mayor dificultad que he tenido al preparar estas conferencias ha sido el resistir la tentación de deducir todas las cosas que sobre la luz les costó tanto aprender en la escuela —tales como el comportamiento de la luz cuando pasa por un borde y arroja una sombra (llamado difracción)— pero puesto que la mayoría de Vds. no ha observado cuidadosamente estos fenómenos, no me ocuparé de ellos. Sin embargo, puedo garantizarles (de otra manera, los ejemplos que voy a mostrarles serían engañosos) que *cada* fenómeno luminoso que ha sido observado con detalle puede ser explicado por la teoría de la electrodinámica cuántica, aunque sólo

vaya a describirles los más sencillos y comunes de entre ellos.

Comenzaremos con un espejo y el problema de determinar cómo se refleja la luz en él (ver Fig. 19). En S tenemos una fuente que emite luz de un color con una intensidad muy baja (usemos de nuevo luz roja). Esta fuente emite un fotón cada vez. En P, colocamos un fotomultiplicador para detectar fotones. Coloquémoslo a la misma altura que la fuente —el dibujar flechas es más sencillo si todo es simétrico—. Queremos calcular la probabilidad de que el detector haga un click después de que la fuente haya emitido un fotón. Puesto que es posible que un fotón vaya directamente hacia el detector, coloquemos una pantalla en Q para evitar que ocurra esto.

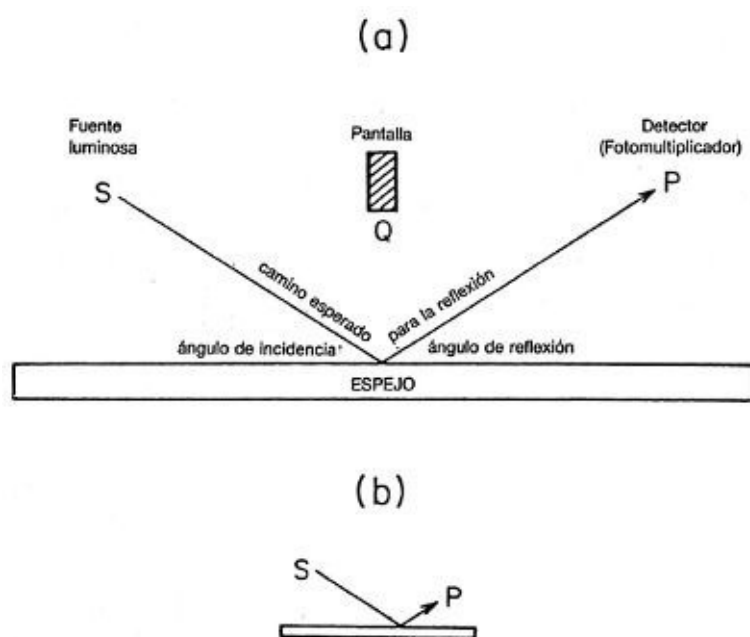


FIGURA 19. El punto de vista clásico del mundo dice que un espejo reflejará la luz allí donde el ángulo de incidencia iguale al de reflexión, incluso si la fuente y el detector están en niveles diferentes, como en (b).

Bien, esperamos que toda la luz que alcance el detector se refleje en el centro del espejo, porque es el sitio en donde el ángulo de incidencia es igual al ángulo de reflexión. Y parece bastante obvio que las partes del espejo cercanas a ambos extremos tengan tanto que ver con la reflexión como con el precio del queso ¿verdad?

Aunque puedan *pensar* que las partes del espejo cercanas a los extremos no tienen nada que ver con la reflexión de la luz que va de la fuente al detector, miremos lo que la teoría cuántica tiene que decir. Regla: la probabilidad de que un suceso determinado tenga lugar es el cuadrado de la flecha final que se obtiene dibujando una flecha para cada camino por el que el suceso puede tener lugar y luego combinando («sumando») las flechas.

En el experimento en el que se medía la reflexión parcial de la luz por dos superficies, existían dos caminos por los que el fotón podía ir de la fuente al detector. En este experimento, existen millones de caminos por los que puede ir el fotón: puede incidir en la parte izquierda del espejo, en A o B (por ejemplo), y reflejarse hacia el detector (ver Fig. 20); puede reflejarse en la parte que piensa que debería de hacerlo, en G; o, puede incidir en la parte derecha del espejo, en K o M, y reflejarse hacia el detector. Pueden pensar que estoy loco porque en la mayoría de los caminos que les he mencionado para que el fotón se reflejase en el espejo, los ángulos no eran iguales. Pero *no* estoy loco, porque ¡esa es la forma en que la luz viaja realmente! ¿Cómo puede ser?

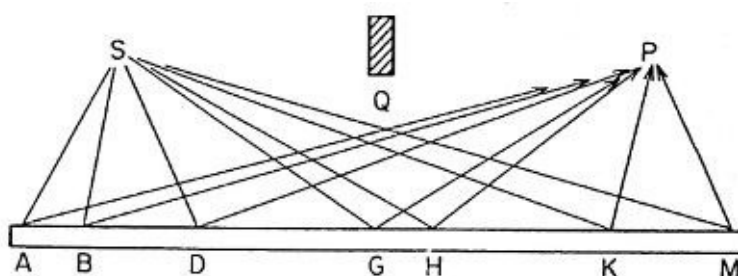


FIGURA 20. El punto de vista cuántico del mundo establece que la luz se refleja con la misma amplitud desde cualquier parte del espejo, desde A a M.

Para simplificar el problema, supongamos que el espejo consiste solo en una tira larga de izquierda a derecha —también es conveniente, por el momento, olvidar que el espejo sobresale del papel (ver Fig. 21)—. Mientras que existen, en realidad, millones de sitios en donde la luz puede reflejarse en esta tira de espejo, hagamos una aproximación temporal dividiendo el espejo en un número finito de pequeños cuadrados, y considerando sólo un camino para cada cuadrado —nuestro cálculo se hace más preciso (pero más arduo también) al hacer los cuadrados más pequeños y considerar más caminos.



FIGURA 21. Para calcular con más facilidad a donde va la luz, consideramos, temporalmente, sólo una tira de espejo dividida en cuadrados, con un camino por cada cuadrado. Esta simplificación en modo alguno disminuye la realización de un análisis preciso de la situación.

Ahora dibujemos una pequeña flecha para cada camino por el que la luz, en esta situación, puede viajar. Cada flechita tiene una cierta longitud y dirección. Consideremos la longitud en primer lugar. Podrían pensar que la flecha que dibujamos representando el camino que va al centro del espejo, a G, es con mucho la más larga (puesto que parece existir una gran probabilidad de que cualquier fotón que llegue al detector lo haga de esta manera), y que

las flechas para los caminos desde los extremos del espejo deben de ser muy cortas. No, no; no debemos hacer una regla tan arbitraria. La regla correcta — lo que realmente ocurre— es mucho más simple: un fotón que llega al detector tiene casi la misma posibilidad de hacerlo por *cualquier* camino, por lo que las pequeñas flechas tienen todas casi idéntica longitud. (Existen, en realidad, algunas ligeras variaciones de longitud debido a los distintos ángulos y distancias implicadas, pero son tan nimias que voy a ignorarlas). Por consiguiente digamos que cada pequeña flecha que dibujamos tiene una longitud arbitraria uniforme —haré su longitud muy pequeña porque existen muchas de estas flechas que representan los muchos caminos en que la luz puede viajar (ver Fig. 22).



FIGURA 22. Cada uno de los caminos por los que puede viajar la luz será representado en nuestros cálculos por una flecha de longitud patrón arbitraria, tal y como se muestra.

Aunque es válido suponer que la longitud de todas las flechas es casi la misma, las direcciones son claramente distintas porque sus tiempos son diferentes —como recordarán de la primera conferencia, la dirección de una flecha particular está determinada por la posición final de un cronógrafo imaginario que cronometra a un fotón en su movimiento a lo largo de un camino particular—. Cuando un fotón va hacia la izquierda del espejo, hacia A, y luego al detector, tarda claramente más tiempo que un fotón que va al detector reflejándose en el centro del espejo, en G (ver Fig. 23). O, imaginen por un momento que tienen prisa y tienen que correr desde la fuente hasta el espejo y luego al detector. Saben que ciertamente no es una buena idea el ir hasta A y luego todo el camino hasta el detector; sería mucho más rápido tocar el espejo en algún lugar próximo a su centro.

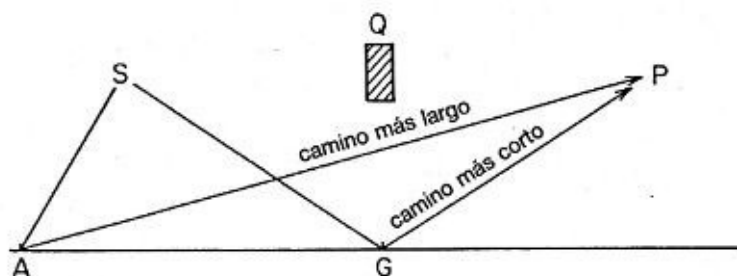


FIGURA 23. Mientras que la longitud de cada flecha es esencialmente la misma, la dirección será diferente porque el tiempo que tarda el fotón es distinto para cada camino. Claramente, tarda más en ir desde S a P por A que desde S a P por G.

Para ayudarnos a calcular la dirección de cada flecha, voy a dibujar una

gráfica debajo del esquema representativo del espejo (ver Fig. 24). En línea con cada parte del espejo en donde se puede reflejar la luz, voy a señalar, en verticales, cuánto tiempo se tardaría si la luz viajase por ese camino. Cuanto más tiempo se emplee, más alto estará el punto en la gráfica. Comenzando por la izquierda, el tiempo que tarda un fotón en hacer el camino que se refleja en A es bastante grande, de manera que dibujamos un punto bastante alto en la gráfica. Al desplazarnos hacia el centro del espejo, el tiempo que emplea un fotón en recorrer el camino, en la forma en que nos movemos, se reduce, de modo que dibujamos cada punto sucesivamente más bajo que el anterior. Después de pasar el centro del espejo, el tiempo que emplea un fotón en cada uno de los sucesivos caminos es cada vez más grande, así que representamos nuestros correspondientes puntos cada vez más altos. Para visualizarlo mejor unamos los puntos: forman una curva simétrica que empieza alta, baja y luego sube de nuevo.

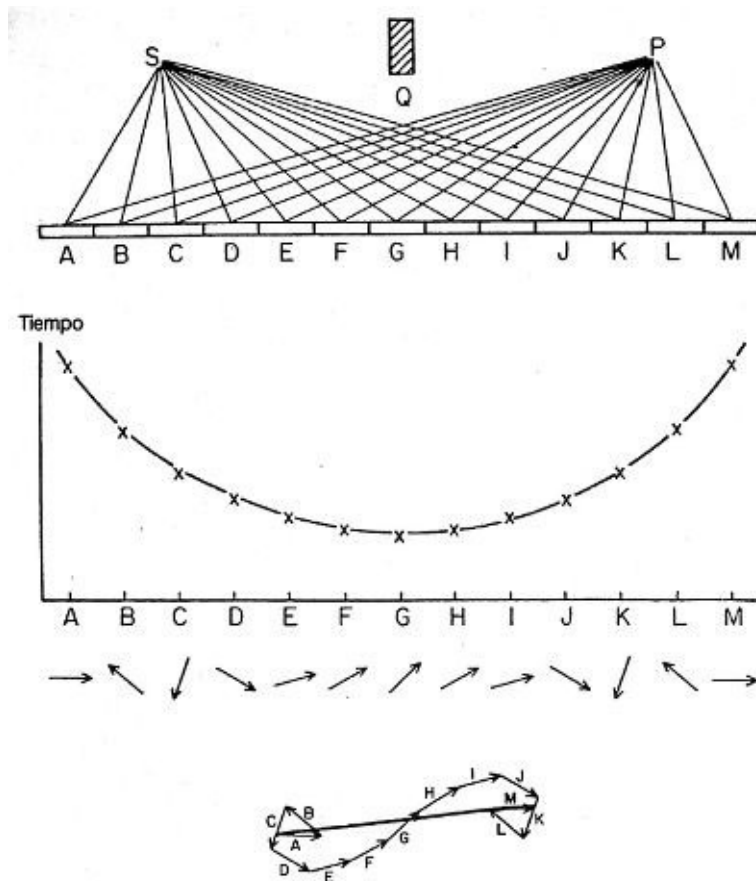


FIGURA 24. Cada camino por el que puede viajar la luz (en esta situación simplificada) se muestra en la parte superior; con un punto en la gráfica de debajo mostrando el tiempo que tarda el fotón en ir desde la fuente a ese punto del espejo y luego al fotomultiplicador. Debajo de esta gráfica está la dirección de cada flecha, y en la parte inferior el resultado de sumar todas las flechas. Es evidente que la mayor contribución a la longitud de la flecha final proviene de las flechas E hasta I, cuyas direcciones son casi iguales porque el tiempo de sus caminos es casi el mismo. También ocurre que aquí es donde el tiempo es mínimo. Es por tanto bastante aproximado

el decir que la luz va por donde el tiempo es mínimo.

Bien, ¿qué relación tiene esto con la dirección de las flechitas? La dirección de una flecha específica corresponde al tiempo que tardaría un fotón en ir desde la fuente al detector siguiendo ese camino específico. Dibujemos las flechas comenzando por la izquierda. El camino A es el que lleva más tiempo, su flecha señala en una dirección (Fig. 24). La flecha del camino B señala en dirección diferente porque su tiempo es diferente. En el centro del espejo, las flechas F, G y H señalan casi en la misma dirección porque sus tiempos son casi los mismos. Después de pasar el centro del espejo, vemos que cada camino a la derecha se corresponde con un camino a la izquierda cuyo tiempo es exactamente el mismo (esto es consecuencia de colocar la fuente y el detector a la misma altura, y el camino por G exactamente en el medio). Así la flecha para el camino J, por ejemplo, tiene la misma dirección que la flecha para el camino D.

Sumemos ahora las pequeñas flechas (Fig. 24). Comenzando con la flecha A, unimos las flechas entre sí, cabeza con cola. Bien, si fuésemos a dar un paseo utilizando cada pequeña flecha como un paso, no iríamos muy lejos al principio porque la dirección de un paso a otro es muy distinta. Pero al cabo de un rato las flechas empiezan a señalar generalmente en la misma dirección. Finalmente, próximo al final de nuestro paseo, la dirección entre pasos es de nuevo bastante diferente, de modo que hacemos unas cuantas eses más.

Todo lo que tenemos que hacer ahora es dibujar la flecha final. Simplemente conectamos la cola de la primera flechita con la cabeza de la última y vemos cuanto hemos progresado en línea recta en nuestro paseo (Fig. 24). ¡Y he aquí que obtenemos una flecha final de tamaño considerable! ¡La teoría de la electrodinámica cuántica predice que la luz, sin duda, se refleja en el espejo!

Bien, investiguemos. ¿Qué determina la longitud de la flecha final? Notamos un número de cosas. Primero, los extremos del espejo no son importantes: allí, las flechitas van erráticas y no nos conducen a ninguna parte. Si cortase los extremos del espejo —partes en las que Vds. instintivamente sabían que yo estaba malgastando mi tiempo entreteniéndome con ellas— apenas afectaría a la longitud final de mi flecha.

Entonces, ¿cuál es la parte de espejo que contribuye sustancialmente a la longitud de la flecha final? Es la parte en donde las flechas señalan todas en casi la misma dirección, porque su *tiempo* es casi el *mismo*. Si miran el gráfico que muestra el tiempo para cada camino (Fig. 24), verán que el tiempo

es casi el mismo para los caminos de la parte inferior de la curva donde el *tiempo es mínimo*.

Resumiendo, donde el tiempo es mínimo es donde también el tiempo para los caminos próximos es casi el mismo; esto es, donde las flechitas señalan en casi la misma dirección y contribuyen de manera substancial a la longitud, es donde se determina la probabilidad de un fotón reflejándose en un espejo. Y esta es la razón por la que, aproximadamente, podemos continuar con la cruda visión del mundo que establece que la luz sólo va donde el *tiempo es menor* (y es fácil probar que donde el tiempo es menor, el ángulo de incidencia es igual al ángulo de reflexión, pero no tengo tiempo de demostrarlo).

Por tanto la teoría de la electrodinámica cuántica da la respuesta correcta —el centro del espejo es la parte importante para la reflexión— pero este resultado correcto surge a expensas de creer que la luz se refleja en todo el espejo, y de tener que sumar un manojito de pequeñas flechas cuyo único propósito es cancelarse. Todo esto puede parecerles una pérdida de tiempo —algún juego tonto sólo para matemáticos—. Después de todo no parece «física real» el tener algo que ¡sólo se cancela!

Comprobemos la idea de que realmente *hay* reflexión a lo largo de todo el espejo haciendo otro experimento. Primero, cortemos la mayor parte del espejo y dejemos alrededor de una cuarta parte, la de la zona izquierda. Todavía nos queda un trozo bastante grande de espejo, pero en el sitio equivocado. En el experimento anterior las flechas del lado izquierdo del espejo señalaban en direcciones muy distintas entre sí debido a la gran diferencia de tiempo entre caminos vecinos (Fig. 24). En este experimento voy a realizar un cálculo más detallado tomando intervalos mucho más pequeños en la parte izquierda del espejo —lo suficientemente estrechos como para que no exista gran diferencia de tiempo entre caminos vecinos (ver Fig. 25)—. Con esta imagen más detallada, vemos que algunas flechas señalan más o menos hacia la derecha; otras más o menos hacia la izquierda. Si sumamos *todas* las flechas, tendremos un conjunto de flechas distribuidas esencialmente en un círculo que no conducen a parte alguna.

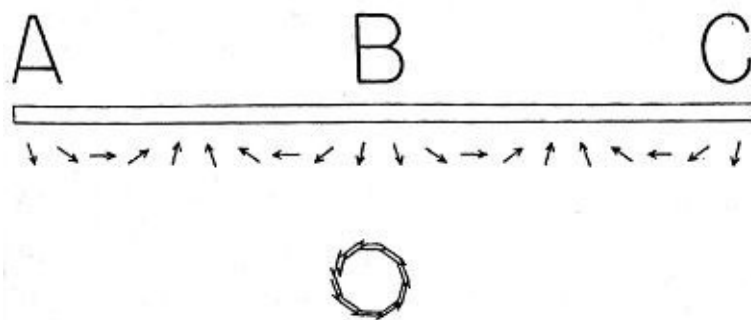


FIGURA 25. Para comprobar la idea de que realmente existe reflexión en los extremos del espejo (pero que se cancela) hacemos un experimento con un trozo grande de espejo colocado en sitio equivocado para la reflexión de S a P. Este trozo de espejo se divide en secciones mucho más pequeñas, de manera que el tiempo entre un camino y el siguiente no sea muy diferente. Cuando se suman todas las flechas, no conducen a ninguna parte: forman un círculo y su suma es muy próxima a nada.

Pero supongamos que rayamos cuidadosamente el espejo en aquellas zonas cuyas flechas se inclinan hacia una determinada dirección —digamos hacia la izquierda— de manera que sólo permanecen aquellos sitios en donde las flechas señalan, en general, en la otra dirección (ver Fig. 26). Cuando sumamos sólo las flechas que señalan más o menos hacia la derecha, obtenemos una serie de depresiones y una flecha final substancial —¡de acuerdo con la teoría deberíamos de tener ahora una gran reflexión!— y de hecho la tenemos, ¡la teoría es correcta! Tal espejo se denomina una red de difracción y funciona que es un encanto.

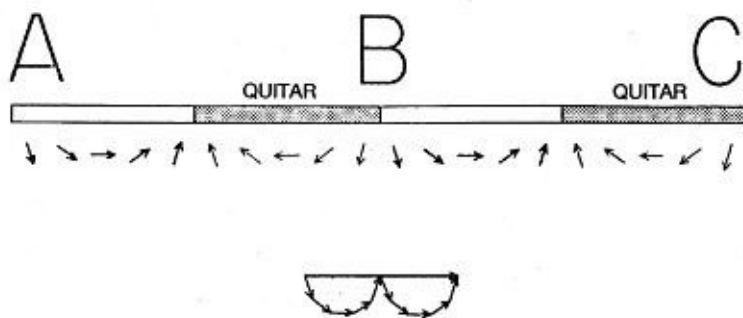


FIGURA 26. Si sólo se suman las flechas con predominio hacia una dirección concreta —tal como hacia la derecha—, mientras que se desprecian las demás (arañando el espejo en los lugares correspondientes), el espejo colocado en el lugar equivocado refleja una cantidad considerable de luz. Un espejo así grabado se denomina red de difracción.

Es maravilloso, Vd. puede coger un trozo de espejo donde no esperaba ninguna reflexión, raspar parte de él y ¡refleja!^[4].

La rejilla particular que acabo de mostrarles ha sido diseñada para la luz roja. No funcionará con luz azul; tendríamos que hacer una nueva rejilla con las tiras cortadas con menor espaciado porque como les dije en la primera conferencia, la manecilla del cronógrafo gira más deprisa cuando sigue a un

fotón azul que cuando sigue a uno rojo. De manera que los cortes diseñados específicamente para la velocidad de giro «roja» no están situados en los lugares correctos para la luz azul; las flechas se acodan y las rejillas no funcionan muy bien. Pero ocurre que si, como por accidente, desplazamos el fotomultiplicador hacia abajo formando un ángulo diferente, la red de difracción hecha para la luz roja funciona ahora para la luz azul. Es un accidente afortunado, una consecuencia de la geometría involucrada (ver Fig. 27).

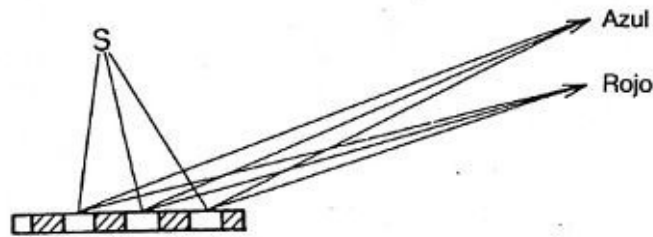


FIGURA 27. Una red de difracción con surcos a la distancia adecuada para la luz roja, también funciona para otros colores si el detector se coloca en un lugar diferente. En consecuencia, dependiendo del ángulo, es posible ver colores diferentes reflejándose en una superficie con surcos —tal como un disco de gramófono.

Si se hace incidir luz blanca sobre la rejilla, la luz roja aparece en un sitio, la naranja ligeramente por encima, seguida de la amarilla, verde y azul — todos los colores del arco iris—. A menudo se pueden ver los colores donde existe una serie de surcos juntos —por ejemplo, cuando sostienen un disco de gramófono (o mejor, un vídeo disco) bajo luz brillante y ángulo correcto—. Quizá Vds. hayan visto esas maravillosas señales plateadas (aquí en la soleada California a menudo se ven en la parte posterior de los coches): cuando se mueve el coche, se ven colores muy brillantes cambiando del rojo al azul. Ahora saben de dónde proceden los colores: están mirando a una red de difracción —un espejo que ha sido arañado en los sitios adecuados—. El sol es la fuente luminosa, y sus ojos son el detector. Puedo continuar fácilmente explicándoles cómo funcionan los láseres y hologramas, pero sé que no todos han visto estas cosas y yo tengo muchas otras cosas de que hablar^[5].

Por consiguiente una red de difracción demuestra que no podemos ignorar las partes de un espejo que no parecen reflejar; si le hiciésemos algunas cosas inteligentes al espejo, podríamos demostrar la realidad de las reflexiones para todas sus partes y producir algunos fenómenos ópticos chocantes.

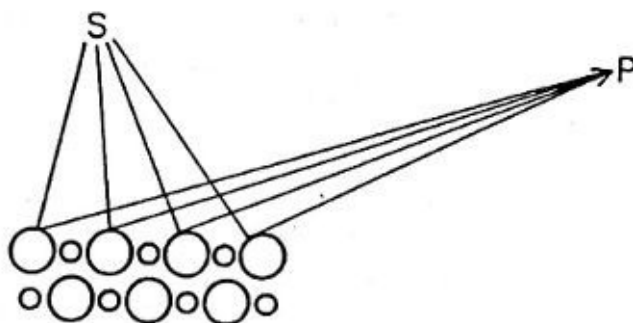


FIGURA 28. La Naturaleza ha creado muchos tipos de redes de difracción en forma de cristales. Un cristal de sal refleja los rayos X (luz para la que la manecilla del cronógrafo imaginario se mueve extremadamente deprisa —quizá 10 000 veces más rápido que para la luz visible—) para determinados ángulos, a partir de los cuales se puede calcular la distribución exacta y el espaciado de los átomos individuales.

Lo que es más importante, al demostrar la realidad de la reflexión por *todas* las partes del espejo demostramos que existe una amplitud —una flecha— para *cada camino* por el que un suceso puede tener lugar. Y a fin de calcular correctamente la probabilidad de un suceso bajo diferentes circunstancias, tenemos que sumar las flechas de *cada camino* por el que el suceso puede ocurrir —¡no sólo las de los caminos que creemos que son importantes!

Bien, me gustaría hablarles de algo más familiar que las redes de difracción: acerca de la luz que va del aire al agua. Esta vez colocaremos el fotomultiplicador debajo del agua —¡suponemos que el experimentador puede hacer eso!—. La fuente luminosa está en el aire en S, y el detector debajo del agua en D (ver Fig. 29). De nuevo, queremos calcular la probabilidad de que un fotón vaya desde la fuente luminosa al detector. Para realizar el cálculo, debemos considerar todos los caminos por los que puede viajar la luz. Cada camino, por el que puede ir la luz, contribuye con una pequeña flecha y, como en el ejemplo previo, todas las flechas tienen casi la misma longitud. Podemos hacer de nuevo una gráfica del tiempo que tarda cada fotón en recorrer cada camino posible. La gráfica será una curva muy parecida a la descrita por la luz reflejándose en un espejo: comienza alta, decrece y sube de nuevo; las contribuciones más importantes provienen de los lugares donde las flechas señalan casi en la misma dirección (donde el tiempo es casi el mismo entre caminos contiguos), es decir, de la parte inferior de la curva. Aquí es donde también el tiempo es mínimo, luego todo lo que tenemos que hacer es descubrir dónde es mínimo el tiempo.

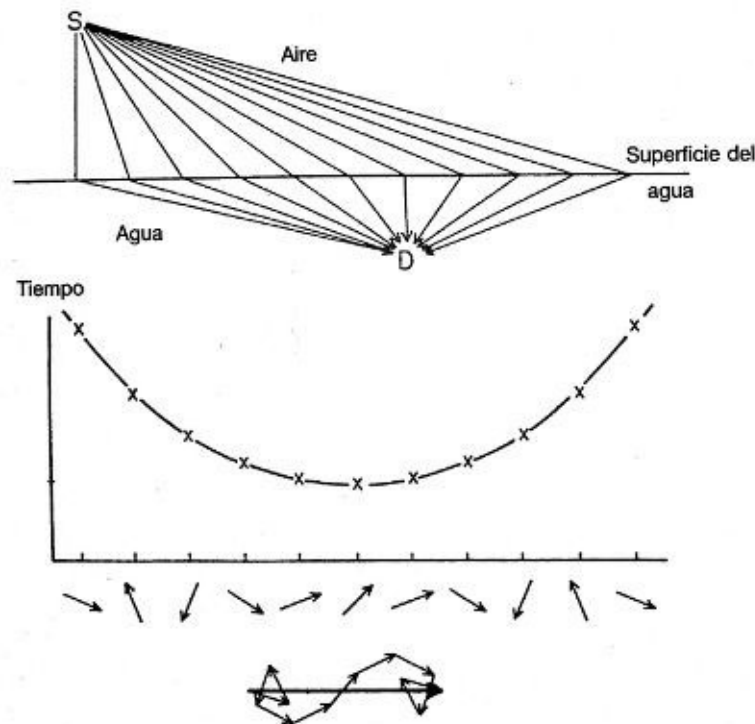


FIGURA 29. La teoría cuántica establece que la luz puede ir desde una fuente en el aire a un detector en el agua por muchos caminos. Si se simplifica el problema, como en el caso del espejo, se puede dibujar una gráfica mostrando los tiempos de cada camino con la dirección de cada flecha debajo. De nuevo, la contribución mayor a la longitud de la flecha final proviene de aquellos caminos cuyas flechas señalan casi en la misma dirección porque sus tiempos son casi iguales; una vez más, aquí es donde el tiempo es mínimo.

Resulta que la luz parece viajar más despacio en el agua de lo que lo hace en el aire (les explicaré la razón en la próxima conferencia) lo que convierte a la distancia a través del agua en «más costosa» por así decir, que la distancia a través del aire. No es difícil imaginar qué camino es el de tiempo menor: Supongamos que Vds. son un guardacostas, sentado en S, y que la chica guapa se está ahogando en D (Fig. 30). Se puede correr en la tierra más deprisa que se puede nadar en el agua. El problema es ¿dónde entrar en el agua para alcanzar más rápidamente a la víctima que se ahoga? ¿Correr hacia el agua en A y luego nadar a toda velocidad? Por supuesto que no. Pero correr en línea recta hacia la víctima y entrar en el agua en J tampoco es el camino más rápido. Mientras que sería tonto para un socorrista analizar y calcular en esas circunstancias, existe una posición calculable en la que el tiempo es mínimo: es un compromiso entre tomar el camino directo, a través de J, y tomar el camino con menor tramo en el agua, a través de N. Y lo mismo ocurre con la luz —el camino de menor tiempo es el que penetra en el agua en un punto entre J y N, tal como L.

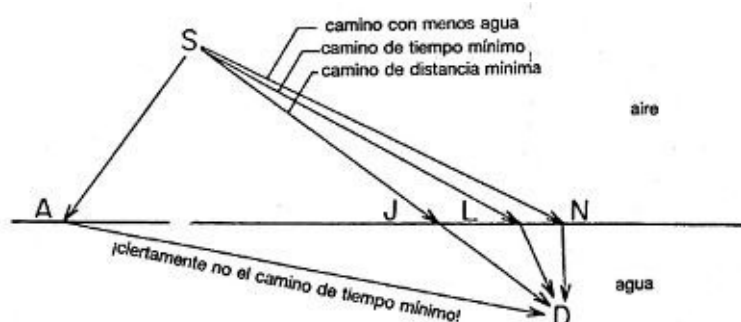


FIGURA 30. Encontrar el camino de tiempo mínimo la luz es análogo a encontrar el camino de tiempo mínimo para un socorrista que corriendo y luego nadando rescatase a una víctima que se está ahogando: el camino más corto tiene demasiada parte con agua; el camino de menor porción de agua tiene demasiada parte con tierra; el camino de menor tiempo es un compromiso entre los dos.

Otro fenómeno de la luz que me gustaría mencionar brevemente es el espejismo. Cuando están conduciendo por una carretera muy caliente, a veces pueden ver lo que parece agua en la carretera. Lo que realmente están viendo es el cielo, y cuando se ve el cielo en la carretera normalmente es porque hay charcos en ella (reflexión parcial de la luz por una única superficie). Pero ¿cómo se puede ver el cielo en la carretera cuando no tiene agua? Lo que necesitan saber es que la luz viaja más despacio en el aire frío que en el caliente, y para que se pueda ver un espejismo, el observador debe de encontrarse en el aire más frío que está por encima del aire caliente de la superficie de la carretera (ver Fig. 31). Cómo es posible ver el cielo mirando hacia *abajo* es algo que se puede comprender encontrando el camino de tiempo mínimo. Les dejo esto para que jueguen en casa —es divertido el pensarlo y muy fácil descubrirlo.



FIGURA 31. Encontrar el camino de tiempo mínimo explica como se produce un espejismo. La luz viaja más deprisa en aire caliente que en aire frío. Parte del cielo parece estar en la carretera porque parte de la luz del cielo llega al ojo proveniente de la carretera. La única ocasión en la que el cielo parece estar en la carretera es cuando lo refleja el agua y así el espejismo hace que parezca que haya agua.

En los ejemplos que les he expuesto de la luz reflejándose en un espejo y la luz atravesando aire y luego agua, estaba haciendo una aproximación: por motivos de simplificación, dibujé los distintos caminos en los que la luz puede viajar como dos líneas rectas —dos líneas rectas que forman un ángulo—. Pero no tenemos que *suponer* que la luz viaja en línea recta cuando está

en un medio uniforme como el aire o el agua; incluso *esto* se puede explicar por el principio general de la teoría cuántica: la probabilidad de un suceso se encuentra sumando las flechas correspondientes a *todos* los caminos por los que puede ocurrir.

Así, en nuestro próximo ejemplo, voy a mostrarles cómo, sumando flechitas, puede parecer que la luz viaja en línea recta. Coloquemos una fuente y un fotomultiplicador en S y P, respectivamente (ver Fig. 32), y consideremos *todos* los caminos por los que puede ir la luz —todo tipo de caminos sinuosos— desde la fuente del detector. Entonces dibujamos una flechita para cada camino y ¡estamos aprendiendo la lección bien!

Para cada camino sinuoso, tal como el camino A, existe otro un poquito más derecho y decididamente más corto —es decir, que lleva mucho menos tiempo—. Pero cuando los caminos son casi rectos —como por ejemplo en C — un camino próximo y más derecho tiene casi el mismo tiempo. Aquí es donde las flechas se suman en lugar de cancelarse; aquí es por donde va la luz.

Es importante señalar que la única flecha que representa el camino en línea recta, a través de D (Fig. 32), no es suficiente para explicar la probabilidad de la luz que va de la fuente al detector. Los caminos próximos, caminos casi rectos —los de C y E por ejemplo— también contribuyen de manera considerable. Luego la luz *realmente* no viaja sólo en línea recta; «olfatea» los caminos próximos y utiliza un pequeño núcleo del espacio cercano (De la misma manera, un espejo tiene que tener un tamaño adecuado para reflejar normalmente: si el espejo es muy pequeño con respecto al núcleo de caminos vecinos, la luz se difunde en muchas direcciones independientemente de donde se coloque el espejo).

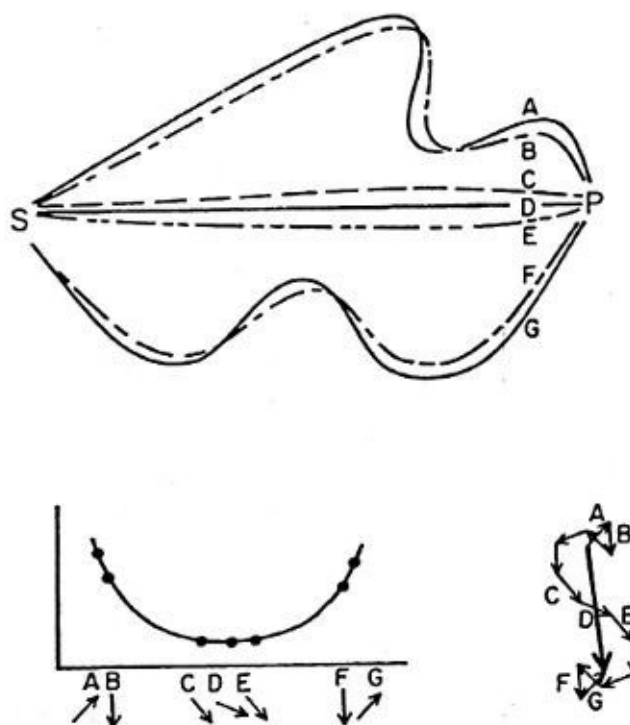


FIGURA 32. La teoría cuántica se puede utilizar para demostrar el por qué la luz parece viajar en línea recta. Cuando se han considerado todos los caminos posibles, cada camino sinuoso tiene un camino vecino de distancia considerablemente inferior y consecuentemente de tiempo mucho menor (y con una flecha en dirección substancialmente distinta). Sólo los caminos próximos al camino recto D tienen flechas señalando casi en la misma dirección, porque sus tiempos son casi los mismos. Únicamente estas flechas son importantes, porque ellas son las que permiten obtener una flecha final grande.

Investiguemos este núcleo de luz con más detenimiento colocando una fuente en S, un fotomultiplicador en P, y un par de bloques entre ellos para evitar que las trayectorias de la luz se alejen demasiado (ver Fig. 33). Ahora, coloquemos un segundo fotomultiplicador en Q, debajo de P, y supongamos de nuevo, por sencillez, que la luz puede ir de S a Q sólo mediante caminos de dobles líneas rectas. Ahora ¿qué ocurre? Cuando la separación entre bloques es lo suficientemente grande como para permitir muchos caminos vecinos hacia P y Q, las flechas de los caminos que llevan a P se suman (porque todos los caminos hacia P suponen el mismo tiempo), mientras que las correspondientes a caminos hacia Q se anulan (porque estos caminos tienen una diferencia considerable de tiempo). Entonces el fotomultiplicador en Q no hace click.

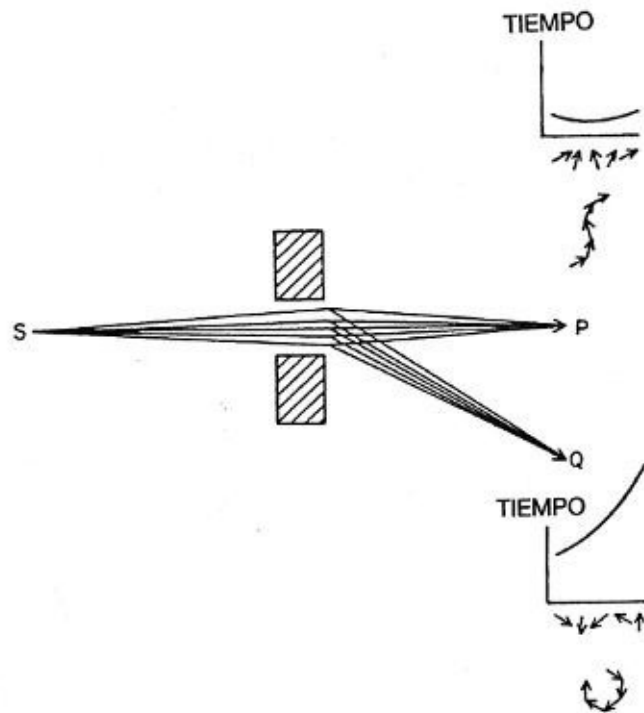


FIGURA 33. La luz viaja no sólo en línea recta sino también a través de los caminos cercanos. Cuando los dos bloques se separan lo suficiente como para permitir la existencia de estos caminos, los fotones normalmente llegan a P, y casi nunca a Q.

Pero si acercamos los bloques entre sí, en un determinado momento ¡el detector en Q comienza a hacer clicks! Cuando el espaciado entre bloques es muy pequeño y sólo existen unos cuantos caminos vecinos, las flechas de los que van hacia Q *también* se suman, porque apenas existe diferencia de tiempo entre ellos (ver Fig. 34). Por supuesto, ambas flechas finales son muy pequeñas de modo que no hay mucha luz en ninguno de los dos caminos al pasar a través de un agujero tan pequeño, ¡pero el detector en Q hace casi tantos clicks como el de P! Por tanto, cuando intentamos constreñir mucho la luz para asegurarnos de que sólo viaja en línea recta, ésta se niega a colaborar y empieza a desperdigarse^[6].

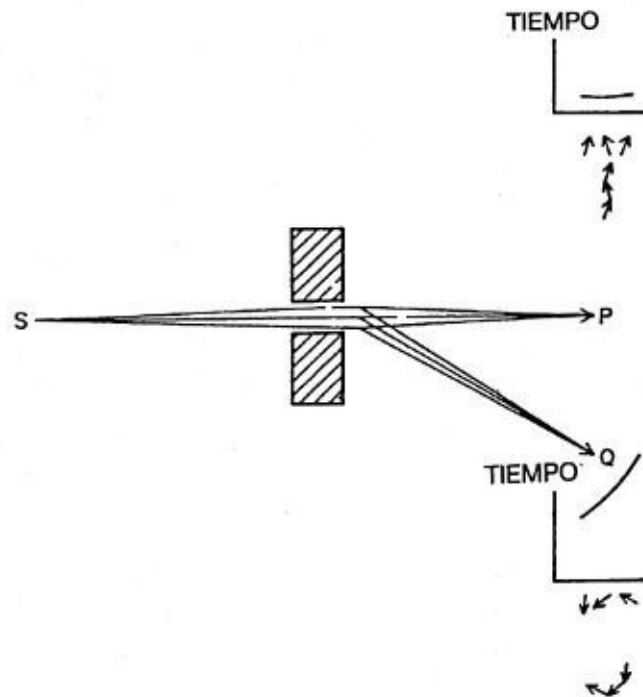


FIGURA 34. Cuando la luz se restringe tanto que sólo unos pocos caminos son posibles y la luz que es capaz de pasar a través de la rendija estrecha va tanto a Q como a P, porque no existen flechas suficientes representando los caminos a Q como para que estos se anulen.

Por consiguiente la idea de que la luz viaja en línea recta es una aproximación conveniente para describir lo que ocurre en nuestro mundo familiar; es similar a la cruda aproximación que establece que cuando la luz se refleja en un espejo, el ángulo de incidencia es igual al ángulo de reflexión.

De la misma manera que fuimos capaces de llevar a cabo un truco ingenioso para hacer que la luz se reflejara con muchos ángulos en un espejo, podemos realizar un truco similar para hacer que la luz vaya de un punto a otro por varios caminos.

En primer lugar simplifiquemos la situación. Voy a dibujar una línea vertical punteada (ver Fig. 35) entre la fuente luminosa y el detector (la línea no significa nada, es una línea artificial) y a establecer que los únicos caminos que vamos a considerar son los de doble línea recta. La gráfica que muestra el tiempo para cada camino tiene el mismo aspecto que en el caso del espejo (pero esta vez la dibujaré de lateral): la curva empieza en A, en la parte superior, y luego retrocede, porque los caminos en la parte central son más cortos y llevan menos tiempo. Finalmente, la curva sale de nuevo.

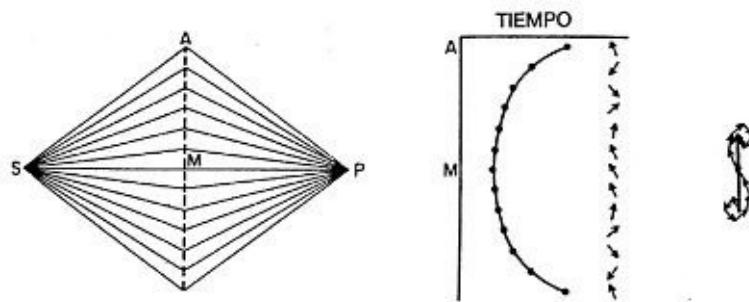


FIGURA 35. El análisis de todos los caminos posibles de S a P se ha simplificado y se incluyen sólo los de doble línea recta (en un único plano). El efecto es el mismo que para el, más complicado, caso real: existe una curva de tiempos con un mínimo, de donde procede la contribución mayor a la flecha final.

Bien, divirtámonos un rato. «Engañemos a la luz» de manera que *todos* los caminos requieran el mismo tiempo. ¿Cómo podemos hacerlo? ¿Cómo podemos hacer que el camino más corto a través de M lleve exactamente el mismo tiempo que el camino más largo a través de A?

Bien, la luz viaja más despacio en el agua que en el aire; también va más lenta en el cristal (¡qué es mucho más fácil de manejar!). Por tanto si colocamos un cristal del espesor adecuado en el camino más corto, a través de M, podemos hacer que el tiempo para este camino sea exactamente el mismo que el del camino A. Los caminos próximos a M, que son un poquito más largos, no requerirán de la misma cantidad de cristal (ver Fig. 36). Cuanto más próximos estemos a A, menos cristal necesitaremos para retardar la luz. Colocando el espesor adecuado de cristal, cuidadosamente calculado, necesario para compensar el tiempo a lo largo de cada camino, podemos repetir el experimento muchas veces. Cuando dibujemos las flechas de cada camino por el que puede viajar la luz, encontraremos que hemos conseguido enderezarlas todas —y hay en realidad *millones* de flechitas— luego el resultado final es una flecha final inesperadamente enorme, ¡sensacionalmente grande! Por supuesto saben lo que estoy describiendo: una lente focalizadora. Arreglando las cosas para que todos los tiempos sean iguales podemos focalizar la luz —podemos hacer que la probabilidad de que la luz llegue a un punto determinado sea muy alta, y que sea muy baja en cualquier otro.

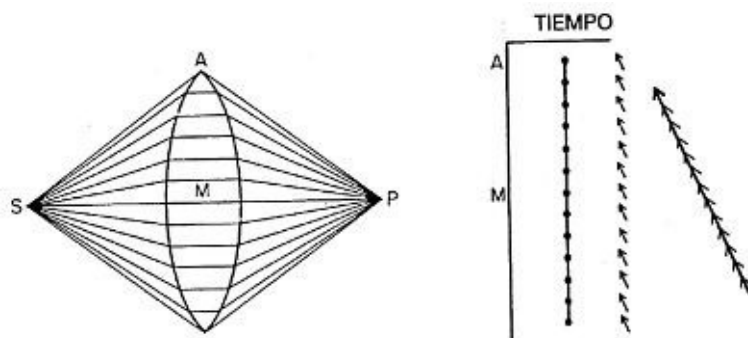


FIGURA 36. Se puede «engañar» a la Naturaleza reduciendo la velocidad de la luz que va por los caminos más cortos: insertando cristal del espesor adecuado de manera que todos los caminos lleven el mismo tiempo. Esto hace que todas las flechas tengan la misma dirección y que produzcan una flecha final enorme —¡muchísima luz!—. Un trozo de cristal diseñado para aumentar mucho la probabilidad de la luz que va desde la fuente a un punto determinado se denomina una lente focalizadora.

He utilizado estos ejemplos para mostrarles cómo la teoría de la electrodinámica cuántica, que parece al principio una idea absurda, sin causalidad, mecanismo o algo real en ella, reproduce efectos que les son familiares: la luz reflejándose en un espejo, la luz desviándose cuando pasa del aire al agua y la luz focalizándose mediante una lente. También reproduce otros ejemplos que pueden o no haber visto, tales como las redes de difracción y cierto número de otras cosas. De hecho, la teoría continúa teniendo éxito explicando todo fenómeno luminoso.

Les he mostrado, con ejemplos, cómo calcular la probabilidad de un suceso que puede tener lugar por *caminos alternativos*: dibujamos una flecha por cada camino en que pueda ocurrir y luego súmanos las flechas. «Sumar flechas» significa que las flechas se colocan cabeza contra cola y se dibuja una «flecha final». El cuadrado de la flecha final resultante representa la probabilidad del suceso.

Para darles una idea más completa de la teoría cuántica, me gustaría mostrarles ahora cómo calculan los físicos las probabilidades de los sucesos compuestos —sucesos que pueden descomponerse en una serie de pasos, o sucesos que consisten en un número de cosas ocurriendo independientemente.

Un ejemplo de suceso compuesto puede conseguirse modificando nuestro primer experimento, en el que enviábamos algunos fotones rojos a una superficie de cristal para medir la reflexión parcial. En lugar de colocar el fotomultiplicador en A (ver Fig. 37) pongamos una pantalla con un agujero que deja que los fotones que llegan al punto A la atraviesen. Luego, pongamos una lámina de cristal en B y coloquemos el fotomultiplicador en C. ¿Cómo determinamos la probabilidad de que un fotón llegue desde la fuente

hasta C?

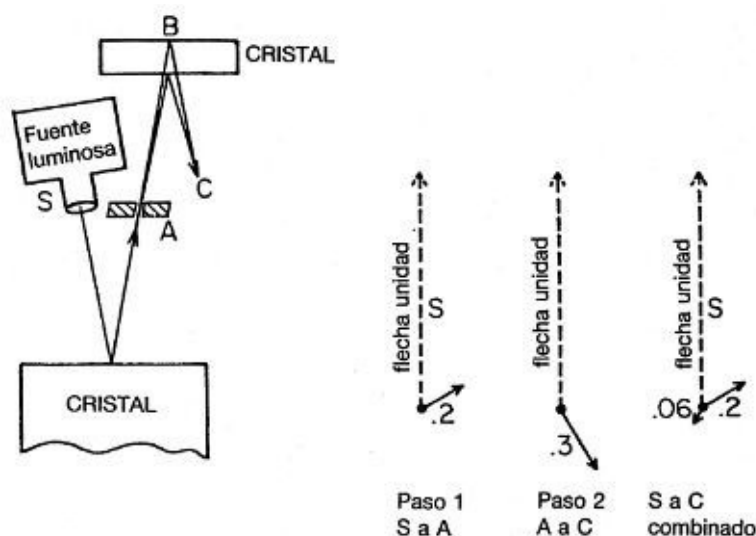


FIGURA 37. Un suceso compuesto puede analizarse como una sucesión de pasos. En este ejemplo, el camino de un fotón que va de S a C puede dividirse en dos pasos: 1) el fotón va de S a A, y 2) el fotón va de A a C. Cada paso se puede analizar por separado y obtener una flecha que se puede considerar de una nueva manera: como una flecha unidad (una flecha de longitud 1 señalando las 12 en punto) que ha experimentado una reducción y un giro. En este ejemplo, la reducción y el giro para el paso 1 es 0,2 y las 2 en punto; la reducción y el giro para el paso 2 es 0,3 y las 5 en punto. Para obtener la amplitud de los dos pasos sucesivos, reducimos y giramos en sucesión: la flecha unidad se reduce y gira para obtener en sucesión: la flecha unidad se reduce y gira para obtener una flecha de longitud 0,2 girada hacia las 2 que a su vez se reduce y gira (como si se tratase de una flecha unidad) en 0,3 y el equivalente a las 5 en punto para obtener una flecha de longitud 0,06 y girada a las 7 en punto. Este proceso de reducciones y giros sucesivos se denomina «multiplicación» de flechas.

Podemos pensar en este suceso como una secuencia de dos pasos. Paso 1: un fotón va desde la fuente al punto A, reflejándose en la superficie del cristal. Paso 2: el fotón va desde el punto A al fotomultiplicador en C, reflejándose en la lámina de cristal en B. Cada paso tiene una flecha final —una «amplitud» (voy a intercambiar las palabras)— que puede calcularse de acuerdo con las reglas que conocemos hasta ahora. La amplitud para el primer paso tiene una longitud de 0,2 (cuyo cuadrado es 0,04, la probabilidad de reflexión por una superficie de cristal) y forma un cierto ángulo —digamos que señala las dos en punto (Fig. 37).

Para calcular la amplitud del segundo paso, colocamos temporalmente la fuente luminosa en A y dirigimos los fotones hacia la lámina de cristal que está por encima. Dibujamos flechas para la reflexión por la superficie frontal y la posterior y las sumamos —digamos que terminamos con una flecha final de longitud 0,3 y que señala las cinco en punto.

Bien, ¿cómo combinamos las dos flechas para dibujar la amplitud del

suceso total? Consideremos cada flecha bajo nuevo aspecto: como una instrucción para una *reducción* y un *giro*.

En este ejemplo, la primera amplitud tiene una longitud de 0,2 y señala hacia las dos. Si comenzamos con una «flecha unidad» —una flecha de longitud 1 señalando hacia arriba— podemos *reducirla* desde 1 a 0,2 y *gírala* desde las doce hasta las dos. La amplitud del segundo paso puede considerarse como una reducción de la flecha unidad desde 1 a 0,3 y un giro desde las doce hasta las cinco.

Ahora, para combinar las amplitudes de los dos pasos, reducimos y giramos en *sucesión*. Primero, reducimos la flecha unidad desde 1 a 0,2 y la giramos desde las 12 hasta las 2; luego reducimos más la flecha, desde 0,2 hasta tres décimas de esto y la giramos lo equivalente de las 12 a las 5 —es decir, la giramos desde las dos hasta las siete en punto—. La flecha resultante tiene una longitud de 0,06 y señala las 7 en punto. Representa una probabilidad de 0,06 al cuadrado, o 0,0036.

Observando las flechas cuidadosamente vemos que el resultado de reducir y girar dos flechas en sucesión es equivalente a sumar sus ángulos (las 2 + las 5) y multiplicar sus longitudes ($0,2 \times 0,3$). Comprender por qué sumados los ángulos es sencillo: el ángulo de una flecha está determinado por la cantidad de giro de la manecilla del cronógrafo imaginario. Por tanto, la cantidad total de giro para los dos pasos en sucesión es simplemente la suma del giro del primer paso con el giro adicional del segundo.

Por qué se denomina a este proceso «multiplicar flechas» requiere un poco más de explicación, pero es interesante. Consideremos la multiplicación, por un momento, desde el punto de vista de los griegos (esto no tiene nada que ver con la conferencia). Los griegos querían usar números que no eran necesariamente enteros, de modo que representaban los números por líneas. Cualquier número se puede expresar como una *transformación* de la línea unidad —extendiéndola o reduciéndola—. Por ejemplo, si la línea A es la línea unidad (ver Fig. 38), la línea B representa 2 y la línea C representa 3.

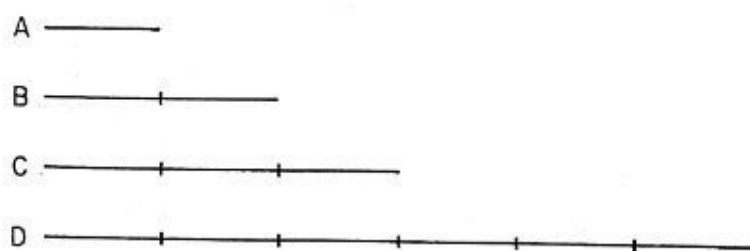


FIGURA 38. Podemos expresar cualquier número como una transformación de la línea unidad

mediante expansiones o reducciones. Si A es la línea unidad, B representa 2 (expansión), y C representa 3 (expansión). La multiplicación de líneas se logra a través de transformaciones sucesivas. Por ejemplo, multiplicar 3 por 2 significa que la línea unidad se expande tres veces y luego 2 veces, dando la respuesta, una expansión de 6 (línea D). Si D es la línea unidad, la línea C representa 1/2 (reducción), la línea B representa 1/3 (reducción) y multiplicar 1/2 por 1/3 significa que la línea unidad D es reducida a 1/2 y luego esto a 1/3, dando como resultado, una reducción a 1/6 (línea A).

Ahora bien, ¿cómo multiplicamos tres veces dos? Aplicamos la transformación *en sucesión*: comenzando con la línea A como unidad, la extendemos dos veces y luego tres veces (o 3 veces y luego 2 veces —el orden es indiferente—). El resultado es la línea D cuya longitud representa 6. ¿Cómo multiplicamos 1/3 por 1/2?, tomando la línea D como línea unidad, primero la reducimos a 1/2 (línea C) y luego ésta a 1/3. El resultado es la línea A que representa 1/6.

Multiplicar flechas funciona de la misma manera (ver Fig. 39). Aplicamos transformaciones en sucesión a la flecha unidad —ocurre que la transformación de una *flecha* implica *dos* operaciones, una reducción y un giro—. Para multiplicar la flecha V por la flecha W, reducimos y giramos la flecha unidad por la cantidad prescrita por V, y luego reducimos y giramos la cantidad prescrita por W —de nuevo, el orden es irrelevante—. Por consiguiente, para multiplicar flechas se sigue la misma regla de transformaciones sucesivas que rige los números habituales^[7].

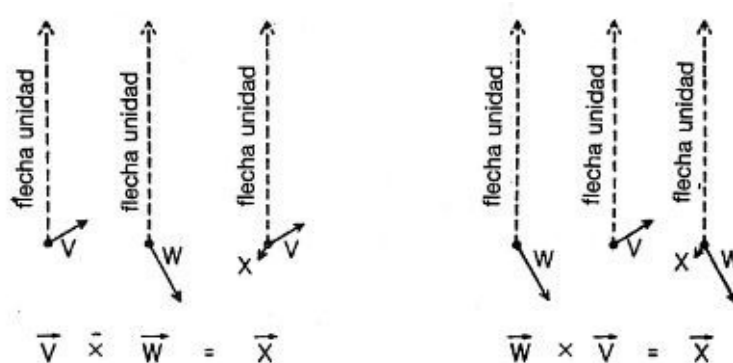


FIGURA 39. Los matemáticos han encontrado que la multiplicación de flechas puede expresarse también como transformaciones sucesivas de la flecha unidad (para nuestros propósitos, una reducción y un giro). Como en la multiplicación normal, el orden es irrelevante: la respuesta, la flecha X, se puede obtener multiplicando la flecha V por la W o la flecha W por la V.

Volvamos al primer experimento de la primera conferencia —la reflexión parcial por una superficie— con esta idea de los pasos en mente (ver Fig. 40). Podemos dividir el camino de reflexión en tres pasos: 1) la luz va desde la fuente hasta el cristal, 2) es reflejada por el cristal, y 3) va del cristal al detector. Cada paso puede ser considerado como una cierta cantidad de reducción y giro de la flecha unidad.

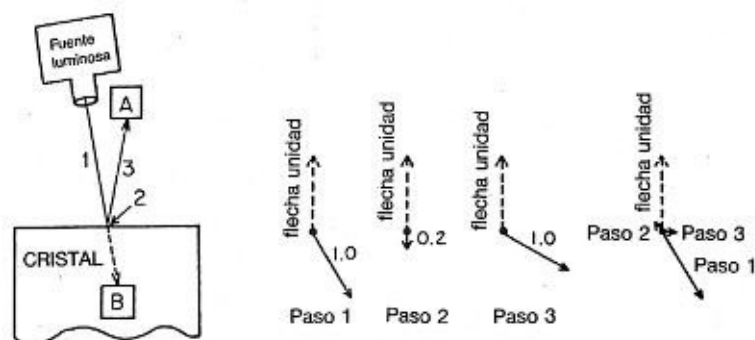


FIGURA 40. La reflexión por una superficie se puede descomponer en tres pasos, cada uno con una reducción y/o un giro de la flecha unidad. El resultado total, una flecha de longitud 0,2 señalando en alguna dirección, es el mismo que antes, pero nuestro método de análisis es ahora más detallado.

Recordarán que en la primera conferencia, no consideramos *todos* los caminos por los que se puede reflejar la luz en un cristal, lo que requiere dibujar y sumar montones y montones de flechitas. Para evitar todo este detalle, di la impresión de que la luz va a un punto particular de la superficie del cristal —que no se dispersa—. Cuando la luz viaja de un punto a otro, en realidad se dispersa (al menos que sea engañada por una lente) y existe una pequeña reducción de la flecha unidad que lleva asociada.

Por el momento, sin embargo, me gustaría ceñirme al punto de vista simplificado de que la luz *no* se dispersa y en consecuencia que es apropiado despreciar esta reducción. También es apropiado suponer que puesto que la luz no se dispersa, cada fotón que abandona la fuente finaliza en A o en B.

Por consiguiente: en el primer paso no hay reducción, pero hay giro —corresponde a la cantidad de giro de la manecilla del cronógrafo imaginario cuando sigue al fotón que va de la fuente a la superficie frontal del cristal—. En este ejemplo, la flecha del primer paso acaba con longitud 1 y un determinado ángulo —digamos señalado a las cinco en punto.

El segundo paso es la reflexión del fotón por el cristal. Aquí hay una reducción considerable —de 1 a 0,2— y un giro de media vuelta. (Estos números ahora parecen arbitrarios: dependen de si la luz es reflejada por un cristal o por cualquier otro material. En la tercera conferencia ¡explicaré también esto!). Luego el segundo paso está representado por una amplitud de longitud 0,2 y una dirección señalando a las 6 (media vuelta).

El último paso es el del fotón yendo del cristal al detector. Aquí, como en el primer paso, no hay reducción, pero hay giro —digamos que la distancia es ligeramente más corta que en el paso 1 y que la flecha señala a las 4 en punto.

Ahora «multiplicamos» las flechas 1, 2 y 3 en sucesión (sumar los ángulos y multiplicar sus longitudes). El efecto total de los tres pasos —1) giro, 2) reducción y media vuelta, vuelta, y 3) giro— es el mismo que el de la primera conferencia: el giro de los pasos 1 y 3 —(las 5 más las 4)— es el mismo que se obtiene cuando dejamos que el cronógrafo funcione a lo largo de todo el recorrido (las nueve en punto); la media vuelta extra del paso 2 hace que la flecha señale la dirección opuesta a la manecilla del cronógrafo, como ocurría en la primera conferencia, y la reducción a 0,2 del segundo paso deja una flecha cuyo cuadrado representa el 4% de reflexión parcial observado para una superficie.

En este experimento hay un problema que no consideramos en la primera conferencia: ¿qué ocurre con los fotones que van a B —los que son transmitidos por la superficie del cristal—?

La amplitud de un fotón que llegue a B debe de estar cerca de 0,98 puesto que $0,98 \times 0,98 = 0,9604$ lo que es muy próximo al 96%. Esta amplitud se puede analizar igualmente descomponiéndola en pasos (ver Fig. 41).

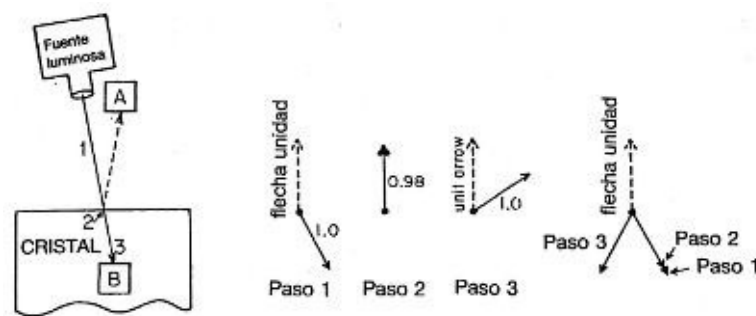


FIGURA 41. La transmisión a través de una superficie también se puede dividir en tres pasos, con una reducción y/o giro en cada uno. Una flecha de longitud 0,98 tiene un cuadrado de aproximadamente 0,96, representando una probabilidad de transmisión del 96% (lo que combinado con la probabilidad de reflexión del 4%, da cuenta del 100% de luz).

El primer paso es el mismo que para el camino hacia A —el fotón va desde la fuente luminosa al cristal— la flecha unidad gira hacia las 5 en punto.

El segundo paso es el del fotón atravesando la superficie del cristal: no existe giro asociado a la transmisión, sólo una ligera reducción —al 0,98.

El tercer paso —el fotón viajando por el interior del cristal— implica un giro adicional pero sin reducción.

El resultado total es una flecha de longitud 0,98 girada en alguna dirección, cuyo cuadrado representa la probabilidad de que un fotón llegue a B —el 96%.

Ahora consideremos de nuevo la reflexión parcial por dos superficies. La reflexión por la superficie frontal es equivalente a la reflexión por una única superficie, luego los tres pasos para la superficie frontal son los mismos que acabamos de ver (Fig. 40).

La reflexión por la superficie posterior se puede descomponer en siete pasos (ver Fig. 42). Implica el girar una cantidad igual al giro de la manecilla del cronógrafo que sigue a un fotón a lo largo de toda la distancia (pasos 1, 3, 5 y 7), una reducción de 0,2 (paso 4) y dos reducciones a 0,98 (pasos 2 y 6). La flecha resultante termina en la misma dirección que antes, pero su longitud es del orden de 0,192 ($0,98 \times 0,2 \times 0,98$), que yo aproximé a 0,2 en la primera conferencia.

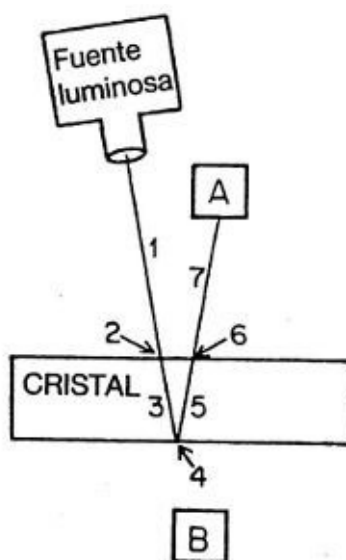


FIGURA 42. La reflexión por la superficie posterior de una lámina de cristal se puede dividir en siete pasos. Los pasos 1, 3, 5 y 7 implican sólo un giro, los pasos 2 y 6 implican reducciones a 0,98 y el paso 4 implica una reducción a 0,2. El resultado es una flecha de longitud 0,192 —que en la primera conferencia se aproximó a 0,2— girada en un ángulo que corresponde a la cantidad total de giro de la manecilla del cronógrafo imaginario.

Resumiendo, aquí están las reglas para la reflexión y transmisión de la luz por el cristal: 1) la reflexión desde el aire hacia el aire (por la superficie frontal) implica una reducción de 0,2 y una media vuelta; 2) la reflexión desde el cristal al cristal (por la superficie posterior) también implica una reducción a 0,2 pero sin giro; y 3) la transmisión desde el aire al cristal o del cristal al aire implica una reducción de 0,98 y sin giro en ambos casos.

Quizás es demasiado, pero no puedo resistir el mostrarles un bonito ejemplo más de cómo funcionan las cosas y cómo son analizadas mediante estas reglas de pasos sucesivos. Movamos el detector a una posición por debajo del cristal, y consideremos algo de lo que no hablamos en la primera

conferencia —la probabilidad de *transmisión* por dos superficies de cristal (ver Fig. 43).

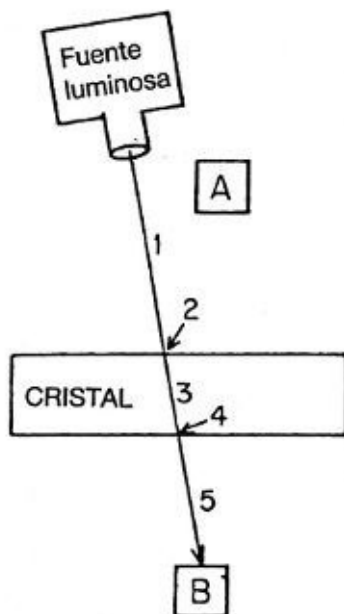


FIGURA 43. La transmisión por dos superficies puede descomponerse en cinco pasos. El paso 2 reduce la flecha unidad a 0,98, el paso 4 reduce la flecha de 0,98 a 0,98 de su valor (aproximadamente 0,96); los pasos 1, 3 y 5 implican sólo un giro. La flecha resultante, de longitud 0,96, tiene un cuadrado de aproximadamente 0,92, representando una probabilidad de transmisión por las dos superficies del 92% (que se corresponde con el esperado 8% de reflexión, correcto «dos veces al día»). Cuando el espesor de la lámina es el adecuado para producir el 16% de reflexión, con un 92% de probabilidad de transmisión, ¡se obtiene el 108% de la luz! ¡Algo está mal en este análisis!

Naturalmente que conocen la respuesta: la probabilidad de que un fotón llegue a B es simplemente el 100% menos la probabilidad de que llegue a A, que ya hemos calculado anteriormente. Así, si encontramos que la probabilidad de llegar a A es el 7%, la probabilidad de llegar a B debe de ser el 93%. Y como la probabilidad de A varía de cero al 16% pasando por el 8% (debido a los distintos espesores del cristal), la probabilidad para B varía del 100% al 84% pasando por el 92%.

Esta es la respuesta correcta, pero estamos esperando calcular *todas* las probabilidades por medio del cuadrado de una flecha final. ¿Cómo calculamos la flecha para la amplitud de transmisión por una lámina de cristal, y cómo se las arregla para variar su longitud de manera tan apropiada que se ajuste, en cada caso, con la correspondiente longitud de A de forma que la probabilidad para A y la probabilidad para B siempre sumen exactamente el 100%? Consideremos los detalles.

Para que un fotón vaya de la fuente al detector, situado debajo del cristal

en B, se precisan cinco pasos. Reduzcamos y giremos la flecha unidad a medida que recorremos el camino.

Los tres primeros pasos son los mismos que los del ejemplo anterior: el fotón va de la fuente al cristal (giro, no reducción); el fotón es transmitido por la superficie frontal (sin giro, reducción a 0,98), el fotón atraviesa el cristal (giro, no reducción).

El cuarto paso —el fotón atraviesa la superficie posterior del cristal— es análogo al segundo paso en lo que se refiere a reducciones y giros: no giros, sino una reducción a 0,98 del 0,98, luego la flecha tiene ahora una longitud de 0,96.

Finalmente, el fotón va de nuevo a través del aire hacia el detector —esto significa más giro, pero no más reducción—. El resultado es una flecha de longitud 0,96 señalando en una dirección determinada por los sucesivos giros de la manecilla del cronógrafo.

Una flecha de longitud 0,96 representa una probabilidad de aproximadamente el 92% (el cuadrado de 0,96) lo que significa que una media de 92 fotones, de cada 100 que salen de la fuente, llegan a B. Esto también significa que el 8% de los fotones son reflejados por las dos superficies y llegan a A. Pero encontramos en nuestra primera conferencia que un 8% de reflexión por las dos superficies es el valor correcto sólo algunas veces («dos veces al día»); que en realidad, la reflexión por dos superficies fluctúa cíclicamente desde cero al 16% al aumentar gradualmente el espesor del cristal. ¿Qué ocurre cuando el cristal tiene el espesor adecuado para que la reflexión parcial sea del 16%? De cada 100 fotones que salen de la fuente, 16 llegan a A y 92 a B, lo que significa que el 108% de la luz ha sido detectada —¡qué horror!—. Algo está mal.

¡Olvidamos considerar *todos* los caminos por los que la luz puede llegar a B! Por ejemplo, puede rebotar en la superficie posterior, atravesar en dirección ascendente el cristal como si fuese hacia A, pero reflejarse luego en la superficie frontal y volver hacia B (ver Fig. 44). Este camino supone nueve pasos. Veamos lo que ocurre a la flecha unidad con cada paso dado por la luz (no se preocupen, ¡sólo se reduce y gira!).

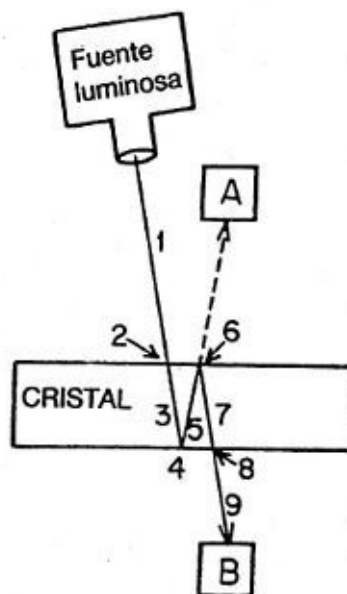


FIGURA 44. A fin de hacer los cálculos más precisos se debe considerar otro camino por el que la luz pueda transmitirse a través de dos superficies. Este camino implica dos reducciones de 0,98 (pasos 2 y 8) y dos reducciones de 0,2 (pasos 4 y 6), resultando en un flecha de longitud 0,0384 (redondeada a 0,04).

Primer paso —el fotón va a través del aire— giro, no reducción. Segundo paso —el fotón penetra en el cristal— no giro, reducción al 0,98. Tercer paso —el fotón va a través del cristal— giro, no reducción. Cuarto paso —reflexión en la superficie posterior— no giro, pero reducción al 0,2 de 0,98 o 0,196. Quinto paso —el fotón asciende a través del cristal— giro, no reducción. Sexto paso —el fotón se refleja en la superficie frontal (es realmente una superficie «posterior», porque el fotón está dentro del cristal)— no giro, pero reducción a 0,2 de 0,196 o 0,0392. Séptimo paso —el fotón retrocede dentro del cristal— más giro, no reducción. Octavo paso —el fotón pasa a través de la superficie posterior— no giro, pero reducción a 0,98 de 0,0392 o 0,0384. Finalmente, el noveno paso —el fotón va hacia el detector a través del aire— giro, no reducción.

El resultado de toda esta reducción y giro es una amplitud de 0,0384 de longitud —digamos 0,04 para los aspectos prácticos— y un ángulo de giro que se corresponde con la cantidad total de giro del cronógrafo cuando sigue al fotón por todo este largo camino. Esta flecha representa un *segundo* camino que puede llevar la luz desde la fuente hasta B. Ahora tenemos dos alternativas, luego *sumemos* las dos flechas —la del camino más directo, cuya longitud es 0,96, y la del camino más largo, cuya longitud es 0,04— y obtendremos la flecha final.

Las dos flechas, en general, no tienen la misma dirección, porque al

cambiar el espesor del cristal cambia la dirección relativa de la flecha de 0,04 con respecto a la de 0.96. Pero vean qué bien funcionan las cosas: la vuelta extra que da el cronógrafo al seguir al fotón durante los pasos 3 y 5 (en su camino hacia A) es igual a la vuelta extra que da al seguir el fotón durante los pasos 5 y 7 (en su camino hacia B). Esto supone que cuando las dos flechas de reflexión se cancelan para dar una flecha final que representa una reflexión nula, las flechas de transmisión se suman para dar una flecha de longitud $0,96 + 0,04$, o 1 —cuando la probabilidad de reflexión es cero, la probabilidad de transmisión es 100% (ver Fig. 45)—. Y cuando las flechas de reflexión se unen para dar una amplitud de 0,4, las flechas de transmisión se oponen, dando una amplitud de longitud $0,96 - 0,04$, o 0,92 —cuando la reflexión se determina como el 16%, la transmisión se calcula como el 84% (o 0,92 al cuadrado)— ¡Ya ven que inteligente es la Naturaleza con sus reglas que aseguran que siempre tendremos el 100% de los fotones considerados!^[8]

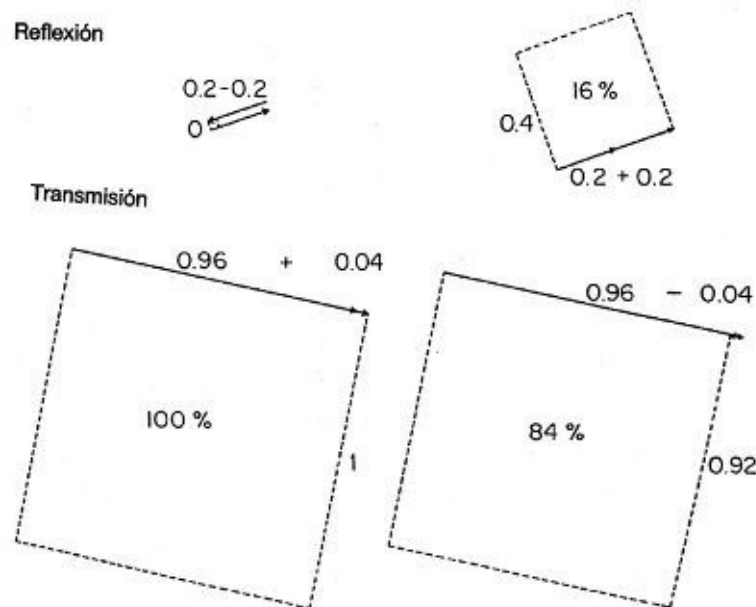


FIGURA 45. La Naturaleza siempre se asegura de que se responde del 100% de la luz. Cuando el espesor es el adecuado para que se acumulen las flechas de transmisión, las flechas de la reflexión se oponen entre sí; cuando las flechas para la reflexión se acumulan, las de la transmisión se oponen.

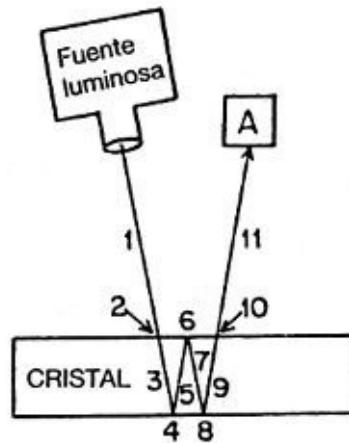


FIGURA 46. Para realizar un cálculo más apropiado se debe considerar aún otros caminos por los que la luz se puede reflejar. En esta figura, tienen lugar reducciones de 0,98 en los pasos 2 y 10, y reducciones de 0,2 en los pasos 4, 6 y 8. El resultado es una flecha con una longitud de aproximadamente 0,008, que representa otra alternativa para la reflexión y que debe por tanto añadirse a las otras flechas que representan la reflexión (0,2 para la superficie frontal y 0,192 para la superficie posterior).

Finalmente, antes de que me vaya, me gustaría decirles que hay una extensión a la regla que nos dice cuándo multiplicar flechas: las flechas se tienen que multiplicar no sólo cuando el suceso consista en una sucesión de pasos, sino también cuando el suceso consista en un número de cosas ocurriendo de modo concomitante —independiente y posiblemente de manera simultánea—. Por ejemplo, supongamos que tenemos dos fuentes, X e Y, y dos detectores, A y B (ver Fig. 47), y queremos calcular la probabilidad del siguiente suceso: el de que cada vez que X e Y emitan un fotón, A y B detecten uno cada uno.

En este ejemplo, los fotones viajan a través del espacio para ir a los detectores —no son ni reflejados ni transmitidos— luego ésta es una buena ocasión para que deje de despreciar el hecho de que la luz se dispersa cuando viaja. Ahora les voy a presentar la *regla completa* para la luz monocromática viajando de un punto a otro a través del espacio —no hay ninguna aproximación aquí y ninguna simplificación—. Esto es todo lo que se necesita conocer sobre la luz monocromática viajando a través del espacio (sin tener en cuenta la polarización): el *ángulo* de la flecha depende de la manecilla del cronógrafo imaginario, que gira un cierto número de veces por pulgada (dependiendo del color del fotón); la *longitud* de la flecha es inversamente proporcional a la distancia que alcanza la luz —en otras palabras, la flecha se reduce según avanza la luz^[9].

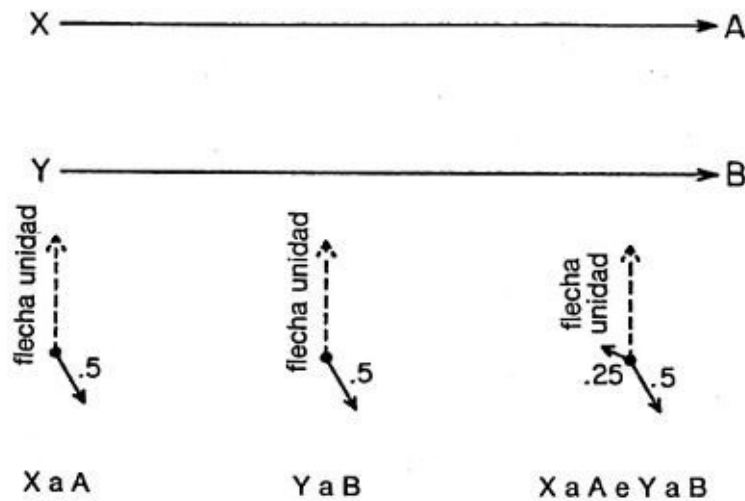


FIGURA 47. Si uno de los caminos por el que un suceso determinado puede tener lugar depende de un número de cosas que ocurren independientemente, la amplitud de este camino se calcula multiplicando las flechas de cada cosa independiente. En este caso el suceso final es: después de que las fuentes X e Y pierden cada una un fotón, los fotomultiplicadores A y B hacen un click. Un camino por el que puede tener lugar este suceso es el de un fotón que vaya de X a A y otro fotón que vaya de Y a B (dos cosas independientes). Para calcular la probabilidad de este «primer camino» las flechas de cada cosa independiente —X a A e Y a B— se multiplican para obtener la amplitud de este camino particular (El análisis continúa en la Fig. 48).

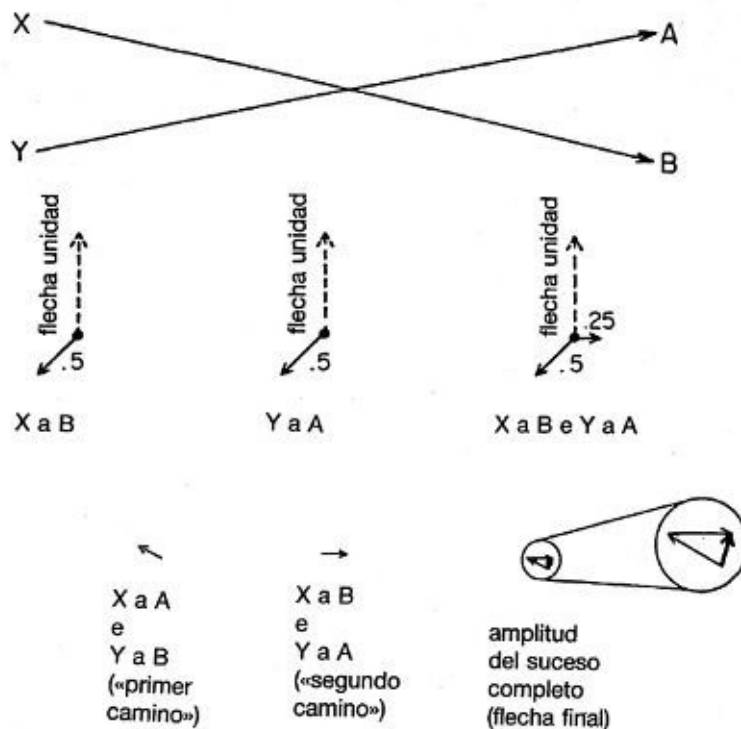


FIGURA 48. El otro camino por el que el suceso descrito en la Figura 47 puede suceder —un fotón que vaya de X a B y otro de Y a A— también depende de que tengan lugar dos cosas independientes, por tanto la amplitud para este «segundo camino» se calcula también multiplicando las flechas de las cosas independientes. Finalmente, se suman las flechas del «primer camino» y del «segundo camino», resultando la flecha final del suceso. La probabilidad de un suceso se representa siempre por una flecha final única —independientemente de cuantas flechas se dibujen, multipliquen o sumen para conseguirlo.

Supongamos que la flecha de X a A tiene una longitud de 0,5 y señala hacia las 5 en punto, lo mismo que la flecha de Y a B (Fig. 47). Multiplicando las flechas entre sí obtenemos una flecha final de longitud 0,25 señalando hacia las 10 en punto.

¡Pero esperen! Este suceso puede ocurrir de otra forma: el fotón de X puede ir a B y el fotón de Y a A. Cada uno de estos subsucesos tiene una amplitud y estas flechas también se deben dibujar y multiplicar para producir una amplitud para este camino particular (ver Fig. 48). Puesto que la cantidad de reducción con respecto a la distancia es mucho menor que el valor del giro, las flechas de X a B y de Y a A tienen esencialmente la misma longitud que las otras flechas, 0,5, pero su giro es muy diferente: la manecilla del cronógrafo gira 36 000 veces por pulgada para la luz roja, de modo que una minúscula diferencia en distancia resulta en una substancial diferencia en cronometraje.

Las amplitudes, para cada camino por el que el suceso puede tener lugar, se suman para obtener la flecha final. Puesto que sus longitudes son esencialmente las mismas, es posible que las flechas se cancelen si sus sentidos son opuestos. Las direcciones relativas de las dos flechas pueden alternarse cambiando la distancia entre las fuentes o los detectores: simplemente alejando o acercando un poquito los detectores se puede hacer que la probabilidad del suceso se amplifique o se anule, de forma análoga al caso de la reflexión parcial por dos superficies^[10].

En este ejemplo, las flechas se multiplicaron y luego se sumaron para producir una flecha final (la amplitud del suceso), cuyo cuadrado es la probabilidad del suceso. Hay que resaltar que independientemente del número de flechas que dibujemos, sumemos o multipliquemos, nuestro objetivo es calcular una *única flecha final del suceso*. A menudo, los estudiantes de física cometen errores al principio porque no tienen en consideración este punto tan importante. Trabajan tanto tiempo analizando sucesos que implican un único fotón, que empiezan a pensar que la flecha está de alguna manera asociada con el fotón. Pero estas flechas son amplitudes de probabilidad, que dan, cuando se elevan al cuadrado, la probabilidad de un suceso completo^[11].

En la próxima conferencia iniciaré el proceso de simplificación y explicación de las propiedades de la materia —a explicar de dónde proviene la reducción de 0,2, el por qué la luz parece ir más lenta a través del cristal o del agua que a través del aire, etc.— porque hasta ahora les he estado engañando: los fotones en realidad no se reflejan en la superficie del cristal;

interaccionan con los electrones *dentro* del cristal. Les mostraré cómo los fotones no hacen otra cosa que ir de un electrón a otro, y cómo la reflexión y la transmisión son realmente el resultado de que un electrón se apodere de un fotón «le rasque la cabeza» por así decir, y emita un *nuevo* fotón. Esta simplificación de todo lo que hemos estado hablando hasta ahora es muy bonita.

Capítulo 3

LOS ELECTRONES Y SUS INTERACCIONES

Esta es la tercera de cuatro conferencias sobre un tema bastante difícil — la teoría de la electrodinámica cuántica— y puesto que obviamente hay aquí más gente esta noche de la que hubo anteriormente, algunos de ustedes no han escuchado las otras dos conferencias y van a encontrar ésta casi incomprensible. Aquellos de ustedes que *hayan* escuchado las otras dos conferencias también encontrarán ésta incomprensible, pero saben que todo está bien: como les expliqué en la primera conferencia, la manera que tenemos de describir la Naturaleza es en general incomprensible para nosotros.

En estas conferencias quiero hablarles de la parte de la física que conocemos mejor, la interacción de la luz y los electrones. La mayoría de los fenómenos que les son familiares tratan de la interacción de la luz y los electrones —toda la química y la biología, por ejemplo—. Los únicos fenómenos que esta teoría no abarca son los de la gravitación y los fenómenos nucleares; todo lo demás está contenido en ella.

Encontramos en nuestra primera conferencia que no disponemos de mecanismos satisfactorios para describir incluso el más sencillo de los fenómenos, como es la reflexión parcial de la luz por el cristal. Tampoco tenemos forma de predecir si un fotón dado será reflejado o transmitido por el cristal. Todo lo que podemos hacer es calcular la *probabilidad* de que un suceso particular tenga lugar —si la luz será reflejada en este caso—. (Esta probabilidad es del orden del 4% cuando la luz incide en ángulo recto sobre una superficie de cristal, la probabilidad de reflexión aumenta cuando la luz incide sobre el cristal con mayor inclinación).

Cuando tratamos con probabilidades bajo circunstancias *ordinarias*, existen las siguientes «reglas de composición»: 1) si algo tiene lugar por *caminos alternativos*, *sumamos* las probabilidades de cada camino; 2) si el suceso ocurre como una *sucesión de pasos* —o depende de un número de cosas que ocurren de manera «concomitante» (independientemente)— *multiplicamos* las probabilidades de cada paso (o cosa).

En el mundo salvaje y maravilloso de la física cuántica, las probabilidades se calculan como *el cuadrado de la longitud de una flecha*: donde en condiciones normales hubiésemos esperado sumar probabilidades nos encontramos «sumando» *flechas*; donde normalmente hubiésemos multiplicado las probabilidades, «multiplicamos» *flechas*. Los resultados peculiares que obtenemos calculando probabilidades de esta manera encajan perfectamente con los resultados de los experimentos. Me encanta el que debamos recurrir a reglas tan peculiares y razonamientos tan extraños para comprender la Naturaleza y disfruto diciéndoselo a la gente. No hay «ruedas y engranajes» detrás de este análisis de la Naturaleza; si quieren comprenderlo éste es el camino que deben tomar.

Antes de entrar en el tema central de esta conferencia, me gustaría mostrarles otro ejemplo de cómo se comporta la luz. De lo que quiero hablarles es de luz muy débil de un color —un fotón cada vez— que va desde una fuente, en S, a un detector, en D (ver Fig. 49). Pongamos una pantalla entre la fuente y el detector y hagamos dos pequeñísimos agujeros, en A y B, separados unos milímetros entre sí. (Si la fuente y el detector están separados 100 cm, los agujeros deben de ser más pequeños que una décima de milímetro). Alineemos A con S y D, y situemos a B en algún sitio cercano a A desalineado de S y D.

Cuando cerramos el agujero en B, obtenemos un cierto número de clicks en D —lo que representa los fotones que llegan a través de A (digamos que el detector hace un click una media de una vez por cada 100 fotones que salen de S, o un 1%)—. Cuando cerramos el agujero de A y abrimos el de B, sabemos por la segunda conferencia que obtenemos casi el mismo número de clicks, en promedio, porque los agujeros son muy pequeños. (Cuando «forzamos» demasiado a la luz, las reglas del mundo ordinario —tal como que la luz viaja en línea recta— no funcionan). Cuando abrimos ambos agujeros obtenemos una respuesta complicada porque está presente la interferencia: Si los agujeros están separados una cierta distancia, obtenemos más clicks que el esperado 2% (el máximo es del orden del 4%); si los dos

agujeros están a una distancia ligeramente mayor, no obtenemos ningún click.

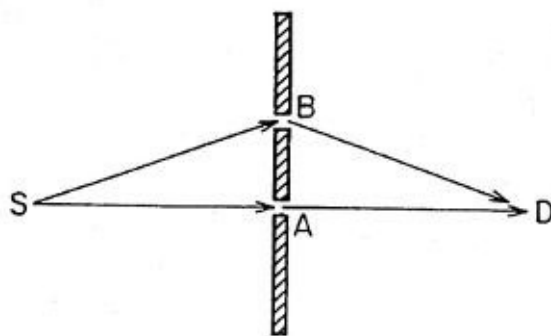


FIGURA 49. Dos minúsculos agujeros (en A y B) en una pantalla que se encuentra entre una fuente S y detector D dejan pasar casi la misma cantidad de luz a su través (en este caso el 1%) cuando se abre uno o el otro. Cuando se abren ambos simultáneamente tiene lugar la «interferencia»: el detector emite clicks desde cero al 4% de las veces, dependiendo de la separación entre A y B—como se muestra en la Figura 51 (a).

Normalmente uno pensaría que la apertura de un segundo agujero *siempre* aumentaría la cantidad de luz que llega al detector, pero esto no es lo que ocurre realmente. Y así, decir que la luz va «por un camino o por otro» es falso. Todavía me recuerdo diciendo, «bien, va por este camino o por aquél» pero cuando digo ésto, tengo que tener en cuenta lo que quiero decir por sumar amplitudes: el fotón tiene una amplitud de ir por un camino, y otra amplitud de ir por otro camino. Si las amplitudes se oponen, la luz no llegará al punto —aunque, como en este caso, ambos agujeros estén abiertos.

Bien, he aquí un nuevo rizo a ese ya extraño comportamiento de la Naturaleza que me gustaría explicarles. Supongamos que situamos unos detectores especiales —uno en A y otro en B (es posible diseñar un detector que pueda decir si el fotón pasó a través de él)— de manera que podamos decir por qué agujero(s) pasa el fotón cuando ambos agujeros están abiertos (ver Fig. 50). Puesto que la probabilidad de que un único fotón vaya de S a D, sólo se ve afectada por la distancia entre agujeros, debe existir alguna forma solapada de que el fotón se divida en dos y luego se unifique de nuevo, ¿verdad? De acuerdo con esta hipótesis, los detectores de A y B deberían dispararse siempre simultáneamente (¿con la mitad de la fuerza quizás?), mientras que el detector en D se dispararía con una probabilidad entre cero y el 4% dependiendo de la distancia entre A y B. Lo que realmente ocurre es esto: los detectores en A y B *nunca* se disparan juntos —o lo hace B o lo hace A—. El fotón no se divide en dos; va por un camino o por otro.

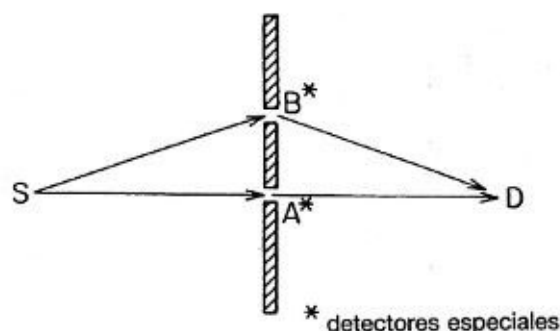


FIGURA 50. Al colocar detectores especiales en A y B para establecer por qué camino ha ido la luz cuando ambos agujeros están abiertos, se ha modificado el experimento. Puesto que un fotón siempre va por un agujero o por otro (cuando se está comprobando los agujeros), existen dos condiciones finales distinguibles: 1) los detectores en A y D se disparan, y 2) los detectores en B y D se disparan. La probabilidad de que ocurra cualquiera de estos sucesos es del 1%. Las probabilidades de los dos sucesos se suman de la manera normal, lo que proporciona un 2% de probabilidad de que el detector en D se dispare —como se muestra en la Figura 51 (b).

Además, bajo estas condiciones, el detector en D se dispara un 2% de las veces —la simple suma de las probabilidades de A y B ($1\% + 1\%$)—. El 2% no se ve afectado por el espaciado entre A y B, ¡la interferencia *desaparece* cuando se colocan detectores en A y B!

La Naturaleza se lo ha preparado tan bien que nunca seremos capaces de saber cómo lo hace: si colocamos instrumentos para descubrir por qué camino va la luz, podemos descubrirlo, perfecto, pero la maravillosa interferencia desaparece. Pero si no tenemos instrumentos que puedan decirnos el camino que lleva la luz, ¡los efectos de interferencia reaparecen! ¡Muy extraño, verdaderamente!

Para entender esta paradoja, déjenme recordarles el principio más importante: para calcular correctamente la probabilidad de un suceso, se debe tener mucho cuidado de *definir con claridad el experimento completo* —en particular, cuáles son las condiciones iniciales y finales del experimento—. Mire al dispositivo antes y después del experimento y busque cambios. Cuando estábamos calculando la probabilidad de que un fotón fuese de S a D sin detectores en A o B, el suceso era, sencillamente, el detector en D hace un click. Cuando el click en D era el único cambio en las condiciones, no había forma de decir por qué camino había ido el fotón y en consecuencia había interferencia.

Cuando colocamos detectores en A y B, cambiamos el problema. Ahora, resulta que hay *dos* sucesos completos —dos conjuntos de condiciones finales— que son distinguibles: 1) los detectores en A y D se disparan, o 2) los detectores en B y D se disparan. Cuando existe un número de condiciones

finales posibles en un experimento, debemos calcular la probabilidad de cada una como un suceso completo y separado.

Para calcular la amplitud de que los detectores en A y D se disparen, multiplicamos las flechas que representan los siguientes pasos: el fotón va de S a A, el fotón va de A a D, el detector en D se dispara. El cuadrado de la flecha final es la probabilidad del suceso —1%— la misma que cuando el agujero B estaba cerrado, porque ambos casos tienen exactamente los mismos pasos. El otro suceso completo es el de los detectores en B y D disparándose. La probabilidad de este suceso se calcula de manera similar y es la misma de antes —alrededor del 1%.

Si queremos saber con qué frecuencia el detector en D se dispara, y no nos importa si fue A o B el que se disparó en el proceso, la probabilidad es simplemente la suma de los dos sucesos —2%—. En principio, si existe algo en el sistema que *pudiésemos haber* observado para decir por qué camino fue el fotón, tendríamos «estados finales» diferentes (condiciones finales distinguibles) y sumaríamos las *probabilidades* —no las amplitudes— de cada estado final^[12].

Les he señalado estas cosas porque cuanto mejor se ve cuán extrañamente se comporta la Naturaleza, más difícil es hacer un modelo que explique cómo ocurre realmente el fenómeno más sencillo. En consecuencia, la física teórica ha abandonado el hacerlo. Vimos en la primera conferencia cómo un suceso se podía dividir en caminos alternativos y cómo se podía «sumar» la flecha final de cada camino. En la segunda conferencia, oímos cómo cada camino podía dividirse en pasos sucesivos, cómo la flecha de cada caso podía ser considerada como la transformación de la flecha unidad, y cómo las flechas de cada paso podían «multiplicarse» mediante sucesivas reducciones y giros. Estamos por tanto familiarizados con todas las reglas necesarias para dibujar y combinar flechas (que representan trocitos del suceso), que permiten obtener una flecha final cuyo cuadrado es la probabilidad de un suceso físico observable.

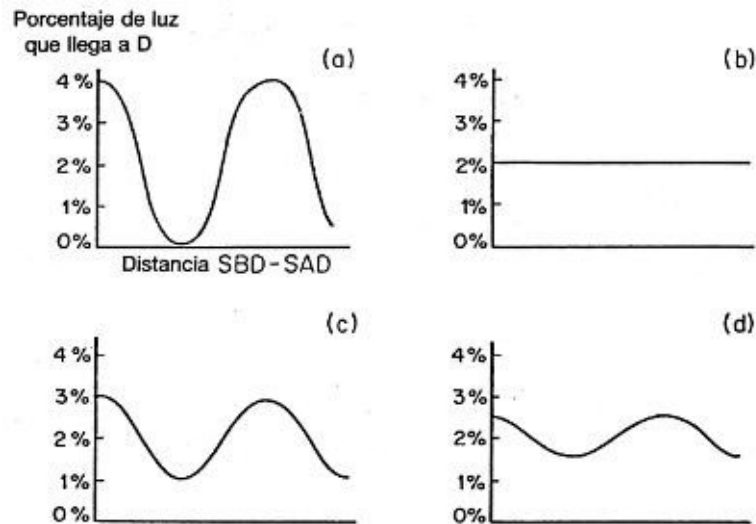


FIGURA 51. Cuando no existen detectores en A o B, existe interferencia —la cantidad de luz varía de cero al 4% (a)—. Cuando hay detectores en A y B que son fiables al 100%, no existe interferencia —la cantidad de luz que alcanza D es una constante, el 2% (b)—. Cuando los detectores en A y B no son fiables al 100% (es decir, cuando a veces no hay nada colocado en A o B que pueda detectarse), existen tres posibles condiciones finales —A y D se disparan, B y D se disparan, y D se dispara solo—. La curva final es una mezcla, formada de contribuciones de cada una de las posibilidades finales. Cuando los detectores en A y B son menos fiables, hay mayor interferencia presente. Así los detectores en el caso (c) son menos fiables que en el caso (d). El principio referente a la interferencia es: La probabilidad de cada una de las distintas condiciones finales posibles debe de ser calculada de manera independiente sumando flechas y elevando al cuadrado la longitud de la flecha final; luego, las diversas probabilidades se suman de la manera normal.

Es natural preguntarse hasta qué punto podemos continuar con este proceso de desdoblamiento de sucesos en subsucesos más y más sencillos. ¿Cuáles son los trocitos más pequeños posibles de los sucesos? ¿Existe un número finito de trocitos que puedan combinarse para formar *todos* los fenómenos que involucran a la luz y los electrones? ¿Existe un número finito de «letras» en este lenguaje de la electrodinámica cuántica que puedan combinarse para formar las «palabras» y «frases» que describen casi cada fenómeno de la Naturaleza?

La respuesta es sí, el número es tres. Sólo existen tres acciones básicas necesarias para producir todos los fenómenos asociados con la luz y los electrones.

Antes de decirles cuáles son estas tres acciones básicas, debería presentarles de manera adecuada a los actores. Los actores son los fotones y los electrones. Los fotones, las partículas de luz, ya han sido tratados con largueza en las dos primeras conferencias. Los electrones fueron descubiertos, en 1895, como partículas: se puede colocar uno en una gota de aceite y medir su carga eléctrica. Gradualmente se hizo obvio que el movimiento de estas

partículas proporcionaba una explicación a la electricidad de los alambres.

Poco después del descubrimiento del electrón se pensó que los átomos eran como pequeños sistemas solares formados por una parte central pesada (llamada el núcleo) y los electrones, que daban vueltas en «órbitas» a la manera en que lo hacen los planetas alrededor del Sol. Si piensan que ésta es la forma en que están constituidos los átomos, están ustedes en 1910. En 1924, Louis de Broglie encontró que existía un carácter ondulatorio asociado a los electrones y, poco después, C. J. Davisson y L. H. Germer de los Laboratorios Bell, bombardearon un cristal de níquel con electrones y demostraron que también éstos se reflejaban con unos ángulos locos (lo mismo que hacían los rayos X) y que estos ángulos podían calcularse a partir de la fórmula de De Broglie para la longitud de onda de un electrón.

Cuando consideramos los fotones a escala macroscópica —mucho mayor que la requerida para un giro del cronógrafo— los fenómenos que observamos se pueden aproximar muy bien mediante reglas tales como que «la luz viaja en línea recta» porque existen caminos suficientes alrededor del camino de tiempo mínimo como para reforzarlo, y caminos suficientes como para anular cualquier otro. Pero cuando el espacio por el que se mueve un fotón se hace muy pequeño (como el de los minúsculos agujeros de una pantalla) estas reglas fallan —descubrimos que la luz no viaja en línea recta, que se crean interferencias por los dos agujeros y así sucesivamente—. La misma situación se presenta para los electrones: cuando se consideran a escala macroscópica, viajan como partículas, por caminos definidos. Pero a pequeña escala, como en el interior de un átomo, el espacio es tan pequeño que no hay camino principal, no hay «órbita»; existen muchos caminos por los que el electrón puede ir, cada uno con una amplitud. El fenómeno de la interferencia se hace muy importante, y tenemos que sumar flechas para predecir hacia dónde es probable que vaya el electrón.

Es muy interesante notar que los electrones aparecieron primero como partículas y que su carácter ondulatorio se descubriese posteriormente. Por otro lado, exceptuando a Newton que se equivocó y pensó que la luz era «corpuscular», la luz parecía al principio ser ondas y sus características como partículas aparecieron posteriormente. De hecho, ambos objetos se comportan a veces como ondas y a veces como partículas. A fin de librarnos de tener que inventar palabras nuevas como «ondículas» hemos decidido llamar a estos objetos «partículas», pero todos sabemos que obedecen las reglas de dibujo y combinación de flechas que ya les he explicado. Parece que *todas* las

«partículas» de la Naturaleza —quarks, gluones, neutrinos y demás (que trataremos en la próxima conferencia)— se comportan de esta manera mecanocuántica.

Por consiguiente, ahora les presento las tres acciones básicas, a partir de las cuales se obtienen todos los fenómenos de la luz y los electrones.

- ACCIÓN N.º 1: Un fotón va de un sitio a otro.
- ACCIÓN N.º 2: Un electrón va de un sitio a otro.
- ACCIÓN N.º 3: Un electrón emite o absorbe un fotón.

Cada una de estas acciones tiene una amplitud —una flecha— que se puede calcular siguiendo ciertas reglas. En un momento les diré estas reglas, o leyes, a partir de las cuales podemos construir el mundo entero (¡exceptuando el núcleo y la gravitación, como siempre!).

Bien, el escenario en donde tienen lugar estas acciones no es el espacio, es el espacio y el tiempo. Hasta ahora, había despreciado los problemas concernientes al tiempo, tales como cuándo exactamente sale el fotón de la fuente y llega al detector. Aunque el espacio es realmente tridimensional, voy a reducirlo a una dimensión en las gráficas que voy a dibujar. Pondré la situación particular de un objeto en el espacio en el eje horizontal, y el tiempo en el vertical.

El primer suceso que voy a representar en el espacio y el tiempo —o espacio-tiempo, como puedo denominarlo inadvertidamente— es una pelota de béisbol que está quieta (ver Fig. 52). El jueves por la mañana, que denominaré T_0 , la pelota de béisbol ocupa un cierto espacio que denominaré X_0 . Un poco después, en T_1 , ocupa el mismo espacio porque está quieta. Momentos después, en T_2 , la pelota todavía está en X_0 . Por tanto, el diagrama de la pelota de béisbol en reposo, es una banda vertical, en sentido ascendente, con la pelota ocupando todo su interior.

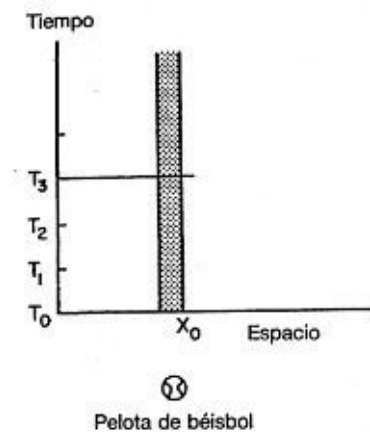


FIGURA 52. El escenario en donde tienen lugar todas las acciones del universo es el espacio-tiempo. Consistiendo normalmente de cuatro dimensiones (tres para el espacio y una para el tiempo), el espacio-tiempo se representará aquí en dos dimensiones —una para el espacio, en la dimensión horizontal, y otra para el tiempo, en la vertical—. Cada vez que miremos la pelota de béisbol (por ejemplo en el instante T_3) estará en el mismo lugar. Esto crea una «banda de pelota» en sentido ascendente al avanzar el tiempo.

¿Qué ocurre si tenemos una pelota de béisbol, flotando en el espacio ingrávido exterior, dirigiéndose hacia una pared? Bien, el jueves por la mañana (T_0) está en X_0 (ver Fig. 53) pero un poco más tarde no está en el mismo lugar —se ha movido un poquito, hacia X_1 —. Al continuar la pelota su deriva, crea una «banda de béisbol» inclinada sobre el diagrama espacio-tiempo. Cuando la pelota golpea la pared (que está quieta y por tanto es una banda vertical), retrocede exactamente al punto del espacio del que viene (X_0) pero con un punto de tiempo diferente T_6 .

Por lo que se refiere a la escala de tiempos, es más conveniente representar el tiempo no en segundos, sino en unidades mucho más pequeñas. Puesto que estamos tratando con fotones y electrones, que se mueven muy rápidamente, voy a representar con un ángulo de 45° a algo moviéndose con la velocidad de la luz. Por ejemplo, para una partícula moviéndose a la velocidad de la luz desde $X_1 T_1$ a $X_2 T_2$, la distancia horizontal entre X_1 y X_2 es la misma que la distancia vertical entre T_1 y T_2 (ver Fig. 54). El factor por el que se ha reducido el tiempo (para hacer que un ángulo de 45° represente a una partícula yendo a la velocidad de la luz) se llama c , y encontrarán que c está revoloteando continuamente en las fórmulas de Einstein —es el resultado de una elección desafortunada del segundo como unidad de tiempo, en lugar del tiempo que tarda la luz en viajar un metro.

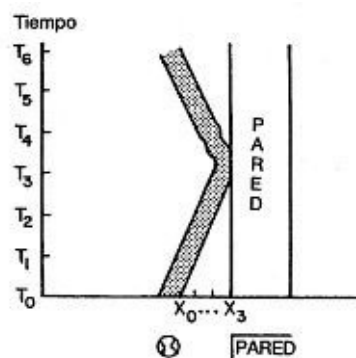


FIGURA 53. Una pelota de béisbol impelida hacia una pared en ángulo recto y luego retrocediendo hacia su situación original (mostrada debajo de la gráfica) se mueve en una dirección y aparece como una «banda de pelota» sesgada. En los instantes T_1 y T_2 , la pelota está muy próxima a la pared, en T_3 golpea la pared y comienza a retroceder.

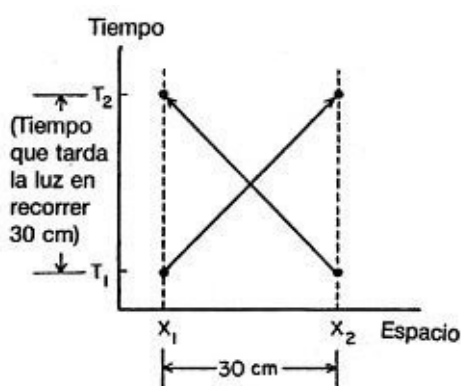


FIGURA 54. La escala de tiempo que utilizaré en estas gráficas mostrará a las partículas que se mueven con la velocidad de la luz, viajando con un ángulo de 45 grados en el espacio-tiempo. El tiempo que tarda la luz en recorrer 30 cm —desde X_1 a X_2 o desde X_2 a X_1 — es del orden de una mil-millonésima de segundo.

Ahora, consideremos detalladamente la primera acción básica —un fotón va de un sitio a otro—. Dibujaré esta acción, sin motivo alguno, con una línea serpenteante de A a B. Debería tener más cuidado: debería decir, un fotón que se sabe que está en un sitio dado en un momento dado tiene una cierta amplitud de llegar a otro lugar en otro momento. En mi gráfica espacio-tiempo (ver Fig. 55), el fotón en el punto A —en X_1 y T_1 — tiene una amplitud de llegar al punto B —en X_2 y T_2 —. A esta amplitud la llamaré $P(A \text{ a } B)$.

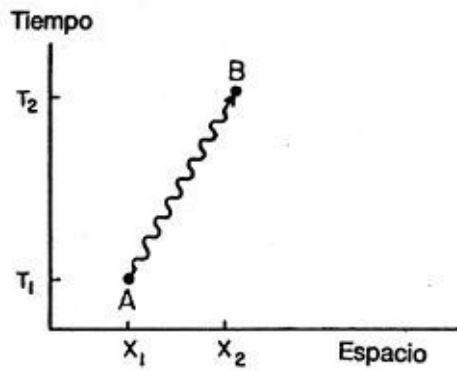


FIGURA 55. Un fotón (representado por una línea ondulante) tiene una cierta amplitud de ir de un punto A del espacio-tiempo a otro B. Esta amplitud, que denominaré $P(A \text{ a } B)$, se calcula a partir de una fórmula que depende sólo de la diferencia de posiciones $-(X_2 - X_1)-$ y de la diferencia de tiempos $-(T_2 - T_1)-$. De hecho, es una función sencilla que es la inversa de la diferencia de sus cuadrados —un «intervalo», I , que puede escribirse como $(X_2 - X_1)$ y $(T_2 - T_1)$.

Existe una fórmula para el tamaño de esta flecha $P(A \text{ a } B)$. Es una de las grandes leyes de la Naturaleza y es muy sencilla. Depende de la diferencia en *distancia* y de la diferencia en *tiempo* de estos puntos. Estas diferencias pueden expresarse matemáticamente^[13] por $(X_2 - X_1)$ y $(T_2 - T_1)$.

La contribución mayor a $P(A \text{ a } B)$ tiene lugar a la velocidad convencional de la luz —cuando $(X_2 - X_1)$ es igual a $(T_2 - T_1)$ — que es donde se espera que ocurriese, pero también hay una amplitud para la luz que va más deprisa (o más despacio) que la velocidad convencional de la luz. Han visto en estas conferencias que la luz no viaja sólo en línea recta, ahora, encontrarán ¡que no sólo va a la velocidad de la luz!

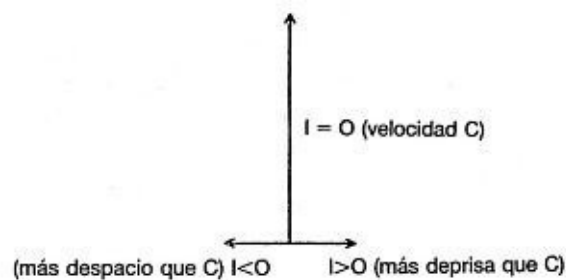


FIGURA 56. Cuando la luz viaja a la velocidad C , el «intervalo», I , se hace cero y existe una contribución grande en la dirección de las 12 en punto. Cuando I es mayor que cero, hay una pequeña contribución, en la dirección de las tres en punto, inversamente proporcional a I ; cuando I es inferior a cero y existe una contribución similar en la dirección de las nueve en punto. Por consiguiente, la luz tiene una amplitud para viajar más deprisa o más despacio que a la velocidad C , pero estas amplitudes se cancelan a largas distancias.

Puede resultarles sorprendente que exista una amplitud para un fotón que vaya a velocidades superiores o inferiores a la de la velocidad convencional c . Las amplitudes de estas posibilidades son muy pequeñas comparadas con la

contribución de la velocidad c , de hecho se anulan cuando la luz viaja largas distancias. Sin embargo, cuando las distancias son cortas —como en muchos de los diagramas que dibujaré— estas otras posibilidades son de importancia vital y deben considerarse.

Por consiguiente, ésta es la primera acción básica, la primera ley básica de la física —el fotón va de un punto a otro—. Esto explica toda la óptica: es decir, ¡la teoría completa de la luz! Bien, no del todo: he dejado fuera la polarización (como siempre) y la interacción de la luz con la materia, lo que me lleva a la segunda ley.

La segunda acción fundamental para la electrodinámica cuántica es: Un electrón va desde el punto A al punto B en el espacio-tiempo. (Por el momento, imaginaremos a este electrón como un electrón falso, simplificado, sin polarización —lo que los físicos denominamos un electrón de «espín cero»—. En realidad, los electrones tienen un tipo de polarización, que no añade nada nuevo a las ideas esenciales, sólo complica algo las fórmulas). La fórmula para la amplitud de esta acción, que denominaré $E(A \text{ a } B)$ también depende de $(X_2 - X_1)$ y $(T_2 - T_1)$ (combinadas de la misma manera que se ha descrito en la nota 13), así como de un número que denominaré « n », número que una vez determinado, permite que todos nuestros cálculos concuerden con los experimentos. (Veremos más adelante cómo determinamos el valor de n). Es una fórmula bastante complicada y lamento no saber cómo explicarla en términos más sencillos. Sin embargo, les interesará saber que la fórmula para $P(A \text{ a } B)$ —un fotón viajando de un sitio a otro en el espacio-tiempo— es la misma que para $E(A \text{ a } B)$ —un electrón viajando de un sitio a otro— si n se hace igual a cero^[14].

La tercera acción básica es: un electrón emite o absorbe un fotón —no importa la diferencia—. Denominaré a esta acción una «unión» o «acoplamiento». Para distinguir los electrones de los fotones en los diagramas, dibujaré cada electrón que viaja por el espacio-tiempo por una línea recta. Cada acoplamiento, en consecuencia, es una unión entre dos líneas rectas y una línea ondulante (ver Fig. 58). No existe una fórmula complicada para la amplitud de un electrón que emita o absorba un fotón; no depende de nada —¡es sólo un número!—. Al número de la unión lo denominaré j —su valor es aproximadamente $-0,1$: una reducción de aproximadamente una décima parte y un giro de media vuelta^[15].

Bien, esto es todo lo que concierne a estas acciones básicas —exceptuando alguna ligera complicación debida a la polarización que siempre

hemos dejado de lado—. Nuestro siguiente trabajo es reunir estas tres acciones para representar circunstancias de algún modo más complejas.

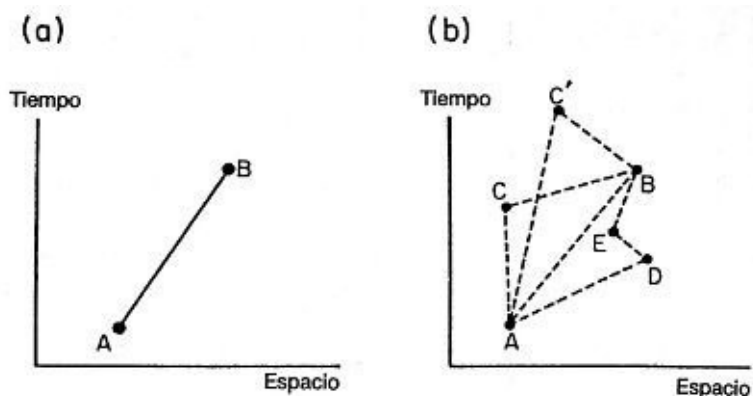


FIGURA 57. Un electrón tiene una amplitud para ir de un punto a otro en el espacio-tiempo, que denominaré « $E(A \text{ a } B)$ ». Aunque representaré $E(A \text{ a } B)$ como una línea recta entre dos puntos (a), podemos considerarla como la suma de muchas amplitudes (b) —entre ellas, la amplitud de que un electrón cambie su dirección en los puntos C o C' en un camino de «dos saltos», la amplitud de que cambie de dirección en D y E en un camino a «tres saltos»— además del camino directo de A a B. El número de veces que un electrón puede cambiar de dirección va desde cero a infinito, y los puntos en los que el electrón puede cambiar de dirección en su camino de A a B en el espacio-tiempo son infinitos. Todo está incluido en $E(A \text{ a } B)$.

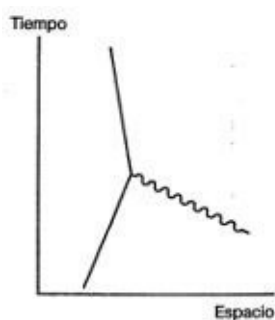


FIGURA 58. Un electrón, representado por una línea recta, tiene una cierta amplitud de emitir o absorber un fotón, representado por una línea ondulante. Puesto que la amplitud de emisión o absorción es la misma, denominaré a ambos casos un «acoplamiento». La amplitud de un acoplamiento es un número que denominaré j ; es de aproximadamente $-0,1$ para el electrón (a este número se le denomina en ocasiones la «carga»).

Como primer ejemplo, calculemos la probabilidad de que dos electrones, situados en los puntos 1 y 2 del espacio-tiempo, acaben en los puntos 3 y 4 (ver Fig. 59). Este suceso puede tener lugar por distintos caminos. El primer camino es el del electrón en 1 yendo a 3 —calculado poniendo 1 y 3 en la fórmula $E(A \text{ a } B)$ que escribiré como $E(1 \text{ a } 3)$ — y el electrón en 2 yendo a 4 —calculado como $E(2 \text{ a } 4)$ —. Estos son dos «subsucos» que ocurren concomitantemente, por consiguiente, las dos flechas se multiplican para dar una flecha para este primer camino, por el que puede tener lugar el suceso. Consecuentemente, escribimos la fórmula para la «flecha del primer camino» como $E(1 \text{ a } 3) \times E(2 \text{ a } 4)$.

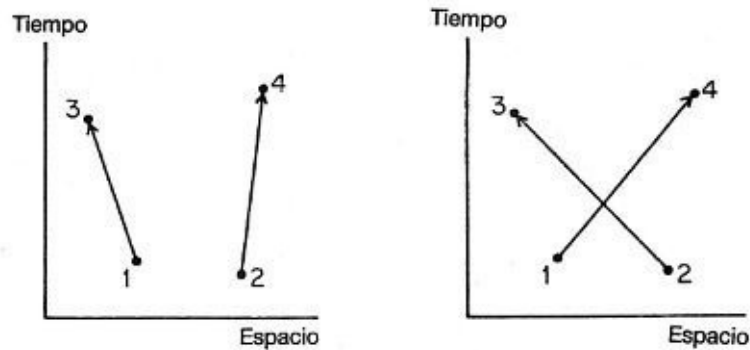


FIGURA 59. Para calcular la probabilidad de que los electrones de los puntos 1 y 2 del espacio-tiempo acaben en los puntos 3 y 4, calculamos la flecha del «primer camino», 1 yendo a 3 y 2 yendo a 4, con la fórmula $E(A \text{ a } B)$; luego calculamos la flecha del «segundo camino», 1 yendo a 4 y 2 yendo a 3 (un «cruce»). Finalmente, sumamos las flechas del «primer camino» y del «segundo camino» para obtener una buena aproximación de la flecha final. (Esto es verdad para el electrón ficticio simplificado de «espín cero». Si hubiésemos incluido la polarización del electrón, deberíamos haber restado —en lugar de sumado— las dos flechas).

Otro camino por el que puede ocurrir este suceso es el del electrón en 1 yendo a 4 y el de 2 yendo a 3 —de nuevo, dos subsucesos concomitantes—. La «flecha del segundo camino» es $E(1 \text{ a } 4) \times E(2 \text{ a } 3)$ y la sumamos a la flecha del «primer camino»^[16].

Esta es una buena aproximación para la amplitud de este suceso. Para realizar un cálculo más preciso, que se ajuste más a los resultados experimentales, debemos considerar otros caminos por los que puede transcurrir el suceso. Por ejemplo, para cada uno de estos dos caminos principales, podría suceder que uno de electrones fuese a descargarse a un sitio nuevo y maravilloso y emitiese un fotón (ver Fig. 60). Mientras tanto, el otro electrón podría ir a otro sitio y absorber el fotón. Calcular la amplitud del primero de estos caminos nuevos supone multiplicar las amplitudes para: un electrón que va desde 1 al sitio nuevo y maravilloso, 5 (donde emite un fotón) y que luego va desde 5 a 3; el otro electrón que va desde 2 al otro lugar, 6 (donde absorbe el fotón), y que luego va de 6 a 4. Debemos recordar incluir la amplitud de que el fotón vaya desde 5 a 6. Voy a escribir la amplitud, para este camino por el que el suceso puede tener lugar, a la manera de las matemáticas superiores, y ustedes pueden hacerlo conmigo: $E(1 \text{ a } 5) \times j \times E(5 \text{ a } 3) \times E(2 \text{ a } 6) \times j + E(6 \text{ a } 4) + P(5 \text{ a } 6)$ —montones de reducciones y giros—. (Les dejo que deduzcan la notación para el otro caso, cuando el electrón en 1 acaba en 4, y del 2 en 3)^[17].

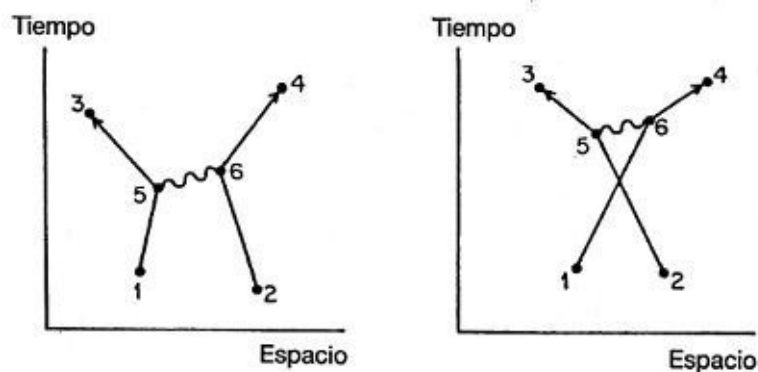


FIGURA 60. Dos caminos «camino alternativo» por los que el suceso de la Fig. 59 puede tener lugar: en cada caso se emite un fotón en 5 y se absorbe en 6. Las condiciones finales de ambas alternativas son las mismas que en los otros casos —dos electrones entran y dos electrones salen— y estos resultados son indistinguibles de los de otras alternativas. En consecuencia, las flechas para estos «camino alternativo» se deben sumar a las flechas de la Fig. 59 para obtener una mejor aproximación a la flecha final del suceso.

Pero, esperen: las posiciones 5 y 6 pueden estar en cualquier lugar en el espacio y tiempo —sí, en cualquier lugar— y las flechas de *todas* estas posiciones se tienen que calcular y sumar. Ya ven, está empezando a convertirse en mucho trabajo. No es que las reglas sean tan difíciles, es como jugar al ajedrez: las reglas son sencillas, pero se usan una y otra vez. Por tanto, nuestra dificultad en el cálculo surge de tener que agrupar tanta cantidad de flechas. Este es el motivo por el que se necesitan cuatro años de estudios universitarios para aprender a realizar esto de manera eficiente —¡y estamos considerando un problema *sencillo*! (Cuando los problemas se vuelven demasiado difíciles, ¡los metemos en un computador!).

Me gustaría señalarles algo a cerca de los fotones que se emiten y absorben: si el punto 6 es posterior al 5 podemos decir que el fotón se emitió en 5 y se absorbió en 6 (ver Fig. 61). Si el punto 6 ocurre antes que el 5, podríamos preferir decir que el fotón se emitió en 6 y se absorbió en 5, pero de la misma manera podemos decir que ¡el fotón va retrocediendo en el tiempo! Sin embargo, no tenemos que preocuparnos por el camino por el que ha ido el fotón en el espacio-tiempo; todo está incluido en la fórmula $P(5 \text{ a } 6)$, y decimos que el fotón se «ha intercambiado». ¿No es bonito ver qué simple es la Naturaleza?^[18]

Bien, además del fotón que se puede intercambiar entre 5 y 6, se puede intercambiar otro —entre dos puntos, 7 y 8 (ver Fig. 62)—. No estoy demasiado cansado para escribir todas las acciones básicas cuyas flechas tienen que multiplicarse, pero —como habrán notado— cada línea recta supone un $E(A \text{ a } B)$, cada línea ondulante un $P(A \text{ a } B)$, y cada acoplamiento

un j . Por tanto, hay seis $E(A \text{ a } B)$, dos $P(A \text{ a } B)$ y cuatro j —¡para cada posible 5, 6, 7 y 8!— ¡Esto supone miles de millones de pequeñas flechas que tienen que multiplicarse y sumarse entre sí!

Parece que calcular la amplitud de este sencillo suceso es un negocio ruinoso, pero cuando se es un estudiante universitario se debe conseguir el título, de manera que se continúa el trabajo.

Pero *hay* esperanza de éxito. Reside en el número mágico j . En los dos primeros caminos por los que puede ocurrir el suceso, no aparece j en los cálculos; el siguiente camino tiene $j \times j$ y el último que consideramos tiene $j \times j \times j \times j$. Puesto que $j \times j$ es menor que 0,01, significa que la longitud de la flecha para este camino es, en general, menor que el 1% del valor de la flecha de los dos primeros caminos; una flecha con $j \times j \times j \times j$ es menos del 1% —una parte en 10 000— comparada con las flechas sin j . Si tiene suficiente tiempo en el computador, puede calcular las posibilidades que involucran j^6 —una parte en un millón— e igualar la precisión de los experimentos. Esta es la forma en que se realizan los cálculos para los sucesos sencillos. Esta es la manera en que funciona ¡no hay más!

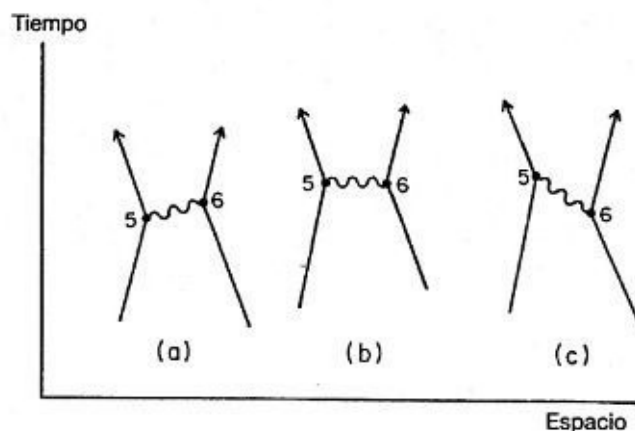


FIGURA 61. Puesto que la luz tiene una cierta amplitud de ir más rápida o más lenta que la velocidad convencional de la luz, los fotones de estos tres ejemplos pueden considerarse como emitidos en el punto 5 y absorbidos en el punto 6, incluso aunque el fotón del ejemplo (b) se haya emitido al mismo tiempo que se ha absorbido, y el fotón en (c) se haya emitido después de que se haya absorbido —una situación para la que Vds. hubiesen preferido decir que se había emitido en 6 y absorbido en 5—; de otro modo, el fotón tiene que ir ¡hacia atrás en el tiempo! Por lo que respecta a los cálculos (y a la Naturaleza) todo es lo mismo (y todo es posible), de manera que decimos sencillamente que se ha «intercambiado» un fotón e introducimos las posiciones del espacio-tiempo en la fórmula $P(A \text{ a } B)$.

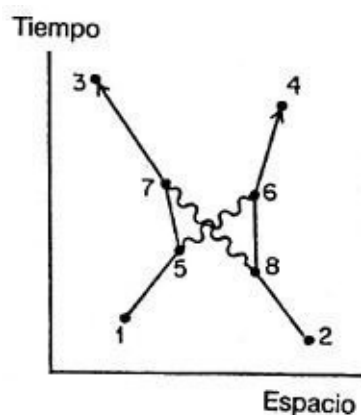


FIGURA 62. Y otro camino más por el que puede tener lugar el suceso de la Fig. 59 es mediante el intercambio de dos fotones. Para este camino son posibles muchos diagramas (como veremos con más detalle posteriormente), aquí se muestra uno de ellos. La flecha de este camino involucra todos los posibles puntos intermedios 5, 6, 7 y 8 y se calcula con gran dificultad. Dado que j es menor que 0,1, la longitud de esta flecha es, en general, menor que una parte en 10 000 (porque hay cuatro acoplamientos involucrados) comparada con las flechas del «primer camino» y «segundo camino» de la Fig. 59 que no contenían ninguna j .

Consideremos ahora otro suceso. Comenzamos con un fotón y un electrón, y acabamos con un fotón y un electrón. Un camino por el que el suceso puede tener lugar es: un fotón es absorbido por un electrón, el electrón continúa por un ratito y un nuevo fotón emerge. Este proceso se denomina difusión de la luz. Cuando hacemos los diagramas y cálculos para la difusión, debemos incluir algunas posibilidades peculiares (ver Fig. 63). Por ejemplo, el electrón puede emitir un fotón *antes* de absorber otro (b). Incluso más extraña es la posibilidad (c), que el electrón emita un fotón, que a continuación *retroceda en el tiempo* para absorber otro fotón, y que después avance de nuevo en el tiempo. El camino de un electrón «moviéndose hacia atrás» puede ser tan largo que parezca real en un experimento físico en el laboratorio. Su comportamiento está incluido en estos diagramas y en la ecuación de E(A a B).

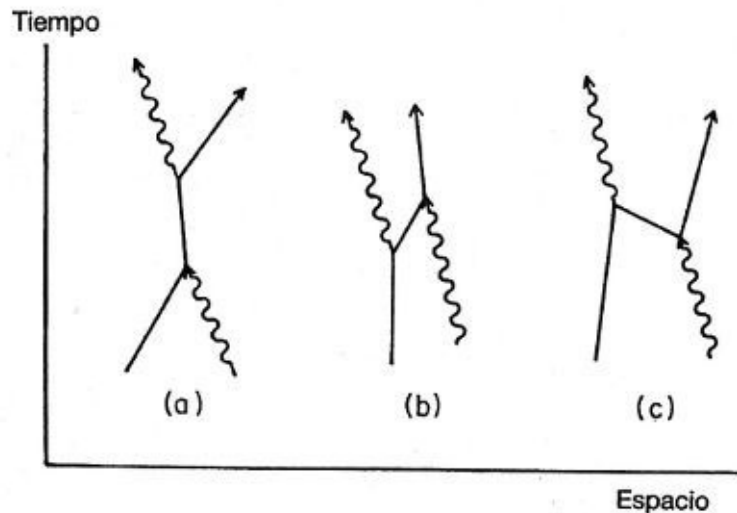


FIGURA 63. La difusión de la luz implica un fotón incidiendo sobre un electrón y emitiendo un fotón —no necesariamente en este orden—, como vemos en el ejemplo (b). El ejemplo (c) muestra una rara posibilidad aunque posible: el electrón emite un fotón, retrocede con rapidez en el tiempo para absorber un fotón, y luego continúa hacia adelante en el tiempo.

El electrón con movimiento de retroceso, cuando se le representa con el tiempo avanzado, parece igual que un electrón ordinario, excepto que se ve atraído por los electrones normales —decimos que tiene una «carga positiva»—. (Si hubiese incluido los efectos de la polarización, sería obvio el por qué para los electrones retrocediendo, el signo de j está invertido, haciendo que la carga aparezca como positiva). Por esta razón se denomina un «positrón». El positrón es una partícula hermana del electrón y es un ejemplo de una «anti-partícula»^[19].

Este fenómeno es general. Cada partícula de la Naturaleza tiene una amplitud de movimiento hacia atrás en el tiempo, y por tanto tiene una anti-partícula. Cuando una partícula y su anti-partícula colisionan, se aniquilan y forman otras partículas. (Para positrones y electrones aniquilándose, normalmente es uno o dos fotones). ¿Y que ocurre con los fotones? Los fotones aparecen exactamente igual en todos los aspectos cuando retroceden en el tiempo —como vimos anteriormente— luego ellos mismos son su anti-partícula. ¡Ya ven qué inteligentes somos haciendo de una excepción parte de la regla!

Me gustaría mostrarles qué aspecto tiene para nosotros este electrón retrocediendo, cuando nosotros avanzamos en el tiempo. Voy a dividir el diagrama en bloques de tiempo, T_0 a T_{10} (ver Fig. 64) mediante una serie de líneas paralelas para visualizar mejor. Empezamos en T_0 con un electrón moviéndose hacia un fotón, que se está moviendo en sentido opuesto. De

repente —en T_3 — el fotón se convierte en dos partículas, un positrón y un electrón. El positrón no dura mucho tiempo: pronto corre hacia el electrón — en T_5 —, donde se aniquilan y producen un nuevo fotón. Mientras tanto, el electrón, creado anteriormente por el fotón original, continúa atravesando el espacio-tiempo.

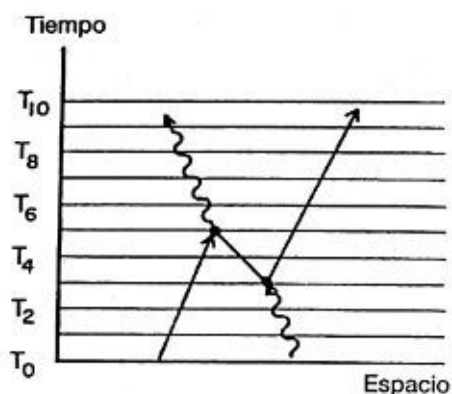


FIGURA 64. Considerando el ejemplo (c) de la Fig. 63 sólo hacia adelante en el tiempo (como nos vemos forzados en el laboratorio), de T_0 a T_1 vemos al electrón y al fotón moviéndose uno hacia el otro. De repente, en T_3 , el fotón «se desintegra» y aparecen dos partículas —un electrón y un nuevo tipo de partícula (denominada un «positrón») que es un electrón yendo hacia atrás en el tiempo y que parece moverse hacia el propio electrón inicial—. En T_5 , el positrón se aniquila con el electrón inicial para producir un nuevo fotón. Entretanto, el electrón creado por el primer fotón continúa hacia adelante en el espacio-tiempo. Esta secuencia de acontecimientos se ha observado en el laboratorio, y se incluye automáticamente en la fórmula $E(A \text{ a } B)$ sin ninguna modificación.

De lo siguiente que me gustaría hablar es de un electrón en un átomo. A fin de comprender el comportamiento de los electrones en los átomos, debemos añadir otra característica, el núcleo —la parte pesada en el centro de un átomo que contiene al menos un protón (un protón es una «Caja de Pandora» que abriremos en la siguiente conferencia)—. No les daré las leyes correctas del comportamiento del núcleo en esta conferencia; son demasiado complicadas. Pero, para el caso en que el núcleo está en reposo, podemos aproximar su comportamiento por el de una partícula con una amplitud para desplazarse de un punto a otro del espacio-tiempo, de acuerdo con la fórmula para $E(A \text{ a } B)$, pero con un valor de n mucho más elevado. Puesto que el núcleo es muy pesado en comparación con el electrón, podemos considerarlo aquí, de forma aproximada, como permaneciendo básicamente en el mismo lugar al avanzar en el tiempo.

El átomo más sencillo, llamado hidrógeno, consta de un protón y un electrón. Intercambiando fotones el protón mantiene al electrón en sus proximidades, bailando a su alrededor (ver Fig. 65)^[20]. Átomos con más de un protón y sus correspondiente número de electrones también difunden la luz

(los átomos del aire difunden la luz del Sol y hacen el cielo azul), ¡pero los diagramas para estos átomos supondrían tal cantidad de líneas rectas y onduladas que serían un completo lío!

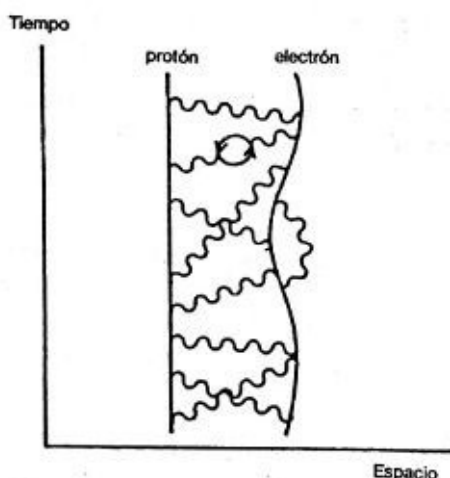


FIGURA 65. Un electrón se mantiene dentro de un cierto rango de distancia del núcleo de un átomo intercambiando fotones con un protón (una «Caja de Pandora» que consideraremos en el Capítulo 4). De momento, el protón puede aproximarse por una partícula estacionaria. Aquí se muestra un átomo de hidrógeno que consiste en un protón y un electrón que intercambian fotones.

Ahora me gustaría mostrarles un diagrama de un electrón en un átomo de hidrógeno difundiendo la luz (ver Fig. 66). Mientras que el electrón y el núcleo intercambian fotones, un fotón, que procede del exterior del átomo, incide sobre éste, golpea al electrón y es absorbido; luego, se emite un nuevo fotón. (Como siempre, existen otras posibilidades a considerar, tal como que el nuevo fotón sea emitido antes de que el viejo se vea absorbido). La amplitud total para todos los caminos por los que un electrón puede difundir un fotón se pueden resumir en una única flecha, con una cierta cantidad de reducción y giro (más adelante llamaremos a esta flecha «S»). Esta cantidad depende del núcleo y de la distribución de los electrones en los átomos, y es diferente para cada material.

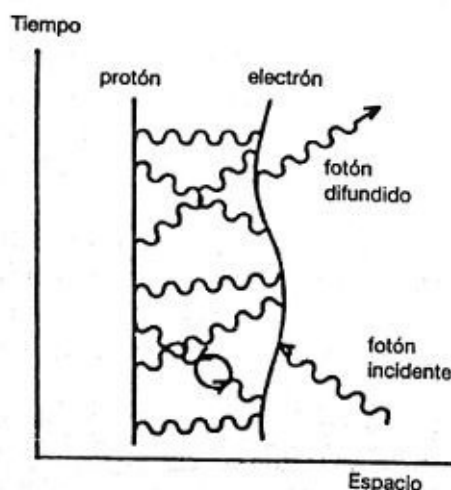


FIGURA 66. La difusión de la luz por un electrón en un átomo es el fenómeno que explica la reflexión parcial de la luz en una lámina de cristal. El diagrama muestra uno de los caminos por los que este suceso puede tener lugar en un átomo de hidrógeno.

Ahora consideremos de nuevo la reflexión parcial de la luz por una lámina de cristal. ¿Cómo tiene lugar? Les hablé de la luz siendo reflejada por la superficie frontal y posterior. Esta idea de las superficies es una simplificación que realicé a fin de mantener inicialmente las cosas a un nivel sencillo. La luz no se ve afectada en realidad por las superficies. Un fotón incidente es difundido por los electrones de los átomos de cristal, y es un *nuevo* fotón el que vuelve al detector. Es interesante que en lugar de sumar todos los miles de millones de flechas diminutas que representan la amplitud para que todos los electrones dentro del cristal difundan un fotón incidente, solamente sumemos dos flechas —una para la reflexión por la «superficie frontal» y otra para la «superficie posterior»— y obtenemos el mismo resultado. Veamos por qué.

Para discutir la reflexión por una lámina desde nuestro nuevo punto de vista, debemos de tener en cuenta la dimensión del tiempo. Anteriormente, cuando hablamos de luz de una fuente monocromática, utilizábamos un cronógrafo imaginario que cronometraba el fotón en su movimiento —la manecilla de este cronógrafo determinaba el ángulo de la amplitud de un camino dado—. En la fórmula de $P(A \text{ a } B)$ (la amplitud para que un fotón vaya de un punto a otro) no se menciona giro alguno. ¿Qué ocurrió con el cronógrafo? ¿Qué ocurrió con el giro?

En la primera conferencia dije sencillamente que la fuente luminosa era monocromática. Para analizar de modo correcto la reflexión parcial de una lámina, necesitamos saber algo sobre las fuentes monocromáticas de la luz. La amplitud de emisión de un fotón por una fuente varía, en general, con el

tiempo: al transcurrir éste, cambia el ángulo de amplitud de emisión de un fotón por la fuente. Una fuente de luz blanca —muchos colores mezclados— emite fotones caóticamente: el ángulo de la amplitud cambia de forma abrupta e irregular, sin continuidad. Pero cuando construimos una fuente *monocromática*, estamos haciendo un dispositivo cuidadosamente diseñado para que se pueda calcular fácilmente la amplitud de emisión de un fotón cada cierto tiempo: cambia su ángulo a velocidad *constante*, como una manecilla de cronógrafo. (En realidad, esta flecha gira a la misma velocidad que el cronógrafo imaginario que utilizamos antes, pero en sentido opuesto —ver Fig. 67—).

La velocidad de giro depende del color de la luz: la amplitud para una fuente de luz azul gira casi dos veces más deprisa que para una luz roja, igual que antes. Por tanto, el contador de tiempo que utilizamos para el «cronógrafo imaginario» era la fuente monocromática: en realidad, el ángulo de la amplitud para un camino dado depende del *tiempo* en que se ha emitido el fotón por la fuente.

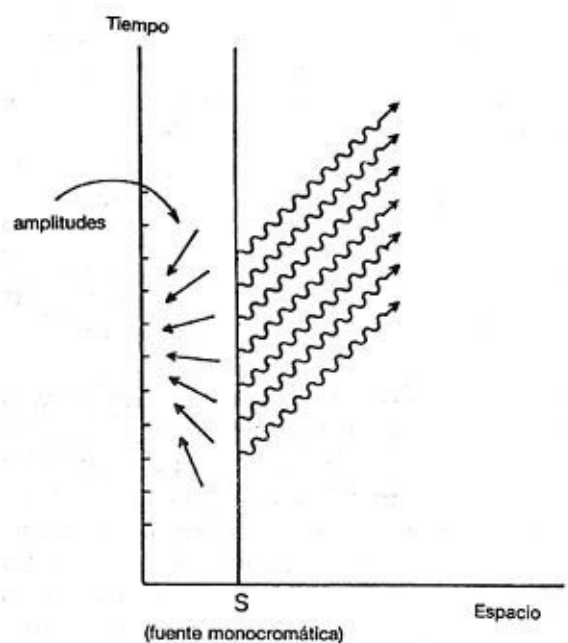


FIGURA 67. Una fuente monocromática es un aparato preciosamente construido que emite un fotón con un camino totalmente predecible: la amplitud de emisión de un fotón a un tiempo determinado gira en el sentido contrario a las agujas del reloj al avanzar el tiempo. En consecuencia, la amplitud de emisión de un fotón por la fuente en un tiempo posterior tiene menor ángulo. Se supondrá que toda la luz emitida por la fuente viaja a la velocidad de la luz c (puesto que las distancias son largas).

Una vez emitido el fotón, no hay giros de la flecha mientras va de un punto a otro del espacio-tiempo. Aunque la fórmula $P(A \text{ a } B)$ dice que hay una amplitud para la luz que va de un sitio a otro a velocidad *distinta* de c , la

distancia entre la fuente y el detector en nuestro experimento es relativamente grande (comparada con un átomo), luego la única contribución a la longitud de $P(A \text{ a } B)$ que cuenta es la que procede de la de la velocidad c .

Para comentar nuestro nuevo cálculo de la reflexión parcial, empecemos por definir completamente el suceso: el detector en A hace un click en un *tiempo* dado, T . A continuación dividamos la lámina de cristal en un número de secciones muy delgadas —digamos, seis (ver Fig. 68a)—. Del análisis que realizamos en la segunda conferencia en el que encontramos que casi toda la luz es reflejada por el centro del espejo, sabemos que aunque cada electrón difunde la luz en todas las direcciones, cuando se suman todas las flechas de cada sección, el único sitio en donde *no* se cancelan es aquel donde la luz va directa a la zona central de la sección y difunde en una de las dos direcciones —directamente de vuelta al detector o hacia abajo a través del cristal—. La flecha final del suceso vendrá determinada entonces por la suma de seis flechas representando la difusión de la luz por los seis puntos medios — X_1 a X_6 — distribuidos verticalmente a través del cristal.

De acuerdo, calculemos la flecha para cada uno de estos caminos por los que puede ir la luz —vía los seis puntos, X_1 a X_6 —. Hay cuatro pasos involucrados en cada camino (lo que significa cuatro flechas a multiplicar):

- PASO N.º 1: La fuente emite un fotón en un determinado momento.
- PASO N.º 2: El fotón va desde la fuente a uno de los puntos del cristal.
- PASO N.º 3: El fotón es difundido por un electrón en ese punto.
- PASO N.º 4: Un nuevo fotón emprende el camino hacia el detector.

Diremos que las amplitudes de los pasos 2 y 4 (el fotón va hacia o viene de un punto del cristal) no incluyen reducciones o giros, porque podemos suponer que nada de luz se pierde o se dispersa entre la fuente y el cristal o entre el cristal y el detector. Para el paso 3 (un electrón difunde un fotón) la amplitud de difusión es una constante —una reducción y un giro de un cierto valor, S — y es la misma en todas partes del cristal. (Esta cantidad es, como mencioné antes, diferente para cada material. Para el cristal, el giro de S el 90°). Por tanto, de las cuatro flechas a multiplicar, sólo la flecha del paso 1 —la amplitud para que un fotón sea emitido por la fuente cada determinado tiempo— es diferente para cada alternativa.

El instante en el que tendría que haberse emitido un fotón para llegar al detector A en el instante T (ver Fig. 68b) no es el mismo para los seis caminos

distintos. Un fotón difundido por X_2 debería haberse emitido ligeramente *antes* que un fotón difundido por X_1 , porque su camino es más largo. Por tanto la flecha en T_2 está ligeramente más girada que la flecha en T_1 , porque la amplitud para que una fuente monocromática emita un fotón en un tiempo determinado gira en el sentido de las agujas del reloj al avanzar del tiempo. Lo mismo ocurre con cada flecha hasta T_6 : las seis flechas tienen la misma longitud, pero están giradas con ángulos distintos —es decir, señalan en distintas direcciones— porque representan un fotón emitido por la fuente en tiempos distintos.

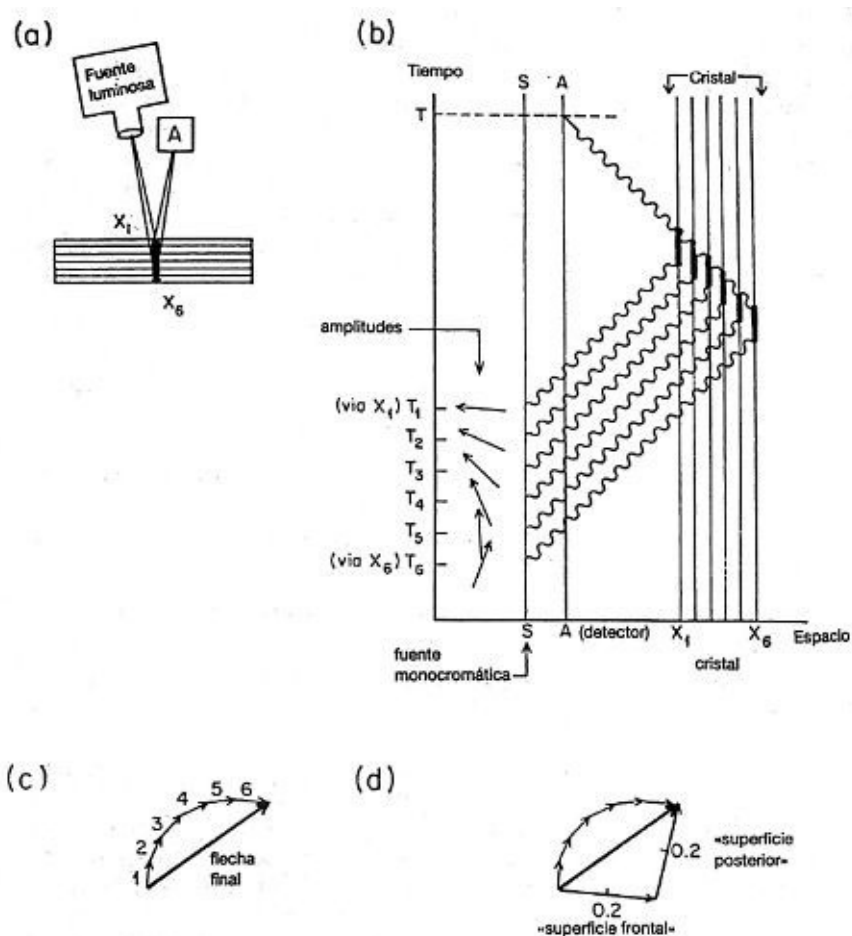


FIGURA 68. Comenzamos nuestro nuevo análisis de la reflexión parcial dividiendo una lámina de cristal en un número de secciones (aquí, seis) y considerando los distintos caminos que pueden llevar la luz desde la fuente al cristal y a su vuelta al detector en A. Los únicos puntos importantes en el cristal (donde las amplitudes de difusión de la luz no se anulan) están situados en el centro de cada sección; X_1 a X_6 se muestran en (a) en su posición física dentro del cristal, y en (b) como líneas verticales en la gráfica del espacio-tiempo. El suceso, cuya probabilidad estamos calculando, es: el detector en A emite un click cada cierto tiempo, T . Luego el suceso aparece como un punto (donde se intersectan A y T) en la gráfica del espacio-tiempo.

Para cada uno de los caminos en que puede tener lugar el suceso, tiene que haber cuatro pasos en sucesión, por lo que se tienen que multiplicar cuatro flechas. Los pasos se muestran en (b): 1) el fotón abandona la fuente en un determinado tiempo (las flechas de T_1 a T_6 representan la amplitud

de que esto ocurra en seis tiempos diferentes); 2) el fotón va desde la fuente a uno de los puntos del cristal (las seis alternativas se representan por líneas ondulantes moviéndose hacia la derecha); 3) un electrón en uno de los puntos difunde un fotón (mostrado como la línea corta vertical); y 4) un nuevo fotón va al detector y llega en el tiempo acordado, T (mostrado por una línea ondulante hacia la izquierda). Las amplitudes de los pasos 2, 3 y 4 son las mismas para las seis vías alternativas, mientras que las amplitudes para el paso 1 son diferentes: comparado con un fotón difundido por un electrón en la parte frontal del cristal (en X_1), un fotón difundido a mayor profundidad en el cristal —en X_2 , por ejemplo— debe abandonar antes la fuente, en T_2 .

Cuando hayamos terminado de multiplicar las cuatro flechas para cada alternativa, las flechas resultantes, mostradas en (c), son más cortas que las de (b); ambas han sido giradas 90° (de acuerdo con las características de difusión de los electrones en el cristal). Cuando estas seis flechas se suman en orden, forman un arco; la flecha final es su cuerda. La misma flecha final se puede obtener dibujando dos flechas radio, mostradas en (d) y «restándolas» (girando la flecha de la «superficie frontal» en dirección opuesta y sumándole a ésta la flecha de la «superficie posterior»). Este atajo se utilizó como una simplificación en la primera conferencia.

Después de reducir el tamaño de la flecha en T_1 por las cantidades prescritas en los pasos 2, 3 y 4 —y girarla los 90° prescritos en el paso 3— acabamos con la flecha 1 (ver Fig. 68c). Lo mismo ocurre con las flechas 2 y 6. Por consiguiente, las flechas 1 a 6 tienen todas la misma longitud (acortada) y están giradas entre sí la misma cantidad que las flechas T_1 a T_6 .

A continuación, sumamos las flechas 1 a 6. Al conectar las flechas ordenadamente del 1 al 6, obtenemos algo parecido a un arco, o parte de un círculo. La flecha final es la cuerda de este arco. La longitud de la flecha final aumenta con el espesor del cristal —un cristal más grueso significa más secciones, más flechas y por tanto más parte de un círculo— hasta que se alcanza un semicírculo (y la flecha final es el diámetro). Por tanto, la longitud de la flecha final *decrece* al continuar aumentando el espesor del cristal y el círculo llega a completarse para comenzar un nuevo ciclo. El cuadrado de esta longitud es la probabilidad del suceso, y varía a lo largo del ciclo desde cero al 16%.

Existe un truco matemático que se puede utilizar para conseguir la misma solución (ver Fig. 68d): Si dibujamos una flecha desde el centro del «círculo» a la cola de la flecha 1 y a la cabeza de la flecha 6, obtenemos dos radios. Si la flecha-radio del centro a la flecha 1 se gira 180° («se resta»), entonces puede combinarse con la otra flecha-radio ¡y darnos la misma flecha final! Esto es lo que hacía en la primera conferencia: estos dos radios son las dos flechas que dije representaban la reflexión por la «superficie frontal» y la «superficie posterior». Cada una tiene la famosa longitud de 0,2.^[21]

Por consiguiente, podemos conseguir la respuesta correcta para la probabilidad de reflexión parcial imaginando (falsamente) que toda la

reflexión proviene sólo de las superficies frontal y posterior. En este sencillo análisis intuitivo, las flechas de la «superficie frontal» y la «superficie posterior» son construcciones matemáticas que nos dan la respuesta correcta, mientras que el análisis que acabo de hacer —con el dibujo del espacio-tiempo y las flechas formando parte de un círculo— es una representación más aproximada de lo que realmente está ocurriendo: la reflexión parcial es la difusión de la luz por los electrones del *interior* del cristal.

Y ahora, ¿qué ocurre con la luz que atraviesa la lámina del cristal? Primero, existe una amplitud de que el fotón vaya, derecho a través del cristal, sin golpear a ningún electrón (ver Fig. 69a). Esta es la flecha más importante en términos de longitud. Pero hay otros seis caminos por los que un fotón puede alcanzar el detector colocado debajo del cristal: un fotón puede incidir en X_1 y difundir el nuevo fotón hacia B; un fotón puede incidir en X_2 y difundir el nuevo fotón hacia B, y así sucesivamente. Estas seis flechas tienen todas la misma longitud que las flechas que forman el «círculo» en el ejemplo anterior: su longitud está basada en la misma amplitud S , que la de que un electrón del cristal difunda un fotón. Pero esta vez, las seis flechas señalan en la misma dirección, porque la longitud de los seis caminos que implican una difusión es la misma. La dirección de estas flechas secundarias forma ángulo recto con la flecha principal para el caso de sustancias transparentes como el cristal. Cuando estas flechas secundarias se suman a la flecha principal, resulta una flecha final de la misma longitud que la principal, pero girada hacia una dirección ligeramente distinta. Cuanto más grueso es el cristal, menos flechas secundarias existen, y más girada se encuentra la flecha final. Así es como funciona en realidad una lente focalizadora: se puede conseguir que la flecha final para cada camino señale en la misma dirección insertando un espesor extra de cristal en los caminos más cortos.

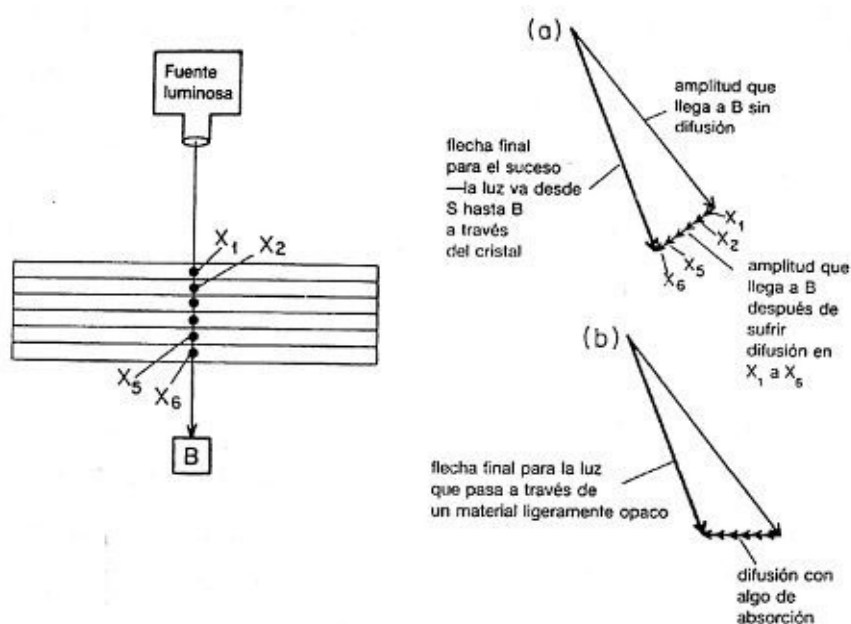


FIGURA 69. La mayor amplitud para la luz que es transmitida, a través de una lámina de cristal y hacia el detector en B, proviene de la zona que representa la no difusión de los electrones en el interior del cristal y que se muestra en (a). A esta flecha sumamos otras seis que representan la difusión de la luz por cada una de las secciones, representadas por los puntos X_1 a X_6 . Estas seis flechas tienen la misma longitud (porque la amplitud de difusión es la misma para cualquier parte del cristal) y señalan en la misma dirección (porque la longitud de cada camino desde la fuente hasta el punto B pasando por cualquier punto X es la misma). Una vez sumadas las flechas pequeñas a la grande, encontramos que la flecha final para la transmisión de la luz a través de una lámina de cristal está más girada de lo que hubiésemos esperado si la luz hubiese incidido directamente. Por esta razón parece que la luz tarda más en atravesar el cristal que en atravesar el vacío o el aire. La cantidad de giro de la flecha final causada por los electrones de un material se denomina el «índice de refracción». Para materiales transparentes, las flechas pequeñas forman ángulo recto con la flecha principal (en realidad se curvan cuando se incluye la difusión doble y triple, impidiendo que la flecha final sea más larga que la flecha principal: la Naturaleza siempre ha funcionado de manera que nunca obtendremos más luz de la que incide). Para materiales parcialmente opacos —que absorben luz hasta cierto punto— las flechas pequeñas señalan hacia la flecha principal, resultando en una flecha final significativamente más corta de lo esperado, como se muestra en (b). Esta flecha final más corta representa una probabilidad reducida de que el fotón sea transmitido por un material parcialmente opaco.

El mismo efecto tendría lugar si los fotones viajaran más despacio a través del cristal que a través del aire: existiría un giro extra de la flecha final. Esta es la razón por la que dije antes que la luz parece ir más despacio a través del cristal (o del agua) que a través del aire. En realidad, la «lentitud» de la luz se debe al giro extra causado por los átomos del cristal (o del agua) difundiendo la luz. El grado de giro extra de la flecha final cuando la luz atraviesa una sustancia determinada se denomina su «índice de refracción»^[22].

Para sustancias que absorben la luz, las flechas secundarias forman ángulos agudos con las flechas principales (ver Fig. 69b). Esto ocasiona que

la flecha final sea más corta que la flecha principal, indicando que la probabilidad de que un fotón pase a través de un cristal parcialmente opaco es menor que para un cristal transparente.

Así que todos los fenómenos y los números arbitrarios mencionados en las dos primeras conferencias —tales como la reflexión parcial de la luz con una amplitud de 0,2, la «lentitud» de la luz en el agua y el cristal, y demás—, se explican con más detalle gracias a tres acciones básicas. Tres acciones que, de hecho, explican también casi todo lo demás.

Es difícil creer que casi toda la aparente extensa variedad de la Naturaleza resulte de la monotonía de combinar repetidamente estas tres acciones básicas. Pero lo es. Esbozaré un poco cómo surge esta variedad.

Podemos empezar con los fotones (ver Fig. 70). ¿Cuál es la probabilidad de que dos fotones, en los puntos 1 y 2 del espacio-tiempo, vayan a dos detectores en los puntos 3 y 4? Existen dos caminos principales por los que este suceso puede tener lugar y cada uno depende de dos cosas que ocurren concomitantemente: los fotones pueden ir directamente — $P(1 \text{ a } 3) \times P(2 \text{ a } 4)$ — o pueden «cruzarse» — $P(1 \text{ a } 4) \times P(2 \text{ a } 3)$ —. Las amplitudes resultantes para estas dos probabilidades se suman, y existe interferencia (como vimos en la segunda conferencia) haciendo que la flecha final varíe en longitud, dependiendo de la situación relativa de los puntos en el espacio-tiempo.

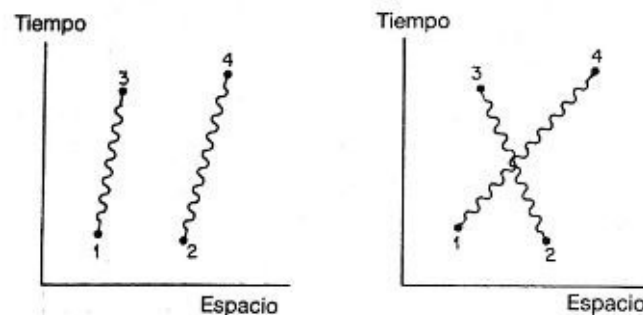


FIGURA 70. Los fotones situados en los puntos 1 y 2 del espacio-tiempo tienen una amplitud de alcanzar los puntos 3 y 4 que puede ser aproximada considerando los dos caminos principales en los que el suceso puede tener lugar: $P(1 \text{ a } 3) \times P(2 \text{ a } 4)$ y $P(1 \text{ a } 4) \times P(2 \text{ a } 3)$, mostrados aquí. Dependiendo de las situaciones relativas de los puntos 1, 2, 3 y 4 existirán diversos grados de interferencia.

¿Qué ocurre si hacemos que 3 y 4 sean el mismo punto del espacio-tiempo (ver Fig. 71)? Digamos que ambos fotones acaban en el punto 3 y vemos cómo afecta esto a la probabilidad del suceso. Ahora tenemos $P(1 \text{ a } 3) \times P(2 \text{ a } 3)$ y $P(2 \text{ a } 3) \times P(1 \text{ a } 3)$, lo que da dos flechas idénticas. Cuando se suman, su longitud es el doble de cualquiera de ellas, y produce una flecha final cuyo cuadrado es cuatro veces el cuadrado de una sola flecha. Como las dos flechas

son idénticas, siempre están «alineadas». En otras palabras, la interferencia no fluctúa de acuerdo con la separación relativa entre los puntos 1 y 2; siempre es positiva. Si no hubiésemos pensado en la interferencia siempre positiva de los dos fotones, hubiésemos pensado que deberíamos haber obtenido, como promedio, la probabilidad cuatro veces mayor. Cuando se ven involucrados varios fotones, la más esperada probabilidad aumenta todavía más.

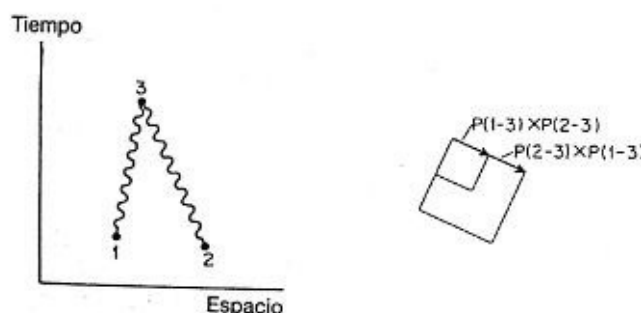


FIGURA 71. Cuando se hacen coincidir los puntos 4 y 3, las dos flechas — $P(1 \text{ a } 3) \times P(2 \text{ a } 3)$ y $P(2 \text{ a } 3) \times P(1 \text{ a } 3)$ — son iguales en longitud y dirección. Cuando se suman siempre se «alinean» y forman una flecha de longitud doble de cualquiera de las otras flechas, con un cuadrado cuatro veces mayor. En consecuencia los fotones tienden a ir al mismo punto del espacio-tiempo. Este efecto se amplifica aún más con más fotones. Esta es la base de la operación de un láser.

Esto tiene como resultado un número de efectos prácticos. Podemos decir que los fotones tienden a tener la misma condición o «estado» (la forma en que varía la amplitud de encontrar uno en el espacio). La posibilidad de que un átomo emita un fotón se ve reforzada si ya están presentes algunos fotones (en un estado en que el átomo los pueda emitir). Este fenómeno de «emisión estimulada» fue descubierto por Einstein cuando presentó la teoría cuántica proponiendo el modelo fotónico de la luz. Los láseres funcionan sobre la base de este fenómeno.

Si hacemos la misma comparación con nuestros falsos electrones de espín cero, ocurriría lo mismo. Pero, en el mundo real, en donde los electrones se hayan polarizados, ocurre algo muy diferente: las dos flechas, $E(1 \text{ a } 3) \times E(2 \text{ a } 4)$ y $E(1 \text{ a } 4) \times E(2 \text{ a } 3)$ se restan —una de ellas se gira 180° antes de sumarse—. Cuando los puntos 3 y 4 son los mismos, las dos flechas tienen la misma longitud y dirección y en consecuencia se cancelan cuando se restan (ver Fig. 72). Esto significa que a los electrones, a diferencia de los fotones, no les gusta ir al mismo sitio; se evitan entre sí como la plaga, no puede haber dos electrones con la misma polarización en el mismo punto del espacio-tiempo —es el denominado «principio de exclusión».

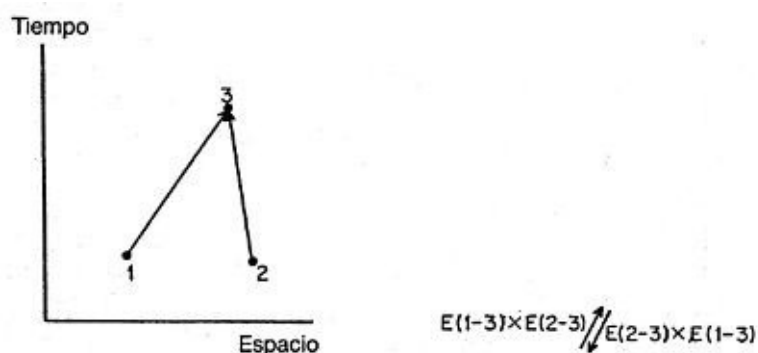


FIGURA 72. Si dos electrones (con la misma polarización) intentan ir al mismo punto del espacio-tiempo, la interferencia es siempre negativa debido a los efectos de polarización: las dos flechas idénticas $E(1 \text{ a } 3) \times E(2 \text{ a } 3)$ y $E(2 \text{ a } 3) \times E(1 \text{ a } 3)$ se restan para dar una flecha final de longitud nula. La aversión de dos electrones a ocupar el mismo lugar del espacio-tiempo se denomina «Principio de Exclusión» y es el responsable de la gran variedad de átomos en el universo.

Este principio de exclusión resulta ser el origen de la gran variedad de propiedades químicas de los átomos. Un protón intercambiando fotones con un electrón girando a su alrededor se denomina un átomo de hidrógeno. Dos protones en el mismo núcleo intercambiando fotones con dos electrones (polarizados en sentidos opuestos) se denomina un átomo de helio. Ya ven, los químicos tienen una forma complicada de contar: en lugar de decir, «uno, dos, tres, cuatro, cinco protones» dicen «hidrógeno, helio, litio, berilio, boro».

Sólo existen dos estados de polarización disponibles para los electrones, luego en un átomo con tres protones en el núcleo intercambiando fotones con tres electrones —una condición que se denomina un átomo de litio— el tercer electrón se encuentra más alejado del núcleo que los otros dos (que han usado el espacio disponible más cercano), e intercambia menos fotones. Esta es la causa de que el electrón pueda desprenderse con facilidad de su núcleo bajo la influencia de los fotones de otros átomos. Un gran número de átomos de este tipo, próximos entre sí, pierden con facilidad su tercer electrón individual para formar un mar de electrones circulando de átomo en átomo. Este mar de electrones reacciona ante cualquier pequeña fuerza eléctrica (fotones), generando una corriente de electrones —estoy describiendo al metal litio conduciendo electricidad—. Los átomos de hidrógeno y helio no ceden sus electrones a otros átomos. Son «aislantes».

Todos los átomos —más de un centenar de tipos diferentes— están constituidos por un cierto número de protones intercambiando fotones con un número igual de electrones. Las formas en que se agrupan son complicadas y ofrecen una enorme variedad de propiedades: algunos son metales, otros aislantes, algunos son gases, otros cristales; forman cosas blandas, cosas

duras, cosas coloreadas, y cosas transparentes, una increíble cornucopia de variedad y excitación que procede del principio de exclusión y de la repetición una y otra vez y otra vez más de las tres acciones tan sencillas $P(A \text{ a } B)$, $E(A \text{ a } B)$, y j . (Si los electrones del mundo no estuvieran polarizados, todos los átomos tendrían propiedades muy similares: los electrones se agruparían, próximos al núcleo de sus átomos y no serían atraídos con facilidad por otros átomos para producir reacciones químicas).

Se podrían preguntar Vds. cómo acciones tan sencillas pueden producir un mundo tan complejo. Se debe a que los fenómenos que vemos en el mundo son el resultado de un enorme entretejido, de un tremendo número de intercambios e interferencias de fotones. El conocer las tres acciones fundamentales es sólo un pequeño inicio hacia el análisis de *cualquier* situación real, donde existen tal multitud de intercambios fotónicos que es imposible calcularlos —se tiene que adquirir experiencia para conocer cuáles son las posibilidades más importantes—. Entonces inventamos ideas tales como «índice de refracción» o «compresibilidad» o «valencia» que nos ayudan a calcular de manera aproximada cuando existe una enorme cantidad de detalles ocurriendo por debajo. Es algo análogo a saber las reglas del ajedrez —que son sencillas y fundamentales— comparado con ser capaz de jugar bien al ajedrez, lo que implica el comprender el carácter de cada posición y la naturaleza de varias situaciones —lo que es mucho más avanzado y difícil.

Las ramas de la física que tratan de cuestiones tales como el porqué el hierro (con 26 protones) es magnético mientras que el cobre (con 29) no lo es, o el porqué un gas es transparente y el otro no, se denominan «física del estado sólido» o «física de los fluidos» o «física honesta». La rama de física que encontró estas tres sencillas acciones (la parte más fácil) se denomina «física fundamental» —¡robamos el nombre para que así los otros físicos se sintiesen incómodos!—. Los problemas más interesantes de la actualidad —y ciertamente los más prácticos— se encuentran obviamente en la física del estado sólido. Pero alguien dijo que no hay nada tan práctico como una buena teoría. ¡Y la teoría de la electrodinámica cuántica es definitivamente una buena teoría!

Finalmente, me gustaría volver a ese número 1,00115965221, el número que en la primera conferencia les dije que había sido medido y calculado tan cuidadosamente. El número representa la respuesta de un electrón a un campo magnético externo —algo denominado el «momento magnético»—. Cuando

Dirac elaboró por primera vez las reglas para calcular este número, utilizó la fórmula $E(A \text{ a } B)$ y obtuvo una respuesta muy sencilla que consideraremos en nuestras unidades como 1. El diagrama de esta primera aproximación del momento magnético de un electrón es muy sencillo —un electrón va de un lugar a otro del espacio-tiempo y se acopla con un fotón de un imán (ver Fig. 73).

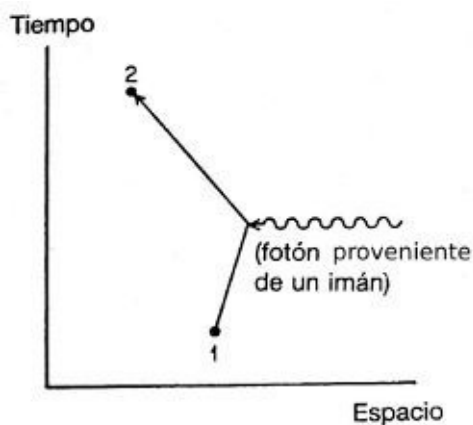


FIGURA 73. El diagrama para el cálculo de Dirac del momento magnético de un electrón es muy sencillo. El valor representado por este diagrama se dirá que es 1.

Al cabo de varios años se descubrió que este valor no era exactamente 1, sino ligeramente superior —algo como 1,00116—. Esta corrección fue calculada por vez primera, en 1948, por Schwinger como $j \times j$ dividido por 2π , y se debía a un camino alternativo por el que el electrón podía ir de un sitio a otro: en lugar de ir directamente de un punto a otro, el electrón viajaba durante un tiempo y de repente emitía un fotón; luego (¡horror!) absorbía su propio fotón (ver Fig. 74). Quizá haya algo de «inmoral» en esto, ¡pero el electrón lo hace! Para determinar la flecha de esta alternativa, tenemos que construir una flecha para cada punto del espacio-tiempo en donde se pueda emitir el fotón y para cada punto en donde pueda ser absorbido. Así que habrá dos extras $E(A \text{ a } B)$, un $P(A \text{ a } B)$ y dos extras j todos multiplicándose entre sí. Los estudiantes aprenden a hacer este sencillo cálculo en su curso elemental de electrodinámica cuántica, durante su segundo año de carrera.

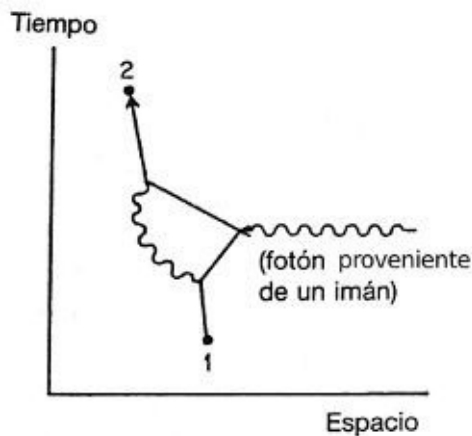


FIGURA 74. Experimentos del laboratorio muestran que el valor real del momento magnético del electrón no es 1, sino un poquito más. Esto se debe a la existencia de alternativas; el electrón puede emitir un fotón y luego absorberlo —requiriendo dos extra $E(A \text{ a } B)$, un $P(A \text{ a } B)$, y dos extra j —. Schwinger calculó que el ajuste que tiene en cuenta esta alternativa es $j \times j$ dividido por 2π . Puesto que esta alternativa es experimentalmente indistinguible del camino original por el que el electrón puede ir —un electrón está inicialmente en el punto 1 y termina en el punto 2— las flechas de las dos vías alternativas se suman, y no existe interferencia.

Pero esperen: los experimentos han medido el comportamiento de un electrón con tanta precisión que tenemos que considerar todavía otras posibilidades en nuestros cálculos —todos los caminos por los que el electrón puede ir de un sitio a otro con *cuatro* acoplamientos extras (ver Fig. 75)—. Hay tres caminos por los que el electrón puede emitir y absorber dos fotones. También hay una nueva e interesante posibilidad (mostrada a la derecha de la Fig. 75): se emite un fotón; forma un par positrón-electrón, y —de nuevo, si Vds. suprimen sus objeciones «morales»— el electrón y el positrón se aniquilan, creando un nuevo fotón que es absorbido finalmente por el electrón. ¡Esta posibilidad también tiene que considerarse!

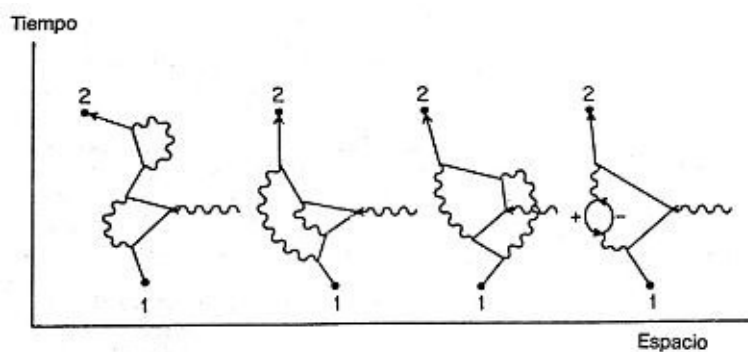


FIGURA 75. Los experimentos de laboratorio han logrado tal precisión que se han tenido que calcular más alternativas, involucrando cuatro acoplamientos extra (sobre todos los posibles puntos intermedios del espacio-tiempo), algunos de los cuales se muestran aquí. La alternativa de la derecha implica un fotón desintegrándose en un par positrón-electrón (como se describe en la Fig. 64), que se aniquila para formar un nuevo fotón, que es finalmente absorbido por el electrón.

La determinación del siguiente término supuso dos años de trabajo a dos

grupos «independientes» de físicos y luego otro año el descubrir que había un error —los experimentos habían medido un valor ligeramente diferente— y durante un tiempo pareció que, por vez primera, la teoría no concordaba con los experimentos, pero no: existía un error de cálculo. ¿Cómo pudieron cometer dos grupos el mismo error? Resultó que hacia el final de sus cálculos los dos grupos compararon notas y limaron las diferencias; es decir, no eran realmente independientes.

El término con *seis j* extras supone todavía más caminos posibles por los que puede ocurrir el suceso; les dibujaré ahora algunos de ellos (ver Fig. 76). Llevó veinte años el que apareciese este refinamiento de la precisión en el valor teórico del momento magnético del electrón. Mientras tanto, los experimentadores llevaron a cabo experimentos más detallados y añadieron algunos dígitos más a su número —y la teoría todavía concordaba.

Por tanto, para hacer nuestros cálculos construimos estos diagramas, escribimos lo que representan matemáticamente y sumamos las amplitudes —un proceso lineal, de «libro de cocina»—. En consecuencia puede ser realizado por las computadoras. Ahora que tenemos super-extra calculadoras, hemos empezado a calcular el término con ocho *j* extras. En la actualidad, el número teórico es 1,00115965246; experimentalmente es 1,00115965221, más o menos 4 en el último decimal. Parte de la incertidumbre del valor teórico (del orden de 4 en el último decimal) se debe al redondeo de los números por el computador; la mayor parte (alrededor de 20) se debe al hecho de que el valor de *j* no se conoce exactamente. El término con ocho *j* extras implica del orden de diez mil diagramas con quinientos términos cada uno —un cálculo fantástico— y se está haciendo en la actualidad.

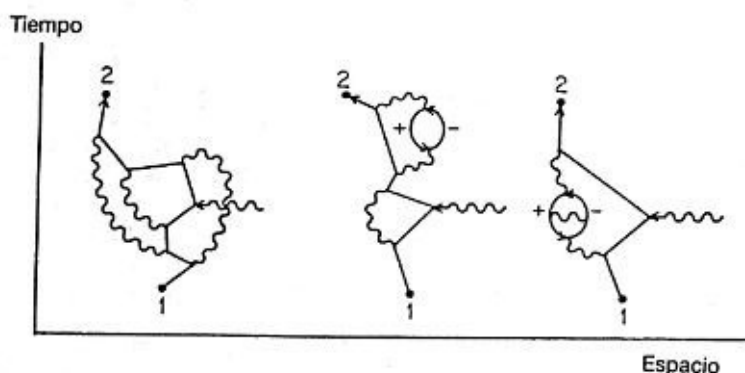


FIGURA 76. Los cálculos en la actualidad se encaminan a obtener un valor teórico aún más preciso. La siguiente contribución a la amplitud, que representa todas las posibilidades con seis acoplamientos extra, implica del orden de 70 diagramas, tres de los cuales se muestran aquí. En 1983, el valor teórico era 1,00115965246 con una incertidumbre de alrededor de 20 en los dos últimos dígitos; el valor experimental era 1,00115965221, con una incertidumbre del orden de 4 en

el último dígito. Esta precisión es equivalente a medir la distancia de Los Angeles a Nueva York, superior a las 3000 millas, con un error del orden del espesor de un cabello humano.

Estoy seguro de que en unos cuantos años más los valores teórico y experimental del momento magnético de un electrón habrán sido calculados con más dígitos. Naturalmente, no estoy seguro de si todavía coincidirán los dos valores. Esto nunca se puede decir mientras no se hagan los cálculos y experimentos.

Y así hemos cerrado el círculo y retornado al número que escogí al principio de estas conferencias para «intimidarles». Espero que comprendan ahora mucho mejor el significado de este número: representa el grado extraordinario al que hemos llegado comprobando que la extraña teoría de la electrodinámica cuántica es realmente correcta.

A lo largo de estas conferencias me ha encantado demostrarles que el precio de lograr una teoría tan precisa ha sido la erosión de nuestro sentido común. Debemos aceptar unos comportamientos un tanto peculiares: la amplificación y supresión de probabilidades, la luz reflejándose en todas las partes de un espejo, la luz viajando por caminos distintos al de la línea recta, los fotones viajando más deprisa o más despacio que la velocidad convencional de la luz, los electrones retrocediendo en el tiempo, los fotones desintegrándose de repente en pares electrón-positrón, etc. Debemos hacerlo para apreciar lo que ocurre realmente en la Naturaleza, subyaciendo a todos los fenómenos que vemos en el mundo.

Con la excepción del detalle técnico de la polarización, les he descrito el marco que nos permite entender todos estos fenómenos. Dibujábamos *amplitud* para cada camino en que podía tener lugar un suceso y las sumábamos cuando, en circunstancias normales esperábamos que se sumasen las probabilidades; multiplicábamos amplitudes cuando esperábamos que se multiplicasen las probabilidades. Pensar todo en términos de amplitudes puede causar dificultades, inicialmente, dada su abstracción, pero al cabo de un tiempo, uno se acostumbra a este extraño lenguaje. Detrás de tantos fenómenos que observamos cada día sólo hay tres acciones básicas: una se describe por el sencillo número de acoplamiento, j ; las otras dos por las funciones — $P(A \text{ a } B)$ y $E(A \text{ a } B)$ — ambas íntimamente relacionadas. Esto es todo, y de aquí surgen todas las demás leyes de la física.

Sin embargo, antes de que termine esta conferencia, me gustaría hacer algunos comentarios adicionales. Se puede entender el carácter y el espíritu de la electrodinámica cuántica sin incluir el detalle técnico de la polarización.

Pero estoy seguro de que todos Vds. se sentirán incómodos a menos que diga algo acerca de lo que he estado dejando a un lado. Ocurre que los fotones vienen en cuatro variedades diferentes, denominadas polarizaciones, que están relacionadas geométricamente con las direcciones del espacio y el tiempo. Así, hay fotones polarizados en las direcciones X, Y, Z y T. (Quizá hayan oído algo como que la luz tiene sólo dos estados de polarización —por ejemplo, un fotón yendo en la dirección X o Y—. Bien, lo han adivinado: en situaciones en las que el fotón recorre una gran distancia y parece ir a la velocidad de la luz, las amplitudes de los términos Z y T se cancelan. Pero para fotones virtuales yendo entre un protón y un electrón en un átomo, es la componente T la que resulta más importante).

De manera análoga, un electrón puede estar en una de las cuatro condiciones que también, pero de manera más sutil, se hayan relacionadas con la geometría. Podemos denotar esas condiciones por 1, 2, 3 y 4. Calcular la amplitud de un electrón yendo del punto A al B en el espacio-tiempo, se vuelve de alguna manera más complicado porque ahora nos podemos hacer preguntas como «¿cuál es la amplitud de que un electrón liberado en la condición 2 en el punto A llegue al punto B en la condición 3?». Las dieciséis combinaciones posibles —provenientes de las cuatro condiciones diferentes en las que puede acabar en B— están relacionadas, de manera matemáticamente sencilla, con la fórmula $E(A \text{ a } B)$ de la que les he hablado.

Para un fotón no se necesita semejante modificación. Así, un fotón polarizado en la dirección X en A sigue polarizado en la dirección X en B, llegando con la amplitud $P(A \text{ a } B)$.

La polarización produce gran número de distintos acoplamientos posibles. Podríamos preguntar, por ejemplo, «¿cuál es la amplitud para que un electrón en la condición 2 absorba un fotón polarizado en la dirección X y en consecuencia se convierta en un electrón en la condición 3?». No todas las condiciones posibles de electrones polarizados y fotones se acoplan, pero las que lo hacen tienen la misma amplitud j aunque, en ocasiones, con un giro adicional de la flecha con un valor múltiplo de 90° .

Todas estas posibilidades para los distintos tipos de polarización y la naturaleza de los acoplamientos, se pueden deducir de manera muy elegante y bella a partir de los principios de la electrodinámica cuántica y de dos suposiciones más: 1) el resultado de un experimento no se ve afectado si el aparato con el que se está realizando se gira en alguna dirección y 2) tampoco existen diferencias si el aparato se encuentra en una nave espacial moviéndose

a una velocidad arbitraria. (Este es el principio de la relatividad).

Este elegante análisis general muestra que cada partícula debe estar en alguna de las clases posibles de polarización que denominamos espín 0, espín 1/2, espín 1, espín 3/2, espín 2 y así sucesivamente. Las distintas clases se comportan de manera diferente. Una partícula de espín 0 es la más sencilla — sólo tiene una componente— y no se ve efectivamente polarizada, (Los falsos electrones y fotones que hemos estado considerando en estas conferencias son partículas de espín 0. Hasta ahora, no se han encontrado partículas fundamentales de espín 0). Un electrón real es un ejemplo de partícula de espín 1/2 y un fotón real es un ejemplo de partícula de espín 1. Tanto las partículas de espín 1/2 como de espín 1 tienen cuatro componentes. Los otros tipos tendrán más componentes; así, las partículas de espín 2 tendrán 10 componentes.

He dicho que la relación entre relatividad y polarización es sencilla y elegante, ¡pero no estoy seguro de que pueda explicarla de manera sencilla y elegante! (Necesitaría al menos una conferencia adicional para poder hacerlo). Aunque los detalles de la polarización no son esenciales para comprender el espíritu y carácter de la electrodinámica cuántica, naturalmente lo son para el cálculo correcto de cualquier proceso real, y a menudo tienen efectos muy importantes. He estado centrando estas conferencias en interacciones relativamente sencillas, a distancias muy pequeñas, entre electrones y fotones, en las que solo se ven implicadas unas pocas partículas. Pero me gustaría hacer una o dos observaciones sobre cómo tienen lugar estas interacciones en el universo, en donde se intercambia un número muy grande de fotones. A tan gran escala, el cálculo de las flechas se vuelve muy complicado.

Existen, sin embargo, algunas situaciones que no son tan difíciles de analizar. Hay circunstancias, por ejemplo, en donde la amplitud de emisión de un fotón por una fuente es independiente de si se ha emitido otro fotón. Esto puede ocurrir cuando la fuente es muy pesada (el núcleo de un átomo), o cuando se mueve del mismo modo un número muy grande de electrones, por ejemplo de un lado a otro en la antena de una estación de radiodifusión, o alrededor de los enrollamientos de un electroimán. Bajo estas condiciones se emiten un gran número de fotones, todos del mismo tipo. La amplitud de absorción de un fotón por un electrón, en tales circunstancias, es independiente de que él o cualquier otro electrón haya absorbido otros fotones previamente. En consecuencia, el comportamiento total puede darse mediante

la amplitud de absorción de un fotón por un electrón, que depende sólo de la posición del electrón en el espacio y el tiempo. Los físicos utilizan palabras comunes para describir esta circunstancia. Dicen que el electrón se está moviendo en un campo externo. Para los físicos la palabra «campo» quiere decir «una cantidad que depende de la posición en el espacio y en el tiempo». La temperatura del aire nos provee con un buen ejemplo: varía dependiendo de dónde y cuándo se haya hecho la medida. Cuando se tiene en cuenta la polarización, existen más componentes del campo. (Existen cuatro componentes —correspondientes a la amplitud de absorción de cada uno de los distintos tipos de polarización (X, Y, Z, T) en que pueda estar el fotón— que se denominan de manera técnica, potencial electromagnético escalar y vectorial. A partir de sus combinaciones, la física clásica ha derivado componentes más convenientes denominadas campo eléctrico y magnético).

En una situación en donde el campo eléctrico y magnético varíen con suficiente lentitud, la amplitud de que un electrón viaje una distancia muy larga depende del camino que escoja. Como vimos con anterioridad en el caso de la luz, los caminos más importantes son aquéllos en que los ángulos de la amplitud de los caminos próximos son casi los mismos. El resultado es que la partícula no viaja necesariamente en línea recta.

Esto nos retrotrae a la física clásica, en donde se supone que existen campos y que los electrones se mueven a través de ellos de manera tal que hacen mínima una determinada cantidad. (Los físicos denominan a esta cantidad «acción» y formulan esta regla como «el principio de mínima acción»). Este es un ejemplo de cómo las reglas de la electrodinámica cuántica producen fenómenos a gran escala. Podríamos extendernos en muchas direcciones, pero tenemos que limitar de algún modo la panorámica de estas conferencias. Sólo quiero recordarles que los efectos que vemos a gran escala y los extraños fenómenos que vemos a pequeña escala están ambos producidos por la interacción de electrones y fotones y todos se describen, en última instancia, por la teoría de la electrodinámica cuántica.

Capítulo 4

CABOS SUELTOS

Voy a dividir esta conferencia en dos partes. En primer lugar voy a hablar de problemas asociados con la propia teoría de la electrodinámica cuántica, suponiendo que todo lo que existe en el mundo son electrones y fotones. Luego hablaré sobre la relación de la electrodinámica cuántica con el resto de la física.

La característica más chocante de la teoría de la electrodinámica cuántica es el loco marco de las amplitudes ¡qué Vds. pueden pensar que indica problemas de algún tipo! Sin embargo, los físicos han estado jugando con amplitudes desde hace más de cincuenta años y han terminado acostumbrándose a ellas. Además, todas las nuevas partículas y los nuevos fenómenos que somos capaces de obtener encajan perfectamente con lo que se puede deducir de ese marco de amplitudes, en el que la probabilidad de un suceso es el cuadrado de una flecha final cuya longitud viene determinada por la combinación de flechas en formas divertidas (con interferencias y cosas por el estilo). Por tanto, este marco de amplitudes *no presenta ninguna duda desde el punto de vista experimental*: puede Vd. tener todas las preocupaciones filosóficas que quiera en lo que se refiere al significado de las amplitudes (si, de hecho, significan algo) pero como la física es una ciencia experimental y el marco concuerda con los experimentos, esto es suficiente para nosotros por el momento.

Existe un conjunto de problemas asociados con la teoría de la electrodinámica cuántica que tiene que ver con la mejora del método de cálculo de la suma de todas las flechitas —técnicas diversas de las que se dispone en circunstancias diferentes— y que llevan tres o cuatro años el

dominarlas a los estudiantes graduados. Puesto que son problemas técnicos, no voy a discutirlos aquí. Es sólo una cuestión de mejorar continuamente las técnicas para analizar lo que la teoría tiene realmente que decir en circunstancias diferentes.

Pero existe un problema adicional característico de la propia teoría de la electrodinámica cuántica, que llevó veinte años resolverlo. Está relacionado con los electrones y fotones idealizados y los números n y j .

Si los electrones fuesen ideales y se moviesen de un punto a otro del espacio-tiempo, sólo por el camino directo (mostrado a la izquierda en la Fig. 77) no existiría problema: n sería simplemente la masa de un electrón (que podemos determinar por observación) y j simplemente su «carga» (la amplitud de acoplamiento de un electrón y un fotón). También puede ser determinada experimentalmente.

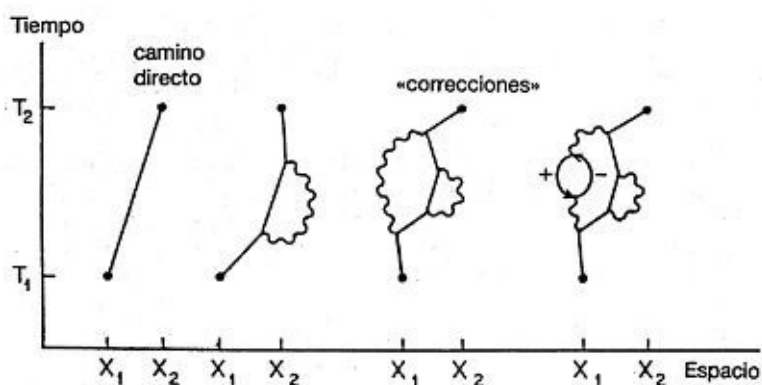


FIGURA 77. Cuando calculamos la amplitud de que un electrón vaya de un punto a otro del espacio-tiempo, utilizamos la fórmula $E(A \text{ a } B)$ para el camino directo. (Luego hacemos «correcciones» que incluyen uno o más fotones emitiéndose y absorbiéndose). $E(A \text{ a } B)$ depende de $(X_2 - X_1)$, $(T_2 - T_1)$ y n , un número que introducimos en las fórmulas para que la respuesta fuese la correcta. El número n se denomina la «masa en reposo» de un electrón «ideal», y no puede ser medida experimentalmente porque la masa en reposo de un electrón real, m , incluye todas las «correcciones». Existía una cierta dificultad para calcular el n que debía de utilizarse en $E(A \text{ a } B)$, y llevó veinte años el remontarla.

Pero no existen esos electrones ideales. La masa que observamos en el laboratorio es la de un electrón *real*, que emite y absorbe sus propios fotones de vez en cuando, y que por tanto depende de la amplitud de acoplamiento, j . Y la carga que observamos está entre la de un electrón *real* y un fotón *real* — que puede formar un par electrón-positrón de vez en cuando —, dependiendo por consiguiente de $E(A \text{ a } B)$, que involucra a n (ver Fig. 78). Puesto que la masa y carga de un electrón se ven afectadas por estas y otras alternativas, la masa experimental medida, m , y la carga experimentalmente medida, e , del electrón son diferentes de los números que usamos en nuestros cálculos, n y j .

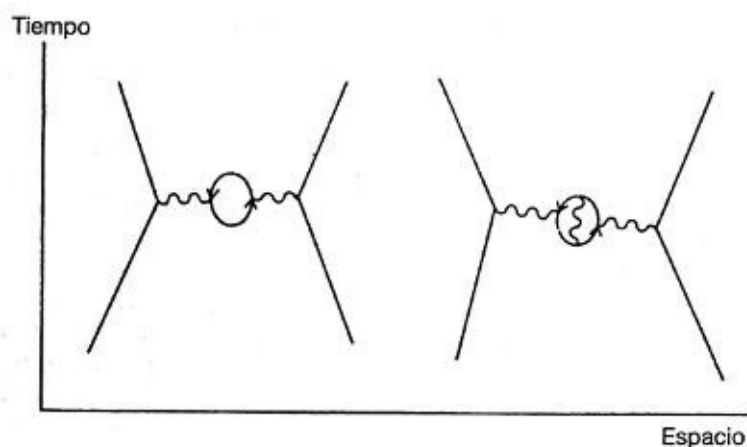


FIGURA 78. La amplitud experimental de acoplamiento entre el electrón y el fotón, un número misterioso, e , es un número determinado experimentalmente que incluye todas las «correcciones» para un fotón yendo de un punto a otro del espacio-tiempo, dos de las cuales se muestran aquí. Cuando hacemos los cálculos, necesitamos un número, j , que no incluye estas correcciones, sino solamente la del fotón yendo directamente de un punto a otro. Existe una dificultad para calcular este valor de j que es análoga a la existente para calcular el valor de n .

Si existiese una relación matemática definida entre n y j por un lado, y m y e por otro, tampoco existiría problema: simplemente calcularíamos los valores de n y j necesarios para empezar, a fin de terminar con los valores observados, m y e . (Si nuestros cálculos no concordasen con m y e , bailaríamos los valores iniciales de n y j hasta que lo hiciesen).

Veamos cómo calculamos m realmente. Escribimos una serie de términos que son análogos a la serie que vimos para el momento magnético del electrón: el primer término no tiene acoplamiento —sólo $E(A \text{ a } B)$ — y representa un electrón ideal viajando directamente de un punto a otro del espacio-tiempo. El segundo término tiene dos acoplamientos y representa un fotón que es emitido y absorbido. Luego viene el término con cuatro, seis y ocho acoplamientos y así sucesivamente (algunas de estas «correcciones» se muestran en la Fig. 77).

Cuando calculemos los términos con acoplamientos, debemos considerar (como siempre) todos los posibles puntos en donde los acoplamientos pueden tener lugar, incluido el caso en el que los dos puntos del acoplamiento se superponen —con distancia cero entre ellos—. El problema es que cuando intentamos calcular todas las situaciones hasta la de distancia cero, la ecuación explota ante nosotros y da respuestas sin sentido —cosas como infinito—. Esto originó muchas molestias cuando la teoría de la electrodinámica cuántica apareció por primera vez. ¡La gente obtenía infinito para cada problema que intentaba calcular! (Uno debería descender hasta la distancia cero a fin de ser matemáticamente consistente, pero aquí es donde ni

n ni j tienen ningún sentido; y donde está el problema).

Bien, en lugar de incluir todos los posibles puntos de acoplamiento hasta la distancia cero, si se *detienen* los cálculos cuando la distancia entre los puntos de acoplamiento es muy pequeña —digamos, 10^{-16} cm— existen entonces unos valores definidos para n y j que se pueden usar de manera que la masa calculada coincida con la masa m observada en los experimentos y la carga calculada concuerde con la carga observada, e . Pero, aquí está la trampa: si alguien más hace los cálculos y los detiene a distancia diferente —digamos 10^{-40} cm—, ¡los valores de n y j que precisa para obtener los mismos valores de m y e son *diferentes*!

Veinte años después, en 1949, Hans Bethe y Victor Weisskopf notaron algo: si dos personas que paraban los cálculos a distancias diferentes para determinar n y j a partir de los mismos m y e , calculaban luego la respuesta a *otro* problema —cada uno utilizando los valores apropiados pero diferentes de n y j —, cuando incluían todas las flechas de todos los términos ¡sus respuestas a este otro problema eran casi idénticas! De hecho, cuanto más se aproximaban a la distancia cero en los cálculos de n y j , ¡mejor concordaban las respuestas finales para los otros problemas! Schwinger, Tomonaga y yo mismo inventamos, de manera independiente, formas de hacer cálculos definidos para confirmar que esto era verdad (obtuvimos premios por esto). ¡La gente pudo finalmente calcular con la teoría de la electrodinámica cuántica!

Por consiguiente parece que las *únicas* cosas que dependen de las pequeñas distancias entre los puntos de acoplamiento son los valores de n y j —*números teóricos que no son observables directamente de ninguna manera*—; todo lo demás, lo que puede ser observado, parece no verse afectado.

El juego de capas que utilizamos para encontrar n y j se denomina técnicamente «renormalización». Pero independientemente de lo inteligente que sea la palabra, es lo que yo llamaría ¡un proceso de profundización! El haber tenido que recurrir a tal malabarismo nos ha impedido demostrar que la teoría de la electrodinámica cuántica es matemáticamente autoconsistente. Es sorprendente que todavía no se haya podido demostrar de una manera u otra que la teoría es autoconsistente; sospecho que la renormalización no sea legítima desde el punto de vista matemático. Lo que *es* cierto es que no tenemos una buena forma matemática de describir la electrodinámica cuántica: semejante montón de palabras para describir la relación entre n y j y m y e no son buenas matemáticas^[23].

Existe un problema más profundo y bonito asociado a la constante de acoplamiento experimental, e —la amplitud de que un electrón real emita o absorba un fotón real—. Es un número sencillo cuyo valor, próximo a $-0,08542455$, ha sido determinado experimentalmente. (Mis amigos físicos no reconocerán este número porque prefieren recordar la inversa de su cuadrado: aproximadamente $137,03597$ con una incertidumbre de 2 en la última cifra decimal. Esta constante ha sido un misterio desde su descubrimiento, hace más de cincuenta años, y todos los buenos físicos teóricos colocan este número en su pizarra y se preocupan por él).

A Vds. les gustaría saber inmediatamente de dónde sale este número de acoplamiento: ¿está relacionado con π , o quizá con la base de los logaritmos neperianos? Nadie lo sabe. Es uno de los condenados misterios *más grandes* de la física: un *número mágico* que aparece sin que el hombre entienda cómo. Podrían decir que «la mano de Dios» escribió ese número, y que «nosotros no sabemos cómo cogió su lápiz». Sabemos al son de qué música bailar para medir experimentalmente con gran precisión este número, pero no sabemos el son para obtener este número en un computador —¡sin introducirlo en secreto!

Una buena teoría diría que e es la raíz cuadrada de 3 dividida por 2π cuadrado, o algo similar. De vez en cuando, se han hecho sugerencias acerca de qué era e , pero ninguna ha resultado útil. Primero, Arthur Eddington demostró, por pura lógica, que el número que les gustaba a los físicos era exactamente 136, el valor experimental de la época. Luego, cuando experimentos más precisos demostraron que el número estaba más cercano a 137, Eddington descubrió un ligero error en su anterior razonamiento y mostró, de nuevo por pura lógica, ¡qué el número tenía que ser el entero 137! De vez en cuando, alguien advertía que una determinada combinación de π y e (la base del logaritmo neperiano) y doses y cincos producía la misteriosa constante de acoplamiento, pero es un hecho no apreciado en su totalidad por la gente que juega con la aritmética que es sorprendente la *gran* cantidad de números que se pueden obtener a partir de píses y ees y similares. Por tanto, a lo largo de la historia de la física moderna han aparecido artículo tras artículo de personas que han obtenido e con varias cifras decimales para que la siguiente ronda de experimentos mejorados produjese un valor en desacuerdo con él.

Pero aunque hoy tenemos que recurrir a un proceso de profundización para calcular j , es posible que, algún día, se encuentre una conexión

matemáticamente válida entre j y e . En este caso sin duda que aparecerían otro montón de artículos que nos dirían cómo calcular j «con nuestras propias manos» por así decirlo, proponiendo que j es 1 dividido por $4 \times \pi$ o algo similar.

Esto muestra todos los problemas asociados con la electrodinámica cuántica.

Cuando planeé estas conferencias, intenté concentrarme sólo en la parte de la física que conocemos muy bien —describirla al completo y no decir, nada más—. Pero ahora que hemos llegado hasta aquí y siendo un profesor (lo que significa tener la costumbre de no ser capaz de dejar de hablar a la hora debida), no puedo resistir el contarles algo del resto de la física.

Primero, debo decirles inmediatamente que ninguna de las partes del resto de la física ha sido tan comprobada como la electrodinámica: algunas de las cosas que voy a decirles son buenas suposiciones, algunas otras, teorías parcialmente elaboradas y otras, pura especulación. Por lo tanto, esta presentación va a parecer algo así como un relativo lío, y comparada con las otras conferencias será incompleta y falta de muchos detalles. Sin embargo, ocurre que la estructura de la teoría de la QED sirve como base excelente para la descripción de otros fenómenos del resto de la física.

Comenzaré hablando de protones y neutrones, lo que constituye el núcleo de los átomos. Cuando se descubrieron por vez primera los protones y neutrones se pensó que eran partículas sencillas, pero pronto se evidenció que no eran sencillas —en el sentido de que la amplitud, para ir de un punto a otro podía explicarse por la fórmula $E(A \text{ a } B)$, pero con un número diferente para n en ella—. Por ejemplo, el protón tiene un momento magnético que si se calcula de la misma manera que para el electrón, debería aproximarse a 1. Pero de hecho, experimentalmente resulta un valor completamente loco ¡-2,79! En consecuencia, se vio pronto que algo ocurría dentro del protón que no se tenía en cuenta en las ecuaciones de la electrodinámica cuántica. Y el neutrón, que no debería tener interacción magnética si fuese realmente neutro, ¡tiene un momento magnético de aproximadamente -1,93! Por tanto, se sabía desde hacía bastante tiempo que algo raro ocurría también dentro del neutrón.

Existía también el problema de qué mantiene a los neutrones y protones juntos dentro del núcleo. Se apreció en seguida que no podía ser debido al intercambio de fotones, porque las fuerzas que mantenían unido el núcleo eran mucho más fuertes —la energía requerida para romper un núcleo es

mucho mayor que la que se requiere para extraer un electrón de un átomo, en una proporción análoga al mucho más elevado poder de destrucción de la bomba atómica frente al de la dinamita—: la explosión de la dinamita es una redistribución de la disposición de los electrones, mientras que la explosión de una bomba atómica es una redistribución de la disposición de protones y neutrones.

Para descubrir algo más sobre lo que mantiene unido a los núcleos, se realizaron muchos experimentos en los que se hacía incidir, sobre los núcleos, protones con energías cada vez más elevadas. Se esperaba que sólo se obtendrían protones y neutrones. Pero cuando la energía fue lo suficientemente elevada, aparecieron nuevas partículas. Primero fueron los piones, luego las landas, y sigmas, y ros, y luego agotaron el alfabeto. Después aparecieron partículas con números (sus masas), tales como sigma 1190 y sigma 1386. Pronto se hizo evidente que el número de partículas en el mundo no tenía fin y dependía de la cantidad de energía empleada para desintegrar el núcleo. Existen en la actualidad más de cuatrocientas partículas. No podemos aceptar cuatrocientas partículas ¡es demasiado complicado!^[24]

Grandes inventores como Murray Gell-Mann casi se vuelven locos tratando de descubrir las leyes de comportamiento de estas partículas, y a principios de los setenta aparecieron con la teoría cuántica de las interacciones fuertes (o «cromodinámica cuántica»), cuyos actores principales eran las partículas denominadas «quarks». Todas las partículas compuestas de quarks se dividen en dos clases: algunas, como el protón y el neutrón, están formadas por tres quarks (y tienen el horrible nombre de «bariones»); otras, como los piones están constituidas por un quark y un antiquark (y se denominan «mesones»).

Déjenme construir una tabla de partículas fundamentales tal como aparecen en la actualidad (ver Fig. 79). Comenzaré por las partículas que van de punto a punto según la fórmula $E(A \rightarrow B)$ —modificada por el mismo tipo de reglas de polarización del electrón—, llamadas «partículas de espín 1/2». La primera de estas partículas es el electrón y su número másico es 0,511 en unidades que usaremos todo el tiempo y que denominaremos MeV^[25].

	partículas de espín $\frac{1}{2}$	fotón
	electrón e .511	0
	quark d ~10	-1/3
	quark u ~10	+2/3
		acoplo

nombre	electrón
símbolo	e
masa (MeV)	.511

FIGURA 79. A nuestra lista de todas las partículas del mundo se inicia con las partículas de «espín 1/2»: el electrón (con una masa de 0,511 MeV), y dos «aromas» de quark, d y u (ambos con una masa de alrededor de 10 MeV). Los electrones y quarks tienen una «carga» —es decir, se acoplan con fotones con los siguientes valores (en términos de la constante de acoplamiento $-j$) = -1 , $-1/3$ y $+2/3$.

Debajo del electrón dejaré un espacio en blanco (que ocuparemos posteriormente), debajo del cual situaré dos tipos de quarks —el d y el u —. La masa de estos quarks no se conoce con exactitud, una buena aproximación es de alrededor de 10 MeV cada uno. (El neutrón es ligeramente superior al protón, lo que parece implicar —como verán en seguida— que el quark d es ligeramente más pesado que el quark u).

Al lado de cada partícula colocaré su carga, o constante de acoplamiento, en términos de $-j$, el número de acoplamiento con fotones invertido de signo. Esto hace que la carga del electrón sea -1 , consistente con el convenio iniciado por Benjamin Franklin y al que desde entonces nos hemos adherido. Para el quark d la amplitud de acoplamiento con un fotón es $-1/3$, y para el quark u es $+2/3$ (si Benjamin Franklin hubiese sabido de los quarks, ¡hubiese tomado la carga del electrón, al menos, como -3 !).

Entonces, la carga del protón es $+1$ y la carga del neutrón es cero. Jugando un poco con los números se puede ver que el protón —hecho de tres quarks— debe tener dos u , y un d mientras que el neutrón —también formado por tres quarks— debe tener dos d y un u (ver Fig. 80).

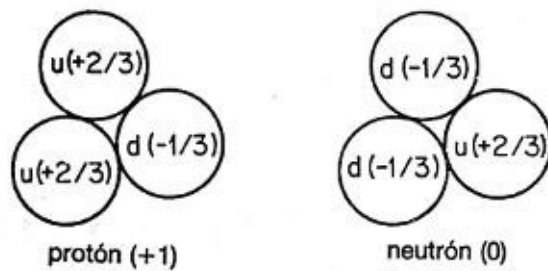


FIGURA 80. Todas las partículas formadas por quarks, de las cuales el protón y el neutrón son los ejemplos más comunes, vienen en una de las dos posibles clases: las constituidas por un quark y un anti-quark, y las constituidas por tres quarks. La carga de los quarks d y u se combina para dar +1 para el protón y cero para el neutrón. El hecho de que el protón y el neutrón estén formados por partículas cargadas moviéndose dentro de ellos da una pista de por qué el protón tiene un momento magnético superior a 1, y por qué el supuestamente neutro neutrón tiene un cierto momento magnético.

¿Qué mantiene unidos a los quarks? ¿Los fotones que se mueven entre ellos? (Puesto que un quark d tiene una carga de $-1/3$ y un quark u de $+2/3$, los quarks, además de electrones, emiten y absorben fotones). No, estas fuerzas eléctricas son demasiado débiles para hacerlo. Se ha inventado algo más, que al viajar de uno a otro quark los mantiene unidos; algo llamado «gluones»^[26]. Los gluones son un ejemplo de otro tipo de partículas denominado «espín 1» (como son los fotones); viajan de un punto a otro con una amplitud determinada por la misma fórmula que la de los fotones, $P(A \text{ a } B)$. La amplitud para que los gluones sean emitidos o absorbidos por los quarks es un número misterioso, g , que es mucho mayor que j (ver Fig. 81).

	partículas de espín 1	
	fotón	gluón
partículas de espín 1/2		
electrón e .511	0 -1	0
	0	0
quark d 10	-1/3	g
quark u ~ 10	+2/3	g
	acoplos	

nombre — electrón

símbolo — e

masa (MeV) — .511

FIGURA 81. Los «gluones» mantienen a los quarks unidos para formar protones y neutrones, e indirectamente son responsables del hecho de que los protones y neutrones se mantengan juntos en

el núcleo de un átomo. Los gluones mantienen unidos a los quarks mediante fuerzas mucho más fuertes que las eléctricas. La constante de acoplamiento de los gluones, g , es mucho más grande que j , lo que hace mucho más difíciles los cálculos de los términos con acoplamientos: la mayor precisión que se puede esperar de momento es del 10%.

Los diagramas que construimos de quarks intercambiando gluones son muy similares a los dibujos que hicimos de los electrones intercambiando fotones (ver Fig. 82). Tan parecidos, de hecho, que pueden Vds. decir que los físicos carecen de imaginación —¡simplemente han copiado la teoría de la electrodinámica cuántica para las interacciones fuertes!—. Y tendrán razón: es lo que hicimos, pero con un ligero cambio.



FIGURA 82. El diagrama de uno de los caminos por el que dos quarks pueden intercambiar un gluón es tan similar al de dos electrones intercambiando un fotón que Vds. pueden pensar que los físicos han copiado, para la teoría de las «interacciones fuertes» que mantiene unidos a los quarks dentro de los protones y neutrones, la teoría de la electrodinámica cuántica. Bien, lo han hecho —casi.

Los quarks tienen un tipo de polarización adicional que no está relacionada con la geometría. Los idiotas de los físicos, incapaces de encontrar alguna maravillosa palabra griega, bautizaron a este tipo de polarización con el desafortunado nombre de «color», que no tiene nada que ver con el color habitual. En un momento dado, el quark puede estar en una de las tres condiciones o «colores» —R, G o B (¿pueden imaginar que significa?)^[27]—. El «color» de un quark puede combinarse cuando el quark emite o absorbe un gluón. Los gluones vienen en ocho tipos diferentes, de acuerdo con los «colores» con que pueden acoplarse. Por ejemplo, si un quark rojo cambia a verde, emite un gluón rojo-antiverde^[28] —un gluón que toma el rojo de quark y cede el verde— («antiverde» significa que el gluón está transportando al verde en la dirección opuesta). Este gluón puede ser absorbido por el quark verde, que cambia a rojo (ver Fig. 83). Existen ocho posibles gluones distintos, tales como rojo-antirrojo, rojo-antiazul, rojo-antiverde, etc. (pensarán que eran nueve, pero por razones técnicas, faltó uno). La teoría no es muy complicada. La regla completa de los gluones es: los gluones se acoplan con cosas con «color» —sólo se precisa un poco de contabilidad para seguir el rastro de a dónde van los «colores».

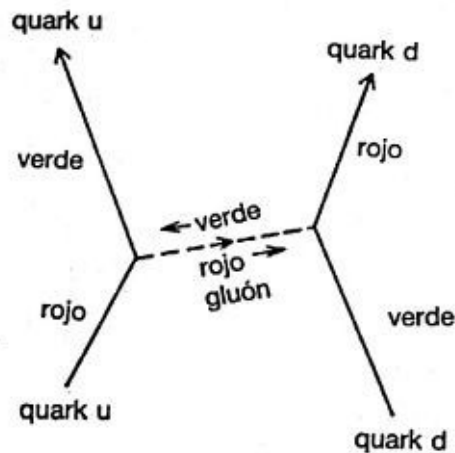


FIGURA 83. La teoría gluónica difiere de la electrodinámica en que los gluones se acoplan con cosas que tienen «color» (con una de las tres posibles condiciones —«rojo», «verde», y «azul»—). Aquí, un quark u cambia a verde emitiendo un gluón rojo-antiverde que es absorbido por un quark d verde que cambia a rojo (si el «color» se transporta hacia atrás en el tiempo, adquiere el prefijo «anti»).

Existe, sin embargo, una interesante posibilidad creada por esta regla: los gluones se pueden acoplar con otros gluones (ver Fig. 84). Por ejemplo, un gluón verde-antiazul encontrándose con un gluón rojo-antiverde resulta en un gluón rojo-antiazul. La teoría de los gluones es muy sencilla —se dibuja el diagrama y se siguen los «colores». La magnitud de los acoplamientos en todos los diagramas viene determinada por la constante de acoplamiento para gluones, g .

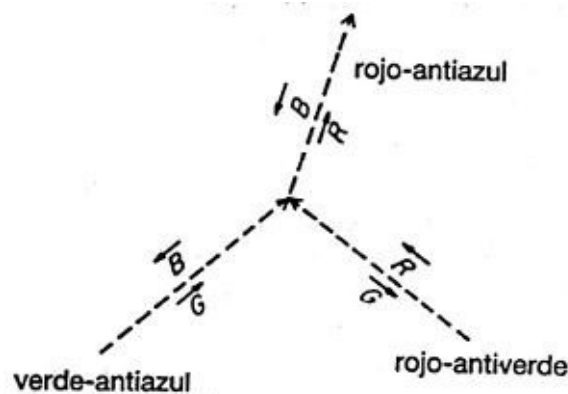


FIGURA 84. Puesto que los gluones tienen también «color» se pueden acoplar entre sí. Aquí un gluón verde-antiazul se acopla con un gluón rojo-antiverde para formar un gluón rojo-antiazul. La teoría gluónica es fácil de entender —sólo hay que seguir los «colores».

La teoría de los gluones no es muy distinta formalmente de la electrodinámica cuántica. Entonces, ¿cómo se compara con los experimentos? Por ejemplo, ¿cómo es el momento magnético observado del protón cuando se compara con el valor calculado de la teoría?

Los experimentos son muy precisos —demuestran que el valor del

momento magnético es $2,79275$ —. En el mejor de los casos, la teoría sólo puede dar $2,7 \pm 0,3$ —si uno es lo suficientemente optimista sobre la precisión de su análisis ¡un error de 10% que es 10 000 veces más impreciso que el experimento! Tenemos una teoría sencilla, definida, que se supone que explica todas las propiedades de los protones y neutrones, y sin embargo no se puede calcular nada con ella porque las matemáticas son demasiado complicadas para nosotros. (Pueden imaginarse lo que estoy tratando de hacer, y no me lleva a ninguna parte). La razón por la que no podemos calcular nada con mayor precisión se debe a la constante de acoplamiento de los gluones g , que es mucho más grande que la de los electrones. Términos con dos, cuatro, e incluso seis acoplamientos, no son meras correcciones menores a la amplitud principal; representan contribuciones considerables que no se pueden ignorar. De manera que hay flechas de tantas posibilidades diferentes que no hemos sido capaces de organizarlas de manera razonable para encontrar cuál es la flecha final.

En los libros se dice que la ciencia es sencilla: se hace una teoría y se compara con los experimentos; si la teoría no funciona, se descarta y se hace una nueva teoría. Aquí tenemos una teoría definida y cientos de experimentos pero ¡no podemos compararlos! Es una situación que nunca antes había existido en la historia de la física. Nos encontramos temporalmente encajonados, incapaces de dar con un método de cálculo. Estamos enterrados bajo una capa de flechitas.

A pesar de nuestras dificultades para realizar cálculos con la teoría, entendemos algunas cosas, cualitativamente, gracias a la cromodinámica cuántica (las interacciones fuertes de quarks y gluones). Los objetos compuestos de quarks, que vemos, son de «color» neutro: los grupos de tres quarks contienen un quark de cada «color» y los pares quark-antiquark tienen la misma amplitud de ser rojo-antirrojo, verde-antiverde o azul-antiazul. También entendemos el por qué los quarks nunca se pueden producir como partículas individuales —el por qué, independientemente de cuanta energía se utilice para hacer chocar un núcleo contra un protón, en lugar de observar la aparición de quarks individuales, vemos un haz de mesones y bariones (pares quark-antiquark y grupos de tres quarks).

La cromodinámica cuántica y la electrodinámica cuántica no son todo en la física. De acuerdo con ellas, un quark no puede cambiar su «aroma»: un quark u siempre será un quark u ; un quark d siempre será un quark d . Pero la Naturaleza en ocasiones se comporta de manera diferente. Existe una forma

de radiactividad que tiene lugar lentamente —del tipo de la que le preocupa a la gente que se pueda escapar de los reactores nucleares— denominada desintegración beta, que implica un neutrón cambiando a protón. Puesto que un neutrón está formado por dos quarks d y un quark u , mientras que un protón se hace con dos u y un d , lo que realmente ocurre es que uno de los quarks tipo d del neutrón se cambia a un quark tipo u (ver Fig. 85). Y así es como ocurre: el quark d emite algo nuevo análogo a un fotón, denominado un W , que se acopla con un electrón y con otra partícula nueva llamada anti-neutrino, un neutrino que va hacia atrás en el tiempo. El neutrino es otra partícula de espín 1/2 (como el electrón y los quarks) pero sin masa ni carga (no interacciona con los fotones). Tampoco interacciona con los gluones, sólo se acopla con los W (ver fig 86).

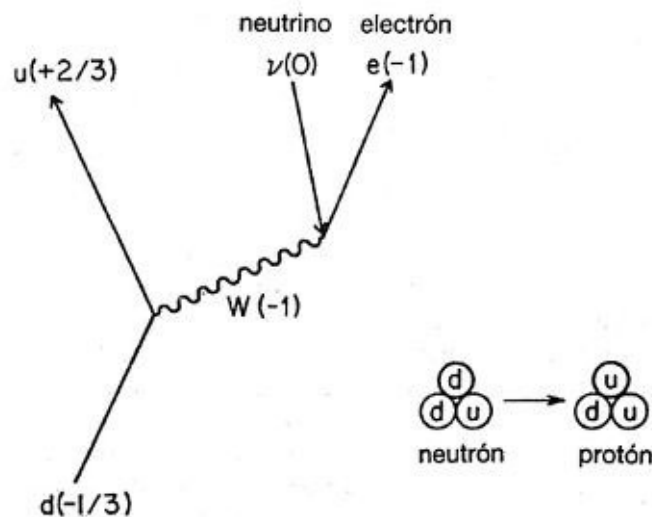


FIGURA 85. Cuando un neutrón se desintegra en un protón (un proceso denominado «desintegración beta») lo único que cambia es el «aroma» de un quark —de d a u emitiéndose un electrón y un anti-neutrino. Este proceso es relativamente lento, por lo que se ha propuesto la existencia de una partícula intermedia (denominada un «bosón intermedio W ») con una masa muy elevada (de alrededor de 80 000 MeV) y una carga de $+1$.

La W es una partícula tipo 1 (como el fotón y el gluón) que cambia el «aroma» de un quark y se lleva su carga —el d , carga $-1/3$, cambia a u , carga $+2/3$ ¡una diferencia de -1 ! (No cambia el «color» del quark)—. Dado que el W^- toma una carga de -1 (y su antipartícula, el W^+ toma una carga de $+1$), también se puede acoplar con un fotón. La desintegración beta dura mucho más que las interacciones de fotones y electrones, de modo que se piensa que el W debe de tener una masa muy elevada (alrededor de 80 000 MeV), a diferencia del fotón y del gluón. No hemos sido capaces de observar el W aislado, debido a la energía tan elevada que requiere el arrancar una partícula con una masa tan grande^[29].

		partículas de espín 1		
	partículas de espín 1/2	fotón	gluón	W
		0	0	~80,000
nombre electrón símbolo e masa (MeV) .511	electrón	-1	0	↕
	e			
	.511			
	neutrino	0	0	↕
	ν_e			
	0			
	quark	-1/3	g	↕
	d			
	~20			
	quark	+2/3	g	↕
	u			
	~20			
		acoplos		

FIGURA 86. El W se acopla con el electrón y el neutrino por un lado y el quark d y el u por el otro.

Existe otra partícula, que podríamos considerarla como una W neutra, denominada Z_0 . La Z_0 no cambia la carga del quark, pero se acopla con un quark d, un quark u, un electrón, un neutrino (ver Fig. 87). Esta interacción tiene el equívoco nombre de «corriente neutra» y produjo mucha excitación cuando se descubrió hace unos pocos años.

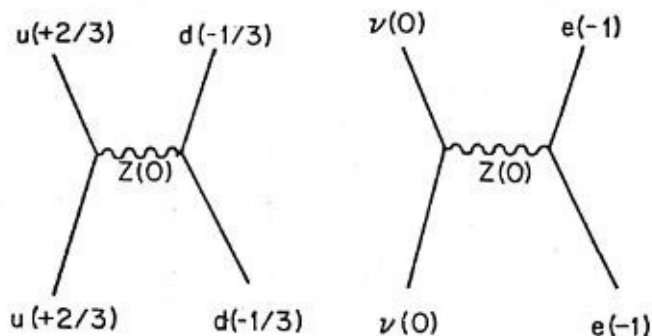


FIGURA 87. Cuando no hay cambio de carga en ninguna de las partículas, el W tampoco tiene carga (se denomina en este caso Z_0). Estas interacciones se denominan «corrientes neutras». Aquí se muestran dos posibilidades.

La teoría de las W es clara y bonita si se permiten tres tipos de acoplamiento entre tres tipos de W (ver Fig. 88). La constante de acoplamiento observada para el W es muy similar a la del fotón —en el entorno de j .

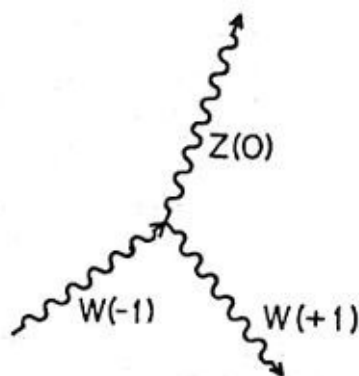


FIGURA 88. Es posible un acoplamiento entre un W_{-1} y su antipartícula (un W_{+1}) y un neutro W (Z_0). La constante de acoplamiento para W es próxima a j , sugiriendo que los W y los fotones pueden ser distintos aspectos de la misma cosa.

Por tanto, existe la posibilidad de que los tres W y el fotón sean aspectos diversos de la misma cosa. Stephen Weinberg y Abdus Salam intentaron combinar la electrodinámica cuántica con lo que se denomina «interacción débil» (interacción con los W) en una única teoría cuántica y lo consiguieron. Pero si Vds. miran los resultados que obtuvieron se puede ver el pegamento^[30] por así decirlo. Está claro que el fotón y los tres W están de alguna manera interconectados, pero al nivel actual de conocimiento es difícil ver la conexión con claridad —todavía se pueden ver las «costuras» en las teorías; aún no se han pulido de forma que las conexiones sean más hermosas y por tanto, probablemente más correctas.

Así que aquí estamos: la teoría cuántica tiene tres tipos principales de interacción —las «interacciones fuertes» de quarks y gluones, las «interacciones débiles» de los W y las «interacciones eléctricas» de los fotones—. Las únicas partículas del mundo (de acuerdo con este esquema) son los quarks (en «aroma» u y d , con tres «colores» cada uno), los gluones (ocho combinaciones de R , G y B), los W (cargados ± 1 y 0), los neutrinos, electrones y fotones —alrededor de veinte partículas diferentes de seis tipos distintos (más sus antipartículas)—. No está mal —alrededor de veinte partículas diferentes— excepto que esto no es todo.

Al hacer incidir sobre los núcleos protones de energía cada vez más alta, siguieron apareciendo nuevas partículas. Una de ellas fue el muón, que es en todos los aspectos idéntico al electrón, salvo que su masa es mucho más elevada —105,8 MeV comparada con 0,511 para el electrón, o alrededor de 206 veces más pesado—. ¡Es como si Dios quisiera probar un número distinto para la masa! Todas las propiedades del muón se pueden describir completamente por la teoría de la electrodinámica —la constante de

acoplamiento j es la misma y $E(A \text{ a } B)$ también; sólo se necesita poner un valor distinto de n ^[31].

Dado que el muón tiene una masa aproximadamente 200 veces más grande que la del electrón, la «manecilla del cronógrafo» para un muón gira 200 veces más deprisa que para un electrón. Esto nos ha permitido verificar si la electrodinámica todavía se comporta como establece la teoría a distancias 200 veces más pequeñas de lo que habíamos sido capaces antes de probar — aunque estas distancias son todavía más de ocho cifras decimales superiores a las distancias a las que la teoría puede tener problemas debido a los infinitos.

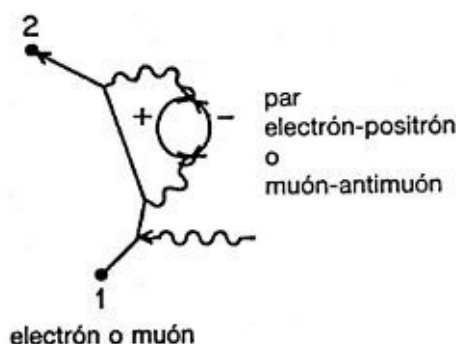


FIGURA 89. En el proceso de bombardear núcleos con protones de energías más y más altas, aparecen nuevas partículas. Una de ellas es el muón, o electrón pesado. La teoría que describe las interacciones del muón es exactamente la misma que para el electrón, excepto que se introduce un número más elevado para n en $E(A \text{ a } B)$. El momento magnético de un muón debería ser ligeramente diferente del electrón a dos alternativas especiales: cuando el electrón emite un fotón que se desintegra en un par electrón-positrón o muón-antimuón, la desintegración crea un par que tiene una masa próxima o muy superior a la del electrón. Por otro lado, cuando el muón emite un fotón que se desintegra en un par muón-antimuón o positrón-electrón, este par tiene una masa próxima o mucho más ligera que la masa del muón. Los experimentos confirman esta ligera diferencia.

Veamos si el muón puede estar implicado en un proceso radiactivo como la desintegración beta: cuando un quark d cambia a un quark u emitiendo un W , ¿puede el W acoplarse con un muón en lugar de con un electrón? Sí (ver Fig. 90). ¿Y qué ocurre con el anti-neutrino? En el caso de acoplamiento del W con un muón, una partícula denominada neutrino-mu ocupa el lugar del neutrino ordinario (que denominaremos ahora un neutrino electrónico). De modo que nuestra tabla de partículas tiene dos partículas adicionales próximas al electrón y al neutrino —el muón y el neutrino-mu.

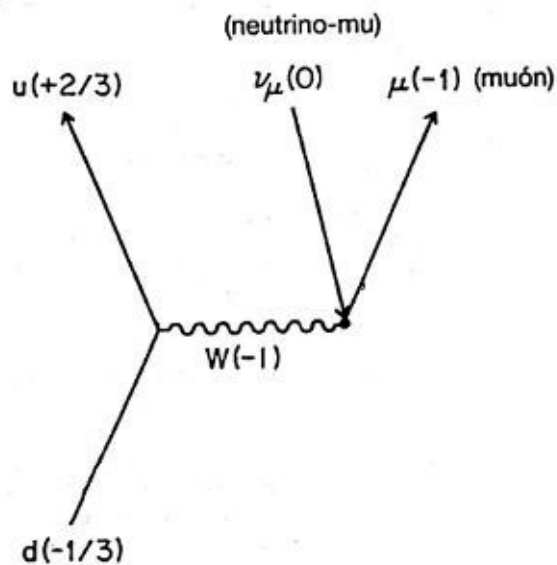


FIGURA 90. El W tiene una cierta amplitud de emisión de un muón en lugar de un electrón. En este caso un neutrino-mu reemplaza a un neutrino electrónico.

¿Y que ocurre con los quarks? Enseguida se supo que las partículas tenían que estar formadas por quarks más pesados que el u o el d . Así, se incluyó un tercer quark, denominado s (por «extraño»^[32] (+)) en la lista de partículas elementales. El quark s tiene una masa de unos 200 MeV, comparada con unos 20 MeV para los quarks u y d .

Durante muchos años pensamos que había tres «aromas» de quarks — u , d y s — pero, en 1974, se descubrió una nueva partícula denominada mesón-pi que no podía estar constituida por tres quarks. Existía también una buena razón teórica de que debía de haber un cuarto quark, acoplado al quark s mediante un W de la misma manera en que se acoplan un quark u y otro d (ver Fig. 91). El «aroma» de este quark se denominó c y yo no tengo agallas para explicarles de dónde proviene la c , pero puede que lo hayan leído en los periódicos^[33]. ¡Los nombres van empeorando cada vez más!

The diagram illustrates the classification of particles by spin and their interactions. It features a table with columns for particle types and their couplings to spin 1/2 and spin 1 particles. A separate box details the properties of an electron.

		partículas de espín 1/2		partículas de espín 1		
		muón μ	electrón e	fotón	gluón	W
		105.8	.511	0	0	~80,000
neutrino ν_μ	0	0	0	-1	0	
neutrino ν_e	0	0	0	0	0	
quark s	~200	-1/3	g			
quark d	~20					
quark c	~1800	+2/3	g			
quark u	~20					

acoplos

nombre — electrón

símbolo — e

masa (MeV) — .511

FIGURA 91. La Naturaleza parece estar repitiendo las partículas de espín 1/2. Además del muón y del neutrino-mu, existen dos nuevos —s y c— que tienen la misma carga pero mayores masas que sus contrapartidas en la siguiente columna.

Esta repetición de partículas con las mismas propiedades pero masas más pesadas, es un completo misterio. ¿Qué es esta extraña duplicidad del esquema? Como el profesor I. I. Rabi dijo cuando se descubrió el muón «¿Quién lo ha ordenado?».

Recientemente ha comenzado otra repetición de la lista. Al ir a energías cada vez más altas, la Naturaleza parece continuar apilando estas partículas como para drogarnos. Tengo que hablarles de ellas porque quiero que vean lo aparentemente complejo que el mundo parece en realidad. Sería muy engañoso si les diese la impresión de que puesto que hemos resuelto el 99% de los fenómenos del mundo mediante electrones y fotones, ¡el otro 1% de los fenómenos requiriese sólo el 1% de partículas adicionales! Resulta que para explicar este 1% final, necesitamos un número diez o veinte veces mayor de partículas adicionales.

De modo que empezamos de nuevo: utilizando energías aún más elevadas en los experimentos, se ha encontrado un electrón todavía más pesado, denominado «tau»; tiene una masa de alrededor de 1800 MeV; ¡tan pesado como dos protones! También se ha inferido un neutrino-tau. Y ahora se ha encontrado una curiosa partícula que implica un nuevo «aroma» para los quarks —esta vez es «b» de «belleza» y tiene una carga de $-1/3$ (ver Fig. 92) —. Bien, ahora, por un momento, les quiero físicos teóricos fundamentales de

primera clase y que predigan algo: se encontrará un nuevo aroma para los quarks, denominado... (de «...», con una carga de ..., una masa de ... MeV) —¡y ciertamente esperamos que sea *verdad* que exista!^[34]

partículas de espín 1/2			partículas de espín 1		
			fotón	gluón	W
tau τ ~1860	muon μ 105.8	electron e .511	0	0	~80,000
neutrino ν_τ 0	neutrino ν_μ 0	neutrino ν_e 0	-1	0	
quark b ~4800	quark s ~200	quark d ~20	0	0	
	quark c ~1800	quark u ~20	-1/3	g	
			+2/3	g	

acoplos

FIGURA 92. ¡Aquí estamos de nuevo! Otra repetición de las partículas de espín 1/2 se ha iniciado a energías aún más elevadas. Esta repetición será completa si se encuentra una partícula con las propiedades adecuadas para implicar la existencia de un nuevo aroma para el quark. Mientras tanto, se están realizando los preparativos para buscar los principios de otra repetición a energías todavía más elevadas. La causa de estas repeticiones es un completo misterio.

Mientras tanto, se están realizando experimentos para ver si el círculo se repite una vez más. En la actualidad se están construyendo máquinas para buscar un electrón aún más pesado que el tau. Si la masa de esta supuesta partícula es 100 000 MeV, no van a ser capaces de producirla. Si está alrededor de 40 000 MeV podrán hacerlo.

Misterios como estos ciclos que se repiten hacen que sea muy interesante ser físico teórico: ¡La Naturaleza nos proporciona unos rompecabezas tan maravillosos! ¿Por qué repite la Naturaleza el electrón, con masas 206 y 3640 veces más grandes?

Me gustaría hacer un último comentario para poder completar totalmente lo referente a las partículas. Cuando un quark *d* acoplándose con un W cambia a un quark *u*, tiene a su vez una pequeña amplitud de cambiar, en su lugar, a un quark *c*. Cuando un quark *u* se transforma en otro *d*, también tiene una pequeña amplitud de convertirse en un quark *s* y otra amplitud aún menor de cambiar a un quark *b* (ver Fig. 93). De modo que el W «tensa» un poco las cosas y permite a los quarks estas proporciones relativas en su amplitud para poder cambiar a otro tipo de quark, es algo completamente desconocido.

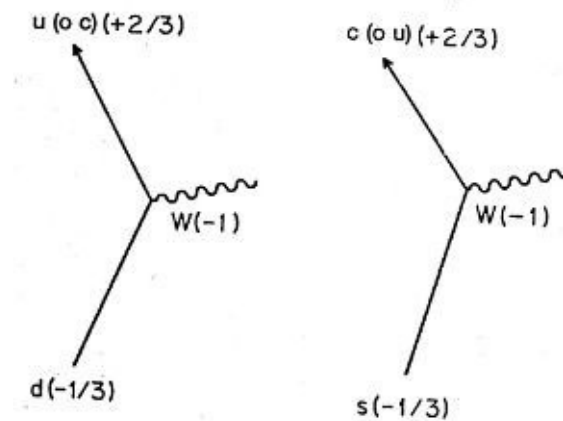


FIGURA 93. Un quark d tiene una pequeña amplitud de cambiar a un quark c en lugar de a un quark u , y un quark s tiene una pequeña amplitud de cambiar a un quark u en lugar de a un quark c , con la emisión en ambos casos de un W . De modo que el W parece ser capaz de cambiar el aroma de un quark de una a otra columna de la tabla (ver Fig. 92).

Así que, aquí está todo sobre el resto de la física cuántica. Es un lío terrible, y podrían decir que la física se ha convertido en un barullo descorazonador. Pero siempre ha tenido esta apariencia. La Naturaleza ha parecido un barullo horrible, pero al avanzar vemos esquemas y aunamos teorías; se aclaran un poco las cosas y se vuelven más sencillas. El lío que les he expuesto es mucho menor que el que hubiese tenido que contar hace diez años, hablándoles de más de un centenar de partículas. Y piensen en el lío de principios de siglo, cuando estaban el calor, magnetismo, la electricidad, la luz, los rayos-X, los rayos ultravioletas, los índices de refracción, los coeficientes de reflexión y otras propiedades de diversas sustancias, cosas todas que se han unificado desde entonces en una teoría, la electrodinámica cuántica.

Me gustaría resaltar algo. Las teorías del resto de la física son muy similares a la teoría de la electrodinámica cuántica: todas implican la interacción de objetos de espín $1/2$ (como electrones y quarks) con objetos de espín 1 (como fotones, gluones o W) dentro de un marco de amplitudes en donde la probabilidad de un suceso es el cuadrado de la longitud de una flecha. ¿Por qué todas las teorías físicas son tan similares en estructura?

Existen un cierto número de posibilidades. La primera, la limitada imaginación de los físicos: cuando observamos un nuevo fenómeno tratamos de encajarlo en el marco que ya tenemos —hasta que no hemos realizado un número suficiente de experimentos, no sabemos que eso no funciona—. De modo que cuando algún físico loco da una conferencia en UCLA en 1983 y dice «Esta es la manera en que ocurren las cosas y miren cuán maravillosamente similares son las teorías», no es porque la Naturaleza sea

realmente similar, es porque los físicos sólo han sido capaces de pensar la condenada misma cosa una y otra vez.

Otra posibilidad es que *sea* la condenada misma cosa una y otra vez —que la Naturaleza sólo tenga una forma de realizar las cosas y que repita su historia de vez en cuando.

Una tercera posibilidad es que las cosas parecen similares porque son aspectos de una misma cosa —una amplia imagen subyacente de la que se pueden desgajar partes que parecen diferentes, análogo a los dedos de una mano—. Muchos físicos están trabajando muy duro tratando de componer una gran imagen que unifique todo en un supermodelo. Es un juego delicioso, pero en la actualidad ninguno de los especuladores coincide con cualquier otro especulador en cuanto a cuál es la gran imagen. Exagero sólo ligeramente cuando digo que la mayoría de estas teorías especulativas no tienen mayor sentido que la suposición que hicieron Vds. sobre la posibilidad de un quark t , ¡y les garantizó que no son mejores que Vds. a la hora de suponer la masa del quark t !

Por ejemplo, parece que el electrón, el neutrino, el quark d , y el quark u van todos juntos —de hecho, los dos primeros se acoplan con el W , como hacen los dos últimos—. En la actualidad se piensa que un quark sólo puede cambiar «colores» o «aromas». Pero quizás un quark podría desintegrarse en un neutrino acoplándose con una partícula desconocida. Bonita idea. ¿Qué ocurriría? Significaría que los protones son inestables.

Alguien construye una teoría: los protones son inestables. ¡Hacen un cálculo y encuentran que ya no existirían protones en el Universo! De manera que juegan con los números, ponen una masa más elevada a la nueva partícula y después de mucho esfuerzo predicen que el protón decaerá a una velocidad ligeramente menor que la medida la última vez cuando se demostró que el protón no decaía.

Cuando aparece un nuevo experimento y mide con más cuidado el protón, las teorías se ajustan para escaparse de la presión. El experimento más reciente mostró que el protón no decae a una velocidad cinco veces menor que la que se predijo en la *última oleada* de teorías. ¿Qué piensan que ocurrió? El fénix surgió de nuevo con una nueva modificación de la teoría que requiere experimentos aún más precisos para comprobarla. Si el protón decae o no es algo desconocido. Demostrar que no decae es muy difícil.

A lo largo de estas conferencias no he discutido la gravitación. La razón es

que la influencia gravitacional entre objetos es *extremadamente* pequeña: es una fuerza que es un 1 seguido de 40 ceros más débil que la fuerza eléctrica entre dos electrones (quizá son 41 ceros). En la materia, casi todas las fuerzas eléctricas se emplean en mantener a los electrones próximos a los núcleos de sus átomos, creando un fino equilibrio mezcla de masas y menos que se cancelan entre sí. Pero con la gravitación, la única fuerza es la atracción y crece y crece según hay más y más átomos hasta que, al final, cuando obtenemos esas grandes masas ponderables que somos, empezamos a poder medir los efectos de la gravedad —sobre los planetas, sobre nosotros mismos y así sucesivamente.

Dado que la fuerza gravitacional es mucho más débil que cualquiera otra de las interacciones, es imposible, en la actualidad, hacer cualquier experimento que sea lo suficientemente delicado como para medir cualquier efecto que requiera la precisión de una teoría cuántica de la gravitación para explicarlo^[35]. Pero aunque no haya manera de probarlas, existen, sin embargo, teorías cuánticas de la gravedad que implican «gravitones» (que podrían aparecer bajo una nueva categoría de polarizaciones, denominada «espín 2») y otras partículas fundamentales (algunas con espín 3/2). La mejor de estas teorías no es capaz de incluir las partículas que encontramos, inventando, no obstante, muchas partículas que no se encuentran. Las teorías cuánticas de la gravitación también tienen infinitos en los términos de acoplamiento, pero el «proceso de profundización», que tanto éxito tiene eliminando los infinitos en la electrodinámica cuántica, no puede eliminarlos en la gravitación. De modo que no sólo no tenemos experimentos con los que probar la teoría cuántica de la gravitación, sino que tampoco tenemos una teoría razonable.

A lo largo de la totalidad de esta historia ha permanecido una característica particularmente insatisfactoria: las masas observadas de las partículas, m . No existe teoría que explique adecuadamente estos números. Utilizamos los números en todas nuestras teorías, pero no los entendemos —ni lo que son, ni de dónde vienen—. Creo que desde un punto de vista *fundamentalista*, es un problema serio y muy interesante.

Lamento si toda esta especulación sobre las nuevas partículas les ha confundido, pero decidí completar mi discusión del resto de la física para demostrarles cómo el *carácter* de estas leyes —el marco de amplitudes, los diagramas que representan las interacciones que hay que calcular, y demás— parece ser el mismo que en la teoría de la electrodinámica cuántica, nuestro

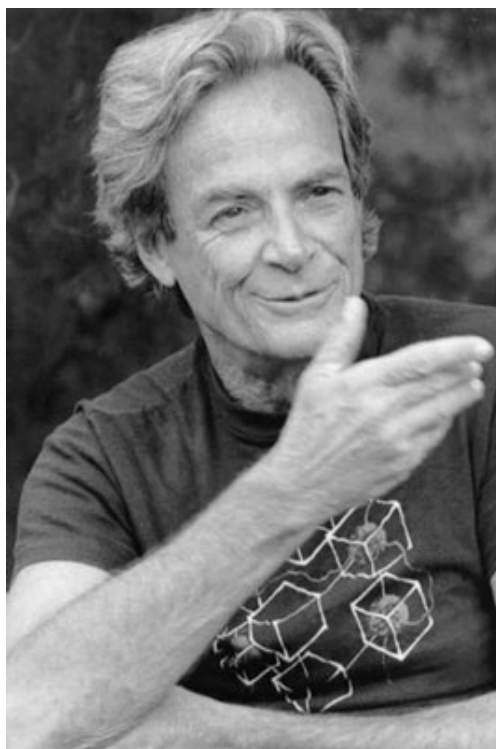
mejor ejemplo de una buena teoría.

Nota añadida al leer las pruebas, noviembre 1984

Desde que se dieron estas conferencias, sospechosos sucesos observados en los experimentos, parecen apuntar que, acaso, sea posible que se descubra pronto otra nueva e inesperada partícula o fenómeno (y en consecuencia no mencionado en estas conferencias).

Nota añadida al leer las pruebas, abril 1985

En este momento, los «sospechosos sucesos» mencionados anteriormente parecen ser una falsa alarma. La situación sin duda habrá cambiado de nuevo en la época en que uds. lean este libro. Las cosas cambian con mayor rapidez en la física que en el negocio de publicación de libros.



RICHARD PHILIPS FEYNMAN (Nueva York, Estados Unidos, 11 de mayo de 1918 - Los Ángeles, California, Estados Unidos, 15 de febrero de 1988). Físico teórico estadounidense. Revisó todo el panorama de la electrodinámica cuántica, y revolucionó el modo en que la ciencia entendía la naturaleza de las ondas y las partículas elementales. En 1965 compartió el Premio Nobel de Física con el estadounidense Julian S. Schwinger y el japonés Tomonaga Shinichiro, científicos que de forma independiente desarrollaron teorías análogas a la de Feynman, aunque la labor de este último destaca por su originalidad y alcance. Las herramientas que ideó para resolver los problemas que se le plantearon, como, por ejemplo, las representaciones gráficas de las interacciones entre partículas conocidas como diagramas de Feynman, o las denominadas integrales de Feynman, permitieron el avance en muchas áreas de la física teórica a lo largo del período iniciado tras la Segunda Guerra Mundial. Descendiente de judíos rusos y polacos, estudió física en el Instituto Tecnológico de Massachusetts, y se doctoró luego en la Universidad de Princeton, donde colaboró en el desarrollo de la física atómica entre 1941 y 1942. Los tres años siguientes lideró el grupo de jóvenes físicos teóricos que colaboraron en el Proyecto Manhattan en el laboratorio secreto de Los Álamos, bajo la dirección de Hans Bethe.

En los años cincuenta justificó, desde el punto de vista de la mecánica cuántica, la teoría macroscópica del físico soviético L. D. Landau, que daba explicación al estado superfluido del helio líquido a temperaturas cercanas al

cero absoluto, estado caracterizado por la extraña ausencia de fuerzas de rozamiento.

En 1968 trabajó en el acelerador de partículas de Stanford, período en el que introdujo la teoría de los partones, hipotéticas partículas localizadas en el núcleo atómico, que daría pie más tarde a la introducción del moderno concepto de quark. Su aportación a la física teórica ha quedado recogida en títulos tales como *Quantum electrodynamics* (1961) y *The theory of fundamental processes* (1961).

Notas

[1] ¿Cómo lo supo? Newton fue un gran hombre: escribió, «Porque puedo pulir un cristal». Pueden preguntarse ¿cómo diablos podía decir que porque podía pulimentarlo no podían existir agujeros y manchas? Newton pulía sus propias lentes y espejos y sabía lo que estaba haciendo al pulir: estaba rayando la superficie de un cristal con polvos cada vez más finos. Según se hacían más finas las rayaduras, la superficie del cristal cambiaba su aspecto desde un gris deslustrado (debido a que la luz era difundida por las grandes rayas) a una claridad transparente (porque las extremadamente finas rayaduras dejaban pasar la luz). Así que vio que era imposible aceptar la suposición de que la luz pudiese verse afectada por irregularidades muy pequeñas tales como arañazos o agujeros y manchas; de hecho, encontró que lo contrario era lo correcto. Las rayaduras más finas y por tanto las manchas igualmente minúsculas, no afectan a la luz. Por lo que la teoría de agujeros y manchas no es adecuada. <<

[2] Resulta muy afortunado para nosotros el que Newton se convenciese de que la luz son «corpúsculos» porque podemos ver por lo que tuvo pasar una mente vigorosa e inteligente al considerar este fenómeno de la reflexión parcial por dos o más superficies y tratar de explicarlo. (Aquellos que creían que la luz eran ondas nunca tuvieron que pelear con ello). Newton argumentaba lo siguiente: Aunque la luz parece reflejarse por la primera superficie, no puede ser así. Si lo fuese, ¿cómo podría entonces ser capturada de nuevo la luz reflejada por la primera superficie cuando el espesor fuese tal que se supone que no existe reflexión alguna? Por consiguiente, la luz debe reflejarse en la segunda superficie. Pero para explicar el hecho de que el espesor del cristal determina el valor de la reflexión parcial, Newton propuso esta idea: la luz que llega a la primera superficie establece una especie de onda o campo que viaja junto con la luz y la predispone a reflejarse o no en la segunda superficie. Denominó a este proceso, que ocurre en ciclos dependiendo del espesor del cristal, «accesos de fácil reflexión o de fácil transmisión».

Existen dos dificultades en esta idea: la primera es el efecto de las superficies adicionales que he descrito en el Texto —cada nueva superficie afecta a la reflexión—. El otro problema es que la luz ciertamente se refleja en la superficie de un lago, que no tiene una segunda superficie, luego la luz *debe* reflejarse en la superficie frontal. En el caso de superficies únicas, Newton decía que la luz tenía una predisposición a reflejarse. ¿Podemos tener una teoría en la que la luz sabe el tipo de superficie en la que está incidiendo y si es la única superficie? Newton no resaltó esas dificultades de su teoría de «accesos de reflexión y transmisión» aunque está claro que sabía que su teoría no era satisfactoria. En la época de Newton, las dificultades de una teoría se discutían con brevedad y se encubrían —un estilo muy distinto del que se usa hoy en la ciencia, en donde se señalan los puntos donde nuestra propia teoría no se ajusta a las observaciones experimentales—. No estoy diciendo nada en contra de Newton; sólo quiero decir algo a favor de cómo nos comunicamos entre nosotros en la ciencia hoy en día. <<

[3] Esta idea utiliza el hecho de que las ondas se pueden combinar o cancelar, y los cálculos basados en este modelo se ajustaban a los resultados de los experimentos de Newton, así como a los realizados durante centenares de años posteriores. Pero cuando se desarrollaron instrumentos suficientemente sensibles como para detectar un único fotón, la teoría ondulatoria predecía que los «clicks» del fotomultiplicador debían de hacerse menos y menos audibles, mientras que lo que ocurría era que mantenían la misma intensidad —simplemente se sucedían con frecuencia cada vez menor—. No existía modelo razonable que pudiese explicar este hecho, por lo que hubo un período durante el cual uno debía ser listo: Había que saber qué experimento se estaba analizando para decir si la luz eran ondas o partículas. Este estado de confusión se denominó la «dualidad onda-partícula» de la luz, y había quién decía de manera chistosa que la luz era ondas los Lunes, Miércoles y Viernes; partículas los Martes, Jueves y Sábados, y los Domingos, ¡había que meditarlo! El propósito de estas conferencias es decirles como se «resolvió» finalmente el rompecabezas. <<

[4] Las zonas del espejo cuyas flechas señalan, en general, hacia la izquierda, también dan lugar a una reflexión fuerte (cuando se eliminan las zonas cuyas flechas señalan en sentido opuesto). Si tanto las zonas con flechas con tendencia hacia la izquierda como aquellas con tendencia hacia la derecha reflejan a la vez, entonces la contribución se cancela. Esto es análogo al caso de la reflexión parcial por dos superficies: mientras que cada superficie es capaz de reflejar por sí misma, si el espesor es tal que las dos superficies contribuyen con flechas señalando en sentidos opuestos, la reflexión se anula.

<<

[5] No puedo resistir sin hablarles de una red de difracción creada por la Naturaleza: los cristales de sal son átomos de sodio y cloro empaquetados con una distribución regular. Su distribución alternada, al igual que nuestra superficie de surcos, actúa como una rejilla cuando la luz del color adecuado (rayos-X en este caso) incide sobre ella. Buscando los lugares específicos donde un detector recoge muchas de estas reflexiones especiales (llamadas difracción), se puede determinar exactamente la distancia entre surcos y en consecuencia la separación entre los átomos (ver Fig. 28). Es una forma preciosa de determinar la estructura de todo tipo de cristales, así como de confirmar que los rayos-X son la misma cosa que la luz. Estos experimentos se realizaron por primera vez en 1914. Fue muy excitante ver, en detalle, por vez primera cómo están empaquetados los átomos en sustancias diferentes.

<<

[6] Este es un ejemplo del «principio de incertidumbre»: existe una especie de «complementariedad» entre el conocimiento de por dónde va la luz entre los bloques y por dónde va después —el conocimiento preciso de ambos es imposible—. Me gustaría situar el principio de incertidumbre en su contexto histórico: Cuando las ideas revolucionarias de la física cuántica comenzaron a llegar, la gente intentaba todavía entenderlas en términos de las viejas ideas pasadas de moda (tales como que la luz viaja en línea recta). Pero en determinado momento las viejas ideas empezaban a fallar, de manera que se ideó una advertencia que decía «Sus viejas ideas son condenadamente malas cuando...». Si se deshacen de todas las ideas pasadas de moda y en su lugar utilizan las ideas que les estoy explicando en estas conferencias —sumando *flechas* para todos los caminos en que un suceso puede tener lugar— ¡no hay necesidad de un principio de incertidumbre! <<

[7] Los matemáticos han tratado de encontrar todos los objetos que puedan existir que obedezcan a las leyes del álgebra ($A + B = B + A$, $A \times B = B \times A$, etc.). Las reglas se hicieron originariamente para los enteros positivos, utilizados para contar cosas como manzanas o personas. Los números se mejoraron tras la invención del cero, fracciones, números irracionales — números que no se pueden explicar como cociente de dos enteros— y números negativos, y continuaron obedeciendo las reglas originales del álgebra. Algunos de los números que inventaron los matemáticos supusieron, al principio, dificultades para las personas —la idea de media persona era difícil de imaginar— pero hoy no existe ninguna dificultad: nadie tiene escrúpulos morales o sentimientos sangrientos incómodos cuando oye que hay una media de 3,2 personas por milla cuadrada en algunas regiones. No intentan imaginar la persona 0,2, en su lugar saben que 3,2 significa que si multiplican 3,2 por 10, obtienen 32. De modo que, algunas cosas, que satisfacen las reglas del álgebra, pueden ser interesantes para los matemáticos incluso aunque no siempre representen la situación real. Las flechas en un plano pueden «sumarse» colocando la cabeza de una flecha sobre la cola de otra, o «multiplicarse» mediante reducciones y giros sucesivos. Puesto que estas flechas obedecen las mismas reglas del álgebra que los números regulares, los matemáticos las denominan números. Pero para distinguirlos de los números ordinarios, los llaman «números complejos». Para aquellos de Vds. que hayan estudiado las matemáticas suficientes para haber llegado a los «números complejos» les podría haber dicho «la probabilidad de un suceso es el valor absoluto del cuadrado de un número complejo. Cuando un suceso puede ocurrir por caminos alternativos, sumen los números complejos; cuando pueda ocurrir sólo por una sucesión de pasos, multipliquen los números complejos». Aunque puede resultar más impresionante de esta manera, no he dicho nada que no dijese antes —sólo he utilizado un lenguaje diferente—. <<

[8] Habrán notado que cambiamos 0,0384 por 0,04 y utilizado 84% como el cuadrado de 0,92 a fin de conseguir el 100% de la luz considerada. Pero cuando se suma *todo*, 0,0384 y 84% no tienen por qué redondearse —todos los trocitos de flecha (representando todos los caminos en que puede ir la luz) se compensan entre sí y dan la respuesta correcta—. Para aquellos de Vds. que gusten de este tipo de cosas, aquí va un ejemplo de otro camino por el que puede viajar la luz desde el detector hasta A —una serie de tres reflexiones (y dos transmisiones), que resultan en una flecha final de longitud $0,98 \times 0,2 \times 0,2 \times 0,2 \times 0,98$ o alrededor de 0,008 —una flecha muy diminuta (ver Fig. 46). Para realizar un cálculo completo de la reflexión parcial por dos superficies, tendrían que añadir esta pequeña flecha, más otra más pequeña aún que representa cinco reflexiones, etc. <<

[9] Esta regla verifica lo que nos enseñan en la escuela —la cantidad de luz transmitida a una cierta distancia varía inversamente con el cuadrado de la distancia— porque una flecha que reduce su tamaño original a la mitad, tiene un cuadrado de valor un cuarto de su longitud inicial. <<

[10] Este fenómeno, denominado el efecto Hanbury-Brown-Twiss, ha sido utilizado para distinguir una fuente única de ondas de radio de otra doble en las profundidades del espacio, incluso cuando las dos fuentes se encontraban extremadamente próximas. <<

[11] Mantener este principio en mente debería de ayudar al estudiante a evitar confusiones con cosas como la «reducción de un paquete de ondas» y magias similares. <<

[12] La historia completa de esta situación es muy interesante: si los detectores en A y B no son perfectos, y detectan los fotones sólo en *algunas* ocasiones, hay *tres* condiciones finales distinguibles: 1) los detectores en A y D se disparan; 2) los detectores en B y D se disparan, y 3) el detector en D se dispara con A y B inalterados (han permanecido en su estado inicial). Las posibilidades para los dos primeros sucesos se calculan de la forma explicada anteriormente (excepto que existirá un paso más —una reducción por la probabilidad de que el detector en A [o en B] se dispare, puesto que los detectores no son perfectos—). Cuando D se dispara solo, no es posible separar ambos casos, y la Naturaleza juega con nosotros y causa interferencia —la misma respuesta peculiar que hubiésemos obtenido si no hubiese habido detectores (excepto que la flecha final se ha reducido en una amplitud equivalente a la de que los detectores *no* se disparen)—. El resultado final es una mezcla, la simple suma de los tres casos (ver Fig. 51). Al aumentar la fiabilidad de los detectores, obtenemos menos interferencia. <<

[13] En estas conferencias estoy dibujando la situación especial de un punto en una dimensión, a lo largo del eje X. Para situar un punto en el espacio tridimensional, se tiene que establecer una «habitación» y medir la distancia del punto hasta el suelo y a cada una de las paredes adyacentes (con ángulos rectos entre sí). Estas tres medidas se pueden llamar X_1 , Y_1 y Z_1 . La distancia real de este punto a un segundo punto con medidas X_2 , Y_2 , Z_2 se puede calcular utilizando un «teorema Pitagórico tridimensional»: el cuadrado de esta distancia real es

$$(X_2 - X_1)^2 + (Y_2 - Y_1)^2 + (Z_2 - Z_1)^2$$

A la diferencia entre *esto* y las diferencias de tiempos al cuadrado—

$$(X_2 - X_1)^2 + (Y_2 - Y_1)^2 + (Z_2 - Z_1)^2 - (T_2 - T_1)^2$$

—se le denomina a veces «Intervalo» o I , y es la combinación de la que, de acuerdo con la teoría de Einstein de la relatividad, debe depender $P(A \text{ a } B)$. La mayor contribución a la flecha final $P(A \text{ a } B)$ viene de donde se supone —de donde la diferencia en distancia igual a la diferencia en tiempo (es decir, cuando I es cero)—. Pero además existe una contribución cuando I no es cero, que es inversamente proporcional a I : señala a las 3 en punto cuando I es positivo (cuando la luz va más deprisa que c), y señala hacia las 9 en punto cuando I es negativo. Estas últimas contribuciones se cancelan en muchas circunstancias (ver Fig. 56). <<

[14] La fórmula para $E(A \text{ a } B)$ es complicada, pero hay una forma interesante de explicar cuánto vale. $E(A \text{ a } B)$ se puede representar como una suma gigantesca de un montón de caminos distintos por los que un electrón puede ir del punto A al punto B en el espacio-tiempo (ver Fig. 57): el electrón puede dar un «vuelo de un salto» yendo directamente de A a B; puede hacer un «vuelo de dos saltos» parando en un punto intermedio C, puede dar un «vuelo de tres saltos» parando en los puntos D y E, y así sucesivamente. En este análisis, la amplitud de cada «salto» —desde un punto F a otro G— es $P(F \text{ a } G)$, la misma amplitud que para un fotón que vaya de F a G. La amplitud de cada «parada» se representa por n^2 , siendo n el número que mencioné antes, el que usamos para que nuestros cálculos resulten correctos.

La fórmula para $E(A \text{ a } B)$ es entonces una serie de términos: $P(A \text{ a } B)$ [el «vuelo de un salto»] + $P(A \text{ a } C) \times n^2 \times P(C \text{ a } B)$ [el «vuelo de dos saltos», parando en C] + $P(A \text{ a } D) \times n^2 \times P(D \text{ a } E) \times n^2 \times P(E \text{ a } B)$ [el «vuelo de tres saltos», parando en D y E] + ... de *todos los posibles puntos intermedios* C, D, E y así sucesivamente.

Nótese que al aumentar n , los caminos indirectos contribuyen en mayor medida a la flecha final. Cuando n es cero (como para el fotón), todos los términos con n desaparecen (porque también ellos son iguales a cero), dejando sólo el primer término que es $P(A \text{ a } B)$. En consecuencia, $E(A \text{ a } B)$ y $P(A \text{ a } B)$ están íntimamente relacionados. <<

[15] Este número, la amplitud para emitir o absorber un fotón, se denomina a veces la «carga» de la partícula. <<

[16] Si hubiese incluido los efectos de la polarización del electrón, la flecha del «segundo camino» debería haberse «restado» —girado 180° y sumado—. (Más detalles sobre el tema aparecerán más adelante en esta conferencia). <<

[17] Las condiciones finales del experimento para estos caminos más complicados son las mismas que para los caminos más sencillos —los electrones situados inicialmente en los puntos 1 y 2 y acabando en los puntos 3 y 4— de modo que no podemos distinguir entre estas alternativas y las dos primeras. En consecuencia debemos sumar las flechas de estos dos caminos a los dos caminos considerados previamente. <<

[18] A un fotón intercambiado de esta manera, que en realidad nunca aparece en las condiciones iniciales o finales del experimento, se le denomina en ocasiones un «fotón virtual». <<

[19] Dirac propuso la existencia de «antielectrones» en 1931; el año siguiente, Carl Anderson los encontró experimentalmente y los llamó «positrones». Hoy, los positrones se generan con facilidad (por ejemplo, haciendo que dos fotones colisionen entre sí) y se mantienen durante semanas en un campo magnético. <<

[20] La amplitud para el intercambio de un fotón es $(-j) \times P(A - B) \times j$ —dos acoplamientos y la amplitud para que un fotón vaya de un sitio a otro—. La amplitud para que un protón se acople con un fotón es $-j$. <<

[21] El radio del arco evidentemente depende de la *longitud* de la flecha de cada sección, que viene determinada, en último lugar, por la amplitud S de que un electrón en un átomo del cristal difunda al fotón. Este radio se puede calcular usando las fórmulas de las tres acciones básicas para la multitud de intercambios de fotones involucrados y sumando las amplitudes. Es un problema muy difícil, pero este radio ha sido calculado para sustancias relativamente sencillas con considerable éxito, y las variaciones del radio, de sustancia a sustancia, se explican bastante bien utilizando estas ideas de la electrodinámica cuántica. Debe decirse, sin embargo, que nunca se ha realizado un cálculo directo, a partir de primeros principios, para una sustancia tan compleja como un cristal. En estos casos, el radio se determina experimentalmente. Para un cristal, se ha obtenido a partir de los experimentos un valor de aproximadamente 0,2 (cuando la luz incide directamente sobre el cristal en ángulo recto). <<

[22] Cada una de las flechas de la reflexión por una sección (que forman un «círculo») tiene la misma longitud que cada una de las flechas que hacen que la flecha final de la transmisión esté más girada. Por tanto, existe una relación entre la reflexión parcial de un material y su índice de refracción.

Parece que la flecha final es más grande que 1, lo que significa: ¡que sale más luz del cristal que la que entra! Parece así porque he despreciado las amplitudes para que un fotón fuese hacia abajo a otra sección, un nuevo fotón se difundiese hacia arriba en otra sección y luego un tercer fotón se difundiese hacia abajo a través del cristal —y otras posibilidades más complicadas—, que son las causantes de que las flechitas se curven manteniendo la longitud de la flecha final entre 0,92 y 1 (por consiguiente, la probabilidad total de que la luz sea reflejada o transmitida por la lámina del cristal es siempre del 100%). <<

[23] Otra forma de describir esta dificultad es decir que quizá la idea de que dos puntos pueden estar infinitamente próximos es errónea —la suposición de que podemos utilizar la geometría hasta el último recoveco es falsa—. Si hacemos la distancia mínima posible entre dos puntos tan pequeña como 10^{-100} cm (la distancia más pequeña involucrada en cualquier experimento, hoy en día, es de alrededor de 10^{-16} cm), los infinitos desaparecen, de acuerdo —pero surgen otras inconsistencias tales como que la probabilidad total de un suceso sume algo más o menos del 100%, o que obtengamos energías negativas en cantidades infinitesimales—. Se ha sugerido que estas consistencias surgen porque no se ha tenido en cuenta los efectos de la gravedad —que son normalmente muy, muy débiles, pero que se vuelven importantes a distancias del orden de 10^{-33} cm—. <<

[24] Aunque en los experimentos a alta energía surgen muchas partículas del núcleo, en los experimentos a baja energía —en condiciones más normales— se encuentra que los núcleos contienen sólo protones y neutrones. <<

[25] Un MeV es muy pequeño —apropiado para tales partículas— alrededor de $1,78 \times 10^{-27}$ gr. <<

[26] Nótese los nombres: «fotón» proviene de la palabra griega para luz; «electrón» viene del griego ámbar, el inicio de la electricidad. Pero según iba progresando la física moderna, los nombres de las partículas han mostrado un interés cada vez menor por el griego clásico, hasta llegar a inventar palabras como «gluones». ¿Pueden imaginar Vds. por qué se llaman «gluones»? De hecho, *d* y *u* representan palabras, pero no quiero confundirles, —un quark *d* no está más «abajo (down)» que está «arriba (up)» un quark *u*. Incidentalmente la *d*-nez o *u*-nez de un quark se llama su «aroma (flavor)». [También se utiliza la expresión «sabor» para *flavor*. *N. de la T.*] <<

^[27] R = red (rojo), G = green (verde), B = blue (azul). (*N. de la T.*) <<

[28] También se denominan: antiverde → magenta, antirrojo → cian, antiazul → amarillo. (*N. de la T.*) <<

[29] Después de haberse pronunciado estas conferencias, se lograron energías lo suficientemente elevadas como para producir un W aislado, encontrándose para su masa un valor muy próximo al valor predicho por la teoría. <<

[30] «glue» en original inglés. (*N. de la T.*) <<

[31] El momento magnético del muón se ha medido con mucha precisión —se ha encontrado ser 1,001165924 (con una incertidumbre de 9 en el último dígito) mientras que el valor para el electrón es 1,00115965221 (con una incertidumbre de 3 en el último dígito)—. Podrían sentir curiosidad de por qué el momento magnético del muón es ligeramente superior al del electrón. Uno de los diagramas que hemos dibujado tenía un electrón emitiendo un fotón que se desintegraba en un par positrón-electrón (ver Fig. 89). Existe también una pequeña amplitud de que el fotón emitido puede formar un par muón-antimuón que es más pesado que el electrón original. Esto es asimétrico, porque el muón emite un fotón, si ese fotón forma un par positrón-electrón, este par es *más ligero* que el muón original. La teoría de la electrodinámica cuántica describe con precisión *todas* las propiedades eléctricas del muón y del electrón. <<

[32] «strange» en el original inglés. (*N. de la T.*) <<

[33] De «charm» (encanto). (*N. de la T.*) <<

[³⁴] Desde que se impartieron estas conferencias, ha aparecido cierta evidencia de la existencia de un quark t con una masa elevada —alrededor de 40 000 MeV. [t de «true» y «truth» *N. de la T.*] <<

[35] Cuando Einstein y otros trataron de unificar la gravitación con la electrodinámica, ambas teorías eran aproximaciones clásicas. En otras palabras, estaban equivocadas. Ninguna de estas teorías tenía el marco de amplitudes que hemos encontrado tan necesario en la actualidad. <<