



62.01. / FÍSICA I

Óptica

Texto complementario

Dr. D. Kurlat – Ing. M. C. Menikheim

1) INTRODUCCION HISTORICA

Recién a principios del siglo XIX se impone la concepción de la luz como un fenómeno generado por las vibraciones de un medio elástico especial denominado éter. La teoría dominante hasta ese momento era la corpuscular. Esta consiste en considerar a la luz formada por corpúsculos regidos por las leyes de la Mecánica. Analicemos a través de este modelo el fenómeno de la refracción de la luz en una lámina de vidrio. En este fenómeno un rayo luminoso que incide sobre la superficie de separación de dos medios, penetra en el segundo (se refracta) desviándose de su dirección original. La desviación del rayo luminoso, se justifica según la teoría corpuscular, porque en la superficie de separación se ejercen sobre los corpúsculos fuerzas normales. Debido a estas fuerzas (ver F_1 y F_2 en la figura 1.1) las partículas son desviadas en su dirección primitiva y aceleradas al cambiar de medio. En la interfase 1-2 según se muestra en la figura 1.1 la partícula se acerca a la normal mientras que en la interfase 2-3 una fuerza de sentido opuesto a la anterior la aleja de la normal.

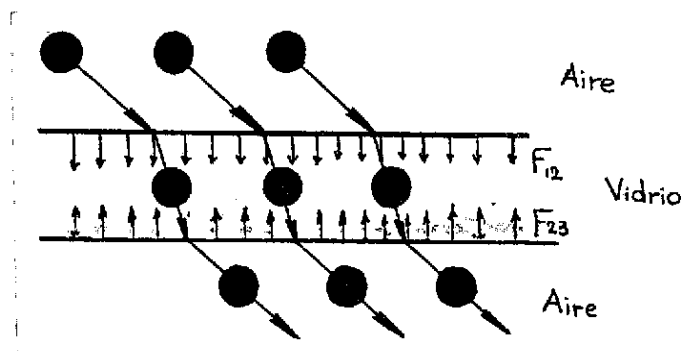


FIGURA 1.1 Explicación de la refracción de la luz según la teoría corpuscular.

En este modelo, los corpúsculos deben moverse con mayor velocidad en el vidrio que en el aire, como lo indican la dirección y el sentido de las fuerzas que actúan en las interfases.

El fenómeno de reflexión luminosa también se explica fácilmente admitiendo que al llegar a la interfase hay corpúsculos que chocan con la misma siguiendo las leyes que rigen el choque de partículas materiales. Como los ángulos de incidencia y de reflexión son iguales, el choque de los corpúsculos contra la interfase debe ser perfectamente elástico. (ver Figura 1.2).

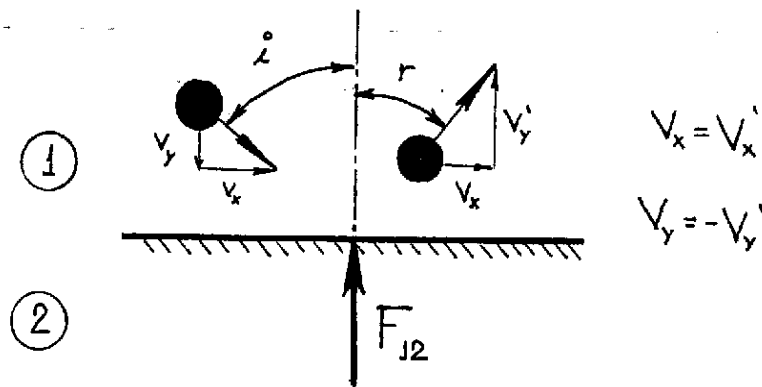


FIGURA 1.2

Entre los defensores de la Teoría Corpuscular se encontraba Isaac Newton, teórico inglés de tal prestigio que hizo incuestionable este modelo por varias generaciones.

Para explicar la dispersión de la luz en un prisma, Newton se vió forzado a introducir una hipótesis adicional a este modelo; las partículas que componen el rayo luminoso no son todas iguales y deben tener, según sus colores, masas diferentes.

Sin embargo, algunos contemporáneos de Newton opusieron algunas objeciones al modelo. Por ejemplo: es muy difícil explicar como una misma fuerza en la interfase de dos medios es capaz de reflejar y refractar a corpúsculos que pertenecen a un mismo haz incidente. Tampoco tienen explicación los fenómenos de birrefringencia y difracción. Pero el paradigma creado por Newton tenía tal fuerza que algunos partidarios de este modelo negaban la realidad de las experiencias que describían estos fenómenos.

El genial físico y matemático holandés Cristian Huygens (1629-1695), propuso un modelo alternativo que consideraba a la luz como vibraciones de un medio especial: el éter. Huygens comparaba las vibraciones del éter con el conocido fenómeno de la transmisión de la cantidad de movimiento en un conjunto de esferas colocadas en fila. (Ver figura 1.3)

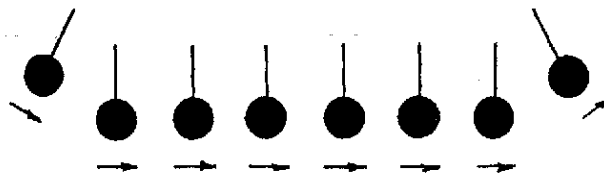


FIGURA 1.3

Al chocar una esfera exterior al sistema, con la primera de la fila, ésta le trasmite la cantidad de movimiento a la segunda y así sucesivamente. Las vibraciones del éter según Huygens

eran longitudinales. Como veremos más adelante, Huygens pudo explicar así, tanto la reflexión como la refracción de la luz y la doble refracción -birrefringencia-. Su teoría contiene, sin embargo, dos hipótesis que debían ser verificadas:

a) las vibraciones eran longitudinales.

b) las partículas del éter debían tener "dureza perfecta".

Pero unido a esta última hipótesis, las partículas del éter deben llenar el espacio y penetrar en diferentes medios. A esta capacidad de las partículas del éter Huygens la denominó "sutilidad".

En el marco de esta teoría se pueden explicar, sólo de forma muy intuitiva, los fenómenos de interferencia y difracción. Con excepción de Euler (1707-1783) la denominada teoría ondulatoria no fue aceptada por la mayoría de los científicos de fines de siglo XVII y XVIII.

A principios del siglo XIX el físico y médico inglés Thomas Young (1773-1829) publicó su célebre experiencia de la interferencia de la luz donde revitalizaba la teoría ondulatoria. Desgraciadamente, contradecir a Newton en aquel entonces era considerado, en Inglaterra, poco menos que una herejía, por lo que Young fue violentamente atacado, aún en forma personal. Más afortunado fue Agustín Jean Fresnel (1788-1827) quien se desempeñaba como ingeniero militar en el Servicio de Puentes y Caminos de Francia. A partir de 1810, desconociendo las publicaciones de Young, Fresnel comenzó a explicar sistemáticamente los hechos experimentales conocidos. Primero explicó la aberración luminosa y, posteriormante, introdujo la hipótesis de las oscilaciones transversales de las partículas del éter (hecho este último que contradecía las hipótesis de Huygens). Con esta suposición, explicó la polarización, fenómeno que había sido descubierto poco tiempo antes.

El éter imaginado por Fresnel era un medio complejo, pues debía tener las propiedades de un sólido (vibraciones transversales) y las de un fluido (sutilidad). Este carácter contradictorio del éter provocó resistencia en los medios científicos de la época. A pesar de ello, en 1818, Fresnel publicó una memoria en los "Comptes Rendús de l'Academie de Sciences de París" en las que predecía matemáticamente la difracción de la luz. La oposición de grandes sabios como Poisson y Biot, se desvaneció al confirmarse experimentalmente las predicciones de Fresnel. En 1819, la academia de ciencias de París premió el trabajo de Fresnel, pero faltaba una prueba experimental decisiva para optar entre la teoría ondulatoria o corpuscular.

A mediados de siglo, la célebre experiencia de Fizeau (1819-1896) en la cual se demostraba que la luz se propaga más rápidamente en el aire que en el agua, orientó la discusión en favor de la teoría ondulatoria. Finalmente en el año 1861, el gran físico inglés J. Clerk Maxwell (1831-1879) llevó a la teoría ondulatoria de la luz a un alto grado de desarrollo al demostrar que la luz no es más que un caso particular de vibración del campo electromagnético.

Con la teoría de Maxwell quedaron unificadas la electricidad, el magnetismo y la óptica bajo un mismo marco teórico. A partir de la publicación de los trabajos de Maxwell, se pensó como explicar las vibraciones del campo electromagnético tomando un modelo mecánico del éter. Se propusieron diversos modelos para este medio hipotético, y se

intentó medir la velocidad de la luz respecto a él.

Pero se presentaron nuevas dificultades, ya que el modelo era muy complejo y las experiencias de medición de la velocidad de la luz fracasaron (Michelson).

En 1905, Albert Einstein (1879-1955) propuso suprimir el éter como concepto teórico en su Teoría de la Relatividad Restringida y en el mismo año da una explicación del fenómeno fotoeléctrico que volvía (en parte) a la teoría corpuscular de Newton.

Durante los veinte años siguientes, la situación de la Física Teórica fue confusa y contradictoria, hasta que a partir de los trabajos de L. de Broglie, W. Heisenberg, M. Born, E. Schrodinger, P. Dirac y otros, se creó una nueva teoría denominada Mecánica Cuántica (o Mecánica Ondulatoria), según la cual la luz no tiene carácter definitivamente corpuscular u ondulatorio, sino que presenta una dualidad onda-partícula. La Mecánica Cuántica ocupa actualmente el rol más importante entre las teorías físicas, y su estudio lo abordarán los alumnos que cursen Física III A.

2) PRINCIPIO DE FERMAT

Una manera particular y muy interesante de estudiar el comportamiento de la luz fue propuesta por primera vez por Herón de Alejandría. Este postuló que al desplazarse desde un punto a otro, la luz sigue la trayectoria mas corta.

Tomemos el ejemplo de un rayo de luz que incide sobre un espejo y supongamos que no sabemos que trayectoria seguirá la luz al reflejarse. Consideremos una fuente luminosa A (ver figura 2.1) y dibujemos algunos de los posibles caminos para llegar de A hasta B después de la reflexión en la superficie del espejo.

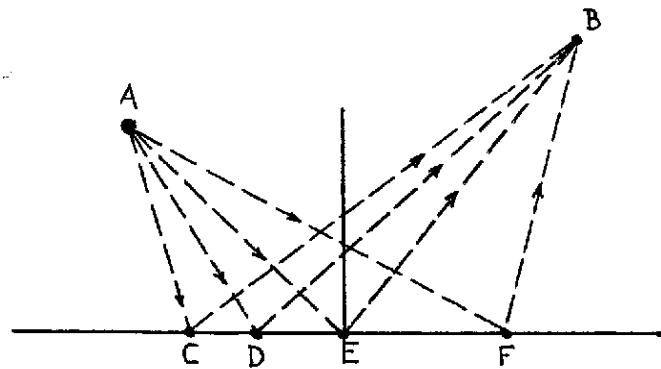


FIGURA 2.1 Posibles trayectorias de un rayo de luz para ir desde A hasta B.

Desde el punto de vista geométrico, todos los caminos (ACB, ADB, AEB, AFB) son teóricamente posibles. Como la trayectoria más corta es la AEB, donde el ángulo de incidencia es igual al ángulo de reflexión y esto se ve confirmado experimentalmente, Herón suponía la validez general de su principio. Sin embargo este principio no se cumple en el caso de la refracción (ver figura 2.2)

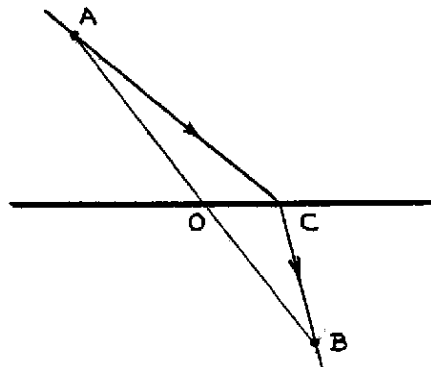


FIGURA 2.2 Refracción de la luz.

Como se puede ver en la figura 2.2, el camino ACB es mayor que el AOB (línea recta entre A y B).

Modificando el principio anterior, Fermat propuso que:

PRINCIPIO
DE FERMAT

LA TRAYECTORIA REAL DE LA LUZ ENTRE DOS PUNTOS
ES LA QUE SE RECORRE EN EL TIEMPO MINIMO.

¿Como se aplica concretamente el principio de Fermat? Tomemos el problema de la reflexión de la luz. Según se indica en la figura 2.3, supongamos que la luz emitida desde el punto A llega al espejo en el punto O (que en un primer análisis puede ser cualquiera).

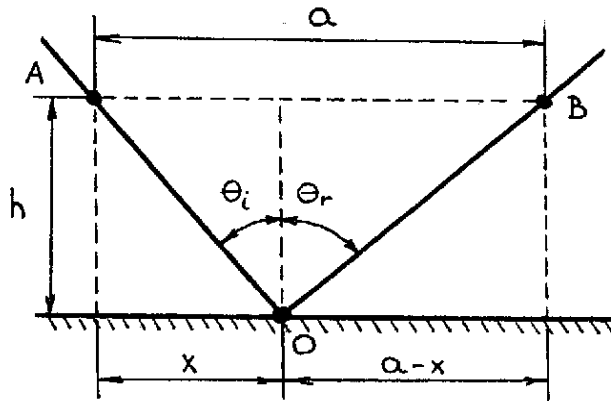


FIGURA 2.3 Ilustración del principio de Fermat para el caso de la reflexión de la luz.

¿Cómo será la relación entre el ángulo de incidencia (θ_i) y el de reflexión (θ_r)?

En términos del principio de Fermat, si los puntos A y B se encuentran entre sí a una distancia "a" y a una altura "h" del espejo como se ve en la figura anterior, tendremos que preguntarnos ¿cuál será el punto O del espejo en donde incide el rayo luminoso? o lo que es lo mismo ¿cuál será la longitud x para que la luz recorra el camino AOB en el tiempo mínimo?

Para responder a este interrogante tratemos de expresar en primer lugar la función $t=f(x)$. Teniendo en cuenta que la luz se propaga con movimiento uniforme, será:

$$(1) \quad t = |\text{desplazamiento}| / |\text{velocidad}|$$

$$\text{luego } t = AO/v + OB/v$$

Pero de acuerdo a la figura $AO = (x^2 + h^2)^{\frac{1}{2}}$ y $OB = ((a-x)^2 + h^2)^{\frac{1}{2}}$

$$(2) \quad t = \frac{(x^2 + h^2)^{\frac{1}{2}}}{v} + \frac{((a-x)^2 + h^2)^{\frac{1}{2}}}{v}$$

si aplicamos la condición de mínimo ($dt/dx=0$) en (2) tendremos la expresión (3):

$$\frac{dt}{dx} = \frac{1}{2} \cdot \frac{2 \cdot x}{(x^2+h^2)^{\frac{1}{2}}} \cdot \frac{1}{v} + \frac{1}{2} \cdot \frac{2 \cdot (a-x) \cdot (-1)}{((a-x)^2+h^2)^{\frac{1}{2}}} \cdot \frac{1}{v} = 0$$

de donde resulta

$$(4) \quad \frac{x}{(x^2+h^2)^{\frac{1}{2}}} - \frac{a-x}{((a-x)^2+h^2)^{\frac{1}{2}}} = 0$$

pero si observamos con cuidado la figura 2.3, el primer término antes mencionado es: $\sin[\theta_i]$ y el segundo $\sin[\theta_r]$, donde surge que: $\sin[\theta_i] = \sin[\theta_r]$ o lo que es lo mismo:

$$(5) \quad \boxed{\theta_i = \theta_r}$$

que constituye La Primera Ley de la Reflexión de la luz (ley empírica, sacada de la experimentación).

PREGUNTAS:

- En la deducción anterior, hay implícita una hipótesis con respecto a la velocidad de la luz, ¿cuál es?
- ¿La condición $dt/dx=0$, es de mínimo?. ¿No podría ser una condición de máximo?. Discuta.
- Demuestre, utilizando el Principio de Fermat, que el rayo incidente, la normal al espejo y el rayo reflejado están en el mismo plano.

De manera análoga podemos deducir las Leyes de la Refracción de la luz (Ley de Snell).

Supongamos que un haz de rayos luminosos que pasa por el punto A de la figura 2.4, incide en la superficie de separación de dos medios (medio 1 y medio 2) en el punto O y se refracta (pasa al medio 2) pasando por el punto B. Trazando la normal a la superficie de separación por O, queremos saber la relación que existe entre θ_i y θ_r (ver figura).

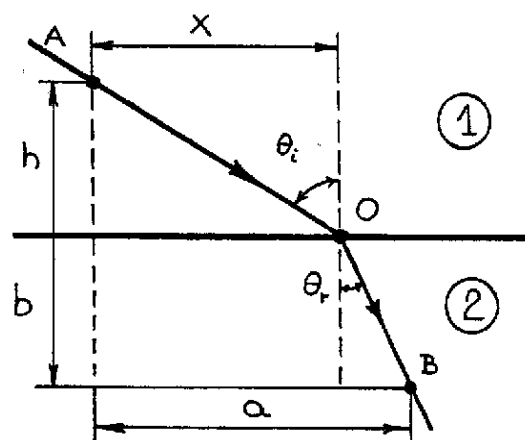


FIGURA 2.4

Tomaremos como hipótesis que al cambiar de medio cambia la velocidad de propagación, o sea, $v_1 \neq v_2$.

En este caso nuestra función $t=f(x)$ será: $t=AO/v_1+OB/v_2$

$$(6) \quad t = \frac{(x^2+h^2)^{\frac{1}{2}}}{v_1} + \frac{((a-x)^2+b^2)^{\frac{1}{2}}}{v_2}$$

buscamos nuevamente la condición de mínimo:

$$(7) \quad \frac{dt}{dx} = \frac{1}{2} \frac{2 \cdot x}{(x^2+h^2)^{\frac{1}{2}}} \cdot \frac{1}{v_1} + \frac{1}{2} \frac{2 \cdot (a-x) \cdot (-1)}{((a-x)^2+b^2)^{\frac{1}{2}}} \cdot \frac{1}{v_2}$$

pero $dt/dx=0$

$$(8) \quad \frac{\text{sen}[\theta_i]}{v_1} = \frac{\text{sen}[\theta_r]}{v_2}$$

pero como Snell estableció la condición que $V_i \cdot n_i = c$, donde c es la velocidad de propagación de la luz en el vacío, V_i es la velocidad de propagación en el medio i , n_i es el índice de refracción del medio i (suponiendo 1 el índice de refracción de la luz en el vacío) resulta ser: $V_1 \cdot n_1 = V_2 \cdot n_2$ ó lo que es lo mismo

$$(9) \quad \frac{V_1}{V_2} = \frac{n_2}{n_1} = \frac{\text{sen}(\theta_i)}{\text{sen}(\theta_r)} \quad \text{esto último hallado en (8)}$$

Ley de Snell

$$(10) \quad \boxed{\frac{V_1}{V_2} = \frac{\text{sen}(\theta_i)}{\text{sen}(\theta_r)} = \frac{n_2}{n_1} = n_{21}}$$

A los índices n_1 y n_2 se los denomina índices de refracción absolutos del medio correspondiente, y a n_{21} índice de refracción relativo del medio 2 (medio en el cual penetra el rayo) al medio 1 (medio del cual proviene el rayo).

Si en cambio el rayo de la luz atraviesa diferentes medios con distinto índice de refracción, la situación sería la siguiente:

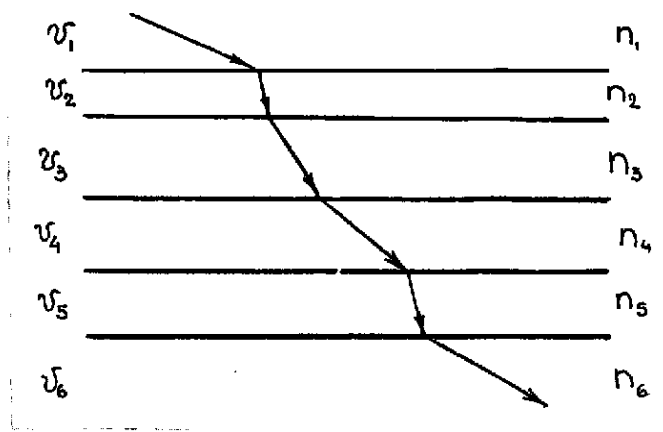


FIGURA 2.5

Al pasar de un medio a otro, el rayo luminoso cambia de trayectoria y de velocidad de propagación. Si llamamos t_i al tiempo que tarda la luz en recorrer el medio i , Si el camino recorrido en el medio i y V_i la velocidad de propagación de la luz en el medio i , entonces: $t_i = S_i / V_i$, y el tiempo total en el cual atravesaría los n medios:

$$(11) \quad t = \sum_{i=1}^n \frac{S_i}{V_i}$$

pero como $V_i = c/n_i$, reemplazando en (11)

$$t = \sum_{i=1}^n \frac{n_i \cdot S_i}{c} = \frac{1}{c} \cdot \sum_{i=1}^n n_i \cdot S_i$$

$$(12) \quad c \cdot t = \sum_{i=1}^n n_i \cdot S_i$$

El producto $c \cdot t$ es el camino que recorrería la luz en el vacío durante el tiempo que tarda en atravesar los n medios. A esta trayectoria se la denomina "camino óptico" y se lo simboliza C.O.

$$(12') \quad \text{C.O.} = \sum_{i=1}^n n_i \cdot S_i$$

Si en el medio atravesado por la luz el índice de refracción cambia de manera continua (como sucede por ejemplo con la luz del sol al atravesar las distintas capas de la atmósfera terrestre) la fórmula expresada en (12) y (12') se transforma en :

$$\text{C.O.} = \int_A^B n(s) \cdot ds$$

De acuerdo a lo propuesto por Fermat (ver página 6) el tiempo que tarda el rayo de luz es mínimo, y como su velocidad de propagación en el vacío es constante, por lo visto en (12'), el camino óptico debe ser mínimo.

2.1) ALGUNOS EJEMPLOS

El principio de Fermat nos permite explicar cualitativamente una serie de fenómenos:

Ejemplo 1: ¿Por qué vemos el sol en el horizonte cuando este se encuentra en realidad por debajo del mismo? (ver figura 2.6)

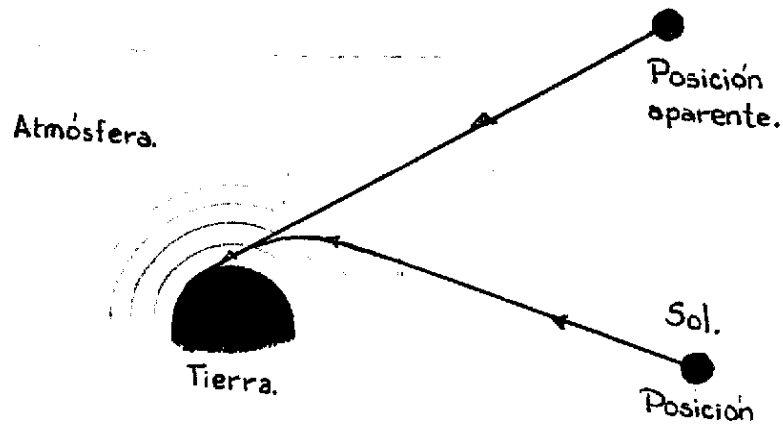


FIGURA 2.6 Deflexión de los rayos al atravesar la atmósfera.

La luz atraviesa más rápidamente las capas menos densas de la atmósfera, luego con el objeto de minimizar el camino óptico se desvía mas abruptamente en las capas más densas. (ver fig. 2.6).

Ejemplo 2: ¿Cómo se producen los espejismos?

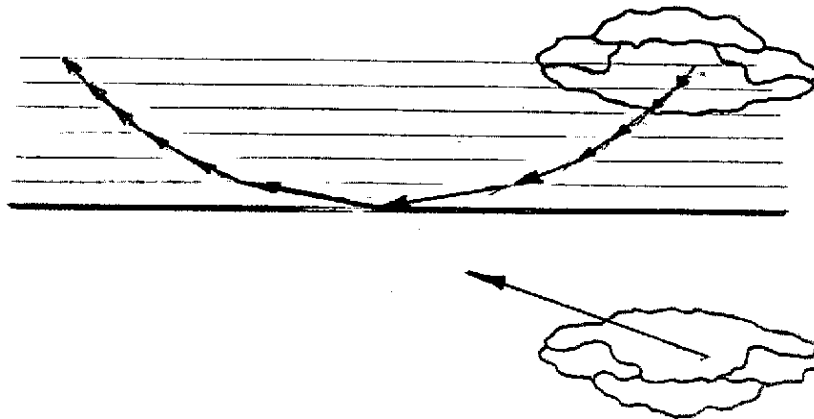


FIGURA 2.7

Cuando la temperatura ambiente es muy elevada, las capas de aire en contacto con el suelo se dilatan. De este modo, las capas menos densas son las que están más cerca del suelo. Es por ello que la luz realiza la mayor trayectoria posible en las capas inferiores y podemos ver como espejos de agua debido a la imagen invertida de los objetos que parecen estar en el piso (ver fig. 2.7).

2.2) VERSION "ACTUALIZADA" DEL PRINCIPIO DE FERMAT.

Hasta ahora hemos admitido la validez del Principio de Fermat, tal cual fue formulado por éste, es decir, suponemos que el camino óptico cumple la condición de mínimo (o que el tiempo empleado es el mínimo posible). Sin embargo, si suponemos por ejemplo tener un espejo cóncavo, el haz de luz reflejado en el espejo sigue una trayectoria real como la ACB y la longitud de

este camino óptico en el aire puede ser mayor que otro cualquiera como el AC'B (ver figura 2.8).

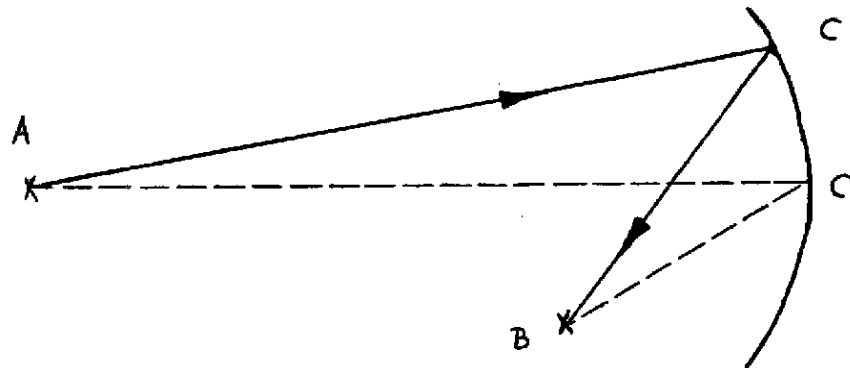


FIGURA 2.8 Reflexión de un espejo cóncavo.

Otro contraejemplo al enunciado tradicional del Principio de Fermat puede ser un espejo elíptico como el representado en la figura 2.9.

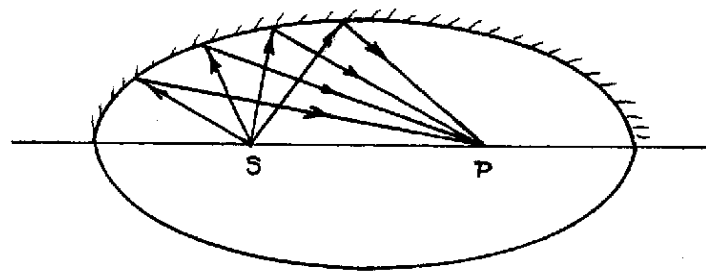


FIGURA 2.9 Reflexión en una cavidad elíptica espejada.

Se observa que si se coloca una fuente luminosa en uno de los focos de la elipse, cualquier rayo de luz que sale de la fuente "s" llega al otro foco p. En este caso todos los caminos ópticos son iguales (por definición de elipse) y no corresponde hablar ni de máximos ni de mínimos.

¿Diremos entonces que el Principio de Fermat es incorrecto tal como fue originariamente formulado? La respuesta es afirmativa. En lugar de postular que el camino óptico es mínimo, se postula:

"EL CAMINO OPTICO DEBE SER ESTACIONARIO",
es decir, debe cumplir con la condición:

$$\delta \int n(s) \cdot ds = 0$$

Esta es la versión "actualizada" del Principio de Fermat. Físicamente esto significa que la luz "elige" seguir un camino óptico que difiere muy poco de todos los caminos próximos a él.

Pero, ¿cómo "sabe" la luz de antemano que debe seguir un camino estacionario? Esta es una cuestión que no puede contestar el Principio de Fermat.

Debe observarse que este enfoque es completamente diferente de describir los hechos de la naturaleza como una secuencia ordenada de causas y efectos.

De acuerdo a Fermat, la luz tiene de "antemano" determinada trayectoria a seguir y ante lo que encuentra a su paso "se adapta" para seguir cumpliendo un principio general.

Seguendo a R. Feynmann (*), diremos que es como si la luz "tanteara" el conjunto de caminos que tienen caminos ópticos cercanos entre sí y descartara a los otros.

(*) Lecturas de Física, R. Feynman, R. Leighton, M. Sands
Pag. 26-12, Fondo Educativo Interamericano S.A. 1971.

3) PRINCIPIO DE HUYGENS(*)

- (*) Antes de ver este tema, repasar el concepto de "frente de onda" (Ver "Texto Complementario - Movimiento Ondulatorio").

Como se ve en la Introducción teórica, de acuerdo con Huygens, la luz consiste en ondas que se propagan en un medio elástico especial denominado "éter". Este medio ocupa todo el espacio, incluso los espacios entre átomos. Según Huygens, el mecanismo de propagación de la onda luminosa en un medio se puede describir de la siguiente manera: supongamos que conocemos la posición del frente de onda en el instante t .

¿Cuál será la forma y posición del mismo frente de onda después de un pequeño intervalo de tiempo δt ?

De acuerdo con el principio de Huygens:

Cada punto del frente de onda original es a su vez una fuente de perturbación del medio.
("perturbación secundaria")

La perturbación secundaria está formada por ondas esféricas con origen en los puntos del frente de ondas, que avanzan con la misma velocidad (v) que la onda primaria (ver fig. 3.1).

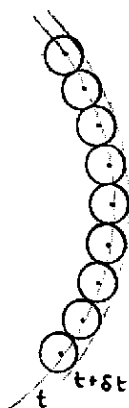


FIGURA 3.1 Perturbación secundaria.

En el instante $t + \delta t$, los frentes de onda de cada una de estas ondas tendrán un radio $v \cdot \delta t$. Huygens postuló que la envolvente de las onditas formarán el nuevo frente de onda de la perturbación original. La aplicación de este principio permite explicar fácilmente los fenómenos de la reflexión y refracción de la luz.

3.1) REFLEXION DE LA LUZ SEGUN HUYGENS.

Supongamos que a la superficie reflectora llega un tren de ondas plano como indica la fig. 3.2 (A distancias suficientemente grandes de la fuente, las ondas esféricas pueden ser consideradas como ondas planas). La secuencia indicada en la figura 3.2 muestra que:

- en el instante t_0 la onda todavía no ha llegado a la superficie.
- en el instante $t_1 > t_0$ el frente de onda alcanza al punto A que se transforma en fuente de ondas secundarias.
- en el instante $t_2 > t_1$ el frente de ondas alcanza al punto B y así sucesivamente.
- La envolvente de las ondas secundarias constituye el frente de onda de la onda reflejada.

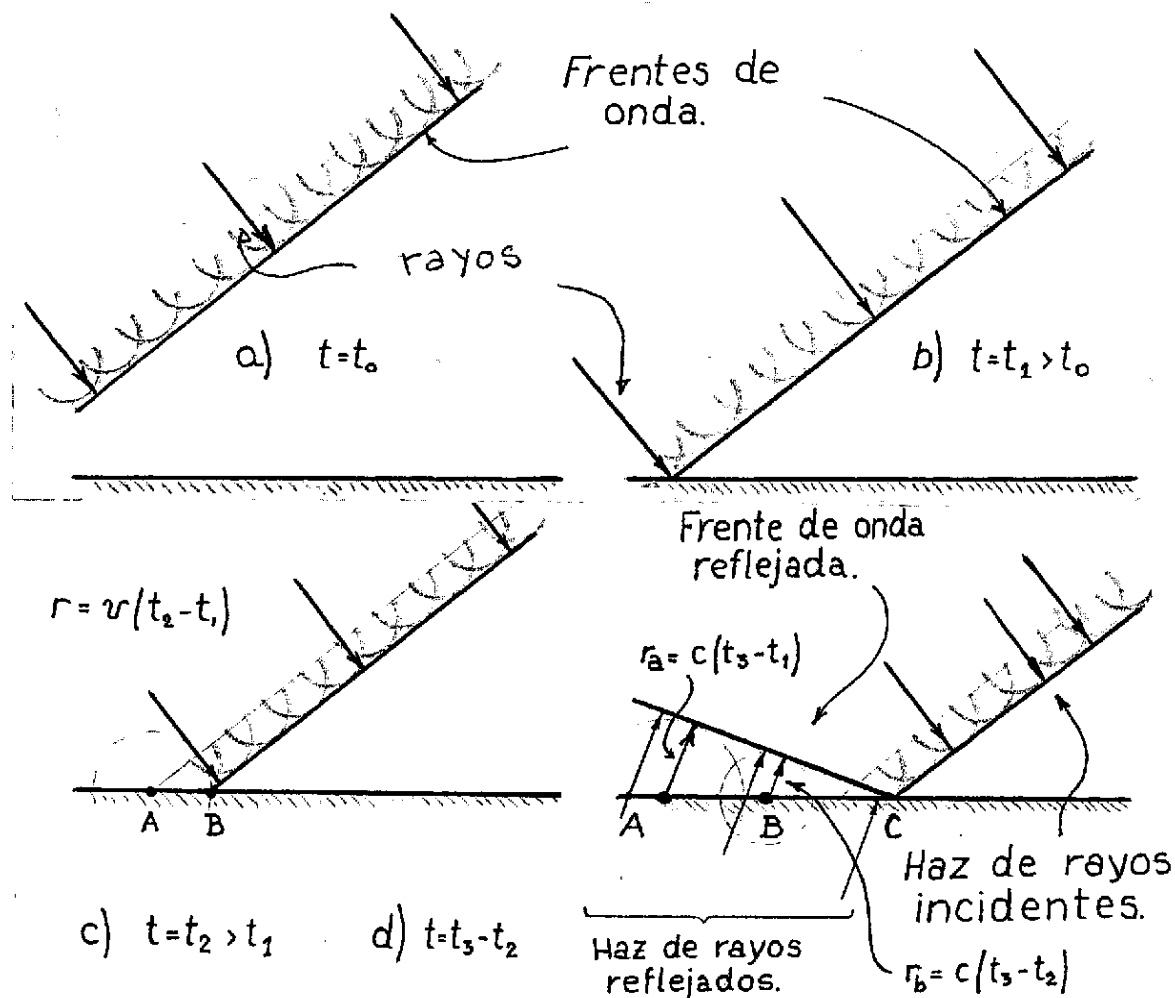


FIGURA 3.2 Secuencia de reflexión.

Cabe preguntarse: ¿Cómo será la relación entre el ángulo de incidencia y el de reflexión?

Según se muestra en la figura 3.3 en donde se sintetiza todo el proceso, cuando la onda llega a A le falta recorrer la distancia BC para llegar al punto C.

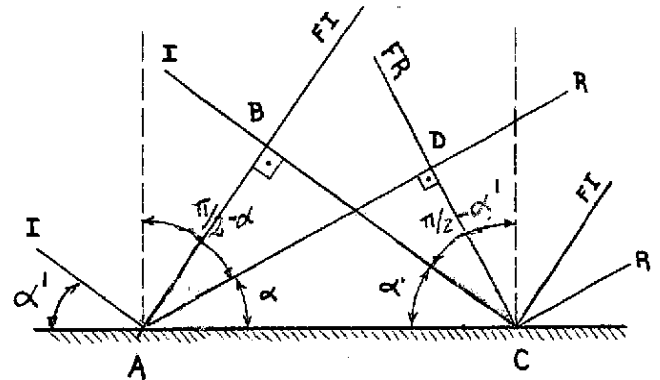
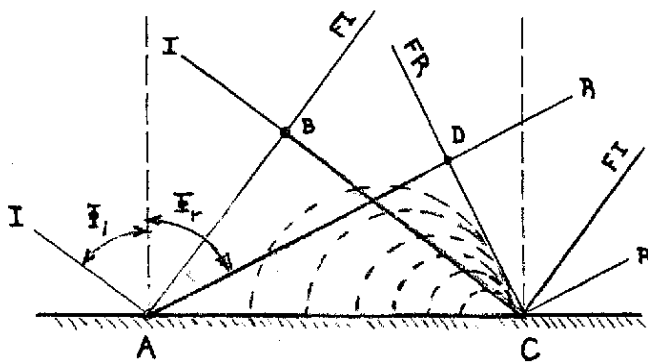
Cuando el frente de onda de la luz incidente llega a C, la onda secundaria reflejada ha recorrido la distancia AD.

Según se observa en la figura 3.3(b) se puede comparar los triángulos ABC y CDA. En ellos BC tiene que ser igual a AD pues en ambos casos el camino recorrido será $v \cdot (t_2 - t_1)$

$$BC = AD = v \cdot (t_2 - t_1)$$

Siendo $(t_2 - t_1)$ el tiempo que debe transcurrir para que la onda

incidente alcance al punto C a partir del instante que llegó a A. Pero en ese mismo lapso de tiempo, la onda secundaria alcanza al punto D habiendo partido de A. También transcurre este intervalo de tiempo para que el rayo que incidía en el punto B llegue al punto C. De todo lo visto se puede concluir que los dos triángulos, el ABC y el CDA, tienen sus lados respectivamente iguales.



(a) FI-Frente de onda incidente. (b) FR-Frente de onda reflejada
I - Rayo incidente. R-Rayo reflejado.

FIGURA 3.3

Luego: $\alpha = \alpha'$ $\pi/2 - \alpha = \pi/2 - \alpha'$
 Pero $\pi/2 - \alpha = \theta_i = \text{ángulo de incidencia}$
 y $\pi/2 - \alpha' = \theta_r = \text{ángulo de reflexión}$
 luego

$$\theta_i = \theta_r$$

Veamos ahora como podemos interpretar el fenómeno de la refracción.

3.2) REFRACCION DE LA LUZ SEGUN HUYGENS.

Procederemos de manera totalmente análoga al caso anterior. La diferencia consiste en que la luz en este caso cambia de medio. Pasa de un medio 1 a otro medio 2, en el cual la velocidad de propagación es diferente (ver la secuencia de la figura 3.4).

Cuando la onda incidente llega al punto A, le falta aún recorrer la distancia BC para llegar al punto C de la figura 3.5, mientras la onda incidente avanza de A a C la onda secundaria refractada lo hace de C a D recorriendo la distancia AD.

Luego: $BC = v_1.t$ $AD = v_2.t$

como el tiempo es el mismo $\frac{AD}{BC} = \frac{v_2}{v_1}$ (13)

como $\theta_i = \pi/2 - \alpha$ y $\cos(\alpha) = \frac{BC}{AC}$; $\sin(\theta_i) = \frac{BC}{AC}$

ó $BC = AC \cdot \sin(\theta_i)$.

análogamente

como $\theta_r = \pi/2 - \beta$ y $\cos(\beta) = \frac{AD}{AC}$; $\sin(\theta_r) = \frac{AD}{AC}$

ó $AD = AC \cdot \sin(\theta_r)$.

luego: $\frac{AD}{BC} = \frac{AC \cdot \sin(\theta_r)}{AC \cdot \sin(\theta_i)} = \frac{\sin(\theta_r)}{\sin(\theta_i)}$ (14)

pero considerando la Ley de Snell y por (13) y (14):

$$\frac{\sin(\theta_r)}{\sin(\theta_i)} = \frac{v_2}{v_1} = \frac{n_1}{n_2} = n_{12} \quad (15)$$

Expresada de una manera más usual $v_1 \sin(\theta_r) = v_2 \sin(\theta_i)$
 Por lo visto en (15) si $n_1 < n_2$, es decir, si la luz pasa de un medio menos denso a uno más denso, $\sin(\theta_i) > \sin(\theta_r)$ por lo tanto $\theta_i > \theta_r$ o sea que el rayo se acerca a la normal. Pero además $v_1 < v_2$. Por consiguiente la trayectoria de Huygens prevé una disminución de la velocidad de la luz al pasar ésta de un medio menos denso a uno más denso. Por el contrario, la trayectoria de Newton (modelo corpuscular de la naturaleza de la luz) estipula que en la interfase solo actúan sobre los corpúsculos fuerzas perpendiculares a la superficie, como se observa en la figura 3.6. Luego las componentes de la velocidad tangenciales a la superficie permanecen constantes.

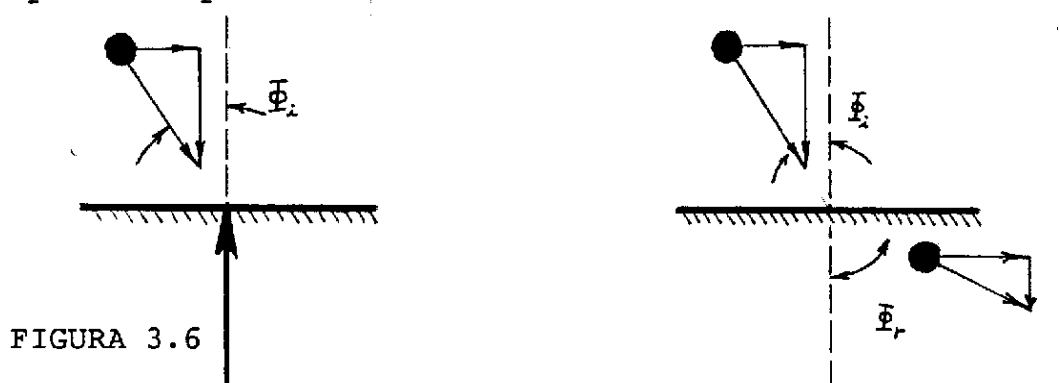


FIGURA 3.6

En consecuencia la velocidad de la luz en el medio más denso debe ser mayor que en el medio menos denso. La única forma de decidir cuál de las dos teorías se ajusta mejor al fenómeno estudiado es medir las velocidades de la luz en dos medios diferentes.

La experiencia realizada por Foucault en 1850 confirmó la predicción de Huygens, desechando la de Newton.

3.3 OBJECIONES AL PRINCIPIO DE HUYGENS

Pueden formularse una serie de objeciones al principio de Huygens tal como fue enunciado:

"Cada punto del frente de onda original es a su vez una fuente de perturbación del medio (perturbación secundaria)".

- En primer lugar no puede precisarse bajo qué condiciones experimentales se cumplirán las leyes de la óptica geométrica.
- Sugiere además la siguiente paradoja: la envolvente de las ondas secundarias se puede formar delante o detrás del frente de onda original. ¿Por qué no se observa una onda que retrocede y solo se observa la onda que avanza?
- El fenómeno de polarización no puede explicarse utilizando el Principio de Huygens.

La solución a estos problemas fue resuelta introduciendo nuevas hipótesis. En la sección que explica el fenómeno de interferencia volveremos sobre estas cuestiones.

4) SUPERPOSICION DE ONDAS:

SUPERPOSICION DE DOS TRENES DE ONDAS QUE SE PROPAGAN EN EL MISMO SENTIDO.

Supongamos tener un fenómeno ondulatorio -puede tratarse de vibraciones mecánicas, de ondas luminosas, de ondas sonoras, etc.- que puede ser descrito por dos trenes de ondas armónicas que se propagan en el sentido de las x positivas, con la misma frecuencia (ω), longitud de onda (λ_i) y velocidad (v), aunque con diferentes fases iniciales (α_i) y diferentes amplitudes máximas (E_{oi}). El fenómeno debe ser analizado como la superposición de dos ondas sinusoidales:

$$E_1 = E_{o1} \cdot \text{sen}(\omega \cdot t + \alpha_1) \quad \text{y} \quad E_2 = E_{o2} \cdot \text{sen}(\omega \cdot t + \alpha_2) \quad (16)$$

donde $\alpha_i(x, \phi) = - (k \cdot x + \phi_i)$ con $i=1,2,3...$ son las fases iniciales de cada punto alcanzado por el tren de ondas.

Sumando las ecuaciones (16) y aplicando la fórmula del $\text{sen}(\alpha + \beta)$ se obtiene:

$$\begin{aligned} E &= E_{o1} [\text{sen}(\omega t) \cdot \cos(\alpha_1) + \cos(\omega t) \cdot \text{sen}(\alpha_1)] + \\ &\quad + E_{o2} [\text{sen}(\omega t) \cdot \cos(\alpha_2) + \cos(\omega t) \cdot \text{sen}(\alpha_2)] \\ \text{o} \\ E &= [E_{o1} \cdot \cos(\alpha_1) + E_{o2} \cdot \cos(\alpha_2)] \cdot \text{sen}(\omega t) + \\ &\quad + [E_{o1} \cdot \text{sen}(\alpha_1) + E_{o2} \cdot \text{sen}(\alpha_2)] \cdot \cos(\omega t) \end{aligned} \quad (17)$$

Si establecemos las siguientes igualdades:

$$\left. \begin{aligned} E_{o1} \cdot \cos(\alpha_1) + E_{o2} \cdot \cos(\alpha_2) &= E_o \cdot \cos(\alpha) \\ E_{o1} \cdot \text{sen}(\alpha_1) + E_{o2} \cdot \text{sen}(\alpha_2) &= E_o \cdot \text{sen}(\alpha) \end{aligned} \right\} \quad (18)$$

y las reemplazamos en (17) resulta:

$$\begin{aligned} E &= E_o \cdot \cos(\alpha) \cdot \text{sen}(\omega t) + E_o \cdot \text{sen}(\alpha) \cdot \cos(\omega t) \\ \text{o} \\ E &= E_o \cdot \text{sen}(\omega t + \alpha) \end{aligned} \quad (19)$$

La expresión (19) de la onda resultante nos señala que se trata de una onda armónica de la misma frecuencia de las componentes, aunque de distinta fase (α) y distinta amplitud (E_o).

4.1.1.) INTENSIDAD DE LA ONDA RESULTANTE.

Como ya se vió en el desarrollo del tema de ondas, la intensidad de una onda es proporcional al cuadrado de su máxima amplitud.

Así $I_1 = C \cdot (E_1)^2$, $I_2 = C \cdot (E_2)^2$ y la intensidad de la onda resultante de la superposición de los dos trenes de onda:

$$I = C \cdot E^2 \quad (20)$$

donde C es una constante de proporcionalidad. Pero por las igualdades de (18) resulta:

$$E^2 = E_{o1}^2 + E_{o2}^2 + 2 \cdot E_{o1} \cdot E_{o2} [\cos(\alpha_1) \cdot \cos(\alpha_2) + \text{sen}(\alpha_1) \cdot \text{sen}(\alpha_2)]$$

$$\text{o sea} \quad E^2 = E_{o1}^2 + E_{o2}^2 + 2 \cdot E_{o1} \cdot E_{o2} \cdot \cos(\alpha_2 - \alpha_1)$$

$$6 \quad I = I_1 + I_2 + 2 \cdot C \cdot E_{o1} \cdot E_{o2} \cdot \cos(\alpha_2 - \alpha_1) \quad (21)$$

Según lo que se observa en la expresión (21), la intensidad de la onda resultante de la superposición de dos trenes de ondas armónicas que se propagan en el mismo sentido, difiere de la suma de las intensidades de las ondas componentes en el término $2 \cdot C \cdot E_{o1} \cdot E_{o2} \cdot \cos(\alpha_2 - \alpha_1)$. A este término se lo denomina TERMINO DE INTERFERENCIA, y a la diferencia $(\alpha_2 - \alpha_1)$ se la llama INTERFERENCIA DE FASE. Por lo visto en (16) será:

$$(\alpha_2 - \alpha_1) = \delta = (k \cdot x_1 + \phi_1) - (k \cdot x_2 + \phi_2)$$

donde x_1 y x_2 son los caminos recorridos por cada onda.

$$\text{Luego} \quad \delta = k \cdot (x_1 - x_2) + (\phi_1 - \phi_2) = (2 \cdot \pi / \lambda \cdot x) + \phi \quad (22)$$

En la expresión (22) se ve que la diferencia de fase se compone de dos términos: uno debido a la diferencia de caminos, y otro debido a la diferencia en las fases iniciales.

Daremos un ejemplo más complicado:

Sean dos fuentes luminosas puntuales S_1 y S_2 que emiten ondas monocromáticas de la misma frecuencia en un medio homogéneo. ¿Qué sucede en un punto P lo suficientemente alejado de las fuentes como para suponer que los frentes de onda que arriban al mismo son planos?

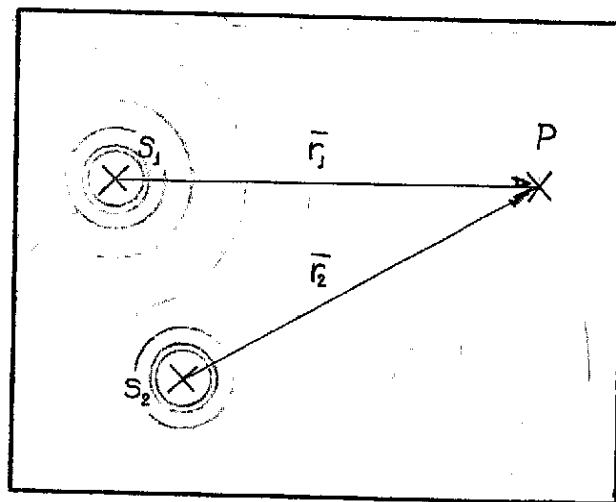


FIGURA 4.1

Por ser iguales las frecuencias de las dos fuentes y tener un medio homogéneo (velocidad de propagación constante), el número de onda $k = (2 \cdot \pi / \lambda)$ es el mismo para ambas. En cuanto a la posición ya no podemos considerar una magnitud escalar x , sino una magnitud vectorial $\vec{r}$ que tenga en cuenta la variación de la posición en el plano.

Cada una de las ondas que llega al punto P estarán entonces representadas por un vector $\vec{E}$ tal como:

$$\vec{E}_1 = E_{o1} \cdot \text{sen}(k \cdot r_1 - \omega t + \phi_1) \quad \text{y} \quad \vec{E}_2 = E_{o2} \cdot \text{sen}(k \cdot r_2 - \omega t + \phi_2) \quad (23)$$

en donde $\vec{E} = \vec{E}_1 + \vec{E}_2$ será el vector resultante de la composición de las dos ondas que llegan al punto.

Como la intensidad es proporcional al cuadrado de la amplitud, la expresión de la intensidad será:

$$I = C.E^2 = C . (E_1+E_2)^2 = C.E_1^2 + C.E_2^2 + 2.\vec{E_1}.\vec{E_2}.C$$

pero $C . E_1^2 = C . E_{01}^2 = I_1$ y $C . E_2^2 = C . E_{02}^2 = I_2$

$$\text{de donde } I = I_1 + I_2 + 2.C.E_1.E_2 \quad (24)$$

Como $\vec{E_1}.\vec{E_2} = E_{01} . E_{02} . \cos(\delta)$ y

$$\cos(\delta) = \cos(\alpha_1 - \alpha_2) = \cos[(k.r_1 - .t + \phi_1) - (k.r_2 - .t + \phi_2)]$$

$$\cos(\delta) = \cos(k . r + \Phi) \quad (25)$$

con $r = r_1 - r_2$ y $\Phi = \phi_1 - \phi_2$

reemplazando lo hallado en (24) será

$$I = I_1 + I_2 + 2 . E_{01} . E_{02} . \cos(k . r + \Phi)$$

(26)

Como $r = f(t)$ y $\Phi = f(t)$ hemos deducido la expresión de la INTENSIDAD INSTANTANEA.

4.1.3.) ANALISIS DE LOS RESULTADOS EXPERIMENTALES

¿Pero es en realidad la intensidad instantánea la que percibimos?. Debido al tiempo que tarda el cerebro procesar la información recibida, a pesar del uso de equipos sofisticados, lo que percibimos es la intensidad promedio de un intervalo de tiempo. Debemos usar promedios temporales. El promedio temporal de cualquier función $f(t)$ periódica se define matemáticamente de la siguiente manera:

$$\langle f(t) \rangle = \frac{1}{T} \int_0^T f(t) . dt$$

donde T es el período de la función considerada.

Como el valor medio del $\sin^2(\alpha) = \cos^2(\omega t) = \frac{1}{2}$,
 entonces será $\langle E_1^2 \rangle = E_{01}^2/2$; $\langle E_2^2 \rangle = E_{02}^2/2$
 y como además $\langle \sin(\omega t) . \cos(\omega t) \rangle = 0$
 resulta:

$$\langle \vec{E_1} . \vec{E_2} \rangle = \frac{1}{2} . E_{01} . E_{02} . \cos(k . \vec{r_1} + \phi_1 - k . \vec{r_2} - \phi_2) \quad (27)$$

pero $k . \vec{r_1} - k . \vec{r_2} + \phi_1 - \phi_2 = \delta$ es lo que en (22) llamamos diferencia de fase. De acuerdo con lo percibido experimentalmente denominaremos TERMINO DE INTERFERENCIA (según lo visto en 21) a:

$$I_{12} = 2 \cdot C \cdot \langle E_1 \cdot E_2 \rangle = c \cdot E_{o1} \cdot E_{o2} \cdot \cos(\delta) \quad (28)$$

de acuerdo a lo hallado en (27)

$$\left. \begin{aligned} I_{12} &= C \cdot E_{o1} \cdot E_{o2} \cdot \cos(\delta) \\ I_1 &= C \cdot \langle E_1^2 \rangle = (C \cdot E_{o1}^2)/2 \\ I_2 &= C \cdot \langle E_2^2 \rangle = (C \cdot E_{o2}^2)/2 \end{aligned} \right\} \quad (29)$$

Las ecuaciones (29) permiten analizar los resultados experimentales. Para analizar con más simplicidad supongamos que las amplitudes de las dos ondas componentes y la de la resultante de su composición son iguales. En ese caso:

$$E_{o1} = E_{o2} = E_o \quad \text{con lo cual} \quad I_1 = I_2 = (C \cdot E_o^2)/2 \quad (30)$$

De acuerdo con lo hallado en (30) y reemplazando estos resultados en (29) y en (26)

$$I = I_1 + I_2 + I_{12}$$

$$I = (C \cdot E_o^2)/2 + (C \cdot E_o^2)/2 - 2 \cdot C \cdot E_o^2 \cdot \cos(\delta)$$

entonces

$$I = C \cdot E_o^2 \cdot [1 - \cos(\delta)] \quad (31)$$

¿Está en coincidencia este resultado con nuestra experiencia?. De acuerdo con la teoría ondulatoria de la luz que estamos analizando, la intensidad de la onda resultante de la composición de dos trenes de ondas puede ser menor ó a lo sumo igual que la suma de las intensidades de cada fuente. En particular si $\cos(\delta)=1$ en ese caso la intensidad resultante es nula!!!

Hemos llegado a una situación paradójica, ya que no concuerda con nuestra experiencia. Lo que observamos es que al superponer entre sí dos focos luminosos (por ejemplo dos lámparas o dos linternas) la intensidad resultante siempre aumenta.

5) EL CONCEPTO DE COHERENCIA.

Debemos examinar con atención los principios fundamentales de la teoría ondulatoria de modo que permita analizar los resultados experimentales. Para ello nos preguntaremos ¿Cómo se produce el fenómeno de la luz natural?

Independientemente de las distintas formas de emisión de la luz (provenga del sol, de un filamento incandescente, de una llama, de un tubo de neón, etc.) podemos decir que una fuente luminosa típica contiene un gran número de átomos excitados. Cada uno de dichos átomos irradia un tren de ondas durante un período muy corto de tiempo, del orden de 10^{-8} segundos y vibrando en diferentes planos como indica la figura 5.1.

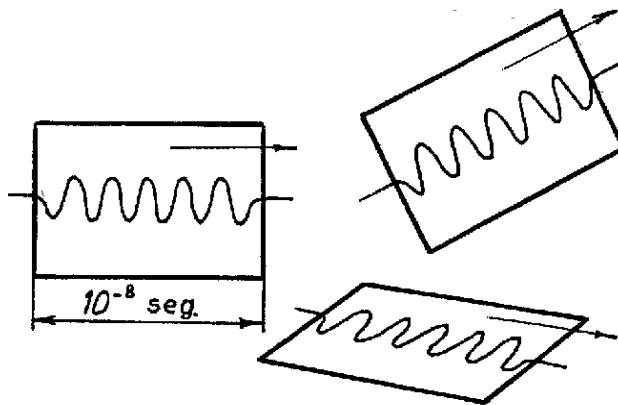


FIGURA 5.1 Trenes de ondas emitidos por átomos excitados.

Además, cada átomo emite grupos de ondas que ni siquiera son monocromáticas, sino que tienen frecuencia variable.

¿QUE SUPUESTO HEMOS HECHO QUE NO COINCIDE CON LA REALIDAD?

Uno es que las ondas fueran monocromáticas. Pero esto prácticamente se puede obtener con mucha facilidad, ya que la luz emitida por una lámpara de sodio es casi monocromática y se consigue fácilmente poniendo, ante la llama de un mechero, un isopo embebido en una solución de cloruro de sodio (sal de cocina). Pero a pesar de tener una fuente monocromática, si colocamos dos llamas la intensidad aumentará siempre. ¿El término de interferencia no existe en este caso?

El otro supuesto que se hizo es que la diferencia de fases se mantenía constante en el tiempo. ¿Es aplicable esta condición a cualquier fuente luminosa? Como se vió anteriormente las fuentes de luz "naturales" se caracterizan por emitir ondas de manera aleatoria. Es evidente entonces que la condición de constancia de la diferencia de fases no se va a cumplir espontáneamente para períodos de tiempo mayores a 10^{-8} segundos.

En un dispositivo experimental que nos asegure además del monocromatismo la constancia de la diferencia de fase en el tiempo, se cumple perfectamente la ley desarrollada en (31).

CUANDO DOS O MAS ONDAS MONOCROMATICAS SE PROPAGAN MANTENIENDO UNA DIFERENCIA DE FASE CONSTANTE DURANTE EL TIEMPO DE OBSERVACION, SE DICE QUE LAS OSCILACIONES SON COHERENTES.

SI LAS OSCILACIONES CAMBIAN DE MANERA ALEATORIA DURANTE EL TIEMPO DE OBSERVACION SE DICE QUE LAS OSCILACIONES SON INCOHERENTES.

Se ve que los tiempos de observación son mucho mayores que los tiempos de emisión atómica y por supuesto mucho mayores que el período de oscilación de la onda luminosa.

5.1) INTENSIDAD RESULTANTE DE ONDAS COHERENTES.

Por lo visto en (31) en el caso de la superposición de ondas coherentes, la intensidad de la onda resultante será:

$$I = I_1 + I_2 + I_{12} = C \cdot E_0^2 \cdot [1 - \cos(\delta)]$$

5.2) INTENSIDAD RESULTANTE DE ONDAS INCOHERENTES.

Cuando las fuentes emiten de manera incoherente, tanto los números de onda (k_1, k_2) como las fases instantáneas (ϕ_1, ϕ_2) son funciones del tiempo.

La función

$$\cos(\delta) = \cos(k_1 \cdot r_1 - k_2 \cdot r_2 + \phi_1 - \phi_2) \quad (32)$$

ya no es una constante, sino que resulta una función aleatoria del tiempo. Para determinar la intensidad resultante deberíamos calcular cómo influye dicha función en el valor medio temporal de la intensidad.

La diferencia de fase cambia de manera completamente aleatoria, tomando los valores posibles desde 0 a (2π) .

Intuitivamente se puede ver que si la probabilidad de tomar cualquier valor es la misma, la función coseno toma todos los valores posibles entre 1 y -1. En este caso el valor medio temporal del término de interferencia será nulo.

$$\langle I_{12} \rangle = 0$$

LA INTENSIDAD RESULTANTE DE ONDAS INCOHERENTES SERA:

$$\langle I \rangle = \langle I_1 \rangle + \langle I_2 \rangle$$

Este es resultado que está de acuerdo a nuestra experiencia cotidiana. (Una demostración rigurosa de $\langle I_{12} \rangle = 0$ está fuera de los alcances de este curso, pues para llegar a ese resultado se debe hacer uso de la teoría de funciones aleatorias).

6) LASER - UNA FUENTE DE LUZ COHERENTE

Es nuestro propósito dar una explicación simplificada de este tema ya que los conocimientos necesarios para una comprensión profunda del mismo están fuera del alcance de este curso.

En Química se han abordado los Principios de la Teoría Atómica y por lo tanto podemos partir del concepto de niveles de energía (atómicos o electrónicos) para desarrollar el tema.

A partir de este concepto Albert Einstein (1917) consideró dos niveles de energía cualesquiera, por ejemplo n y m . ¿Qué sucede cuando una radiación de frecuencia apropiada incide sobre un átomo del nivel n ? En principio el átomo puede absorberla y pasar del estado n al m .

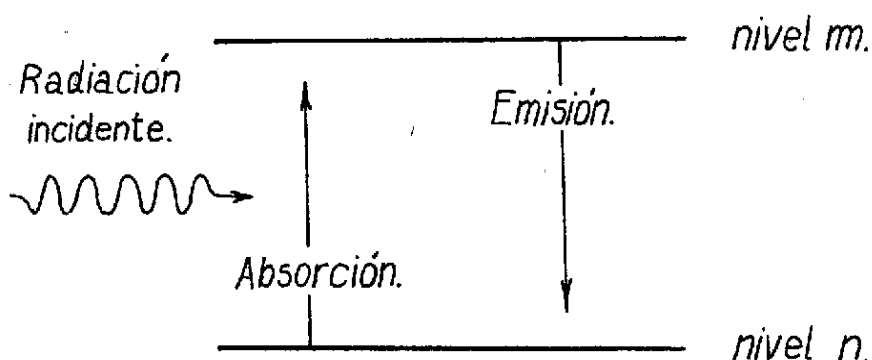


FIGURA 6.1 Diagrama de niveles de energía

El número de átomos que efectúan esta transición será proporcional al número inicial de átomos en el nivel n y a la intensidad de la radiación incidente (I_w)

$$R_{n-m} = N_n \cdot B_{nm} \cdot I_w$$

siendo: R_{n-m} El número de átomos que efectúan la transición de n a m por unidad de tiempo.

N_n Es el número de átomos inicialmente en el estado de energía n .

B_{nm} Es un coeficiente de proporcionalidad.

El proceso inverso, el paso de m a n , puede producirse de dos maneras distintas. En primer lugar, aún si no hubiera radiación presente, existiría la probabilidad que un átomo excitado cayera al estado inferior emitiendo radiación. Dado que no hay ninguna causa aparente que lo provoque a este fenómeno se lo denomina "emisión espontánea". El número de átomos que efectúan esta transición será proporcional al número de átomos que inicialmente están en el estado m ($A_{mn} \cdot N_m$).

Hasta aquí la descripción del fenómeno se encontraba dentro del marco de la teoría conocida en ese momento. Sin embargo A. Einstein introdujo una nueva hipótesis: la emisión de luz también puede ser influenciada por la radiación incidente. A este tipo de emisión generada por la radiación incidente sobre el átomo la llamó "emisión inducida" y supuso que el número de átomos que la efectuaba por unidad de tiempo era proporcional a la intensidad incidente (I_w) y al número de átomos excitados.

Por consiguiente el número de átomos que efectúa la transición de m a n por unidad de tiempo será:

$$\begin{aligned} R_{m \rightarrow n} &= A_{mn} \cdot N_m + B_{mn} \cdot N_m \cdot I_w \\ &= N_m \cdot (A_{mn} + B_{mn} \cdot I_w) \end{aligned}$$

En el equilibrio

$$\begin{aligned} R_{m \rightarrow n} &= R_{n \rightarrow m} \\ N_m \cdot (A_{mn} + B_{mn} \cdot I_w) &= N_n \cdot B_{nm} \cdot I_w \end{aligned}$$

La relación N_m/N_n y la función I_w se pueden encontrar empleando conceptos de Física Estadística y Mecánica Cuántica.

Del razonamiento anterior se deduce una consecuencia muy importante. Si se consigue colocar (no importa por el momento mediante que mecanismo) átomos en un nivel superior ("bombeo óptico") la radiación luminosa externa los inducirá a cambiar a un nivel inferior emitiendo luz. Más importante aún es el hecho (que se puede demostrar) que bajo determinadas condiciones la radiación emitida (inducida) es coherente con la radiación que la provoca.

Teóricamente en un sistema con el nivel m mucho más poblado que el n , bastará con que un solo átomo efectúe la transición de m a n emitiendo radiación para que se produzca una especie de emisión "en cadena" dado que la radiación emitida inducirá a otro átomo a emitir y así sucesivamente. Este es el fundamento del LASER, cuyo nombre deriva de la expresión Light Amplification by Stimulated Emission of Radiation.

Una analogía mecánica simplificada de este fenómeno puede ser la siguiente. Supongamos tener un estanque con un conjunto de boyas que en su parte superior poseen piedras. Las piedras pueden caer si las boyas vibran. Si por algún motivo una de las piedras cae al estanque, esto provocará ondas que al llegar a las boyas más próximas harán caer las piedras al agua con lo cual se seguirán produciendo nuevas ondas. En esta analogía los átomos serán las piedras y el agua agitada la luz inducida. Esta similitud es una simplificación muy grande del fenómeno LASER dado que, entre otras cosas, nunca las ondas del agua serán coherentes entre sí.

Analicemos ahora los elementos componentes de un laser:

a- Medio Activo: se necesita tener un medio que posea los átomos que van a ser inducidos a emitir. A ese medio se lo denomina activo. Existen diferentes medios activos: gases o mezcla de gases, cristales o vidrios adoptados con iones especiales (laser de estado sólido), líquidos, semiconductores.

b- Mecanismo de estimulación: es necesario llevar los átomos a un nivel superior a tal fin de ponerlos en disposición de emitir. En el primer laser que se construyó (T. Maiman 1960) se utilizaba un tubo fluorescente arrollado alrededor del medio activo. En el interior del tubo se provocaba una descarga eléctrica uniendo los electrodos del interior del tubo a una fuente. Entonces en el interior del tubo se producía un pulso de luz muy intenso. Esta energía luminosa era absorbida por los átomos del medio activo (cristal de rubí) que de esta manera eran excitados a un nivel superior.

c- Una vez producida la primera avalancha luminosa se hace necesario impedir que la radiación emitida se pierda en el medio ambiente y que de esa manera se interrumpa el proceso. Para

ello se confina al medio activo entre dos espejos ("resonador óptico"). De esta manera se van produciendo reflexiones que van incrementando el efecto inicial.

Se obtiene así un haz de luz extremadamente coherente y muy confinado en el espacio.

7) INTERFERENCIA.

DIREMOS QUE ESTAMOS EN PRESENCIA DE UN FENOMENO DE INTERFERENCIA CUANDO LAS FUENTES QUE SUPERPONEN SUS EFECTOS SON COHERENTES , ES DECIR, CUANDO $\Delta\phi \neq 0$.

Es importante ante todo aclarar una cuestión de terminología que es, generalmente, una fuente de confusión.

En muchos libros de texto aparecen frases del siguiente tipo: "...para que se produzca interferencia es necesario que se cumpla con la condición de coherencia (diferencia de fase constante)...". Pero en los mismos libros se puede ver que se denomina interferencia entre dos (o más) trenes de ondas, a la superposición entre ellas con su correspondiente onda resultante. Si por interferencia se entiende la interacción de ondas, debe tenerse en cuenta que ésta se produce independientemente de la coherencia o no de las fuentes. En cambio, si por interferencia entendemos la condición $\Delta\phi \neq 0$, entonces es evidente que necesitamos la coherencia como condición necesaria.

Cuando los libros de texto describen las clásicas experiencias de interferencia, se refieren a la aparición de diagramas (o patrones) estables en el espacio y en el tiempo, y pueden ser registrados ya sea visualmente o con equipos adecuados, y en este caso nuevamente es indispensable la condición de coherencia antes dicha.

7.1.) EXPERIENCIA DE YOUNG.

Young estudió la interferencia de dos fuentes puntuales coherentes con una famosa experiencia cuyo esquema se muestra en la figura 7.1.

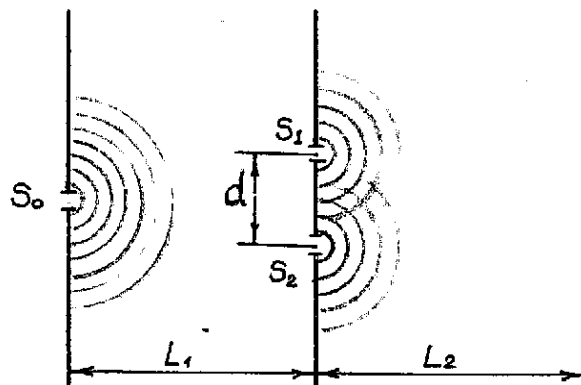


FIGURA 7.1

En principio se deben cumplir las siguientes hipótesis experimentales:

- * La fuente S_0 debe ser "puntual y monocromática".
- * Los orificios S_1 y S_2 puntuales y próximos ($d \ll L_2$).
- * $L_1 \gg d$.
- * Medio isótropo y de índice de refracción igual a 1.

- * Luz linealmente polarizada en un mismo plano (plano de oscilación único).

Así, los trenes de onda que parten de S_1 y de S_2 serán coherentes pues, en la primera aproximación, se cumplen las siguientes condiciones:

- * Las dos ondas secundarias generadas en S_1 y S_2 provienen de una única fuente S_0 , que por ser $L_1 \gg d$, llegan a S_1 y S_2 con frentes de ondas planos (al mismo tiempo). Las fases iniciales ϕ_1 y ϕ_2 son entonces las mismas para ambas ondas "secundarias" y esta constancia se mantiene durante todo el tiempo.
- * Las otras características de las ondas "secundarias" (amplitud, longitud de onda, frecuencia, plano de vibración) son invariables en el tiempo.

En base a este análisis podemos afirmar que las condiciones de coherencia se logran con este dispositivo. Podemos entonces, efectuar el siguiente análisis matemático del fenómeno:

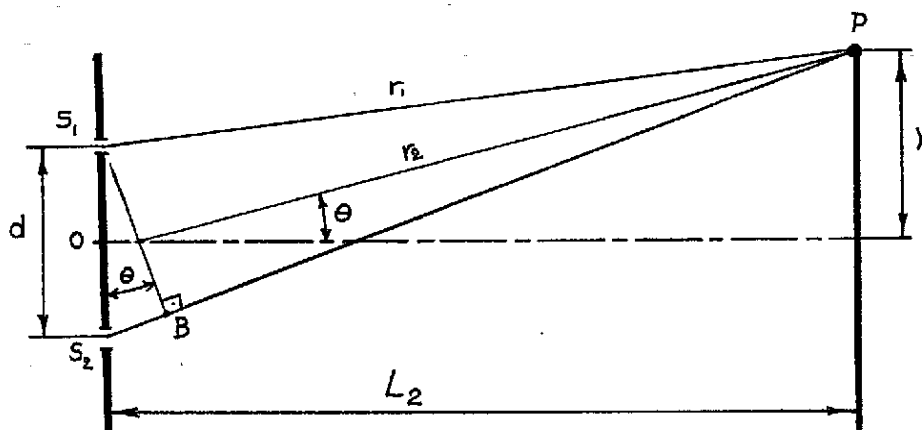


FIGURA 7.2

Dado que $L_2 \gg S_1S_2$, los segmentos $\overline{S_2P}$, $\overline{OP}$ y $\overline{S_1P}$ son aproximadamente paralelos entre sí (ver figura 7.2). Como podemos apreciar en la figura, si bien ϕ_1 y ϕ_2 son iguales, $r_1 \neq r_2$. Al punto P , arbitrariamente elegido en la pantalla, llegarán las ondas con una diferencia de fase δ debido a una diferencia de camino Δr tal que:

$$\delta = k \cdot \Delta r = (2 \cdot \pi \cdot \Delta r) / \lambda \quad (33)$$

Si trazamos por S_2 el segmento $\overline{S_2B}$ perpendicular común a todas las "paralelas" ($\overline{S_1P}$, $\overline{S_2P}$ y $\overline{OP}$), queda determinado el triángulo rectángulo S_2PS_1 . Analizando el gráfico desde el punto de vista físico, las paralelas $\overline{S_1P}$ y $\overline{S_2P}$ simbolizan un rayo luminoso que, saliendo de cada fuente, llega al punto P arbitrario de la pantalla. La perpendicular común $\overline{S_2B}$ es entonces un frente de onda común a ambas fuentes. Los caminos ópticos a partir de ese frente de onda son iguales. Luego el segmento $\overline{S_1B} = \Delta r$ es la diferencia de camino óptico buscada $r_1 - r_2 = \Delta r$.

Pero geométicamente en el triángulo S1S2B es $S1B = S1S2 \cdot \sin(\theta)$, entonces:

$$\Delta r = d \cdot \sin(\theta) \quad (34)$$

reemplazando lo hallado en (34) en (33) resulta:

$$\delta = \frac{2 \cdot \pi}{\lambda} \cdot d \cdot \sin(\theta) \quad (35)$$

Para que se produzca interferencia constructiva las ondas deben llegar al punto P en fase. Esto se cumplirá cuando los caminos ópticos sean iguales o cuando produzca una diferencia de fase de un número entero de veces 2π .

$$\delta = \Delta \alpha = 2 \cdot m \cdot \pi \quad \text{ó} \quad \Delta r = m \cdot \lambda \quad \Rightarrow \quad 2 \cdot m \cdot \pi = \frac{2 \cdot \pi}{\lambda} \cdot d \cdot \sin(\theta)$$

$$\text{Luego: } \sin(\theta) = \frac{m \cdot \lambda}{d} \quad (36)$$

Pero la condición escrita en (36) está referida al ángulo θ y si se observa la pantalla en lugar de visualizar los ángulos, resulta mas sencillo ver las posiciones del punto P de la pantalla.

En la figura 7.2 se observa que, por ser $S1B \perp OP$, el triángulo S1S2B es semejante al OPA de donde:

$$\sin(\theta) = \frac{y}{(L^2 + y^2)^{\frac{1}{2}}}$$

Podemos entonces afirmar que tendremos máxima iluminación en la pantalla para los puntos de la misma que cumplan con la condición:

$$\frac{y}{(L^2 + y^2)^{\frac{1}{2}}} = \frac{m \cdot \lambda}{d}$$

para $\theta < 10^\circ$ se cumple $\sin(\theta) \approx \text{tg}(\theta) \approx \theta \approx \frac{y}{L}$

Condición de máximo
para θ pequeño:

$$y = \frac{m \cdot \lambda}{d} \cdot L \quad (38)$$

Con este análisis se determinó donde se encuentran las franjas luminosas, y donde las franjas oscuras en la pantalla.

Análogamente, la condición de mínimo o interferencia destructiva será:

$$\delta = \Delta \alpha = (2 \cdot m + 1) \cdot \pi \quad \text{ó} \quad \Delta r = (2 \cdot m + 1) \cdot \frac{\lambda}{2} \quad \Rightarrow \quad \sin(\theta) = \frac{(2 \cdot m + 1) \cdot \lambda}{2 \cdot d}$$

La condición general de mínimo:
$$\frac{y}{(L^2 + y^2)^{\frac{1}{2}}} = \frac{(2.m+1).\lambda}{2.d}$$

Condición de mínimo
para θ pequeño:

$$y = \frac{(2.m+1).\lambda}{2.d} \cdot L2 \quad (39)$$

y la distancia entre franjas oscuras resulta igual a la de franjas luminosas vista en (38). LAS FRANJAS LUMINOSAS Y LAS OSCURAS ESTAN IGUALMENTE ESPACIADAS.

Ahora determinaremos la intensidad luminosa que llega a cada uno de los puntos de la pantalla. Como las fuentes son coherentes, para cada uno de ellos los rayos provenientes de cada fuente mantendrán una diferencia de fase constante $\delta = \alpha = \text{cte}$ y como la frecuencia es única e igual para ambas fuentes, puede aplicar el METODO FASORIAL DE FRESNEL.

7.1.1) ANALISIS FASORIAL DE LA SUPERPOSICION DE DOS FUENTES COHERENTES.

Siendo E_1 y E_2 la amplitud de cada onda en el punto P y $\Delta \alpha$ (ó δ) el desfase entre ambas, la amplitud de la onda resultante E la podemos hallar sumando los fasores representados en la figura 7.3.

Será $E_x = E_0 + E_0 \cdot \cos(\delta)$

y $E_y = E_0 \cdot \sin(\delta)$

ya que $E_1 = E_2 = E_0$, las amplitudes de ambas ondas son iguales.

Luego $E = (E_x^2 + E_y^2)^{\frac{1}{2}}$

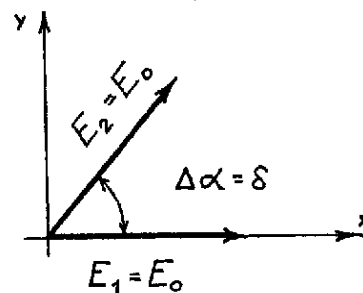


FIGURA 7.3

o sea $E = \left[E_0^2 + E_0^2 \cdot \cos^2(\delta) + 2 \cdot E_0^2 \cdot \cos(\delta) + E_0^2 \cdot \sin^2(\delta) \right]^{\frac{1}{2}}$

como $\sin^2(\delta) + \cos^2(\delta) = 1$ entonces la expresión de E resulta:

o sea $E = \left[2 \cdot E_0^2 + 2 \cdot E_0^2 \cdot \cos(\delta) \right]^{\frac{1}{2}}$

$$E = 2 \cdot E_0 \cdot \left[1 + \cos(\delta) \right]^{\frac{1}{2}} \quad (40)$$

Como la intensidad es directamente proporcional al cuadrado de la amplitud, será:

$$I = C \cdot 2 \cdot E_0^2 [1 + \cos(\delta)] \quad \text{ó} \quad I = 2 \cdot I_0 \cdot [1 + \cos(\delta)]^2 \quad (41)$$

donde I_0 es la intensidad de cada fuente. Aplicando la fórmula trigonométrica del ángulo mitad, la expresión (41) se transforma en:

$$I = 4 \cdot I_0 \cdot \cos^2 \left[\frac{\delta}{2} \right]^{\frac{1}{2}} = 4 \cdot I_0 \cdot \cos^2 \left[k \cdot \frac{\Delta r}{2} \right]^{\frac{1}{2}}$$

y reemplazando $k=(2\pi/\lambda)$ y $\Delta r=d.\text{sen}(\theta)$ ó $\Delta r=(d.y)/L_2$ para valores pequeños de θ resulta:

INTENSIDAD EN LA PANTALLA

PARA θ PEQUEÑO

$$I = 4 \cdot I_0 \cdot \cos^2 \left[\frac{\pi \cdot d \cdot y}{L_2 \cdot \lambda} \right] \quad (42)$$

La representación gráfica de $I=f(y)$ es la que se muestra en la figura 7.4 y coincide con lo hallado en (37), (38), (39) y con las franjas iluminadas y oscuras que se observan experimentalmente.

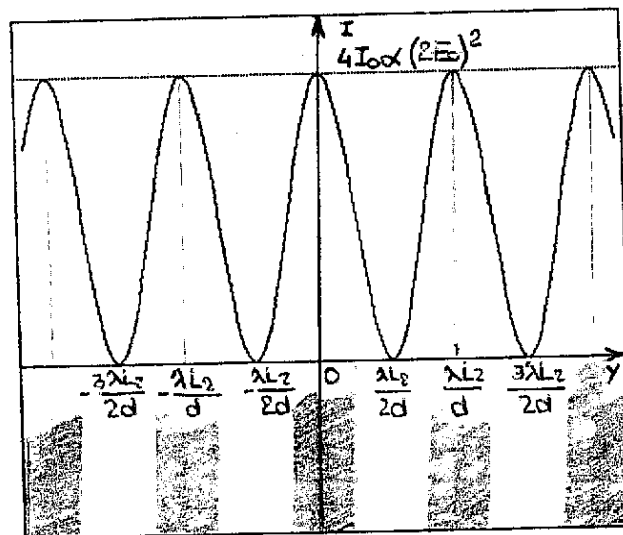


FIGURA 7.4

En la figura 7.4, las franjas rayadas son las que se ven iluminadas en la pantalla, y las franjas lisas son las de oscuridad.

A pesar de haber llegado a un resultado satisfactorio para dos fuentes, es evidente la complejidad trigonométrica que presenta el problema de encontrar la función resultante de la interferencia de fuentes en una pantalla a medida que se aumenta el número de dichas fuentes. Pero como a los fines experimentales, sólo interesa poder determinar dónde se encuentran las franjas claras u oscuras y no cuál es exactamente la función $I=f(y)$, analizaremos solamente los valores característicos de la función I a través de un estudio fasorial simplificado.

7.2) ANALISIS DE LOS DIAGRAMAS DE INTENSIDAD PRODUCIDOS POR INTERFERENCIA LUMINOSA DE FUENTES PUNTUALES.

7.2.1) INTERFERENCIA DE DOS FUENTES PUNTUALES.

En realidad este caso es el que acabamos de estudiar en la experiencia de Young, pero veamos como llegamos a los mismos resultados sin necesidad de conocer la función $I=f(y)$.

Sean dos fuentes puntuales coherentes S_1 y S_2 separadas entre sí por una distancia d como se muestra en la figura 7.5

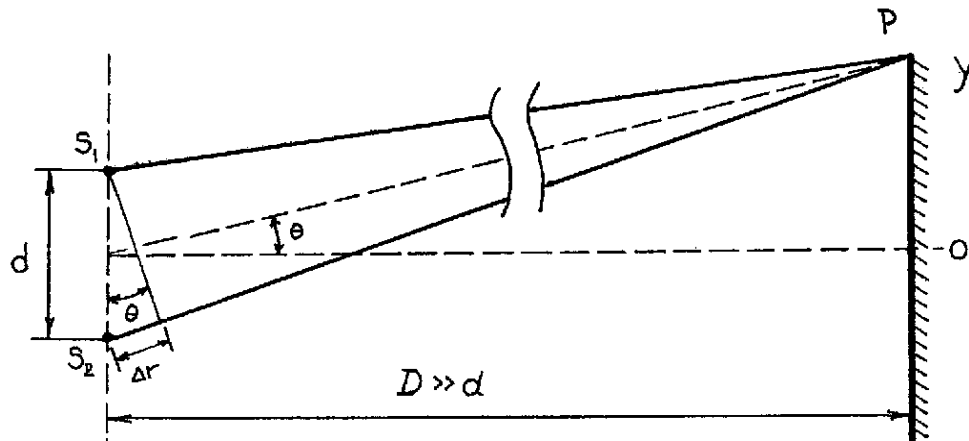
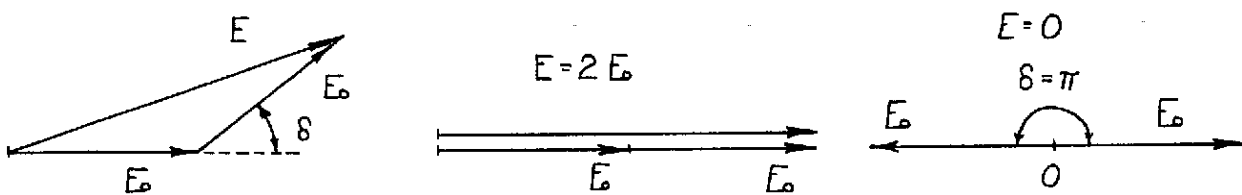


FIGURA 7.5

Las ondas procedentes de S_1 y S_2 llegan a los diferentes puntos P de la pantalla con un desfase $\Delta\phi$ (ó δ) dado por la diferencia de camino Δr tal como ya hemos analizado. La amplitud E en cada caso se obtendrá de la suma vectorial de las amplitudes de las ondas componentes (ver figura 7.6).



- (a) Para cualquier δ . (b) Para $\delta=2.m.\pi$ (c) Para $\delta=(2.m+1).\pi$

FIGURA 7.6 Amplitud de la onda resultante.

Si se observa la figura 7.6 se llega a la siguiente conclusión:

$$\delta = 2.m.\pi \quad \Rightarrow \quad E = 2.E_0 \quad \Rightarrow \quad I = C.(2.E_0)^2 = I_{\max}$$

$$\delta = (2.m=1) \cdot \pi \Rightarrow E = 0 \quad \Rightarrow \quad I = 0$$

$$\delta = k \cdot \Delta r = \frac{2 \cdot \pi}{\lambda} \cdot d \cdot \sin(\theta) = \frac{2 \cdot \pi}{\lambda} \cdot \frac{y \cdot d}{D}$$

Luego los resultados son los mismos que antes (ver figura 7.7 y comparar con la figura 7.4)

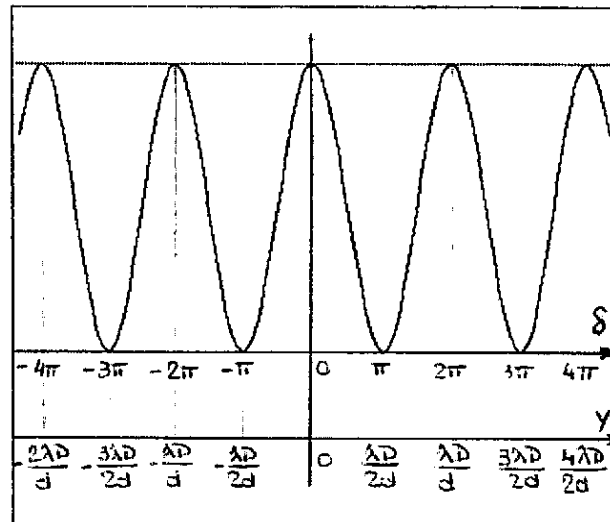


FIGURA 7.7

7.2.2) INTERFERENCIA DE TRES FUENTES PUNTUALES.

Sean tres fuentes puntuales coherentes S_1 , S_2 y S_3 separadas entre sí por la misma distancia d como se indica en la figura 7.8. Las ondas procedentes de S_1 , S_2 y S_3 llegan a un punto P cualquiera de la pantalla con una diferencia de fase $\Delta\alpha$ (ó $\delta\alpha$) constante dada por la diferencia de camino Δr que se muestra la figura 7.8.

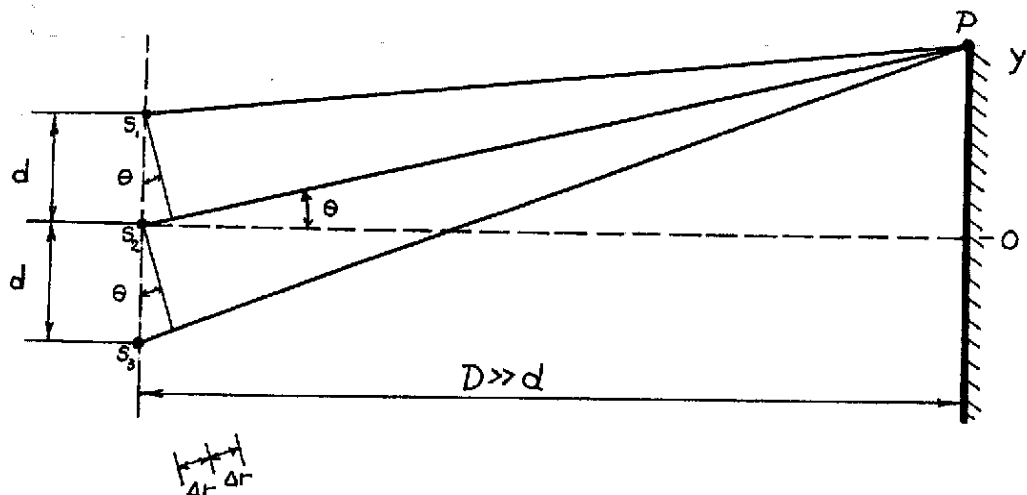
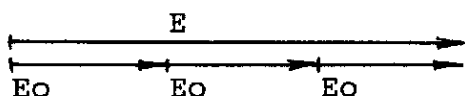


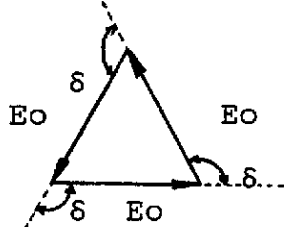
FIGURA 7.8 Interferencia de tres fuentes puntuales.

Los casos particulares a analizar serán todos los

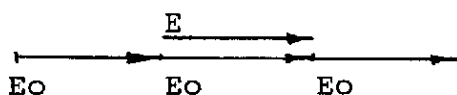
anteriores más aquellos que hagan cero la amplitud resultante y son los que se muestran en la figura 7.9.



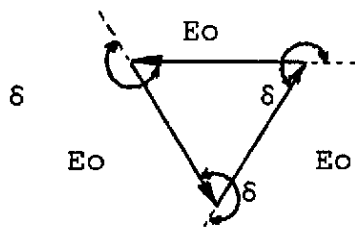
(a) Para $\delta = 2.m.\pi \Rightarrow E = 3.E_0 \quad I = C.(3.E_0)^2$ Máximo



(b) Para $\delta = (2\pi/3) + 2.m.\pi \Rightarrow E=0 \Rightarrow I=0$ Primer mínimo



(c) Para $\delta = (2.m+1).\pi \Rightarrow E=E_0 \Rightarrow I = C.E_0^2$ Máximo relativo



(d) Para $\delta = (4\pi/3) + 2.m.\pi \Rightarrow E=0 \Rightarrow I=0$ Segundo mínimo

FIGURA 7.9 Condiciones de interferencia para tres fuentes puntuales.

Generalizando las condiciones de máximo y de mínimo para tres fuentes puntuales :

Si $\delta = n . (2\pi/3)$ con n entero no múltiplo de 3 tenemos un mínimo ó como $y = (D\delta/2\pi d)$.

Si $y = (n\lambda D/3d)$ con n entero no múltiplo de 3 $\Rightarrow I = 0$
(mínimo de interferencia)

Si $\delta = 2.m.\pi$ ó $y = (m\lambda D/d)$ $I = c.(3E_0)^2 = I_{\max}$
(máximo de interferencia)

Si $\delta = (2.m+1).\pi$ $y = (2.m+1).(\lambda D/2.d)$ $I = C.E_0^2$
(máximo relativo)

Con estos resultados se pueden realizar los diagramas aproximados de $I=f(y)$ como se muestra en la figura 7.10.

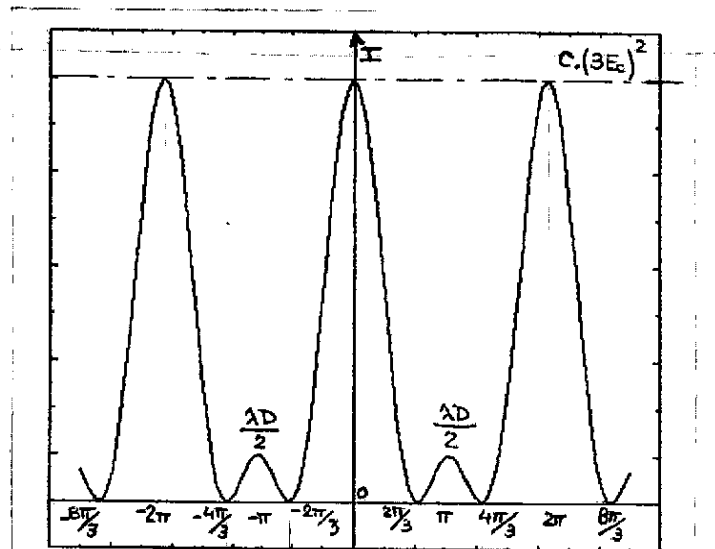


FIGURA 7.10 Intensidad en la pantalla producida por la interferencia de tres fuentes puntuales.

7.2.3) INTERFERENCIA DE N FUENTES PUNTUALES.

Generalizando los resultados obtenidos en los casos anteriores, sean n fuentes puntuales separadas entre sí una distancia d constante y una pantalla suficientemente alejada $D \gg d$ como se muestra en la figura 7.11.

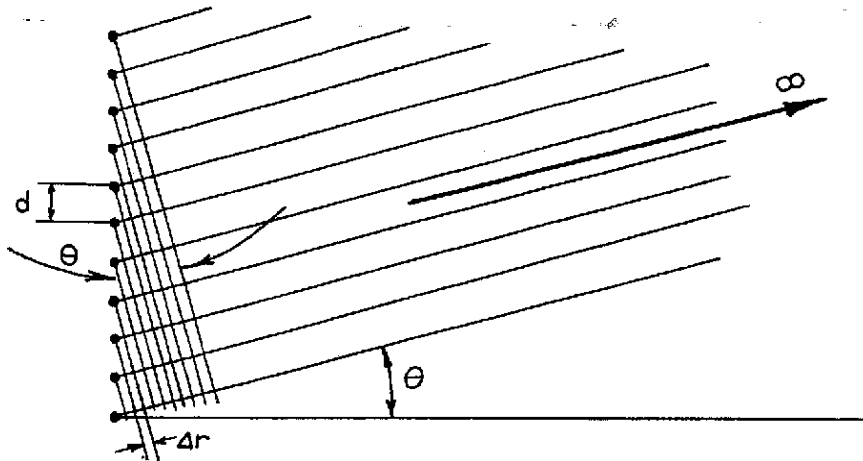


FIGURA 7.11 Interferencia de n fuentes puntuales.

En cada punto de la pantalla se tendrá la superposición de N ondas originadas en cada una de las n fuentes. La diferencia de fase entre dos fuentes sucesivas debido a la diferencia de camino Δr será constante como se muestra en la figura

$$\delta = k \cdot \Delta r = \frac{2\pi}{\lambda} \cdot d \cdot \sin(\theta) \quad \text{como siempre.}$$

Para los mínimos en la pantalla, la amplitud de la onda resultante de la onda debe ser cero. La condición del primer mínimo será que se cierre el polígono de n lados resultante de llevar uno a continuación de otro los fasores de cada fuente con el defasaje δ constante como se muestra en la figura 7.12, para que su suma de por resultado un fasor nulo ($A=0$).

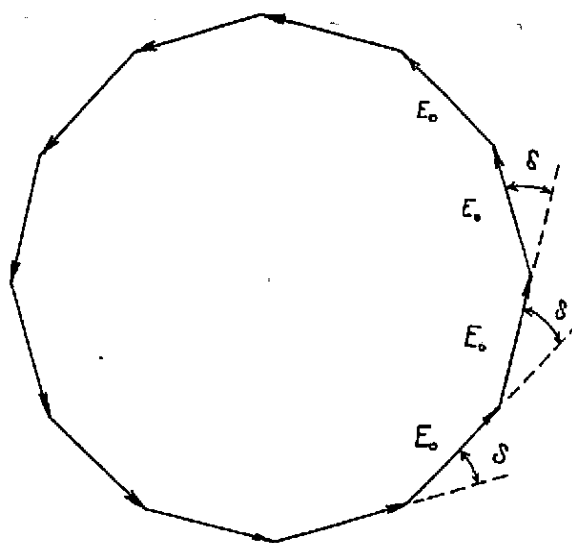


FIGURA 7.12

Como el polígono resultante de N lados es un polígono regular, la condición que debe cumplir para encontrar el primer mínimo será:

$$\delta = \frac{2\pi}{N} \quad \text{Primer mínimo.}$$

Pero, tendremos igualmente un mínimo ($A=0$) cada vez que obtengamos una poligonal cerrada.

Como vimos para dos fuentes los mínimos los teníamos en:

$$\frac{2\pi}{2} ; 3 \cdot \frac{2\pi}{2} ; 5 \cdot \frac{2\pi}{2} ; 7 \cdot \frac{2\pi}{2} \quad \text{o sea } m \cdot \frac{2\pi}{2}$$

siendo m números enteros no múltiplos de 2.

De manera análoga, para 3 fuentes los mínimos estaban en:

$$\frac{2\pi}{3} ; 2 \cdot \frac{2\pi}{3} ; 4 \cdot \frac{2\pi}{3} ; 5 \cdot \frac{2\pi}{3} \quad \text{o sea } m \cdot \frac{2\pi}{3}$$

siendo m números enteros no múltiplos de 3.

Generalizando para N fuentes, los mínimos serán para:

$$\delta = m \cdot \frac{2\pi}{n}$$

siendo m números enteros no múltiplos de n ó $y = \frac{D \cdot \lambda \cdot \delta}{2 \cdot \pi \cdot d}$
para pequeños valores de θ .

CONDICION DE MINIMO $\sin(\theta) = m \cdot \frac{\lambda}{N \cdot d}$ ó $y = \frac{m \cdot \lambda \cdot D}{2 \cdot \pi \cdot d}$ para $\theta \rightarrow 0$

para pequeños valores de θ , con m enteros no múltiplos de n .

CONDICION DE MAXIMO $\delta = 2n\pi$; $\sin(\theta) = \frac{n \cdot \lambda}{d}$ ó $y = \frac{n \cdot \lambda \cdot D}{d}$

Estos resultados permiten graficar aproximadamente la función $I=f(y)$ como se indica en la figura 7.13.

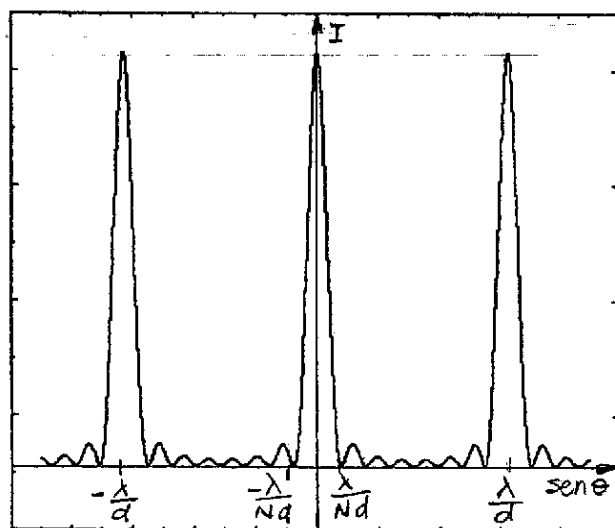


FIGURA 7.13.

El diagrama de la figura 7.13 corresponde a la interferencia de 8 fuentes puntuales.

Como se observa la separación de máximo principal a máximo principal es independiente del número de fuentes, solo depende de la separación d entre fuentes.

El ancho del máximo principal en cambio es de $\Delta\delta = 4\pi/N$ ó $\sin(\theta) = 2\lambda/Nd$ ó $\Delta y = 2\lambda d/Nd$, con lo que se concluye que la franja brillante se estrecha a medida que aumenta el número de fuentes y viceversa.

Además de máximo principal a máximo principal hay $N-1$ ceros y $N-2$ máximos relativos.

Por otro lado el valor máximo de la intensidad se mantiene constante e igual a $I = C \cdot (NE_0)^2 = I_{\text{máx}}$; vemos entonces que a medida que aumenta el número de fuentes aumenta también la intensidad máxima de la pantalla.

8) DIFRACCION DE LA LUZ.

Cuando se interpone en la trayectoria de un haz de rayos luminosos un diafragma o un orificio o ranura cuyo ancho ó diámetro no es muy superior a la longitud de onda del mismo, se produce el fenómeno de difracción. Es decir, la luz ya no continua su trayectoria rectilínea sino que se desvía iluminando la parte posterior de la barrera como se indica en la figura 7.1. Decimos entonces que la luz no produce sombra neta (es como si bordeara el objeto).

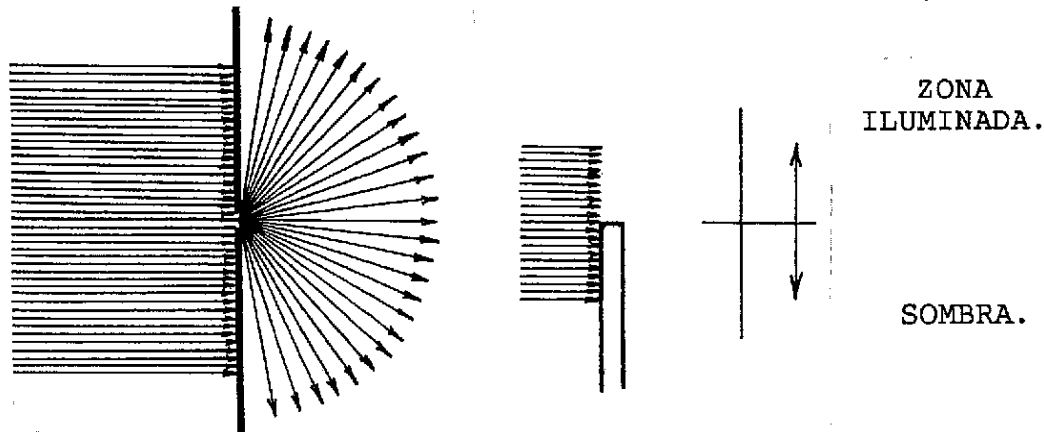


FIGURA 8.1 Difracción de la luz.

El fenómeno se explica aceptando que de acuerdo con el Principio de Huygens, cada uno de los puntos de la ranura se comporta como centro emisor de nuevas ondas y detrás de las ranuras tendremos entonces la superposición de las ondas originadas por todas estas fuentes.

La difracción es un fenómeno poco frecuente en la vida diaria debido a que el tamaño del orificio necesario para que se produzca debe ser muy pequeño. Sin embargo se deben a la difracción los siguientes fenómenos:

- * Que el perfil de una montaña aparezca iluminado en sus bordes pocos segundos antes de que el sol se eleve detrás de la misma.

- * Que las manchas producidas en el piso al atravesar la luz solar las copas frondosas de los árboles de un bosque sean siempre circulares.

- * Que si miramos el sol entrecerrando los ojos veamos una serie de rayas negras. (Son imágenes de las pestañas producidas por la difracción de la luz al penetrar en nuestra pupila).

- * Etc.

A continuación daremos una idea aproximada del estudio matemático de la difracción, porque su tratamiento riguroso es muy complejo.

8.1) DIFRACCION POR UNA RANURA.

Sea una ranura de ancho " a " alcanzada por un frente de onda

plano. Si la ranura está suficientemente alejada de la pantalla y a es lo suficientemente estrecha como para suponer que los infinitos rayos que salen de ella e inciden en un punto de la pantalla son paralelos entre sí, el problema se simplifica y la difracción que se observa se denomina de Fraunhofer en honor al científico que la estudió (ver figura 8.2).

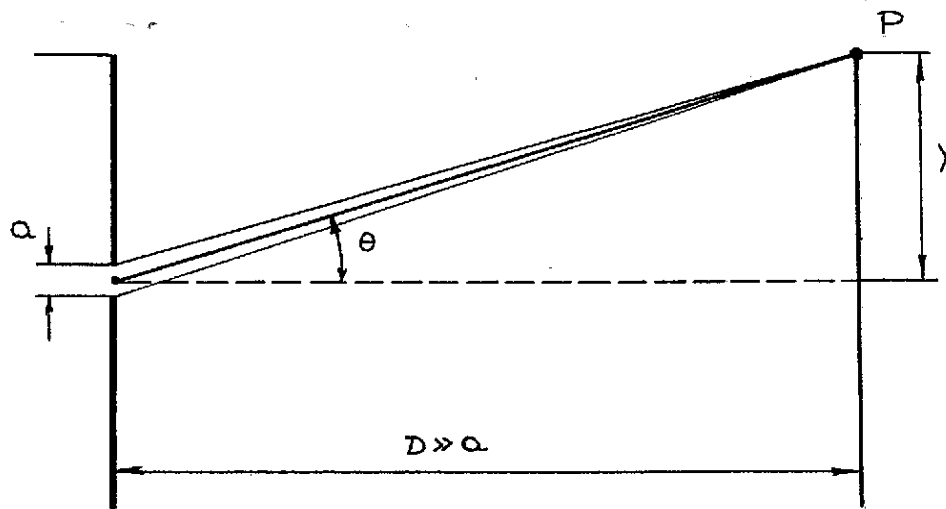


FIGURA 8.2 Difracción de Fraunhofer.

La solución entonces para obtener la distribución de intensidades en la pantalla es asimilarlo al caso de interferencia de N fuentes puntuales coherentes en que:

$N \rightarrow$ cantidad de fuentes.

$d \rightarrow$ separación entre fuentes.

$a = N \cdot d$ ancho de la ranura.

En el último punto del tema anterior vimos que la condición de mínimos de interferencia era:

$$\sin(\theta) = \frac{n \cdot \lambda}{N \cdot d} \quad \text{ó} \quad y = \frac{n \cdot \lambda \cdot D}{N \cdot d}$$

luego, en el caso del fenómeno de difracción $N \cdot d = a$, por lo tanto $\sin(\theta) = \frac{n \cdot \lambda}{a}$ ó $y = \frac{n \cdot \lambda \cdot D}{a}$ es la CONDICION DE MINIMO DE DIFRACCION CON $n=1,2,3,\dots$.

En cambio si $n=0$ se tiene el máximo central de difracción, el cual es un máximo absoluto que no puede ser alcanzado por ningún otro punto de la pantalla.

Esto se debe a que todos los infinitos rayos que llegan a él realizan el mismo camino óptico.

En cualquier otro punto de la pantalla, los rayos recorren un camino óptico diferente y es así como algunos interfieren destructivamente mientras otros lo hacen constructivamente.

El diagrama de intensidades es aproximadamente como se indica en la figura 8.3.

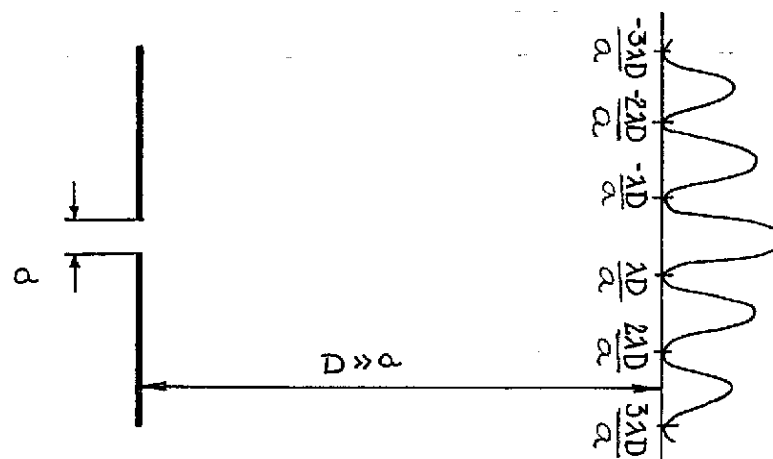


FIGURA 8.3 Intensidad en una pantalla iluminada por una ranura.

El estudio matemático riguroso del fenómeno de difracción es muy complicado, pero el resultado semicuantitativo al que arribamos explica las características fundamentales del fenómeno.

Vemos así que el ancho del máximo central aumenta a medida que aumenta o disminuye el ancho de la ranura a . Si el ancho de la ranura aumentara, el ancho del máximo central disminuye. Todos estos fenómenos son comprobables experimentalmente. En lugar de utilizar una fuente muy lejana, se suele colocar una fuente puntual en el plano focal de una lente convergente (colimador) que transforma el frente de onda esférico de la fuente en ondas planas (rayos paralelos), y del mismo modo los rayos difractados pueden recogerse en otra lente convergente observando la figura de difracción en el plano focal de la segunda lente (ver figura 8.4).

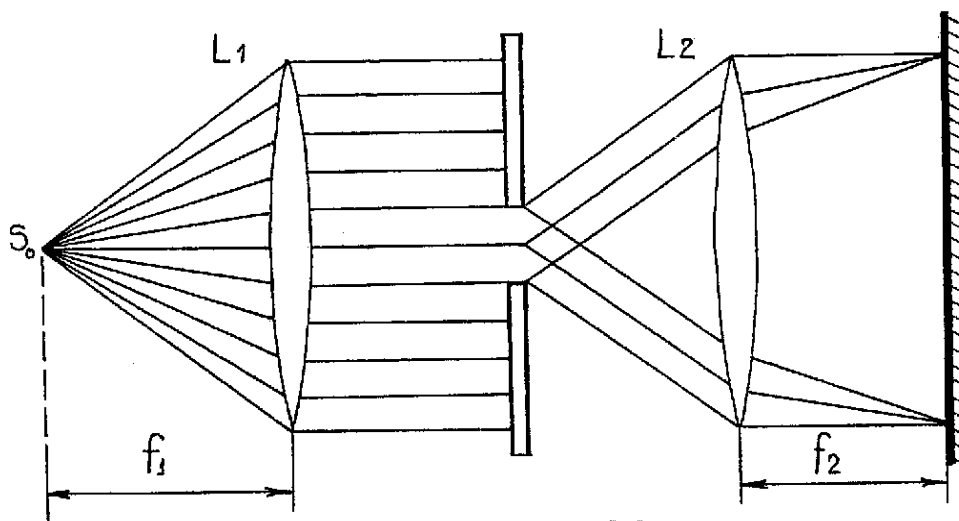


FIGURA 8.4

8.2) RED DE DIFRACCION.

Una red de difracción plana está constituida por un gran número de ranuras iguales y equidistantes en un mismo plano. Cada ranura se denomina raya o línea de la red. Se puede hacer una red de difracción sacando parte de los hilos de un trozo de seda, dejando sólo los hilos de un único sentido de la trama.

Se denomina constante de la red de difracción a la distancia que separa los puntos homólogos de dos ranuras sucesivas. Es decir, si c es el número de ranuras por unidad de longitud se cumpliría que $C = 1/d$, siendo C la constante de la red, y d la separación entre dos ranuras sucesivas.

Queremos analizar el diagrama de intensidades que se puede ver en una pantalla lo suficientemente alejada de la red, trataremos de hacerlo en forma aproximada.

Los diagramas de difracción de cada una de las ranuras que conforman la red son coincidentes, ya que todas las ranuras tienen el mismo ancho a . Pero además de la difracción propia de cada ranura, los rayos provenientes de las N ranuras que conforman la red al llegar a un punto P arbitrario de una pantalla lo suficientemente alejada interfieren entre sí como lo hacían las fuentes puntuales.

Luego EN UNA RED DE DIFRACCION SE SUPERPONEN LOS FENOMENOS DE INTERFERENCIA Y DIFRACCION.

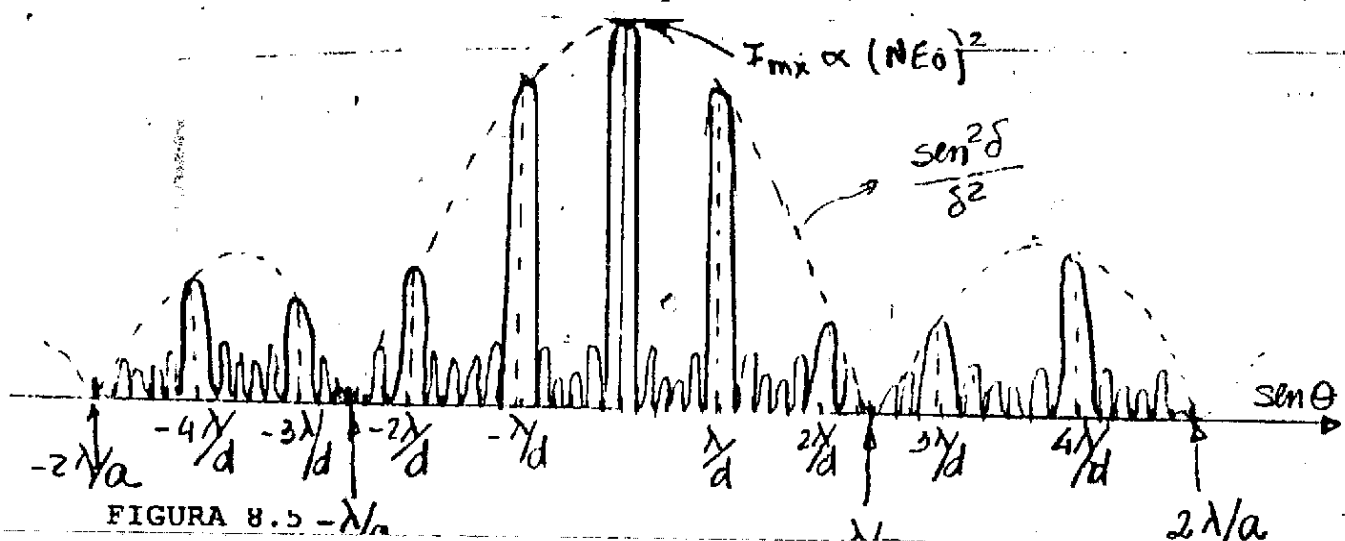
Las posiciones de los máximos son entonces independientes del número de ranuras y del ancho de cada ranura, ya que surgen del fenómeno de interferencia.

CONDICION DE MAXIMO: $\sin(\theta) = n\lambda/d$ y $y = n\lambda D/d$
(para $n=1,2,3,\dots$)

Condición de mínimos de interferencia: $\sin(\theta) = n\lambda/Nd$ y $y = n\lambda D/Nd$ con n distinto de múltiplo de N , siendo N el número de ranuras de la red y d la separación entre ranuras.

El ancho de los máximos principales será como en el fenómeno de interferencia de $2\lambda D/Nd$. Es decir, cuanto mayor sea el número de ranuras, mas angostos serán los máximos principales.

Además el fenómeno de interferencia se encuentra como modulado por el de difracción como puede observarse en la figura 8.5.



CEaD
2013