



Benemérita Universidad Autónoma de Puebla

Facultad de Ciencias Físico-Matemáticas
Licenciatura en Física

Estructuras hipercomplejas en teoría clásica de campos

Tesis presentada al
Colegio de Física

como requisito parcial para la obtención del grado de
Licenciado en Física

por
Oscar Meza Aldama

Asesorado por
Dr. Roberto Cartas Fuentevilla

Puebla, Pue.
Abril de 2015



Benemérita Universidad Autónoma de Puebla

Facultad de Ciencias Físico-Matemáticas
Licenciatura en Física

Estructuras hipercomplejas en teoría clásica de campos

Tesis presentada al
Colegio de Física

como requisito parcial para la obtención del grado de
Licenciado en Física

por
Oscar Meza Aldama

Asesorado por
Dr. Roberto Cartas Fuentesvilla

Puebla, Pue.
Abril de 2015

Título: Estructuras hipercomplejas en teoría clásica de campos.

Estudiante: Oscar Meza Aldama

Dr. J. Jesús Toscano Chávez
Presidente

Dr. Gerardo F. Torres del Castillo
Secretario

Dr. J. Lorenzo Díaz Cruz
Vocal

Dr. Roberto Cartas Fuentevilla
Asesor

Dr. Cupatitzio Ramírez Romero
Suplente

A mi familia: mis padres y mi hermano, que siempre estuvieron a mi lado a pesar de estar lejos.

Agradecimientos

A la Vicerrectoría de Investigación y Estudios de Posgrado de la Benemérita Universidad Autónoma de Puebla, por el apoyo económico que brindó para la realización de esta tesis. A los profesores de la FCFM, por todas sus enseñanzas y su apoyo, en especial al Dr. Torres, el Dr. Arturo Fernández y al Dr. Lorenzo. Al Dr. Cartas, del Instituto de Física, por la dirección de esta tesis. A mis amigos, en especial a Pablo, Jonathan, Leopoldo y Jaz. Finalmente, a mi familia, tanto la de Puebla como la de Torreón, pero sobre todo a mis padres y a mi hermano, que siempre me han dado su apoyo y cariño por todos los medios posibles.

Resumen

En este trabajo se propone la utilización de los números hiperbólicos en teoría del campo como una alternativa o, más bien, como una extensión de los números reales y complejos. Se estudia la estructura del vacío de una teoría formulada sobre estos números, y el respectivo rompimiento espontáneo de simetría, el cual soluciona el problema de masas no reales y genera la aparición de un bosón de Goldstone, al igual que en el caso complejo usual.

Introducción

Los números hiperbólicos y, en general, los bicomplejos, aparecieron por primera vez en un artículo de James Cockle en el año de 1848 [1]. A pesar de haber aparecido hace tanto tiempo y haber sido estudiados ampliamente por los matemáticos desde entonces, las aplicaciones físicas de estos números no han aparecido sino hasta décadas recientes, hecho extraño tomando en cuenta que dan una interpretación geométrica muy directa de, por ejemplo, las transformaciones de Lorentz en 1+1 dimensiones. Esto se debe, probablemente, a que los números hiperbólicos no poseen la estructura usual de campo, a la que los físicos estamos acostumbrados por el uso de sólo números reales y complejos. Sin embargo, trabajos recientes (ver [2]) han mostrado que muchas estructuras usadas en física, usualmente construidas sobre el conjunto de los números complejos, admiten expresiones en términos de números hiperbólicos; véanse, además, [3], [4] y [5] para estudiar mecánica cuántica relativista, la construcción de espinores, y una teoría de gravitoelectromagnetismo, respectivamente, todo sobre el álgebra de los números hiperbólicos.

En este trabajo se estudia la utilización de los números bicomplejos en la teoría de campo más sencilla: la de un campo escalar. En particular, veremos los efectos que provoca una extensión hiperbólica sobre la forma del vacío del modelo $\lambda\phi^4$, el cual es el modelo más sencillo que incorpora la posibilidad de tener escenarios con rompimiento espontáneo de simetría.

En el primer capítulo damos una introducción y una recopilación de resultados y propiedades generales de algunos sistemas de números, conocidos como números hipercomplejos, en particular de los números hiperbólicos. Las definiciones de esta parte son rigurosas, se construyen los sistemas de números hipercomplejos como álgebras reales de dimensión finita, en particular estudiamos las bidimensionales, y demostramos que esencialmente sólo hay tres de ellas. Damos luego un resumen de aplicaciones de los números hipercomplejos bidimensionales que han aparecido recientemente en artículos de física.

En el segundo, damos una introducción a las teorías clásicas de campos. Iniciamos con la definición del concepto de campo; luego repasamos ciertas propiedades de la teoría especial de la relatividad, en esta parte lo esencial es introducir la notación tensorial usada en la mayoría de los textos de teoría de campos y definir el grupo de Lorentz como el grupo de transformaciones ortogonales en el espacio de Minkowski. Luego estudiamos las propiedades que debe tener una teoría física (clásica) de campos, en particular aquellos conocidos como campos escalares, de los cuales resumimos sus propiedades más importantes y sus respectivos rompimientos espontáneos de simetría (global).

En el tercero, empezamos con una motivación en la cual la simetría hiperbólica surge de una manera puramente algebraica mediante la manipulación del potencial; después nos centramos en el estudio de las transformaciones hiperbólicas como un posible grupo de simetría de una teoría de campos. Se dan las definiciones precisas de números hiperbólicos, pero desde un punto de vista más práctico (no tan formal) y resumiendo sólo las propiedades que usaremos en el desarrollo del trabajo, tratando siempre de resaltar el paralelismo que tienen con los números complejos usuales (aunque también se indican explícitamente sus diferencias cuando surgen). Así como las fases complejas aparecen como las transformaciones del grupo $U(1)$, aparecerán, en el caso hiperbólico, un cierto tipo de “fases hiperbólicas”, los cuales se expresarán, al igual que en el caso complejo, como la exponencial de alguna unidad compleja por un parámetro real. Se extenderá posteriormente este conjunto de números al

conjunto más general de los números bicomplejos, el cual es, esencialmente, el producto directo del conjunto de números hiperbólicos por el de números complejos. A continuación se estudia el modelo $\lambda\phi^4$ del campo escalar hiperbólico. La conclusión más importante de esta sección, que marca un punto de quiebre entre el campo complejo y el hiperbólico, es que éste no se puede interpretar como dos campos reales interactuantes, a diferencia de aquel. Finalmente, en la última sección, trabajamos sobre la deformación bicompleja del modelo $\lambda\phi^4$; ésta es la sección más importante de todas. Veremos que la “norma de un número bicomplejo nos lleva a generalizar de los números reales a cierto subconjunto de los bicomplejos, cuyos elementos son siempre hermíticos. Así, permitiremos que los parámetros de nuestra teoría, en particular el parámetro de masa, sea un elemento de este conjunto de números hermíticos. Al final se estudian los rompimientos espontáneos de las dos simetrías (compleja usual y compleja hiperbólica) y se hacen comentarios sobre la geometría del vacío.

Índice general

Resumen	I
Introducción	II
1. Números complejos y sus aplicaciones en física	1
1.1. Clasificación de los sistemas de números complejos	1
1.2. Propiedades de los números parabólicos e hiperbólicos	2
1.3. Forma exponencial, grupos de Lorentz y de Galileo	7
1.4. Ecuaciones de onda sobre $\mathbb{D}$	9
2. Teoría del campo escalar	13
2.1. Características generales de los campos clásicos	13
2.2. El campo escalar real	19
2.3. El campo complejo	21
3. Simetría hiperbólica en la teoría del campo escalar	24
3.1. Motivación	24
3.2. Transformaciones hiperbólicas como un grupo de simetría	25
3.3. El modelo $\lambda\phi^4$ bicomplejo	28
Conclusiones	34
Bibliografía	35

Estructuras hipercomplejas en teoría clásica de campos

Oscar Meza Aldama

Abril de 2015

Capítulo 1

Números complejos y sus aplicaciones en física

1.1. Clasificación de los sistemas de números complejos

Sea V un espacio vectorial real; V es un álgebra [6] (real) si hay una operación definida en V , llamada multiplicación, tal que a cada par de vectores $u, v \in V$ les asigna otro vector $uv \in V$, y cumple las siguientes condiciones:

1. $(av)w = a(vw) = v(aw)$
2. $(u + v)w = uw + vw$
3. $u(v + w) = uv + uw$,

donde a es cualquier escalar real, y u, v, w son elementos cualesquiera de V . Si, además, se cumple que

4. $(uv)w = u(vw)$,

entonces V es un álgebra real asociativa. Note que las propiedades 2, 3 y 4 anteriores, junto con las de la suma de vectores que posee V por ser un espacio vectorial, le dan a V estructura de anillo asociativo. Así, un álgebra asociativa se puede definir, alternativamente (ver [7]), como un anillo asociativo que también es un espacio vectorial y que cumple con la propiedad 1 de arriba.

Se define [10] un sistema de números hipercomplejos como un álgebra (no necesariamente asociativa) real, de dimensión finita (vista como espacio vectorial), y con elemento unidad (vista como anillo). Primeramente, nos fijamos en sistemas de números complejos, es decir, aquellos que tienen dimensión dos como espacios vectoriales reales. Como bien sabemos, cualquier espacio vectorial real bidimensional es isomorfo a $\mathbb{R}^2$, así que podemos representar sus elementos por pares ordenados (a, b) , con $a, b \in \mathbb{R}$. Usualmente [8], el número complejo elíptico ¹ (a, b) se identifica con $a + ib$, donde i es la

¹Aquí, adelantando un poco la terminología, les llamamos números complejos elípticos a los complejos usuales (los que surgen de buscar una solución a la ecuación $x^2 = -1$). Posteriormente veremos que son un caso particular de los números hipercomplejos bidimensionales.

unidad imaginaria que cumple $i^2 = -1$. Aquí hacemos algo similar y escribimos, por comodidad y familiaridad con la notación, el número complejo (a, b) como $a + \epsilon b$, donde ϵ es cierta unidad imaginaria; es decir, ϵ no es un número real. Además, $\{1, \epsilon\}$ es una base de todo el sistema. Es fácil probar que, salvo isomorfismos, existen solamente tres sistemas de números complejos. Como la multiplicación es cerrada y $\{1, \epsilon\}$ es una base, podemos escribir, sin pérdida de generalidad,

$$\epsilon^2 = \alpha + 2\epsilon\beta, \quad (1.1)$$

donde α y β pueden ser cualesquiera números reales. Reacomodando y “completando el cuadrado”, tenemos

$$\begin{aligned} \epsilon^2 - 2\epsilon\beta &= \alpha \\ \epsilon^2 - 2\epsilon\beta + \beta^2 &= \alpha + \beta^2 \\ (\epsilon - \beta)^2 &= \alpha + \beta^2 \end{aligned} \quad (1.2)$$

Ahora, tenemos tres posibilidades: $\alpha + \beta^2 = 0$, $\alpha + \beta^2 < 0$ y $\alpha + \beta^2 > 0$. Para la primera de ellas, definimos $h \equiv \epsilon - \beta$, así que $h^2 = 0$. A este sistema de números complejos, cuya unidad imaginaria es nilpotente, se le llama sistema de números complejos parabólicos, o duales².

En el segundo caso, definimos $i \equiv \frac{\epsilon - \beta}{\sqrt{-\alpha - \beta^2}}$, entonces $i^2 = -1$. El sistema así obtenido no es otro más que el de los números complejos usuales, al cual, como ya se indicó en un pie de página, le llamamos sistema de números complejos elípticos.

Finalmente, en el tercer caso definimos $j \equiv \frac{\epsilon - \beta}{\sqrt{\alpha + \beta^2}}$, de manera que $j^2 = 1$. Éste es el sistema de números hiperbólicos, o dobles, el cual será nuestro tema principal de estudio. De aquí en adelante, usaremos el símbolo $\mathbb{D}$ para denotar al conjunto de los números hiperbólicos.

Así que sólo existen tres sistemas de números complejos: parabólicos, elípticos e hiperbólicos, obtenidos mediante la incorporación de una unidad imaginaria que cumpla $h^2 = 0$, $i^2 = -1$ ó $j^2 = +1$, respectivamente. La prueba anterior apareció en textos de matemáticas desde el siglo pasado, en trabajos de Kantor y Solodovnikov [10], y de Yaglom [11]. La sección siguiente es un resumen de las primeras secciones del texto de Yaglom. Las dos posteriores están basadas en los trabajos de Hucks [2] y Kocik [9].

1.2. Propiedades de los números parabólicos e hiperbólicos

Las propiedades de los números elípticos son bien conocidas, así que aquí nos concentraremos en los parabólicos e hiperbólicos.

Las operaciones aritméticas básicas de los números parabólicos están dadas por las siguientes fórmulas:

$$(\alpha_1 + h\beta_1) + (\alpha_2 + h\beta_2) = (\alpha_1 + \alpha_2) + h(\beta_1 + \beta_2), \quad (1.3)$$

$$(\alpha_1 + h\beta_1)(\alpha_2 + h\beta_2) = \alpha_1\alpha_2 + h(\alpha_1\beta_2 + \alpha_2\beta_1). \quad (1.4)$$

²Algunos autores (ver [9]) les llaman números duales a los que surgen de tomar la tercera posibilidad, $\alpha + \beta^2 > 0$; para evitar confusiones, trataremos de llamar a los distintos sistemas por los nombres alusivos a cónicas, es decir, elípticos, parabólicos e hiperbólicos.

El conjugado (parabólico) de $z = \alpha + h\beta$ se define como $\bar{z} \equiv \alpha - h\beta$. Su norma³ cuadrada, usando la fórmula anterior, está dada por

$$|z|^2 \equiv z\bar{z} = \alpha^2, \quad (1.5)$$

por lo tanto tenemos la siguiente propiedad

$$|z| = \frac{z + \bar{z}}{2} = \alpha. \quad (1.6)$$

Note que la norma de un número parabólico puede tener cualquier valor real (positivo, negativo o cero). Ahora definimos la división entre números parabólicos de manera similar a como se hace con los elípticos:

$$\frac{z_1}{z_2} \equiv \frac{z_1\bar{z}_2}{z_2\bar{z}_2} = \frac{\alpha_1}{\alpha_2} + h \frac{-\alpha_1\beta_2 + \alpha_2\beta_1}{\alpha_2^2}. \quad (1.7)$$

De ahí se ve que la división está bien definida sólo si el denominador tiene parte real, la cual coincide con su norma, distinta de cero, es decir, si $|z| \neq 0$.

Un número parabólico $h\alpha$ de norma nula se puede caracterizar por el hecho de que existe otro número, digamos $h\beta$, tal que su producto con $h\alpha$ es igual a cero:

$$(h\alpha)(h\beta) = (\alpha\beta)h^2 = 0. \quad (1.8)$$

Tales números se llaman *divisores de cero*. Por otro lado, los números con $|z| \neq 0$ se pueden escribir en “forma polar” de la siguiente manera:

$$z = \alpha + h\beta = \alpha(1 + h\omega), \quad (1.9)$$

donde, de nuevo, α es la norma del número parabólico, y al cociente $\omega \equiv \frac{\beta}{\alpha}$ se le llama *argumento* de z . Al igual que con los números elípticos, la forma polar resulta muy útil cuando se quieren multiplicar o dividir dos números. En efecto, tenemos que

$$\alpha_1(1 + h\omega_1)\alpha_2(1 + h\omega_2) = \alpha_1\alpha_2(1 + h(\omega_1 + \omega_2)), \quad (1.10)$$

$$\frac{\alpha_1(1 + h\omega_1)}{\alpha_2(1 + h\omega_2)} = \frac{\alpha_1}{\alpha_2}(1 + h(\omega_1 - \omega_2)). \quad (1.11)$$

Es decir que la norma del producto (cociente) de dos números parabólicos es igual al producto (cociente) de las normas de los factores, y su argumento es igual a la suma (diferencia) de los argumentos de los factores. Finalmente, de lo anterior se deduce el análogo a la fórmula de De Moivre

$$[\alpha(1 + h\omega)]^n = \alpha^n(1 + hn\omega). \quad (1.12)$$

Consideremos ahora los números hiperbólicos. Sea $z = \alpha + j\beta$, su conjugado se define como $\bar{z} \equiv \alpha - j\beta$ y su norma cuadrada es

$$|z|^2 \equiv z\bar{z} = \alpha^2 - \beta^2. \quad (1.13)$$

La norma se define como $\sqrt{|z|^2} = \sqrt{|\alpha^2 - \beta^2|}$; el signo de la raíz se toma como el signo de la componente de mayor valor absoluto. La suma, resta, multiplicación y división están dadas por las

³Veremos que esta “norma” no es definida positiva, así que en realidad no es una norma en el sentido estricto. Un término más adecuado sería módulo (que es el que se usa en textos de matemática pura), pero aquí nos quedamos, por comodidad, con el primer nombre. La misma situación se nos presentará un poco más adelante cuando definamos la “norma” de un número hiperbólico.

siguiente fórmulas:

$$(\alpha_1 + j\beta_1) \pm (\alpha_2 + j\beta_2) = (\alpha_1 \pm \alpha_2) + j(\beta_1 \pm \beta_2), \quad (1.14)$$

$$(\alpha_1 + j\beta_1)(\alpha_2 + j\beta_2) = (\alpha_1\alpha_2 + \beta_1\beta_2) + j(\alpha_1\beta_2 + \alpha_2\beta_1), \quad (1.15)$$

$$\frac{(\alpha_1 + j\beta_1)}{(\alpha_2 + j\beta_2)} = \frac{\alpha_1\alpha_2 - \beta_1\beta_2}{\alpha_2^2 - \beta_2^2} + j\frac{-\alpha_1\beta_2 + \alpha_2\beta_1}{\alpha_2^2 - \beta_2^2}, \quad (1.16)$$

de donde se ve que la división está definida sólo si $|z| \neq 0$. Los números de norma nula siempre son de la forma $x \pm jx$; a éstos se les llama *divisores de cero*.

Los números hiperbólicos de norma distinta de cero también se pueden expresar en “forma polar” de la siguiente manera. Sea $\rho \equiv \pm\sqrt{|\alpha^2 - \beta^2|}$, donde, de nuevo, el signo se toma de aquella componente con el mayor valor absoluto. Escribamos

$$z = \alpha + j\beta = \rho \left(\frac{\alpha}{\rho} + j\frac{\beta}{\rho} \right). \quad (1.17)$$

Debido a la definición de ρ , tenemos que $\left(\frac{\alpha}{\rho}\right)^2 - \left(\frac{\beta}{\rho}\right)^2 = \pm 1$, donde la fracción de mayor valor absoluto es siempre positiva. Entonces, recordando la relación $\cosh^2 x - \sinh^2 x = 1$, podemos escribir

$$\frac{\alpha}{\rho} = \cosh \theta, \quad \frac{\beta}{\rho} = \sinh \theta \quad \text{si } |z| > 0 \quad (1.18)$$

o

$$\frac{\alpha}{\rho} = \sinh \theta, \quad \frac{\beta}{\rho} = \cosh \theta \quad \text{si } |z| < 0. \quad (1.19)$$

Y por consecuencia,

$$z = \begin{cases} \rho(\cosh \theta + j \sinh \theta) & \text{si } |z| > 0, \\ \rho(\sinh \theta + j \cosh \theta) & \text{si } |z| < 0. \end{cases} \quad (1.20)$$

El número θ se llama el *argumento* del número hiperbólico z . La forma anterior para los números hiperbólicos resulta muy conveniente cuando se quieren multiplicar dos de ellos, al igual que sucede con los elípticos y los parabólicos. Tenemos que

$$\begin{aligned} z_1 z_2 &= \rho_1(\cosh \theta_1 + j \sinh \theta_1) \rho_2(\cosh \theta_2 + j \sinh \theta_2) \\ &= \rho_1 \rho_2 [\cosh(\theta_1 + \theta_2) + j \sinh(\theta_1 + \theta_2)] \end{aligned} \quad (1.21)$$

y

$$\begin{aligned} \frac{z_1}{z_2} &= \frac{\rho_1(\cosh \theta_1 + j \sinh \theta_1)}{\rho_2(\cosh \theta_2 + j \sinh \theta_2)} \\ &= \frac{\rho_1}{\rho_2} [\cosh(\theta_1 - \theta_2) + j \sinh(\theta_1 - \theta_2)], \end{aligned} \quad (1.22)$$

con expresiones similares para las otras posibilidades dadas por la ecuación (1.20). En palabras, tenemos que la norma del producto de dos números hiperbólicos es igual al producto de las normas de los factores, y su argumento es igual a la suma de los argumentos de los factores.

Introduzcamos la nueva base $\{\pi, \bar{\pi}\}$ de los números hiperbólicos, cuyos elementos se definen como⁴

$$\pi \equiv \frac{1}{2}(1 + j) \quad , \quad \bar{\pi} \equiv \frac{1}{2}(1 - j). \quad (1.23)$$

⁴En [9], se define esta base como $\pi \equiv \frac{1}{\sqrt{2}}(1 + j)$ y $\bar{\pi} \equiv \frac{1}{\sqrt{2}}(1 - j)$. El factor $\frac{1}{\sqrt{2}}$ da la apariencia de que estos números están “normalizados”, pero esto es falso ya que la norma de π y $\bar{\pi}$ es igual a cero. Además, el factor $\frac{1}{2}$ que aquí usamos resulta en expresiones algebraicas más simples, como se verá más adelante.

Es directo verificar las siguientes propiedades:

$$\begin{aligned}\pi^2 &= \pi, \\ \bar{\pi}^2 &= \bar{\pi}, \\ \pi\bar{\pi} &= 0, \\ \pi + \bar{\pi} &= 1,\end{aligned}\tag{1.24}$$

las cuales nos dicen que los elementos de la base nula forman un conjunto completo de operadores de proyección para los números hiperbólicos.

Como $\{\pi, \bar{\pi}\}$ es base de $\mathbb{D}$, podemos escribir cualquier $z \in \mathbb{D}$ como

$$z = \alpha\pi + \beta\bar{\pi},\tag{1.25}$$

con $\alpha, \beta \in \mathbb{R}$. La base nula permite multiplicar números hiperbólicos de una manera muy sencilla:

$$z_1 z_2 = (\alpha_1\pi + \beta_1\bar{\pi})(\alpha_2\pi + \beta_2\bar{\pi}) = \alpha_1\alpha_2\pi^2 + \beta_1\beta_2\bar{\pi}^2 + (\alpha_1\beta_2 + \alpha_2\beta_1)\pi\bar{\pi} = \alpha_1\alpha_2\pi + \beta_1\beta_2\bar{\pi},\tag{1.26}$$

es decir que las componentes del producto son simplemente los productos de las componentes de los factores. De la fórmula anterior y la propiedad $\pi + \bar{\pi} = 1$ se ve que el inverso de $z \in \mathbb{D}$, en términos de la base nula, es

$$z^{-1} = \alpha^{-1}\pi + \beta^{-1}\bar{\pi} = \frac{1}{\alpha}\pi + \frac{1}{\beta}\bar{\pi},\tag{1.27}$$

de lo cual volvemos a ver que el inverso está definido si y sólo si la parte real de z no es proporcional a su parte hiperbólica. El conjugado de $z \in \mathbb{D}$ se obtiene intercambiando sus componentes, es decir, si $z = \alpha\pi + \beta\bar{\pi}$, entonces

$$\bar{z} = \beta\pi + \alpha\bar{\pi},\tag{1.28}$$

por lo tanto, la norma de z está dada por

$$|z|^2 = z\bar{z} = (\alpha\pi + \beta\bar{\pi})(\beta\pi + \alpha\bar{\pi}) = \alpha\beta(\pi + \bar{\pi}) = \alpha\beta.\tag{1.29}$$

Para los números complejos elípticos, es usual escribir una función $f : \mathbb{R} \rightarrow \mathbb{C}$ como $f(t) = u(t) + iv(t)$, separando partes reales e imaginarias; o bien, una función $g : \mathbb{C} \rightarrow \mathbb{C}$ como $g(x, y) = u(x, y) + iv(x, y)$. En lo que sigue será útil hacer algo similar para una función con valores en $\mathbb{D}$, pero no separando en partes real e hiperbólica, sino en sus componentes respecto a la base nula. Esta separación resulta muy útil cuando se quiere extender el dominio de una función de los números reales a los hiperbólicos, tal y como veremos a continuación.

Sea f una función $f : \mathbb{R} \rightarrow \mathbb{D}$, la cual podemos escribir como

$$f(t) = u(t)\pi + v(t)\bar{\pi},\tag{1.30}$$

y supongamos que tiene el desarrollo en serie siguiente:

$$f(t) = \sum_{n=0}^{\infty} \frac{f^{(n)}(0)}{n!} t^n = \sum_{n=0}^{\infty} \frac{u^{(n)}(0)\pi + v^{(n)}(0)\bar{\pi}}{n!} t^n.\tag{1.31}$$

Podemos extender el dominio de f a $\mathbb{D}$ mediante este desarrollo simplemente permitiendo que el argumento de f sea ahora un número hiperbólico:

$$f(z) = \sum_{n=0}^{\infty} \frac{u^{(n)}(0)\pi + v^{(n)}(0)\bar{\pi}}{n!} z^n.\tag{1.32}$$

Expresando z en términos de $\{\pi, \bar{\pi}\}$ y usando (1.26), tenemos que

$$z^n = \alpha^n \pi + \beta^n \bar{\pi}; \quad (1.33)$$

por lo tanto

$$\begin{aligned} f(z) &= \sum_{n=0}^{\infty} \frac{1}{n!} \left[u^{(n)}(0) \alpha^n \pi + v^{(n)}(0) \beta^n \bar{\pi} \right] \\ &= \left(\sum_{n=0}^{\infty} \frac{u^{(n)}(0)}{n!} \alpha^n \right) \pi + \left(\sum_{n=0}^{\infty} \frac{v^{(n)}(0)}{n!} \beta^n \right) \bar{\pi} \\ &= u(\alpha) \pi + v(\beta) \bar{\pi} \\ &= f(\alpha) \pi + f(\beta) \bar{\pi}, \end{aligned} \quad (1.34)$$

es decir que tenemos la siguiente propiedad de “linealidad”:

$$f(\alpha \pi + \beta \bar{\pi}) = f(\alpha) \pi + f(\beta) \bar{\pi}, \quad (1.35)$$

con $\alpha, \beta \in \mathbb{R}$.

Sea $z = t + jx$, es fácil verificar que, con respecto a la base nula, z se representa por el par ordenado (x_+, x_-) , donde $x_{\pm} = t \pm x$. Usando esto, definimos los siguientes operadores diferenciales:

$$\begin{aligned} \partial &\equiv 2 \frac{\partial}{\partial z} = 2 \left(\frac{\partial}{\partial t} + j \frac{\partial}{\partial x} \right) = 2 \left(\frac{\partial}{\partial t} + \frac{\partial}{\partial x}, \frac{\partial}{\partial t} - \frac{\partial}{\partial x} \right) \equiv 2(\partial_+, \partial_-), \\ \bar{\partial} &\equiv 2 \frac{\partial}{\partial \bar{z}} = 2 \left(\frac{\partial}{\partial t} - j \frac{\partial}{\partial x} \right) = 2 \left(\frac{\partial}{\partial t} - \frac{\partial}{\partial x}, \frac{\partial}{\partial t} + \frac{\partial}{\partial x} \right) \equiv 2(\partial_-, \partial_+). \end{aligned} \quad (1.36)$$

Una condición necesaria para que una función $\psi : \mathbb{D} \rightarrow \mathbb{D}$ sea holomorfa (en el sentido hiperbólico) es que dependa sólo de la variable hiperbólica z y no de su conjugada hiperbólica. Esto se puede escribir como

$$\bar{\partial} \psi = 0, \quad (1.37)$$

lo cual, haciendo $\psi(z) = u(z) + jv(z)$ y usando (1.36), implica

$$\left(\frac{\partial u}{\partial t} - \frac{\partial v}{\partial x} \right) + j \left(\frac{\partial v}{\partial t} + \frac{\partial u}{\partial x} \right) = 0; \quad (1.38)$$

igualando partes reales e hiperbólicas de ambos lados de la ecuación anterior obtenemos

$$\frac{\partial u}{\partial t} = \frac{\partial v}{\partial x} \quad , \quad \frac{\partial v}{\partial t} = \frac{\partial u}{\partial x}, \quad (1.39)$$

lo cual es el análogo de las condiciones de Cauchy-Riemann para una función de variable compleja usual. Las dos ecuaciones anteriores pueden combinarse y escribirse como

$$\frac{\partial^2 u}{\partial t^2} - \frac{\partial^2 u}{\partial x^2} = 0, \quad (1.40)$$

$$\frac{\partial^2 v}{\partial t^2} - \frac{\partial^2 v}{\partial x^2} = 0. \quad (1.41)$$

Vemos que una condición necesaria para que ψ sea analítica es que sus componentes satisfagan la ecuación de onda (unidimensional). En variable compleja usual, lo análogo es que si una función es analítica en cierto dominio D del plano complejo entonces sus componentes son armónicas en D .

1.3. Forma exponencial, grupos de Lorentz y de Galileo

Sea ϵ alguna de las unidades imaginarias definidas arriba. Podemos definir el número $e^{\epsilon\theta}$, con $\theta \in \mathbb{R}$, mediante el desarrollo en serie de potencias siguiente:

$$e^{\epsilon\theta} \equiv \sum_{n=0}^{\infty} \frac{(\epsilon\theta)^n}{n!} = 1 + \epsilon\theta + \epsilon^2 \frac{\theta^2}{2} + \epsilon^3 \frac{\theta^3}{3!} + \dots \quad (1.42)$$

Con $\epsilon = i$ recuperamos la conocida fórmula de Euler:

$$\begin{aligned} e^{i\theta} &= 1 + i\theta - \frac{\theta^2}{2} - i\frac{\theta^3}{3!} + \frac{\theta^4}{4!} + \dots \\ &= \left(1 - \frac{\theta^2}{2} + \frac{\theta^4}{4!} + \dots\right) + i\left(\theta - \frac{\theta^3}{3!} + \frac{\theta^5}{5!} + \dots\right) \\ &= \cos \theta + i \sin \theta. \end{aligned} \quad (1.43)$$

Si ahora tomamos $\epsilon = h$ y recordamos que $h^2 = 0$, tenemos simplemente

$$e^{h\theta} = 1 + h\theta. \quad (1.44)$$

Mientras que, para $\epsilon = j$, obtenemos

$$\begin{aligned} e^{j\theta} &= 1 + j\theta + \frac{\theta^2}{2} + j\frac{\theta^3}{3!} + \frac{\theta^4}{4!} + \dots \\ &= \left(1 + \frac{\theta^2}{2} + \frac{\theta^4}{4!} + \dots\right) + j\left(\theta + \frac{\theta^3}{3!} + \frac{\theta^5}{5!} + \dots\right) \\ &= \cosh \theta + j \sinh \theta, \end{aligned} \quad (1.45)$$

lo cual es una buena justificación para llamar “hiperbólico” a este sistema numérico.

A los elementos de la forma $e^{\epsilon\theta}$ les llamaremos fases (o versores) elípticos, parabólicos o hiperbólicos, según corresponda, y probaremos que representan rotaciones en $\mathbb{R}^2$, transformaciones de Galileo y transformaciones de Lorentz en 1 + 1 dimensiones, respectivamente. Para esto, combinemos las coordenadas espacial y temporal en una sola coordenada compleja $z = t + \epsilon x$, es decir, hacemos la identificación

$$(t, x) \leftrightarrow z \equiv t + \epsilon x. \quad (1.46)$$

(En cuanto a las dimensiones físicas, tenemos dos opciones: podemos trabajar en unidades naturales, en las cuales t y x tienen las mismas dimensiones y, por lo tanto, la unidad compleja ϵ es adimensional; o podemos usar el Sistema Internacional y suponer que ϵ tiene dimensiones de velocidad inversa; elegiremos la primera.) Veamos el efecto que tiene el multiplicar una fase parabólica por la coordenada z . Usando (1.44), tenemos

$$z' \equiv e^{h\theta} z = (1 + h\theta)(t + hx) = t + h(x + \theta t). \quad (1.47)$$

Con la identificación (1.46) vemos que $z' \leftrightarrow (t, x + \theta t)$, es decir,

$$\begin{cases} t' = t \\ x' = x + \theta t \end{cases}, \quad (1.48)$$

lo cual corresponde a una transformación de Galileo con velocidad $-\theta$. Note que las fases son unitarias en el sentido hiperbólico:

$$e^{h\theta} e^{-h\theta} = (1 + h\theta)(1 - h\theta) = 1 + h\theta - h\theta - h^2\theta^2 = 1, \quad (1.49)$$

de lo que se deduce que el inverso de $e^{h\theta}$ siempre está bien definido y es igual al conjugado parabólico de $e^{-h\theta}$. Además tenemos la propiedad asociativa

$$e^{h\theta_1} e^{h(\theta_2+\theta_3)} = e^{h(\theta_1+\theta_2)} e^{h\theta_3}, \quad (1.50)$$

así que los versores parabólicos forman un grupo, el cual es identificable con (isomorfo a) el grupo de Galileo en 1 + 1 dimensiones.

Ahora fijémonos en las transformaciones hiperbólicas. Usando (1.45),

$$z' \equiv e^{j\theta} z = (\cosh \theta + j \sinh \theta)(t + jx) = (t \cosh \theta + x \sinh \theta) + j(x \cosh \theta + t \sinh \theta). \quad (1.51)$$

Para hacer la fórmula anterior más clara, definamos

$$\cosh \theta \equiv \gamma \equiv \frac{1}{\sqrt{1 - \beta^2}}. \quad (1.52)$$

Recordando la relación $\cosh^2 \theta - \sinh^2 \theta = 1$,

$$\sinh^2 \theta = \cosh^2 \theta - 1 = \frac{1}{1 - \beta^2} - 1 = \frac{\beta^2}{1 - \beta^2}; \quad (1.53)$$

$$\sinh \theta = \frac{\beta}{\sqrt{1 - \beta^2}} = \beta \gamma. \quad (1.54)$$

En términos de γ y β tenemos $z' = \gamma(t + \beta x) + j\gamma(\beta t + x)$ o, usando (1.46),

$$\begin{cases} t' = \gamma(t + \beta x) \\ x' = \gamma(x + \beta t) \end{cases}, \quad (1.55)$$

lo cual corresponde a una transformación de Lorentz (propia y ortocrona) con velocidad $v = -\beta$. También tenemos las propiedades

$$e^{j\theta} e^{-j\theta} = 1, \quad (1.56)$$

$$e^{j\theta_1} e^{j(\theta_2+\theta_3)} = e^{j(\theta_1+\theta_2)} e^{j\theta_3}, \quad (1.57)$$

que muestran que las fases hiperbólicas forman un grupo, el cual identificamos con $SO_+^\uparrow(1,1)$, es decir, el conjunto de todas las transformaciones de Lorentz en 1 + 1 dimensiones que no incluyen una reflexión con respecto al eje espacial o temporal.

Finalmente, combinemos dos coordenadas espaciales en una sola coordenada elíptica, $z = x + iy$. Multiplicando por una fase,

$$e^{i\theta} z = (\cos \theta + i \sin \theta)(x + iy) = (x \cos \theta - y \sin \theta) + i(x \sin \theta + y \cos \theta); \quad (1.58)$$

obtenemos una rotación por un ángulo θ en el espacio euclídeo bidimensional. Es bien conocido el hecho de que el conjunto de fases elípticas forma el grupo llamado $U(1)$, el cual es el grupo de simetría interna de la electrodinámica cuántica, y que es identificable topológica y geoméricamente con S^1 . Este hecho y las similitudes que hay con las fases parabólicas e hiperbólicas hacen pensar en la posibilidad de que estas últimas puedan ser consideradas como elementos del grupo de simetría interna de alguna teoría de campos, y que sería interesante estudiar la geometría y topología tanto del grupo mismo como de la variedad del vacío de la teoría resultante. Esto se hará en detalle en el capítulo 3 para el caso hiperbólico.

1.4. Ecuaciones de onda sobre $\mathbb{D}$

Considere la ecuación de Schrödinger (en unidades naturales, en particular $\hbar = 1$, para una partícula de masa $m = 1$),

$$-\frac{1}{2}\nabla^2\Psi + V\Psi = i\frac{\partial}{\partial t}\Psi. \quad (1.59)$$

Si escribimos $\Psi = e^{R+iS}$, la ecuación anterior es equivalente al sistema

$$\begin{cases} -\frac{1}{2}\nabla^2 R - \frac{1}{2}(\nabla R)^2 + \frac{1}{2}(\nabla S)^2 + V = -\frac{\partial S}{\partial t} \\ -\frac{1}{2}\nabla^2 S - \nabla R \cdot \nabla S = \frac{\partial R}{\partial t} \end{cases}. \quad (1.60)$$

Ahora considere el análogo de la ecuación (1.59) sobre los números hiperbólicos

$$-\frac{1}{2}\nabla^2\Psi + V\Psi = j\frac{\partial}{\partial t}\Psi. \quad (1.61)$$

Si ahora hacemos

$$\Psi = e^{R+jS}, \quad (1.62)$$

obtendremos

$$\begin{cases} -\frac{1}{2}\nabla^2 R - \frac{1}{2}(\nabla R)^2 - \frac{1}{2}(\nabla S)^2 + V = \frac{\partial S}{\partial t} \\ -\frac{1}{2}\nabla^2 S - \nabla R \cdot \nabla S = \frac{\partial R}{\partial t} \end{cases}. \quad (1.63)$$

Las definiciones de la densidad y la corriente de probabilidad son similares al caso complejo usual⁵ (ver [12]):

$$\rho \equiv |\Psi|^2 = \Psi\Psi^*, \quad (1.64)$$

$$\vec{J} \equiv \frac{j}{2}(\Psi\nabla\Psi^* - \Psi^*\nabla\Psi). \quad (1.65)$$

Usando (1.62) obtenemos $\rho = e^{2R} > 0$, es decir que en el caso hiperbólico también tenemos una probabilidad definida positiva. Además tenemos

$$\vec{J} = -e^{2R}\nabla S, \quad (1.66)$$

entonces

$$\nabla \cdot \vec{J} = -e^{2R}(\nabla^2 S + 2\nabla S \cdot \nabla R), \quad (1.67)$$

por lo tanto

$$\nabla \cdot \vec{J} + \frac{\partial \rho}{\partial t} = -2e^{2R} \left(\frac{1}{2}\nabla^2 S + \nabla S \cdot \nabla R - \frac{\partial R}{\partial t} \right); \quad (1.68)$$

de ahí vemos que la segunda de las ecuaciones (1.63) es equivalente a la ecuación de continuidad (ya que si aquella ecuación se cumple entonces el miembro derecho de la ecuación anterior es igual a cero), es decir que la ecuación hiperbólica de Schrödinger implica la conservación local de probabilidad (tal y como sucede en el caso elíptico).

Si escribimos la ecuación de Schrödinger para la función de onda $\Psi' = e^{R'+iS'}$, con un potencial

⁵Sólo en esta sección, usaremos el asterisco para denotar conjugación en el sentido hiperbólico. También, usamos una flecha sobre una variable, contrariamente a las más populares negritas, para indicar que se trata de un vector.

U , y su análogo hiperbólico para $\Psi = e^{R+jS}$ con un potencial W , se ve que los sistemas (1.60) y (1.63) son equivalente si

$$W = U + 2\frac{\partial S'}{\partial t} + (\nabla S')^2, \quad (1.69)$$

es decir que las dos versiones difieren sólo en la interpretación del potencial. Ahora supongamos que la función de onda hiperbólica es separable:

$$\Psi(\vec{r}, t) = \psi(\vec{r})\varphi(t). \quad (1.70)$$

Sustituyendo en (1.61) obtenemos, para φ ,

$$\frac{j}{\phi} \frac{\partial \varphi}{\partial t} = -E, \quad (1.71)$$

donde $-E$ es una constante de separación. La solución a la ecuación anterior es, salvo una constante multiplicativa,

$$\varphi = e^{-jEt}. \quad (1.72)$$

(Note que para obtener un signo negativo en el exponente tuvimos que poner “a mano” el signo menos en la constante de separación.) Por otro lado, para ψ se llega al análogo de la ecuación de Schrödinger independiente del tiempo:

$$-\frac{1}{2}\nabla^2\psi + V\psi = -E\psi; \quad (1.73)$$

comparando con la versión elíptica

$$-\frac{1}{2}\nabla^2\psi' + U\psi' = E\psi', \quad (1.74)$$

vemos que coinciden si $V = U - 2E$, es decir que difieren sólo en la interpretación de la energía.

De todo lo anterior se ve que la ecuación de Schrödinger se puede proponer y resolver sobre $\mathbb{D}$ de una manera similar a como se hace en $\mathbb{C}$, además de que la interpretación de la función de onda como una amplitud de densidad de probabilidad es igualmente consistente. Sin embargo, como se hace notar en [9], las dos formulaciones no son completamente equivalentes. Considere, por ejemplo, el experimento de Young de la doble rendija. La intensidad I que incide sobre la pantalla detectora es (ver, por ejemplo, el capítulo 9 de [13] para revisar el análogo óptico del cálculo que lleva a esta fórmula)

$$I \propto \cos^2\left(\frac{\delta}{2}\right) \quad (1.75)$$

donde $\delta = \frac{a}{L}x$, a es la distancia entre las rendijas, L es la distancia a la pantalla, con $a \ll L$, y x es la distancia desde el centro de la pantalla. Un cálculo similar sobre $\mathbb{D}$ muestra que

$$I \propto \cosh^2\left(\frac{\delta}{2}\right), \quad (1.76)$$

lo cual tiene un *mínimo* (al contrario de lo expresado en [9]) en $x = 0$ y se incrementa sin límite cuando se incrementa el valor absoluto de x . Es claro de la gráfica de (1.75) que el resultado será el clásico patrón de interferencia, mientras que para (1.76) se obtendrá una intensidad pequeña en el centro del detector, la cual se irá aumentando mientras nos alejamos del centro y se hará muy grande en los extremos del mismo. De hecho, si hacemos la pantalla lo suficientemente larga obtendremos una intensidad tan grande como queramos; esta predicción es claramente errónea. La mecánica cuántica sobre $\mathbb{D}$ es, por lo tanto, no equivalente a su contraparte original sobre $\mathbb{C}$, y dará resultados incorrectos

para sistemas cuyo espacio de configuración tenga un grupo fundamental no trivial, como es el caso del experimento de Young.

Ahora pasemos a la ecuación de Dirac,

$$(i\gamma^\mu \partial_\mu - m)\psi = 0. \quad (1.77)$$

Trabajaremos en $1+1$ dimensiones, es decir que en la ecuación anterior $\mu = 0, 1$. Además utilizaremos la representación quirral (ver [14]), en la cual

$$\gamma^0 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad \gamma^1 = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \quad (1.78)$$

y, por definición,

$$\gamma^5 \equiv -\gamma^0 \gamma^1 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}. \quad (1.79)$$

El espinor ψ que aparece en (1.77) es de la forma

$$\psi = \begin{pmatrix} \psi_R \\ \psi_L \end{pmatrix}. \quad (1.80)$$

y depende de una sola coordenada espacio-temporal hiperbólica. Hagamos la identificación

$$\begin{pmatrix} \psi_R \\ \psi_L \end{pmatrix} \leftrightarrow (\psi_R, \psi_L) \equiv \psi' \quad (1.81)$$

del espinor de Dirac con el objeto hiperbólico ψ' expresado con respecto a la base nula, es decir, $\psi' = \psi_R \pi + \psi_L \bar{\pi}$. Denotemos por $*$ el *operador* de conjugación. Note que

$$\gamma^0 \psi = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} \psi_L \\ \psi_R \end{pmatrix} = \begin{pmatrix} \psi_R \\ \psi_L \end{pmatrix} \leftrightarrow (\psi_R, \psi_L) = *(\psi_L, \psi_R) = *\psi', \quad (1.82)$$

$$\gamma^5 \psi = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} \psi_L \\ \psi_R \end{pmatrix} = \begin{pmatrix} \psi_L \\ -\psi_R \end{pmatrix} \leftrightarrow (\psi_L, -\psi_R) = (1, -1)(\psi_L, \psi_R) = j\psi', \quad (1.83)$$

y, como es de esperarse debido a que $\gamma^1 = -\gamma^0 \gamma^5$,

$$\begin{aligned} \gamma^1 \psi &= \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \begin{pmatrix} \psi_L \\ \psi_R \end{pmatrix} = \begin{pmatrix} \psi_R \\ -\psi_L \end{pmatrix} \leftrightarrow (\psi_R, -\psi_L) \\ &= (1, -1)(\psi_R, \psi_L) = (1, -1) * (\psi_L, \psi_R) = j * \psi' = - * j \psi'. \end{aligned} \quad (1.84)$$

Es decir que, bajo la identificación (1.81), la acción de γ^0 , γ^1 y γ^5 sobre ψ corresponde a la acción de $*$, $- * j$ y j , respectivamente, sobre ψ' .

Regresando a la ecuación de Dirac, si tomamos $m = 0$ obtenemos

$$\gamma^\mu \partial_\mu \psi = 0, \quad (1.85)$$

o bien, desarrollando la suma sobre μ y usando (1.81),

$$\begin{aligned} 0 &= \begin{pmatrix} 0 & \partial_0 + \partial_1 \\ \partial_0 - \partial_1 & 0 \end{pmatrix} \begin{pmatrix} \psi_R \\ \psi_L \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} \partial_0 - \partial_1 & 0 \\ 0 & \partial_0 + \partial_1 \end{pmatrix} \begin{pmatrix} \psi_R \\ \psi_L \end{pmatrix} \\ &= \gamma^0 (\partial_0 - \gamma^5 \partial_1) \psi' \leftrightarrow *(\partial_0 - j \partial_1) \psi' = * \bar{\partial} \psi' = \partial \psi'^*. \end{aligned} \quad (1.86)$$

Tomando conjugado hiperbólico y quitando la prima,

$$\bar{\partial}\psi = 0, \tag{1.87}$$

lo cual nos dice que ψ' cumple con la ecuación de Dirac sin masa en $1 + 1$ dimensiones siempre que sea holomorfa (y que sea lo suficientemente bien comportada), es decir, siempre que dependa sólo de la coordenada z y no de su conjugada $\bar{z}$, lo cual es una simplificación muy grande con respecto al caso complejo usual, en donde la expresión $\gamma^\mu \partial_\mu \psi = 0$ nos impone ciertas condiciones extra sobre las componentes de ψ .

Capítulo 2

Teoría del campo escalar

En este capítulo damos una introducción a la teoría clásica de campos escalares reales y complejos y resumiremos sus propiedades esenciales.

2.1. Características generales de los campos clásicos

Un campo se define como una función que le asigna cierta magnitud a cada punto del espacio físico. Algunos ejemplos, de los más conocidos en mecánica clásica, son el potencial electrostático (escalar) y los campos eléctrico y gravitacional (vectoriales). Debemos también especificar a qué nos referimos cuando decimos “espacio físico”; en particular nos interesan los campos relativistas, es decir, aquellos que se transforman de una manera bien definida bajo transformaciones de Lorentz y rotaciones en el espacio tridimensional. Al trabajar en un escenario relativista (nos referimos sólo a la relatividad *especial*), nuestro espacio físico será el espacio-tiempo plano tetradimensional, llamado también espacio de Minkowski y denotado por $\mathcal{M}^4$. A continuación repasaremos algunas propiedades de la teoría de la relatividad especial.

A cada elemento x de $\mathcal{M}^4$ le llamaremos punto del espacio-tiempo, o evento, o suceso, y le asignaremos el cuadrivector $(x^0, x^1, x^2, x^3)^T$, donde $x^0 \equiv ct$, c es la velocidad de la luz y $\vec{x} \equiv (x^1, x^2, x^3)^T$ es el vector de posición del evento en un sistema de ejes rectangulares cartesianos. (La T en las expresiones anteriores denota transposición matricial, entonces estamos pensando en los cuadrivectores de posición como vectores columna.) Sabemos [15] que las coordenadas x' de un suceso medidas en un sistema¹ S' , el cual se mueve con una rapidez constante v en dirección x positiva con respecto a otro sistema S , están dadas por

$$\begin{aligned}x'^0 &= \gamma(x^0 - \beta x^1), \\x'^1 &= \gamma(x^1 - \beta x^0), \\x'^2 &= x^2, \\x'^3 &= x^3,\end{aligned}\tag{2.1}$$

¹En este contexto, cuando decimos “sistema” nos referimos a un sistema inercial de referencia. La definición de sistema inercial (la cual no daremos aquí) se puede dar de varias maneras, con distintos grados de rigor, pero en general se entiende que un sistema inercial es aquel en que una partícula libre se mueve con velocidad constante.

donde $\gamma \equiv 1/\sqrt{1-\beta^2}$, $\beta \equiv \frac{v}{c}$, y hemos supuesto que los ejes primados y no primados correspondientes son paralelos y que los orígenes de ambos sistemas coinciden en $x^0 = x'^0 = 0$, es decir, $t = t' = 0$. Las ecuaciones anteriores son llamadas usualmente transformaciones de Lorentz. Note que, definiendo la matriz $(\Lambda^\mu{}_\nu)$ mediante

$$\Lambda \equiv \begin{pmatrix} \gamma & -\beta\gamma & 0 & 0 \\ -\beta\gamma & \gamma & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}, \quad (2.2)$$

la transformación de Lorentz (2.1) anterior se puede escribir en notación matricial como $x' = \Lambda x$, o, en componentes y usando el convenio de suma sobre índices repetidos (el cual usaremos de aquí en adelante sin mención explícita), $x'^\mu = \Lambda^\mu{}_\nu x^\nu$. Los subíndices y superíndices griegos (μ, ν, ρ , etc.) siempre correrán de 0 a 3, mientras que los índices latinos lo harán sólo de 1 a 3, es decir, sólo sobre las componentes espaciales.

Se define el cuadrado de un cuadvivector como

$$x^2 \equiv (x^0)^2 - \vec{x}^2 = (x^0)^2 - (x^1)^2 - (x^2)^2 - (x^3)^2. \quad (2.3)$$

La expresión anterior se puede escribir de una manera distinta, a veces más útil. Primero se define el tensor métrico (de Minkowski) g como²

$$(g_{\mu\nu}) \equiv \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}. \quad (2.4)$$

Las componentes del inverso de g se escriben usualmente con superíndices, $g^{\mu\nu}$; evidentemente el inverso de g es él mismo, por lo tanto se tiene

$$(g^{\mu\nu}) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}. \quad (2.5)$$

Y, como uno es el inverso del otro, $g^{\mu\nu}g_{\nu\rho} = \delta^\mu_\rho$. Ahora, sea x^μ un cuadvivector³; definimos $x_\mu \equiv g_{\mu\nu}x^\nu$. En general, el tensor métrico $g_{\mu\nu}$ se usa para bajar índices y su inverso $g^{\mu\nu}$, para subirlos. Con esta nueva notación, el cuadrado de un cuadvivector se puede escribir simplemente como $x^2 = x^\mu x_\mu$.

Note que las transformaciones (2.1) dejan invariante el cuadrado de un cuadvivector: $x'^2 = x^2$. Esta es de hecho la propiedad fundamental de las transformaciones de Lorentz, la cual usaremos para definir transformaciones más generales. Definimos (ver [18]) el *grupo de Lorentz* como el conjunto de todas las transformaciones lineales, las cuales representamos mediante cierta matriz $\Lambda^\mu{}_\nu$, que actúan sobre el espacio de los cuadvivectores mediante

$$x^\mu \rightarrow x'^\mu = \Lambda^\mu{}_\nu x^\nu, \quad (2.6)$$

tales que

$$x'^\mu x'_\mu = x^\mu x_\mu. \quad (2.7)$$

²Usamos la convención adoptada en la mayoría de los textos de física de partículas y teoría cuántica de campos (ver, por ejemplo, [14] o [16]) de tomar el tensor métrico con signatura -2 , contraria a la de la mayoría de los libros de relatividad (ver, por ejemplo, [17]), en los cuales se define $g_{\mu\nu} \equiv \text{diag}(-1, 1, 1, 1)$.

³Aquí hay un poco de abuso de notación; estrictamente, x^μ representa la componente μ -ésima del cuadvivector x , pero es costumbre usar x^μ para denotar al cuadvivector completo. Las ambigüedades rara vez surgen y el lector debe estar atento a qué significado se le da a cada símbolo en cada contexto.

La expresión anterior se puede escribir como

$$x'^{\mu}x'_{\mu} = x'^{\mu}g_{\mu\nu}x'^{\nu} = \Lambda^{\mu}_{\rho}x^{\rho}g_{\mu\nu}\Lambda^{\nu}_{\sigma}x^{\sigma} = x^{\rho}g_{\rho\sigma}x^{\sigma}. \quad (2.8)$$

De la última igualdad, usando que el cuadvivector x es arbitrario, se obtiene que se debe cumplir

$$\Lambda^{\mu}_{\rho}g_{\mu\nu}\Lambda^{\nu}_{\sigma} = g_{\rho\sigma}. \quad (2.9)$$

De hecho, Λ^{μ}_{ν} representa una transformación de Lorentz si y sólo si se cumple (2.9). Es importante escribir esta condición ya que muchas propiedades surgirán de ella, como veremos a continuación.

Es fácil comprobar que el grupo de Lorentz efectivamente es un grupo, con la multiplicación dada por la composición de transformaciones. En efecto, sean Λ^{μ}_{ν} y Λ'^{μ}_{ν} , y sea $\Lambda''^{\mu}_{\nu} = \Lambda^{\mu}_{\rho}\Lambda'^{\rho}_{\nu}$; tenemos que

$$\begin{aligned} \Lambda''^{\mu}_{\rho}g_{\mu\nu}\Lambda''^{\nu}_{\sigma} &= (\Lambda^{\mu}_{\kappa}\Lambda'^{\kappa}_{\rho})g_{\mu\nu}(\Lambda^{\nu}_{\lambda}\Lambda'^{\lambda}_{\sigma}) = \Lambda'^{\kappa}_{\rho}(\Lambda^{\mu}_{\kappa}g_{\mu\nu}\Lambda^{\nu}_{\lambda})\Lambda'^{\lambda}_{\sigma} \\ &= \Lambda'^{\kappa}_{\rho}g_{\kappa\lambda}\Lambda'^{\lambda}_{\sigma} = g_{\rho\sigma}, \end{aligned} \quad (2.10)$$

o sea que se cumple (2.9) para Λ'' y por lo tanto ésta representa una transformación de Lorentz, por lo tanto la operación es cerrada. La asociatividad se cumple porque la composición de mapeos es asociativa. El elemento neutro está dado por la transformación identidad, la cual está representada por la matriz cuyas entradas son δ^{μ}_{ν} ; ésta es en efecto una transformación de Lorentz ya que $\delta^{\mu}_{\rho}g_{\mu\nu}\delta^{\nu}_{\sigma} = g_{\rho\sigma}$. Finalmente, notemos que (2.9) se puede escribir como $\Lambda_{\nu\rho}\Lambda^{\nu}_{\sigma} = g_{\rho\sigma}$; contrayendo con $g^{\rho\lambda}$ obtenemos $\Lambda_{\nu}^{\lambda}\Lambda^{\nu}_{\sigma} = \delta^{\lambda}_{\sigma}$, por lo tanto tenemos que la inversa de cualquier transformación de Lorentz existe y está dada por

$$(\Lambda^{-1})^{\mu}_{\nu} = \Lambda^{\mu}_{\nu}. \quad (2.11)$$

Se cumplen así todas las propiedades de grupo.

En forma matricial, (2.9) se escribe como $\Lambda^T g \Lambda = g$. Tomando determinante a ambos lados, usando que el determinante de la transpuesta de una matriz es igual al determinante de la matriz original y que el determinante de g es igual a -1 , tenemos $\det^2 \Lambda = 1$, por lo tanto

$$\det \Lambda = \pm 1. \quad (2.12)$$

A las transformaciones que cumplen $\det \Lambda = +1$ se les llama transformaciones (de Lorentz) *propias*, y a las que cumplen $\det \Lambda = -1$ se les dice *impropias*.

En (2.9), hay dos índices libres: ρ y σ . Escogiendo ambos iguales a cero, y recordando que $g_{00} = 1$ y $g_{ii} = -1$ (sin suma sobre i), tenemos

$$(\Lambda^0_0)^2 - \Lambda^i_0 \Lambda^i_0 = 1 \quad (\text{suma sobre } i) \quad (2.13)$$

$$\Rightarrow (\Lambda^0_0)^2 = 1 + (\Lambda^1_0)^2 + (\Lambda^2_0)^2 + (\Lambda^3_0)^2. \quad (2.14)$$

Como $(\Lambda^i_0)^2$ es mayor o igual que cero para cada i , tenemos que $(\Lambda^0_0)^2 \geq 1$, lo cual implica que $|\Lambda^0_0| \geq 1$, que a su vez nos deja dos opciones:

$$\Lambda^0_0 \geq 1 \quad \text{ó} \quad \Lambda^0_0 \leq -1. \quad (2.15)$$

A las que cumplen con la primera condición se les llama transformaciones *ortocronas*.

Debido a las propiedades anteriores, el grupo de Lorentz está formado por (al menos) cuatro componentes disconexas entre sí. De ahora en adelante, cuando digamos grupo de Lorentz nos referiremos al subgrupo de las transformaciones propias y ortocronas. Las otras componentes, claramente, no son subgrupos, ya que no contienen a la transformación identidad. Al grupo completo usualmente se le llama grupo de Lorentz extendido con paridad y reversión temporal.

Como mencionamos anteriormente, nos interesan campos que se transformen de una manera bien definida bajo la acción del grupo de Lorentz. El caso más simple es el de un campo escalar. En este contexto, el calificativo “escalar” quiere decir que el campo se transforma de manera trivial, es decir, que el valor del campo en cada punto del espacio-tiempo es el mismo para todos los observadores. Considere dos observadores $\mathcal{O}$ y $\mathcal{O}'$; $\mathcal{O}$, el primero de ellos denota el campo por ϕ y el segundo por ϕ' . Además, a cada punto del espacio-tiempo, $\mathcal{O}$ le asigna las coordenadas x y $\mathcal{O}'$ le asigna las coordenadas x' , donde, como ya sabemos, los cuadvectores primados y no primados están relacionados mediante la fórmula⁴ $x' = \Lambda^{-1}x$. Que el campo sea escalar quiere decir simplemente que

$$\phi'(x') = \phi(x), \quad (2.16)$$

es decir que las funciones ϕ y ϕ' se encargan de asignarle el mismo valor del campo a cada punto del espacio-tiempo.

La formulación de una teoría clásica de campos se inicia, como es frecuente en la mecánica clásica, definiendo la acción (o simplemente la lagrangiana) de nuestro sistema físico (véanse, por ejemplo, las secciones (9.3) y (9.4) de [20] para ver la construcción de la electrodinámica cuántica y cromodinámica cuántica, respectivamente, a partir de las lagrangianas correspondientes). En teoría de campos (formulada sobre el espacio de Minkowski de 4 dimensiones), la acción se escribe como

$$S = \int d^4x \mathcal{L}(\phi(x), \partial_\mu \phi(x)), \quad (2.17)$$

donde $\partial_\mu \equiv \frac{\partial}{\partial x^\mu}$ y $\mathcal{L}$ es llamada *densidad lagrangiana*, o simplemente lagrangiana, la cual depende usualmente sólo de los campos y de sus primeras derivadas⁵, y ϕ denota, en general, un conjunto de campos de algún tipo.

La acción (o la lagrangiana) es la cantidad física fundamental, a partir de la cual se pueden obtener todas las otras variables mecánicas, por lo tanto el primer problema que hay que resolver en teoría de campos es postular una acción apropiada. Hay algunos criterios que nos dicen qué propiedades debe tener una acción para describir una teoría física real, los cuales se traducen en algunas características que debe tener la lagrangiana correspondiente. Algunos de estos criterios son los siguientes⁶:

1. La acción S debe ser invariante de Lorentz. Esto va de acuerdo con la ley de la relatividad especial que nos dice que las leyes de la física deben lucir igual en todos los marcos inerciales de referencia. La acción será invariante (de Lorentz) si la lagrangiana en sí es invariante, ya que el “elemento de volumen” d^4x se transforma con el jacobiano de la transformación, que en este caso es igual al determinante de la matriz que representa a la transformación de Lorentz, el cual, como ya vimos, es igual a +1 (para transformaciones propias). Así que cada uno de los términos de la lagrangiana debe ser un escalar de Lorentz.

2. La acción debe ser *local*. Este es en realidad un principio físico muy importante y general; lo que nos quiere decir es que el comportamiento de una partícula debe depender sólo de sus alrededores cercanos, y no de las condiciones de porciones del espacio-tiempo muy lejanos a ella. En teoría de campos, la localidad se asegura proponiendo que la acción sea de la forma (2.17); un ejemplo de una

⁴Aquí seguimos la convención usada en la mayoría de los libros de texto de teoría de campo que consiste en considerar las transformaciones de Lorentz como transformaciones activas.

⁵En principio, $\mathcal{L}$ también podría depender de las coordenadas x^μ e incluso de derivadas de más alto orden de los campos, sin embargo, esto no es común y las teorías físicas fundamentales tienen lagrangianas con una dependencia sólo de los campos y sus primeras derivadas, como en la ecuación (2.17).

⁶Cabe mencionar que estos no son los criterios más generales posibles (por ejemplo, en [22] se construye un modelo que viola invariancia de Lorentz), sin embargo han sido muy útiles para la formulación de teorías sencillas de campos escalares y son los que usaremos nosotros.

teoría no local sería la dada por una acción de la forma

$$S = \int \int d^4x d^4y \mathcal{L}(\phi(x), \phi(y), \partial_\mu \phi(x), \partial_\mu \phi(y)), \quad (2.18)$$

en donde la lagrangiana tendría términos que dependen de puntos lejanos del espacio-tiempo, los cuales contribuirían a la acción.

3. La lagrangiana debe cumplir con el criterio de renormalizabilidad⁷. Esto requiere un poco más de explicación. Normalmente se utilizan unidades *naturales* (que nosotros mismos usaremos a partir de ahora), en las cuales se define $c = \hbar = 1$, es decir que las velocidades y los momentos angulares son adimensionales. Esto a su vez implica que $|d| = |t| = |E|^{-1} = |m|^{-1}$, donde $|d|$ representa una cantidad con unidades de longitud, $|t|$ de tiempo, $|E|$ de energía y $|m|$ de masa. Así, por ejemplo, el operador diferencial $\partial_\mu = \frac{\partial}{\partial x^\mu}$ tiene unidades de energía (o de masa). Recuerde que la acción tiene unidades de momento angular (esto se puede ver de la definición $S = \int L dt$ que se da en mecánica clásica para sistemas discretos) y es, por lo tanto, adimensional, así que el producto $\mathcal{L} d^4x$ es adimensional. Pero d^4x tiene unidades de masa a la menos cuatro, por lo tanto la lagrangiana debe tener unidades de masa a la cuarta potencia. Ahora, el criterio de renormalizabilidad de Dyson [19] nos dice que una teoría es renormalizable si los coeficientes de la lagrangiana tienen unidades de masa a una potencia no negativa. Veremos ejemplos de esto en las siguientes secciones.

4. La lagrangiana, y por consecuencia la acción, debe respetar ciertas simetrías. El punto 1 en realidad es un caso especial, sólo que ahí nos referimos sólo a simetrías espacio-temporales. En general, los campos que aparecen en una lagrangiana lo hacen de tal forma que ésta es invariante ante ciertas transformaciones de los campos. A tales transformaciones, que modifican la configuración de los campos ϕ pero no la de las coordenadas espacio-temporales x^μ , se les conoce como transformaciones *internas*, y a las correspondientes simetrías se les llama simetrías internas. Este es uno de los criterios más usados al proponer una lagrangiana para una teoría. En particular, para los campos escalares, nos interesarán las simetrías de reflexión y la simetría bajo el grupo $U(1)$. Las simetrías, por lo general, se eligen de manera que la fenomenología de la teoría resultante concuerde con los datos experimentales.

Supongamos que tenemos una acción que es invariante ante las transformaciones de algún grupo continuo. La versión infinitesimal de una de esas transformaciones se puede escribir como

$$\phi(x) \rightarrow \phi'(x) = \phi(x) + \delta\phi(x), \quad (2.19)$$

donde $\delta\phi(x)$ es una deformación infinitesimal del campo. La acción es invariante siempre y cuando la correspondiente variación de la lagrangiana sea una cuatridivergencia:

$$\delta\mathcal{L} = \partial_\mu \mathcal{I}^\mu(x). \quad (2.20)$$

En general, tomando en cuenta que $\mathcal{L}$ depende sólo de ϕ y sus primeras parciales, tenemos

$$\begin{aligned} \delta\mathcal{L} &= \frac{\partial\mathcal{L}}{\partial\phi} \delta\phi + \frac{\partial\mathcal{L}}{\partial\partial_\mu\phi} \delta(\partial_\mu\phi) = \frac{\partial\mathcal{L}}{\partial\phi} \delta\phi + \partial_\mu \left(\frac{\partial\mathcal{L}}{\partial\partial_\mu\phi} \delta\phi \right) - \left(\partial_\mu \frac{\partial\mathcal{L}}{\partial\partial_\mu\phi} \right) \delta\phi \\ &= \left(\frac{\partial\mathcal{L}}{\partial\phi} - \partial_\mu \frac{\partial\mathcal{L}}{\partial\partial_\mu\phi} \right) \delta\phi + \partial_\mu \left(\frac{\partial\mathcal{L}}{\partial\partial_\mu\phi} \delta\phi \right). \end{aligned} \quad (2.21)$$

Usando las ecuaciones de movimiento, el primer término se anula y tenemos que, igualando (2.20) y (2.21),

$$\partial_\mu \left(\frac{\partial\mathcal{L}}{\partial\partial_\mu\phi} \delta\phi \right) = \partial_\mu \mathcal{I}^\mu. \quad (2.22)$$

⁷Este es realmente un criterio cuántico y no tiene análogo clásico. No daremos una definición precisa de renormalización aquí, lo único que nos interesa es introducir el criterio de conteo de potencias. Ver en [14] o [18] las definiciones precisas y generales.

O bien, definiendo

$$j^\mu \equiv \frac{\partial \mathcal{L}}{\partial \partial_\mu \phi} \delta \phi - \mathcal{I}^\mu, \quad (2.23)$$

tenemos

$$\partial_\mu j^\mu = 0. \quad (2.24)$$

Si tenemos más de un campo (es decir, más de un grado de libertad), en el primer término del miembro derecho de (2.23) debe haber una sumatoria sobre cada uno de ellos.

Al cuadrivector j^μ se le llama corriente de Noether, o corriente conservada, o simplemente corriente. A partir de la ecuación anterior podemos encontrar una cantidad conservada. En efecto, definimos la carga Q como

$$Q \equiv \int j^0(x) d^3x, \quad (2.25)$$

donde la integral se extiende sobre todo el espacio. Note que, al haber integrado sobre las tres coordenadas espaciales, Q dependerá sólo de la coordenada temporal. Ésta es una constante de movimiento ya que su derivada se anula:

$$\frac{dQ}{dt} = \frac{d}{dt} \int j^0(x) d^3x = \int \partial_0 j^0(x) d^3x = - \int \nabla \cdot \vec{j}(x) d^3x = 0, \quad (2.26)$$

suponiendo que $\vec{j}$ tiende a 0 en el infinito.

El argumento anterior es llamado *teorema de Noether*⁸ (ver [20]), y nos permite obtener cantidades conservadas a partir de las simetrías de una lagrangiana. Es importante recalcar que, para eliminar el primer término de (2.21) fue necesario usar las ecuaciones de movimiento, así que la carga Q es constante sólo a lo largo de la “trayectoria” seguida por el sistema.

Como un ejemplo, supongamos que tenemos una acción invariante ante traslaciones infinitesimales:

$$x^\mu \rightarrow x'^\mu = x^\mu + \epsilon^\mu, \quad (2.27)$$

con ϵ^μ un cuadrivector constante infinitesimal arbitrario. Haciendo una expansión en serie de Taylor y quedándonos con los primeros dos términos obtenemos la variación del campo

$$\begin{aligned} \delta \phi(x) &= \phi(x') - \phi(x) = \phi(x + \epsilon) - \phi(x) \approx \phi(x) + \epsilon^\mu \partial_\mu \phi(x) - \phi(x) \\ &= \epsilon^\mu \partial_\mu \phi(x) = \epsilon_\nu \partial^\nu \phi(x). \end{aligned} \quad (2.28)$$

Es decir que la variación de un campo escalar, bajo una traslación infinitesimal de las coordenadas, está dada por $\delta \phi(x) = \epsilon^\mu \partial_\mu \phi(x)$. La lagrangiana es también un escalar, así que esperamos que se transforme de la misma forma. Escribimos entonces

$$\delta \mathcal{L} = \epsilon^\mu \partial_\mu \mathcal{L} = \partial_\mu (\epsilon^\mu \mathcal{L}). \quad (2.29)$$

Comparando la expresión anterior con (2.20), tenemos $\mathcal{I}^\mu = \epsilon^\mu \mathcal{L} = \epsilon_\nu g^{\mu\nu} \mathcal{L}$. Así, sustituyendo en la definición de corriente conservada,

$$j^\mu = \frac{\partial \mathcal{L}}{\partial \partial_\mu \phi} \epsilon_\nu \partial^\nu \phi - \epsilon_\nu g^{\mu\nu} \mathcal{L}. \quad (2.30)$$

Definimos el tensor de energía-momento (ver [20]), también llamado tensor de energía-impulso, como

$$T^{\mu\nu} \equiv \frac{\partial \mathcal{L}}{\partial \partial_\mu \phi} \partial^\nu \phi - g^{\mu\nu} \mathcal{L}. \quad (2.31)$$

⁸En honor a Emmy Noether, quien lo formuló en 1915.

Por lo tanto, $j^\mu = \epsilon_\nu T^{\mu\nu}$. Ahora, sabemos que la integral espacial de j^0 es constante en el tiempo, pero $j^0 = \epsilon_\mu T^{0\mu}$, así que, como las ϵ_μ son constantes arbitrarias, cada componente $T^{0\mu}$ integrada sobre todo el espacio debe ser constante. La primera de estas componentes se identifica con el hamiltoniano H del sistema:

$$H = \int T^{00} d^3x = \int \left(\frac{\partial \mathcal{L}}{\partial \dot{\phi}} \dot{\phi} - \mathcal{L} \right) d^3x, \quad (2.32)$$

donde $\dot{\phi} \equiv \partial_0 \phi = \partial^0 \phi$. Las otras tres componentes se identifican con las componentes del momento lineal:

$$P^i = \int T^{0i} d^3x = \int \left(\frac{\partial \mathcal{L}}{\partial \dot{\phi}} \partial^i \phi \right) d^3x. \quad (2.33)$$

Definiendo la densidad de momento canónico

$$\pi \equiv \frac{\partial \mathcal{L}}{\partial \dot{\phi}} \quad (2.34)$$

y la densidad hamiltoniana (o simplemente hamiltoniana)

$$\mathcal{H} \equiv \pi \dot{\phi} - \mathcal{L}, \quad (2.35)$$

las cargas asociadas a la invariancia bajo traslaciones temporales y espaciales se pueden escribir, respectivamente, como

$$H = \int \mathcal{H} d^3x \quad (2.36)$$

y

$$P^i = \int \pi \partial^i \phi d^3x. \quad (2.37)$$

2.2. El campo escalar real

Construyamos ahora una lagrangiana para un campo escalar real ϕ . Queremos tener simetría bajo la reflexión $\phi \rightarrow -\phi$, así que debemos usar sólo potencias pares de los campos. Una posibilidad es tomar simplemente

$$\mathcal{L}(\phi, \partial_\mu \phi) = \frac{1}{2} \partial^\mu \phi \partial_\mu \phi - \frac{1}{2} m^2 \phi^2, \quad m^2 > 0. \quad (2.38)$$

Al término $\frac{1}{2} \partial^\mu \phi \partial_\mu \phi$ se le llama término cinético (ya que es el análogo de $\frac{1}{2} m v^2$ para el caso discreto); todos los otros términos se identifican con un potencial con el signo cambiado. Es decir que en este caso el potencial, llamémoslo V , está dado por $V(\phi) = \frac{1}{2} m^2 \phi^2$. Es por eso que se pide que $m^2 > 0$; si se tuviera $m^2 < 0$ el potencial no estaría acotado por abajo. Podemos reescribir la lagrangiana como

$$\mathcal{L} = \frac{1}{2} \partial^\mu \phi \partial_\mu \phi - V(\phi), \quad V(\phi) = \frac{1}{2} m^2 \phi^2. \quad (2.39)$$

Analizando el término cinético de (2.38), llegamos a la conclusión de que el campo ϕ tiene unidades de masa. Por lo tanto, términos de la forma $a(\partial_\mu \phi \partial^\mu \phi)^2$ (o con potencias más altas) están descartados debido al criterio de renormalizabilidad, ya que el coeficiente a tendría unidades de masa a la menos cuatro. Por la misma razón no tenemos términos de la forma $b \phi \partial_\mu \phi \partial^\mu \phi$.

Sustituyendo la lagrangiana en la ecuación de Euler-Lagrange obtenemos la ecuación de movimiento

$$(\square + m^2)\phi = 0, \quad (2.40)$$

donde $\square \equiv \partial_\mu \partial^\mu$ es el llamado operador d'Alembertiano. La ecuación anterior es la llamada ecuación de Klein-Gordon (KG), en honor a Oskar Klein y Walter Gordon, quienes la propusieron en 1926, aunque fue descubierta por primera vez por Schrödinger en 1925, sin embargo la descartó porque no describía correctamente el espectro del átomo de hidrógeno.

Note que la ecuación de movimiento que se obtiene a partir de (2.38) tiene soluciones de onda plana:

$$\phi = A \cos(-\vec{k} \cdot \vec{x} + \omega t) + B \sin(-\vec{k} \cdot \vec{x} + \omega t), \quad (2.41)$$

donde $\vec{k}$ es un tri-vector de onda arbitrario y $\omega \equiv \sqrt{\vec{k}^2 + m^2}$. Como la ecuación KG es lineal, el principio de superposición es válido y las ondas planas se suman linealmente, es decir que no se dispersan. Por esta razón se dice que (2.38) es la lagrangiana de una teoría *libre*. Cuando la ecuación de movimiento es no lineal, entonces sí habrá dispersión y se dice que la teoría es *interactuante*.

En particular, al término $\frac{1}{2}m^2\phi^2$ se le llama término de masa. La razón viene de que, cuando la teoría se cuantiza, las partículas correspondientes al campo ϕ resultan tener masa m . Sin embargo, se puede dar un argumento semiclásico para convenirse de lo mismo. Para una partícula relativista de masa m , tenemos la relación $E^2 - \vec{p}^2 = m^2$ (donde E es la energía de la partícula, y $\vec{p}$ es su tri-momento); haciendo la sustitución canónica $\vec{p} \rightarrow -i\nabla$ y $E \rightarrow i\frac{\partial}{\partial t}$ y multiplicando por ϕ , y actuando sobre ϕ con el operador resultante, se obtiene precisamente la ecuación KG. Entonces el término $m\phi^2$ proviene de la masa invariante relativista, así como del término $\frac{1}{2}m^2\phi^2$ de la lagrangiana. Es por eso que a este último se le llama término de masa, y se dice también que m es el parámetro de masa o, simplemente, la masa. El término de masa tiene signo positivo en el potencial y, por lo tanto, signo negativo en la lagrangiana. Si esto no sucede así, es decir, si tenemos $\mathcal{L} = \frac{1}{2}\partial^\mu\phi\partial_\mu\phi + \frac{1}{2}m^2\phi^2$, entonces el segundo término no se puede interpretar como un término de masa.

Ahora pasemos a un caso un poco más interesante: el de una lagrangiana autointeractuante,

$$\mathcal{L} = \frac{1}{2}\partial^\mu\phi\partial_\mu\phi - \frac{1}{2}m^2\phi^2 - \frac{1}{4!}\lambda\phi^4; \quad (2.42)$$

en este caso, para que el potencial $V(\phi) = \frac{1}{2}m^2\phi^2 + \frac{1}{4!}\lambda\phi^4$ esté acotado inferiormente, necesitamos $\lambda > 0$, mientras que m^2 puede tener cualquier valor real. De hecho, si $m^2 < 0$, tendremos el caso descrito arriba, en el que el segundo término de la lagrangiana tendrá el signo incorrecto y no corresponderá a un término de masa. Este escenario, además, nos permitirá estudiar el fenómeno llamado rompimiento espontáneo de simetría.

En general, se dice que una simetría está rota cuando el vacío (es decir, el estado de energía

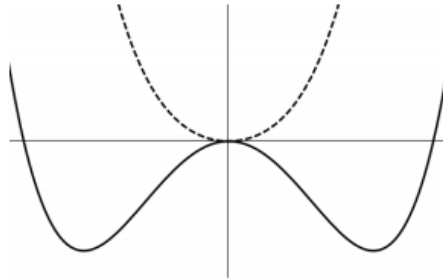


Figura 2.1: La línea continua muestra el potencial del campo escalar real cuando $m^2 > 0$ y la línea punteada, cuando $m^2 < 0$.

mínima) de una teoría no respeta las simetrías de la lagrangiana. En este caso, con $m^2 < 0$, la gráfica del potencial es la línea punteada de la figura 2.1, de donde se ve que el vacío está formado por dos puntos simétricamente colocados con respecto al eje vertical. La lagrangiana tiene simetría de reflexión, pero si escogemos uno de los dos puntos, el vacío ya no tendrá esa simetría.

Reescribamos el potencial como $V = -\frac{1}{2}\mu^2\phi^2 + \frac{1}{4!}\lambda\phi^4$, con $\mu^2 = -m^2 > 0$. Para encontrar una expresión analítica para el vacío, minimicemos V igualando a cero su primera derivada

$$0 = \frac{dV}{d\phi} = \left(-\mu^2\phi + \frac{1}{6}\lambda\phi^3 \right) \Big|_{\phi_{min}} = \phi_{min}(-\mu^2 + \frac{1}{6}\lambda\phi_{min}^2). \quad (2.43)$$

El valor $\phi_{min} = 0$ corresponde a un máximo local; los mínimos entonces están dados por

$$\phi_{min} = \pm v \quad , \quad v \equiv \frac{6\mu^2}{\lambda}. \quad (2.44)$$

Elegimos el valor positivo $\phi_{min} = +v$ y desarrollamos la lagrangiana alrededor de ese mínimo, es decir, hacemos el “corrimiento”

$$\phi \equiv \eta + v, \quad (2.45)$$

donde ahora η es nuestro nuevo campo dinámico, definido de tal manera que cuando $\eta = 0$, V sea mínimo; y luego sustituimos en la lagrangiana, quedándonos (ignorando términos constantes)

$$\mathcal{L} = \frac{1}{2}\partial^\mu\eta\partial_\mu\eta - \mu^2\eta^2 - \sqrt{\frac{\lambda\mu^2}{6}}\eta^3 - \frac{1}{4!}\lambda\eta^4. \quad (2.46)$$

Ahora el término de masa tiene el signo correcto. La lagrangiana ya no tiene una invariancia explícita bajo la transformación de reflexión, por lo tanto se dice que la simetría está *oculta*. En la siguiente sección, generalizamos al caso de una simetría continua.

2.3. El campo complejo

Ahora consideramos la siguiente lagrangiana de un campo escalar complejo (elíptico):

$$\mathcal{L} = \partial_\mu\phi^*\partial^\mu\phi - m^2\phi^*\phi \quad , \quad m^2 > 0. \quad (2.47)$$

De nuevo, a $\partial_\mu\phi^*\partial^\mu\phi$ y $m^2\phi^*\phi$ se les llama término cinético y término de masa, respectivamente. El potencial es ahora $V(\phi^*,\phi) = m^2\phi^*\phi$. La lagrangiana anterior es, claramente, invariante ante transformaciones de fase:

$$\begin{cases} \phi \rightarrow \phi' = e^{i\alpha}\phi \\ \phi^* \rightarrow \phi^{*'} = e^{-i\alpha}\phi^* \end{cases} \quad , \quad (2.48)$$

donde α es cualquier número real constante. Las fases complejas forman (o representan) un grupo continuo de transformaciones, denotado por $U(1)$, así que podemos encontrar una corriente de Noether y, de allí, una cantidad conservada. Si α es infinitesimal, podemos hacer un desarrollo en serie y, quedándonos sólo con los primeros dos términos, escribir las variaciones de los campos como

$$\begin{cases} \delta\phi \equiv \phi' - \phi = (1 + i\alpha)\phi - \phi = i\alpha\phi \\ \delta\phi^* \equiv \phi^{*'} - \phi^* = (1 - i\alpha)\phi^* - \phi^* = -i\alpha\phi^* \end{cases} \quad . \quad (2.49)$$

Como $\mathcal{L}$ es invariante, tenemos $\delta\mathcal{L} = 0$. Sustituyendo esto en las fórmulas de la primera sección, obtenemos la corriente (ignorando la α , que es una constante)

$$j^\mu = i(\phi\partial^\mu\phi^* - \phi^*\partial^\mu\phi). \quad (2.50)$$

Podemos comprobar directamente que j^μ es una corriente conservada, es decir, que $\partial_\mu j^\mu = 0$. Para esto, primero obtengamos las ecuaciones de movimiento, las cuales resultan ser idénticas a la ecuación KG de la sección anterior:

$$(\square + m^2)\phi = 0, \quad (2.51)$$

$$(\square + m^2)\phi^* = 0. \quad (2.52)$$

Multiplicando, por la izquierda, la primera ecuación por ϕ^* , y la segunda por ϕ , se obtiene

$$\phi^*\square\phi + m^2\phi^*\phi = 0, \quad (2.53)$$

$$\phi\square\phi^* + m^2\phi\phi^* = 0. \quad (2.54)$$

Restando (2.54) de (2.53), y recordando que $\square = \partial_\mu\partial^\mu$,

$$0 = \phi^*\partial_\mu\partial^\mu\phi - \phi\partial_\mu\partial^\mu\phi^* = \partial_\mu(\phi^*\partial^\mu\phi - \phi\partial^\mu\phi^*) = i\partial_\mu j^\mu, \quad (2.55)$$

que es lo que se quería probar. Note que en la demostración fue necesario usar las ecuaciones de movimiento. La carga correspondiente es simplemente la integral espacial de la componente temporal de j :

$$Q = i \int (\phi\dot{\phi}^* - \phi^*\dot{\phi})d^3x, \quad (2.56)$$

la cual puede se interpreta como la carga eléctrica. Por esa razón, a los campos complejos a veces se les llama campos cargados, mientras que a los reales, que no poseen esta carga conservada, se les llama campos neutros.

Ahora pasemos a una lagrangiana interactuante, con signo incorrecto en el parámetro de masa,

$$\mathcal{L} = \partial_\mu\phi^*\partial^\mu\phi - V(\phi, \phi^*) \quad , \quad V(\phi, \phi^*) = -\mu^2\phi^*\phi + \frac{\lambda}{4}(\phi^*\phi)^2, \quad (2.57)$$

la cual es invariante bajo transformaciones de fase, igual que antes. Ahora los mínimos del potencial están dados por

$$|\phi|_{min}^2 = v^2 \quad , \quad v \equiv \sqrt{\frac{2\mu^2}{\lambda}}; \quad (2.58)$$

por lo tanto

$$\phi_{min} = e^{i\theta}v \quad , \quad \theta \in \mathbb{R}. \quad (2.59)$$

Note que el vacío retiene la simetría bajo las transformaciones del grupo $U(1)$. Ahora tenemos una infinidad de mínimos posibles, uno distinto para cada valor de θ en el intervalo $[-\pi, \pi)$. La elección más simple es $\theta = 0$; así, escribiendo ϕ en términos de sus partes real e imaginaria,

$$\phi \equiv \frac{1}{\sqrt{2}}(\phi_1 + i\phi_2), \quad (2.60)$$

hacemos el corrimiento⁹

$$\begin{cases} \phi_1 \equiv \chi_1 + \sqrt{2}v \\ \phi_2 \equiv \chi_2 \end{cases}, \quad (2.61)$$

⁹Note el factor de $\sqrt{2}$ en la primera definición; en las referencias [20] y [21] no se pone, y los resultados a los que llegan son falsos.

donde χ_1 y χ_2 son, evidentemente, campos reales. Sustituyendo, obtenemos la lagrangiana

$$\begin{aligned} \mathcal{L} = & \left(\partial_\mu \chi_1 \partial^\mu \chi_1 - \mu^2 \chi_1^2 - \frac{\lambda}{16} \chi_1^4 \right) + \left(\partial_\mu \chi_2 \partial^\mu \chi_2 - \frac{\lambda}{16} \chi_2^4 \right) \\ & - \frac{\lambda}{8} \chi_1^2 \chi_2^2 - \frac{1}{2} \sqrt{\lambda \mu^2} (\chi_1^2 + \chi_2^2) \chi_1. \end{aligned} \quad (2.62)$$

Note que el campo χ_2 ha perdido su masa, mientras que χ_1 sigue siendo masivo. Se recomienda ver [20] para una interpretación de este resultado.

Capítulo 3

Simetría hiperbólica en la teoría del campo escalar

3.1. Motivación

La (densidad) lagrangiana de un campo escalar real autointeractuante se suele escribir como

$$\mathcal{L} = \frac{1}{2} \partial_\mu \phi \partial^\mu \phi - V(\phi) \quad , \quad V(\phi) = \frac{m^2}{2} \phi^2 + \frac{\lambda}{4!} \phi^4, \quad (3.1)$$

donde el parámetro de autointeracción λ se supone positivo para que el potencial V esté acotado por debajo. Debido a que $V(\phi)$ contiene sólo potencias pares de ϕ , es invariante bajo la transformación de reflexión $\phi \rightarrow -\phi$.

Podemos reescribir el potencial de la siguiente manera:

$$\begin{aligned} V &= \frac{m^2}{2} (i\phi)(-i\phi) + \frac{\lambda}{4!} [(i\phi)(-i\phi)]^2 \\ &= \frac{m^2}{2} \varphi \bar{\varphi} + \frac{\lambda}{4!} (\varphi \bar{\varphi})^2, \end{aligned} \quad (3.2)$$

donde hemos definido

$$\varphi \equiv i\phi \quad (3.3)$$

y la barra denota conjugación compleja. Note que $\varphi = -\bar{\varphi}$, ya que φ es puramente imaginario. Si consideramos a φ y $\bar{\varphi}$ como grados de libertad independientes, tenemos ahora que $V = V(\varphi, \bar{\varphi})$ es una función real que depende de dos variables puramente imaginarias. La gráfica de V se muestra¹ en la figura 1, en la cual resaltamos los potenciales usuales del modelo $\lambda\phi^4$. Se puede ver que los mínimos de V forman una hipérbola horizontal; de hecho, las curvas de nivel de esta superficie son siempre hipérbolas.

¹Por motivos de visualización, las gráficas como esta se muestran rotadas por un ángulo de 45°

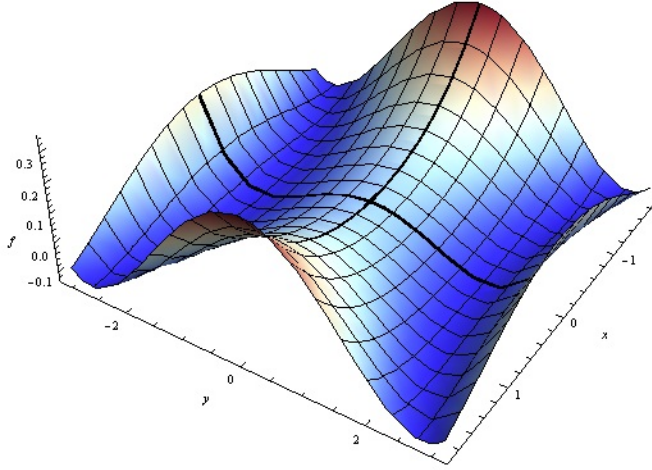


Figura 3.1: Potencial como función de la pareja $(\varphi, \bar{\varphi})$

Resulta interesante que esta simetría hiperbólica (es decir, la invariancia del valor de $\varphi\bar{\varphi}$ a lo largo de hipérbolas) surja de la “complejificación” (3.3) del campo escalar; ésto es una de nuestras motivaciones para estudiar las estructuras hiperbólicas.

3.2. Transformaciones hiperbólicas como un grupo de simetría

En esta sección repasamos las definiciones y propiedades de los números hiperbólicos que nos serán útiles en nuestros desarrollos. Además introducimos el sistema extendido de los números *bicomplejos*. Como ya mencionamos, un análisis más extenso se puede encontrar en [10] y [11].

En analogía con el conjunto $\mathbb{C}$ de los números complejos, se define $\mathbb{D}$, el conjunto de los números *hiperbólicos* (o “dobles”, de ahí la $\mathbb{D}$), como [9]

$$\mathbb{D} \equiv \{x + jy | x, y \in \mathbb{R}\}, \quad (3.4)$$

donde la “unidad hiperbólica” j se define como un número que tiene la propiedad $j^2 = +1$, pero $j \neq \pm 1$. Sea $z \in \mathbb{D}$, con $z = x + jy$, su conjugado (hiperbólico) es $\bar{z} \equiv x - jy$, y su norma está dada por $|z|^2 \equiv z\bar{z} = x^2 - y^2$. De ahí se ve que la norma de un número hiperbólico no es definida positiva; ésto es una diferencia importante con respecto a los números complejos usuales. El inverso de z es

$$z^{-1} \equiv \frac{\bar{z}}{|z|^2}. \quad (3.5)$$

De ahí que z tiene inverso si y sólo si su parte real no es igual a su parte hiperbólica ni al negativo de ésta. Como no todos los elementos de $\mathbb{D}$ tienen inverso, este conjunto no es un campo, sino un anillo conmutativo.

Resulta muy útil definir la “forma polar” de $z = x + jy \in \mathbb{D}$ como

$$z \equiv \rho e^{j\theta} \equiv \rho(\cosh \theta + j \sinh \theta), \quad (3.6)$$

con $\rho \in \mathbb{R}$ si $|z|^2 > 0$ y $\rho \in j \cdot \mathbb{R}$ si $|z|^2 < 0$, y además $\tanh \theta \equiv \frac{y}{x}$. Note que si $z = x \pm jx$ (es decir, $|z| = 0$) entonces $\tanh \theta = \pm 1 \Leftrightarrow \theta = \pm \infty$, así que los números hiperbólicos de norma nula no admiten una descomposición polar.

Para dos elementos z_1, z_2 de $\mathbb{D}$, no nulos (en el sentido de que su norma es distinta de cero), con $z_1 = \rho_1 e^{j\theta_1}$ y $z_2 = \rho_2 e^{j\theta_2}$, se cumple que su producto es $z_1 z_2 = \rho_1 \rho_2 e^{j(\theta_1 + \theta_2)}$. En particular, tenemos

$$e^{j\theta_1} e^{j\theta_2} = e^{j(\theta_1 + \theta_2)} \quad (3.7)$$

y

$$(e^{j\theta})^{-1} = e^{-j\theta}. \quad (3.8)$$

Los elementos de $\mathbb{D}$ de la forma $e^{j\theta}$ (con $\theta \in \mathbb{R}$) se denominan *versores* (hiperbólicos) y son el análogo de las fases unitarias complejas $e^{i\theta}$, aunque tienen diferencias sustanciales (véase figura 3.2). El conjugado de un versor es $e^{\bar{j}\theta} = e^{-j\theta} = \cosh \theta - j \sinh \theta$.

Ahora definimos los números *bicomplejos*, o complejos hiperbólicos, como [23]

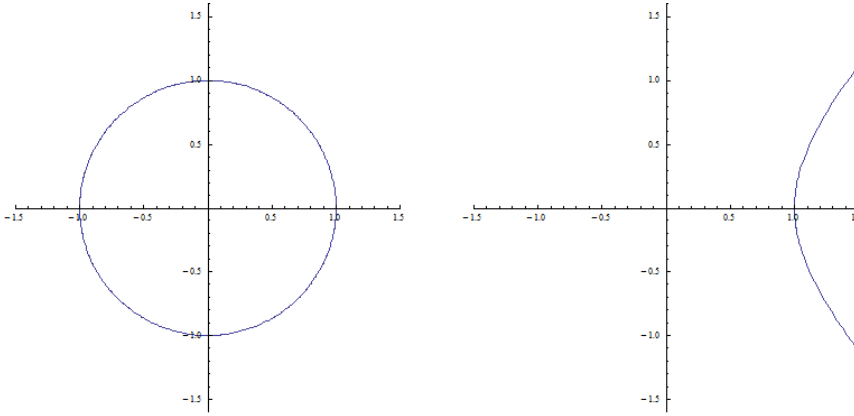


Figura 3.2: Fases complejas (izquierda), las cuales forman una variedad compacta; y versores hiperbólicos (derecha), los cuales forman una no compacta.

$$\mathbb{H} \equiv \{\xi + j\zeta | \xi, \zeta \in \mathbb{C}\} = \{x + iy + j(v + iw) | x, y, v, w \in \mathbb{R}\}, \quad (3.9)$$

con $i^2 = -1$ y $j^2 = +1$. La conjugación bicompleja ahora afecta tanto a la i como a la j (es posible definir otros tipos de conjugación que sólo afecten a una de las unidades complejas, véase por ejemplo [23], pero no serán necesarios para nuestros fines). Así, $\bar{z} \equiv x - iy - j(v - iw)$, de manera que

$$|z|^2 \equiv z\bar{z} = x^2 + y^2 - v^2 - w^2 + 2ij(xw - yv). \quad (3.10)$$

Note que, en general, $|z|^2$ no es un número real, pero es hermitico (es decir, invariante ante conjugación). Podemos ahora definir las “fases” bicomplejas $e^{i\alpha + j\beta}$ y sus conjugadas $e^{-i\alpha - j\beta}$. Explícitamente, tenemos

$$e^{i\alpha + j\beta} \equiv e^{i\alpha} e^{j\beta} = \cos \alpha \cosh \beta + i \sin \alpha \cosh \beta + j \cos \alpha \sinh \beta + ij \sin \alpha \sinh \beta. \quad (3.11)$$

Al igual que $\mathbb{D}$, $\mathbb{H}$ forma un anillo conmutativo. También tenemos que $|z|^2 = 0$ si y sólo si $z = x + iy \pm j(x + iy)$.

Sea $z = x + jy$, con $x^2 > y^2$ ($x^2 < y^2$); como vimos, $|z|^2 = x^2 - y^2$. Si definimos $z' \equiv e^{j\alpha}z$, es claro que $|z'|^2 = z^2$. Así que el conjunto de todos los números hiperbólicos de la forma $e^{j\alpha}z$, con $z \in \mathbb{D}$ y α corriendo sobre todos los reales, forma una rama de la hipérbola $x^2 - y^2 = \text{const}$, la cual es equilátera, horizontal (vertical) y centrada en el origen (véase figura 3.2). El hecho de que los versores transformen z a otros puntos de sólo una rama de la hipérbola es consecuencia de que el coseno hiperbólico nunca toma valores negativos. La transformación más general es de la forma $z' = \eta e^{j\alpha}z$, donde $\eta = \pm 1$ y $\alpha \in \mathbb{R}$. Con esto ya se pueden cubrir las dos ramas de la hipérbola.

Es claro que el conjunto de todos los versores hiperbólicos forma un grupo multiplicativo: la operación es cerrada debido a (3.7), la identidad es el número $e^{j0} = 1$, el inverso está dado por (3.8), y la asociatividad resulta de la asociatividad de los reales y de (3.7). La inclusión del factor $\eta = \pm 1$, introducido en el párrafo anterior, no afecta toda esta estructura. A este grupo (el que incluye el factor η) lo denotaremos $H(1)$ y le llamaremos grupo hiperbólico. Si no incluimos η , tenemos el grupo que denotaremos por $H^+(1)$. Cuando también tomamos en cuenta las fases complejas, tenemos elementos de la forma $e^{i\alpha}e^{j\beta}$, que forman el grupo $U(1) \times H(1)$.

Ahora podemos escribir y estudiar la lagrangiana de un campo escalar $\varphi \in \mathbb{D}$:

$$\mathcal{L} = \partial_\mu \varphi \partial^\mu \bar{\varphi} - m^2 \varphi \bar{\varphi} - \frac{\lambda}{2} (\varphi \bar{\varphi})^2, \quad (3.12)$$

donde la barra ahora indica conjugación en el sentido de j . La lagrangiana anterior es invariante de Lorentz, además de tener simetría interna bajo la acción del grupo $H(1)$. Escribiendo $\varphi = \varphi_1 + j\varphi_2$, donde $\varphi_1, \varphi_2 \in \mathbb{R}$, tenemos que el potencial está dado por

$$V(\varphi_1, \varphi_2) = m^2(\varphi_1^2 - \varphi_2^2) + \frac{\lambda}{2}(\varphi_1^2 - \varphi_2^2)^2 \quad (3.13)$$

Los puntos extremos de esta función son $\varphi_1 = \varphi_2 = 0$ y la región $m^2 \pm \lambda(\varphi_1^2 - \varphi_2^2) = 0$. Como la traza de la matriz hessiana de V , evaluada en $\varphi_1 = \varphi_2 = 0$, es igual a cero, éste es un punto silla (véase [26]). Por lo tanto los mínimos de V están en la segunda región. Supongamos que $m^2 < 0$; tenemos que el vacío está dado por

$$|\varphi|_{min}^2 = \frac{m^2}{\lambda} \equiv K^2, \quad (3.14)$$

el cual retiene la simetría bajo las transformaciones

$$\begin{cases} \varphi \rightarrow \varphi' = \eta e^{j\alpha} \varphi \\ \bar{\varphi} \rightarrow \bar{\varphi}' = \eta e^{-j\alpha} \bar{\varphi} \end{cases}, \quad (3.15)$$

con $\eta = \pm 1$ y $\alpha \in \mathbb{R}$. Para romper esta simetría, necesitamos escoger valores de φ_1 y φ_2 que satisfagan (3.14) y expresar la lagrangiana alrededor de este punto. Elegimos, por sencillez, $\eta = 1$ y $\alpha = 0$, obteniendo $\varphi_{1min} = K$ y $\varphi_{2min} = 0$. Definimos los campos reales χ_1 y χ_2 por

$$\begin{cases} \varphi_1 \equiv \chi_1 + K \\ \varphi_2 \equiv \chi_2 \end{cases}. \quad (3.16)$$

En términos de estas nuevas variables la lagrangiana queda como

$$\mathcal{L} = (\partial_\mu \chi_1 \partial^\mu \chi_1 - 2m^2 \chi_1^2 - \frac{\lambda}{2} \chi_1^4) + (-\partial_\mu \chi_2 \partial^\mu \chi_2 - \frac{\lambda}{2} \chi_2^4) + \lambda \chi_1^2 \chi_2^2 - 2\lambda K \chi_1 (\chi_1^2 - \chi_2^2). \quad (3.17)$$

Obtenemos un término cinético negativo para χ_2 , lo que podría dar, más adelante, problemas con su propagador (podría dar lugar a propagación acausal). Este comportamiento es independiente del

signo de m^2 .

De lo anterior vemos que el campo escalar hiperbólico es distinto al campo complejo usual en el sentido de que no se puede interpretar como dos campos reales interactuantes². Por lo tanto, si queremos evitar los términos cinéticos negativos, al romper la simetría hiperbólica no debemos “partir” φ en sus partes real e hiperbólica, como lo hacemos con el campo complejo, sino que nuestras variables dinámicas deben ser el campo y su conjugado hiperbólico.

Sea ψ_0 el valor esperado del vacío; tenemos las dos posibilidades $\psi_0 = K$ y $\psi_0 = jK$, dependiendo de si $m^2 < 0$ ó $m^2 > 0$, respectivamente. En cualquier caso, definimos los campos (hiperbólicos) ψ y su conjugado $\bar{\psi}$ por

$$\begin{cases} \varphi \equiv \psi + \psi_0 \\ \bar{\varphi} \equiv \bar{\psi} + \bar{\psi}_0 \end{cases} \quad (3.18)$$

Sustituyendo en la lagrangiana obtenemos

$$\begin{aligned} \mathcal{L} = & \partial_\mu \psi \partial^\mu \bar{\psi} - (m^2 + 2\lambda\psi_0\bar{\psi}_0)\psi\bar{\psi} - \frac{\lambda}{2}(\psi\bar{\psi})^2 - (m^2 + \lambda\psi_0\bar{\psi}_0)(\bar{\psi}_0\psi + \psi_0\bar{\psi}) \\ & - \frac{\lambda}{2}(\bar{\psi}_0^2\psi^2 + \psi_0^2\bar{\psi}^2) - \lambda\psi\bar{\psi}(\bar{\psi}_0\psi + \psi_0\bar{\psi}). \end{aligned} \quad (3.19)$$

Note que los términos del último renglón pueden ser hiperbólicos. También, en la expresión anterior, al parecer han sobrevivido términos lineales en los campos, sin embargo esos términos también se anulan debido a la restricción del vacío.

3.3. El modelo $\lambda\phi^4$ bicomplejo

Ahora estudiemos la lagrangiana

$$\mathcal{L} = \frac{1}{2}\partial_i \bar{z} \partial^i z - \frac{m^2}{2}\bar{z}z - \frac{\lambda}{4!}(\bar{z}z)^2, \quad (3.20)$$

donde ahora $z \in \mathbb{H}$ y el índice i puede correr desde 0 hasta $d-1$, donde d es la dimensión del espacio-tiempo sobre el cual formulemos la teoría. Recordemos que, en general, la norma de un número bicomplejo $z = x + iy + jv + i jw$ está dada por

$$|z|^2 = x^2 + y^2 - v^2 - w^2 + 2ij(xw - yv). \quad (3.21)$$

Las partes imaginarias e hiperbólicas puras ya no aparecen en la expresión, y las únicas sobrevivientes son reales e híbridas. En su totalidad, la expresión anterior es hermítica, es decir, es invariante bajo conjugación bicompleja. Esto nos sugiere que generalicemos del conjunto de los números reales al conjunto de los números bicomplejos hermíticos, es decir, aquellos con partes tanto real como híbrida. Así, además de admitir que en la lagrangiana aparezcan términos no reales, supondremos que la masa del campo bicomplejo puede ser a su vez hermítica en el sentido ya descrito.

Los mínimos del potencial los tomaremos como aquellos puntos en que el gradiente del potencial $V(z, \bar{z})$ se hace igual a cero. Así, obtenemos la restricción usual

$$\bar{z}_0 z_0 = -\frac{6m^2}{\lambda}, \quad (3.22)$$

²Los dos últimos términos de (3.17) sí representan interacciones entre χ_1 y χ_2 , pero el término cinético negativo de χ_2 no nos permite interpretarlo como un campo escalar real usual. Es en ese sentido en el que (3.17) no representa una lagrangiana de dos campos reales interactuantes.

donde z_0 y $\bar{z}_0$ representan los valores esperados en el vacío del campo bicomplejo z y su conjugado bicomplejo $\bar{z}$, respectivamente. Ahora hacemos el corrimiento

$$\begin{cases} z \rightarrow z + z_0 \\ \bar{z} \rightarrow \bar{z} + \bar{z}_0 \end{cases}, \quad (3.23)$$

y desarrollamos la lagrangiana alrededor de este mínimo. Usando la restricción que define al vacío la lagrangiana queda simplemente como

$$\mathcal{L} = \frac{1}{2} \partial_\mu \bar{z} \partial^\mu z - V(\bar{z}, z), \quad (3.24)$$

con

$$V(\bar{z}, z) = \frac{m^2}{2} \bar{z} z + \frac{\lambda}{4!} [(\bar{z}_0 z)^2 + (z_0 \bar{z})^2] + \frac{\lambda}{12} \bar{z} z (\bar{z}_0 z + z_0 \bar{z}) + \frac{\lambda}{4!} (\bar{z} z)^2. \quad (3.25)$$

Los dos primeros términos del potencial son cuadráticos en los campos, así que los términos de masa saldrán de allí y sólo de allí. Por otro lado, los términos tercero y cuarto contiene potencias cúbicas y cuárticas, respectivamente, de los campos, así que estos términos representan interacciones entre ellos. En particular nos interesarán los términos de masa.

Por simplicidad, para no trabajar con las 4 componentes del campo bicomplejo, podemos eliminar dos de ellas mediante las siguientes relaciones de proporcionalidad:

$$\begin{aligned} x &= \beta w, \\ y &= \beta v, \end{aligned} \quad (3.26)$$

de manera que

$$z = \beta(w + iv) + j(v + iw). \quad (3.27)$$

La expresión más general de los campos se obtendrá haciendo una “doble rotación” (elíptica e hiperbólica): $z \rightarrow e^{i\theta} e^{j\chi} z$. Escribamos entonces

$$\begin{aligned} z &= e^{i\theta} e^{j\chi} [\beta(w + iv) + j(v + iw)], \\ z_0 &= e^{i\theta_0} e^{j\chi_0} [\beta(w_0 + iv_0) + j(v_0 + iw_0)]. \end{aligned} \quad (3.28)$$

La norma cuadrada de z_0 es entonces

$$|z_0|^2 = \bar{z}_0 z_0 = (\beta^2 - 1)(v_0^2 + w_0^2) - 2ij\beta(v_0^2 - w_0^2). \quad (3.29)$$

Ahora la expresión que define al vacío,

$$z_0 \bar{z}_0 = -\frac{6m^2}{\lambda}, \quad (3.30)$$

debido a que tenemos contribuciones tanto reales como híbridas, es equivalente a las dos condiciones siguientes:

$$v_0^2 + w_0^2 = \frac{6m_1^2}{(1 - \beta^2)\lambda}, \quad (3.31)$$

$$v_0^2 - w_0^2 = \frac{3m_2^2}{\beta\lambda}, \quad (3.32)$$

donde hemos escrito $m^2 = m_1^2 + ijm_2^2$, es decir que permitimos que el parámetro de masa sea ahora un número bicomplejo hermítico. Por el contrario, seguiremos tomando el parámetro de autointeracción λ real y positivo.

Como dijimos, nos interesan en particular los términos de masa. Empecemos calculando el término $(\bar{z}_0 z_0)^2$. Usando la expresión (3.28), tenemos que

$$\begin{aligned}
(\bar{z}_0 z_0)^2 = & v^2 \{ [v_0^2 ((\beta^2 - 1)^2 - 4\beta^2) - w_0^2 (\beta^2 + 1)^2] \cos 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& - 2(\beta^4 - 1)v_0 w_0 \sin 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& + 4\beta(\beta^2 + 1)v_0 w_0 \cos 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \\
& + 4\beta(\beta^2 - 1)v_0^2 \sin 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \} \\
& + w^2 \{ [w_0^2 ((\beta^2 - 1)^2 - 4\beta^2) - v_0^2 (\beta^2 + 1)^2] \cos 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& + 2(\beta^4 - 1)v_0 w_0 \sin 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& + 4\beta(\beta^2 + 1)v_0 w_0 \cos 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \\
& - 4\beta(\beta^2 - 1)w_0^2 \sin 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \} \\
& + vv \{ 4v_0 w_0 (\beta^2 + 1)^2 \cos 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& + 2(\beta^4 - 1)(v_0^2 - w_0^2) \sin 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& - 4\beta(\beta^2 + 1)(v_0^2 + w_0^2) \cos 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \} \\
& + ijv^2 \{ -4\beta(\beta^2 - 1)v_0^2 \cos 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& + 4\beta(\beta^2 + 1)v_0 w_0 \sin 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& + 2(\beta^4 - 1)v_0 w_0 \cos 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \\
& + [v_0^2 ((\beta^2 - 1)^2 - 4\beta^2) - w_0^2 (\beta^2 + 1)^2] \sin 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \} \\
& + ijw^2 \{ 4\beta(\beta^2 - 1)w_0^2 \cos 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& + 4\beta(\beta^2 + 1)v_0 w_0 \sin 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& - 2(\beta^4 - 1)v_0 w_0 \cos 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \\
& + [w_0^2 ((\beta^2 - 1)^2 - 4\beta^2) - v_0^2 (\beta^2 + 1)^2] \sin 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \} \\
& + ijvw \{ -4\beta(\beta^2 + 1)(v_0^2 + w_0^2) \sin 2(\theta - \theta_0) \cosh 2(\chi - \chi_0) \\
& + 2(\beta^4 - 1)(w_0^2 - v_0^2) \cos 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \\
& + 4v_0 w_0 (\beta^2 + 1)^2 \sin 2(\theta - \theta_0) \sinh 2(\chi - \chi_0) \} \\
& + I + J,
\end{aligned} \tag{3.33}$$

donde I y J representan todos los términos puramente imaginarios y puramente hiperbólicos, los cuales no nos interesan, ya que lo que queremos es la suma $(z_0 \bar{z})^2 + (\bar{z}_0 z)^2$, pero un sumando es el bicomplejo conjugado del otro, así que al sumarlos los términos en i y en j se anularán. Debemos elegir adecuadamente los parámetros θ , χ , θ_0 , χ_0 , v_0 y w_0 para eliminar el *mixing* (es decir, los términos cuadráticos proporcionales a vw) en la expresión anterior. Se ve que para eliminar los términos cruzados, podemos elegir $\theta_0 = \theta = 0$, $\chi_0 = \chi = 0$ y $v_0 w_0 = 0$. Esta última condición es de interés particular, ya que nos implica que alguno de los parámetros v_0 ó w_0 debe anularse. Esto a su vez nos permitirá calcular el valor que debe tener la constante de proporcionalidad original β en función solamente de los parámetros de masa real e híbrida de la lagrangiana original. En efecto, sumando y restando (3.32) de (3.31) obtenemos, respectivamente,

$$2v_0^2 = \frac{6m_1^2}{(1 - \beta^2)\lambda} + \frac{3m_2^2}{\beta\lambda} \tag{3.34}$$

y

$$2w_0^2 = \frac{6m_1^2}{(1 - \beta^2)\lambda} - \frac{3m_2^2}{\beta\lambda}. \tag{3.35}$$

Multiplicando estas dos ecuaciones y sacando raíz cuadrada, obtenemos que, salvo un signo,

$$v_0 w_0 = \frac{1}{2} \sqrt{\left(\frac{6m_1^2}{(1-\beta^2)\lambda}\right)^2 - \left(\frac{3m_2^2}{\beta\lambda}\right)^2}. \quad (3.36)$$

Y como este producto debe ser igual a cero, entonces el radicando se debe anular, es decir que debemos tener

$$\frac{6m_1^2}{(1-\beta^2)\lambda} = \frac{3m_2^2}{\beta\lambda}. \quad (3.37)$$

Despejando β , obtenemos finalmente

$$\beta = \frac{-m_1^2 \pm \sqrt{m_1^4 + m_2^4}}{m_2^2}; \quad (3.38)$$

además, la ecuación (3.37) implica que

$$w_0 = 0 \quad , \quad v_0^2 = \frac{6m_1^2}{(1-\beta^2)\lambda} = \frac{3m_2^2}{\beta\lambda}. \quad (3.39)$$

Usando estas condiciones obtenemos

$$[(\bar{z}_0 z)^2 + (z_0 \bar{z})^2] = 2v_0^2[(\beta^2 - 1)^2 - 4\beta^2]v^2 - 2v_0^2(\beta^2 + 1)^2 w^2 + 8ij\beta(1 - \beta^2)v_0^2 v^2. \quad (3.40)$$

Así, se puede ver que el único término de masa sobreviviente es el de masa real de w :

$$\frac{m^2}{2} \bar{z}z + \frac{\lambda}{4!} [(\bar{z}_0 z)^2 + (z_0 \bar{z})^2] = [(\beta^2 - 1)m_1^2 - 2\beta m_2^2]w^2. \quad (3.41)$$

Entonces identificamos la masa del campo escalar real w como $m_w = 2(\beta^2 - 1)m_1^2 - 4\beta m_2^2$. Si elegimos el signo positivo en la ecuación (3.38), obtenemos que

$$m_w = 4 \left(m_1^2 - \sqrt{m_1^4 + m_2^4} \right) \left(1 + \frac{m_1^2}{m_2^2} \right), \quad (3.42)$$

lo cual es negativo para cualquier valor real distinto de cero de m_2 . Por otro lado, eligiendo el signo negativo en (3.38), obtenemos

$$m_w = 4 \left(m_1^2 + \sqrt{m_1^4 + m_2^4} \right) \left(1 + \frac{m_1^2}{m_2^2} \right), \quad (3.43)$$

lo cual es positivo siempre que $m_1^2 > 0$.

En resumen, iniciando con una masa bicompleja hermítica para el campo z , hemos llegado a la aparición de un bosón de Goldstone, a saber, v , el cual pierde su masa después del rompimiento de simetría y la elección apropiada del vacío; mientras que el campo w ha perdido la componente híbrida de su masa y, con una elección adecuada del parámetro β , ha quedado con masa real positiva.

Finalmente, podemos fijarnos en la geometría del vacío. Rotando el punto del vacío por fases hiperbólicas y elípticas obtenemos un número bicomplejo con todas sus partes (real, imaginaria, hiperbólica y bicompleja). Sin embargo, la variedad vacío es bidimensional, ya que sólo tenemos dos parámetros para las rotaciones. Es decir que nuestro vacío es una superficie bidimensional encajada en un espacio euclídeo de cuatro dimensiones. Para graficarlo, hacemos, por ejemplo, la proyección $z = x + iy + ju + iw \rightarrow (x, y, u)$, es decir, simplemente suprimimos una de las componentes. Tal

proyección se muestra en la figura 3.3, donde se puede apreciar una autointersección, la cual debe ser producto de la proyección y no una característica intrínseca de la variedad. Otra vista “más lejana” del vacío se muestra en la figura 3.4, en la cual se ve que es homomorfo a dos planos (que se intersectan) con un disco removido en el centro de cada uno de ellos.

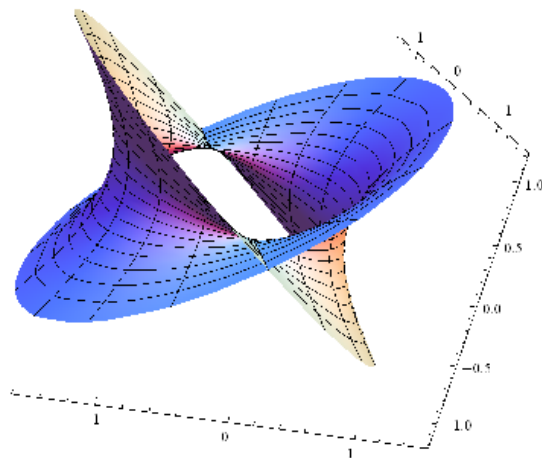


Figura 3.3: Vacío del modelo bajo la deformación hiperbólica.

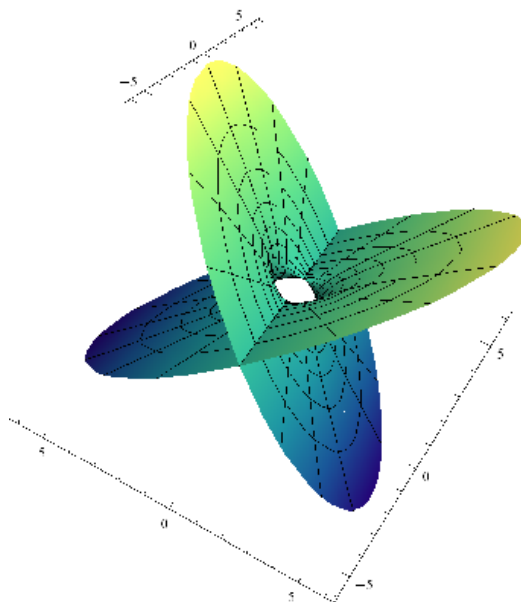


Figura 3.4: Una vista más amplia del vacío. Se aprecia que, topológicamente, consiste en dos planos “agujerados” que se intersectan.

De las figuras se puede ver, a pesar de ser proyecciones, que la superficie generada es en cierto sentido el producto directo de una hipérbola con un círculo. Las hipérbolas se aprecian en los “costados” de la figura 3.3, mientras que los círculos se pueden observar más claramente en la forma de los planos de la figura 3.4. El agujero que aparece en ambas figuras, al centro de la superficie, nos dice que el grupo fundamental de esta es no trivial, es decir que en este sentido “hereda” la topología de su espacio factor $U(1)$, cuyo primer grupo de homotopía es (isomorfo a) $\mathbb{Z}$.

Por otro lado, el grupo $H(1)$, visto como variedad, es una hipérbola (de dos hojas), que tiene la propiedad $\pi_0 = 2$, sin embargo, cada parte conexa es isomorfa a la recta real, la cual es simplemente conexa, es decir, tiene grupo fundamental trivial. Entonces, bajo la extensión bicompleja del modelo, esta propiedad de la variedad vacío se complica y, como ya dijimos, da lugar a una variedad que no es simplemente conexa.

Conclusiones

Los números hiperbólicos, al menos desde un punto de vista matemático, poseen interés propio, al igual que los complejos. Ambos son casos especiales de álgebras reales bidimensionales (ver [11]). Los números hiperbólicos, sin embargo, presentan varias estructuras diferentes a los elípticos. Por un lado, están las propiedades algebraicas: los números hiperbólicos no poseen una norma que sea positiva definida y, de hecho, hay elementos de $\mathbb{H}$ no nulos que tienen “norma” nula, por lo tanto la división no es posible en general para los números hiperbólicos. Sin embargo, al ser la lagrangiana un polinomio en las variables de campo, la falta de la operación de división no afecta de ningún modo la construcción de la teoría. Por lo tanto debería considerarse seriamente, en general, la posibilidad de construir densidades lagrangianas de campos cuyos valores estén en algún anillo, es decir, sobre un conjunto de números que admitan mínimamente adición y multiplicación (una propuesta a lo largo de las mismas líneas se encuentra en [9] para la mecánica cuántica). Este trabajo es un pequeño paso en tal dirección.

Otro aspecto interesante de nuestro modelo es la no compacidad del grupo de simetría. Las dos grandes teorías del siglo pasado son sin duda el modelo estándar y la relatividad general. El primero de ellos, el cual describe con mucha precisión el mundo microscópico, está basado en el grupo de norma $SU(3)_C \times SU(2)_L \times U(1)_Y$ (ver [14]), el cual es un grupo compacto. Este grupo da lugar a una rica fenomenología que se comprueba día con día en los grandes aceleradores. La segunda teoría, por su lado, tiene como grupo de simetría el grupo de los difeomorfismos (ver [17]), el cual forma por sí mismo una variedad no compacta. Así, si se quiere ver la relatividad general como una teoría cuántica de campos, será inevitable el tener que estudiar las propiedades de los grupos no compactos. De nuevo, este trabajo es un paso en tal dirección.

Como un último aspecto interesante, quisiera recalcar el proceso de elección del mínimo del potencial de la última sección. En general, sólo podemos hablar de mínimos cuando tratamos con conjuntos de números reales, ya que éstos poseen una estructura de orden bien definida. Al igual que con los complejos, a los números hiperbólicos no es posible darles una estructura de orden (tampoco a los bicomplejos), es decir, una relación entre sus elementos que sea reflexiva, antisimétrica y transitiva. Así que, estrictamente, la elección de un mínimo para el potencial, el cual es una función no real, no es posible en general. Sin embargo, como se comentó en el cuerpo de la tesis, trabajamos con una generalización de los números reales, a saber, los bicomplejos hermíticos, es decir, aquellos de la forma $a + ijb$, con $a, b \in \mathbb{R}$. Los reales son un subconjunto de éstos y ambos sistemas comparten la característica de que sus elementos son invariantes ante conjugación. Así, pues, tratamos una función hermítica como lo haríamos con una real, y la manera más natural de generalizar los mínimos de tal función es definiéndolos como aquellos en donde el gradiente de la función es igual a cero. Como se vio, esto nos permitió seguir con el cálculo y llegar a los resultados finales.

Bibliografía

- [1] J. Cockle, London-Edinburgh-Dublin Philosophical Magazine. **33**, 37 (1848).
- [2] J. Hucks, Journal of Mathematical Physics. **34**, 5986 (1993).
- [3] S. Ulrych, Physics Letters B. **625**, 313 (2005).
- [4] S. Ulrych, Physics Letters B. **632**, 417 (2005).
- [5] S. Ulrych, Physics Letters B. **633**, 631 (2006).
- [6] J. B. Fraleigh, *A First Course in Abstract Algebra*. Addison-Wesley-Longman, 6a. edición, 2000.
- [7] I. N. Herstein, *Topics in Algebra*. John Wiley & Sons, 2da. edición, 1975.
- [8] J. W. Brown, R. V. Churchill, *Complex Variables and Applications*. McGraw-Hill, 8a. edición, 2009.
- [9] J. Kocik, International Journal of Theoretical Physics. **38**, 2221 (1999).
- [10] I. L. Kantor, A. S. Solodovnikov, *Hypercomplex Numbers. An Elementary Introduction to Algebras*. Springer-Verlag, 1989.
- [11] I. M. Yaglom, *Complex Numbers in Geometry*. Academic, New York, 1968.
- [12] D. J. Griffiths, *Introduction to Quantum Mechanics*. Pearson-Prentice Hall, 2a. edición, 2005.
- [13] E. Hecht, *Optics*. Addison Wesley, 4a. edición, 2002.
- [14] M. E. Peskin y D. V. Schroeder, *An introduction to Quantum Field Theory*. Westview Press, 1995.
- [15] W. Rindler, *Relativity. Special, General and Cosmological*. Oxford University Press, 2006.
- [16] D. J. Griffiths, *Introduction to Elementary Particles*. Wiley-Vch, 2a. edición, 2008.
- [17] S. Carroll, *Spacetime and Geometry. An introduction to General Relativity*. Addison-Wesley, 2004.
- [18] M. Srednicki, *Quantum Field Theory*. Cambridge University Press, 2007.
- [19] F. J. Dyson, Physical Review. **75**, 1736 (1949).
- [20] G. Costa, G. Fogli, *Symmetries and Group Theory in Particle Physics. An Introduction to Space-time and Internal Symmetries*. Springer, 2012.

-
- [21] F. Mandl, G. Shaw, *Quantum Field Theory*. John Wiley & Sons, Ltd. 2a. edición, 2010.
 - [22] D. Colladay, V. A. Kostelecky, Physical Review D. **58**, 116002 (1998).
 - [23] S. Ulrych, Physics Letters B. **612**, 89 (2005).
 - [24] M. Nakahara, *Geometry, Topology, and Physics*. IoP Publishing, 2da. edición, 2003.
 - [25] J. R. Munkres, *Topology*. Prentice Hall, 2da. edición, 2000.
 - [26] J. E. Marsden, A. J. Tromba, *Vector Calculus*. Addison Wesley, 5a. edición, 2004.